



Proceeding of

the 11th Seminar on

Reliability Theory and its Applications

Department of Statistics
University of Isfahan

14–15 May 2025



In the Name of God



Proceeding of
the 11th Seminar on
Reliability Theory and its Applications

Department of Statistics

University of Isfahan

Isfahan, Iran

14-15 May, 2025

This book contains the proceeding of the 11th 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 11th Seminar on Reliability Theory and its Applications

By: Somayeh Ashrafi

Cover Designer: Ali Maarefvand

Date: September, 2025

Preface

We are delighted to host the “11th Seminar on Reliability Theory and its Applications” on 14-15 May 2025 in the Department of Statistics at the University of Isfahan. On behalf of the organizing committee, we warmly welcome all participants and trust that this seminar will provide engaging discussions and the exchange of valuable scientific ideas. We are grateful to the many individuals and institutions whose support has made this event possible, including our colleagues in other universities, researchers, and postgraduate students whose contributions enrich this event. In particular, we wish to thank the Research Council of the University of Isfahan, the “Center of Excellence on Ordered Data, Reliability, and Dependency”, the Islamic World Science Citation Center (ISC), the Iranian Statistical Society, the scientific and the organizing committees of the seminar, our peer reviewers, and the staff of the Department of Statistics at the University of Isfahan for their kind cooperation.

Mahdi Tavangar (Scientific Chair)

Somayeh Ashrafi (Organizing Chair)

September, 2025

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. Amini Seresht, E., Bu-Ali Sina University
3. Asadi, M., University of Isfahan
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. Mansourvar, Z., University of Isfahan
13. Razmkhah, M., Ferdowsi University of Mashhad
14. Sharafi, M., Razi University
15. Tavangar, M., University of Isfahan (Chair)
16. Torabi, H., Yazd University
17. Zarezadeh, S., Shiraz University

Honorary Scientific Advisory Committee

1. Ahmadi, J., Ferdowsi University of Mashhad
2. Khaledi, B., Razi University

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

Organizing Committee

1. Ahmadi, J., Ferdowsi University of Mashhad
2. Asadi, M., University of Isfahan
3. Ashrafi, S., University of Isfahan (Chair)
4. Jabari, H., Ferdowsi University of Mashhad
5. Mansourvar, Z., University of Isfahan
6. Tavangar, M., University of Isfahan
7. Zamanzade, E., University of Isfahan

Contents

An Approach to Modeling the Mean Residual Time Ahmadi, R., Jamali, M.	1
Stochastic Comparisons of Extremes with Dependent and Heterogeneous Observations from a General Two Parameters Distribution Al-jabbar, A., Kelkinnama, M.	7
Best Unbiased Prediction of Order Statistics in Generalized Extreme Value Distribution Almasi, I., Esmailzadeh, N.	15
A New Test Based on Linear Ordered Rank Statistics Amini-Seresht, E., Izadi, M.	30
Preponed Preventive Maintenance for k-out-of-n Systems Subject to External Shocks Ashrafi, S., Asadi, M.	35
Analysing Random Maintenance Data by Repair Alert Model for Discrete lifetimes Atlehkhani, M., Doostparast, M.	45
A Q-Learning Approach to Maintenance Policy Optimization for Series Systems with Degradation and Common Shocks Farahmand, A. H., Fartash, F., Yadollahi, M. R., Zarandi, S.	53
Optimizing Maintenance Scheduling in a Scrap-Based Steel Production Line Using Reinforcement Learning Haerian, N., Gholami, Z., Safari, A., Haghighi, F.	65
Efficiency of Ranked Set Sampling in Mean Residual Lifetime Estimation Jabari Koopaei, L., Zamanzade, E.	81

Mission Abort Strategy for Multi-Component Systems Using Survival Signature	
Karimi, A., Tavangar, M.	93
Burg Entropy: Extension and an Application to Binary Classification of Malaria Cell Images	
Kharazmi, O., Asadi, A.	106
Binary Classification of Heart Disease Data Based on Cumulative Residual Discrete χ^2_α Divergence	
Kharazmi, O., Bahrehvar, M.	116
Causal Inference on the Differences of the Restricted Mean Residual Lifetime	
Mansourvar, Z.	125
Analyzing the Doubly Geometric Process with Flexible Interarrival Times: A Tool for Modeling Failure and Reliability Data	
Mohammadi, Z.	133
A New Design of a Combined Control Chart Based on Exponentially Weighted Moving Average for Normal Quality Variables	
Moradi, N., Salehi, E.	145
Optimized Maintenance Planning Under Random Delayed Interventions: An RL-Based Approach	
Rahimi, G., Karimi, M., Rafeighi, M., Vesali, N., Safari, A., Haghighi, F.	155
Reliability of Complex k-out-of-$n:d$ Systems	
Razmkhah, M., Khatib Astaneh, B.	172
Analyzing Rayleigh-Poisson Model for Censored Data with Informative Random Removals	
Sharafi, M., Fallah Hasan, A.	179



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



An Approach to Modeling the Mean Residual Time

Ahmadi, R.¹ Jamali, M.²

¹ Iran University of Science and Technology, Iran

² Faculty of Mathematics, Statistics and Computer Science, University of Tabriz, Tabriz, Iran

Abstract

There has been increasing interest in employing the mean residual lifetime (MRL) as a key index for failure prognosis. But, it is not flexible enough to accommodate condition monitoring information required for maintenance decision-making. The primary innovation of this work lies in computing the mean residual time based on the expected time to reach a defective state, rather than waiting for complete system failure. Further in this study, the characterization problem of the proposed reliability index is addressed and its monotonicity properties are analyzed.

Keywords: Mean residual lifetime, Failure prognosis, Parallel system configuration, System inspection

1 Introduction

For multi-component systems one of the most important system structures is the parallel structure. Nuclear reactor safety systems, emergency core cooling systems, fire detectors, and protective device are examples of systems with parallel structure [14, 1]. It is well known that, for a parallel configuration with n components, the system works

¹Ahmadi_Reza@iust.ac.ir

²mahsajamali_90@yahoo.com

whenever at least one of its components works. In this regard, Jamali et al. [10] propose a generic MRL model for the reliability analysis of a load-sharing coherent system by employing the signature technique and a generalized Farlie-Gumbel-Morgenstern (FGM) copula function. Therefore, utilizing mean residual life (MRL) information provides a more accurate assessment of system reliability over time[8].

A reliability index used in a parallel system is the *Mean Residual Life (MRL)*. Knowing that the system has up to a specific age, MRL provides an idea of how long this system can be expected to survive. Denoting by T_f the system lifetime, MRL is

$$\text{MRL}(t) = \mathbb{E}[T_f - t \mid T_f > t], \quad t \geq 0 \quad (1.1)$$

which is the expected remaining life of the system given it has survived up to t . Although MRL is known as a key tool in reliability and survival analysis, it is not flexible enough to accommodate actual situations. To get more accurate estimates, authors attempt to upgrade it through incorporating system information. System information can be related, for instance, to the system degradation [9], to the components failure process [6, 5], or to covariates [7, 2, 3]. In a multi-component system, when the number of failed components exceeds a determined defective threshold, the system can be less efficient, causing losses of production, energy, capital or even environmental damages [13, 15].

In this paper, a parallel system consisting of n identical components is studied. Before the system breaks down, there will be signs of reduced performance resulting from the number of failed components. We assume that, if d components are not functioning ($d < n$), the system is less efficient and it is working with a reduced performance. If the number of failed components does not exceed d , the system is properly working.

In this paper, we develop a detailed analytical framework for evaluating the state-dependent mean residual lifetime (SDMRL) function for parallel systems with identical components. Initially, an explicit integral representation of the SDMRL, $\bar{m}_{1:n}(t, d, j)$, is established. Then, the monotonicity of the SDMRL with respect to t and the parameter d is considered.

Additionally, we express the SDMRL for parallel systems in terms of the mean residual lifetime (MRL) of equivalent series systems. An inverse relationship is also derived, allowing the MRL of a series system to be represented through a combination of SDMRLs of parallel systems with varying configurations. These results offer deeper insights into the dynamic behavior of system reliability under progressive component failures.

2 Framework of the model

2.1 Model assumptions

A parallel system consisting of n identical components is analyzed. The following assumptions are given:

1. Each component is identical and operates independently. The failure time of each component is modeled as a continuous random variable with probability density function f and cumulative distribution function F .

2. At time 0, the system is completely new.
3. If d components are not functioning ($d < n$), the system is in a defective state. That means the system continues to operate but with diminished performance.

2.2 Features of the SDMRL

Let $X(t)$ be the number of failed components at time t , and $T_{[1]}, T_{[2]}, \dots, T_{[n]}$ denote the order statistics of lifetimes $T_1, T_2, \dots, T_n$. According to assumption (3), the d -th order statistic $T_{[d]}$ (for $d = 1, 2, \dots, n-1$) can be interpreted as the time at which the system enters in a defective state.

Then, given $X(t) = j$ with $j < d$, the state-dependent mean residual time to reach $T_{[d]}$ is represented as

$$\bar{m}_{1:n}(t, d, j) = \mathbb{E}[T_{[d]} - t \mid X(t) = j]. \quad (2.1)$$

Notation: $1 : n$ indicates that a system with n components operates correctly if at least 1 out of n components are operating. Intuitively, the SDMRL of a series system can be written as $\bar{m}_{n:n}(t, 1, 0)$.

Lemma 2.1. *Let $\bar{m}_{1:n}(t, d, j)$ be the state-dependent mean residual time given in Eq. (2) with $j < d$ for a parallel system of n identical components. Then*

$$\bar{m}_{1:n}(t, d, j) = \int_t^\infty \sum_{k=j}^{d-1} B(n-k, n-j, \bar{G}(v, t)) dv \quad (2.2)$$

where

$$\bar{G}(v, t) = \frac{\bar{F}(v)}{\bar{F}(t)}, \quad v \geq t. \quad (2.3)$$

$$\mathcal{B}(n-k, n-j, \bar{G}(v, t)) = \binom{n-j}{n-k} \bar{G}(v, t)^{n-k} (1 - \bar{G}(v, t))^{k-j}.$$

Proof. For proof see [4]. □

For fixed t and j , it is trivial that $\bar{m}_{1:n}(t, d, j)$ is increasing in d . Lemmas 2.2 and 2.3 analyze the monotony of $\bar{m}_{1:n}(t, d, j)$ given in Eq. (2.2) with respect to t and j .

Lemma 2.2. *For fixed t and d , the SDMRL $\bar{m}_{1:n}(t, d, j)$ given by Eq. (2.2) is a decreasing function of j .*

Proof. For proof see [4]. □

Lemma 2.3. *Under the assumption of the increasing failure rate (IFR) of the lifetime of the components, for fixed n , j and d , $\bar{m}_{1:n}(t, d, j)$ given in Eq. (2.1) is a decreasing function of t .*

Proof. For proof see [4]. □

Example 2.4. Figure 1 demonstrates the state-dependent mean residual time as a function of time t , for different values of the number of failed components $j \in \{5, 10, 15\}$ and Weibull shape parameters $a \in \{1.1, 1.2, 1.3\}$, with fixed $d = 25$. This expectation has been computed using a parallel system consisting of 30 identical components, assuming that the system reduces its performance when the number of failed components equals $d = 25$. The components' lifetime follows a Weibull distribution with shape parameter a and scale parameter 1. As we can check in Figure 1, the state-dependent mean residual time $m_{1:30}(t, 25, j)$ is decreasing with respect to t for different values of the shape parameter a and numbers of failed components $j \in \{5, 10, 15\}$.

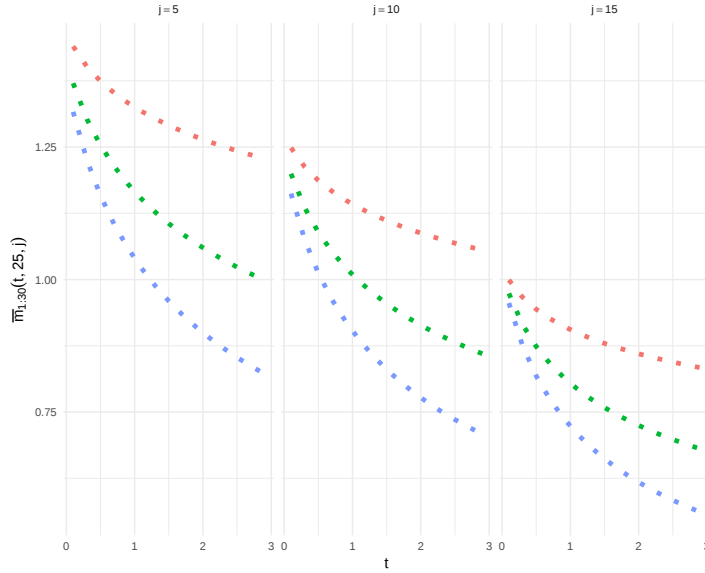


Figure 1: Comparison of the SDMRL function $\bar{m}_1^{30}(t, 25, j)$ trends over time for different shape parameters $a \in \{1.1, 1.2, 1.3\}$ and numbers of failed components j in a parallel system with a Weibull lifetime distribution. Red, green and blue color correspond to $a \in \{1.1, 1.2, 1.3\}$ respectively.

Example 2.5. Figure 2 shows the state-dependent mean residual time versus t . This expectation has been computed using a parallel system consisting of $n = 30$ identical components and assuming that the system exhibits a performance reduction when the number of failed components is equal $d \in \{20, 25, 30\}$. The components' lifetime follows a Weibull distribution with shape parameter $a \in \{1.1, 1.2, 1.3\}$ and scale parameter equals to 1.

We further assume that, at time t , there are $j = 5$ failed components in the system. As we can check in Figure 2, $\bar{m}_{1:30}(t, d, 5)$ is decreasing with respect to t .

Before addressing the characterization problem we provide the following auxiliary result.

Remark 2.6. Let $\bar{m}_{n:n}(t, 1, 0)$ be the mean residual life at time t of a series system consisting of n identical and independent components with lifetime distribution F . That is,

$$\bar{m}_{n:n}(t, 1, 0) = \int_t^\infty \bar{G}(v, t)^n dv, \quad (2.4)$$

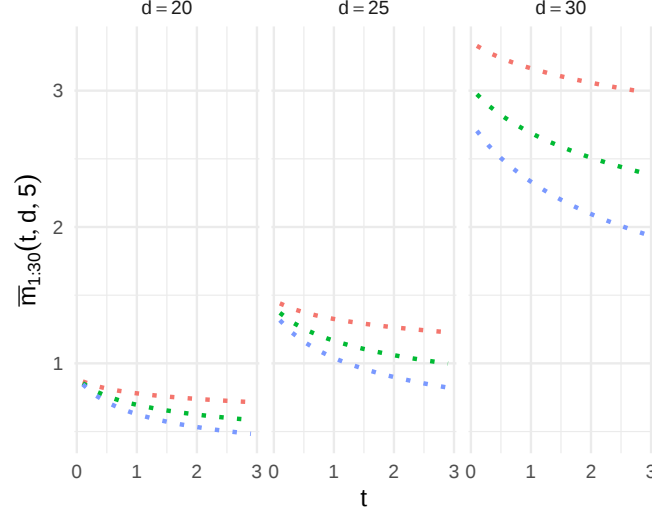


Figure 2: Behavior of the SDMRL function $\bar{m}_{1:30}(t, d, 5)$ over time for a Weibull distribution with shape parameters $a \in \{1.1, 1.2, 1.3\}$ and different threshold values $d \in \{20, 25, 30\}$, with fixed $j = 5$. Red, green and blue color correspond to $a \in \{1.1, 1.2, 1.3\}$ respectively.

where $\bar{G}(\cdot)$ is given by Eq. (2.3). Next result addresses the characterization problem of the SDMRL with respect to an ordinary MRL and vice versa.

2.3 Characterization

Lemma 2.7. *Let $\bar{m}_{n:n}(t, 1, 0)$ be the mean residual lifetime at time t of a series system with n identical and independent components. Then the state-dependent mean residual time given by Eq. (2.1) can be expressed in terms of $\bar{m}_{n:n}(t, 1, 0)$ as follows:*

$$\bar{m}_{1:n}(t, d, j) = \sum_{k=1}^{d-j} \binom{n-j}{n-d+k} \binom{k+n-d-1}{k-1} (-1)^{k+1} \bar{m}_{k+n-d+k+n-d}(t, 1, 0).$$

Proof. For proof see [4] . □

Lemma 2.7 allows to inversely represent the ordinary MRL of a series system with respect to the SDMRLs of parallel systems. More specifically, for $k = 0, 1, \dots, n-j-1$ we have:

$$\bar{m}_{n-j:n-j-k}(t, 1, 0) = \sum_{w=1}^{k+1} \frac{\binom{n-j-w}{k-w+1}}{\binom{n-j}{n-j-k}} \bar{m}_{1:n}(t, j+w, j)$$

Proof. For proof see [4] . □

References

- [1] Ahmadi, R. (2020), Reliability modeling and maintenance planning for a parallel system, *Journal of Reliability Engineering*, **45**(3), 123–135.

- [2] Ahmadi, R. (2021), A novel MRL-based policy for systems under imperfect inspection, *Reliability Engineering & System Safety*, **205**, 107–117.
- [3] Ahmadi, R. et al. (2022), Maintenance optimization based on expected time to defectiveness, *Journal of Risk and Reliability*, **236**(1), 89–102.
- [4] Ahmadi, R., Castro, I.T. and Bautista, L. (2024), Reliability modeling and maintenance planning for a parallel system with respect to the state-dependent mean residual time, *Journal of the Operational Research Society*, **75**(2), 297–313.
- [5] Asadi, M. and Bayramoglu, K. (2005), On some stochastic comparisons of residual lifetimes, *Statistics and Probability Letters*, **71**(2), 123–131.
- [6] Bairamov, I.G., Bayramoglu, K. and Ozkut, M. (2002), Mean residual life and failure rate functions under covariates, *Journal of Applied Probability*, **39**, 48–58.
- [7] Banjevic, D., Jardine, A.K.S., Makis, V. and Ennis, M. (2005), A control-limit policy and software for condition-based maintenance optimization, *INFOR*, **43**(1), 32–44.
- [8] Dong, Y., Huynh, K.T. and Castanier, B. (2020), A novel age-based policy using residual life information, *Reliability Engineering & System Safety*, **200**, 106943.
- [9] Huynh, K.T., Barros, A. and Berenguer, C. (2013), A decision model for condition-based maintenance optimization, *IEEE Transactions on Reliability*, **62**(2), 439–450.
- [10] Jamali, M., Ahmadi, R., and Bevrani, H. (2025), Signature-based Reliability and Maintenance Planning for a Load-sharing Coherent System Operating in a Random Environment, *Journal of Reliability Engineering*, **50**(2), 200–220.
- [11] Levitin, G. (2008), *Optimal Allocation in Redundant and Fault Tolerant Systems*, Springer, Berlin.
- [12] Lisnianski, A. and Levitin, G. (2003), *Multi-State System Reliability: Assessment, Optimization and Applications*, World Scientific Publishing Co., Singapore.
- [13] Santos, M. and Cavalcante, C.A.V. (2018), Modeling defective states in complex systems: Environmental and economic impact, *Journal of Cleaner Production*, **198**, 972–981.
- [14] Smidt, R. and Johnson, P., Performance degradation in phased array radar systems, *IEEE Transactions on Aerospace and Electronic Systems*, **43**(4), 1235–1242.
- [15] Zhang, X. and Zeng, Y. (2014), On the performance of parallel multi-component systems, *International Journal of Systems Science*, **50**(6), 789–799.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Stochastic Comparisons of Extremes with Dependent and Heterogeneous Observations from a General Two Parameters Distribution

Al-jabbar, A.¹ Kelkinnama, M.²

Department of Mathematical Sciences, Isfahan University of Technology, Isfahan, Iran

Abstract

In this paper, we carry out stochastic comparisons of lifetimes of two series (and two parallel) systems consisting of dependent and heterogeneous components. The underlying dependence model for component lifetimes are supposed to be an Archimedean copula. Also, we assume a general two parameters distribution for the lifetime distribution of components. Sufficient conditions are developed to get the usual stochastic ordering between two systems. Example is given to illustrate the theoretical findings.

Keywords: Stochastic comparisons, Majorization, Archimedean copula

1 Introduction

The k th smallest variable of the set of random variables $X_1, \dots, X_n$ is denoted by order statistic $X_{k:n}$, $k = 1, \dots, n$. The role of order statistics in various branches of statistics such as reliability theory, survival analysis and actuarial science, is clear for all. In reliability theory, $X_{k:n}$ is the lifetime of a well-known $(n - k + 1)$ -out-of- n system. This system

¹amal@math.iut.ac.ir

² m.kelkinnama@iut.ac.ir

works if at least $n - k + 1$ of its n components function and specially includes series system $(X_{1:n})$ and parallel system $(X_{n:n})$. Therefore, the study of stochastic properties of order statistics aims the study of lifetimes of $(n - k + 1)$ -out-of- n system.

Numerous researchers have examined stochastic comparisons between order statistics from specific families of distributions when the underlying random variables have a heterogeneity in the parameters vector. These comparisons were initially made for common and well-known distributions such as exponential, gamma, and Weibull and were then studied by many researchers for other more flexible distributions such as exponentiated Weibull, a generalized modified Weibull model, negative binomial model, Chen's distribution, Gompertz-Makeham model etc. Also, some authors have extended the results to a larger class of distributions, such as location-scale family, proportional odds model, exponentiated location-scale models, etc. Most early papers in this field performed random comparisons assuming independence of components. In practical situations, the components of a system may have a dependence structure which result in a set of statistically dependent observations. The dependence condition of the components are mostly investigated with the help of copulas. In the last decade, many researches focuses on stochastic comparison on order statistics assuming the dependency among random variables, among them we refer to Fang et al. [4] for scaled proportional hazards model, Barmalzan et al. [1] for extended exponential distribution, Das et al. [3] for exponentiated location-scale models, Chowdhury et al. [2] for Weibull-G distribution, and Samanta et al. [6] for extended Weibull distribution.

In a more general framework, Nadeb et al. [5] assumed that the component distributions have a general form as $\bar{F}(x, \lambda)$, and under the Archimedean copula obtained some results on comparing series (parallel) systems with heterogeneous components. In obtaining the results, they assumed some conditions on the behavior of $\bar{F}$ w.r.t λ . In this subject, Sharahili [7] also obtained some other results.

In this paper, we assume a general two parameters distribution for the components' lifetimes (say $F(x; \alpha, \theta)$) and obtain sufficient conditions for satisfying the usual stochastic ordering between two series (and two parallel) systems with two vectors of parameters. In fact, we aim to generalize the previous two mentioned works from one parameter family to two parameters one.

2 Preliminaries

In this section, we review some definitions, lemmas and notations which will be used to present the main results of the article.

Definition 2.1. Let X and Y be two non-negative random variables with cumulative distribution functions (cdf) F and G and then survival (reliability) functions $\bar{F} = (1 - F)$ and $\bar{G} = (1 - G)$, respectively. We say that X is smaller than Y in the usual stochastic order, denoted by $X \leq_{st} Y$, if $\bar{F}(t) \leq \bar{G}(t)$ or equivalently $F(t) \geq G(t)$, $\forall t \in [0, \infty)$.

Definition 2.2. For two real vectors $\mathbf{x} = (x_1, \dots, x_n)$ and $\mathbf{y} = (y_1, \dots, y_n)$, let $x_{1:n} \leq x_{2:n} \leq \dots \leq x_{n:n}$ and $y_{1:n} \leq y_{2:n} \leq \dots \leq y_{n:n}$ be the increasing arrangement of them, respectively. Then the vector $\mathbf{x}$ is said to

1 . weakly supermajorize $\mathbf{y}$ (written as $\mathbf{x} \succeq^w \mathbf{y}$) if $\sum_{i=1}^j x_{i:n} \leq \sum_{i=1}^j y_{i:n}$ for $j = 1, \dots, n$.

2 . weakly submajorize $\mathbf{y}$ (written as $\mathbf{x} \preceq_w \mathbf{y}$) if $\sum_{i=j}^n x_{i:n} \geq \sum_{i=j}^n y_{i:n}$ for $j = 1, \dots, n$.

Lemma 2.3. *Let ϕ be a real-valued function, defined and continuous on $D = \{\mathbf{x} : x_1 \geq x_2 \geq \dots \geq x_n\}$ and continuously differentiable on the interior of D . Denote the partial derivative of ϕ with respect to its k th argument by $\phi_{(k)}(z) = \frac{\partial \phi(z)}{\partial z_k}$. Then,*

- i. $\phi(\mathbf{x}) \leq \phi(\mathbf{y})$ whenever $\mathbf{x} \preceq_w \mathbf{y}$ on D if and only if $\phi_{(1)}(z) \geq \phi_{(2)}(z) \geq \dots \geq \phi_{(n)}(z) \geq 0$, i.e, the gradient $\nabla \phi(z) \in D_+ = \{\mathbf{x} : x_1 \geq x_2 \geq \dots \geq x_n \geq 0\}$, for all z in the interior of D .
- ii. $\phi(\mathbf{x}) \leq \phi(\mathbf{y})$ whenever $\mathbf{x} \preceq^w \mathbf{y}$ on D if and only if $0 \geq \phi_{(1)}(z) \geq \phi_{(2)}(z) \geq \dots \geq \phi_{(n)}(z)$, i.e, the gradient $\nabla \phi(z) \in D_- = \{\mathbf{x} : 0 \geq x_1 \geq x_2 \geq \dots \geq x_n\}$, for all z in the interior of D .

Consider a random vector $\mathbf{X} = (X_1, X_2, \dots, X_n)$ with joint distribution function F , joint survival function $\bar{F}$, univariate cdfs $F_1, \dots, F_n$ and univariate survival functions $\bar{F}_1, \dots, \bar{F}_n$. We have

$$F(x_1, \dots, x_n) = C(F_1(x_1), \dots, F_n(x_n)); \bar{F}(x_1, \dots, x_n) = \bar{C}(\bar{F}_1(x_1), \dots, \bar{F}_n(x_n)),$$

where C and $\bar{C}$ are called the copula and survival copula of random vector $\mathbf{X}$, respectively. An important class of copulas is Archimedean copula, which its definition is given at the sequel.

Definition 2.4. For a decreasing and continuous function $\psi : [0, +\infty] \mapsto [0, 1]$ such that $\psi(0) = 1$ and $\psi(+\infty) = 0$, let $\phi = \psi^{-1}$ be the pseudo-inverse. Then

$$C_\psi(u_1, \dots, u_n) = \psi(\phi(u_1) + \dots + \phi(u_n)), u_i \in (0, 1), i = 1, \dots, n,$$

is said to be an Archimedean Copula with generator ψ if $(-1)^k \psi^k(x) \geq 0$ for $k = 0, \dots, n-2$ and $(-1)^{n-2} \psi^{n-2}(x)$ is decreasing and convex.

The following lemmas are useful to establish the main results. First, recall that a function f is said to be super-additive if $f(x+y) \geq f(x) + f(y)$, for all x and y in the domain of f .

Lemma 2.5. *For two n -dimensional Archimedean Copulas C_{ψ_1} and C_{ψ_2} , if $\phi_2 \circ \psi_1$ is super-additive, then $C_{\psi_1}(\mathbf{u}) \leq C_{\psi_2}(\mathbf{u})$ for all $\mathbf{u} = (u_1, \dots, u_n) \in [0, 1]^n$.*

For convenience, from now on, we denote

$$I^+ = \{(x_1, \dots, x_n) : 0 < x_1 \leq x_2 \leq \dots \leq x_n\}, \quad D^+ = \{(x_1, \dots, x_n) : x_1 \geq x_2 \geq \dots \geq x_n > 0\}$$

3 Main results

Let $\mathbf{X} = (X_1, \dots, X_n)$ be a random vector, we denote $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi)$ for the sample follows a two parameters distriybtion F , with $\boldsymbol{\alpha} = (\alpha_1, \dots, \alpha_n)$, and $\boldsymbol{\theta} = (\theta_1, \dots, \theta_n)$ as the parameter vectors, and ψ is generator of the associated Archimedean copula.

3.1 Maximum of sample

In this subsection, we carry out stochastic comparison of the largest order statistics from dependent and heterogeneous observations having a general two parameters distribution.

Theorem 3.1. *Let $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ and $\mathbf{Y} \sim F(\boldsymbol{\mu}, \boldsymbol{\theta}, \psi_2)$, Suppose that*

- (i) $\boldsymbol{\alpha}, \boldsymbol{\mu} \in I^+, \boldsymbol{\theta} \in D^+$,
- (ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive,
- (iii) F is increasing in α and decreasing in θ , F is log-concave in α and $\frac{\partial F / \partial \alpha}{F}$ is increasing in θ ,

then $\boldsymbol{\alpha} \succeq_w \boldsymbol{\mu}$ implies $X_{n:n} \geq_{st} Y_{n:n}$.

Proof. The respective distribution functions of $X_{n:n}$ and $Y_{n:n}$ are equal to

$$F_{X_{n:n}}(x) = \psi_1\left(\sum_{k=1}^n \phi_1[F(x; \alpha_k, \theta_k)]\right).$$

and

$$F_{Y_{n:n}}(x) = \psi_2\left(\sum_{k=1}^n \phi_2[F(x; \mu_k, \theta_k)]\right).$$

The super-additivity of $\phi_2 \circ \psi_1$ implies that

$$\psi_1\left(\sum_{k=1}^n \phi_1[F(x; \mu_k, \theta_k)]\right) \leq \psi_2\left(\sum_{k=1}^n \phi_2[F(x; \mu_k, \theta_k)]\right)$$

then, it suffices to prove that

$$\psi_1\left(\sum_{k=1}^n \phi_1[F(x; \alpha_k, \theta_k)]\right) \leq \psi_1\left(\sum_{k=1}^n \phi_1[F(x; \mu_k, \theta_k)]\right). \quad (3.1)$$

Let $\psi_1(\sum_{k=1}^n \phi_1[F(x; \alpha_k, \theta_k)]) = W(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$, which is differentiable. Based on the assumption (i), according to Lemma 2.3, we then only need to show that for all $1 \leq i < j \leq n$

$$0 \leq \frac{\partial}{\partial \alpha_j} W(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1) \leq \frac{\partial}{\partial \alpha_i} W(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1). \quad (3.2)$$

Taking the derivative of $W(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ with respect to α_i , $i = 1, \dots, n$ we have

$$\begin{aligned} \frac{\partial}{\partial \alpha_i} W(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1) &= \psi_1' \left(\sum_{k=1}^n \phi_1[F(x; \alpha_k, \theta_k)] \right) \frac{\psi_1(\phi_1[F(x; \alpha_i, \theta_i)])}{\psi_1'(\phi_1[F(x; \alpha_i, \theta_i)])} \\ &\quad \times \frac{\partial F(x; \alpha_i, \theta_i) / \partial \alpha_i}{F(x; \alpha_i, \theta_i)} \geq 0 \end{aligned}$$

where the inequality follows from the assumption (iii) of the theorem (F is increasing in α for all $\theta > 0$).

From the log-convexity of ψ_1 , decreasing behavior of ϕ_1 , and decreasing trend of F in θ for all $\alpha > 0$, we have that $\frac{\psi_1(\phi_1[F(x; \alpha_i, \theta_i)])}{\psi'_1(\phi_1[F(x; \alpha_i, \theta_i)])}$ is non-positive, increasing in $\alpha > 0$ for all fixed θ and decreasing in $\theta > 0$ for all fixed $\alpha > 0$. Therefore, for $1 \leq i < j \leq n$ we have

$$0 \geq \frac{\psi_1(\phi_1[F(x; \alpha_j, \theta_j)])}{\psi'_1(\phi_1[F(x; \alpha_j, \theta_j)])} \geq \frac{\psi_1(\phi_1[F(x; \alpha_i, \theta_i)])}{\psi'_1(\phi_1[F(x; \alpha_i, \theta_i)])} \quad (3.3)$$

Thus, it can be obtained that

$$\begin{aligned} & \frac{\partial}{\partial \alpha_j} W(\boldsymbol{\alpha}, \theta, \psi_1) - \frac{\partial}{\partial \alpha_i} W(\boldsymbol{\alpha}, \theta, \psi_1) \\ &=_{\text{sgn}} \left[\frac{\psi_1(\phi_1[F(x, \alpha_i, \theta_i)])}{\psi'_1(\phi_1[F(x, \alpha_i, \theta_i)])} \times \frac{\partial F(x, \alpha_i, \theta_i)/\partial \alpha_i}{F(x, \alpha_i, \theta_i)} - \frac{\psi_1(\phi_1[F(x, \alpha_j, \theta_j)])}{\psi'_1(\phi_1[F(x, \alpha_j, \theta_j)])} \times \frac{\partial F(x, \alpha_j, \theta_j)/\partial \alpha_j}{F(x, \alpha_j, \theta_j)} \right] \\ &\leq \frac{\psi_1(\phi_1[F(x, \alpha_j, \theta_j)])}{\psi'_1(\phi_1[F(x, \alpha_j, \theta_j)])} \left[\frac{\partial F(x, \alpha_i, \theta_i)/\partial \alpha_i}{F(x, \alpha_i, \theta_i)} - \frac{\partial F(x, \alpha_j, \theta_j)/\partial \alpha_j}{F(x, \alpha_j, \theta_j)} \right] \\ &=_{\text{sgn}} \left[\frac{\partial F(x, \alpha_j, \theta_j)/\partial \alpha_j}{F(x, \alpha_j, \theta_j)} - \frac{\partial F(x, \alpha_i, \theta_i)/\partial \alpha_i}{F(x, \alpha_i, \theta_i)} \right] \\ &\leq 0 \end{aligned}$$

where the first inequality follows from (3.3), and the last one follows from assumption (iii). Hence (3.2) is proved, and from Lemma 2.3, $\boldsymbol{\alpha} \succeq_w \boldsymbol{\mu}$ implies that (3.1) is satisfied and the proof is completed. $\square$

Next, we turn to the case when the vectors of second parameter are majorized. We omit the proof of the following theorem because similarity.

Theorem 3.2. *Let $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ and $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\sigma}, \psi_2)$, Suppose that*

- (i) $\boldsymbol{\theta}, \boldsymbol{\sigma} \in D^+$, $\boldsymbol{\alpha} \in I^+$
- (ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive,
- (iii) F is increasing in α and decreasing in θ , F is log-concave in θ and $\frac{\partial F/\partial \theta}{F}$ is increasing in α ,

then $\boldsymbol{\theta} \succeq^w \boldsymbol{\sigma}$ implies $X_{n:n} \geq_{st} Y_{n:n}$.

Combining Theorem 3.1 and 3.2 a more general result is achieved as follows.

Theorem 3.3. *Let $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ and $\mathbf{Y} \sim F(\boldsymbol{\mu}, \boldsymbol{\sigma}, \psi_2)$, Suppose that*

- (i) $\boldsymbol{\theta}, \boldsymbol{\sigma} \in D^+$, $\boldsymbol{\alpha}, \boldsymbol{\mu} \in I^+$
- (ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive,
- (iii) F is increasing in α , decreasing in θ , F is log-concave in θ and α , $\frac{\partial F/\partial \alpha}{F}$ is increasing in θ , $\frac{\partial F/\partial \theta}{F}$ is increasing in α ,

then $\boldsymbol{\theta} \succeq^w \boldsymbol{\sigma}$ and $\boldsymbol{\alpha} \succeq_w \boldsymbol{\mu}$ imply $X_{n:n} \geq_{st} Y_{n:n}$.

3.2 Minimum of sample

In this subsection, we carry out stochastic comparison of smallest order statistics from dependent and heterogeneous observations from a general two parameters distribution, with the heterogeneity in the model parameters.

Theorem 3.4. *Let $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ and $\mathbf{Y} \sim F(\boldsymbol{\mu}, \boldsymbol{\theta}, \psi_2)$, Suppose that*

- (i) $\boldsymbol{\alpha}, \boldsymbol{\mu} \in I^+$, $\boldsymbol{\theta} \in D^+$
- (ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive,
- (iii) F is increasing in α and decreasing in θ , $1 - F$ is log-concave in α and $\frac{\partial F / \partial \alpha}{1 - F}$ is decreasing in θ ,

then $\boldsymbol{\alpha} \succeq_w \boldsymbol{\mu}$ implies $X_{1:n} \leq_{st} Y_{1:n}$.

Theorem 3.5. *Let $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ and $\mathbf{Y} \sim F(\boldsymbol{\alpha}, \boldsymbol{\sigma}, \psi_2)$, Suppose that*

- (i) $\boldsymbol{\theta}, \boldsymbol{\sigma} \in D^+$, $\boldsymbol{\alpha} \in I^+$
- (ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive
- (iii) F is increasing in α and decreasing in θ , $1 - F$ is log-concave in θ and $\frac{\partial F / \partial \theta}{1 - F}$ is decreasing in α ,

then $\boldsymbol{\theta} \succeq^w \boldsymbol{\sigma}$ implies $X_{1:n} \leq_{st} Y_{1:n}$.

Theorem 3.6. *Let $\mathbf{X} \sim F(\boldsymbol{\alpha}, \boldsymbol{\theta}, \psi_1)$ and $\mathbf{Y} \sim F(\boldsymbol{\mu}, \boldsymbol{\sigma}, \psi_2)$, Suppose that*

- (i) $\boldsymbol{\theta}, \boldsymbol{\sigma} \in D^+$, $\boldsymbol{\alpha}, \boldsymbol{\mu} \in I^+$
- (ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive
- (iii) F is increasing in α and decreasing in θ , $1 - F$ is log-concave in θ and α , $\frac{\partial F / \partial \alpha}{1 - F}$ is decreasing in θ , and $\frac{\partial F / \partial \theta}{1 - F}$ is decreasing in α ,

then $\boldsymbol{\theta} \succeq^w \boldsymbol{\sigma}$ and $\boldsymbol{\alpha} \succeq_w \boldsymbol{\mu}$ imply $X_{1:n} \leq_{st} Y_{1:n}$.

4 Special case

In this section, we present a case that satisfy the distributional conditions of the general model discussed in the previous section. Wang et al. [8] carried out stochastic comparisons of lifetimes of series and parallel systems with dependent heterogeneous Type II half logistic-resilience scale (TIIHL-RS) with cdf $F(x; \alpha, \theta) = \frac{2G^\theta(\alpha x)}{1 + G^\theta(\alpha x)}$, $x \geq 0, \alpha > 0, \theta > 0$, where G is the baseline cdf with g as the corresponding baseline density function. Here, we investigate the conditions in Theorem 3.6 for this two parameters distribution. First, we see that

$$\frac{\partial F(x; \alpha, \theta)}{\partial \theta} = \frac{2G^\theta(\alpha x) \ln(G(\alpha x))}{(1 + G^\theta(\alpha x))^2} \leq 0, \quad \frac{\partial F(x; \alpha, \theta)}{\partial \alpha} = \frac{2\theta x g(\alpha x) G^{\theta-1}(\alpha x)}{(1 + G^\theta(\alpha x))^2} \geq 0,$$

hence $F(x; \alpha, \theta)$ is increasing in α and decreasing in θ for each $x > 0$. On the other hand, we have

$$\frac{\partial \ln(F(x; \alpha, \theta))}{\partial \alpha} = x \frac{\theta}{1 + G^\theta(\alpha x)} r(\alpha x) \quad (4.1)$$

where, $r(\alpha x) = \frac{g(\alpha x)}{G(\alpha x)}$. Assuming that the function $r(\cdot)$ decreases, it can be deduced that (4.1) is decreasing in α , i.e. F is log-concave in α . Also, it can be easily shown that (4.1) is increasing in θ .

Also, we have

$$\frac{\partial \ln(F(x; \alpha, \theta))}{\partial \theta} = \frac{\ln G(\alpha x)}{G^\theta(\alpha x)} \quad (4.2)$$

which can be shown easily that it is decreasing in θ and increasing in α .

Hence, the set of conditions (iii) in Theorems 3.1 and 3.2 are satisfied and we get the following proposition, which coincides with the obtained result in [8].

Theorem 4.1. *Let $\mathbf{X} \sim TIIHL - RS(\alpha, \theta, \psi_1)$ and $\mathbf{Y} \sim TIIHL - RS(\mu, \sigma, \psi_2)$. Suppose that*

(i) $\theta, \sigma \in D^+$, $\alpha, \mu \in I^+$

(ii) ψ_1 or ψ_2 is log-convex and $\phi_2 \circ \psi_1$ is super-additive,

(iii) $r(y) = \frac{g(y)}{G(y)}$ is decreasing in y ,

then $\theta \succeq^w \sigma$ and $\alpha \succeq_w \mu$ imply $X_{n:n} \geq_{st} Y_{n:n}$.

References

- [1] Barmalzan, G., Ayat, S. M., Balakrishnan, N., & Roozegar, R. (2020). Stochastic comparisons of series and parallel systems with dependent heterogeneous extended exponential components under Archimedean copula. *Journal of computational and applied mathematics*, 380, 112965.
- [2] Chowdhury, S., Kundu, A., & Mishra, S. K. (2023). Ordering properties of the smallest order statistic from Weibull-G random variables. *Communications in Statistics-Theory and Methods*, 52(18), 6525-6541.
- [3] Das, S., & Kayal, S. (2020). Ordering extremes of exponentiated location-scale models with dependent and heterogeneous random samples. *Metrika*, 83(8), 869-893.
- [4] Fang, R., Li, C., & Li, X. (2018). Ordering results on extremes of scaled random variables with dependence and proportional hazards. *Statistics*, 52(2), 458-478.
- [5] Nadeb, H., Torabi, H., & 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.

- [6] Samanta, R. J., Das, S., & Balakrishnan, N. (2023). Orderings of extremes among dependent extended Weibull random variables. *Probability in the Engineering and Informational Sciences*, 1-28.
- [7] Shrahili, M. (2024). Stochastic ordering results on extreme order statistics from dependent samples with Archimedean copula. *Journal of Inequalities and Applications*, 2024(1), 120.
- [8] Wang, J., Yan, R., & Lu, B. (2020). Stochastic comparisons of parallel and series systems with type II half logistic-resilience scale components. *Mathematics*, 8(4), 470.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Best Unbiased Prediction of Order Statistics in Generalized Extreme Value Distribution

Almasi, I. ¹ Esmailzadeh, N. ²

Department of Statistics, Faculty of Basic Science, Razi University, Kermanshah, Iran.

Abstract

Assume that some order statistics from a Generalized Extreme value distribution are missing. We aim to impute the missing values. We compute the Best Unbiased Prediction (BUP) of these missing values. Additionally, we examine and compare the performance of the derived predictor with several alternatives through comprehensive simulations for various extreme value indexes, sample sizes, and the number of missing values. We observe that generally, the method BUP demonstrates better performance.

Keywords: Extreme value distribution, Missing order statistics, Extreme value index, Prediction

1 Introduction

Extreme distributions represent a fascinating family of distributions in the literature, encompassing Fréchet, Gumbel, and Weibull skew and heavy tail distributions. The Fréchet, Gumbel, and Weibull distribution functions can, in turn, be nested within a one-parameter family of distributions via the von Mises-Jenkinson parameterization [8, 19].

¹i.almasi@razi.ac.ir

²n.esmailzadeh@razi.ac.ir

Notably, the Generalized Extreme Value (GEV) distribution has a distribution function defined by

$$G_\gamma(x) = \begin{cases} \exp(-(1 + \gamma x)^{-1/\gamma}) & \gamma \neq 0, 1 + \gamma x > 0 \\ \exp(-e^{-x}) & \gamma = 0 \end{cases} \quad (1.1)$$

The parameter γ , known as extreme value index (EVI), determines the shape of the GEV distribution. If $\gamma > 0$, G_γ represents a heavy-tailed distribution and is said to be in the domain of attraction of the Fréchet distribution. If $\gamma = 0$, the Gumbel domain of attraction represents the simplest case for inference. Lastly, if $\gamma < 0$, then $x < -1/\gamma$, the distribution function is said to be in the domain of attraction of the Weibull distribution. This domain includes short-tailed distributions with a finite right endpoint.

The reality of data collection often involves incomplete observations due to censoring, a phenomenon where the precise value of an observation remains unknown. Order statistics, representing the ordered values within a sample, offer a powerful analytical tool for such data. The versatility of order statistics in analyzing extreme value distributions with censoring is evident in its widespread applications across diverse fields. In reliability analysis, order statistics are fundamental for estimating lifetime distributions and predicting component failures. The analysis of failure times often involves censoring, making order statistics indispensable for accurate estimation. In hydrology, order statistics are used to model extreme rainfall events, design flood defenses, and assess the risk of flooding [19, 21]. Long-term rainfall data are often incomplete due to missing observations or limitations in historical records, necessitating the use of methods that handle censoring effectively. Financial applications involve modeling extreme market fluctuations, estimating Value at Risk, and managing financial risk [21]. Financial data often exhibit heavy tails and may contain censored observations, requiring the use of robust methods based on order statistics. Environmental science uses order statistics to analyze air pollution levels, assess environmental risks, and predict extreme weather events [21]. Environmental data often involve censoring due to limitations in monitoring equipment or the infrequent occurrence of extreme events. Medical research applies order statistics to model survival times, analyze censored data in clinical trials, and assess treatment effectiveness [7]. Clinical trial data frequently involve censoring due to patients dropping out of the study or the study ending before all patients experience the event of interest.

Let $X_1, \dots, X_n$ are random variables of GEV with a parameter γ , and let $Y_1 \leq Y_2 \leq \dots \leq Y_n$ be the corresponding order statistics. Assume that some of the order statistics values are lost (or censored), that is we only observed the data set $\mathbf{Y} = (Y_1, \dots, Y_m)$ that $0 \leq m \leq n$. The problem of predicting the censored order statistics and interpolate the missing order statistics are fundamental interest. A few methods have been proposed for interpolation of order statistics. Razmkhah [20] compared some interpolation (prediction or reconstruction) methods of missing order statistics. Khatib *et al.* [10] developed Bayesian interpolation under the assumption that the underlying distribution is exponential and that scale parameter has a conjugate prior. Best unbiased prediction of missing order statistics of a stable distribution, based on conditional expected value has been studied in [2, 3].

In this paper, we address the issue of predicting missing order statistics in the generalized extreme value distribution. We use Best Unbiased Prediction (Interpolation or reconstruction) (BUP) to generate imputed values for missing orders. For Fréchet, Gumbel and Weibull, we derive the BUP. Additionally, we examine the Conditional Median Prediction (CMP), which has been explored by several authors, including Raqab [17, 16, 1]. We also consider the Conditional Average Prediction (CAP) studied by Razmkhah *et al.* [20] and Asgharzadeh *et al.* [5].

We also examine the performance of derived predictor in three methods through comprehensive simulation for various extreme value indexes, sample sizes and the number of missing values.

The structure of the paper is outlined as follows. In Section 2, necessary and sufficient conditions for the existence of conditional moments of order statistics of GEV random variables are introduced. In Section 3, using the results obtained in Section 2, we attempt to interpolate the missed order statistics. Finally, in Section 4, we examine the performance of methods mentioned in section 3 through comprehensive simulation.

2 Existence conditional of order statistics Moments

Let $Y_i, i = 1, \dots, n$ be order statistics from a GEV distribution with a common probability density function (pdf) and cumulative distribution function (cdf), denoted as $g_\gamma(x)$ and $G_\gamma(x)$, respectively.

For an Extreme random variable $X \sim G_\gamma(x)$ with $\gamma \in R$, if $\gamma > 0$, moments $E|X|^k$ of order $k > 1/\gamma$ are not finite (See [15]). Formally, heavy-tailed distributions are those whose tail probability, $1 - G_{\gamma>0}(x)$, is larger than that of an exponential distribution, (See [18]). Thus, noting that if $\gamma > 0$,

$$\lim_{\lambda \rightarrow \infty} \lambda^{1/\gamma} P(X > \lambda) = \gamma^{-1/\gamma}.$$

Lemma 2.1. *Let $X_1, \dots, X_n$ be random variables of $G_\gamma(x)$ with the extreme value index, $\gamma > 0$. Let $Y_1 \leq Y_2 \leq \dots \leq Y_n$ be the corresponding order statistics, then*

$$\lim_{\lambda \rightarrow \infty} \lambda^{(n-t+1)/\gamma} P(Y_t > \lambda) = \binom{n}{n-t+1} \gamma^{-(n-t+1)/\gamma} \quad (2.1)$$

Proof. It can be show that

$$P(Y_t > \lambda) = \sum_{i=n-t+1}^n (-1)^{i-(n-t+1)} \binom{i-1}{n-t} \sum_{1 \leq h_1 < \dots < h_i \leq n} P(X_{h_1} > \lambda, \dots, X_{h_i} > \lambda).$$

Refer to [22] for the proof.

From independence and regularly varying property, equations (1), we have

$$\begin{aligned} \sum_{1 \leq h_1 < \dots < h_i \leq n} P(X_{h_1} > \lambda, \dots, X_{h_i} > \lambda) &= \sum_{1 \leq h_1 < \dots < h_i \leq n} \prod_{j=1}^i P(X_{h_j} > \lambda) \\ &= \binom{n}{i} (P(X > \lambda))^i. \end{aligned}$$

Therefore,

$$P(Y_t > \lambda) = \sum_{i=n-t+1}^n (-1)^{i-(n-t+1)} \binom{i-1}{n-t} \binom{n}{i} (P(X > \lambda))^i$$

So,

$$\lim_{\lambda \rightarrow \infty} \lambda^{(n-t+1)/\gamma} P(Y_t > \lambda) = \sum_{i=n-t+1}^n \left[(-1)^{i-(n-t+1)} \binom{i-1}{n-t} \binom{n}{i} \lim_{\lambda \rightarrow \infty} \lambda^{(n-t+1)/\gamma} (P(X > \lambda))^i \right].$$

Thus

$$\lim_{\lambda \rightarrow \infty} \lambda^{(n-t+1)/\gamma} P(Y_t > \lambda) = \begin{cases} \binom{n}{n-t+1} \gamma^{-(n-t+1)/\gamma} & i = n - t + 1 \\ 0 & i > n - t + 1 \end{cases}$$

□

Suppose that we have observed the data set $\mathbf{Y} = (Y_1, \dots, Y_m)$, where $m < n$. A necessary and sufficient condition for the existence of order statistics moments is introduced in the following lemma.

Lemma 2.2. *Let k be a non-negative number, Y_m and Y_t , $m, t \in \{1, \dots, n\}$, $m < t$ be order statistics from GEV with the extreme value index, $\gamma > 0$. Then $E(Y_t^k | y_1, \dots, y_m)$ is finite if and only if $t < n + 1 - k\gamma$.*

Proof. Using the Markov property of order statistics (Theorem 2.4.3 page 24, Arnold et al. 1992), it is enough to show that $E(|Y_t|^k | y_m) < +\infty$. Now, we have

$$\begin{aligned} E(|Y_t|^k | y_m) &= \int_0^{+\infty} P(|Y_t|^k > \lambda | y_m) d\lambda \\ &= \int_0^\epsilon P(|Y_t|^k > \lambda | y_m) d\lambda + \int_\epsilon^{+\infty} P(|Y_t|^k > \lambda | y_m) d\lambda \\ &= k \int_0^{\epsilon^{1/k}} u^{k-1} P(|Y_t| > u | y_m) du \\ &\quad + k \int_{\epsilon^{1/k}}^{+\infty} u^{k-1} P(|Y_t| > u | y_m) du \quad (\text{by the change of variable } u = \lambda^{1/k}), \end{aligned}$$

where $\epsilon > 0$. The first term is finite. Using (2.1) and Theorem 2.4.1 in [6], second term is finite if and only if

$$\int_{\epsilon^{1/k}}^{+\infty} u^{k-1} u^{-(n-t+1)/\gamma} du < +\infty,$$

which is finite if and only if $t < n + 1 - \gamma k$.

□

3 Prediction in Generalized Extreme Value distribution

In this section, we described three methods for prediction of missed order statistics. Then, in the subsequent subsections, the predictions are derived for Fréchet, Gumbel and Weibull distribution.

Conditional Median Prediction (CMP) approach has been studied by several authors, including [17, 16, 1]. We denote the CMP by $\hat{Y}_{t,\text{CMP}}$. We say that $\hat{Y}_{t,\text{CMP}}$ is CMP of Y_t if the condition

$$P(Y_t \leq \hat{Y}_{t,\text{CMP}} | Y_1, \dots, Y_m) = P(Y_t \geq \hat{Y}_{t,\text{CMP}} | Y_1, \dots, Y_m)$$

is satisfied. The CMP can be computed as follows,

$$\hat{Y}_{t,\text{CMP}} = G_\gamma^{-1}\{G_\gamma(Y_m) + \text{med}(V)(1 - G_\gamma(Y_m))\}, \quad (3.1)$$

where $V \sim \text{Beta}(t - m, n + 1 - t)$ and $\text{med}(V)$ represents the median of V .

The Conditional Average Prediction (CAP) of Y_t has been studied by Razmkhah *et al.* (2010) and Asgharzadeh *et al.* (2012). In Eq.(3.1), they use the average value instead of median. The CAP is defined and computed as follows,

$$\hat{Y}_{t,\text{CAP}} = G_\gamma^{-1}\left\{\frac{t - m}{n - m + 1} + \frac{n - t + 1}{n - m + 1}G_\gamma(y_m)\right\}. \quad (3.2)$$

It is logical to consider $E(Y_t | y_1, \dots, y_m)$ as a predictor of Y_t having observed $y_1, \dots, y_m$. Let us recall that the conditional density function of y_t given y_m takes the following form

$$g_\gamma(y_t | y_m) = \frac{(1 - G_\gamma(y_m))^{m-n}}{\text{Beta}(n_1 + 1, n_2 + 1)} \left(G_\gamma(y_t) - G_\gamma(y_m)\right)^{n_1} \left(1 - G_\gamma(y_t)\right)^{n_2} g_\gamma(y_t), \quad (3.3)$$

for $y_t \geq y_m$, where $n_1 = t - m - 1$, $n_2 = n - t$.

By the change of variable $V = \frac{G_\gamma(Y_t) - G_\gamma(y_m)}{1 - G_\gamma(y_m)}$ and using (3.3), the Best Unbiased Prediction (BUP) of t 'th order statistic ($t = m + 1, \dots, n$), if it exists, is defined and computed as follows,

$$\begin{aligned} \hat{Y}_t &= E(Y_t | y_1, \dots, y_m) \\ &= E\left(G_\gamma^{-1}\left(G_\gamma(y_m) + V(1 - G_\gamma(y_m))\right)\right) \end{aligned} \quad (3.4)$$

where $V \sim \text{Beta}(n_1 + 1, n_2 + 1)$,

$$\hat{Y}_{t,\text{BUP}} = \frac{(1 - G_\gamma(y_m))^{m-n}}{\text{Beta}(n_1 + 1, n_2 + 1)} \sum_{i=0}^{n_1} \sum_{j=0}^{n_2} \left[\binom{n_1}{i} \binom{n_2}{j} (-1)^{n_1+j-i} G_\gamma(y_m)^{n_1-i} \right. \quad (3.5)$$

$$\left. \int_{G_\gamma(y_m)}^1 G_\gamma^{-1}(u) u^{i+j} du \right]. \quad (3.6)$$

3.1 Prediction in Gumbel distribution

By applying the wellknown formula $(1 + \gamma x)^{1/\gamma} \rightarrow \exp(x)$ as $\gamma \rightarrow 0$, the Gumbel distribution function obtained as follows

$$G_{\gamma=0}(x) = \exp(-e^{-x}), \quad x \in R.$$

Using (3.4), the BUP of the t th order statistic is as follows,

$$\hat{Y}_{t,\text{BUP}} = y_m - E(\mathbf{Log}[1 - e^{y_m} \mathbf{Log}[V(e^{e^{-y_m}} - 1) + 1]])$$

where $V \sim \text{Beta}(t - m, n + 1 - t)$.

Furthermore, using (3.5), the BUP of the t th order statistic are as follows,

$$\hat{Y}_{t,\text{BUP}} = \frac{(1 - G_0(y_m))^{m-n}}{\text{Beta}(n_1 + 1, n_2 + 1)} \sum_{i=0}^{n_1} \sum_{j=0}^{n_2} \binom{n_1}{i} \binom{n_2}{j} (-1)^{n_1+j-i+1} G_0(y_m)^{n_1-i} \int_{G_0(y_m)}^1 \mathbf{Log}(-\mathbf{Log}(u)) u^{i+j} du$$

where the last part of above equation is

$$\frac{\text{Exp}(-y_m - (i + j + 1)e^{-e^{-y_m}}) - \gamma_0 - E_1((i + j + 1)e^{-e^{-y_m}}) - \text{Log}(i + j + 1))}{1 + i + j}$$

where $E_n(z) = -\int_1^\infty \frac{e^{-zt}}{t^n} dt$, gives the exponential integral function.

By Eq.(3.1), the CMP is computed as follows,

$$\hat{Y}_{t,\text{CMP}} = y_m - \mathbf{Log}\left(1 - e^{y_m} \mathbf{Log}\left(e^{e^{-y_m}} \frac{t - m - \frac{1}{3}}{n - m + \frac{1}{3}} + \frac{n - t + \frac{2}{3}}{n - m + \frac{1}{3}}\right)\right).$$

By Eq.(3.2), the CAP is defined and computed as follows,

$$\hat{Y}_{t,\text{CAP}} = y_m - \mathbf{Log}\left(1 - e^{y_m} \mathbf{Log}\left(e^{e^{-y_m}} \frac{t - m}{n - m + 1} + \frac{n - t + 1}{n - m + 1}\right)\right).$$

3.2 Prediction in Fréchet distribution

By applying the reparameterization $\alpha = 1/\gamma$ and the change of variable $z = 1 + \gamma x$ in (1.1), if $\gamma > 0$, then the random variable Z follows a standard Fréchet distribution, and its cdf is

$$G_\alpha(z) = \exp(-z^{-\alpha}), \quad z \geq 0, \quad \alpha > 0,$$

where α is the unknown shape parameter.

Using (3.5), the BUP of the t th order statistic are as follows,

$$\hat{Y}_{t,\text{BUP}} = \frac{(1 - G_\alpha(y_m))^{m-n}}{\text{Beta}(t - m, n - t + 1)} \sum_{i=0}^{n_1} \sum_{j=0}^{n_2} \binom{n_1}{i} \binom{n_2}{j} (-1)^{n_1-i+j} G_\alpha(y_m)^{n_1-i}$$

$$\int_{G_\alpha(y_m)}^1 \left(-\mathbf{Log}(u) \right)^{-1/\alpha} u^{i+j} du$$

where the last part of above equation if $\alpha \in N$ is

$$(1+i+j)^{1/\alpha-1} \Gamma(1-1/\alpha) - E_{1/\alpha}((1+i+j)y_m^{-\alpha}) y_m^{-\alpha+1}$$

where $E_n(z) = -\int_1^\infty \frac{e^{-zt}}{t^n} dt$.

By Eq.(3.1), the CMP is defined and computed as follows,

$$\hat{Y}_{t,\text{CMP}} = \left(y_m^{-\alpha} - \mathbf{Log}\left(\frac{t-m-\frac{1}{3}}{n-m+\frac{1}{3}} e^{y_m^{-\alpha}} + \frac{n-t+\frac{2}{3}}{n-m+\frac{1}{3}} \right) \right)^{-1/\alpha}.$$

By Eq.(3.2), the CAP is defined and computed as follows,

$$\hat{Y}_{t,\text{CAP}} = \left(y_m^{-\alpha} - \mathbf{Log}\left(\frac{t-m}{n-m+1} e^{y_m^{-\alpha}} + \frac{n-t+1}{n-m+1} \right) \right)^{-1/\alpha}.$$

3.3 Prediction in Weibull distribution

By applying the reparameterization $\alpha = -1/\gamma$ and the change of variable $z = -1 - \gamma x$ in (1.1), if $\gamma < 0$, then the random variable Z follows a standard Weibull distribution and its cdf is

$$G_\alpha(x) = \exp(-(-z)^\alpha), \quad z \leq 0, \quad \alpha > 0,$$

where α is the unknown shape parameter.

Using (3.5), the BUP interpolators of the t th order statistic are as follows,

$$\hat{Y}_{t,\text{BUP}} = \frac{(1 - G_{-\alpha}(y_m))^{m-n}}{\text{Beta}(t-m, n-t+1)} \sum_{i=0}^{n_1} \sum_{j=0}^{n_2} \binom{n_1}{i} \binom{n_2}{j} (-1)^{n_1+j-i+1} G_{-\alpha}(y_m)^{n_1-i}$$

$$\int_{G_{-\alpha}(y_m)}^1 \left(-\mathbf{Log}(u) \right)^{1/\alpha} u^{i+j} du$$

where the last part of above equation if $\alpha \in N$ is

$$(1+i+j)^{-1/\alpha-1} \left(\Gamma(1+a/\alpha) - \Gamma(1+a/\alpha, (i+j+1)(-ym)^\alpha) \right)$$

where $E_n(z) = -\int_1^\infty \frac{e^{-zt}}{t^n} dt$, where the principal value of the integral is taken.

By Eq.(3.1), the CMP is defined and computed as follows,

$$\hat{Y}_{t,\text{CMP}} = -\left((-y_m)^\alpha - \mathbf{Log}\left(\frac{t-m-\frac{1}{3}}{n-m+\frac{1}{3}} e^{(-y_m)^\alpha} + \frac{n-t+\frac{2}{3}}{n-m+\frac{1}{3}} \right) \right)^{1/\alpha}.$$

By Eq.(3.2), the CAP is defined and computed as follows,

$$\hat{Y}_{t,\text{CAP}} = -\left((-y_m)^\alpha - \mathbf{Log}\left(\frac{t-m}{n-m+1} e^{(-y_m)^\alpha} + \frac{n-t+1}{n-m+1} \right) \right)^{1/\alpha}.$$

4 Simulation Study

In this section, we examine the performance of three methods mentioned in section 3 through comprehensive simulation. We consider various settings for sample size n , the number of available samples m , and the parameter γ . The sample sizes we consider are $n = 10, 20, 50$ and 100 , with the corresponding values of $m = 7, 16, 45$ and 90 , respectively. The values of $\gamma = 0.25$ and 0.5 for Fréchet distribution, while $\gamma = -0.25$ and -0.5 apply to Weibull distribution. For each setting and distribution, we repeat the simulation 1000 times and calculate the mean and MSE of predictions. The results are presented in tables 1, 2 and 4. The scatter-plots of predictions versus true observations are presented in figures 1, 2 and 3, for Gumbel, Fréchet and Weibull distributions, respectively, where $n = 20$ and $m = 16$. To conserve space and because of the similarity in the structure of the graphs, we have not included them for other settings.

From the results, we observed that generally, the method BUP demonstrates better performance in almost all settings. As we expected, the predictions improve when t is closer to m .

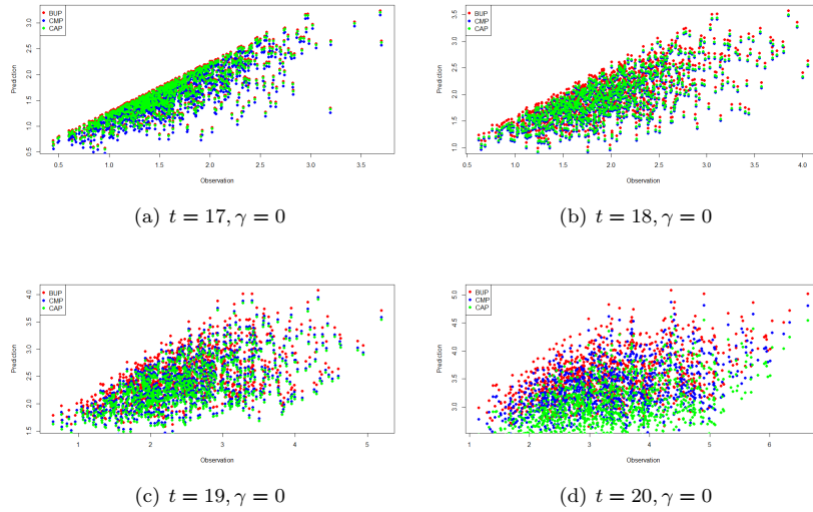


Figure 1: Plots of predictions vs. observations for Gumbel distribution, $n = 20, m = 16$

Table 1: Means and MSEs for Gumbel distribution

n = 10, m=7			n = 20, m=16			n = 50, m=45			n=100, m=90		
	Mean	MSE		Mean	MSE		Mean	MSE		Mean	MSE
y_8	1.276		y_{17}	1.679		y_{46}	2.364		y_{91}	2.301	
$\hat{Y}_{BUP}$	1.277	0.133		1.672	0.077		2.373	0.189		2.303	0.026
$\hat{Y}_{CMP}$	1.156	0.146		1.583	0.0862		2.287	0.054		2.267	0.012
$\hat{Y}_{CAP}$	1.234	0.133		1.646	0.078		2.338	0.049		2.298	0.011
y_9	1.824		y_{18}	2.048		y_{47}	2.657		y_{92}	2.392	
$\hat{Y}_{BUP}$	1.824	0.399		2.033	0.198		2.638	0.132		2.397	0.102
$\hat{Y}_{CMP}$	1.686	0.418		1.934	0.211		2.566	0.129		2.359	0.0249
$\hat{Y}_{CAP}$	1.686	0.418		1.962	0.205		2.598	0.124		2.385	0.024
y_{10}	2.940		y_{19}	2.581		y_{48}	3.007		y_{93}	2.483	
$\hat{Y}_{BUP}$	2.897	1.495		2.544	0.472		2.993	0.252		2.460	0.205
$\hat{Y}_{CMP}$	2.685	1.559		2.419	0.497		2.897	0.2446		2.436	0.049
$\hat{Y}_{CAP}$	2.452	1.732		2.377	0.512		2.896	0.244		2.457	0.047
			y_{20}	3.631		y_{49}	3.510		y_{94}	2.643	
$\hat{Y}_{BUP}$				3.572	1.440		3.493	0.469		2.641	0.260
$\hat{Y}_{CMP}$				3.367	1.507		3.378	0.477		2.586	0.065
$\hat{Y}_{CAP}$				3.099	1.720		3.312	0.499		2.600	0.064
						y_{50}	4.479		y_{95}	2.764	
$\hat{Y}_{BUP}$							4.493	1.636		2.754	0.324
$\hat{Y}_{CMP}$							4.272	1.489		2.729	0.0836
$\hat{Y}_{CAP}$							3.982	1.694		2.735	0.083
									y_{96}	2.960	
$\hat{Y}_{BUP}$										2.965	0.309
$\hat{Y}_{CMP}$										2.917	0.136
$\hat{Y}_{CAP}$										2.910	0.137
									y_{97}	3.240	
$\hat{Y}_{BUP}$										3.264	0.380
$\hat{Y}_{CMP}$										3.193	0.185
$\hat{Y}_{CAP}$										3.168	0.188
									y_{98}	3.639	
$\hat{Y}_{BUP}$										3.622	0.399
$\hat{Y}_{CMP}$										3.558	0.272
$\hat{Y}_{CAP}$										3.502	0.284
									y_{99}	4.096	
$\hat{Y}_{BUP}$										4.126	0.496
$\hat{Y}_{CMP}$										4.047	0.474
$\hat{Y}_{CAP}$										3.926	0.500
									y_{100}	4.954	
$\hat{Y}_{BUP}$										5.166	0.849
$\hat{Y}_{CMP}$										4.977	0.792
$\hat{Y}_{CAP}$										4.633	0.895

Table 2: Means and MSEs for Fréchet distribution

n=10, m=7, $\gamma = 0.25$			n=10, m=7, $\gamma = 0.5$			n=20, m=16, $\gamma = 0.25$			n=20, m=16, $\gamma = 0.5$		
	Mean	MSE		Mean	MSE		Mean	MSE		Mean	MSE
y_8	1.395		y_8	2.004		y_{17}	1.526		y_{17}	2.372	
$\hat{Y}_{BUP}$	1.393	0.021		1.996	0.222		1.532	0.011		2.389	0.125
$\hat{Y}_{CMP}$	1.510	0.034		1.857	0.253		1.612	0.018		2.238	0.1424
$\hat{Y}_{CAP}$	1.582	0.056		1.927	0.238		1.674	0.033		2.298	0.130
y_9	1.634		y_9	2.826		y_{18}	1.681		y_{18}	2.895	
$\hat{Y}_{BUP}$	1.615	0.107		2.740	2.373		1.683	0.038		2.907	0.531
$\hat{Y}_{CMP}$	2.024	0.262		2.360	2.720		1.959	0.116		2.573	0.656
$\hat{Y}_{CAP}$	2.024	0.262		2.360	2.720		1.987	0.132		2.600	0.640
y_{10}	2.204		y_{10}	5.800		y_{19}	1.924		y_{19}	3.860	
$\hat{Y}_{BUP}$	2.183	0.851		5.629	66.225		1.936	0.121		3.928	2.780
$\hat{Y}_{CMP}$	2.981	1.471		3.307	73.941		2.457	0.407		3.0597	3.507
$\hat{Y}_{CAP}$	2.751	1.167		3.079	75.127		2.416	0.364		3.0195	3.573
						y_{20}	2.571		y_{20}	7.687	
$\hat{Y}_{BUP}$							2.598	1.031		7.958	88.574
$\hat{Y}_{CMP}$							3.400	1.718		3.991	102.649
$\hat{Y}_{CAP}$							3.132	1.347		3.726	104.68

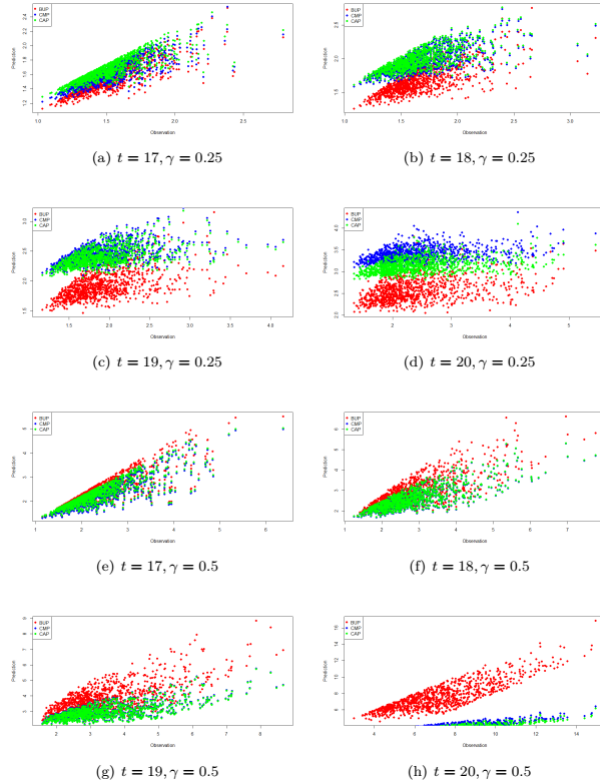

 Figure 2: Plots of predictions vs. observations for Fréchet distribution, $n = 20, m = 16$

Table 3: Means and MSEs for Fréchet distribution

n=50, m=45, $\gamma = 0.25$			n=50, m=45, $\gamma = 0.5$			n=100, m=90, $\gamma = 0.25$			n=100, m=90, $\gamma = 0.5$		
	Mean	MSE		Mean	MSE		Mean	MSE		Mean	MSE
y_{46}	3.352		y_{46}	1.817		y_{91}	1.781		y_{91}	3.128	
$\hat{Y}_{BUP}$	3.355	0.140		1.818	0.009		1.799	0.045		3.119	0.115
$\hat{Y}_{CMP}$	3.140	0.190		1.868	0.012		1.806	0.003		3.028	0.044
$\hat{Y}_{CAP}$	3.191	0.171		1.920	0.020		1.837	0.006		3.058	0.039
y_{47}	3.843		y_{47}	1.942		y_{92}	1.820		y_{92}	3.200	
$\hat{Y}_{BUP}$	3.855	0.495		1.944	0.027		1.817	0.078		3.206	0.289
$\hat{Y}_{CMP}$	3.387	0.729		2.126	0.062		1.911	0.014		3.037	0.104
$\hat{Y}_{CAP}$	3.419	0.702		2.159	0.075		1.937	0.020		3.062	0.096
y_{48}	4.611		y_{48}	2.118		y_{93}	1.865		y_{93}	3.328	
$\hat{Y}_{BUP}$	4.650	1.756		2.127	0.074		1.863	0.173		3.343	0.443
$\hat{Y}_{CMP}$	3.711	2.651		2.462	0.194		2.028	0.037		3.089	0.197
$\hat{Y}_{CAP}$	3.711	2.651		2.462	0.194		2.049	0.045		3.109	0.187
y_{49}	6.029		y_{49}	2.404		y_{94}	1.919		y_{94}	3.491	
$\hat{Y}_{BUP}$	6.233	6.539		2.437	0.193		1.923	0.239		3.491	0.599
$\hat{Y}_{CMP}$	4.187	10.137		2.948	0.492		2.166	0.077		3.174	0.308
$\hat{Y}_{CAP}$	4.122	10.381		2.882	0.424		2.180	0.084		3.188	0.300
y_{50}	12.096		y_{50}	3.179		y_{95}	1.996		y_{95}	3.762	
$\hat{Y}_{BUP}$	12.529	643.362		3.257	1.902		1.986	0.258		3.835	0.732
$\hat{Y}_{CMP}$	5.109	692.509		3.881	2.403		2.330	0.137		3.315	0.515
$\hat{Y}_{CAP}$	4.820	696.628		3.589	2.079		2.336	0.141		3.321	0.510
						y_{96}	2.105		y_{96}	4.159	
$\hat{Y}_{BUP}$							2.115	0.275		4.218	0.748
$\hat{Y}_{CMP}$							2.534	0.224		3.518	0.934
$\hat{Y}_{CAP}$							2.528	0.219		3.512	0.943
						y_{97}	2.259		y_{97}	4.755	
$\hat{Y}_{BUP}$							2.281	0.269		4.947	0.957
$\hat{Y}_{CMP}$							2.791	0.358		3.796	1.843
$\hat{Y}_{CAP}$							2.766	0.332		3.771	1.892
						y_{89}	2.483		y_{98}	5.696	
$\hat{Y}_{BUP}$							2.515	0.269		6.026	1.209
$\hat{Y}_{CMP}$							3.131	0.563		4.170	3.793
$\hat{Y}_{CAP}$							3.075	0.493		4.114	3.966
						y_{99}	2.840		y_{99}	7.967	
$\hat{Y}_{BUP}$							2.870	0.336		8.301	1.302
$\hat{Y}_{CMP}$							3.620	0.891		4.724	13.502
$\hat{Y}_{CAP}$							3.499	0.716		4.604	14.297
						y_{100}	3.535		y_{100}	17.124	
$\hat{Y}_{BUP}$							3.844	0.853		17.257	1.389
$\hat{Y}_{CMP}$							4.556	1.793		5.758	136.514
$\hat{Y}_{CAP}$							4.211	1.208		5.414	144.443

Table 4: Means and MSEs for Weibull distribution

n=10, m=7, $\gamma = -0.25$			n=10, m=7, $\gamma = -0.5$			n=20, m=16, $\gamma = -0.25$			n=20, m=16, $\gamma = -0.5$		
	Mean	MSE		Mean	MSE		Mean	MSE		Mean	MSE
y_8	-0.5561		y_8	-0.7385		y_{17}	-0.6661		y_{17}	-0.4507	
$\hat{Y}_{BUP}$	-0.5535	0.0081		-0.7363	0.0038		-0.6685	0.0019		-0.4538	0.0034
$\hat{Y}_{CMP}$	-0.5785	0.0087		-0.7556	0.0041		-0.6820	0.0022		-0.4705	0.0038
$\hat{Y}_{CAP}$	-0.5563	0.0081		-0.7410	0.0038		-0.6713	0.0020		-0.4558	0.0034
y_9	-0.4424		y_9	-0.6537		y_{18}	-0.6078		y_{18}	-0.3774	
$\hat{Y}_{BUP}$	-0.4303	0.0147		-0.6451	0.0091		-0.6124	0.0042		-0.3835	0.0059
$\hat{Y}_{CMP}$	-0.4399	0.0146		-0.6595	0.0091		-0.6243	0.0044		-0.3940	0.0061
$\hat{Y}_{CAP}$	-0.4399	0.0146		-0.6595	0.0091		-0.6200	0.0043		-0.3886	0.0060
y_{10}	-0.2881		y_{10}	-0.5159		y_{19}	-0.5329		y_{19}	-0.2947	
$\hat{Y}_{BUP}$	-0.2794	0.0203		-0.5093	0.0189		-0.5411	0.0080		-0.3026	0.0086
$\hat{Y}_{CMP}$	-0.2691	0.0207		-0.5161	0.0188		-0.5506	0.0082		-0.3062	0.0087
$\hat{Y}_{CAP}$	-0.3025	0.0204		-0.5471	0.0198		-0.5563	0.0085		-0.3127	0.0089
						y_{20}	-0.4229		y_{20}	-0.1936	
$\hat{Y}_{BUP}$							-0.4298	0.0127		-0.1992	0.0091
$\hat{Y}_{CMP}$							-0.4345	0.0128		-0.1906	0.0090
$\hat{Y}_{CAP}$							-0.4647	0.0143		-0.2180	0.0096

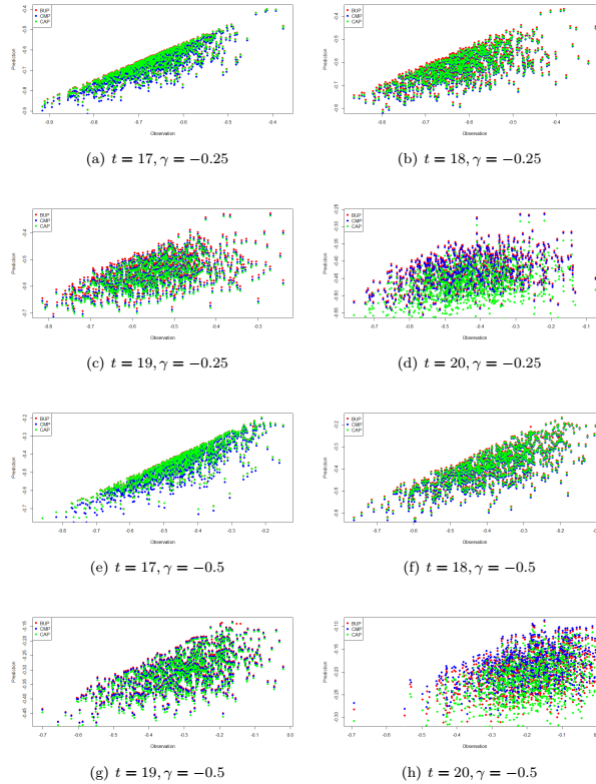
Figure 3: Plots of predictions vs. observations for Weibull distribution, $n = 20, m = 16$

Table 5: Means and MSEs for Weibull distribution

n=50, m=45, $\gamma = -0.25$			n=50, m=45, $\gamma = -0.5$			n=100, m=90, $\gamma = -0.25$			n=100, m=90, $\gamma = -0.5$		
	Mean	MSE		Mean	MSE		Mean	MSE		Mean	MSE
y_{46}	-0.560		y_{46}	-0.318		y_{91}	-0.564		y_{91}	-0.314	
$\hat{Y}_{BUP}$	-0.561	0.001		-0.312	0.006		-0.563	0.002		-0.314	0.001
$\hat{Y}_{CMP}$	-0.570	0.001		-0.329	0.001		-0.569			-0.319	
$\hat{Y}_{CAP}$	-0.563	0.001		-0.320	0.001		-0.565			-0.315	
y_{47}	-0.526		y_{47}	-0.282		y_{92}	-0.548		y_{92}	-0.299	
$\hat{Y}_{BUP}$	-0.527	0.004		-0.309	0.078		-0.548	0.006		-0.293	0.019
$\hat{Y}_{CMP}$	-0.535	0.002		-0.289	0.002		-0.553	0.001		-0.302	
$\hat{Y}_{CAP}$	-0.531	0.002		-0.285	0.002		-0.550			-0.298	
y_{48}	-0.487		y_{48}	-0.243		y_{93}	-0.533		y_{93}	-0.282	
$\hat{Y}_{BUP}$	-0.485	0.007		-0.194	0.142		-0.535	0.023		-0.297	0.037
$\hat{Y}_{CMP}$	-0.493	0.003		-0.245	0.003		-0.536	0.001		-0.284	0.001
$\hat{Y}_{CAP}$	-0.493	0.003		-0.245	0.003		-0.534	0.001		-0.281	0.001
y_{49}	-0.429		y_{49}	-0.192		y_{94}	-0.515		y_{94}	-0.264	
$\hat{Y}_{BUP}$	-0.430	0.007		-0.216	0.069		-0.509	0.041		-0.243	0.036
$\hat{Y}_{CMP}$	-0.437	0.005		-0.193	0.004		-0.518	0.001		-0.265	0.001
$\hat{Y}_{CAP}$	-0.444	0.006		-0.199	0.004		-0.516	0.001		-0.263	0.001
y_{50}	-0.343		y_{50}	-0.128		y_{95}	-0.494		y_{95}	-0.245	
$\hat{Y}_{BUP}$	-0.343	0.009		-0.121	0.008		-0.501	0.035		-0.256	0.042
$\hat{Y}_{CMP}$	-0.346	0.009		-0.121	0.004		-0.497	0.001		-0.244	0.001
$\hat{Y}_{CAP}$	-0.372	0.009		-0.140	0.004		-0.496	0.001		-0.244	0.001
						y_{96}	-0.469		y_{96}	-0.222	
$\hat{Y}_{BUP}$							-0.465	0.036		-0.196	0.033
$\hat{Y}_{CMP}$							-0.472	0.002		-0.221	0.001
$\hat{Y}_{CAP}$							-0.473	0.002		-0.222	0.001
						y_{97}	-0.441		y_{97}	-0.196	
$\hat{Y}_{BUP}$							-0.440	0.034		-0.204	0.036
$\hat{Y}_{CMP}$							-0.444	0.002		-0.195	0.002
$\hat{Y}_{CAP}$							-0.447	0.002		-0.197	0.002
						y_{98}	-0.405		y_{98}	-0.167	
$\hat{Y}_{BUP}$							-0.399	0.029		-0.150	0.017
$\hat{Y}_{CMP}$							-0.409	0.003		-0.165	0.002
$\hat{Y}_{CAP}$							-0.415	0.003		-0.170	0.002
						y_{99}	-0.356		y_{99}	-0.132	
$\hat{Y}_{BUP}$							-0.355	0.013		-0.132	0.017
$\hat{Y}_{CMP}$							-0.363	0.004		-0.130	0.002
$\hat{Y}_{CAP}$							-0.374	0.005		-0.138	0.002
						y_{100}	-0.284		y_{100}	-0.087	
$\hat{Y}_{BUP}$							-0.287	0.010		-0.087	0.002
$\hat{Y}_{CMP}$							-0.288	0.007		-0.082	0.002
$\hat{Y}_{CAP}$							-0.314	0.007		-0.097	0.002

References

- [1] Ahmadi, J., Jozani, M. J., Marchand, É., Parsian, A. (2009). Prediction of k -records from a general class of distributions under balanced type loss functions, *Metrika*, 70(1):19-33.
- [2] Almasi, I., Mohammadpour, A., and Mohammadi, M. (2017). Best linear unbiased interpolation of order statistics. *Communications in Statistics-Simulation and Computation* 46 (5), 4161-4171.
- [3] Almasi, I., Mohammadpour, A., and Mohammadi, M. (2018) Best Unbiased Prediction of Order Statistics in Stable Distributions, *Communications in Statistics-Simulation and Computation* 47 (8), 2424-2435.
- [4] Arnold, B. C., Balakrishnan, N., Nagaraja, H. N. (1992) *A first course in order statistics*, vol. 54, Siam.
- [5] Asgharzadeh, A., Ahmadi, J., Mirzazadeh, G, Z., Valiollahi, R. (2012). Reconstruction of the past failure times for the proportional reversed hazard rate model, *Journal of Statistical Computation and Simulation*. 82(3):475-489.
- [6] Balakrishnan, N., A. C. Cohen. (1991). *Order Statistics and Inference: Estimation Methods*, Academic Press, San Diego.
- [7] Ballenberger N, Lluís A, von Mutius E, Illi S, Schaub B. Novel statistical approaches for non-normal censored immunological data: analysis of cytokine and gene expression data. *PLoS One*. 2012;7(10):e46423. doi: 10.1371/journal.pone.0046423. Epub 2012 Oct 26. PMID: 23110049; PMCID: PMC3482200.
- [8] Jenkinson, A.F.(1955). The frequency distribution of the annual maximum (or minimum) values of meteorological elements. *Quarterly Journal of the Royal Meteorological Society*, 81, 158-171.
- [9] Kerman, J. (2011). A closed-form approximation for the median of the beta distribution. arXiv:1111.0433v1.
- [10] Khatib, B., Ahmadi, J., Razmkhah, M. (2011). Bayesian reconstruction of the missing failure times in exponential distribution, *Journal of Statistical Computation and Simulation* . 83(3):501-517.
- [11] Lih, J. L. (2004). Bivariate Ranked Set Sampling. A Dissertation in Applied Statistics, University of California, Riverside, CA.
- [12] Liu, Huashuai, Yang, Fan, and Wang, Hongchuan. (2023). "Research on Threshold Selection Method in Wave Extreme Value Analysis". *Water*. <https://doi.org/10.3390/w15203648>
- [13] Mohammadi, M., Mohammadpour, A. (2014). Estimating the parameters of an α -stable distribution using the existence of moments of order statistics. *Statistics and Probability Letters* 90:78-84.

- [14] Nolan, J. P. (2001). Maximum likelihood estimation of stable parameters, *Lévy processes: Theory and applications*, 379-400.
- [15] Haande Haan, L., Ferreira, A. (2006). Extreme Value Theory: An Introduction. Springer.
- [16] Raqab, M. Z., Ahmadi, J., Doostparast, M. (2007). Statistical inference based on record data from pareto model, *Statistics*. 41(2):105-118.
- [17] Raqab, M. Z., Nagaraja, H. N. (1995). On some predictors of future order statistics, *Metron*. 53(1-2):185-204.
- [18] Maria Jacob Cludia Neves, Danica Vukadinovi and Greetham, (2020). Forecasting and Assessing Risk of Individual Electricity Peaks. pp, 39-60.
- [19] Misesvon Mises, R.(1936). La distribution de la plus grande de n valeurs. Rev. math. Union interbalcanique 1(1).
- [20] Razmkhah, M., Khatib, B., Ahmadi, J. (2010). Reconstruction of order statistics in exponential distribution, *Journal of the Iranian Statistical Society*. 9:21-40.
- [21] Raymond-Belzile, Lo. NaN. Contributions to Likelihood-Based Modelling of Extreme Values. <https://doi.org/10.5075/epfl-thesis-9685>.
- [22] Samorodnitsky, G., (1986). Extrema of skewed stable process. Stochastic Process. Appl. 30, 1739.
- [23] Samoradnitsky, G., Taqqu, M. S. (1994). Stable non-Gaussian random processes: stochastic models with infinite variance (Vol. 1). CRC Press.
- [24] Zolotarev, V. M. (1986). *One-dimensional stable distributions*, vol. 65, American Mathematical Soc.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



A New Test Based on Linear Ordered Rank Statistics

Amini-Seresht, E.¹ Izadi, M.²

¹ Department of Statistics, Faculty of Science, Bu-Ali Sina University, Hamedan, Iran

² Department of Statistics, Razi University, Kermanshah, Iran

Abstract

In this paper we consider the problem of hypothesis testing the equality of two distribution function against the distribution functions are ordered with respect to usual stochastic order and the hazard rate order. The proposed statistical test is based on the linear combination of ordered rank statistics. The exact null and alternative distributions of the proposed test statistics are derived and are then used to generate the critical value.

Keywords: Non-parametric test, Stochastic ordering, Rank statistic, PHR model, Order statistics

1 Introduction

In the reliability context, it is often of interest to compare two populations and there are many non-parametric tests for comparing two populations based on complete independent random samples. In reliability studies, one may be interested in inferring the lifetimes of the two coherent system whither the components of a system A have the longer lifetime that those of the system B . Therefore, the non-parametric tests often very important to be used in testing two samples observed lifetime data are coming from a known statistical

¹e.amini@basu.ac.ir

model against a stochastically ordered alternative. In this paper we propose a new non-parametric test based on the linear ordered rank statistics.

Let $X_1, \dots, X_m$ and $Y_1, \dots, Y_n$ be two random samples from unknown continuous distribution functions F and G , with sizes m and n , respectively. Moreover, we also assume that X_i 's, $i = 1, \dots, m$ and Y_j 's, $j = 1, \dots, n$ are mutually stochastically independent and $X_{1:m} \leq \dots \leq X_{m:m}$ and $Y_{1:n} \leq \dots \leq Y_{n:n}$ are the corresponding order statistics. Banerjee [2] based on empirical distribution function obtained a new goodness of fit test for a one-sample and two-sample problem under the alternative hypothesis of usual stochastic ordering. In this paper, we extend the results in Banerjee [2] for the case of two sample problem based on the linear combination of the ordered rank statistics under the alternative hypothesis of the hazard rate order and we derive the exact distribution of proposed test statistic under H_0 and H_1 , where

$$H_0 : \overline{F}(x) = \overline{G}(x) \quad \text{versus} \quad H_1 : \overline{F}(x) = \overline{G}^\theta(x),$$

and $\theta > 1$ is a positive real number.

Davidov [3] consider a family of simple rank statistics as $\sum_{i=1}^n \Psi(R_i)$ for various choices of the function Ψ and proposed some test statistics by the stochastic ordering of the ranks. One may be defined a generic form of this type of test statistics as the form $\mathcal{A} = \sum_{i=1}^n \alpha_i \Psi(R_i)$, where $\alpha_1, \dots, \alpha_n$ are some real numbers (weights) which may the statistical properties of it also be of independent interest. We may conjecture that the test statistics as the form $\mathcal{A}$ more powerful than the existing test statistics of this type in the literature, such as Wilcoxon's rank sum statistic, van der Waerden's statistic and Gibbon's statistics (see [1]). Therefore, it will be of interesting to examine whether this conjecture is true, on the other word whether the test statistic $\mathcal{A}$ will result in a more powerful test. Motivated by this, suppose that $\Psi(R_i) = R_{i:m}$, $i = 1, \dots, m$. We consider the new test statistic as the form $\mathcal{A}_{n,m} = \sum_{i=1}^m \alpha_i R_{i:m}$, where $R_{i:m} = \text{Rank}(X_{i:m}) = \sum_{j=1}^{m+n} I(X_{i:m} - Z_j)$ will be called the rank of $X_{i:m}$ in the combined increasing arrangement of X_i 's, $i = 1, \dots, m$ and Y_j 's, $j = 1, \dots, n$, where $Z_1, \dots, Z_{n+m}$ be the random variables in the combined sample and $I(A)$ is an indicator function assuming the value 1 if $A \geq 0$ and 0 otherwise. We derive the exact distribution of the null and alternative hypothesis of $\mathcal{A}_{n,m}$ and provide the critical values for different sample sizes.

2 The test statistic

In this section we construct the non-parametric test statistic. This investigation will be based on the empirical distributions. Let $X_1, \dots, X_m$ and $Y_1, \dots, Y_n$ be independent random samples from the distributions F and G , respectively. The empirical distribution functions of F and G denote by F_m and G_n , respectively, and are defined by

$$F_m(x) = \frac{1}{m} \sum_{i=1}^m I(x - X_i) \quad \text{and} \quad G_n(x) = \frac{1}{n} \sum_{i=1}^n I(x - Y_i).$$

Denote by $\mathcal{A}_{n,m}$ for the estimated of below the (F_m, G_n) curve in the unit square. To construct the test statistic, we use the Simpsons one-third rule to compute $\mathcal{A}_{n,m}$. Based

on $\mathcal{A}_{n,m}$ we define the measure of deviation from H_0 to H_1 by

$$\mathcal{A}_{n,m} = \sum_{i=1}^m \alpha_i G_n \left(F_m^{-1} \left(\frac{i}{m} \right) \right), \quad (2.1)$$

where the coefficients are $\alpha_1 = \frac{1}{3m}$, $\alpha_i = \frac{2}{3m}$, for $i = 2, 4, \dots, m-2$, $\alpha_i = \frac{4}{3m}$, $i = 3, 5, \dots, m-1$ and $\alpha_m = \frac{1}{3m}$. Notice that under H_0 , $\mathcal{A}_{n,m}$ will be identical with the $X = Y$ line in the unit box and $\mathcal{A}_{n,m} = \frac{1}{m} \sum_{i=1}^m i \alpha_i$. Under H_1 , $\mathcal{A}_{n,m} \geq K$ for a large value $K > 0$. Note that, $F_m(X_{i:m}) = \frac{1}{m} \sum_{i=1}^m I(X_{i:m} - X_i) = \frac{i}{m}$, hence, $F_m^{-1} \left(\frac{i}{m} \right) = X_{i:m}$. Therefore, (2.1) can be rewritten as

$$\mathcal{A}_{n,m} = \sum_{i=1}^m \alpha_i G_n(X_{i:m}).$$

It is easy to see that

$$\begin{aligned} G_n(X_{i:m}) &= \frac{\{\text{number of random variables in the sample} \leq X_{i:m}\}}{n} \\ &= \frac{\{\text{rank of } X_{i:m}\} - i}{n} \\ &= \frac{R_{i:m} - i}{n}, \end{aligned} \quad (2.2)$$

from (2.2), $\mathcal{A}_{n,m}$ can be rewritten as

$$\begin{aligned} \mathcal{A}_{n,m} &= \frac{1}{n} \sum_{i=1}^m \alpha_i R_{i:m} - \frac{1}{n} \sum_{i=1}^m i \alpha_i \\ &= \frac{1}{n} \sum_{i=1}^m \alpha_i R_{i:m} - \frac{m}{n} \mathbb{E}_{H_0}(\mathcal{A}_{n,m}). \end{aligned} \quad (2.3)$$

To obtain the exact distribution of the test statistics $\mathcal{A}_{n,m}$ under the null hypothesis and alternative hypothesis, we first try to propose a new convenient form of $\mathcal{A}_{n,m}$. It is worth noting that $R_{i:m}$ (the rank of $X_{i:m}$) is equivalent to saying that it is the number of Z_j 's, $j = 1, \dots, m+n$ less than $X_{i:m}$ in the combined sample. Therefore, from this observation we have the following proposition to derive a generic expression of $\mathcal{A}_{n,m}$.

Proposition 2.1. *Let $X_1, \dots, X_m$ and $Y_1, \dots, Y_n$ be independent random samples from two absolutely continuous distribution functions F and G , respectively. Suppose that $R_{j:m}$ is the rank of $X_{j:m}$ in the combined increasing arrangement of X_i 's, $i = 1, \dots, m$ and Y_j 's, $j = 1, \dots, n$. Then*

$$R_{j:m} = \sum_{r=1}^j C_r + j$$

where $C_r = \sum_{i=1}^n [I(X_{r:m} - Y_i) - I(Y_i - X_{r-1:m})]$, $r = 2, \dots, m$ is the number of Y_i 's in $(X_{r-1:m}, X_{r:m}]$ and $C_1 = \sum_{i=1}^n I(X_{1:m} - Y_i)$ is the number of Y_i 's less than $X_{1:m}$.

3 Exact distribution of $\mathcal{A}_{n,m}$ under the hypothesis

In this section the exact and asymptotic distribution of the proposed test statistic under the null and alternative hypothesis will be derived. The exact distribution has been established through the use of beta distribution.

In order to evaluate the distribution of $\mathcal{A}_{n,m}$ under the null hypothesis, we consider two independent random samples of different size that the distribution functions F and G of the samples $X_1, \dots, X_m$ and $Y_1, \dots, Y_n$ are equal. First, we will derive the joint probability mass function of random variables $S_1, \dots, S_m$.

Theorem 3.1. *The joint distribution of $S_1, \dots, S_m$, under the null hypothesis, is given by*

$$\mathbb{P}_{H_0}(S_1 = s_1, \dots, S_m = s_m) = m! \prod_{i=1}^{m-1} B(s_i + i, s_{i+1} - s_i + 1) B(s_m + m, m - s_m + 1)$$

Note that to decide to reject the null hypothesis for large values of $\mathcal{A}_{m,n}$ and the level of significant α , or, equivalently, to calculate the p -value of the observed value of $\mathcal{A}_{m,n}$, from Theorem 3.2, the critical value K can be obtained as the following

$$\begin{aligned} \mathbb{P}_{H_0}(\mathcal{A}_{m,n} > K) &= \sum_{\substack{(s_1, \dots, s_m) \\ \alpha_1 s_1 + \dots + \alpha_m s_m > nK}} \mathbb{P}_{H_0}(S_1 = s_1, \dots, S_m = s_m) \\ &= \sum_{\substack{(s_1, \dots, s_m) \\ \alpha_1 s_1 + \dots + \alpha_m s_m > nK}} m! \prod_{i=1}^{m-1} B(s_i + i, s_{i+1} - s_i + 1) \\ &\quad \times B(s_m + m, m - s_m + 1). \end{aligned}$$

Next, we study the power of the proposed test. It worth mentioned here that in general, the power of a test is the probability of rejecting the null hypothesis when it is false and usually expressed as $1 - \beta$, where β is the Type II error and is the probability of accepting the null hypothesis when the alternative hypothesis is true. Likewise Theorem 3.2, the following theorem provides the joint probability mass function of random variables $S_1, \dots, S_m$ under H_1 .

Theorem 3.2. *The joint distribution of $S_1, \dots, S_m$, under the alternative hypothesis, is given by $P_{H_1}(S_1 = s_1, \dots, S_m = s_m) = m! \theta^m \sum_{k_1=0}^{\theta-1} \dots \sum_{k_{m-1}=0}^{\theta-1} \left\{ (-1)^{\sum_{r=1}^{m-1} k_r} \prod_{r=1}^{m-1} \binom{\theta-1}{k_r} \right.$*
 $\times \prod_{r=1}^{m-1} B(s_r + \sum_{r=1}^j k_r + r, s_r - s_{r-1} + 1)$
 $\times B(s_m + \sum_{r=1}^{m-1} k_r + m, m + \theta - s_m) \Big\}.$

This work presents the new test statistic for equality of two distribution, the results of the present work are not yet complete and we are currently working on this problem and try to generate critical values and the corresponding exact levels of significance and hope to update the present work by new findings.

References

- [1] Sidak, Z., Sen, P. K., & Hajek, J. (1999). *Theory of rank tests*. Academic Press.
- [2] Banerjee, S. (2008). A distribution free goodness of fit test for stochastically ordered alternatives. *Statistics and Probability Letters*, **78**, 2868-2875.
- [3] Davidov, O. (2012). Ordered inference, rank statistics and combining p-values: a new perspective. *Statistical Methodology*, **9(3)**, 456-465.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Preponed Preventive Maintenance for k -out-of- n Systems Subject to External Shocks

Ashrafi, S.¹ Asadi, M.²

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

Abstract

This paper investigates an $(n - k + 1)$ -out-of- n system exposed to shocks arriving via a counting process, with each shock potentially causing component failures. A system begins operation at time $t = 0$, and a preventive maintenance (PM) is scheduled for time T_p . The system is inspected at a time T_s before T_p . We implement a preponed PM strategy, where the decision to perform maintenance at a scheduled time depends on the observed number of failed components at T_s . If the number of failures exceeds a certain threshold, the PM is performed at T_s ; otherwise, it is performed at the scheduled time T_p . For this PM policy, the average cost per unit time is calculated, and the optimal decision parameters that minimize this average cost are determined. Numerical examples demonstrate the applicability and effectiveness of the proposed maintenance strategy.

Keywords: Nonhomogeneous Poisson process, Mean cost function, Emergency repair.

¹s.ashrafi@sci.ui.ac.ir

²m.asadi@sci.ui.ac.ir

1 Introduction

Reliability engineering strives to maintain technical systems at their optimal working conditions to ensure they perform reliably throughout their expected lifespan. One of the standard maintenance techniques is age-based maintenance, which analyzes failure time data based on a set overhaul schedule on calendar time. Age-based maintenance encompasses emergency repair (ER) and preventive maintenance (PM). ER involves unplanned interventions to rectify system failures, demanding skilled specialists and an immediate supply of spare parts. In contrast, PM is a scheduled maintenance technique aimed at averting catastrophic breakdowns, albeit interrupting production and incurring additional expenses. In recent years, many authors have investigated the PM of systems under different scenarios. For example, [5, 6, 8, 7] studied the PM models for the complex systems subject to shocks.

In a recent work, Finkelstein et al. [4] considered a single-unit system that is inspected at a predetermined time, and based on the state of the system, it is decided whether to perform PM at that time or postpone it to another time. Asadi [1] explored the postponed PM policy for complex systems. Ashrafi and Asadi [3] studied the postponed PM policy for complex systems where the inspection time is random.

In engineering and industrial applications, systems often face external or internal shocks and vibrations, which can significantly impair their performance. Minimising these effects is crucial for maintaining system safety and reliability. However, the unpredictable nature of shock events creates substantial challenges for engineers and reliability analysts. Shock models are essential tools for evaluating the physical effects of damage on engineered systems. These models incorporate key parameters, including shock amplitude, frequency, and duration, enabling engineers to simulate and analyze their effects on system components. By measuring these factors, potential failure modes can be recognised, aiding in the creation of more resilient designs. Recently, Asadi [2] investigated the postponement of PM policy for systems that are subject to shocks, and as a result of each shock, exactly one component fails. In this study, we examine an $(n - k + 1)$ -out-of- n system affected by shocks arriving via a counting process, with some components potentially failing after each shock. The system begins operation at time $t = 0$, and a PM is scheduled for time T_p . We propose a preponed PM policy in which the system is inspected at time T_s . If the number of failed components at T_s is equal to or exceeds a critical value m ($m \in \{1, \dots, k\}$), the PM action is performed at T_s , unless it is carried out at the scheduled time T_p . Additionally, if the system fails before T_s , an emergency repair is conducted. By considering the related costs for PM action, ER action, and the cost of renewing the failed components, we obtain the mean cost function as a function of (m, T_s, T_p) . The aim is to obtain the optimal parameters that minimize the mean cost per unit of time. To illustrate the application of the proposed model, a numerical example is provided.

2 System Description

In this paper, we consider an $(n - k + 1)$ -out-of- n system that is subject to shocks arriving according to a counting process $\{\xi(t), t > 0\}$. Assume that, as a result of each shock, some of the components may fail. Let W_i denote the number of failed components at the occurrence time of i th shock. It can be seen that the reliability function of this system is obtained as

$$P(T > t) = P(N(t) < k) = \sum_{i=0}^{\infty} G_i(k-1)P(\xi(t) = i)$$

where G_i denotes the distribution function of $\sum_{r=1}^i W_r$ (See [9] for any details). In this paper, we assume that at the time of the first shock occurrence, each component may fail with probability p_1 . Additionally, if a component is not failed by the first shock, it may fail with probability p_2 at the time of the second shock; similarly, if it is not failed by the first $(r-1)$ shocks, it may fail with probability p_r at the time of the r th shock. Let $q_r = 1 - p_r$ and $E_r, r = 1, 2, \dots$ denote the event that a component fails at the occurrence of the r th shock. Under this assumption, the probability that a component fails before or at the i th shock is

$$\begin{aligned} P(\cup_{r=1}^i E_r) &= P(E_1) + P(E_2) + \dots + P(E_i) \\ &= P(E_1) + P(E_2|E'_1)P(E'_1) + \dots + P(E_i|\cap_{r=1}^{i-1} E'_r)p(\cap_{r=1}^{i-1} E'_r) \\ &= p_1 + p_2q_1 + \dots + p_i(\prod_{r=1}^{i-1} q_r) = 1 - \prod_{r=1}^i q_r \end{aligned}$$

Also, the probability that one component fail between $(i+1)$ th shock and j th shock is

$$\begin{aligned} P(\cup_{r=i+1}^j E_r) &= P(E_{i+1}) + \dots + P(E_j) \\ &= P(E_{i+1}|\cap_{r=1}^i E'_r)P(\cap_{r=1}^i E'_r) + \dots + P(E_j|\cap_{r=1}^{j-1} E'_r)p(\cap_{r=1}^{j-1} E'_r) \\ &= p_{i+1}(\prod_{r=1}^i q_r) + \dots + p_j(\prod_{r=1}^{j-1} q_r) = \prod_{r=1}^i q_r - \prod_{r=1}^j q_r. \end{aligned}$$

Therefore, the probability that one component fails after the j th shock is $\prod_{r=1}^j q_r$. Then, the number of failed components until i th shock, i.e., $\sum_{r=1}^i W_r$, has binomial distribution with parameters $(n, 1 - \prod_{r=1}^i q_r)$ and $(\sum_{r=1}^i W_r, \sum_{r=i+1}^j W_r)$ has trinomial distribution with parameters $(n, 1 - \prod_{r=1}^i q_r, \prod_{r=1}^i q_r - \prod_{r=1}^j q_r)$. Hence, the joint distribution of $(\sum_{r=1}^i W_r, \sum_{r=i+1}^j W_r)$ can be written as

$$\begin{aligned} P\left(\sum_{r=1}^i W_r = s_1, \sum_{r=i+1}^j W_r = s_2\right) \\ = \frac{n!}{s_1!(s_2)!(n - s_1 - s_2)!} (1 - \prod_{r=1}^i q_r)^{s_1} (\prod_{r=1}^i q_r - \prod_{r=1}^j q_r)^{s_2} (\prod_{r=1}^j q_r)^{n-s_1-s_2} \end{aligned}$$

where $s_1, s_2 = 0, \dots, n$. Also, it can be seen that

$$P\left(\sum_{r=1}^i W_r = s\right) = \frac{n!}{s!(n-s)!} (1 - \prod_{r=1}^i q_i)^s (\prod_{r=1}^i q_r)^{n-s}$$

3 PM Model

We investigate a preponed preventive maintenance strategy for this system as follows. We assume that the system begins operating at time $t = 0$, and a PM time is scheduled at T_p . The system is inspected at time T_s , $T_s < T_p$. If the number of failed components at T_s is equal to or more than a critical value m ($m \in \{1, \dots, k\}$), the PM action is performed at T_s . Otherwise, the PM is carried out at time T_p . Also, if the system fails before T_s , an emergency repair is performed on the system. Assume that at failure time of the system, or time T_s with the number of failed components equal or more than m , or at time T_p , whichever comes first and is considered as the renewal time of the system the failed components are replaced with new ones which imposes the cost c_0 , the non-failed components are repaired with cost c_1 ($c_1 < c_0$) to become the same as a new component. Additionally, the cost c_{PM} is considered for preventive maintenance on the system, and the cost c_{ER} is considered for emergency repair action ($c_{PM} \ll c_{ER}$). To calculate the system's cost function, we need to determine the average number of failed components at the moment of system failure. Since more than one component can fail due to each shock, the total number of failed components at the system's failure may exceed k . Consider an $(n-k+1)$ -out-of- n system. If $M(k, n)$ denotes the mean number of failed components at the failure time of the system, then

$$M(k, n) = \frac{\sum_{l=k}^{n-1} l \binom{n}{l} B(n-l) (\sum_{r=1}^{k-1} \binom{l}{r} B(r) + 1) + n \sum_{r=1}^{k-1} \binom{n}{r} B(r) + n}{B(n)}$$

where

$$B(r) = \sum_{j=1}^r \sum_{i=0}^j \binom{j}{i} (-1)^i (j-i)^r$$

Proof. First, note that the total number of ways in which n components can fail across different shocks, with at least one component failing in each shock, is

$$B(n) \equiv \sum_{j=1}^n \sum_{i=0}^j \binom{j}{i} (-1)^i (j-i)^n. \quad (3.1)$$

See Lemma 1 of Zarezadeh et al. (2016) for details. Below, we calculate the number of ways the components can fail such that the total failed components at the system's failure time is l , with $k \leq l \leq n$. First, assume that $l < n$. We select l components from the total of n components in $\binom{n}{l}$ possible ways. In this scenario, two cases may occur:

- (I) The first is that all l components fail in one shock and the remaining $(n - l)$ components fail in different shocks. The number of ways of component failures in this case is $\binom{n}{l}B(n - l)$.
- (II) The second case is that r components ($r = 1, \dots, k-1$) fail in different shocks, $(l-r)$ remaining components fail in one shock and the remaining $(n - l)$ components can fail in different shocks. In this case, the number of ways of component failures is

$$\binom{n}{l} \sum_{r=1}^{k-1} \binom{l}{r} B(r) * B(n - l)$$

Therefore, for $l < n$, the total number of ways of component failures in which l components can fail in different shocks, such that in each shock at least one component fails, is

$$\binom{n}{l} B(n - l) + \binom{n}{l} \sum_{r=1}^{k-1} \binom{l}{r} B(r) * B(n - l)$$

Now, suppose that $l = n$, meaning all components fail at the system's failure time. Using the same reasoning as above, the number of ways components can fail such that the number of failed components at the system's failure time is n , is $(\sum_{r=1}^{k-1} \binom{l}{r} B(r) + 1)$. Hence, the mean number of failed components at the failure time of the system ($M(k, n)$) can be obtained as

$$M(k, n) = \frac{\sum_{l=k}^{n-1} l \binom{n}{l} B(n - l) (\sum_{r=1}^{k-1} \binom{l}{r} B(r) + 1) + (n \sum_{r=1}^{k-1} \binom{l}{r} B(r) + n)}{B(n)}$$

This completes the proof of the Lemma. □

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

- (I) The system fails before inspection time T_s . In this case, the mean cost function can be achieved as

$$C_1(T_s) = (M(k, n)c_0 + (n - M(k, n))c_1 + c_{ER})P(T < T_s)$$

and the mean length of one renewal cycle is obtained as

$$L_1(T_s) = P(T < T_s)E(T|T < T_s) = T_s P(T < T_s) - \int_0^{T_s} P(T < t)dt$$

where

$$P(T < T_s) = 1 - P(N(T_s) < k) = 1 - \sum_{i=0}^{\infty} P(\sum_{j=0}^i W_j \leq k-1) P(\xi(T_s) = i)$$

Under the assumption that shocks arrive according to the nonhomogeneous Poisson process (NHPP) with mean value function (m.v.f.) $\Lambda(t)$, we have

$$P(T < T_s) = 1 - \sum_{i=0}^{\infty} \frac{(\Lambda(T_s))^i \exp(-\Lambda(T_s))}{i!} \sum_{j=0}^{k-1} \frac{n!}{j!(n-j)!} (1 - \prod_{r=1}^i q_r)^j (\prod_{r=1}^i q_r)^{n-j}$$

- (II) The system works at time T_s and the number of failed components at T_s is more than or equal to m . In this case, the mean cost function can be obtained as

$$C_2(m, T_s) = \sum_{l=m}^{k-1} (lc_0 + (n-l)c_1 + c_{PM})P(N(T_s) = l)$$

and the mean length of one renewal cycle is obtained as

$$L_2(m, T_s) = T_s P(T > T_s, N(T_s) \geq m) = T_s \sum_{l=m}^{k-1} P(N(T_s) = l)$$

where

$$P(N(T_s) = l) = \sum_{i=0}^{\infty} P\left(\sum_{r=0}^i W_r = l\right) P(\xi(T_s) = i)$$

Under the assumption that shocks arrive according to the NHPP with m.v.f. $\Lambda(t)$, it can be seen that

$$P(N(T_s) = l) = \sum_{i=0}^{\infty} \frac{(\Lambda(T_s))^i \exp(-\Lambda(T_s))}{i!} \frac{n!}{l!(n-l)!} \left(1 - \prod_{r=1}^i q_r\right)^l \left(\prod_{r=1}^i q_r\right)^{n-l}$$

- (III) The number of failed components at T_s is less than m and the system fails before T_p . In this case, the mean cost function can be obtained as

$$C_3(m, T_s, T_p) = (M(k, n)c_0 + (n - M(k, n))c_1 + c_{ER})P(N(T_s) < m, T < T_p)$$

and the mean length of one renewal cycle is obtained as

$$\begin{aligned} L_3(m, T_s, T_p) &= E(T | N(T_s) < m, T < T_p) P(N(T_s) < m, T < T_p) \\ &= T_p P(N(T_s) < m, T < T_p) - \int_{T_s}^{T_p} P(N(T_s) < m, T < t) dt \end{aligned}$$

where

$$\begin{aligned} P(N(T_s) < m, T < T_p) &= P(N(T_s) < m) - P(N(T_s) < m, T > T_p) \\ &= P(N(T_s) < m) - P(N(T_s) < m, N(T_p) < k) \\ &= \sum_{i=0}^{\infty} P\left(\sum_{j=0}^i W_j \leq m-1\right) P(\xi(T_s) = i) \\ &\quad - \sum_{i=0}^{\infty} \sum_{j=i}^{\infty} P\left(\sum_{r=0}^i W_r \leq m-1, \sum_{r=0}^j W_r \leq k-1\right) P(\xi(T_s) = i, \xi(T_p) = j) \end{aligned}$$

Under the assumption that shocks arrive according to the NHPP with m.v.f. $\Lambda(t)$, it can be seen that

$$\begin{aligned}
& P(N(T_s) < m, T < T_p) \\
&= \sum_{i=0}^{\infty} \frac{(\Lambda(T_s))^i \exp(-\Lambda(T_s))}{i!} \sum_{j=0}^{m-1} \frac{n!}{j!(n-j)!} (1 - \prod_{r=1}^i q_r)^j (\prod_{r=1}^i q_r)^{n-j} \\
&- \sum_{i=0}^{\infty} \sum_{j=i}^{\infty} \frac{(\Lambda(T_s))^i (\Lambda(T_p) - \Lambda(T_s))^{j-i} \exp(-\Lambda(T_p))}{i!(j-i)!} \\
&\quad \left(\sum_{s_1=0}^{m-1} \sum_{s_2=s_1}^{k-1} \frac{n!}{s_1!(s_2-s_1)!(n-s_2)!} (1 - \prod_{r=1}^i q_r)^{s_1} (\prod_{r=1}^i q_r - \prod_{r=1}^j q_r)^{s_2-s_1} (\prod_{r=1}^j q_r)^{n-s_2} \right)
\end{aligned}$$

- (V) The number of failed components at T_s is less than m and the system works at T_p . In this case, the mean cost function can be obtained as

$$\begin{aligned}
& C_4(m, T_s, T_p) \\
&= \sum_{l=0}^{k-1} (lc_0 + (n-l)c_1 + c_{PM}) P(N(T_s) < m, N(T_p) = l) \\
&= \sum_{l=0}^{k-1} (lc_0 + (n-l)c_1 + c_{PM}) \sum_{i=0}^{\infty} \sum_{j=i}^{\infty} P(\sum_{r=1}^i W_r \leq m-1, \sum_{r=1}^j W_r = l) \\
&\quad \times P(\xi(T_s) = i, \xi(T_p) = j)
\end{aligned}$$

and the mean length of one renewal cycle is obtained as

$$L_4(m, T_s, T_p) = T_p \sum_{l=0}^{k-1} P(N(T_s) < m, N(T_p) = l)$$

where

$$\begin{aligned}
& P(N(T_s) < m, N(T_p) = l) \\
&= \sum_{i=0}^{\infty} \sum_{j=i}^{\infty} P(\sum_{r=0}^i W_r \leq m-1, \sum_{r=0}^j W_r = l) P(\xi(T_s) = i, \xi(T_p) = j)
\end{aligned}$$

Under the assumption that shocks arrive according to the NHPP with m.v.f. $\Lambda(t)$, it can be seen that

$$\begin{aligned}
& P(N(T_s) < m, N(T_p) = l) \\
&= \sum_{i=0}^{\infty} \sum_{j=i}^{\infty} \frac{(\Lambda(T_s))^i (\Lambda(T_p) - \Lambda(T_s))^{j-i} \exp(-\Lambda(T_p))}{i!(j-i)!} \\
&\quad \left(\sum_{s=0}^{\min\{m-1, l\}} \frac{n!}{s!(l-s)!(n-l)!} (1 - \prod_{r=1}^i q_r)^s (\prod_{r=1}^i q_r - \prod_{r=1}^j q_r)^{l-s} (\prod_{r=1}^j q_r)^{n-l} \right)
\end{aligned}$$

Therefore, the mean cost per unit of time can be obtained as

$$\zeta(m, T_s, T_p) = \frac{C(m, T_s, T_p)}{L(m, T_s, T_p)}$$

where

$$C(m, T_s, T_p) = C_1(T_s) + C_2(m, T_s) + C_3(m, T_s, T_p) + C_4(m, T_s, T_p)$$

and

$$L(m, T_s, T_p) = L_1(T_s) + L_2(m, T_s) + L_3(m, T_s, T_p) + L_4(m, T_s, T_p)$$

The aim is to obtain m^* , and T_p^* that minimize the mean cost per unit of time.

4 Numerical Results

We consider a 2-out-of-5 system that is subject to shocks, where, as a result of each shock, some components may fail. We assume that shocks arrive according to the NHPP with m.v.f. $\Lambda(t) = \frac{t^2}{100}$. Further, we assume that $q_i = 0.9$, $i = 1, \dots, 5$. This means that the probability of each component failing due to the shocks is 0.1. Consider the cost parameters $c_0 = 2$, $c_1 = 0.5$, $c_{PM} = 5$ and $c_{ER} = 20$. In Table 1, the optimal values T_p^* and related optimum costs $\zeta(m, T_s, T_p^*)$ are given for different values of $T_s = 10, 15, 20, 25$ and $m = 1, 2, 3, 4$. It can be concluded from the results of the table that:

- For each fixed T_s , the optimal time to perform PM, T_p^* , decreases as m increases, indicating that increasing m results in a decrease in T_p^* .
- For each fixed $m = 1, 2, 3$, when the inspection time T_s increases, the optimal PM time T_p^* also increases and for $m = 4$, the optimal PM time T_p^* is fixed and does not depend on the inspection time T_s .
- For $T_s = 10, 15$, the mean cost per unit of time, $\zeta(m, 10, T_p^*)$, reaches its minimum at $m = 3$, and for $T_s = 20, 25$, it reaches its minimum at $m = 2$.

Table 1: The optimal values of T_{PM}^* and related optimum costs for selected T_s .

	$T_s = 10$		$T_s = 15$	
m	T_p^*	$\zeta(m, 10, T_p^*)$	T_p^*	$\zeta(m, 15, T_p^*)$
1	23.24278	0.59077	24.50609	0.57185
2	22.12057	0.56033	23.13513	0.54948
3	21.60003	0.55738	22.07028	0.54909
4	21.48830	0.56105	21.48830	0.56105
	$T_s = 20$		$T_s = 25$	
m	T_p^*	$\zeta(m, 20, T_p^*)$	T_p^*	$\zeta(m, 25, T_p^*)$
1	26.45269	0.55415	29.52992	0.57232
2	25.13218	0.54506	28.34077	0.56936
3	23.61553	0.54662	26.70812	0.57054
4	21.48830	0.56105	21.48830	0.56105

The plots of the average cost function $\zeta(m, T_s, T_p)$ as a function of T_p for $T_s = 15$, and various values of m is presented in Figure 1: The case $m = 1$ is shown in the upper-left plot, $m = 2$ is depicted in the upper-right plot, $m = 3$ in the lower-left plot, and the case of $m = 4$ is shown in the lower-right plot. It is evident from the plots that the optimal time T_p^* is a decreasing function of m under the given cost parameters when $T_s = 15$.

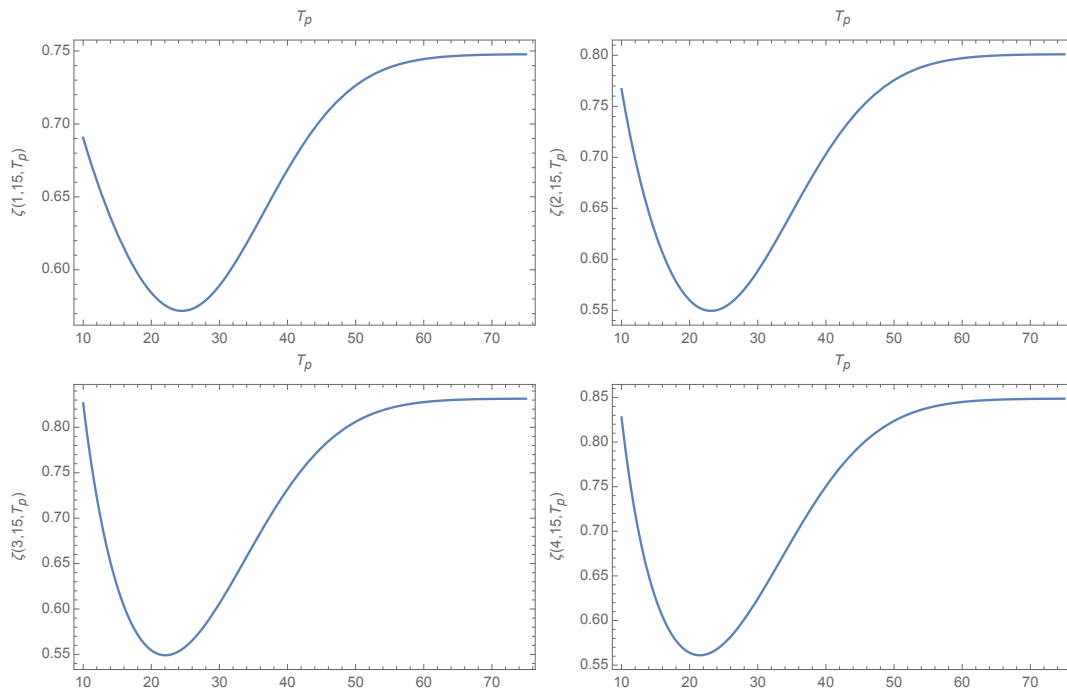


Figure 1: The plots of the average cost function $\zeta(m, T_s, T_p)$ as a function of T_p for $T_s = 15$, and various values of m : $m = 1$ (the upper-left), $m = 2$ (the upper-right), $m = 3$ (the lower-left), and $m = 4$ (the lower-right).

References

- [1] Asadi, M. (2024a). Preventive maintenance for coherent systems considering postponed replacement. *Applied Stochastic Models in Business and Industry*, **40**(4), 875-894.
- [2] Asadi, M. (2024b). An opportunistic maintenance model for multi-component systems experiencing shocks. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, 1748006X251338435.
- [3] Ashrafi, S., & Asadi, M. (2025). Optimal preventive maintenance of multi-component systems under random inspection. *Reliability Engineering & System Safety*, **257**, 110809.
- [4] 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.

- [5] 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.
- [6] Finkelstein, M., & Gertsbakh, I. (2016). On preventive maintenance of systems subject to shocks. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, **230(2)**, 220-227.
- [7] Zarezadeh, S., & Asadi, M. (2019). Coherent systems subject to multiple shocks with applications to preventative maintenance. *Reliability Engineering & System Safety*, **185**, 124-132.
- [8] Zarezadeh, S., & Ashrafi, S. (2019). On preventive maintenance of networks with components subject to external shocks. *Reliability Engineering & System Safety*, **191**, 106559.
- [9] Zarezadeh, S., Ashrafi, S. & Asadi, M. (2016). A shock model based approach to network reliability. *IEEE Transactions on Reliability*, **65(2)**, 992-1000.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Analysing Random Maintenance Data by Repair Alert Model for Discrete lifetimes

Atlekhani, M.¹ Doostparast, M.²

¹ Department of Statistics, Faculty of Mathematical Sciences and Statistics, Malayer University, Malayer, Iran

² Department of Statistics, Ferdowsi University of Mashhad, Iran

Abstract

This paper presents an examination of repair alert models for analyzing discrete lifetimes coming from maintenance programs. The study considers the negative binomial distribution for component lifetimes. Furthermore, we address the challenges associated with parameter estimation for the negative binomial distribution and compute some approximate confidence intervals for the selected distribution based on the Akaike Information Criterion. Findings are illustrated by analyzing a real dataset.

Keywords: Maximum likelihood estimation, Random sign censoring, Repair alert model

1 Introduction

Reliability is a fundamental concept in engineering, statistics, and system management, defined as the probability that a system, product, or process will perform its intended

¹m.atlekhani@malayeru.ac.ir

²doustparast@um.ac.ir

function under specified conditions over a defined period. In simpler terms, reliability indicates a system's capability to operate effectively without failure or defects. A crucial aspect of reliability is Reliability-Centered Maintenance (RCM), which is a systematic approach focused on ensuring the reliability and functionality of equipment and systems throughout their life cycle. The primary objective of RCM is to identify the most effective maintenance strategies that can prevent failures and minimize downtime while optimizing costs. 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 [4, 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]). To address this challenge, Cook [5] introduced the concept of *random sign censoring*, where X and Z are dependent, yet the marginal DF of X remains identifiable.

Definition 1.1 ([7, 4]). Let (X, Z) be a pair of lifetimes. Then $\{(Y = \min(X, Z), \delta = I(Z < X))\}$ is called a random signs censoring (RSC) of X by Z if the event $\{Z < X\}$ is (stochastically) independent of X .

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 Random Sign Censoring Data (RSCD). The formal definition of the RAM is provided.

Definition 1.2. [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 $x \in (0, +\infty)$,

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

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] examined the RAM in situations where the discrete lifetime distribution X belongs to a broad class known as the telescopic family. They modified Definition 1.2 for discrete lifetimes. For sake of brevity, the supports of X and Z are assumed to be $\mathbb{N} := \{0, 1, 2, \dots\}$, i. e. $S_x = S_z = \mathbb{N}$. Following [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 $x \in \mathbb{N}$

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

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

The sub-DFs and the conditional DFs of X and Z are, respectively, given by

$$\tilde{F}_X(x) = F_X(x), \quad (1.4)$$

$$F_X^*(x) = (1 - q)F_X(x), \quad (1.5)$$

$$\tilde{F}_Z(z) = F_X(z) + \sum_{k=z+1}^{\infty} \frac{G(z)}{G(k)} P(X = k), \quad (1.6)$$

$$F_Z^*(z) = q\tilde{F}_Z(z), \quad (1.7)$$

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

Note that the conditional PMF of Z is derived as

$$\begin{aligned} \tilde{f}_Z(z) &= P(Z = z | Z < X) \\ &= g(z) \sum_{k=z}^{\infty} \frac{P(X = k)}{G(k)}, \quad \text{for all } z \in \mathbb{N}, \end{aligned} \quad (1.8)$$

where $g(z) = G(z) - G(z - 1)$. 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 RSCD. Our focus will be on the RAM when the lifetimes are represented as discrete random variables, with the distribution of device lifetimes is the negative binomial distribution. 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 DRAM in the negative binomial distribution is studied in detail. In Section 3, the maximum likelihood estimates (MLEs) of the parameters and approximate confidence intervals are also derived. In Section 4, a real data set on ARC-1 VHF, reported by [8], is analyzed. Section 5 provides conclusions.

2 DRAM in the negative binomial distribution

In this section, we assume that the DF of lifetime X under the DRAM following the NB distribution.

Proposition 2.1. *Under the DRAM, assume that $X \sim NB(r, p)$. Then, the sub-PMFs of X and Z are, respectively, given by*

$$\begin{aligned} f_X^*(x; p) &= P(X = x, X < Z) \\ &= (1 - q)P(X = x) \\ &= (1 - q) \frac{\Gamma(x + r)}{x! \Gamma(r)} p^x (1 - p)^r, \quad \forall x \in \mathbb{N}. \end{aligned} \quad (2.1)$$

The second equality follows by the fact that X and the event $\{X < Z\}$ are independent. The sub-PMF of Z as

$$\begin{aligned} f_Z^*(z; \theta) &= P(Z = z, Z < X) \\ &= P(Z < X)P(Z = z | Z < X) \\ &= q\tilde{f}_Z(z) = qg(z) \sum_{k=z}^{\infty} \frac{P(X = k)}{G(k)} \\ &= qg(z) \sum_{k=z}^{\infty} \frac{\Gamma(k + r)}{k! \Gamma(r) G(k)} p^k (1 - p)^r, \quad \forall z \in \mathbb{N}. \end{aligned} \quad (2.2)$$

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

$$\begin{aligned} f_Z^*(z; p) &= qg(z) \sum_{k=z}^{\infty} \frac{P(X = k)}{G(k)} \\ &= q(z^p - (z - 1)^p) \sum_{k=z}^{\infty} \frac{\Gamma(k + r)}{k! \Gamma(r) k^p} p^k (1 - p)^r, \quad \forall z \in \mathbb{N}. \end{aligned} \quad (2.3)$$

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 negative binomial distribution, from Equations (2.1), (2.3) and (3.1) we obtain

$$\begin{aligned}
L(p, r, q; \mathbf{x}, \mathbf{z}) &= \prod_{i=1}^m \left((1-q) \frac{\Gamma(x_i + r)}{x_i! \Gamma(r)} p^{x_i} (1-p)^r \right) \\
&\quad \prod_{j=1}^n \left(q (z_j^p - (z_j - 1)^p) \sum_{k=z_j}^{\infty} \frac{\Gamma(k+r)}{k! \Gamma(r) k^p} p^k (1-p)^r \right) \\
&= p^{\sum_{i=1}^m x_i} (1-q)^m q^n \frac{(1-p)^{r(m+n)}}{(\Gamma(r))^{m+n}} \prod_{i=1}^m \left(\frac{\Gamma(x_i + r)}{x_i!} \right) \\
&\quad \prod_{j=1}^n \left((z_j^p - (z_j - 1)^p) \sum_{k=z_j}^{\infty} \frac{\Gamma(k+r)}{k! k^p} p^k \right). \tag{3.2}
\end{aligned}$$

The LLF is then

$$\begin{aligned}
l(p, r, q; \mathbf{x}, \mathbf{z}) &= \ln L(p, r, q; \mathbf{x}, \mathbf{z}) \\
&= n \ln q + m \ln(1-q) + \left(\sum_{i=1}^m x_i \right) \ln p + r(m+n) \ln(1-p) \\
&\quad - (m+n) \ln \Gamma(r) + \sum_{i=1}^m (\ln \Gamma(x_i + r) - \ln x_i!) + \sum_{j=1}^n \ln (z_j^p - (z_j - 1)^p) \\
&\quad + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{\Gamma(k+r)}{k! k^p} p^k \right). \tag{3.3}
\end{aligned}$$

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

$$\begin{aligned}
l_P(p, r; \mathbf{x}, \mathbf{z}) &= \ln L(p, r, \hat{q}; \mathbf{x}, \mathbf{z}) \\
&= n \ln \left(\frac{n}{m+n} \right) + m \ln \left(\frac{m}{m+n} \right) \\
&\quad + \left(\sum_{i=1}^m x_i \right) \ln p + r(m+n) \ln(1-p) - (m+n) \ln \Gamma(r) \\
&\quad + \sum_{i=1}^m (\ln \Gamma(x_i + r) - \ln x_i!) + \sum_{j=1}^n \ln (z_j^p - (z_j - 1)^p) \\
&\quad + \sum_{j=1}^n \ln \left(\sum_{k=z_j}^{\infty} \frac{\Gamma(k+r)}{k! k^p} p^k \right). \tag{3.4}
\end{aligned}$$

The ML estimate of p is obtained by maximizing (3.4). To do this, one may use numerical methods; see Section 4.

3.1 Approximate confidence intervals

Under some regular conditions, the ML estimate of the parameter vector $\boldsymbol{\theta} = (\theta_1, \dots, \theta_r)$ has asymptotically the r -variate normal distribution with mean $\boldsymbol{\theta}$ and the covariance

matrix $\Sigma = \mathbf{I}^{-1}$, where $\mathbf{I}$ is the Fisher information (FI) matrix (see [6, p.476]). Indeed, $\mathbf{I} = -E_{\theta}(\mathbf{H})$, where $\mathbf{H} = [[\partial^2/\partial\theta_i\partial\theta_j \log L]]_{1 \leq i,j \leq r}$ is the Hessian matrix of the LLF. Also, one can use the observed FI instead of the FI $\mathbf{I}$. The observed FI is defined by $i(\hat{\theta}) = -\hat{\mathbf{H}}$, where $\hat{\mathbf{H}} = [[\partial^2/\partial\theta_i\partial\theta_j \log L]]|_{\theta=\hat{\theta}}$. Therefore, under the regular conditions, we have

$$\hat{\Sigma}^{-1}(\hat{\theta} - \theta) \rightarrow N_r(\mathbf{0}, \mathbf{J}), \quad (3.5)$$

where J is the $r \times r$ identity matrix. In the rest, the approximate confidence intervals are derived under a member of the PS family. From Equation (3.5), the approximate $100(1 - \gamma)\%$ equal-tailed confidence intervals for θ_i is obtained as

$$\hat{\theta}_i \pm z_{\gamma/2} \sqrt{\widehat{\text{var}}(\hat{\theta}_i)}, \quad 1 \leq i \leq r \quad (3.6)$$

respectively, where z_{γ} stands for the upper $100\gamma\%$ th percentile of the standard normal distribution. The details are illustrated in the next section.

4 Case study

Table 1 presents 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).

For the analysis this dataset, the negative binomial distribution was considered for modeling the lifetime X . Numerical computations were performed using *R* software (version 4.4.0). The source code is available from the authors upon request. The Akaike Information Criterion (AIC) was employed to compare the fitted models. The AIC is defined as: $\text{AIC} = -2\widehat{LLF} + 2M$, 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 $p \in (0.9917, 0.9929)$.

5 Conclusion

In this paper, we examine 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. Notably, we considered NB distribution known. Furthermore, we addressed the challenges associated with parameter estimation

Table 1: ARC-1 VHF communication transmitter-receivers of a single commercial number of hours to failure

x	16, 224, 16, 80, 128, 168, 144, 176, 176, 568, 392, 576, 128, 56, 112, 160, 384, 600, 40, 416, 408, 384, 256, 246, 184, 440, 64, 104, 168, 408, 304, 16, 72, 8, 88, 160, 48, 168, 80, 512, 208, 194, 136, 224, 32, 504, 40, 120, 320, 48, 256, 216, 168, 184, 144, 224, 488, 304, 40, 160, 488, 120, 208, 32, 112, 288, 336, 256, 40, 296, 60, 208, 440, 104, 528, 384, 264, 360, 80, 96, 360, 232, 40, 112, 120, 32, 56, 280, 104, 168, 56, 72, 64, 40, 480, 152, 48, 56, 328, 192, 168, 168, 114, 280, 128, 416, 392, 160, 144, 208, 96, 536, 400, 80, 40, 112, 160, 104, 224, 336, 616, 224, 40, 32, 192, 126, 392, 288, 248, 120, 328, 464, 448, 616, 168, 112, 448, 296, 328, 56, 80, 72, 56, 608, 144, 408, 16, 560, 144, 612, 80, 16, 424, 264, 256, 528, 56, 256, 112, 544, 552, 72, 184, 240, 128, 40, 600, 96, 24, 184, 272, 152, 328, 480, 96, 296, 592, 400, 8, 280, 72, 168, 40, 152, 488, 480, 40, 576, 392, 552, 112, 288, 168, 352, 160, 272, 320, 80, 296, 248, 184, 264, 96, 224, 592, 176, 256, 344, 360, 184, 152, 208, 160, 176, 72, 584, 144, 176.
z	368, 136, 512, 136, 472, 96, 144, 112, 104, 104, 344, 246, 72, 80, 312, 24, 128, 304, 16, 320, 560, 168, 120, 616, 24, 176, 16, 24, 32, 232, 32, 112, 56, 184, 40, 256, 160, 456, 48, 24, 200, 72, 168, 288, 112, 80, 584, 368, 272, 208, 144, 208, 114, 480, 114, 392, 120, 48, 104, 272, 64, 112, 96, 64, 360, 136, 168, 176, 256, 112, 104, 272, 320, 8, 440, 224, 280, 8, 56, 216, 120, 256, 104, 104, 8, 304, 240, 88, 248, 472, 304, 88, 200, 392, 168, 72, 40, 88, 176, 216, 152, 184, 400, 424, 88, 152, 184.

for the NB distribution. Based on the Akaike Information Criterion the negative binomial 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 $p \in (0.9917, 0.9929)$. This paper may be extended in various direction. For example, RSCD may be analyzed by the Bayesian approach. Prediction of the component lifetime X based on the observed RSCD is also worth for 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*,

Table 2: ML estimates of parameters based on ARC-1 VHF data set

Distribution	$G(t)$	$MLestimate$	$\widehat{LLF}$	AIC
negative binomial	t^p	$r = 1, \hat{p} = 0.9962845$	-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
		$r = 6, \hat{p} = 0.97676$	-2430.464	4862.928
		$r = 10, \hat{p} = 0.9616$	-2646.737	5295.474
		$r = 20, \hat{p} = 0.9257$	-3197.791	6397.582

to appear.

- [4] Cook, R. M. (1993). The total time on test statistics and age-dependent censoring. *Statistics and Probability Letters*, **18**, 307–312.
- [5] Cook, (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.
- [6] Lehmann, E. L., Casella, G. (1998). *Theory of Point Estimation*, 2-th Edition. Springer-Verlag, New York.
- [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.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



A Q-Learning Approach to Maintenance Policy Optimization for Series Systems with Degradation and Common Shocks

Farahmand, A. H.¹ Fartash, F.² Yadollahi, M. R.³ Zarandi, S.⁴

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

Abstract

This paper addresses the optimization of maintenance strategies for a two-component series system subjected to both natural degradation and common shock-induced degradation. The system is modeled with degradation processes for each component, where natural degradation follows a homogeneous gamma process, and shock-induced degradation is governed by a homogeneous Poisson process with shared shocks affecting both components. The objective is to develop an optimal maintenance policy using Q-learning, a reinforcement learning technique, to minimize maintenance costs while preventing system failure. The maintenance strategy includes corrective, preventive, and opportunistic maintenance actions, with the decision-making process based on the degradation levels of both components observed during periodic inspections. The proposed method is evaluated through simulations to assess the effectiveness of the Q-learning-based policy in reducing overall costs and ensuring system reliability. Results show that the optimal policy balances the costs of inspections, maintenance actions, and system downtime, providing significant cost savings compared to traditional maintenance approaches. The

¹amir.farahmand@ut.ac.ir

²fateme.fartash@ut.ac.ir

³mhryadollahi@ut.ac.ir

⁴zarandi.sina@ut.ac.ir

approach demonstrates the potential of reinforcement learning in dynamic maintenance optimization for complex systems subject to degradation and external shocks.

Keywords: Condition-based maintenance, Opportunistic maintenance, Common cause failure, Degradation modeling, Q-learning.

1 Introduction

The increasing complexity of modern equipment poses significant challenges to reliability and maintenance, highlighting the critical concern of common cause failures (CCF), which involve the simultaneous failure of multiple components leading to system-wide breakdowns. A prominent example is the 2009 Air France Flight AF447 crash, attributed to the simultaneous failure of three pitot tubes. While various models, such as shock models, have attempted to characterize CCF, a universally accepted definition and estimation method remain elusive, often due to the difficulty in incorporating gradual performance deterioration.

Effective maintenance is crucial for ensuring system reliability. The cost of unexpected downtime in US factories alone reached nearly \$50 billion in 2013, and maintenance costs can consume up to %70 of production budgets, with an estimated %33 of that being wasted. This underscores the growing industry trend toward more efficient maintenance strategies that optimize maintenance timing and thresholds to minimize costs and maximize system availability.

Traditional maintenance approaches, such as post-failure or age-based strategies, suffer from significant limitations. Post-failure maintenance often results in substantial consequential damage, while age-based strategies can be overly conservative, leading to unnecessary actions like premature component replacements. The advent of advanced sensing and monitoring technologies has enabled condition-based maintenance (CBM), which triggers maintenance "when needed" or just before failure. CBM offers significant cost reductions and extends equipment lifespan across diverse sectors, including mechanical, aviation, and electrical engineering. However, conventional CBM strategies typically focus on individual components or single-component systems, neglecting the crucial interdependencies between components in complex systems. Applying such single-component strategies to multi-component systems without considering these dependencies can lead to suboptimal outcomes, driving increasing research into CBM policies for complex, interdependent systems.

Recent studies have explored CBM policies for two-component systems. For instance, Broek et al. [1] proposed a CBM policy for a two-unit system with economic dependency, while Wang et al. [3] investigated the joint optimization of condition-based replacement, age-based replacement, and spare parts inventory for a two-unit series system. However, these studies often consider only one type of dependency. Incorporating multiple dependencies can significantly increase model complexity, making them practically unsolvable or analytically intractable.

Dependencies among components can manifest in various forms, including economic

and structural dependencies. Economic dependency arises when combined maintenance of multiple components is more cost-effective than individual maintenance due to shared setup costs. Structural dependency is common in continuously operating systems like aircraft, power plants, and chemical processing facilities. To address these dependencies, opportunistic maintenance (OM) has emerged as a valuable strategy. OM capitalizes on specific opportunities, such as scheduled downtime or maintenance windows, to perform maintenance actions, rather than adhering to rigid pre-planned schedules.

Several studies have focused on optimizing OM strategies for multi-component systems. For example, Chen et al. [3] explored condition-based OM policies for continuously degrading multi-component systems with real-time monitoring and periodic inspections. Shi et al. [4] proposed an OM strategy for a series multi-machine production system, where maintenance on other components is performed when one component reaches its preventive maintenance threshold. While many OM studies assume progressive system degradation, they often overlook the impact of CCF, potentially leading to inaccurate reliability analyses and suboptimal maintenance scheduling.

This paper addresses this gap by considering a two-component series system with economic dependence [5], where simultaneous maintenance of multiple components reduces costs due to shared setup expenses. Specifically, we examine a two-component series system subject to both natural degradation and common shocks. Our maintenance strategy comprises four actions: corrective maintenance (CM), preventive maintenance (PM), opportunistic maintenance (OM), and do nothing. When one component undergoes PM or CM, the other component is considered for OM if appropriate. All maintenance actions are assumed to be perfect, restoring components to an as-good-as-new condition. We employ Q-learning to model the maintenance problem and determine the optimal actions, minimizing cost based on the learned Q-table.

2 Method

2.1 System Overview

The system consists of two components connected in series. Both components are subject to individual degradation and common shock processes. At time $t = 0$, both components were in perfect conditions and didn't have any degradation. The degradation of each component is modeled through a gamma process while exposed to common shocks arriving according to a homogeneous Poisson process.

2.2 Natural Degradation

The natural degradation $W_i(t)$ of component i follows a homogeneous gamma process. The cumulative distribution function (CDF) of $W_i(t)$ is given by:

$$G_{W_i}(x_i, t) = P\{W_i(t) \leq x_i\}$$

For any time interval δ , the increment $W_i(t + \delta) - W_i(t)$ follows a gamma distribution with shape parameter $\alpha_i \delta$ and scale parameter β_i . The probability density function (PDF)

is:

$$g_{\alpha_i\delta,\beta_i}(x) = \frac{\beta_i^{\alpha_i\delta}}{\Gamma(\alpha_i\delta)} x^{\alpha_i\delta-1} e^{-\beta_i x} I_{\{x \geq 0\}}$$

2.3 Shock-Induced Degradation

The system is exposed to common shocks that arrive according to a homogeneous Poisson process (HPP) with rate λ . Each shock causes sudden damage to both components, with the damage size for component i denoted by Y_i . The cumulative shock-induced degradation for component i at time t is:

$$\theta_i(t) = \sum_{j=1}^{N(t)} Y_{ij}$$

where $N(t)$ is the number of shocks by time t .

2.4 Total Degradation and Failure Criteria

The total degradation of component i at time t is the sum of natural degradation and shock-induced degradation:

$$X_i(t) = W_i(t) + \theta_i(t)$$

The system fails if the total degradation of either component exceeds its predefined failure threshold D_i .

3 Maintenance Strategy

A condition-based opportunistic maintenance policy is used with a two-component series system. Maintenance actions are ignited based on apprehended degradation levels of the components. The system experiences common-cause failures (CCF) due to shared shock processes. A key economic dependency explored in this paper is the presence of a fixed cost associated with any maintenance activity, regardless of how many components are involved. This fixed cost makes it economically advantageous to combine maintenance tasks, leading to the consideration of opportunistic maintenance (OM). OM, in conjunction with corrective maintenance (CM) and preventive maintenance (PM), forms the complete set of maintenance operations analyzed here.

3.1 Maintenance Actions

The following maintenance actions are considered:

- **Corrective Maintenance (CM):** Performed when a component's degradation exceeds its failure threshold D_i .

- **Preventive Maintenance (PM):** Performed when a component's degradation reaches a preventive threshold M_i .
- **Opportunistic Maintenance (OM):** Performed on one component when the other undergoes PM or CM, if its degradation exceeds an opportunistic threshold O_i .
- **Do Nothing:** No maintenance is performed if the degradation levels of both components are below their respective thresholds.

All maintenance actions are considered as perfect maintenance, restoring the system to an as-good-as-new condition.

3.2 Costs in Maintenance Decisions

The costs associated with maintenance decisions include:

- C : Inspection cost at each inspection time.
- C_S : Set-up cost when the system is stopped.
- C_M, C_P, C_O : Costs for corrective, preventive, and opportunistic maintenance for component i .
- C_T : Cost of sending a maintenance team.

Since we are dealing with economic dependency in this problem, establishing a logical connection between the various repair costs was one of the main challenges we faced. Another issue that caused many challenges was the completeness of the repairs because this made it difficult for the agent to initially distinguish between the different actions, and this issue also had to be addressed by selecting appropriate costs.

4 Reinforcement Learning

4.1 Discretization

In many real-world scenarios, reinforcement learning agents operate in environments with continuous state and action spaces. However, the standard Q-learning algorithm is designed for discrete spaces. To apply Q-learning to continuous problems, **discretization** is a common technique used to approximate the continuous space with a finite set of discrete values.

The core idea is to partition the continuous range of each state and action variable into a finite number of intervals, or "bins." Each bin represents a discrete state or action, and the agent interacts with the environment using these discretized values. For instance, consider the degradation levels of two components, S_1 and S_2 , as depicted in the provided image. Instead of treating the degradation levels as continuous values, we can discretize them into five levels: Healthy1, Healthy2, Healthy3, Healthy4, and Failed State.

Specifically, the discretization scheme used in this example is defined as follows:

- **Healthy1 (0):** Low degradation ($deg \leq 2.5$).
- **Healthy2 (1):** Moderate degradation ($2.5 < deg \leq 5$).
- **Healthy3 (2):** Higher degradation ($5 < deg \leq 7.5$).
- **Healthy4 (3):** Significant degradation ($7.5 < deg < 10$).
- **Failed State (4):** Critical degradation ($deg > 10$).

This discretization scheme maps the continuous degradation levels to discrete integers from 0 to 4. The thresholds D_1 and D_2 in the original description implicitly refer to the boundaries of these discrete intervals. A failure state is triggered when the degradation of either component exceeds the threshold corresponding to the "Failed State (4)."

By employing discretization, the Q-learning algorithm can now operate on a finite state space, where each state is represented by a combination of the discrete degradation levels of the two components. For example, a state could be $(S_1 = 1, S_2 = 2)$, indicating that component S_1 is in the Healthy2 state and component S_2 is in the Healthy3 state. The Q-table then becomes a finite-dimensional table, where each entry $Q(s, a)$ represents the estimated reward for taking action a in the discretized state s .

4.2 Q-learning Overview

The Q-learning algorithm is employed to optimize the maintenance policy. The key steps of the algorithm are:

- Initialize the Q-table (states as rows, actions as columns).
- Define hyperparameters: learning rate α , discount factor γ , and exploration rate ϵ .
- Start in an initial state S .
- Repeat until terminal state:
 - Choose action A using ϵ -greedy policy.
 - Perform action A , observe reward R and new state S' .
 - Update Q-value using the Bellman equation.
 - Set $S \leftarrow S'$.
- Convergence: Q-values stabilize, and the optimal policy is learned.

5 Simulation Study

Define the following maintenance actions clearly:

- **#act = 0** No action required
- **#act = 1** Corrective Maintenance

- **#act = 2** Preventive Maintenance
- **#act = 3** Opportunistic Maintenance

In this section, a condition-based opportunistic maintenance (CBOM) policy is investigated for a multi-component degrading system, incorporating common cause failures. We formulate a maintenance optimization model under the framework of Semi-Markov Decision Processes (SMDP), considering both economic and failure dependencies among the components.

Unlike traditional approaches, where optimization models rely on heuristics or conventional algorithms such as the Iterative Genetic Algorithm (IGA), Common Cause Shock Algorithm (CCSA), and Evolutionary Algorithm (EA), we apply a Q-learning approach to determine the optimal maintenance thresholds. Our approach allows for the dynamic adaptation of the maintenance policy over time, learning from interactions between system degradation and maintenance actions.

5.1 Classical Approaches

In traditional approaches to condition-based opportunistic maintenance (CBOM), various optimization models and algorithms have been proposed to address the challenges of multi-component degrading systems with common cause failure. These methods typically aim to optimize the maintenance decision-making process while considering both the degradation and failure dependencies among components. Some classical approaches include:

- **Iterative Genetic Algorithm (IGA):** This approach uses evolutionary techniques to optimize maintenance thresholds by generating a population of solutions and iteratively evolving them based on fitness functions, which are typically designed to minimize maintenance costs and downtime.
- **Common Cause Shock Algorithm (CCSA):** This algorithm focuses on modeling the impact of common cause failures (CCF) on maintenance decisions. It addresses the challenges of coordinating maintenance actions in systems where multiple components fail simultaneously due to shared shocks.
- **Evolutionary Algorithm (EA):** EA is a broader class of algorithms that uses mechanisms inspired by biological evolution, such as mutation, selection, and crossover, to find optimal maintenance strategies. It is often employed to handle the complexity and non-linearity of multi-component systems and failure dependencies.

These classical algorithms focus on determining the optimal preventive maintenance (PM) and opportunistic maintenance (OM) thresholds for components, seeking to minimize the total cost of maintenance actions while ensuring system reliability. However, each of these methods has its limitations, particularly in their ability to adapt dynamically to changing system conditions. In contrast, our approach leverages Q-learning to optimize the maintenance policy in a more flexible and adaptive manner, addressing these limitations.

5.2 Q-learning Implementation

To evaluate the performance of the Q-learning-based maintenance policy, we simulate the degradation processes and maintenance actions over a series of time periods. The Q-learning algorithm is used to optimize the maintenance thresholds, including preventive maintenance (PM) and opportunistic maintenance (OM) thresholds, based on the state of the system and the degradation levels of the components. The goal is to minimize the overall cost while ensuring that the system remains operational.

We have implemented the following key steps:

- Discretize the degradation levels of both components into finite states.
- Define a set of maintenance actions: corrective maintenance (CM), preventive maintenance (PM), opportunistic maintenance (OM), and do nothing.
- Use the Q-learning algorithm to update the Q-values based on the rewards observed from maintenance actions and system states.

5.3 Convergence of the Algorithm

To demonstrate the effectiveness of the Q-learning approach, we present the convergence behavior of the algorithm. As shown in Figure 2, we observe the decay of the exploration parameter ϵ over time, which allows the algorithm to gradually shift from exploration to exploitation. The Q-values stabilize after a sufficient number of iterations, indicating that the algorithm has converged to an optimal maintenance policy.

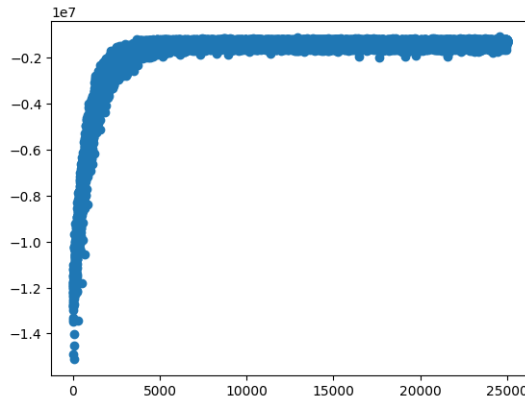


Figure 1: Convergence of the Q-learning algorithm: Epsilon decay and stabilization of Q-values.

5.4 Q-table

The Q-table, which stores the learned values for each state-action pair, is a crucial part of the Q-learning process. After the algorithm has converged, the Q-table reflects the optimal maintenance policy for each state. As illustrated in Figure 3, the table consists

	0	1	2	3	4
0	(0, 0)	(0, 0)	(0, 2)	(0, 2)	(0, 1)
1	(0, 0)	(0, 0)	(2, 3)	(0, 2)	(0, 1)
2	(2, 0)	(2, 3)	(2, 3)	(2, 3)	(0, 1)
3	(2, 0)	(2, 0)	(2, 3)	(2, 2)	(2, 1)
4	(1, 0)	(1, 0)	(1, 0)	(1, 2)	(1, 1)

Figure 2: Q-table after convergence, showing optimal maintenance actions, Where the rows are the first component and columns are second component and for each pair the cell correspond to optimal action for each component in the (first, second) component.

of values that represent the expected rewards for different maintenance actions across various system states.

The results show that the Q-learning approach effectively determines the optimal maintenance thresholds, balancing the costs of maintenance and system downtime. so for each pair of state we have our optimal action in the table instead of creating thresholds but we could have approximate those thresholds as we discuss further.

6 Results

The proposed Q-learning-based maintenance policy was evaluated through simulations and compared with existing maintenance strategies. The results highlight how this approach optimizes maintenance decisions while balancing cost-effectiveness and system reliability.

6.1 Degradation Trends and Maintenance Efficiency

The degradation patterns of system components under different maintenance strategies were analyzed. As shown in Figure 3, degradation levels increase over time, with maintenance actions helping to reset them at key intervention points. The Q-learning policy identifies maintenance moments that prevent excessive degradation while avoiding unnecessary interventions. Key thresholds, including minor degradation, high degradation, and the failure threshold, provide reference points for decision-making.

The comparison between maintenance and non-maintenance cases (Figure 4) shows that degradation accumulates more rapidly without intervention, leading to earlier failures. When maintenance is applied, degradation is managed more effectively, extending the operational lifespan of components. This suggests that a data-driven maintenance policy can help adjust intervention timing to balance reliability and cost.

6.2 Comparison with Existing Policies

We compare maintenance thresholds and degradation levels between our approach and the traditional method based on the values reported in [5]. The traditional approach sets:

$$O_1 = 5.84, \quad M_1 = 8.38, \quad M_2 = 8.96, \quad O_2 = 7.22$$

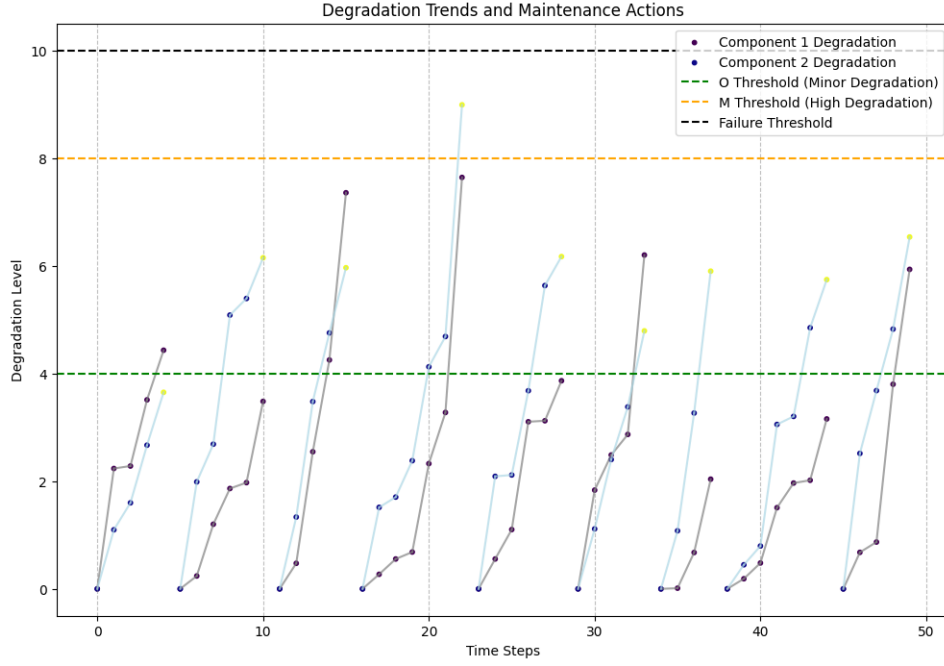


Figure 3: Convergence of the Q-learning algorithm: Epsilon decay and stabilization of Q-values.

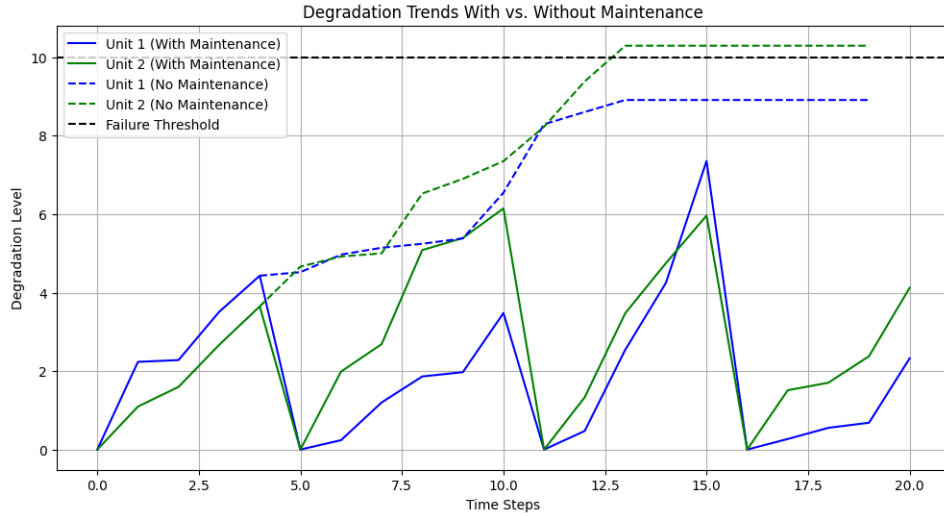


Figure 4: Convergence of the Q-learning algorithm: Epsilon decay and stabilization of Q-values.

In contrast, we approximate our discretization borders as:

$$O_1^* = 4, \quad O_2^* = 4, \quad M_1^* = 8, \quad M_2^* = 8$$

This adjustment is beneficial because:

- The earlier preventive maintenance thresholds (M_1^*, M_2^*) mitigate degradation before it escalates, enhancing reliability.

- Lower opportunistic maintenance limits (O_1^*, O_2^*) ensure proactive intervention, reducing the likelihood of unexpected failures.
- The refined discretization effectively balances preventive and corrective actions, optimizing maintenance efficiency.

Compared to heuristic-based policies, which rely on predefined rules, the Q-learning approach adapts based on observed degradation patterns. While heuristic methods are often straightforward and computationally efficient, their fixed decision rules may not always capture changing system conditions. In contrast, Q-learning refines maintenance decisions over time by learning from past interactions.

6.3 Cost and Time Considerations

A key factor in maintenance decision-making is balancing *inspection frequency and maintenance costs* while ensuring system reliability. Unlike fixed-interval strategies, which schedule maintenance based on predefined thresholds, the Q-learning policy determines intervention points based on real-time system conditions. This leads to:

- **Reduced operational costs**, as maintenance is performed when needed rather than at fixed intervals.
- **Less frequent interventions**, extending the time between major repairs.
- **Better resource allocation**, ensuring maintenance actions are distributed efficiently over time.

These findings align with previous research on reinforcement learning for maintenance, suggesting that adaptive decision-making can complement heuristic approaches, particularly in complex multi-component systems. Rather than replacing heuristic methods, reinforcement learning provides a structured way to optimize maintenance policies. Future work can explore integrating Q-learning with rule-based systems to leverage both adaptability and computational efficiency.

7 Discussion and Conclusion

In this paper, we proposed a novel maintenance policy optimization method based on reinforcement learning (RL) for a multi-component degrading system subjected to common cause failures. By leveraging the Q-learning algorithm, we were able to dynamically optimize the timing and thresholds for preventive maintenance (PM) and opportunistic maintenance (OM) to minimize maintenance costs while maintaining system reliability.

Our approach demonstrates several advantages over traditional methods. Classical maintenance optimization algorithms, such as the Iterative Genetic Algorithm (IGA), Common Cause Shock Algorithm (CCSA), and Evolutionary Algorithm (EA), are often simpler to understand and implement. These methods typically rely on predefined rules or heuristic models, which are relatively easy to set up and execute. However, these classical approaches lack the ability to adapt dynamically to changing system conditions. In

contrast, our RL-based approach is more flexible and capable of learning from interactions between system degradation and maintenance actions over time, which ultimately leads to better optimization.

Despite these advantages, there are some limitations to the proposed method. One key limitation arises from the discretization of the state and action spaces, which is necessary for applying Q-learning to continuous environments. Discretization can lead to a loss of precision, as the continuous nature of system degradation is approximated by discrete intervals. However, this issue can be mitigated by using more advanced RL techniques, such as Deep Q-Networks (DQN) and Double Deep Q-Networks (DDQN), which allow for more precise decision-making while still leveraging the power of RL. These techniques, however, come with an increased computational cost, as they require deep neural networks to approximate the Q-values.

Another significant aspect is that as reinforcement learning techniques evolve, maintenance criteria may become more sophisticated and high-level. Traditional maintenance approaches are often based on simple heuristics or predefined models, but with RL, maintenance strategies can evolve and adapt based on real-time data and system behavior. This shift in maintenance strategy allows for more intelligent decision-making and can lead to substantial improvements in both cost-efficiency and system reliability.

In conclusion, the proposed Q-learning-based maintenance policy offers a significant improvement over traditional methods by enabling dynamic, data-driven decision-making. While the discretization process introduces some loss in precision, this can be addressed through more advanced RL techniques. As reinforcement learning continues to evolve, maintenance criteria will become increasingly advanced, opening the door for smarter, more efficient maintenance strategies in complex systems.

References

- [1] Broek, M. A. U. H., Teunter, R. H., De Jonge, B., & Veldman, J. (2021). Joint condition-based maintenance and load-sharing optimization for two-unit systems with economic dependency. *European Journal of Operational Research*, 295(3), 1119-1131.
- [2] Wang, J., Qiu, Q., & Wang, H. (2021). Joint optimization of condition-based and age-based replacement policy and inventory policy for a two-unit series system. *Reliability Engineering & System Safety*, 205, 107251.
- [3] Chen, Y., Qiu, Q., & Zhao, X. (2022). Condition-based opportunistic maintenance policies with two-phase inspections for continuous-state systems. *Reliability Engineering & System Safety*, 228, 108767.
- [4] Shi, H., Zhang, J., Zio, E., & Zhao, X. (2023). Opportunistic maintenance policies for multi-machine production systems with quality and availability improvement. *Reliability Engineering & System Safety*, 234, 109183.
- [5] Li, M., & Wu, B. (2024). Optimal condition-based opportunistic maintenance policy for two-component systems considering common cause failure. *Reliability Engineering & System Safety*, 250, 110269.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Optimizing Maintenance Scheduling in a Scrap-Based Steel Production Line Using Reinforcement Learning

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

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

Abstract

This paper suggests a Reinforcement Learning (RL)-based optimization method for shredder machine maintenance and inspection planning in a steel production line with scrap as raw material. The method incorporates Remaining Useful Life (RUL) into decision-making for proactive maintenance considering real production conditions like production rate, buffer level, and RUL. The production line is represented as a multi-component system with dependencies and uncertainties, and a simulation platform used to estimate performance. The result indicates that the incorporation of RUL improves maintenance effectiveness in terms of reducing downtime and operation costs. With this approach, it is demonstrated that Reinforcement Learning, based on RUL data, can significantly enhance maintenance control in production systems.

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

¹nhaerian@ut.ac.ir

²gholami.z78@ut.ac.ir

³a.Safari@ut.ac.ir;

⁴fhaghighi@ut.ac.ir

^{1,2}These authors contributed equally to this work

1 Introduction

Steel is one of the most commonly used metal today, playing a significant role in such sectors as construction, automobile, and manufacturing. Steel production can be carried out using iron ore (primary steel) or recycled scrap (secondary steel). Steelmaking using iron demands a substantial amount of energy to produce and also has a significant environmental footprint. It is estimated that the steel industry is responsible for producing around 7-9% of global carbon dioxide (CO₂) emissions, which contributes to the ongoing issue of climate change [1].

Recycling of steel has been a powerful step to counter these environmental problems. Steel recycling involves collecting and processing scrap metal to produce new steel products. Production costs are reduced, energy is conserved, natural resource extraction is reduced, and solid scrap management is improved by utilizing recycled steel. Despite all these benefits, recycling of steel is still not maximized, and cost and quality of final product are some of the issues that have kept it from being massively used [2].

New technology such as shredder machines, remedies these problems in steel production from scraps. Shredder machines shred metallic scrap into tiny pieces, clean out impurities, and prepare the suitable material ready for the production of steel. Scrap is shredded to conserve energy, facilitate easy transport, and enhance the quality of the product. Their performance also contributes to a great extent towards the overall process performance. Hammers in shredders which play a very key role in the process sustain great wear, hence causing them to degrade at a very fast pace. This sort of hammer wear can lower production capacity and have a negative impact on productivity at the production lines. Proper maintenance of such parts is therefore immensely needed in a bid to achieve long-term productivity and lower the cost of running.

A steel production line is made up of interlinked modules that create a complex system, and failure in one module may impact the performance of the others. It is hence challenging to maintain such systems. Additionally, taking down the system for maintenance requires advanced planning to ensure demand versus cost optimization. This complexity has led to research in production line maintenance and planning [3].

1.1 Types of Maintenance Approaches Applied in the Study

In this paper, we distinguish four principal maintenance approaches, each of which can be applied under different circumstances depending on system requirements and available monitoring data:

Time-Based Maintenance (TBM):

TBM is where the maintenance and inspection work are performed at planned, scheduled time intervals regardless of the condition of the equipment. While TBM allows for regular maintenance work to be performed at a standard interval, it results in unnecessary visits when the equipment is not in need of service at due time. As stated by Ferreira Neto et al. [1], In spite of industry interest in Condition-Based Maintenance (CBM) techniques,

TBM continues to be the most common maintenance practice, which involve scheduling inspections at periodical time interval.

Condition-Based Maintenance (CBM):

CBM is dynamic in the aspect that it uses real-time, continuous monitoring of the actual condition of the equipment. The method utilizes a variety of sensors and diagnostic equipment to take measurements of performance parameters and indicators. Maintenance is suggested depending on the condition of the equipment at the time, with the potential for immediate action as soon as certain thresholds are met. CBM not only enhances the operational efficiency but also lowers the cost on unwanted maintenance because it triggers service action only when necessary.

Corrective Maintenance (CM):

CM, takes place only after the equipment fails. While it is straightforward, it can lead to prolonged downtimes, emergency repair costs, and potential collateral damage. In high-throughput industries, such as steel manufacturing, an unscheduled shutdown is frequently more expensive than a planned intervention, making CM less desirable unless failures are rare or minimally disruptive.

Preventive Maintenance (PM):

PM is an overarching term that includes any proactive steps taken to reduce or avert breakdowns before they occur. TBM and CBM may each be regarded as a subtype of PM: TBM follows a strict schedule, whereas CBM is driven by condition data. In most modern industrial setups, PM is crucial for boosting availability, reducing unplanned downtime, and extending equipment life, particularly in systems subject to heavy or variable wear.

Accurately identifying and defining system dynamics and selecting an appropriate modeling technique are the foundation of maintenance optimization. For multi-component systems, dynamic CBM methods based on the Markov Decision Process (MDP) have been suggested [4]. Although model-based methods can solve MDPs, they require a complete model of the environmental components' transition matrices, which could be infeasible for practical applications. When such matrices do not exist and component deterioration must be accounted for, model-free methods such as RL can then be used as an alternative. Mathematical modeling of system behavior is essential in real-time control and decision-making [5]. Aside from component interdependencies in manufacturing systems, uncertainties in production and maintenance also render it challenging to model such multi-component systems.

There are some existing research articles on sophisticated maintenance methods in steel production. In the article by Ferreira Neto et al. [1], a methodology with a Deep Reinforcement Learning (DRL) was constructed to enhance the inspection and preventive maintenance policy of a shredder in a scrap-based steel production process. In their work, the scrap-based steel production line was represented as a multi-component system with dependencies between components, consisting of the first workstation, which includes the shredder machine, an intermediate buffer, and the second workstation where the remaining steel production operations are carried out. The DRL-based maintenance optimization approach recommended the appropriate timing for inspection and maintenance.

nance tasks based on the monitoring condition of the production line, including machine productivity, buffer levels, and production demand. However, their model largely took into account the failure modes of the first workstation and did not directly consider the condition and potential failure of the second workstation.

Unlike these previous studies, our research extends the scope further by actively modeling the second workstation’s failures and integrating RUL information into the decision-making process. RUL is valuable information that provides knowledge about the health of a system. Predicting the time to system failure or knowing how much longer the system can operate, is valuable in best planning [4]. This metric is a key one in system management. This allows the maintenance planners to make more proactive and context-aware maintenance scheduling decisions, which can significantly reduce downtime and overall operating costs.

For clarity, a list of abbreviations used throughout this paper is provided in Table 1.

The rest of this paper is organized as follows. Section 2 explains the mathematical modeling and simulation of the production system, Section 3 describes the RL framework and Section 4 Presents the results and draws conclusions from them.

Table 1: List of abbreviations

Abbreviation	Full Term
TBM	Time-Based Maintenance
CBM	Condition-Based Maintenance
CM	Corrective Maintenance
PM	Preventive Maintenance
MDP	Markov Decision Process
RL	Reinforcement Learning
DRL	Deep Reinforcement Learning
MD	Maintenance Duration
MC	Maintenance Cost

2 Mathematical Modeling and Simulation

2.1 System description

The environment used in this study is based on that of Ferreira Neto et al. [1] with modifications and adjustments to add more practical features and dynamic elements. The study focuses on a scrap-based steel production line consisting of multiple interconnected workstations. The system consists of three main components: a first workstation responsible for separating impurities and shredding scrap, a second workstation where the remaining steelmaking processes take place, and a buffer between them to control fluctuations between production and demand. The buffer level also serves as an indicator to control the production rate. (Figure1)

The first key workstation is the shredder, which processes steel scrap by reducing it to smaller pieces. The scrap is a mixture of metals, plastics, foam, rubber, wood, sand, and

gravel, and is separated into ferrous and non-ferrous materials by a magnetic separator. The ferrous scrap is then transferred to the buffer for later use in steelmaking.

The shredder consists of a large rotor with multiple crushing hammers that are essential for breaking down the scrap. These hammers degrade over time. As the hammers degrade, their contribution to the shredders production decreases. When a hammer reaches the failed state, it no longer contributes to the shredders productivity.

The production rate of the shredder at time t , denoted P_t , depends on the degradation of the hammers. The shredder operates in two possible paces based on the buffer level:

- T1 (high production pace)
- T2 (low production pace)

The degradation process of the hammers follows a Gamma distribution and increases based on the machine's production pace, with the degradation at time $t + \Delta t$ given by:

$$D_i(t + \Delta t) = D_i(t) + G_a(\alpha, \Delta t, \beta) \quad (2.1)$$

Here, α and β are the shape and scale parameters of the Gamma distribution that governs the degradation process, and Δt represents the operational time step. When the shredder operates at rate T1, the parameters are α_{T1} and β_{T1} . During operation at rate T2, the parameters change to α_{T2} and β_{T2} . This relationship between the hammer condition and the production rate is modeled in the following way.

2.2 Mathematical Modelling

The production rate of the shredder P_t at time t is modeled by the following equation:

$$P_t = \begin{cases} c_1 \rho_t \sum_{i=1}^n P_i(t) & \text{if operating in T1} \\ c_2 \rho_t \sum_{i=1}^n P_i(t) & \text{if operating in T2} \end{cases} \quad (2.2)$$

Where:

- $P_i(t)$ is the production rate of hammer i at time t .
- n is the total number of hammers in the shredder.
- c_1 and c_2 are constants that define the production pace: c_1 for the high production pace T1, and c_2 for the low production pace T2.
- ρ_t is the ferrous content of the scrap at time t , which influences the shredders production rate.

This equation reflects how the shredder's production rate is influenced by the degradation of the hammers, composition of the scrap and the production pace. The constants c_1 and c_2 determine how the production pace affects the shredders performance, with higher production intensifying the hammer degradation process.

ρ_t follows a normal distribution. This variability impacts the shredder's performance and adds uncertainty to the system.

The system also uses the intermediate buffer to store the shredded material, which helps maintain the flow of material to the next workstation when the shredder's performance fluctuates. As mentioned earlier, there are two paces (T_1 or T_2) for shredding, depending on the buffer level. Initially, the system starts at T_1 . When the buffer level reaches h_1 , the production pace decreases to T_2 , and if the buffer level reaches h_2 , the pace increases back to T_1 .

As the hammers degrade over time, they pass through four states: new, minor defective, major defective, and failed. Each state reduces the contribution of that hammer to the total production. For example, a new hammer contributes 100%, a minor defective hammer contributes 80%, a major defective hammer contributes 60%, and a failed hammer contributes 0%. These percentages are determined by the degradation level of the hammer. A degradation less than 0.2 is considered new, between 0.2 and 0.3 is minor defective, between 0.3 and 0.4 is major defective, and greater than 0.4 is out of service.

The buffer level at each time step is calculated from the difference between production and demand ($P_t - d_t$). If production exceeds demand, the surplus is stored in the buffer. If demand exceeds the sum of the buffer storage and production, there will be unmet demand, which is accounted for in the cost section.

The second workstation is responsible for the remaining steelmaking processes and consists of various interconnected components. Any failure in this workstation leads to its downtime. To prevent this, periodic maintenance is performed every T_s hours on this workstation, and the duration of that is T_d hours.

Additionally, in case of failure (RUL=0), CM is performed for a duration F_s , which follows an exponential distribution with a mean of λ_{F_s} .

A preventive maintenance is also planned so that if the second workstation stops, whether for periodic inspections or CM, the first one also halts for PM.

During second workstation maintenance periods, the system is stopped, and the demand is considered zero. At other times, the demand d_t follows a normal distribution (F_d) with a mean of μ_{d_t} and standard deviation of σ_{d_t} . The demand can be formulated as:

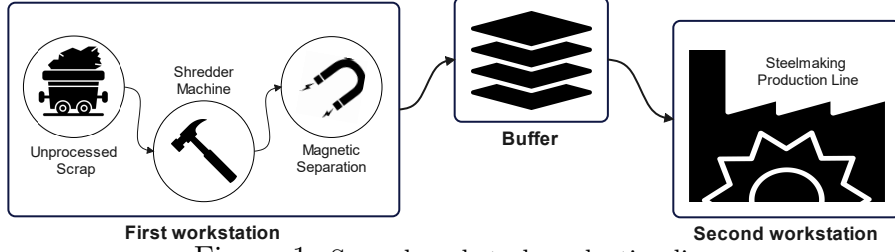
$$d_t = \begin{cases} 0, & \text{during workstation stoppages} \\ \text{Following } F_d, & \text{otherwise} \end{cases} \quad (2.3)$$

For the second workstation, the RUL is also calculated at each time step. The lifetime of the second workstation follows an exponential distribution with mean of λ_{F_f} . The exact RUL for time t is calculated as follows:

$$RUL_t = lifetime - t \quad (2.4)$$

2.3 Maintenance and Cost Modeling

Both PM and CM actions on the first workstation involve stopping the shredder to inspect



and replace defective (minor and major) and failed hammers at a cost of c_r per hammer. Each shutdown incurs a cost c_i for personnel, materials, and tasks like inspections, and cleaning. Additional costs include the inventory carrying cost c_s for storing crushed scrap in the buffer and the cost c_l per unit of unmet demand.

The maintenance duration (MD) varies based on the type of maintenance performed. Accessing the hammers requires a time T_P . Each hammer replacement takes T_u . CM after a failure necessitates additional time T_a , following an exponential distribution (F_a) with a mean of λ_{F_a} . A failure also incurs a cost c_f .

The Maintenance Duration is calculated as follows:

$$MD = \begin{cases} T_P + N_r T_u, & \text{if PM} \\ T_P + N_r T_u + T_a, & \text{if CM} \end{cases} \quad (2.5)$$

where N_r is the number of replaced hammers.

The aim of this paper is to suggest a CBM policy with RL for the shredder machine, that determines the optimal time to inspect based on system monitoring information in order to realize the minimum long-term maintenance costs per unit. Although the second workstation is maintained periodically, its maintenance costs are not actually calculated; only unmet demand is accounted for when production is stopped due to material shortage.

At time t , the operating cost of the production line is given by $c_s b_t + c_l U_t$. The quantity of unmet demand at time t (U_t) can be determined as follows:

$$U_t = \begin{cases} d_t - P_t - b_t, & \text{if } d_t > P_t + b_t \\ 0, & \text{otherwise} \end{cases} \quad (2.6)$$

The Maintenance Cost (MC) associated with a maintenance intervention can be computed as follows:

$$MC = \begin{cases} c_i + c_s \int_0^{MD} b_t dt + c_r N_r + c_l \int_0^{MD} U_t dt, & \text{if PM} \\ c_i + c_s \int_0^{MD} b_t dt + c_r N_r + c_l \int_0^{MD} U_t dt + c_f, & \text{if CM} \end{cases} \quad (2.7)$$

Where $\int U_t$ and $\int b_t$ represent, respectively, the total unmet demand and the total buffer content during the maintenance duration (MD). Let π denote the shredder's maintenance policy. The total cost up to time t is:

$$C(t, \pi) = c_l \int_0^t U_t dt + c_s \int_0^t b_t dt + c_r N_r(t) + c_i PM(t) + (c_f + c_i) CM(t) \quad (2.8)$$

where $N_r(t)$ is the number of replaced hammers up to time t , and $PM(t)$ and $CM(t)$ are the number of PM and CM interventions up to t .

The optimal maintenance policy π^* that minimizes the long-term cost per unit of time is:

$$\pi^* = \min_{\pi} \left\{ \lim_{t \rightarrow \infty} \frac{C(t; \pi)}{t} \right\} \quad (2.9)$$

2.4 Simulation

A discrete event simulation model was developed to simulate the scrap-based steel production process, discretizing continuous operation in terms of a one-hour time step. The model uses the Monte Carlo (MC) method to simulate hammer degradation, heterogeneity of scrap, demand variation, failure incidents, and uncertainty in maintenance. By hammer and machine status update, the model calculates the critical variables like P_t and b_t . The efficiency of the maintenance policy is quantified through the operational and maintenance expenditures accumulated over the simulation period. To provide clarity, the parameters and their values used in the simulation are summarized in Table 2.

3 Reinforcement Learning and Optimization

To determine the optimal maintenance policy, we utilize RL. This approach allows the agent to learn the dynamics of the environment through model-free interactions. The agent can effectively adapt and improve its decision-making over time by sampling and experiencing various trajectories composed of sequences of consecutive states, actions, and rewards. Through trial and error, it learns to optimize costs, which we provide as penalties, to find the optimal policy. As mentioned earlier, the first step is to define the problem as a Markov Decision Process.

3.1 Markov Decision Process Formulation

The problem is formulated as a Markov Decision Process (MDP), which is defined by a tuple (S, A, R) , where:

- S is the set of states,
- A is the set of actions, and
- R is the reward function,

Table 2: 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	10
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	0.1 un./ton.hours
c_l	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	35 ton

3.1.1 State Definition

State definition is a crucial step in formulating the Markov Decision Process (MDP). The agent learns the environment's behavior and determines the optimal action by observing the defined system states. The system states s_t should encompass both machine- and system-level data, represented as a vector containing the necessary information for the agent to understand the shredder's degradation process and production line dynamics.

For the considered system, the instantaneous shredder production rate P_t , buffer level b_t , and RUL_t can represent the system's dynamics. P_t informs the agent about the shredder's degradation, b_t helps understand the dynamics between workstations, and RUL_t indicates the second workstation's operation.

The system state vector s_t is defined as:

$$s_t = [P_t, b_t, RUL_t] \quad (3.1)$$

All variables are accessible and trackable in real life. Flowmeters can measure P_t

and d_t , while a level meter measures b_t . These variables serve as proxies for system and machine deterioration, as b_t and P_t correlate with shredder health, and RUL_t indicates the second station’s operational state.

While continuous, these variables are discretized via state aggregation. Despite information loss during discretization, this is done because we are using a Q-learning based algorithm that requires a discrete state and action space. Limiting consideration to a few discrete deterioration states simplifies decision-making. Practical applications often require discrete states due to economic and technical considerations, as continuous condition monitoring can be infeasible or costly.

State aggregation was performed as follows: P_t is divided into four states: high productivity rate (0), intermediate productivity rate (1), low productivity rate (2), and non-producing (3), indicating shredder failure. b_t has five states: full level (0), upper-intermediate level (1), intermediate level (2), pre-intermediate level (3), and low level (4). The RUL state is determined based on an 80% system lifetime. If RUL_t is 80% or higher of the lifetime, the second workstation is in the high RUL state (0), indicating that it is relatively far from failure. If the RUL is below this value, the second workstation will be in the low RUL state (1), indicating that it is near failure and requires more critical attention. The aggregation creates a three-dimensional state space.

3.1.2 Action Definition

The action taken by the agent determines whether the shredder should be shut down for preventive maintenance (PM). At each time step t , the agent can choose an action a_t as follows:

$$a_t = \begin{cases} 0, & \text{Keep the system running} \\ 1, & \text{Shutdown for maintenance} \end{cases} \quad (3.2)$$

There are two scenarios which are automatically activated by the simulation environment based on industry standards. The first scenario arises when the buffer level hits critical thresholds (K or 0), triggering an automatic shutdown of the first workstation for PM. The second scenario is triggered by a machine failure, which immediately activates CM on the shredder machine. The reward in such cases is calculated to penalize the agent for allowing the failure to occur. The action options are constrained for the first workstation, focusing on developing a maintenance policy specifically for the shredder machine.

3.1.3 Reward Function Definition

The reward function r_t should align with the maintenance objective, as the agent’s mission is to maximize the expected cumulative reward. It is expressed as a stepwise objective function, enabling the agent to learn a suitable maintenance policy. The r_t includes all inspection, maintenance, and operating costs associated with maintenance actions in the

shredder machine and is defined as:

$$r_t = -c_l U_t - c_s b_t - a_t (c_i + c_r(N_r(t) - N_r(t-1))) - W_{CM(t) > CM(t-1)} (c_f + c_i + c_r(N_r(t) - N_r(t-1))) \quad (3.3)$$

At each time step t , the reward function returns a combination of the instantaneous operating cost, including the cost of unmet demand ($c_l U_t$) and inventory carrying cost ($c_s b_t$), with the cost associated with maintenance actions. The action a_t is a binary variable equal to 1 if the agent chooses to perform a PM action and 0 otherwise. When a_t is 1, the cost of that PM action is included in the reward function. W_A is a Bernoulli variable equal to 1 if event A is realized and 0 otherwise. If the cumulative amount of CM actions at time t is greater than at time $(t-1)$, indicating a system breakdown between times $t-1$ and t , $W_{CM(t) > CM(t-1)}$ is 1, and the cost of that CM action is included in the reward function. The cost of replacing defective and failed hammers in both PM and CM is calculated by the difference between $N_r(t)$ and $N_r(t-1)$, the cumulative quantity of hammers replaced at times $t-1$ and t . The negative sign in all terms indicates that maximizing r_t implies minimizing the total cost.

3.2 Q-learning and Optimization

Q-learning is a model-free RL algorithm that aims to find the optimal action-value function, $Q(s, a)$, which represents the expected cumulative reward of taking action a in state s and following the optimal policy thereafter. Since we do not have the model of the environment, meaning that we do not know the transition probability matrix and the behavior of the environment, we need to sample. This means that we allow the agent to move in the environment and explore different paths, and optimize its policy with the information it obtains. Q-learning is an off-policy algorithm, meaning that the policy used for exploration and sample generation is different from the policy being optimized. This algorithm uses the temporal difference (TD) method to update the Q-table for each state-action pair.

The Q-learning algorithm can be described as follows:

1. Initialize the Q-table with arbitrary values.
2. For each episode:
 - (a) Initialize the state s .
 - (b) For each step:
 - i. Select action a in state s using an exploration policy (e.g., ϵ -greedy).
 - ii. Take action a , observe reward r and next state s' .
 - iii. Update the Q-value for state s and action a using the following equation:
$$Q(s, a) = Q(s, a) + \alpha[r + \gamma \max_{a'} Q(s', a') - Q(s, a)] \quad (3.4)$$
 - iv. Set $s = s'$.

4 Results And Conclusion

4.1 Results

The Q-learning algorithm used in this research depends on a tabular approach for approximating action-value functions. As seen in table 3, there are 50,000 training episodes in total, with each consisting of a cap of 500 time steps. The learning rate parameter is set to 0.01 and the discount factor parameter is set at 0.99 in order to ensure that the agent does not overlook due credit for remote rewards. The epsilon-greedy policy is initialized at 0.7, and this decreases gradually by a factor of 0.9995 each episode, to a minimum exploration rate of 0.01. This gives a balance of exploration and exploitation, where the agent first explores and then exploits more and more learned policy.

The training procedure includes 1,000 interactions under the simulation framework. During every interaction, the agent updates the Q-values following the actions that were taken and the rewards it received. The model's action-selection process takes a typical Q-learning approach with the action maximizing the Q-value being selected by a probability under the epsilon-greedy policy.

Figure2 illustrates the moving average reward plot over 50,000 episodes. The initial moving average reward is very low at around -2 million and rises gradually as the agent learns from its experience. The increase indicates the agent is improving over time by learning from the environment and refining its decision-making process.

In the episodes, variations in the reward values are observed, which is typical in RL as the agent experiments with different methods and activities. Such variations mean that even during learning, the agent occasionally performs sub-optimal actions that result in temporary dips in the reward.

As the episodes increase, the moving average reward becomes stable and trends upwards consistently. This indicates that the agent had learned an optimal or near-optimal policy, as evidenced from the higher and more stable rewards. At the end of the 50,000 episodes, the agent consistently improves, and the moving average reward converges to -1.1 million.

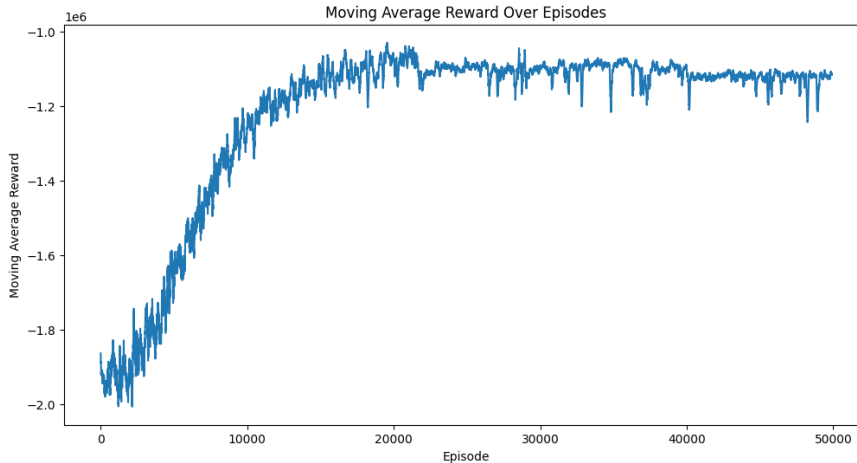


Figure 2: Moving Average Reward

Table 3: Q-learning parameters

Par	Value
Number of episodes	50000
Learning rate	0.01
Discount factor	0.99
<i>epsilon</i>	0.7
<i>epsilon</i> -decay rate	0.9995

Once verified that the training process is stable, the second important step is to verify how well the produced maintenance policy works. That is to determine the optimal maintenance action in each state-space configuration by examining a heatmap.

A Q-value heatmap is an efficient way to represent the performance of an RL agent, particularly for action choice. Each data point in figure 3 is a state-action tuple and is denoted by its associated Q-value, the forecasted future reward for taking some action in some state. The bold value in the heatmap is the maximal Q-value that is associated with the best choice action from what the agent has learned so far.

For each combination (P, B, RUL) , there are two Q-values for “action = 0” (Keep the system running) and “action = 1” (Shutdown for PM). Because a more negative Q-value means higher long-term cost, we look at both the numerical values and the color shading to identify the main patterns:

Lighter shades (yellowish) indicate less negative values, hence lower long-term cost. Such cells usually correspond to “high buffer,” “high RUL,” and “good production.” In these states, continuing to run the system (action 0) tends to be more cost-effective, as long-term costs are relatively small.

Darker shades (orange or red) indicate more negative values and thus higher long-term costs. They typically appear when the buffer is empty, or the RUL is low (high failure risk), or the production is already low. Under these conditions, repair (action 1) is often more economical in the long run—although in some cells you may still see that continuing operation is better, it’s common that unexpectedly severe failures are very costly, which makes earlier maintenance worthwhile.

In high production and high buffer level states, continuing operation is usually more beneficial.

In medium and low production states with low buffer levels, stopping for maintenance might be more beneficial.

When the buffer level is high and the RUL is low, the agent typically decides not to take action. This is because it predicts that the second workstation is likely to fail soon. Initiating maintenance on the first workstation in this situation would result in unnecessary costs.

The weaker the production state and the lower the RUL, the stronger the incentive to carry out repairs.

A key change from Ferreira Neto et al. [1] approach is that now we rely on RUL. This replacement means the agent can leverage (Buffer, RUL) information to decide maintenance timing more effectively. RUL is a better predictor of imminent failure risk than demand.

		Q-value of the state-action pairs										Color scale
		P = 0		P = 1		P = 2		P = 3				
		Action		Action		Action		Action				
B	RUL	0	1	0	1	0	1	0	1			
0	0	-1077073	-1078275	-1077736	-1081775	-1074836	-1082935	-83408.3	-83196.4			
0	1	-1077088	-1077104	-1080075	-1080903	-1071911	-1072792	-36326.1	-42116.8			
1	0	-1083481	-1084446	-1087981	-1086667	-1096232	-1095600	-1064890	-1064756			
1	1	-1085994	-1085232	-1089503	-1088785	-1098926	-1098690	-960270	-960821			
2	0	-1087247	-1084271	-1091270	-1086851	-1100465	-1095532	-1113195	-1113174			
2	1	-1089286	-1086153	-1094603	-1091910	-1102486	-1097171	-1111999	-1111984			
3	0	-1086214	-1084350	-1091964	-1086307	-1097929	-1098928	-1113657	-1113987			
3	1	-1088022	-1087025	-1089611	-1094912	-1093038	-1099238	-1115912	-1116734			
4	0	-1086593	-1085209	-1087212	-1085238	-1098545	-1097189	-1105932	-1106981			
4	1	-1089067	-1087657	-1087208	-1090855	-1093052	-1099843	-1103143	-1099564			

Figure 3: Action-Value Table

Figure 4 shows the RUL over time, with each point as the RUL at a specific step. The RUL starts high, around 700, and gradually decreases as time passes, which would naturally happen as the second workstation ages.

Towards the end of the graph, specifically between step 90 and step 100, the RUL continues to slowly decrease, which shows the second workstation is approaching its end of life. This slow decrease, represents reflects the continuous degradation of the second workstation over time.

After the second workstation failure, the RUL suddenly rises to 350, most likely showing a maintenance action that had been accomplished to recover the second workstation's condition and extend useful life. Such a recovery in RUL is usual after repairs or replacement of parts.

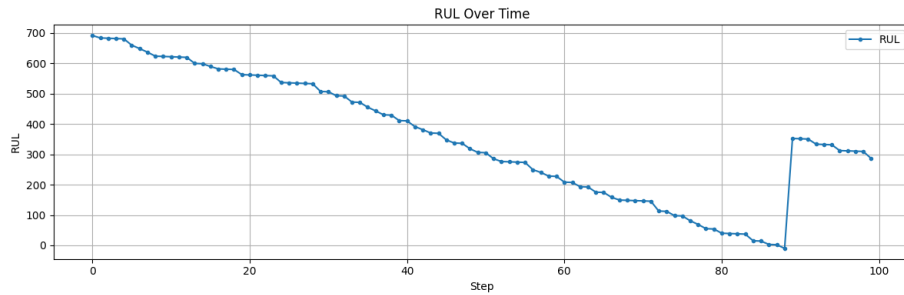


Figure 4: Example of RUL Over Time

In Figure 5, the levels of degradation for the hammers rise progressively. This indicates that the hammers are degrading steadily and can move to the "minor defective", "major defective" or "failed" conditions (degradation moves towards 0.2 and higher). When the degradation drops suddenly, it usually indicates maintenance activities. During maintenance interventions, all hammers with a degradation value higher than 0.2 are replaced. After these activities, the hammers return to a healthier state.

As we can see, at the time marked by the vertical dashed line, maintenance is done. At this time, hammers 3 and 4, whose degradation values were greater than 0.2 were replaced.

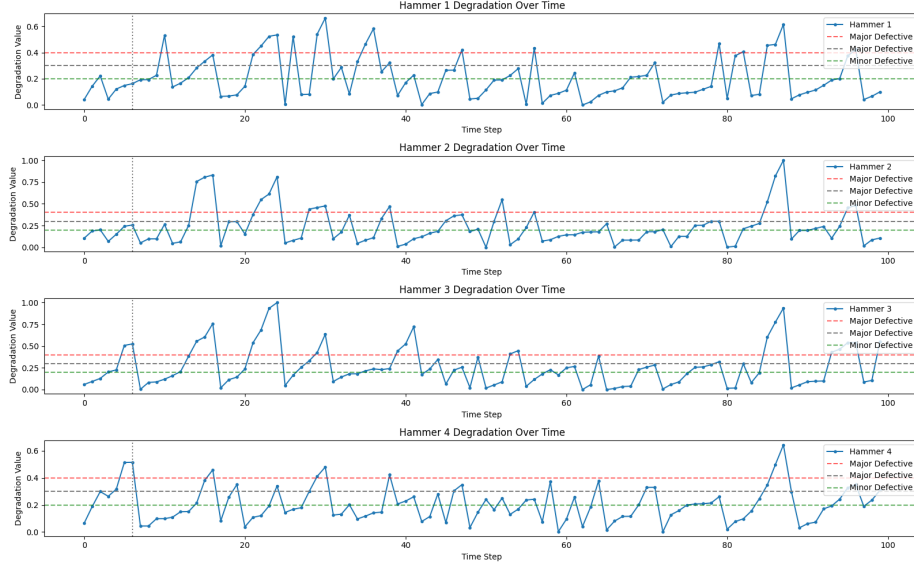


Figure 5: Example of Hammer Degradation Over Time

The following plot 6 shows the relationship between the mean hammer degradation and the production rate of the shredder at time t , denoted by P_t . From the plot, we observe that there is a clear negative correlation between the two variables. As the mean hammer degradation increases, the production rate tends to decrease. This implies that as the hammers degrade over time, their efficiency in performing the shredding activity diminishes, leading to a reduction in production rate. When the hammers are in better condition (lower degradation), the production rate is higher, as expected. Therefore, maintaining the hammers in a well-conditioned state (with minimal degradation) is crucial for achieving a high level of production.

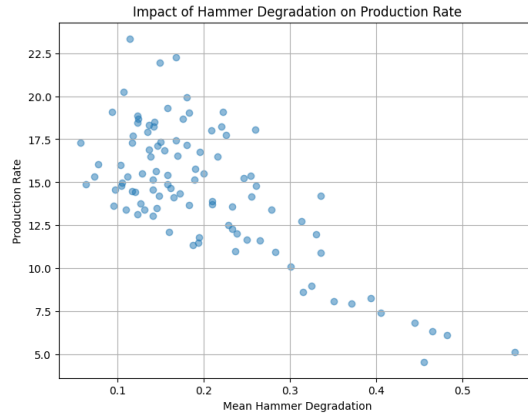


Figure 6: Impact of Hammer Degradation on Production Rate

4.2 Conclusion

This study suggests a Q-learning-based optimization method for determining the optimal inspection and maintenance schedule of a shredder machine in a scrap-based steel produc-

tion line. Unlike the traditional methods, this method employs Q-learning to recommend the best time to perform inspections and maintenance activities based on the real-time monitoring of the key production line variables, including the production rate, buffer level, and RUL of the second workstation.

The scrap-based steel production line is modeled as a complex multi-component system taking into consideration component interdependencies, maintenance time uncertainties, and random production rates. An overall mathematical model of the production line is established, and a simulation model is built to generate data and investigate the performance of the proposed maintenance policy.

The Q-learning algorithm is applied in the simulation model to train the decision-making agent, which learns to select the optimal maintenance actions based on the system states. These system states consist of key factors like production rate, buffer level, and RUL, which influence the decision-making. The Q-learning model is demonstrated to be convergent, and the trained agent exhibits good stability in the prediction of maintenance actions.

This approach allows the agent to efficiently learn an optimal maintenance policy by constantly observing readily available and economic system conditions in real-world applications, i.e., degradation-related indicators. The results show that Q-learning can effectively be applied in dynamic and complex systems like the scrap-based production line for improving maintenance management and overall system performance.

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] Harvey, L. D. D. (2021). Iron and steel recycling: Review, conceptual model, irreducible mining requirements, and energy implications. *Renewable and Sustainable Energy Reviews*, 138, 110553.
- [3] Wang, H. (2002). A survey of maintenance policies of deteriorating systems, *European Journal of Operational Research*, 139(3), 469-489.
- [4] He, R., Tian, Z., Wang, Y., Zuo, M., & Guo, Z. (2023). Condition-based maintenance optimization for multi-component systems considering prognostic information and degraded working efficiency. *Reliability Engineering & System Safety*, 234, 109167.
- [5] Sutton, R.S. and Barto, A.G. (2018), *Reinforcement Learning: An Introduction*, 2nd edition, MIT Press.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Efficiency of Ranked Set Sampling in Mean Residual Lifetime Estimation

Jabari Koopaei, L.¹ Zamanzade, E.²

¹ Department of Statistics, Vali-e-Asr University of Rafsanjan, Rafsanjan, Iran

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

Abstract

The mean residual life (MRL) function is a fundamental concept in survival analysis and reliability. In addition to characterizing the survival distribution, its clear interpretability makes it particularly useful for conveying survival information to non-statisticians. Ranked set sampling (RSS) is a sampling technique designed for situations where obtaining precise measurements of sample units is costly or challenging, while ranking them without relying on exact values is more feasible and cost-effective. However, its practical application is limited by the requirement that each sample unit must be assigned a unique rank. To address this challenge, Frey introduced an innovative variation of RSS, called RSS-t, that integrates tie structure into the ranking process for improved effectiveness. In this paper, we use the RSS-t technique to estimate the MRL function. Then, we compare the proposed estimators with their counterparts in simple random sampling (SRS) and standard RSS using Monte Carlo simulation. Our results show that incorporating tie information enhances statistical inference for the MRL function, reducing the sample size required to achieve a predetermined precision.

Keywords: mean residual life, Monte Carlo simulation, nonparametric estimation, ranked set sampling, ties information

¹l.jabari@sci.ui.ac.ir

²e.zamanzade@sci.ui.ac.ir

1 Introduction

The mean residual life (MRL) is a widely used concept for modeling lifetime data in reliability, survival analysis, and other fields of applied probability and statistics. Suppose X is a non-negative random variable, representing the lifetime of a unit, with finite mean μ and survival function $S(t) = \mathbb{P}(X > t)$. The mean residual lifetime function at time $t \geq 0$ is defined by

$$m(t) = \mathbb{E}(X - t | X > t) = \frac{\int_t^\infty S(x)dx}{S(t)}, \quad (1.1)$$

provided that $S(t) > 0$. If $S(t) = 0$, then $m(t)$ is defined to be zero. The MRL answers an important and easy to interpret question what is the remaining life expectancy of an individual who has survived until time t ? This question is an important topic for many non-statistician scientists because, unlike the survival or hazard functions, it conveys information in units of time rather than a probability scale, which requires careful interpretation.

In recent decades, various methods based on simple random sampling (SRS) technique, have been developed to estimate MRL function. [4] proposed a nonparametric estimation of the MRL function by substituting the survival function with the empirical survival function in Equation (1.1) and analyzing the asymptotic properties of the estimator for a fixed value of t over a finite interval.

In various fields such as medicine, agriculture, and environmental studies, measuring the variable of interest can be costly, time-consuming. However, there is some extra information such as an expert judgment, personal opinions, or concomitant variables with little or no cost, which can be effectively used to obtain a more representative sample of the population. Ranked set sampling (RSS), a sampling technique introduced by [3], is suitable for these situations. One of the challenges of RSS compared to SRS is the need to accurately rank sample units without relying on precise measurements. When the ranker encounters uncertainty in ranking adjacent units within a set, they often randomly select one from the tied units to proceed with RSS. Such situations occur frequently in real-life applications. [2] proposed a novel variation of RSS, called RSS-t, which records and utilizes the structure of tied units for estimation.

In this paper, we propose two RSS-t nonparametric estimators for the MRL function. To compare the performance of these MRL estimators, we conducted a Monte Carlo simulation study with two different classes of tie-generating methods (DPS and TIC). Our results suggest that incorporating tie information into RSS enhances estimator efficiency. In Section 2, we describe RSS and RSS-t in detail. Section 3 describes two MRL estimators using RSS-t. In Section 4, we conduct a comprehensive Monte Carlo simulation to compare the performance of the proposed MRL estimators with their counterparts in SRS and RSS techniques. Some concluding remarks are provided in Section 5.

2 Sampling plan for RSS-t

In this section, we provide a detailed description of how to draw samples using RSS and RSS-t. To obtain a balanced RSS with a total sample size of $n = m \times r$, where m (set size)

and r (cycle size) must be determined before sampling, the following steps are carried out:

[itemsep= 0.8ex]Select m simple random samples with size m from the population. Rank each set from smallest to largest, either virtually or by any cost-effective method. Actual measurements are taken only on the i th smallest units in each i th set, where $i = 1, \dots, m$. Repeat the above steps r times.

The final sample is denoted as $\{X_{[i]j}, i = 1, \dots, m; j = 1, \dots, r\}$, where $X_{[i]j}$ is the i th ordered judgment unit in the j th cycle. The square brackets in the subscription indicate that in-set ranking may contain errors due to subjective judgment or ranking based on an auxiliary variable. This is often referred to as imperfect RSS in the literature. If the ranking process is error-free, the square bracket is replaced with round one, and $X_{(i)j}$ becomes the i th order statistics from j th cycle with set size m , where known as perfect RSS.

The RSS-t technique is similar to standard RSS, with the difference that in RSS-t, along with the values of $X_{[i]j}$, a tie information matrix is also recorded for each cycle as follows:

$$\mathbf{T}_j = \begin{bmatrix} I_{[1]j1} & \cdots & I_{[1]jm} \\ \vdots & \vdots & \vdots \\ I_{[m]j1} & \cdots & I_{[m]jm} \end{bmatrix},$$

where $j = 1, \dots, r$. The i th row of the matrix $\mathbf{T}_j$ shows the tie structure corresponding to $X_{[i]j}$, where the indicator variable $I_{[i]jk}$ is one, if the unit with judgment rank i is tied for rank k in the j th cycle for $k = 1, \dots, m$, and zero otherwise. It is clear that the judgment rank i is always tied with itself, so $I_{[i]ji} = 1$ for $i = 1, \dots, m$. In RSS-t, like RSS, ties are randomly broken when the researcher cannot differentiate between two or more adjacent ranks. However, in RSS-t, the structure of tied units is recorded in matrices for later use. Consequently, RSS-t can be considered a generalized form of RSS. If $\mathbf{T}_1 = \dots = \mathbf{T}_r = \mathbf{I}_m$, where $\mathbf{I}_m$ is the identity matrix of size m , then RSS-t simplifies to RSS. Table 1 illustrates the construction of an RSS-t with a cycle size of $r = 2$ and a set size of $m = 4$. In this table, $X_{i,j,k}$ represents the i th ordered judgment unit of the k th set in the j th cycle. Tied units are connected using ampersands, and the final quantified unit in each row, which is randomly selected from the tied units, is indicated in bold.

3 Estimation of the MRL function

In this section, we present a brief overview of the MRL function estimators based on SRS and RSS techniques, and then introduce our proposed estimators using the RSS-t technique.

Let $X_1, \dots, X_n$ be a simple random sample of size n , drawn from a population with the survival function $S(t)$. A basic nonparametric estimator of the MRL function was first defined by [4] by replacing the survival function in equation (1.1) with the empirical survival function, as follows

$$m_{srs}(t) = \frac{\int_t^{+\infty} S_{srs}(x) dx}{S_{srs}(t)} \mathbb{I}(S_{srs}(t) > 0), \quad (3.1)$$

Table 1: An example of RSS-t with $r = 2$ and $m = 4$.

Cycle Set	Ranked units		Quantified unit	Tie information	
			with RSS-t notation	matrix	
1	1	$X_{1,1,1} \& \mathbf{X}_{2,1,1} \& X_{3,1,1}, X_{4,1,1}$	$X_{[1]1}$	$\mathbf{T}_1 =$	$\begin{bmatrix} 1 & 1 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 1 & 1 & 1 & 1 \end{bmatrix}$
	2	$X_{1,1,2}, \mathbf{X}_{2,1,2}, X_{3,1,2} \& X_{4,1,2}$	$X_{[2]1}$		
	3	$X_{1,1,3}, \mathbf{X}_{2,1,3} \& X_{3,1,3}, X_{4,1,3}$	$X_{[3]1}$		
	4	$X_{1,1,4} \& X_{2,1,4} \& X_{3,1,4} \& \mathbf{X}_{4,1,4}$	$X_{[4]1}$		
2	1	$\mathbf{X}_{1,2,1}, X_{2,2,1} \& X_{3,2,1} \& X_{4,2,1}$	$X_{[1]2}$	$\mathbf{T}_2 =$	$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{bmatrix}$
	2	$X_{1,2,2}, \mathbf{X}_{2,2,2}, X_{3,2,2}, X_{4,2,2}$	$X_{[2]2}$		
	3	$X_{1,2,3} \& X_{2,2,3}, \mathbf{X}_{3,2,3} \& X_{4,2,3}$	$X_{[3]2}$		
	4	$X_{1,2,4} \& X_{2,2,4}, X_{3,2,4}, \mathbf{X}_{4,2,4}$	$X_{[4]2}$		

where $S_{srs}(t) = \frac{1}{n} \sum_{i=1}^n \mathbb{I}(X_i > t)$ is the empirical survival function and $\mathbb{I}(A)$ represents the indicator function on the set A . Noting that $\int_t^\infty \mathbb{I}(X > x) dx = (X - t) \mathbb{I}(X > t)$, it is easy to show that $m_{srs}(t)$ can be rewritten as

$$m_{srs}(t) = \frac{1}{\sum_{i=1}^n \mathbb{I}(X_i > t)} \sum_{i=1}^n (X_i - t) \mathbb{I}(X_i > t), \quad (3.2)$$

given that $\sum_{i=1}^n \mathbb{I}(X_i > t) > 0$, otherwise the entire expression is defined as zero.

Let $\{X_{[i]j}; i = 1, \dots, m, j = 1, \dots, r\}$ be a ranked set sample of size $n = r \times m$, obtained from a population with the survival function $S(t)$. [5] proposed a nonparametric estimator of the MRL function based on RSS. They assumed the ranking mechanism is consistent, meaning that the same ranking mechanism is applied to all sets of size m . Under this assumption, the survival function $S(t)$ is given by $S(t) = \frac{1}{m} \sum_{i=1}^m S_{[i]}(t)$, where $S_{[i]}$ represents the survival function of the i th judgment order statistic in the sample of size m . Based on this, they reformulated the MRL function as follows

$$\begin{aligned} m(t) &= \frac{1}{S(t)} \int_t^{+\infty} \frac{1}{m} \sum_{i=1}^m S_{[i]}(x) dx \\ &= \frac{1}{m} \sum_{i=1}^m \frac{\int_t^{+\infty} S_{[i]}(x) dx}{S(t)} \\ &= \frac{1}{m} \sum_{i=1}^m \frac{S_{[i]}(t)}{S(t)} m_{[i]}(t), \end{aligned} \quad (3.3)$$

where $m_{[i]}(t) = \mathbb{E}(X_{[i]} - t | X_{[i]} > t) = \frac{\int_t^\infty S_{[i]}(x) dx}{S_{[i]}(t)}$. By replacing $S_{[i],rss}(t) = \frac{1}{r} \sum_{j=1}^r \mathbb{I}(X_{[i]j} > t)$, $S_{rss}(t) = \frac{1}{m} \sum_{i=1}^m S_{[i],rss}(t)$, and $m_{[i],rss}(t) = \frac{\int_t^{+\infty} S_{[i],rss}(x) dx}{S_{[i],rss}(t)} \mathbb{I}(S_{[i],rss}(t) > 0)$ with $S_{[i]}$, S and $m_{[i]}$, respectively, in

equation (3.3), we can obtain the estimator of the MRL function based on RSS as follows

$$m_{rss}(t) = \frac{1}{m} \sum_{i=1}^m \frac{S_{[i],rss}(t)}{S_{rss}(t)} m_{[i],rss}(t) \mathbb{I}(S_{rss}(t) > 0). \quad (3.4)$$

It is easy to show that $m_{rss}(t)$ can be rewritten as

$$m_{rss}(t) = \frac{1}{\sum_{i=1}^m \sum_{j=1}^r \mathbb{I}(X_{[i]j} > t)} \sum_{i=1}^m \sum_{j=1}^r (X_{[i]j} - t) \mathbb{I}(X_{[i]j} > t), \quad (3.5)$$

provided that $\sum_{i=1}^m \sum_{j=1}^r \mathbb{I}(X_{[i]j} > t) > 0$, otherwise the entire expression is defined as zero.

A common way to incorporate the tie information in the estimation process is to split each tied unit among the corresponding strata based on the ranks for which the unit is tied. If a measured unit is tied for k different ranks, then it is assigned to each of those k strata with a weight of $1/k$ instead of 1. This splitting strategy has been proposed by [2] for estimating the population mean. Based on this approach, the estimator of the survival function $S(t)$ can be expressed as

$$S_{sp}(t) = \frac{1}{m} \sum_{i=1}^m S_{[i],sp}(t), \quad (3.6)$$

where

$$S_{[i],sp}(t) = \frac{1}{n_i} \sum_{l=1}^m \sum_{j=1}^r \frac{I_{[l]ji}}{\sum_{k=1}^m I_{[l]jk}} \mathbb{I}(X_{[l]j} > t), \quad (3.7)$$

with $n_i = \sum_{l=1}^m \sum_{j=1}^r \frac{I_{[l]ji}}{\sum_{k=1}^m I_{[l]jk}}$, and $I_{[l]jk}$ is the indicator variable in the l th row and k th column of matrix $\mathbf{T}_j$. Thus, the first MRL estimator based on RSS-t is given by

$$m_{sp}(t) = \frac{1}{m} \sum_{i=1}^m \frac{S_{[i],sp}(t)}{S_{sp}(t)} m_{[i],sp}(t) \mathbb{I}(S_{sp}(t) > 0), \quad (3.8)$$

where $m_{[i],sp}(t) = \frac{\int_t^{+\infty} S_{[i],sp}(x) dx}{S_{[i],sp}(t)} \mathbb{I}(S_{[i],sp}(t) > 0)$.

The second estimator of the MRL function based on RSS-t is suggested based on the requirement that the survival function of sample units in the judgment strata are expected to adhere to the order restriction

$$S_{[1]}(t) \leq \cdots \leq S_{[m]}(t), \quad (3.9)$$

for all $t \in \mathbb{R}$. This constraint clearly holds under perfect ranking assumption where $S_{[i]}(t)$ follows the survival function of the i th order statistic of a sample of size m . However, sample estimates of these quantities often violate the restriction due to sampling noise,

especially for small sample sizes. To address this issue and improve estimation, we propose an isotonized version of $m_{sp}(t)$, which is given by

$$m_{iso}(t) = \frac{1}{m} \sum_{i=1}^m \frac{S_{[i],iso}(t)}{S_{iso}(t)} m_{[i],iso}(t) \mathbb{I}(S_{iso}(t) > 0), \quad (3.10)$$

where $S_{iso}(t) = \frac{1}{m} \sum_{i=1}^m S_{[i],iso}(t)$, $m_{[i],iso}(t) = \frac{\int_t^{+\infty} S_{[i],iso}(x) dx}{S_{[i],iso}(t)} \mathbb{I}(S_{[i],iso}(t) > 0)$ and

$$S_{[i],iso}(t) = \max_{1 \leq r \leq i} \min_{i \leq k \leq m} \frac{\sum_{l=r}^k n_l S_{[l],sp}(t)}{\sum_{l=r}^k n_l}, \quad i = 1, 2, \dots, m. \quad (3.11)$$

The estimates $S_{[1],iso}, \dots, S_{[m],iso}$, obtained from the equation (3.11), are isotonized estimates of $S_{[1],sp}, \dots, S_{[m],sp}$ with weighted sample sizes $n_1, \dots, n_m$, respectively. These estimates minimize the weighted least square $\sum_{i=1}^m n_i (S_{[i],sp}(t) - S_{[i],iso}(t))^2$ under the restricted space $\{\mathbf{S}(t) \in [0, 1]; S_{[1]}(t) \leq \dots \leq S_{[m]}(t)\}$, and therefore they follow the order restriction (3.9).

4 Monte Carlo simulation

In this section, we compare the performance of the RSS-t estimators of the MRL function to the RSS and SRS estimators, using Monte Carlo simulation. The simulations were run with 100,000 repetitions in R, considering different set sizes and ranking quality. To do this, we examine a gamma distribution with shape parameter 0.5 and scale parameter 1, denoted $G(0.5, 1)$, with an increasing MRL function. We used different set sizes $m \in \{3, 5\}$, and a sample size of $n = 15$. To generate the ranking process in RSS and RSS-t, we applied the linear ranking model proposed by [1]. This model assumes that for each set of size m , the ranking is determined by the perceived value of a variable Y , while the variable of interest is X . The model is described as follows

$$Y = \rho \left(\frac{X - \mu_x}{\sigma_x} \right) + \sqrt{1 - \rho^2} Z,$$

where μ_x and σ_x are the mean and the standard deviation of X , respectively, Z is a random variable follows a standard normal distribution, independent of X , and $\rho \in [0, 1]$ is the correlation coefficient between X and Y , which controls the quality of the ranking process. The parameter ρ takes its values from the set $\{1, 0.8, 0.5\}$. A perfect ranking occurs when $\rho = 1$. A fairly good ranking occurs when $\rho = 0.8$. A weak ranking occurs when $\rho = 0.5$. We used the mean squared error (MSE) to evaluate the accuracy of the proposed estimators and calculated the simulated relative efficiency (RE) with respect to $m_{srs}(t)$. For instance, the relative efficiency of $m_{rss}(t)$ compared to $m_{srs}(t)$ is given by

$$RE(t) = \frac{MSE(m_{srs}(t))}{MSE(m_{rss}(t))},$$

where $MSE(\hat{\theta}) = \frac{1}{100,000} \sum_{l=1}^{100,000} (\hat{\theta}_l - \theta)^2$ is the simulated MSE of an estimator $\hat{\theta}$ of a parameter θ based on 100,000 iterations. According to the definition of relative efficiency, when $RE(t)$ is greater than one, it indicates an advantage of the RSS estimator, while a value smaller than one indicates an advantage of the SRS estimator. We note that the $RE(t)$ is estimated for $t = Q(p)$, where $Q(p)$ represents the p th quantile of the chosen distribution, with p ranging from 0.01 to 0.99.

In the literature, two common methods for generating ties in rankings are Discrete Perceived Size (DPS) and Tied-if-close (TIC). In the DPS method, the perceived value Y is discretized using an operation such as $[Y/c]$, where $[y]$ is the greatest integer less than or equal to y , and $c > 0$ is a user-chosen model parameter. We declare ties between the i th and j th units in the set of size m , if their perceived values are close enough, such that $[Y_i/c] = [Y_j/c]$. In the TIC method, we declare ties between the i th and j th units in the set of size m , if $|Y_i - Y_j| < c$, where $c > 0$ is a user-chosen model parameter. Due to the transitive property of ties, the i th and j th units can be considered as tied, even when $|Y_i - Y_j| \geq c$, given the presence of other perceived values that bridge the gap. Both methods can approximate a wide range of real situations. In the DPS method, we set the c values from the set $\{0.5, 1, 2, 4\}$ and in the TIC method, c takes values from the set $\{0.25, 0.5, 1, 2\}$.

Figures 2 and 1 illustrate the simulation results for $m = 5$ and $m = 3$, respectively, under DPS method. Similarly, Figures 3 and 4 present the simulation results for $m = 5$ and $m = 3$, respectively, under TIC method. By comparing these figures, we observe that the $RE(p)$ generally decreases as set size m decreases. This observation is consistent with the findings reported in most RSS studies. All figures clearly demonstrate that the REs of the three estimators exhibit a decreasing trend as p increases. The top four panels in Figure 2 present the results under perfect ranking ($\rho = 1$). It is clear from these panels that the efficiency improvements gained from utilizing the RSS-t estimators in comparison to its SRS counterpart can be substantial, reaching up to 200% for small values of p . In the second and third rows of Figure 2, the simulation study results are displayed for $\rho = 0.8$ and $\rho = 0.5$, respectively. The trends observed in the REs of these three estimators within these rows closely resemble those of the top panels. The only distinction is that the REs in these panels are lower compared to their counterparts in the case of perfect ranking, as one would anticipate. Among these three estimators, it is evident that $m_{rss}(t)$, which ignores the tie information, generally exhibits weaker performance than the other two estimators which use the tie information, especially when $p < 0.6$. This finding supports the conclusion that incorporating tie information in the RSS technique improves the estimation of the MRL function. Additionally, in the comparison between $m_{sp}(t)$ and its isotonized version, $m_{iso}(t)$, we observe that $m_{iso}(t)$ demonstrates slightly higher efficiency in all considered cases. However, the REs are so close that the two curves overlap in all figures, making them indistinguishable from each other. Therefore, $m_{iso}(t)$ has better performance than $m_{srs}(t)$ and $m_{sp}(t)$ in nearly all the settings, which can be attributed to the fact that the order constraint (3.9) is satisfied in this simulation study. We observe from Figures 3 and 4 that the general pattern of the RE curves for the three estimators is consistent with that presented in Figures 2 and 1.

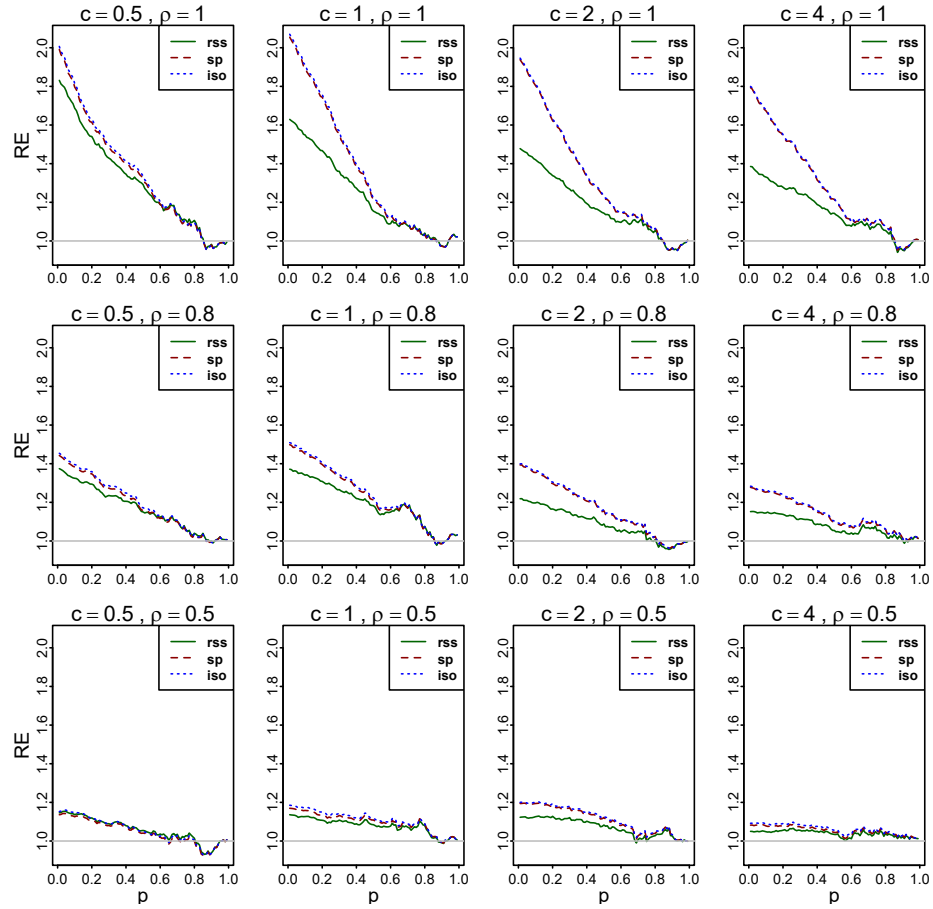


Figure 1: Simulated efficiencies of m_{rss} (green line), m_{sp} (red dash line) and m_{iso} (blue dotted line) with respect to m_{srs} as a function of p , for $G(0.5, 1)$ distribution, under DPS method, when $c \in \{0.5, 1, 2, 4\}$, $\rho \in \{1, 0.8, 0.5\}$, and $\{n, m\} \in \{15, 5\}$.

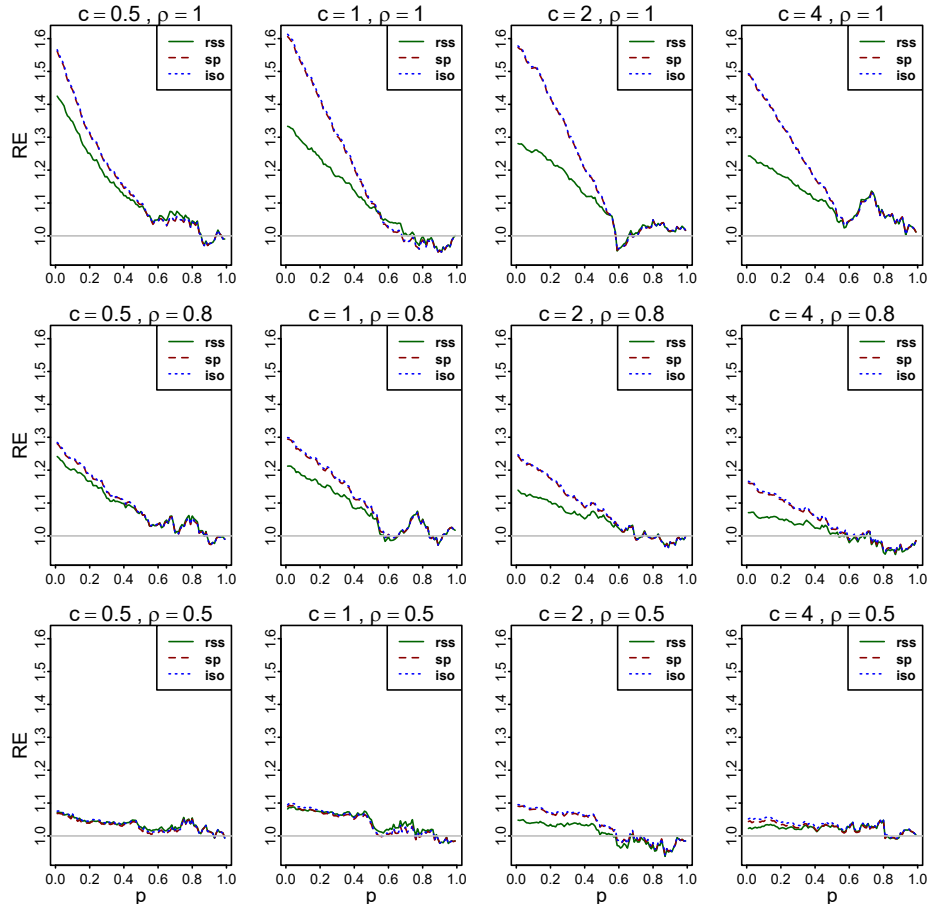


Figure 2: Simulated efficiencies of m_{rss} (green line), m_{sp} (red dash line) and m_{iso} (blue dotted line) with respect to m_{srs} as a function of p , for $G(0.5, 1)$ distribution, under DPS method, when $c \in \{0.5, 1, 2, 4\}$, $\rho \in \{1, 0.8, 0.5\}$, and $\{n, m\} \in \{15, 3\}$.

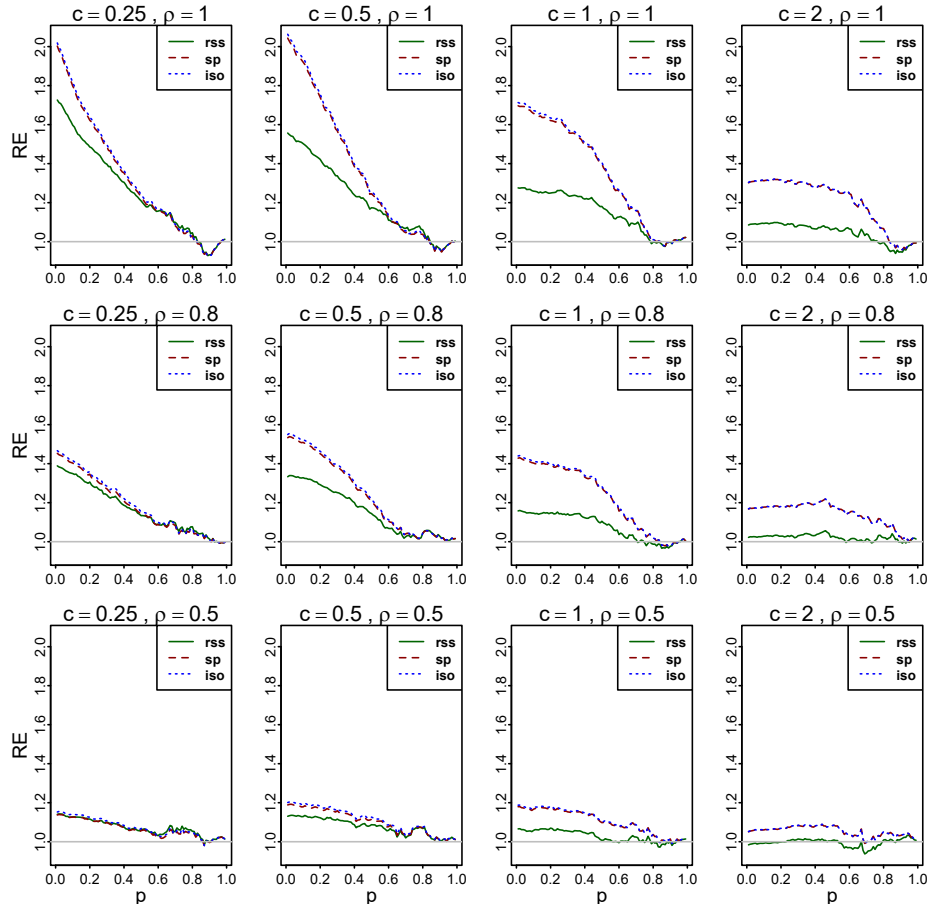


Figure 3: Simulated efficiencies of m_{rss} (green line), m_{sp} (red dash line) and m_{iso} (blue dotted line) with respect to m_{srs} as a function of p , for $G(0.5, 1)$ distribution, under TIC method, when $c \in \{0.25, 0.5, 1, 2\}$, $\rho \in \{1, 0.8, 0.5\}$, and $\{n, m\} \in \{15, 5\}$.

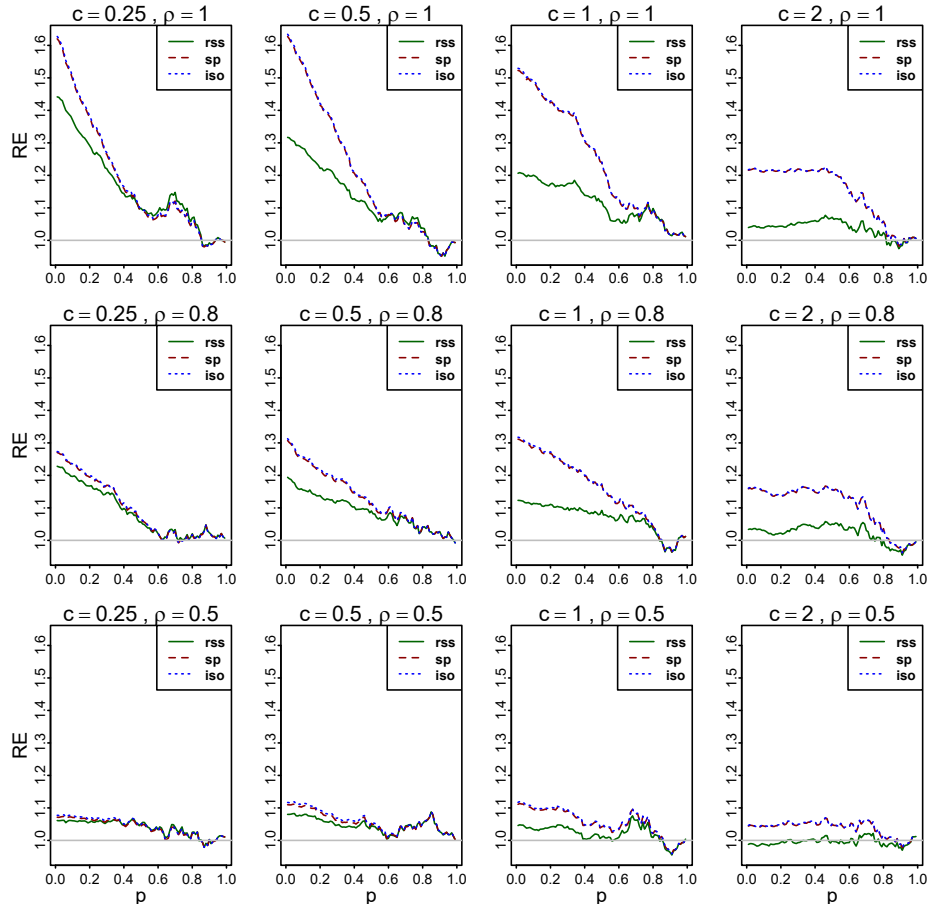


Figure 4: Simulated efficiencies of m_{rss} (green line), m_{sp} (red dash line) and m_{iso} (blue dotted line) with respect to $m_{sr,s}$ as a function of p , for $G(0.5, 1)$ distribution, under TIC method, when $c \in \{0.25, 0.5, 1, 2\}$, $\rho \in \{1, 0.8, 0.5\}$, and $\{n, m\} \in \{15, 3\}$.

5 Conclusion

The mean residual life (MRL) function plays an important role in survival analysis and reliability engineering. This field can be costly due to the need for expensive equipment or long follow-up periods. In such situations, where ranking the variable of interest using a cost-effective method is feasible, employing a cost-effective sampling method, such as ranked set sampling (RSS) or its modifications, instead of simple random sampling (SRS), is recommended. This approach reduces the required sample size to achieve the desired level of precision. In RSS, the researcher must assign a unique judgment rank to each set of size m . If there is uncertainty in ranking certain adjacent units, RSS can only proceed once the tied units are randomly broken. RSS-t, proposed by [2] as a novel variation of RSS, records and utilizes the structure of tied units for use in estimation.

In this paper, we proposed two MRL estimators using the RSS-t technique. Then we compared their performance with counterparts in RSS and SRS through a Monte Carlo simulation study. Our simulation results suggest that incorporating the tie information into RSS improves the efficiency of estimators.

References

- [1] Dell, T. R., and Clutter, J. L. (1972). Ranked set sampling theory with order statistics background. *Biometrics*, 545-555.
- [2] Frey, J. (2012). Nonparametric mean estimation using partially ordered sets. *Environmental and ecological statistics*, **19**, 309-326.
- [3] McIntyre, G. A. (1952). A method for unbiased selective sampling, using ranked sets. *Australian journal of agricultural research*, **3**(4), 385-390.
- [4] Yang, G. L. (1978). Estimation of a biometric function. *The Annals of Statistics*, 112-116.
- [5] Zamanzade, E., Parvardeh, A., and Asadi, M. (2019). Estimation of mean residual life based on ranked set sampling. *Computational Statistics and Data Analysis*, **135**, 35-55.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Mission Abort Strategy for Multi-Component Systems Using Survival Signature

Karimi, A.¹ Tavangar, M.²

Department of Statistics, Faculty of Mathematics and Statistics, University of Isfahan,
Isfahan 81744

Abstract

Survivability of important real-world systems such as satellites, aircraft, and submarines is vital as their failures can lead to significant economic losses and operational risks. An effective approach to enhance system survivability and reduce failure probability is to implement a mission abort policy. The mission is halted and recovery measures are initiated when the risk of system failure exceeds a certain level. This paper discusses a coherent system consisting of independent, heterogeneous components with a mission abort policy. Under survival signatures, we develop a probabilistic model to evaluate system survivability. With the help of elaborate illustrative examples, we compare different abort rules.

Keywords: Reliability, Survival signature, Mission abort policy, Mission success probability, System survival probability.

1 Introduction

Many real-world critical systems are designed to perform specific missions for a predetermined duration. The mission success probability (MSP) is commonly used as a key

¹m.amin.karimi@sci.ui.ac.ir

²m.tavangar@sci.ui.ac.ir

measure to evaluate the mission-wise dependability of such systems (see, e.g., [6, 27]). However, in cases where mission failure can lead to severe economic losses or operational risks, ensuring the survival of the system may be more critical than mission completion. Examples of such critical systems include Unmanned Aerial Vehicles (UAVs), airplanes, and complex data-processing computer systems. Previous studies (e.g., [23, 19]) have sought to improve both the probability of mission success and the systems survival probability.

A mission abort policy serves as an effective strategy for enhancing system survival and minimizing failure risks in high-stakes environments. In this approach, a mission may be terminated when a predefined failure condition is met, such as the failure of a certain number of components, allowing for immediate rescue procedures to mitigate financial and operational losses. Such policies are particularly useful in applications where system survivability is paramount. Earlier research has examined mission abort strategies for various system configurations, including redundant systems [7], k-out-of-n systems [20], and warm standby systems with dynamic task arrivals [34, 2]. More recent studies have explored mission abort policies in stochastic environments influenced by random shocks [13, 14] and scenarios involving multiple mission attempts and rescue operations [15, 17, 18].

In addition, several studies have investigated mission abort policies under different conditions, including systems with minor and major failures [3, 26], multi-state systems [16, 31], and UAVs with optimized route planning and termination strategies [22]. Recent research has also examined mission abort strategies based on degradation processes [25, 8, 32] and systems with incomplete state monitoring [5]. Furthermore, dynamic rescue policies and cost-minimizing mission abort strategies have been analyzed in various contexts [33, 4].

Most existing studies on mission abort policies do not explicitly consider system structure, with notable exceptions such as [20], who analyzed k-out-of-n systems with exponentially distributed component lifetimes. However, mission abort strategies for general coherent systems, especially those with heterogeneous components, remain largely unexplored. Addressing this gap is the primary objective of the present paper.

To achieve this, we introduce a novel signature-based approach that enables the formulation of new optimization problems that have not been previously considered. While our earlier work [10, 11] introduced a signature-based mission abort policy for systems with independent and identically distributed (i.i.d.) components, the assumption of identical component lifetimes is often unrealistic in practice. In this paper, we extend the methodology to accommodate systems with heterogeneous components by utilizing the concept of survival signatures, as introduced by [4] and further developed by [5].

A coherent system is defined as a system whose structure function is increasing in each components state, with all components contributing to system performance [1]. The signature vector $\mathbf{s} = (s_1, s_2, \dots, s_m)$ is a well-established tool for analyzing the reliability of such systems [12, 28]. When all components are i.i.d., the system lifetime distribution can be expressed in terms of its signature. Navarro et al. [21] extended this representation to systems with continuous and exchangeable component lifetimes. Let the coherent system consist of m components with independent and identically distributed (i.i.d.) lifetimes. Let also T denote the system lifetime. Using the concept of signature,

the reliability function $\bar{F}_T(t)$ of the system can be represented as

$$\bar{F}_T(t) = \mathbb{P}(T > t) = \sum_{i=1}^m s_i \mathbb{P}(X_{i:m} > t),$$

where $X_{i:m}$ denotes the i th ordered component lifetime and

$$s_i = \mathbb{P}(T = X_{i:m}), \quad i = 1, 2, \dots, m,$$

is the i th element of the signature vector. Navarro et al. [21] proved that such a representation also holds in the case where the component's lifetimes are continuous and exchangeable random variables.

Karimi et al. [11] employ the survival signature, a generalization of the traditional signature concept that accounts for multiple component types. Let a system consist of m components, categorized into H distinct types, where m_h represents the number of components of type $h \in \{1, 2, \dots, H\}$, such that $\sum_{h=1}^H m_h = m$. The survival signature, denoted as $\Phi(l_1, l_2, \dots, l_H)$, expresses the probability that the system remains functional given exactly l_h operational components of type $h \in \{1, 2, \dots, H\}$. Following Hashemi et al. [9], we use an equivalent dual form:

$$\Phi(l_1, l_2, \dots, l_H) = \Psi(m_1 - l_1, m_2 - l_2, \dots, m_H - l_H).$$

This representation is particularly useful for deriving key reliability metrics, including mission success probability, rescue effectiveness, and system survival probability.

The rest of the paper is structured as follows. Section 2 presents the main results and the proposed probabilistic model. Section 3 provides the relevant optimization procedures. Section 4 provides numerical illustrations and detailed discussions.

2 Mission abort strategy

Consider a coherent system consisting of $m = \sum_{h=1}^H m_h$ components, categorized into H distinct types, where m_h represents the number of components of type h . The system is required to operate within the fixed mission duration $[0, t_M]$. Ensuring the survival of critical systems is a top priority, as failures can lead to significant operational and economic consequences.

To address this challenge, we propose a mission abort policy aimed at increasing system survivability by reducing the probability of catastrophic failure. Under this policy, a mission is aborted if, at a predetermined critical time ξ ($\xi \leq t_M$), the number of failed components exceeds a specified threshold L . Here, L can take any value from the set $1, 2, \dots, m - 1$. In other words, if the total number of component failures at time ξ , denoted as $N(\xi)$, reaches or exceeds L (i.e., $N(\xi) \geq L$), the mission is immediately aborted, and a rescue procedure is initiated to improve system survivability. In this mission abort and rescue policy, both L and ξ serve as decision variables in the corresponding optimization problem. To facilitate the analysis, we employ the systems survival signature,

an effective probabilistic tool for assessing system reliability and deriving the relevant probabilities. The following sections introduce three key methods for identifying.

In the next sections, we will discuss three key methods that are useful for identifying and evaluating the optimal mission abort policy.

2.1 Mission success probability (MSP)

In the context of mission abort policies, the probability of mission success (MSP) is a fundamental measure of system dependability (see, e.g., [13, 14]). This metric plays a crucial role in evaluating the effectiveness of the mission abort policy, as it quantifies the probability of a system successfully completing its designated mission within the given time frame.

To derive the MSP within our proposed approach, two key conditions must be considered. First, the systems lifetime denoted as T , must exceed the mission duration t_M . Second, at the critical time ξ , the total number of failed components, represented as $N(\xi)$, must not exceed $(L - 1)$, where L is an integer parameter ranging from 1) to $m - 1$. In other words, the system can tolerate up to $(L - 1)$ component failures at time ξ without triggering mission termination.

This condition ensures that, among all instances where the system operates beyond t_M , only those meeting the failure threshold at ξ contribute to mission success. Conversely, if the number of failed components surpasses $(L - 1)$ at ξ , the mission is aborted, reducing the probability of complete system failure in the remaining mission duration and thereby enhancing system survivability.

To compute the MSP, we apply the law of total probability. Specifically, we categorize system realizations based on the number of component failures at ξ and evaluate the corresponding conditional probabilities. Thus, the MSP denoted as $M(L, \xi, t_M)$, can be formally expressed as follows.

$$\begin{aligned} M(L, \xi, t_M) &= \mathbb{P}(T > t_M, N(\xi) \leq L - 1) \\ &= \sum_{l=1}^{L-1} \sum_{\mathbf{l} \in B_l} \sum_{\mathbf{r} \in D_l} \mathbb{P}(T > t_M \mid N_i(\xi) = l_i, N_i(t_M) = r_i, i = 1, 2, \dots, H) \\ &\quad \times \mathbb{P}(N_i(t_M) = r_i, i = 1, 2, \dots, H \mid N_i(\xi) = l_i, i = 1, 2, \dots, H) \\ &\quad \times \mathbb{P}(N_i(\xi) = l_i, i = 1, 2, \dots, H) \end{aligned}$$

where $\mathbf{l} = (l_1, l_2, \dots, l_H)$, $\mathbf{r} = (r_1, r_2, \dots, r_H)$, $B_l = \{\mathbf{l} : \sum_{i=1}^H l_i = l\}$, $D_l = \{\mathbf{r} : l_i \leq r_i \leq m_i, i = 1, 2, \dots, H\}$. Hashemi et al. [9] have proven that the first conditional probability in the right-hand side of the last equation is the same as $\Psi(r_1, \dots, r_H)$. Therefore, we have

$$M(L, \xi, t_M) = \sum_{l=1}^{L-1} \sum_{\mathbf{l} \in B_l} \sum_{\mathbf{r} \in D_l} \Psi(r_1, \dots, r_H) \mathbb{P}(N_i(t_M) = r_i, N_i(\xi) = l_i, i = 1, \dots, H). \quad (2.1)$$

Therefore, finally

$$M(L, \xi, t_M) = \sum_{l=1}^{L-1} \sum_{\mathbf{l} \in B_l} \sum_{\mathbf{r} \in D_l} \Psi(r_1, \dots, r_H) \prod_{i=1}^H \binom{m_i}{r_i} \binom{r_i}{l_i} \times \{F_i(\xi)\}^{l_i} \{F_i(t_M) - F_i(\xi)\}^{r_i - l_i} \{\bar{F}_i(t_M)\}^{m_i - r_i}. \quad (2.2)$$

see, for example, [11].

In the special case where all components have a common lifetime distribution $F(t)$, the MSP can be simplified to

$$\begin{aligned} M(L, \xi, t_M) &= \mathbb{P}(T > t_M, N(\xi) \leq L-1) \\ &= \sum_{l=1}^{L-1} \mathbb{P}(T > t_M, N(\xi) = l) \\ &= \sum_{l=1}^{L-1} \sum_{\mathbf{r} \in D_l} \bar{S}_L \binom{m}{r} \binom{r}{l} \{F(\xi)\}^l \{F(t_M) - F(\xi)\}^{r-l} \{\bar{F}(t_M)\}^{m-r}. \end{aligned}$$

where $\bar{S}_L = \sum_{j=L+1}^m s_j$. Moreover, if $\xi = t_M$, then we have

$$\begin{aligned} \mathbb{P}(T > t_M, N(t_M) \leq L-1) &= \sum_{l=0}^{L-1} \mathbb{P}(T > t_M, N(t_M) = l) \\ &= \sum_{l=0}^{L-1} \bar{S}_L \binom{m}{l} \{F(t_M)\}^l \{\bar{F}(t_M)\}^{m-l}. \end{aligned} \quad (2.3)$$

see, for example, [10].

In the context of mission planning and decision-making, the connection between mission abort time, mission duration, and mission success probability is of great significance. By diligently examining and comprehending these interconnections, mission planners can acquire valuable insights into the determinants that impact the achievement of a mission. Such information empowers them to make well-informed decisions concerning the appropriate duration of the mission and, if necessary, to abort it.

Remark 2.1. From Equation (2.2), one can get a simple comparison among the MSPs of two m -component systems consisting of H groups (of independent components) where the corresponding groups in the systems have the same number of components. Let T_1 and T_2 be the system lifetimes with m_h components of H types for $h = 1, 2, \dots, H$, that have the cumulative distribution function (cdf) $F_h(t)$. Let Ψ_i be the survival signature of system j , $j = 1, 2$. If $\Psi_1(r_1, \dots, r_H) \leq \Psi_2(r_1, \dots, r_H)$ for all $r_1, r_2, \dots, r_H$, then $MSP^{(1)} \leq MSP^{(2)}$, where $MSP^{(j)}$ is the mission success probability corresponding to system T_j , $j = 1, 2$, for all cumulative distribution functions $F_1, F_2, \dots, F_H$.

2.2 Rescue procedure

As mentioned previously, in our mission abort model, the rescue procedure (RP) is activated when at least L components malfunction before the critical time ξ , resulting

in a mission abort at ξ (provided that the system has not failed until time ξ). For the survival of the system after abort, its remaining lifetime should be larger than the RP duration $t_{r(\xi)}$. However, in some instances, it is important that the system not just survive the RP, but accomplish it without experiencing many component failures. The latter, among other undesirable consequences, can result, for instance, in more costly maintenance after the completion of the RP. To deal with it, we suggest an innovative approach to the management of the RP by introducing the maximal number of failed components during the RP, $K < m - L$. Under this new condition, which generalizes the unconditional case of just survival, the probability of accomplishing the RP, which is denoted by $R(L, \xi, t_{r(\xi)}, K)$, can be defined as

$$R(L, \xi, t_{r(\xi)}, K) = \mathbb{P}(T > \xi + t_{r(\xi)}, N(\xi + t_{r(\xi)}) - N(\xi) \leq K \mid T > \xi, N(\xi) \geq L). \quad (2.4)$$

This conditional probability can be rephrased as

$$R(L, \xi, t_{r(\xi)}, K) = \frac{\mathbb{P}(T > \xi + t_{r(\xi)}, N(\xi + t_{r(\xi)}) - N(\xi) \leq K, N(\xi) \geq L)}{\mathbb{P}(T > \xi, N(\xi) \geq L)}.$$

First, note that the numerator of the last expression can be written as

$$\begin{aligned} & \mathbb{P}(T > \xi + t_{r(\xi)}, N(\xi + t_{r(\xi)}) - N(\xi) \leq K, N(\xi) \geq L) \\ &= \sum_{l=L}^m \sum_{k=0}^{\min(m-l, K)} \mathbb{P}(T > \xi + t_{r(\xi)}, N(\xi + t_{r(\xi)}) - N(\xi) = k, N(\xi) = l) \end{aligned} \quad (2.5)$$

On the other hand, the denominator of $R(L, \xi, t_{r(\xi)}, K)$ is given as

$$\begin{aligned} \mathbb{P}(T > \xi, N(\xi) \geq L) &= \sum_{l=L}^m \sum_{l \in B_l} \mathbb{P}(T > \xi, N_i(\xi) = l_i) \\ &= \sum_{l=L}^m \sum_{l \in B_l} \Psi(l_1, \dots, l_H) \prod_{i=1}^H \binom{m_i}{l_i} \{F_i(\xi)\}^{l_i} \{\bar{F}_i(\xi)\}^{m_i - l_i}. \end{aligned} \quad (2.6)$$

Ultimately, the probability $R(L, \xi, t_{r(\xi)}, K)$ is determined by dividing expression (2.5) by expression (2.6).

Remark 2.2. When planning to start a rescue procedure for a system, it is also crucial to consider the aborting time at the critical time ξ in connection with the mission's overall completion time, t_M . For instance, it may seem unreasonable to launch a rescue procedure if the critical time ξ approaches t_M .

2.3 System survival probability

System survival probability (SSP) is a crucial indicator used by decision-makers to evaluate the chances that a mission, including any potential rescue procedures, would be successful [13, 14]. A comprehensive mission abort policy that outlines precise protocols

for potential system failures proactively mitigates organizational risk. SSP is used by decision-makers to estimate the probability that a mission will be completed successfully or that a rescue procedure will be carried out.

Considering the results and notations introduced in the previous section, the SSP, defined by $S(L, \xi, t_M, t_{r(\xi)}, K)$, is formulated as follows

$$S(L, \xi, t_M, t_{r(\xi)}, K) = M(L, \xi, t_M) + \mathbb{P}[T > \xi, N(\xi) \geq L]R(L, \xi, t_{r(\xi)}, K). \quad (2.7)$$

A crucial part of the total SSP is represented by the second term on the right-hand side of the Equation (2.7). It represents the probability that the system can be successfully and effectively rescued after aborting a mission.

3 Optimization

In the policy of mission abort, critical parameters and timing must be determined to achieve maximum system survival probability at the lowest associated costs. In striving to achieve these ends, determining the best values for such parameters and their timings is vital. Numerous techniques may be employed to determine such optimum values in maximizing system survival probability as well as reducing costs.

The selection of the optimal abort parameter L and critical time ξ is crucial in the design of an effective mission abort policy. The parameter value L specifies the number of component failures at which a mission abort occurs at the critical time ξ . An optimization method can be employed to determine the most appropriate values for these parameters, as will be discussed below.

3.1 Trade-off between MSP and SSP

A method for obtaining the optimal values of L and ξ is to establish a balance between $M(L, \xi, t_M)$ and $S(L, \xi, t_M, t_{r(\xi)}, K)$, in order to maximize the MSP while providing a desired level of system survivability S^* (Levitin and Finkelstein, 2018a):

$$\max_{L, \xi} M(L, \xi, t_M) \quad s.t. \quad S(L, \xi, t_M, t_{r(\xi)}, K) \geq S^*. \quad (3.1)$$

The optimization problem outlined in the aforementioned equations necessitates determining the optimal values for the discrete parameter L and the continuous parameter ξ using the methodology proposed by [24]. For each discrete value of L , we employ a one-dimensional optimization technique to ascertain the optimal value of ξ . We begin by optimizing ξ for each fixed value of L , as L can only assume a finite set of discrete values. By maintaining L constant and systematically applying the optimization method across the range of ξ , we can identify the optimal ξ corresponding to each specific L . Subsequently, we compare these results across all possible values of L to derive the overall optimal combination of (ξ, L) , having first identified the optimal ξ for each discrete L .

3.2 Cost-Based Optimization

Selecting optimal mission abort parameters L (failure threshold) and ξ (critical abort time) requires considering the costs associated with mission failure (C_F) and system loss (C_L). The system loss probability is $1 - S(L, \xi, t_M, t_r(\xi), K)$, where t_M is the mission time and $t_r(\xi)$ the time to rescue. If the system is lost, both the mission and system fail, incurring a total cost of $C_F + C_L$. If the mission fails but the system is successfully rescued, the incurred cost is C_F . This leads to an expected cost function:

$$C(\xi, L) = (1 - S(L, \xi, t_M, t_r(\xi), K))(C_F + C_L) + (S(L, \xi, t_M, t_r(\xi), K) - M(L, \xi, t_M))C_F. \quad (3.2)$$

where $S(L, \xi, t_M, t_r(\xi), K)$ and $M(L, \xi, t_M)$ represent the system survivability and mission success probabilities, respectively. The optimal (L, ξ) pair minimizes this expected cost and can be found using optimization techniques such as linear or dynamic programming.

4 Numerical illustration, discussion and sensitivity analysis

In this section, we provide numerical illustrations with detailed discussions of our theoretical findings in Sections 2 and 3. All results, including figures and numerical computations, are obtained using Mathematica Program System version 14.

Consider a system with 12 components of three types (as pictured in Figure 1) consisting of 6 components of type 1; 2 components of type 2, and 4 components of type 3. We assume that the lifetimes of components of type 1 are i.i.d. with the cdf F_1 , the lifetimes of components of type 2 are i.i.d. with the cdf F_2 and the lifetimes of components of type 3 are i.i.d. with the cdf F_3 , and components of different types are also independent. Let

$$\begin{aligned} F_1(t) &= 1 - \exp\{-0.5t\}, & t \geq 0, \\ F_2(t) &= 1 - \exp\{-0.5\sqrt{t}\}, & t \geq 0, \\ F_3(t) &= 1 - \exp\{-0.5t^2\}, & t \geq 0. \end{aligned}$$

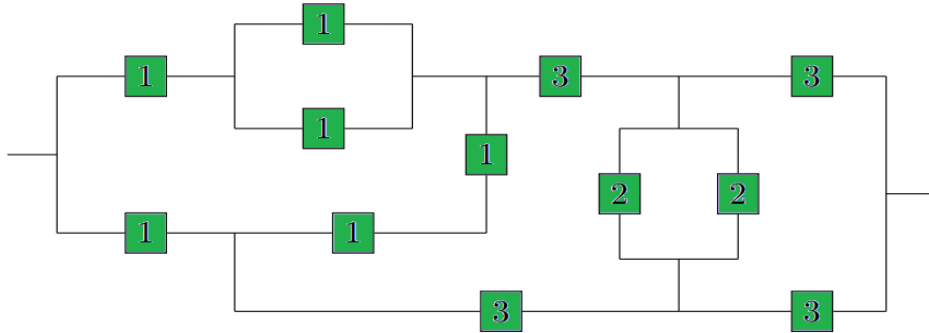


Figure 1: Reliability block diagram of a 12-component coherent system of three types.

Figure 2 provides descriptions of the mission success probability, which has been calculated using Equation (2.2). The complex relationship between two crucial factors,

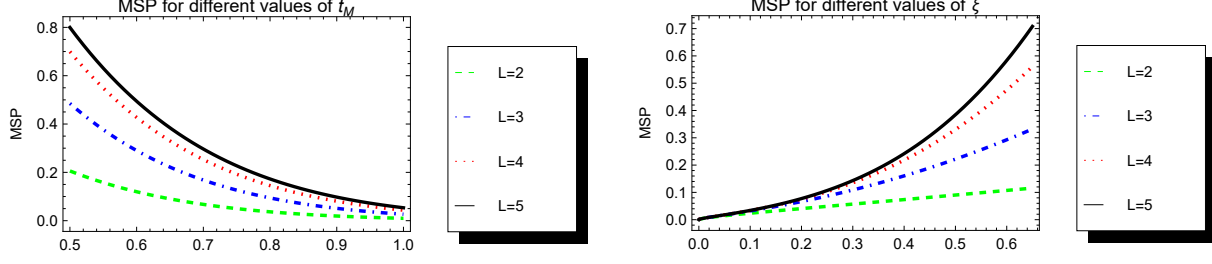


Figure 2: (left) Mission success probability for $L = 2, 3, 4, 5$ and $\xi = 0.5$. (right) Mission success probability for $L = 2, 3, 4, 5$ and $t_M = 0.65$.

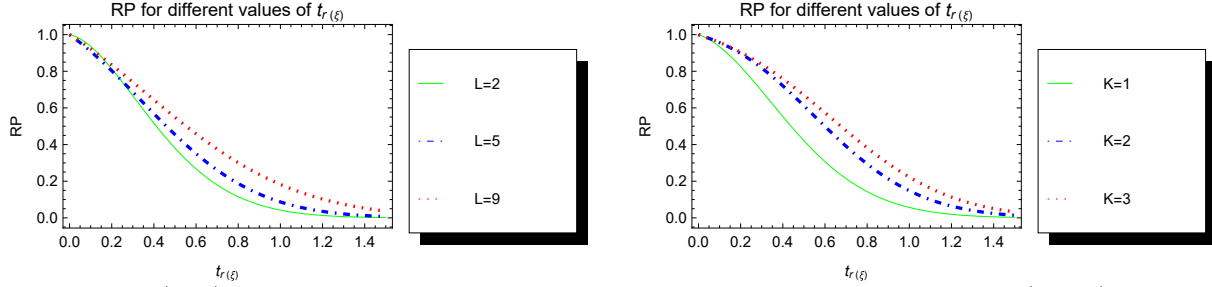


Figure 3: (left) Rescue Procedure for $L = 2, 7, 9$, $K = 3$ and $\xi = 1$. (right) Rescue Procedure for $K = 1, 2, 3$, $L = 3$ and $\xi = 0.1$.

i.e., the mission time t_M and the mission success probability, is explained in Figure 2 (left). It can be seen that the probability of mission success can be enhanced by increasing the value of L (the constraint on the number of failures in $[0, \xi]$). On the other hand, Figure 2 (right) illustrates the relationship between the critical time ξ and the probability of mission success.

Figure 3 provides the probability of completing the rescue procedure with decision parameters L and K , which have been calculated using Equation (2.4). Figure 3 illustrates the impact of extending the rescue period on the probability of completing the rescue procedure. Moreover, the figure highlights the pivotal role of the decision parameter L (the number of failed components at the critical time ξ), leading to a mission abort. It can also be observed (Figure 3 (right)) that the probability of completing the rescue procedure decreases significantly as the rescue time $t_r(\xi)$ increases. Furthermore, the effect of the threshold parameter K , which denotes the largest number of failed components the system can tolerate during the RP, is shown.

We use Equation (2.7) to calculate the system survival probability. We will specifically examine how several important parameters, such as the decision parameter L , threshold parameter K , mission period length t_M , mission abort time ξ , and the rescue procedure length $t_r(\xi)$, affect this probability. By adjusting these parameters, we can see how the probability of system survival changes (increasing or decreasing).

Figure 4 (left) shows the probability of the system survival plotted against the critical time ξ , assuming a constant length of the rescue period (0.1 in this case), where the mission period length is set at 1.5. This figure shows that the probability of system survival initially increases but eventually starts to decline. Furthermore, increasing L decreases the system survival probability. On the other hand, Figure 4 (right) shows

that an increase in the threshold value K results in a greater system survival probability. This implies that for middle values of ξ , higher values for K increase the system survival probability and strengthen the system's robustness. Furthermore, the system survival probability is decreasing with the increase of the rescue time $t_{r(\xi)}$. Nonetheless, the system survival probability behaves differently throughout a range of mission lengths. The probability of system survival tends to increase, for instance, if the mission length is less than 1 and near the critical period ($t_M < 1$). Conversely, a longer mission length causes the system survival probability to decrease.

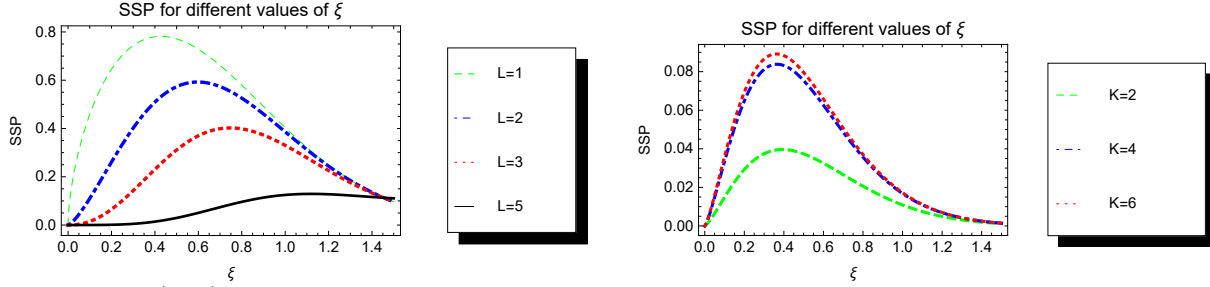


Figure 4: (left) System survival probability for $K = 2$, $L = 1, 2, 3, 5$, $t_{r(\xi)} = 0.1$ and $t_M = 1.5$. (right) System survival probability for $K = 2, 4, 6$, $L = 2$, $t_{r(\xi)} = 1$ and $t_M = 5$.

To minimize the system's overall cost, the cost function in Equation (??) was evaluated using mission and system failure costs of 10 and 120 units, respectively. By analyzing different values of ξ and L , the optimal parameters were found to be $\xi^* = 0.4265$ and $L^* = 1$, resulting in the minimum cost of 35.74 units, as shown in Figure 5.

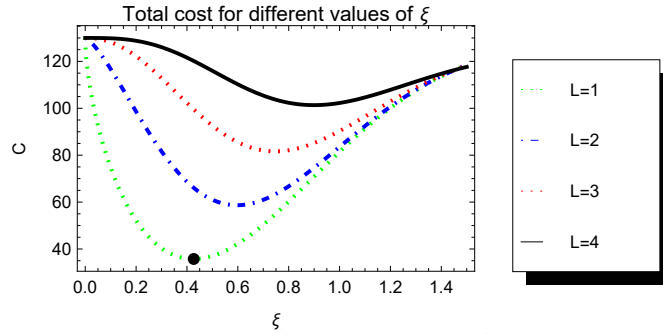


Figure 5: The system total cost.

References

- [1] Barlow, R. and Proschan, F. (1981). *Statistical Theory of Reliability and Life Testing: Probability models*. Hold, Reinhartand Wiston, Inc.
- [2] Cao, S. and Wang, X. (2023). Mission abort strategy for generalized k -out-of- n : F systems considering competing failure criteria. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, 1748006X231170909.

- [3] Cha, J.H., Finkelstein, M. and Levitin, G. (2018). Optimal mission abort policy for partially repairable heterogeneous systems. *European Journal of Operational Research*, **271** (3), 818–825.
- [4] Coolen, F. P. and Coolen-Maturi, T. (2012). Generalizing the signature to systems with multiple types of components. In *Complex systems and dependability* (pp. 115–130). Springer, Berlin, Heidelberg.
- [5] Coolen-Maturi, T., Coolen, F. P. and Balakrishnan, N. (2021). The joint survival signature of coherent systems with shared components. *Reliability Engineering and System Safety*, **207**, 107350.
- [6] Finkelstein, M. (2008). *Failure rate modeling for reliability and risk*. Springer.
- [7] Finkelstein, M (2001). How does our n-component system perform? *IEEE Transactions on Reliability*, **50**, 414–418.
- [8] Gao, K., Xiao, H., Mi, J., Peng, R. and Zhai, Q. (2020, October). Optimal abort policy of a distributed system of computers with Weibull failure time. In *2020 Global Reliability and Prognostics and Health Management (PHM-Shanghai)* (pp. 1–5). IEEE.
- [9] Hashemi, M., Asadi, M. and Zarezadeh, S. (2020). Optimal maintenance policies for coherent systems with multi-type components. *Reliability Engineering and System Safety*, **195**, 106674.
- [10] Karimi, A., Tavangar, M., and Finkelstein, M. (2025a). New optimal mission abort policies for coherent systems using signature. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, 1748006X251324958.
- [11] Karimi, A., Tavangar, M. and Finkelstein and M. (2025b). Mission Abort Policy for Coherent Systems with Heterogeneous Components. *Submitted*.
- [12] Kochar, S.C., Mukerjee, H. and Samaniego, F.J. (1999). The signature of a coherent system and its application to comparisons among systems. *Naval Research Logistics*, **46**, 507–523.
- [13] Levitin, G. and Finkelstein, M. (2018). Optimal mission abort policy for systems in a random environment with variable shock rate. *Reliability Engineering & System Safety*, **169**, 11–17.
- [14] Levitin, G. and Finkelstein, M. (2022). Making mission abort decisions for systems operating in random environment. In *Multicriteria and Optimization Models for Risk, Reliability, and Maintenance Decision Analysis: Recent Advances* (pp. 283–304). Cham: Springer International.
- [15] Levitin, G., Finkelstein, M. and Huang, H. Z. (2019). Optimal abort rules for multi-attempt missions. *Risk Analysis*, **39**(12), 2732–2743.
- [16] Levitin, G., Xing, L. and Dai, Y. (2022). Optimal mission aborting in multistate systems with storage. *Reliability Engineering & System Safety*, **218**, 108086.

- [17] Levitin, G., Xing, L. and Dai, Y. (2023a). Optimal task aborting policy and component activation delay in consecutive multi-attempt missions. *Reliability Engineering & System Safety*, **238**, 109482.
- [18] Levitin, G., Xing, L. and Dai, Y. (2024). Multi-attempt missions with multiple rescue options. *Reliability Engineering & System Safety*, **248**, 110168.
- [19] Liu, Y., Chen, Y. and Jiang, T. (2018). On sequence planning for selective maintenance of multi-state systems under stochastic maintenance durations. *European Journal of Operational Research*, **268** (1), 113–127.
- [20] Myers, A. (2009). Probability of loss assessment of critical k -out-of- n : G systems having a mission abort policy. *IEEE Transactions on Reliability*, **58** (4), 694701.
- [21] Navarro, J., Ruiz, J.M. and Sandoval, C.J. (2005). A note on comparisons among coherent systems with dependent components using signature. *Statistics & Probability Letters*, **72**, 179–185.
- [22] Peng, R. (2018). Joint routing and aborting optimization of cooperative unmanned aerial vehicles. *Reliability Engineering & System Safety*, **177**, 131–137.
- [23] Peng, R., Zhai, Q., Xing, L., Yang, J. (2016). Reliability analysis and optimal structure of series-parallel phased-mission systems subject to fault-level coverage. *IIE Transactions*, **48** (8), 736–746.
- [24] Press, W., Teukolsky S., Vetterling, W. and Flannery, B. (1992). *Numerical recipes in C*. Cambridge University Press, New York.
- [25] Qiu, Q., Cui, L. (2019). Gamma process based optimal mission abort policy. *Reliability Engineering & System Safety*, **190**, 106496.
- [26] Qiu, Q., Cui, L. and Wu, B. (2020). Dynamic mission abort policy for systems operating in a controllable environment with self-healing mechanism. *Reliability Engineering & System Safety*, **203**, 107069.
- [27] Rausand, M., Barros, A. and Hoyland, A. (2021). *System reliability theory: models, statistical methods, and applications*. John Wiley & Sons.
- [28] Samaniego, F.J. and Shaked, M. (2008). Systems with weighted components. *Statistics and Probability Letters*, **78**, 815–823.
- [29] Wei, F., Zhou, S., Ma, X. and Yang, L. (2023). Optimal mission abort planning for partially observable system. In *9th International Symposium on System Security, Safety, and Reliability (ISSSR)* (pp. 71–76).
- [30] Yang, L., Wei, F. and Qiu, Q. (2024). Mission risk control via joint optimization of sampling and abort decisions. *Risk Analysis*. **44** (3), 666–685.

- [31] Zhao, X., Chai, X., Sun, J. and Qiu, Q. (2022). Joint optimization of mission abort and protective device selection policies for multistate systems. *Risk Analysis*, **42** (12), 2823-2834.
- [32] Zhao, X., Sun, J., Qiu, Q. and Chen, K. (2021). Optimal inspection and mission abort policies for systems subject to degradation. *European Journal of Operational Research*, **292** (2), 610–621.
- [33] Zhao, X., Lv, Z., Qiu, Q. and Wu, Y. (2023a). Designing two-level rescue depot location and dynamic rescue policies for unmanned vehicles. *Reliability Engineering & System Safety*, **233**, 109119.
- [34] Zhao, X., Liu, H., Wu, Y. and Qiu, Q. (2023b). Joint optimization of mission abort and system structure considering dynamic tasks. *Reliability Engineering & System Safety*, **234**, 109128.
- [35] Zhao, X., Wang, X., Dai, Y., and Qiu, Q. (2024). Joint optimization of loading, mission abort and rescue site selection policies for UAV. *Reliability Engineering & System Safety*, **244**, 109955.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15. 2025



Burg Entropy: Extension and an Application to Binary Classification of Malaria Cell Images

Kharazmi, O.¹ Asadi, A.²

¹ Department of Statistics, Faculty of Mathematical Sciences, Vali-e-Asr University of Rafsanjan, Rafsanjan, Iran

² Department of Medical Entomology and Vector Control, School of Medical Sciences, Tarbiat Modares University, Tehran, Iran

Abstract

In this paper, we first establish some results for Burg entropy based on escort and generalized escort distributions. Then, we propose relative Burg divergence measure, cumulative residual Burg entropy (CRB), and cumulative residual relative Burg (CRRB) entropy measures. We apply the proposed CRRB measure to malaria cell image analysis, demonstrating its effectiveness in quantifying image similarity and segmentation performance. Comparative results with Otsus thresholding and K -means clustering highlight the efficacy of our approach in classification tasks.

Keywords: Burg Entropy, Shannon entropy, Binary Classification, Malaria cell disease

1 Introduction

Information measures play a fundamental role in quantifying uncertainty and diversity in various domains, providing essential tools for analyzing randomness, complexity, and structural properties of data.

¹omidkharazmi14@gmail.com

²asadiazadeh@modares.ac.ir

Regarding discrete information measures, there are several well-known measures, such as the Shannon entropy [5], Rényi entropy [10], Burg entropy [2], Tsallis entropy [13], Shafee entropy [11], extropy [7], Gini-Simpson index [9], and discrete Fisher information [8]. These measures find extensive application across diverse fields such as information theory, reliability, statistics, economics, and engineering.

Let X be a discrete random variable with probability mass function (PMF) $\mathbf{P} = (p_1, \dots, p_n)$, The Burg entropy associated with X is defined as

$$B(\mathbf{P}) = \sum_{i=1}^n \log p_i, \quad (1.1)$$

where $\log$ stands for the natural logarithm.

In this paper, by considering the discrete Burg entropy see [2], we introduce a relative divergence measure based on Burg entropy. Then, these concepts are extended by incorporating the survival function instead of the probability mass function. Within this framework, we provide some theoretical results, followed by an application to binary image classification of malaria cells. The classification results are compared using two algorithms: K -means and the Otsu method.

The rest of this paper is organized as follows. In Section 2, we first study the discrete Burg entropy for the escort and generalized escort probability mass functions (PMFs). We then propose a relative Burg entropy divergence and establish some of its properties. In Section 3, we introduce Burg entropy and its relative version in terms of the discrete survival function instead of PMFs, referred to as the cumulative residual Burg (CRB) entropy and the cumulative residual relative Burg (CRRB) entropy measures. In Section 4, we apply the proposed CRRB measure to a real-world example in image processing related to malaria cell images, presenting numerical results that demonstrate the CRRB measure as a valuable tool for quantifying similarity between two images. Performance benchmarks against established methods, such as Otsu's thresholding and K -means clustering, highlight the efficacy of the proposed approach in classification tasks. Finally, we conclude with some remarks in Section 5, summarizing our findings and suggesting potential directions for future research.

2 Burg entropy of escort and generalized escort PMFs

In this section, we examine the Burg entropy measure for escort and generalized escort probability mass functions. The corresponding results reveal interesting connections with the Rényi entropy and relative Rényi entropy.

Let $\mathbf{P} = (p_1, \dots, p_n)$ and $\mathbf{Q} = (q_1, \dots, q_n)$ be two probability mass functions. Then, the escort PMF associated with $\mathbf{P}$ and the generalized escort PMF associated with $\mathbf{P}$ and $\mathbf{Q}$ are defined, respectively, as

$$\mathbf{P}_\alpha = \left(\frac{p_1^\alpha}{\sum_{i=1}^n p_i^\alpha}, \dots, \frac{p_n^\alpha}{\sum_{i=1}^n p_i^\alpha} \right); \quad \mathbf{S}_\alpha = \left(\frac{p_1^\alpha q_1^{1-\alpha}}{\sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}}, \dots, \frac{p_n^\alpha q_n^{1-\alpha}}{\sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}} \right).$$

where $\alpha > 0$. The escort distribution holds significant importance in both non-extensive statistical mechanics and coding theory, being intricately tied to entropy measurements introduced by [10] and [13]. [1] examined the connections between coding theory and the evaluation of complexity within non-extensive statistical mechanics, with a specific focus on escort distributions.

Theorem 2.1. *Let $\mathbf{P}$ be a probability mass function (PMF), and let $\mathbf{P}_\alpha$ be its corresponding escort PMF with order α . Then, for $\alpha > 0$, we have*

$$B(\mathbf{P}_\alpha) = \alpha B(\mathbf{P}) + n(\alpha - 1)R_\alpha(\mathbf{P}), \quad (2.1)$$

where $B(\mathbf{P}_\alpha)$ is the Burg entropy defined in (1.1), and $R_\alpha(\mathbf{P})$ is Rnyi entropy is defined as

$$R_\alpha(\mathbf{P}) = \frac{1}{1 - \alpha} \log \sum_{i=1}^n p_i^\alpha, \quad \alpha > 0, \alpha \neq 1. \quad (2.2)$$

Theorem 2.2. *Let $\mathbf{P}$ and $\mathbf{Q}$ be two probability mass functions, and let $\mathbf{S}_\alpha$ be their corresponding generalized escort PMF with order α . Then, for $\alpha > 0$, we have*

$$B(\mathbf{S}_\alpha) = \alpha B(\mathbf{P}) + (1 - \alpha)B(\mathbf{Q}) + n(1 - \alpha)D_\alpha(\mathbf{P}, \mathbf{Q}), \quad (2.3)$$

where $B(\mathbf{S}_\alpha)$ is the Burg entropy defined in (1.1), and $D_\alpha(\mathbf{P}, \mathbf{Q})$ is the Rnyi divergence, defined as

$$D_\alpha(\mathbf{P}, \mathbf{Q}) = \frac{1}{\alpha - 1} \log \sum_{i=1}^n p_i^\alpha q_i^{1-\alpha}, \quad \alpha > 0, \alpha \neq 1. \quad (2.4)$$

Corollary 2.3. *Let $\mathbf{P}$ and $\mathbf{Q}$ be two probability mass functions, and let $\mathbf{S}_\alpha$ be their corresponding generalized escort PMF with order α . Then, for $\alpha > 1$ ($0 < \alpha < 1$), an upper bound (lower bound) for $B(\mathbf{S}_\alpha)$ is given by*

$$B(\mathbf{S}_\alpha) \leq (\geq) \alpha B(\mathbf{P}) + (1 - \alpha)B(\mathbf{Q}).$$

3 Relative Burg entropy divergence

Definition 3.1. Let X and Y be two discrete random variables with probability mass functions $\mathbf{P} = (p_1, \dots, p_n)$ and $\mathbf{Q} = (q_1, \dots, q_n)$, respectively. The relative Burg entropy divergence measure is defined as

$$D_B(\mathbf{P}, \mathbf{Q}) = \sum_{i=1}^n \left[\log \left(\frac{p_i}{q_i} \right) - \left(1 - \frac{q_i}{p_i} \right) \right]. \quad (3.1)$$

Theorem 3.2. *The relative Burg entropy divergence measure defined in (3.1) is non-negative.*

Example 3.3. Let $X \sim \text{Bernoulli}(p)$ and $Y \sim \text{Bernoulli}(q)$, where X takes the value 1 with probability p and 0 with probability $1 - p$, and similarly, Y takes the value 1 with probability q and 0 with probability $1 - q$. The corresponding probability vectors of the

variables X and Y are given by $\mathbf{P} = (p, 1 - p)$ and $\mathbf{Q} = (q, 1 - q)$, respectively. Then, from the definition of the relative Burg entropy divergence in (3.1), we obtain

$$D_B(\mathbf{P}, \mathbf{Q}) = \log\left(\frac{p}{q}\right) - \left(1 - \frac{q}{p}\right) + \log\left(\frac{1-p}{1-q}\right) - \left(1 - \frac{1-q}{1-p}\right). \quad (3.2)$$

Figure 3 presents a 3D surface plot of the relative Burg divergence $D_B(\mathbf{P}, \mathbf{Q})$ for $0 < p, q < 1$, illustrating how the divergence quantifies the difference between the probability distributions characterized by p and q . The plot visually demonstrates the variation of $D_B(\mathbf{P}, \mathbf{Q})$ with respect to changes in p and q , highlighting regions of higher and lower divergence.

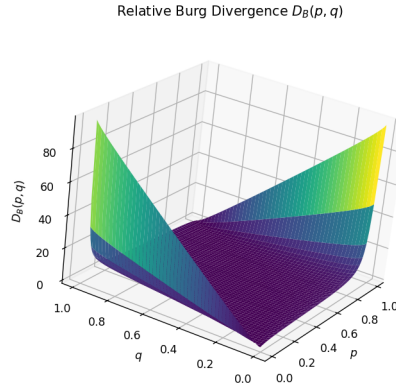


Figure 1: 3D surface plot of the relative Burg divergence $D_B(\mathbf{P}, \mathbf{Q})$ for $0 < p, q < 1$. The divergence measures the difference between the distributions characterized by p and q .

4 Cumulative residual Burg entropy and cumulative residual relative Burg entropy measures

Let X be a discrete random variable with probability mass function $\mathbf{P} = (p_1, \dots, p_n)$ and corresponding survival function

$$\bar{P}(i) = 1 - \sum_{j=0}^{i-1} p_j, \quad p_0 = 0.$$

Then, the cumulative residual Burg (CRB) entropy associated with X is defined as

$$CRB(\bar{P}) = \sum_{i=1}^n \log \bar{P}(i),$$

where $\log$ denotes the natural logarithm.

Example 4.1. Let $X \sim \text{Bernoulli}(p)$. Then, the probability mass function (PMF) is $\mathbf{P} = (p, 1 - p)$, and the survival function is

$$\bar{P}(1) = 1, \quad \bar{P}(2) = 1 - p.$$

Thus, the cumulative residual Burg (CRB) entropy is

$$CRB(\mathbf{P}) = \log \bar{P}(1) + \log \bar{P}(2) = \log(1) + \log(1 - p) = \log(1 - p).$$

Consider two discrete random variables X and Y with PMFs $\mathbf{P} = (p_1, p_2, \dots, p_n)$ and $\mathbf{Q} = (q_1, q_2, \dots, q_n)$, respectively. The survival functions are:

$$\bar{P}(i) = 1 - \sum_{j=0}^{i-1} p_j \quad \text{and} \quad \bar{Q}(i) = 1 - \sum_{j=0}^{i-1} q_j.$$

If the survival functions satisfy the relation $\bar{Q}(i) = (\bar{P}(i))^\beta$, then the hazard rates are related as:

$$h_X(i) = \frac{p_i}{\bar{P}(i)} \quad \text{and} \quad h_Y(i) = \frac{q_i}{\bar{Q}(i)} = \frac{q_i}{(\bar{P}(i))^\beta}.$$

Thus, the hazard rates of X and Y follow a proportional hazards model (PHM), with $\bar{Q}(i)$ being a power transformation of $\bar{P}(i)$.

Theorem 4.2. *Let X and Y be two discrete random variables with probability mass functions $\mathbf{P} = (p_1, p_2, \dots, p_n)$ and $\mathbf{Q} = (q_1, q_2, \dots, q_n)$, respectively. If X and Y follow a proportional hazards model (PHM), then their cumulative residual Burg entropy measures are proportional, i.e.,*

$$CRB(\bar{\mathbf{Q}}) = \beta \cdot CRB(\bar{\mathbf{P}}).$$

Definition 4.3. Let X and Y be two discrete random variables with probability mass functions

$$\mathbf{P} = (p_1, \dots, p_n) \quad \text{and} \quad \mathbf{Q} = (q_1, \dots, q_n),$$

respectively, defined on a common support $S = \{1, 2, \dots, n\}$. The cumulative residual relative Burg (CRRB) divergence is defined as

$$CD_B(\bar{\mathbf{P}}, \bar{\mathbf{Q}}) = \sum_{i=1}^n \left[\log \left(\frac{\bar{P}(i)}{\bar{Q}(i)} \right) - \left(1 - \frac{\bar{Q}(i)}{\bar{P}(i)} \right) \right]. \quad (4.1)$$

where

$$\bar{P}(i) = 1 - \sum_{j=0}^{i-1} p_j \quad \text{and} \quad \bar{Q}(i) = 1 - \sum_{j=0}^{i-1} q_j$$

are the survival functions associated with the PMFs $\mathbf{P}$ and $\mathbf{Q}$, respectively.

Example 4.4. Let $X \sim \text{Bernoulli}(p)$ and $Y \sim \text{Bernoulli}(q)$ with probability vectors $\mathbf{P} = (p, 1 - p)$ and $\mathbf{Q} = (q, 1 - q)$. Their survival functions are

$$\bar{P}(1) = 1, \quad \bar{P}(2) = 1 - p, \quad \bar{Q}(1) = 1, \quad \bar{Q}(2) = 1 - q.$$

Thus, the cumulative residual relative Burg (CRRB) divergence is

$$CD_B(\bar{\mathbf{P}}, \bar{\mathbf{Q}}) = \log \left(\frac{1 - p}{1 - q} \right) - \left(1 - \frac{1 - q}{1 - p} \right).$$

Figure 2 illustrates the 3D surface plot of the CRRB divergence $CD_B(\bar{\mathbf{P}}, \bar{\mathbf{Q}})$ as a function of the parameters p and q for two Bernoulli distributions. This plot visually demonstrates how the divergence changes with varying probability values p and q within the range $0 < p, q < 1$, providing insight into the sensitivity of the CRRB measure to differences between the distributions.

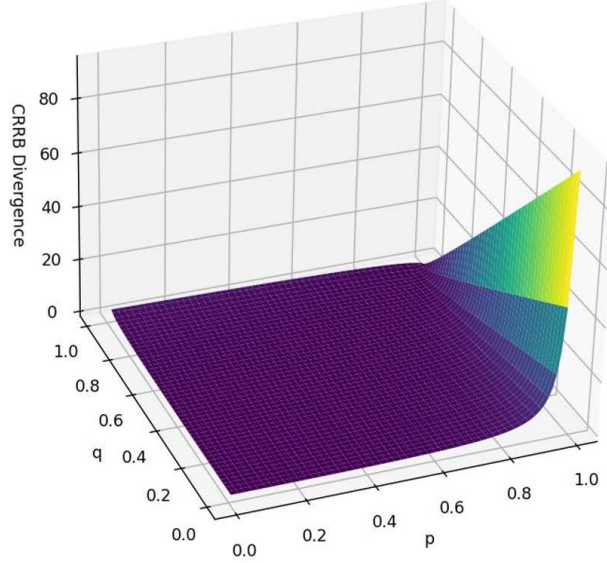


Figure 2: 3D surface plot of the CRRB Divergence $CD_B(\bar{\mathbf{P}}, \bar{\mathbf{Q}})$ as a function of the parameters p and q for two Bernoulli distributions. The plot shows how the divergence varies with changes in the probabilities p and q , where $0 < p, q < 1$.

In the following theorem, we establish some results associated with the CRRB divergence of a mixture hazards model.

Theorem 4.5. *Let X and Y be two discrete random variables with probability mass functions*

$$\mathbf{P} = (p_1, \dots, p_n) \quad \text{and} \quad \mathbf{Q} = (q_1, \dots, q_n),$$

respectively, defined on a common support $\{1, 2, \dots, n\}$. Then, we have

- (i) *The cumulative residual Burg divergence between $\bar{\mathbf{S}}$ and $\bar{\mathbf{P}}$ is given by*

$$CD_B(\bar{\mathbf{S}}, \bar{\mathbf{P}}) = (1 - \alpha) \left[CRB(\bar{\mathbf{Q}}) - CRB(\bar{\mathbf{P}}) \right] - n + \sum_{i=1}^n \left(\frac{\bar{P}(i)}{\bar{Q}(i)} \right)^{1-\alpha}.$$

- (ii) *The cumulative residual Burg divergence between $\bar{\mathbf{S}}$ and $\bar{\mathbf{Q}}$ is given by*

$$CD_B(\bar{\mathbf{S}}, \bar{\mathbf{Q}}) = \alpha \left[CRB(\bar{\mathbf{P}}) - CRB(\bar{\mathbf{Q}}) \right] - n + \sum_{i=1}^n \left(\frac{\bar{Q}(i)}{\bar{P}(i)} \right)^{\alpha}.$$

where

$$\bar{\mathbf{S}} = \bar{\mathbf{P}}^{\alpha} \bar{\mathbf{Q}}^{1-\alpha}, \quad 0 \leq \alpha \leq 1,$$

is the survival function of the mixture hazards model associated with survival functions $\bar{\mathbf{P}}$ and $\bar{\mathbf{Q}}$.

5 Application of malaria cell images

In this section, we selected two images from a dataset available on the official NIH website: <https://ceb.nlm.nih.gov/repositories/malaria-datasets/>. The dataset consists of images depicting both uninfected and parasitized cells, with a total of 27,558 images. The first image has dimensions of 130×130 pixels, while the second image is 121×124 pixels.

Each image is segmented using the proposed CRRB divergence, and the results are then compared to their corresponding ground truth images. The performance of the proposed algorithm in binary classification and prediction accuracy is assessed using a set of evaluation metrics, including Accuracy, Recall, F1-Score, Dice Coefficient, Jaccard Index, and Adjusted Rand Index (ARI). To demonstrate its effectiveness, the results are compared with two well-established segmentation techniques, namely, Otsu and K-means, highlighting the superior reliability and accuracy of the proposed approach. We implement our analysis using the *OpenCV* library within the *Python* environment [3]. Figures 4 and 3 illustrate the original colored images along with their grayscale versions and corresponding histograms of pixel intensities. These visual representations provide insights into the distribution of pixel values, which play a crucial role in quantifying image similarity.

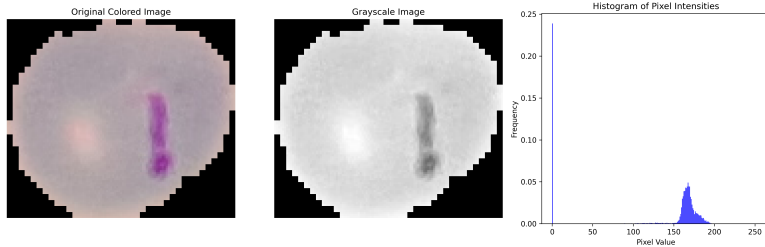


Figure 3: Image A: Original colored image, grayscale version, and histogram of pixel intensities.

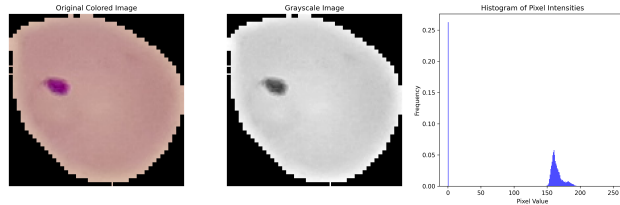


Figure 4: Image B: Original colored image, grayscale version, and histogram of pixel intensities.

Figures 5 and 6 present a comparison between the Ground Truth segmentation and the results obtained using the proposed CRRB divergence, Otsus method, and K-means segmentation. These comparisons highlight the effectiveness of each method in capturing structural details and demonstrate the advantages of the proposed approach in preserving image features.

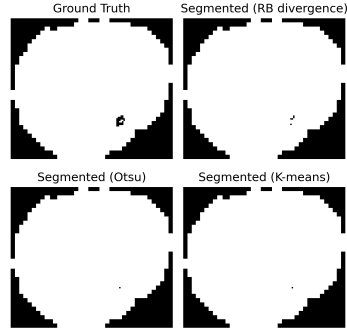


Figure 5: Image B: Comparison of Ground Truth with algorithm I_S^* , Otsu, and K-means segmentation methods

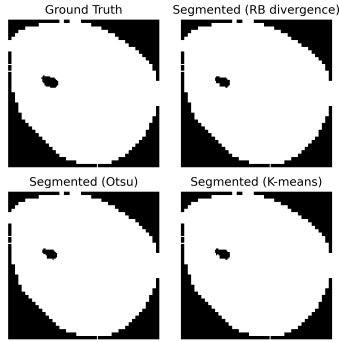


Figure 6: Image A: Comparison of Ground Truth with algorithm I_S^* , Otsu, and K-means segmentation methods

Table 1: Comparison of segmentation methods for two images

Metric	Image 1			Image 2		
	CRI Threshold	Otsu	K-means	CRI Threshold	Otsu	K-means
Threshold (t)	87	83.18	—	84	80.71	—
Accuracy (%)	99.77	99.74	99.74	99.87	99.84	99.86
Recall (%)	100.00	100.00	100.00	100.00	100.00	100.00
Precision (%)	99.70	99.66	99.66	99.83	99.79	99.81
Specificity (%)	99.06	98.92	98.92	99.53	99.42	99.48
F1-Score	0.9985	0.9983	0.9983	0.9991	0.9989	0.9991
Dice Coefficient	0.9985	0.9983	0.9983	0.9991	0.9989	0.9991
Jaccard Index (%)	99.70	99.66	99.66	99.83	99.79	99.81
ARI (%)	99.03	98.88	98.88	99.47	99.35	99.42

Table 1 shows a comparison of three segmentation methods CRI Threshold, Otsu, and K-means based on several performance metrics for two images. The results show that the CRI Threshold method performs better than both Otsu and K-means in all

metrics, including accuracy, precision, specificity, F1-score, Dice coefficient, Jaccard index, and ARI. Specifically, the CRI Threshold method has higher precision and specificity, meaning it is better at avoiding false positives and accurately separating the object from the background. Additionally, it shows better values for F1-score, Dice coefficient, and Jaccard index, indicating that the segmented regions are closer to the ground truth. These results suggest that the CRI Threshold method is a more effective and reliable option for image segmentation compared to the Otsu and K-means methods.

6 Conclusion

In this paper, we have proposed new developments of Burg entropy and have introduced its relative divergence measure. We have extended these concepts to the cumulative residual framework, leading to the definitions of the cumulative residual Burg (CRB) entropy and the cumulative residual relative Burg (CRRB) entropy measures. Furthermore, we have applied the proposed CRRB measure to a real-world image processing problem involving malaria cell images. The numerical results have demonstrated the effectiveness of the CRRB measure in quantifying image similarity and improving segmentation performance. We have compared our proposed approach with established methods such as Otsus thresholding and K -means clustering, showing that the CRRB measure provides competitive classification accuracy. The results have highlighted the potential of our method as a valuable tool in biomedical image analysis.

References

- [1] Bercher, J. F. (2009). Source coding with escort distributions and Rényi entropy bounds. *Physics Letters A*, 373(36), 3235-3238.
- [2] Burg, J. P. (1975). *Maximum entropy spectral analysis*. Stanford University.
- [3] Kaehler, A. & Bradski, G. (2016). *Learning OpenCV 3: Computer Vision in C++ with the OpenCV Library*. O'Reilly Media, Inc.
- [4] Kharazmi, O. & Balakrishnan, N. (2021). Discrete Versions of Jensen-Fisher, Fisher and Bayes-Fisher Information Measures of Finite Mixture Distributions. *Entropy*, 23(3), 363.
- [5] Kharazmi, O., Balakrishnan, N., & Ozonur, D. (2023a). Jensen-discrete information generating function with an application to image processing. *Soft Computing*, 27(8), 4543-4552.
- [6] Kharazmi, O., Contreras-Reyes, J. E., & Balakrishnan, N. (2023b). Jensen-Fisher information and Jensen-Shannon entropy measures based on complementary discrete distributions with an application to Conway's game of life. *Physica D: Nonlinear Phenomena*, 453, 133822.

- [7] Lad, F., Sanfilippo, G., & Agro, G. (2015). Extropy: Complementary dual of entropy. *Statistical Science*, 40-58.
- [8] Martin, M. T., Pennini, F., & Plastino, A. (1999). Fisher's information and the analysis of complex signals. *Physics Letters A*, 256(2-3), 173-180.
- [9] Rao, C. R. (1982). Diversity and dissimilarity coefficients: a unified approach. *Theoretical population biology*, 21(1), 24-43.
- [10] Rényi, A. (1961). On measures of entropy and information. In *Proceedings of the fourth Berkeley symposium on mathematical statistics and probability, volume 1: contributions to the theory of statistics (Vol. 4, pp. 547-562)*. University of California Press.
- [11] Shafee, F. (2007). Lambert function and a new non-extensive form of entropy. *IMA journal of applied mathematics*, 72(6), 785-800.
- [12] Shannon, C. E. (1948). A mathematical theory of communication. *The Bell system technical journal*, 27(3), 379-423.
- [13] Tsallis, C. (1988). Possible generalization of Boltzmann-Gibbs statistics. *Journal of statistical physic*, 52, 479-687.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15. 2025



Binary Classification of Heart Disease Data Based on Cumulative Residual Discrete χ^2_α Divergence

Kharazmi, O.¹ Bahrehvar, M.²

¹ Department of Statistics, Faculty of Mathematical Sciences, Vali-e-Asr University of Rafsanjan, Rafsanjan, Iran

² Department of Industrial Engineering, Tabriz University, Tabriz, Iran

Abstract

In this paper, we propose cumulative residual discrete χ^2 divergence and cumulative residual discrete generalized χ^2_α divergence measures, and establish some theoretical results for three mixture survival models. We further apply the proposed cumulative residual discrete χ^2_α divergence measure to a binary classification task using heart disease data. Our numerical experiments demonstrate the effectiveness of the proposed measure as a robust tool for quantifying classification performance. The performance of our approach is benchmarked against ten machine learning models, including Logistic regression, decision tree, random forest, support vector machine (SVM), K-nearest neighbors (KNN), nave Bayes, linear discriminant analysis (LDA), AdaBoost, and gradient boosting, based on accuracy, precision, recall, and F1-score.

Keywords: χ^2 divergence, Survival function, Classification, Machine learning

¹omidkharazmi14@gmail.com

²minabahrehvar@gmail.com

1 Introduction

Information divergence measures, in both continuous and discrete random variable cases, have widespread applications across various scientific fields such as electrical engineering, physics, economics, statistics, and biology. Two of the most important and widely studied information divergence measures are the KullbackLeibler divergence and the chi-square divergence see [2]. In information theory literature, it has been shown that these two measures are closely related. Information measures in the discrete case are easier to apply compared to the continuous case because, in the continuous setting, estimation and discretization are required. When dealing with discrete information sources, these additional steps are not necessary.

As mentioned chi-square divergence is one of the most important divergence measures in information theory and statistics. Kharazmi and Balakrishnan [3] have presented results related to the chi-square divergence in the continuous case in terms of survival functions. In this paper, we apply a similar idea in the case of discrete random variables, focusing on the problem of classification based on mixture models. We provide results related to two well-known forms of the chi-square divergence in the discrete case, expressed in terms of the survival function, and extend them to mixture models. Finally, we present an application of one of these measures in the context of binary classification.

Let X and Y be two discrete random variables with probability mass functions $\mathbf{P} = (p_1, \dots, p_n)$ and $\mathbf{Q} = (q_1, \dots, q_n)$, respectively. The chi-square divergence and its generalization, the generalized chi-square divergence (also referred to as the discrete χ_α^2 divergence), are defined as follows:

The chi-square divergence is given by

$$\chi^2(\mathbf{P}, \mathbf{Q}) = \sum_{i=1}^n \frac{(p_i - q_i)^2}{p_i}. \quad (1.1)$$

The generalized chi-square divergence is defined as

$$\chi_\alpha^2(\mathbf{P}, \mathbf{Q}) = \frac{\alpha + 1}{2} \sum_{i=1}^n \frac{(p_i - q_i)^2}{p_i^{1-\alpha}}, \quad \alpha \geq 0. \quad (1.2)$$

For more details, one may refer to [4]. These measures quantify the divergence between two probability distributions, with χ^2 being a special case of χ_α^2 when $\alpha = 0$.

The rest of this paper is organized as follows. In Section 2, we first propose cumulative residual discrete χ^2 divergence and cumulative residual discrete generalized χ_α^2 divergence measures and then establish some results associate with three mixture survival models. In Section 3, we apply the proposed cumulative residual discrete χ_α^2 divergence measure to a binary classification task using heart disease data. We present numerical results demonstrating the effectiveness of the proposed measure as a valuable tool for quantifying classification performance. This approach is compared against ten machine learning models. The comparison is conducted in terms of four key performance metrics: Accuracy, Precision, Recall, and F1-score. Finally, we conclude with some remarks in Section 4, summarizing our findings and suggesting potential directions for future research.

2 Cumulative Residual Discrete χ_α^2 Divergence Measure

In this section, we first propose the cumulative residual discrete χ^2 divergence and cumulative residual discrete generalized χ_α^2 divergence measures. We then establish several theoretical results associated with three mixture survival models, including arithmetic mixture, geometric mixture, and harmonic mixture survival functions.

Definition 2.1. Let X and Y be two discrete random variables with probability mass functions

$$\mathbf{P} = (p_1, \dots, p_n) \quad \text{and} \quad \mathbf{Q} = (q_1, \dots, q_n),$$

respectively, defined on a common support $S = \{1, 2, \dots, n\}$. The cumulative residual discrete χ^2 divergence is defined as

$$C\chi^2(\bar{\mathbf{P}}, \bar{\mathbf{Q}}) = \sum_{i=1}^n \frac{(\bar{P}(i) - \bar{Q}(i))^2}{\bar{P}(i)}, \quad (2.1)$$

where

$$\bar{P}(i) = 1 - \sum_{j=0}^{i-1} p_j \quad \text{and} \quad \bar{Q}(i) = 1 - \sum_{j=0}^{i-1} q_j, \quad p_0 = 0.$$

are survival functions related to the PMFs $\mathbf{P}$ and $\mathbf{Q}$, respectively.

Definition 2.2. Let X and Y be two discrete random variables with probability mass functions

$$\mathbf{P} = (p_1, \dots, p_n) \quad \text{and} \quad \mathbf{Q} = (q_1, \dots, q_n),$$

respectively, defined on a common support $S = \{1, 2, \dots, n\}$. The cumulative residual discrete χ_α^2 divergence for $\alpha \geq 0$ is defined as

$$C\chi_\alpha^2(\bar{\mathbf{P}}, \bar{\mathbf{Q}}) = \frac{\alpha + 1}{2} \sum_{i=1}^n \frac{(\bar{P}(i) - \bar{Q}(i))^2}{(\bar{P}(i))^{1-\alpha}}, \quad (2.2)$$

where

$$\bar{P}(i) = 1 - \sum_{j=0}^{i-1} p_j \quad \text{and} \quad \bar{Q}(i) = 1 - \sum_{j=0}^{i-1} q_j.$$

Let X_1 and X_2 be two discrete non-negative random variables with survival functions $\bar{\mathbf{P}}_1$ and $\bar{\mathbf{P}}_2$, respectively. The survival functions corresponding to the arithmetic, geometric, and harmonic mixture models are defined as

$$\bar{\mathbf{P}}_A = w \bar{\mathbf{P}}_1 + (1 - w) \bar{\mathbf{P}}_2,$$

$$\bar{\mathbf{P}}_G = \bar{\mathbf{P}}_1^w \bar{\mathbf{P}}_2^{1-w},$$

and

$$\bar{\mathbf{P}}_H = \frac{\bar{\mathbf{P}}_1 \bar{\mathbf{P}}_2}{w \bar{\mathbf{P}}_2 + (1 - w) \bar{\mathbf{P}}_1},$$

respectively, for $0 < w < 1$. For more details, one may refer to [1] and [2].

The $C\chi_\alpha^2$ divergence between the arithmetic mixture model $\bar{\mathbf{P}}_A$ and the component $\bar{\mathbf{P}}_1$ is given by

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_A) = \frac{\alpha + 1}{2} \sum_{i=1}^n \frac{(\bar{P}_1(i) - w\bar{P}_1(i) - (1 - w)\bar{P}_2(i))^2}{\bar{P}_1^{1-\alpha}(i)}.$$

Expanding and simplifying the expression, we obtain

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_A) = \frac{(\alpha + 1)(1 - w)^2}{2} \sum_{i=1}^n \frac{(\bar{P}_1(i) - \bar{P}_2(i))^2}{\bar{P}_1^{1-\alpha}(i)}.$$

Rewriting in terms of the divergence between $\bar{\mathbf{P}}_1$ and $\bar{\mathbf{P}}_2$, we get

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_A) = (1 - w)^2 C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_2).$$

Similarly, we find that

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_G) = \frac{\alpha + 1}{2} \sum_{i=1}^n \bar{P}_1^{\alpha+1}(i) \left(1 - \left(\frac{\bar{P}_2(i)}{\bar{P}_1(i)} \right)^{1-w} \right)^2.$$

Additionally, the explicit form of the $C\chi_\alpha^2$ divergence for the harmonic mixture is given by

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_H) = \frac{(1 + \alpha)(1 - w)^2}{2} \sum_{i=1}^n \frac{\bar{P}_1^{\alpha+1}(i)}{[(1 - w)\bar{P}_1(i) + w\bar{P}_2(i)]^2} (\bar{P}_1(i) - \bar{P}_2(i))^2.$$

An upper bound for $C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_H)$ is given as

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_H) \leq C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_2)$$

Theorem 2.3. *From the above results, we have*

$$C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_H) \leq \frac{C\chi_\alpha^2(\bar{\mathbf{P}}_1, \bar{\mathbf{P}}_A)}{(1 - w)^2}.$$

3 Application for binary classification of heart data

In this section, we apply the proposed cumulative residual discrete χ_α^2 divergence measure to a binary classification task using heart disease data <https://www.kaggle.com/code/desalegngeb/heart-disease-predictions>. This approach is compared against ten machine learning models, including logistic regression, decision tree, random forest, support vector machine (SVM), K-nearest neighbors (KNN), naive Bayes, linear discriminant analysis (LDA), AdaBoost, and gradient boosting. The comparison is conducted in terms of four key performance metrics: accuracy, precision, recall, and F1-score.

3.1 Variables & Data Preparation

The analysis uses maximum heart rate achieved during exercise (**thalach**) as the sole predictive feature, derived from the UCI Heart Disease dataset. This continuous variable (measured in beats per minute) is standardized via **StandardScaler** to ensure comparability across models. The binary target variable (**target**: 0=healthy, 1=disease) represents clinical diagnosis outcomes. The single-feature design intentionally isolates heart rate’s predictive power while providing a controlled benchmark environment for model comparisons.

3.2 Methodology Overview

A proposed $C\chi^2_\alpha$ divergence threshold method calculates optimal heart rate cutoffs (141.58 bpm) through histogram analysis, maximizing distributional separation between classes. This baseline approach is compared against 10 machine learning models using standard metrics (Accuracy, Precision, Recall, F1-score). Figure 2 illustrates the optimal maximum heart rate threshold selection using the $C\chi^2_\alpha$ divergence, where the dashed red line (labeled **threshold**) marks the point of maximum separation between healthy and diseased distributions.

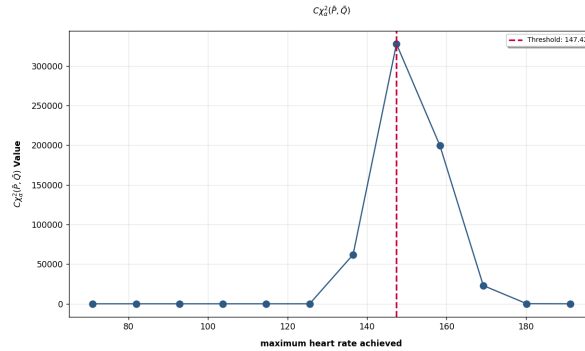


Figure 1: Optimal maximum heart rate threshold selection via $C\chi^2_\alpha$ divergence. Dashed red line (**threshold**) maximizes separation between healthy/diseased distributions.

Figure 1 presents the classification performance of the threshold-based method along with baseline machine learning models (logistic regression, decision tree, and random forest), where the values on the diagonal indicate the correctly predicted cases. Similarly, Figure 3 illustrates the performance of advanced models (SVM, KNN, Naive Bayes, LDA/QDA, AdaBoost, and gradient boosting), with the pronounced diagonal dominance highlighting their strong predictive capabilities.

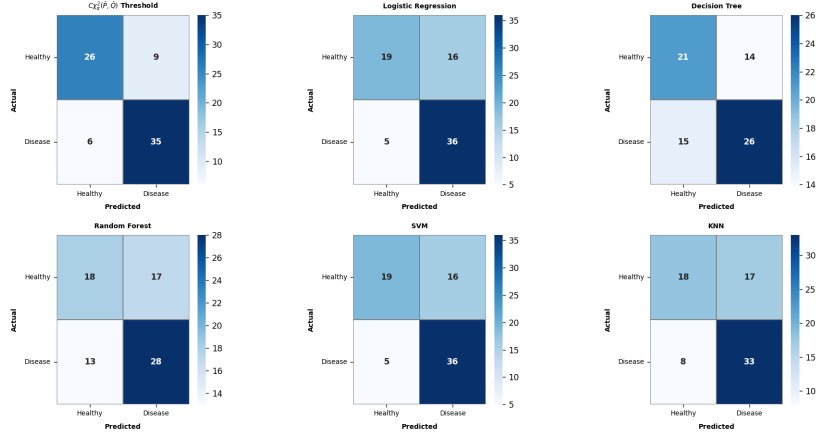


Figure 2: Classification performance of threshold-based method and baseline ML models (logistic regression, decision tree, random forest). Values on the diagonal represent correct predictions.

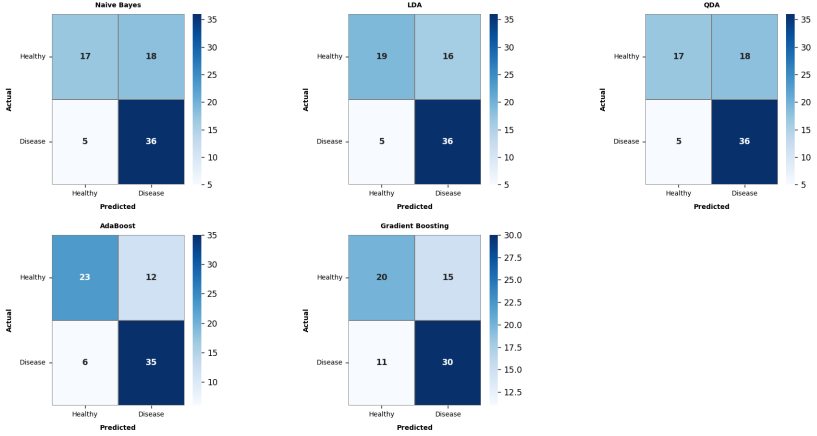


Figure 3: Classification performance of advanced models (SVM, KNN, naive Bayes, LDA/QDA, AdaBoost, gradient boosting). All models show strong diagonal dominance indicating good predictive capability.

Figure 4 compares the ROC curves for the proposed $C\chi^2_\alpha$ threshold-based approach and several machine learning models (Logistic Regression, Decision Tree, Random Forest, SVM, KNN, Naive Bayes, LDA, QDA, AdaBoost, and Gradient Boosting). Notably, the $C\chi^2_\alpha$ threshold achieves the highest AUC (0.80), closely followed by AdaBoost (0.76), indicating strong discriminative performance in distinguishing healthy and diseased samples.

Similarly, Figure 3 illustrates the Precision-Recall relationship across models, with curves nearer to the top-right corner indicating a more effective balance between capturing true positives and minimizing false alarms.

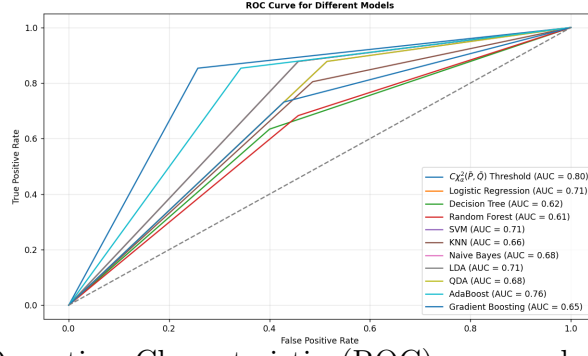


Figure 4: Receiver Operating Characteristic (ROC) curves showing trade-off between True Positive Rate and False Positive Rate. Models closer to the top-left corner (higher AUC) demonstrate better diagnostic performance.

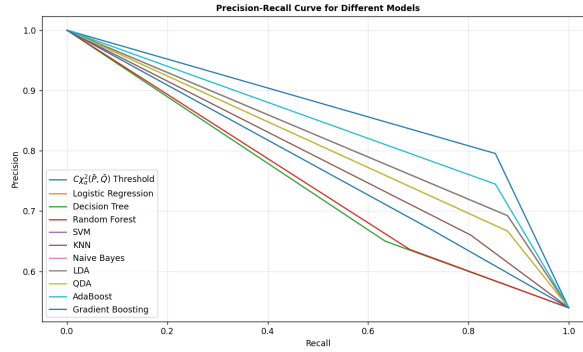


Figure 5: Precision-Recall relationship across models. Curves closer to the top-right corner indicate better balance between identifying true positives and minimizing false alarms.

Table 1 presents the evaluation metrics (Accuracy, Precision, Recall, and F1-score) for the cumulative residual discrete χ^2_α divergence method, as well as for various machine learning models applied to the heart disease binary classification task. The proposed $C\chi^2_\alpha$ method outperforms most of the other models in terms of Accuracy (0.803) and F1-score (0.824), demonstrating its potential as a strong tool for classification. The Logistic Regression, Support Vector Machine (SVM), and Linear Discriminant Analysis (LDA) models show consistent performance, with accuracy scores around 0.724 and F1-scores around 0.774. Decision Tree and Random Forest, on the other hand, exhibit relatively lower performance across all metrics, particularly with Accuracy and F1-score values of 0.618 and 0.642, respectively. Among the other models, AdaBoost performs well, with an Accuracy of 0.763 and F1-score of 0.795, while Gradient Boosting shows the lowest performance in comparison, with an Accuracy of 0.658 and F1-score of 0.698. In general, the $C\chi^2_\alpha$ method proves to be a competitive approach, providing reliable and effective results for heart disease classification tasks.

Table 1: Evaluation Metrics for Different Models

Model	Accuracy	Precision	Recall	F1
$C\chi_\alpha^2(\bar{P}, \bar{Q})$ Method	0.803	0.795	0.854	0.824
Logistic Regression	0.724	0.692	0.878	0.774
Decision Tree	0.618	0.650	0.634	0.642
Random Forest	0.618	0.636	0.683	0.659
SVM	0.724	0.692	0.878	0.774
KNN	0.671	0.660	0.805	0.725
Naive Bayes	0.697	0.667	0.878	0.758
LDA	0.724	0.692	0.878	0.774
QDA	0.697	0.667	0.878	0.758
AdaBoost	0.763	0.745	0.854	0.795
Gradient Boosting	0.658	0.667	0.732	0.698

4 Conclusion

In this paper, we have proposed cumulative residual discrete χ^2 divergence and cumulative residual discrete generalized χ_α^2 divergence measures, and established some theoretical results for three mixture survival models. We have further applied the proposed cumulative residual discrete χ_α^2 divergence measure to a binary classification task using heart disease data. It is shown that the proposed measure serves as a robust tool for quantifying classification performance. Moreover, by benchmarking our approach against ten machine learning models including logistic regression, decision tree, random forest, SVM, KNN, naive Bayes, LDA, AdaBoost, and gradient boosting based on accuracy, precision, recall, and F1-score, we have demonstrated its competitive performance. Future research could explore extensions of these divergence measures to more complex models and additional real-world applications.

References

- [1] Balakrishnan, N., & Kharazmi, O. (2022). Cumulative past Fisher information measure and its extensions. *Brazilian Journal of Probability and Statistics*, 36(3), 540-559.
- [2] Cover, T. M. & Thomas, J. A. (2006). *Elements of Information Theory*. John Wiley & Sons, Inc.
- [3] Kharazmi, O., & Balakrishnan, N. (2021). Cumulative residual and relative cumulative residual Fisher information and their properties. *IEEE Transactions on Information Theory*, 67(10), 6306-6312.

- [4] Kharazmi, O., & Balakrishnan, N. (2023). On Jensen- χ^2_α divergence measure. *Probability in the Engineering and Informational Sciences*, 1, 25.
- [5] Shannon, C. E. (1948). A mathematical theory of communication. *The Bell system technical journal*, 27(3), 379-423.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Causal Inference on the Differences of the Restricted Mean Residual Lifetime

Mansourvar, Z.¹

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

Abstract

In biomedical studies, the difference in restricted mean residual lifetime can be used as an appropriate quantity for comparing different treatment groups with respect to their survival times. In observational studies where the factor of interest is not randomized, covariate adjustment is needed to take into account imbalances in confounding factors. An estimator is obtained for the average causal treatment difference using the restricted mean residual lifetime as target parameter. Specifically, treatment-specific additive hazards model is used to account for confounding factors. Large sample properties of the proposed estimator is also established.

Keywords: Additive hazards model, Average causal effect, Restricted mean residual lifetime

1 Introduction

In medical research, it is often of interest to compare survival times between two treatment groups with respect to their mean residual life function which is the remaining life expectancy of a subject who has survived up to a certain time. For a non-negative survival

¹z.mansourvar@sci.ui.ac.ir

time T with finite expectation, the mean residual life function at time $t \geq 0$ is defined as $m(t) = E(T - t | T > t)$. However, in the presence of right censoring in studies of survival time, the mean residual life function of T may not be estimable at some points. As an alternative to the mean residual life function, it may be more appropriate to use the restricted mean residual life function which for fixed $L > 0$ is $m_L(t) = E(T_L - t | T_L > t)$ ($0 < t < L$), where $T_L = \min(T, L)$. This function is less sensitive to heavily right-skewed survival time distributions than the overall mean residual life function. The restricted mean residual life function of $m_L(t)$ can also be expressed via the survival function of $S(t)$ by

$$m_L(t) = \frac{1}{S(t)} \int_t^L S(u) du. \quad (1.1)$$

If the treatments were assigned to subjects randomly, the restricted mean residual lifetime can be estimated for each treatment group by (1.1) and then the difference between the treatment-specific estimates can be used as the basis for comparing treatment groups. However in observational studies where treatment is not randomized, the distribution of prognostic factors is imbalanced between treatment groups. Thus a proper regression model needs to be selected to adjust for the effects of prognostic factors.

Various authors proposed methods to compare restricted mean lifetime between treatments which is the life expectancy restricted to some time L . [4] examined the restricted mean lifetime with adjustment for covariates and presented a method to compare two groups with respect to survival. To incorporate covariates into the analysis, he considered a proportional hazards model in that separate baseline hazards were postulated but the covariates were assumed to have the same multiplicative effect on the hazard for both groups. The baseline hazards were assumed to take a piecewise exponential model. [3] proposed two estimators for the average causal treatment difference in restricted mean lifetime that account for treatment imbalances in prognostic factors assuming a proportional hazards relationship.

As in survival data analysis, it is important to measure survival expectancy progressively, therefore the restricted mean residual lifetime can be a more useful quantity than restricted mean lifetime. In randomized studies, restricted mean residual life models have been developed under right censored data. For example, [7] studied the proportional restricted mean residual life model, and [8] considered the additive-multiplicative restricted mean residual life model. In observational studies, [6] proposed a method to estimate group-specific differences in restricted mean residual lifetime in settings where covariate adjustment is needed under right censoring. To model the association between the survival time distribution and covariates, they considered the Aalen additive hazards model [1]. In this paper, we propose a method to estimate treatment differences in restricted mean residual lifetime under the additive hazards model [5] which is flexible as it allows for group-specific baseline hazards and regression coefficients.

The rest of the paper is organized as follows. We set up the notation and state the required assumptions needed to formalize the problem of interest in Section 2. The proposed method is described in Section 3 along with asymptotic properties of the resulting estimator with theoretical proofs.

2 Notation and data structure

Suppose we are interested in comparing survival times of two treatment groups ($A = 0, 1$) in terms of restricted mean residual lifetime up to time L . We imagine a set-up where there is a non-randomized binary exposure A and a p -dimensional vector of additional prognostic factors X . Let C denote the censoring time and assume that survival time T and C are conditionally independent given (A, X) . Let (T_i, C_i, A_i, X_i) , where $i = 1, \dots, n$ be independent replicates of (T, C, A, X) . In the presence of censoring, we only observe $U = \min(T, C)$, and whether it is the failure or the censoring that has occurred, recorded by the indicator variable $\Delta = I(T \leq C)$. When right censoring is present, the observed data set thus consists of n independent copies of $\{U_i, \Delta_i, A_i, X_i\}$ and we observe survival times only in the interval $[0, \tau]$ where $\tau < \infty$ is the endpoint of the study.

We use the notation of counterfactual outcomes described by [11, 12], to define the average causal treatment effect. Let T^0 denote the lifetime of a randomly selected individual from the population under study that would have been received treatment 0, and similarly T^1 the lifetime of the same individual that would have been received treatment 1. The variables T^0 and T^1 are referred to as counterfactual outcomes, but only one of those outcomes is actually observed for an individual. Nevertheless, the average causal exposure effect is defined as

$$\delta(t) = \frac{1}{P(T^1 > t)} \int_t^L P(T^1 > u) du - \frac{1}{P(T^0 > t)} \int_t^L P(T^0 > u) du.$$

Actually, each individual will receive only one treatment, 0 or 1, and the observed survival time T is equal to $T = T^0 I(A = 0) + T^1 I(A = 1)$. It can be assumed that the survival time (T^0, T^1) are conditionally independent of the treatment assignment A given X . Under this assumption, we have $P(T^a > t \mid A = a, X) = P(T > t \mid A = a, X)$ for $a = 0, 1$. Then the population-averaged causal exposure effect may be calculated using the G-computation formula ([10]; [9]) for the distribution, $P(T \leq t \mid \hat{A} = a)$. In our setting, the G-computation formula reads

$$P(T \leq t \mid A = a) = \int P(T \leq t \mid A = a, X = x) dF_X(x),$$

where F_X denotes the marginal distribution function of X . Therefore, the corresponding survival function is given by

$$S(t \mid A = a) = \int P(T > t \mid A = a, X = x) dF_X(x). \quad (2.1)$$

Therefore, the average causal treatment effect can be written as

$$\delta(t) = \frac{1}{S(t \mid A = 1)} \int_t^L S(u \mid A = 1) du - \frac{1}{S(t \mid A = 0)} \int_t^L S(u \mid A = 0) du. \quad (2.2)$$

In order to estimate $\delta(t)$ under a right-censored survival data, we need to model the conditional survival function of T given treatment and covariates and derive estimates for the parameters of this model. We set $S_a(t) = S(t | A = a)$ and $S_a(t | X) = P(T > t | A = a, X = x) = \exp\{-\Lambda(t | a, x)\}$ where $\Lambda(t | a, x) = \int_0^t \lambda(u | a, x) du$ is the cumulative conditional hazard function. It is natural to estimate $S_a(t)$ by $\hat{S}_a(t) = n^{-1} \sum_{i=1}^n \hat{S}_a(t | X_i)$ where $\hat{S}_a(t | X_i)$ is an estimator for $S_a(t | X_i)$ and the averaging is over all covariate vectors $X_i, i = 1, \dots, n$, across both treatment groups. Then an estimator for $\delta(t)$ is given by

$$\hat{\delta}(t) = \frac{1}{\hat{S}_1(t)} \int_t^L \hat{S}_1(u) du - \frac{1}{\hat{S}_0(t)} \int_t^L \hat{S}_0(u) du. \quad (2.3)$$

We now turn to the estimation of the effect measure $\delta(t)$ based on data where there are independent replicates from the additive hazards model.

3 Estimation and large sample properties

3.1 Estimation

We consider a more general model where both baseline hazards function and regression coefficients of X are allowed to vary by treatment group, given by

$$\lambda(t | X_i, A_i = a) = \lambda_{0a}(t) + \beta_a^T X_i, \quad a = 0, 1, \quad (3.1)$$

where $\lambda_{0a}(t)$ is an unspecified baseline hazard function for group $A = a$, and β_a is a p -vector of unknown regression parameters. Under model (3.1), the causal survival function defined by (2.1) can be obtained as

$$S_a(t) = S(t | A = a) = \exp\{-\Lambda_{0a}(t)\} \int \exp\{-\beta_a^T X t\} dF_X(x), \quad a = 0, 1$$

where $\Lambda_{0a}(t) = \int_0^t \lambda_{0a}(u) du$ is the cumulative baseline hazards function. Thus, in view of (2.2), it is seen that the average causal exposure effect $\delta(t)$ is a function of $\Lambda_{0a}(t)$ and β_a .

In the counting process framework of [2], the observed data can be replaced with $N_{ia}(t)$ and $Y_{ia}(t)$ of functions of time t , where $N_{ia}(t) = I(A_i = a)I(U_i \leq t, \Delta_i = 1)$ is the counting process and $Y_{ia}(t) = I(A_i = a)I(U_i \geq t)$ is at risk process for the i th individual in treatment group $a = 0, 1$. With this notation, by using the approach of [5], we can estimate β_a and $\Lambda_{0a}(t)$ by

$$\hat{\beta}_a = \left\{ \sum_{i=1}^n \int_0^\tau Y_{ia}(t) \{X_i - \bar{X}_a(t)\}^{\otimes 2} dt \right\}^{-1} \left\{ \sum_{i=1}^n \int_0^\tau \{X_i - \bar{X}_a(t)\} dN_{ia}(t) \right\},$$

where $\bar{X}_a(t) = \sum_{i=1}^n Y_{ia}(t)X_i / \sum_{i=1}^n Y_{ia}(t)$, and

$$\hat{\Lambda}_{0a}(t) = \int_0^t \frac{\sum_{i=1}^n \{dN_{ia}(u) - Y_{ia}(u)\hat{\beta}_a^T X_i du\}}{\sum_{i=1}^n Y_{ia}(u)}.$$

Therefore, the estimator for the average causal exposure effect, $\hat{\delta}(t)$, can be obtained by equation (2.3), where

$$\hat{S}_a(t) = \exp\{-\hat{\Lambda}_{0a}(t)\} n^{-1} \sum_{i=1}^n \exp\{-\hat{\beta}_a^T X_i t\}, \quad a = 0, 1.$$

3.2 Large sample properties

It follows from [5] that $\hat{\beta}_a$ is consistent to β_a , and $\hat{\Lambda}_{0a}(t)$ converges in probability to $\Lambda_{0a}(t)$ uniformly in $t \in [0, \tau]$. Furthermore,

$$n^{1/2}(\hat{\beta}_a - \beta_a) = n^{-1/2} \Omega_a^{-1} \sum_{i=1}^n \epsilon_{ia} + o_p(1), \quad (3.2)$$

$$n^{1/2}(\hat{\Lambda}_{0a}(t) - \Lambda_{0a}(t)) = n^{-1/2} \sum_{i=1}^n Q_{ia}(t) + o_p(1), \quad (3.3)$$

where

$$\Omega_a = E \left[\int_0^\tau Y_{ia}(t) \{X_i - \bar{X}_a(t)\}^{\otimes 2} dt \right],$$

$$\epsilon_{ia} = \int_0^\tau \{X_i - \bar{X}_a(t)\} dM_{ia}(t),$$

$$dM_{ia}(t) = dN_{ia}(t) - Y_{ia}(t) \{d\Lambda_{0a}(t) + \beta_a^T X_i dt\},$$

$$Q_{ia}(t) = \int_0^t \left(\sum_{i=1}^n Y_{ia}(u) \right)^{-1} dM_{ia}(u) - \Omega_a^{-1} \epsilon_{ia} \int_0^t \bar{X}_a(u) du.$$

We now give the asymptotic properties of the proposed estimator for $\delta(t)$ which can be obtained through the asymptotic properties of the influence function of $n^{1/2}\{\hat{\delta}(t) - \delta(t)\}$. To derive the influence function for $n^{1/2}\{\hat{\delta}(t) - \delta(t)\}$, it follows that

$$\begin{aligned} n^{1/2}\{\hat{\delta}(t) - \delta(t)\} &= n^{1/2} \left\{ \frac{1}{\hat{S}_1(t)} \int_t^L \hat{S}_1(u) du - \frac{1}{S_1(t)} \int_t^L S_1(u) du \right\} \\ &\quad - n^{1/2} \left\{ \frac{1}{\hat{S}_0(t)} \int_t^L \hat{S}_0(u) du - \frac{1}{S_0(t)} \int_t^L S_0(u) du \right\}. \end{aligned}$$

It can also be written that, for $a = 0, 1$,

$$\frac{1}{\hat{S}_a(t)} \int_t^L \hat{S}_a(u) du - \frac{1}{S_a(t)} \int_t^L S_a(u) du,$$

is equal to

$$\frac{1}{\hat{S}_a(t)} \int_t^L \left\{ \hat{S}_a(u) - S_a(u) \right\} du - \frac{\left\{ \hat{S}_a(t) - S_a(t) \right\}}{S_a(t) \hat{S}_a(t)} \int_t^L S_a(u) du. \quad (3.4)$$

Therefore, we first need to derive the influence function for $n^{1/2}\{\hat{S}_a(t) - S_a(t)\}$. Define

$$S_{a1}(t) = \exp\{-\Lambda_{0a}(t)\}, \quad S_{a2}(t) = \int \exp\{-\beta_a^T X t\} dF_X(x),$$

so that $S_a(t) = S(t \mid A = a) = S_{a1}(t) S_{a2}(t)$, and also define

$$\hat{S}_{a1}(t) = \exp\{-\hat{\Lambda}_{0a}(t)\}, \quad \hat{S}_{a2}(t) = n^{-1} \sum_{i=1}^n \exp\{-\hat{\beta}_a^T X t\},$$

where $\hat{S}_a(t) = \hat{S}_{a1}(t) \hat{S}_{a2}(t)$. Then it easily follows that

$$n^{1/2}\{\hat{S}_a(t) - S_a(t)\} = n^{1/2}\{\hat{S}_{a1}(t) - S_{a1}(t)\}\hat{S}_{a2}(t) + n^{1/2}\{\hat{S}_{a2}(t) - S_{a2}(t)\}S_{a1}(t). \quad (3.5)$$

Taking the Taylor series expansion of $\hat{S}_{a1}(t)$, together with the consistency of $\hat{\Lambda}_{0a}(t)$ yields that

$$\begin{aligned} n^{1/2}\{\hat{S}_{a1}(t) - S_{a1}(t)\} &= -n^{1/2} S_{a1}(t) \left[\hat{\Lambda}_{0a}(t) - \Lambda_{0a}(t) \right] + o_p(1) \\ &= n^{-1/2} \sum_{i=1}^n \epsilon_i^{S_{a1}}(t) + o_p(1), \end{aligned} \quad (3.6)$$

where

$$\epsilon_i^{S_{a1}}(t) = -S_{a1}(t) Q_{ia},$$

and Q_{ia} is defined as the influence function of $n^{1/2}(\hat{\Lambda}_{0a}(t) - \Lambda_{0a}(t))$, given in (3.3).

By a Taylor series expansion of $\hat{S}_{a2}(t)$ along with the consistency of $\hat{\beta}_a$ and the uniform strong law of large numbers, it further follows that

$$n^{1/2}\{\hat{S}_{a2}(t) - S_{a2}(t)\} = n^{-1/2} \sum_{i=1}^n \epsilon_i^{S_{a2}}(t) + o_p(1), \quad (3.7)$$

where

$$\epsilon_i^{S_{a2}}(t) = \exp\{-\beta_a^T X_i t\} - \mu(t) \epsilon_{ia} - S_{a2}(t),$$

with $\mu(t)$ being the limit in probability of $n^{-1} \sum_{i=1}^n \exp\{-\beta_a^T X_i t\} X_i t$, and ϵ_{ia} is defined as the influence function of $n^{1/2}(\hat{\beta}_a - \beta_a)$, given in (3.2).

Therefore, combining the results from equations (3.5), (3.6) and (3.7) follows that $n^{1/2}\{\hat{S}_a(t) - S_a(t)\}$ can be decomposed as a sum of independent and identically distributed terms as

$$n^{1/2}\{\hat{S}_a(t) - S_a(t)\} = n^{-1/2} \sum_{i=1}^n \epsilon_i^{S_a}(t) + o_p(1), \quad (3.8)$$

where

$$\epsilon_i^{S_a}(t) = S_{a2}(t) \epsilon_i^{S_{a1}}(t) + S_{a1}(t) \epsilon_i^{S_{a2}}(t).$$

Thus, Using expressions (3.4) and (3.8), asymptotic approximation for $n^{1/2}\{\hat{\delta}(t) - \delta(t)\}$ can be displayed by

$$n^{1/2}\{\hat{\delta}(t) - \delta(t)\} = n^{-1/2} \sum_{i=1}^n \epsilon_i^{\delta}(t) + o_p(1),$$

where $\epsilon_i^{\delta}(t) = \epsilon_i^{\delta_1}(t) - \epsilon_i^{\delta_0}(t)$ with

$$\epsilon_i^{\delta_a}(t) = \frac{1}{S_a(t)} \int_t^L \epsilon_i^{S_a}(u) du - \frac{\epsilon_i^{S_a}(t)}{S_a(t)^2} \int_t^L S_a(u) du, \quad \text{for } a = 0, 1,$$

where $\epsilon_i^{\delta}(t)$, $i = 1, \dots, n$, are independent and identically distributed zero-mean random variables. Accordingly, $n^{1/2}\{\hat{\delta}(t) - \delta(t)\}$ converges weakly to a zero-mean Gaussian process whose covariance function $E\{\epsilon_i^{\delta}(s) \epsilon_i^{\delta}(t)\}$ can be consistently estimated by $n^{-1} \sum_{i=1}^n \{\hat{\epsilon}_i^{\delta}(s) \hat{\epsilon}_i^{\delta}(t)\}$ where $\hat{\epsilon}_i^{\delta}(t)$ is defined analogously to $\epsilon_i^{\delta}(t)$, with unknown quantities substituted by their empirical counterparts.

References

- [1] Aalen, O. (1980). A model for nonparametric regression analysis of counting processes. In *Mathematical Statistics and Probability Theory: Proceedings, Sixth International Conference, Wisa (Poland), 1978* (pp. 1-25). New York, NY: Springer New York.
- [2] Andersen, P. K., Borgan, O., Gill, R. D., and Keiding, N. (2012). *Statistical models based on counting processes*. Springer Science and Business Media.
- [3] Chen, P. Y., and Tsiatis, A. A. (2001). Causal inference on the difference of the restricted mean lifetime between two groups. *Biometrics*, **57**(4), 1030-1038.

- [4] Karrison, T. (1987). Restricted mean life with adjustment for covariates. *Journal of the American Statistical Association*, **82**(400), 1169-1176.
- [5] Lin, D. Y., and Ying, Z. (1994). Semiparametric analysis of the additive risk model. *Biometrika*, **81**(1), 61-71.
- [6] Mansourvar, Z., and Martinussen, T. (2017). Estimation of average causal effect using the restricted mean residual lifetime as effect measure. *Lifetime Data Analysis*, **23**, 426-438.
- [7] Mansourvar, Z., Martinussen, T., and Scheike, T. H. (2015). Semiparametric regression for restricted mean residual life under right censoring. *Journal of Applied Statistics*, **42**(12), 2597-2613.
- [8] Mansourvar, Z., Martinussen, T., and Scheike, T. H. (2016). An additive multiplicative restricted mean residual life model. *Scandinavian Journal of Statistics*, **43**(2), 487-504.
- [9] Pearl, J. (2000). Models, reasoning and inference. *Cambridge, UK: Cambridge University Press*, **19**(2), 3.
- [10] Robins, J. (1986). A new approach to causal inference in mortality studies with a sustained exposure period application to control of the healthy worker survivor effect. *Mathematical modelling*, **7**(9-12), 1393-1512.
- [11] Rubin, D. B. (1974). Estimating causal effects of treatments in randomized and nonrandomized studies. *Journal of educational Psychology*, **66**(5), 688.
- [12] Rubin, D. B. (1978). Bayesian inference for causal effects: The role of randomization. *The Annals of statistics*, **6**(1), 34-58.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Analyzing the Doubly Geometric Process with Flexible Interarrival Times: A Tool for Modeling Failure and Reliability Data

Mohammadi, Z.¹

Department of Statistics, Faculty of Science, Jahrom University, Jahrom, Iran

Abstract

The doubly geometric process is widely applied in practical problems, including reliability and failure data analysis. Recently, a lot of scholars have been interested in the doubly geometric process. It is an advanced stochastic model that extends the traditional geometric process by incorporating two distinct scaling factors. Although this process offers greater flexibility in data modeling than the geometric process, it involves more complexity in parameter estimation. This study focuses on estimating the parameters for the doubly geometric process with log-normal interarrival times. The unknown parameters of the process are estimated using the maximum likelihood technique. The estimator's asymptotic properties are also obtained. Simulation results and practical applications are presented to validate the effectiveness of the proposed process.

Keywords: Geometric process, Doubly geometric process, Log-normal distribution, Maximum likelihood estimation

¹zohreh.moh@gmail.com

1 Introduction

For the first time in 1988, the geometric process, denoted by GP in brief, was introduced by Lam [11] as the generalization of the renewal process (RP). The GP is used in many practical problems, such as reliability analysis, maintenance modeling, and failure data analysis. This process is beneficial when the event times are monotone.

Lam's definition [11] states that a series of non-negative random variables $\{X_k, k = 1, 2, \dots\}$ is a GP if the cumulative distribution function (or CDF) of X_k is presented as $F(a^{k-1}x)$ for $k = 1, 2, \dots$, where F is the distribution of X_1 and a is a positive constant. It is a renewal process if $a = 1$, stochastically decreasing if $a \geq 1$, and stochastically increasing if $0 < a \leq 1$. Braun et al. [5], Lam [13], and Lam et al. [16] have all worked on the theoretical aspects of the GP. The reader is directed to Lam [14] and its references for further information.

Parameter estimation and hypothesis testing on the GP have been the subject of several papers. For instance, Lam [12] initially estimates μ , σ^2 , and $\ln a$ using the least squares approach, where μ and σ^2 stand for the expectance and variance of the random variables, respectively. The asymptotic normality of the estimators is carried out by Lam and Zhu [17]. Yang et al. [22] discuss the consistency and asymptotic normality of a GP's ratio parameter by estimating it using the maximum likelihood estimation technique.

The statistical conclusions for the GP are examined assuming that the distributions in the definition above are distinct. For example, Aydoğdu et al. [2] used both the maximum likelihood (ML) approach and the modified maximum likelihood (MML) method to estimate the process parameters while taking into account the GP with a Weibull distribution. ML estimators for GP using the Gamma distribution were obtained by Chan et al. [6]. Lam and Chan [15] examined the statistical inference of the process using the ML and modified moment approaches, taking into account the GP with the log-normal distribution. Biçer [3], Biçer et al. [4], kara et al. [9, 10], and Usta [20] all examined the statistical inference problem for the GP based on the assumption that it adheres to the power Lindley, Rayleigh, inverse Gaussian, Gamma, and inverse Rayleigh distributions. Further information may be found in Usta [21] and its references.

Although the GP is a widely used and influential model, it faces challenges in accurately capturing interarrival times with non-monotonic patterns. In order to overcome these limitations, numerous authors have worked hard to develop the GP. These include the quasi-renewal process suggested by Wang and Pham [23], the α -series process suggested by Braun et al. [5], and the threshold GP model recommended by Chan et al. [7] that was fitted to the 2003 SARS data from four regions. A further illustration is the doubly geometric process (DGP) that was suggested by Wu [24]. The DGP is a sophisticated stochastic model that incorporates two scaling factors to account for trends in both the timing and magnitude of events. The definition of this process is as follows,

Definition 1.1. (Doubly geometric process [24]) When a sequence of non-negative random variables $\{X_k, k = 1, 2, \dots\}$ are independent and the cumulative distribution function of X_k is given by $F(a^{k-1}x^{h(k)})$ for $k = 1, 2, \dots$, where a is a positive constant, $h(k)$ is a function of k , and the likelihood of the parameters in $h(k)$ has a known closed form, $h(1) = 1$ and $h(k) > 0$ for $k \in \mathbb{N}$, then $\{X_k, k = 1, 2, \dots\}$ is referred to as a doubly geometric process.

Wu [24] evaluated the DGP on 10 real-world data sets using several versions of $h(k)$, including $(1 + \log(k))^b$, b^{k-1} , $b^{\log(k)}$, and $1 + b \log(k)$. They then computed the processes' corrected Akaike Information Criterion (AIC_c). He discovered that the DGP with used $h(k) = (1 + \log(k))^b$ performed better than the others. Hence, $h(k)$ is specified as $(1 + \log(k))^b$, where $\log$ is the base-10 logarithm and b is a real value.

The statistical inference issue for the DGP was examined by Pekaip et al. [18, 19] under the assumptions that the distribution of the initial interarrival time has an exponential and a Weibull distribution, respectively.

Estimating the model parameters can be difficult since the DGP offers a more adaptable model for a wider variety of applications than the GP. The estimation of parameters for the model specified by the DGP is the main topic of this study, particularly in cases when the distribution of the initial interarrival time is log-normal. The log-normal distribution is chosen for its critical role in reliability engineering and analysis, particularly due to its ability to model failure times for various products and systems. It is especially useful in scenarios where the underlying processes are multiplicative. As far as we are aware, this is the first work to deal with the DGP's parameter estimate under the presumption that the initial interarrival time has a log-normal distribution.

The reminder of the paper is organized as follows: The doubly geometric process with log-normal interarrival times is introduced, and its primary statistical characteristics are described in Section 2. Section 3 presents the asymptotic joint distribution of the maximum likelihood estimators of the model and distribution parameters, which are obtained by deriving the likelihood equations. Lastly, to assess the effectiveness of the suggested DGP model, Sections 4 and 5 include simulated tests and real-world data applications.

2 Doubly geometric process with log-normal interarrival times

In this section, we introduce a DGP (as defined in Definition 1.1) where the distribution of the first interarrival time follows a log-normal distribution with parameters μ and σ^2 , referred to as DGPLN, and drive its basic properties.

Based on Definition 1.1, we define the DGP with ratio parameter a , $h(k) = (1 + \log k)^b$, where $\log$ represents the base-10 logarithm, and b is a real number, with the log-normal distribution as the distribution of the first interarrival time. Therefore, the probability density function (pdf) of X_1 is

$$f_{X_1}(x) = \frac{1}{x\sigma\sqrt{2\pi}} e^{-\frac{(\ln x - \mu)^2}{2\sigma^2}} \quad x > 0 \quad \mu \in \mathbb{R} \quad \sigma > 0. \quad (2.1)$$

From the definition of the DGP, the pdf of X_k is

$$f_{X_k}(x) = \frac{(1 + \log k)^b}{x\sigma\sqrt{2\pi}} e^{-\frac{((k-1)\ln a + (1+\log k)^b \ln x - \mu)^2}{2\sigma^2}} \quad x > 0, \quad \text{for } k = 1, 2, \dots \quad (2.2)$$

where $b, \mu \in \mathbb{R}$ and $a, \sigma > 0$. For $k = 1, 2, \dots$, the expected value and the variance of X_k are:

$$E(X_k) = a^{\frac{-(k-1)}{(1+\log k)^b}} e^{\frac{\mu}{(1+\log k)^b} + \frac{\sigma^2}{2(1+\log k)^{2b}}},$$

and

$$\text{Var}(X_k) = a^{\frac{-2(k-1)}{(1+\log k)^b}} e^{\frac{2\mu}{(1+\log k)^b} + \frac{\sigma^2}{(1+\log k)^{2b}}} \left(1 - e^{\frac{\sigma^2}{(1+\log k)^{2b}}} \right).$$

Also, the coefficient of variation (CV) of X_k is

$$\gamma_k = \frac{\sqrt{\text{Var}(X_k)}}{E(X_k)} = \sqrt{1 - e^{\frac{\sigma^2}{(1+\log k)^{2b}}}} \quad \text{for } k = 1, 2, \dots,$$

which implies that the CVs change over k 's.

3 Likelihood equations

In this section, we will estimate the parameters of the DGP suggested in the previous section using the maximum likelihood method.

Let $\{x_1, x_2, \dots, x_n\}$ be the gap times from n successive failures have been observed on a system, that is, the failures occurred at times $x_1, x_1 + x_2, \dots, \sum_{k=1}^n x_k$, and the first interarrival time X_1 has the log-normal distribution with parameters μ and σ^2 . Then, the log-likelihood function of the DGP with ratio parameter a and $h(k) = (1 + \log k)^b$ is given by:

$$\begin{aligned} l(\boldsymbol{\theta}) &= \ln L(\boldsymbol{\theta}) = b \sum_{k=1}^n \ln(1 + \log k) - \sum_{k=1}^n \ln x_k - \frac{n}{2} \ln 2\pi\sigma^2 \\ &\quad - \frac{1}{2\sigma^2} \sum_{k=1}^n \left((1 + \log k)^b \ln x_k + (k-1) \ln a - \mu \right)^2, \end{aligned} \quad (3.1)$$

where $\boldsymbol{\theta}$ is the vector of parameters with components $\theta_1 = a, \theta_2 = b, \theta_3 = \mu, \theta_4 = \sigma^2$, and the parameter space $\Theta = \{(\theta_1, \theta_2, \theta_3, \theta_4)^T; \theta_2, \theta_3 \in \mathbb{R}, \theta_1, \theta_4 > 0\}$. The maximum likelihood estimator of $\boldsymbol{\theta}$ is obtained by maximizing $l(\boldsymbol{\theta})$ over Θ .

We denote the maximum likelihood estimate (MLE) by $\hat{\boldsymbol{\theta}}$ and obtain it as a solution of the following equations.

$$\begin{aligned} \frac{\partial l(\boldsymbol{\theta})}{\partial \theta_1} &= \sum_{k=1}^n (k-1) \left((1 + \log k)^b \ln x_k + (k-1) \ln a - \mu \right) = 0, \\ \frac{\partial l(\boldsymbol{\theta})}{\partial \theta_2} &= - \sum_{k=1}^n \left((1 + \log k)^b \ln x_k + (k-1) \ln a - \mu \right) \times \ln x_k \\ &\quad \times (1 + \log k)^b \times \ln(1 + \log k) + \sigma^2 \sum_{k=1}^n \ln(1 + \log k) = 0, \\ \frac{\partial l(\boldsymbol{\theta})}{\partial \theta_3} &= \frac{n(n-1)}{2} \ln a - n\mu + \sum_{k=1}^n (1 + \log k)^b \ln x_k = 0, \end{aligned}$$

$$\frac{\partial l(\boldsymbol{\theta})}{\partial \theta_4} = -n\sigma^2 + \sum_{k=1}^n ((1 + \log k)^b \ln x_k + (k-1) \ln a - \mu)^2 = 0.$$

There is no closed-form solution for the ML estimates, so the equations must be solved numerically. The numerical procedure will be implemented in the R programming language, using the BB package.

We prove the ML estimator's asymptotic normality and consistency in the following assertion.

Proposition 3.1. *The estimator $\hat{\boldsymbol{\theta}}$ is strongly consistent for estimating $\boldsymbol{\theta}$ and satisfy the asymptotic normality*

$$(\hat{\boldsymbol{\theta}} - \boldsymbol{\theta}) \xrightarrow{d} N_4(\mathbf{0}, \Sigma), \quad (3.2)$$

where $\Sigma = I^{-1}$, and I is a 4×4 symmetric matrix, named the Fisher information matrix, with elements given by:

$$I_{11} = \frac{n(n-1)(2n-1)}{6a^2\sigma^2}, \quad I_{12} = \frac{1}{a\sigma^2} \sum_{k=1}^n (k-1)(\mu - (k-1) \ln a) \ln(1 + \log k),$$

$$I_{13} = \frac{-n(n-1)}{2a\sigma^2}, \quad I_{14} = 0,$$

$$I_{22} = \frac{1}{\sigma^2} \sum_{k=1}^n (\ln(1 + \log k))^2 (2\sigma^2 + \mu^2 + (k-1)^2(\ln a)^2 - 2(k-1)\mu \ln a),$$

$$I_{23} = \frac{1}{\sigma^2} \sum_{k=1}^n ((k-1) \ln a - \mu) \ln(1 + \log k), \quad I_{24} = \frac{-1}{\sigma^2} \sum_{k=1}^n \ln(1 + \log k),$$

$$I_{33} = \frac{n}{\sigma^2}, \quad I_{34} = 0, \quad I_{44} = \frac{n}{2\sigma^4}.$$

Proof. The proof of the proposition is given in the Appendix. $\square$

It is possible to deduce that each parameter's asymptotic distribution is asymptotically normal as $n \rightarrow \infty$ based on the aforementioned premise. $\hat{a}$, $\hat{b}$, $\hat{\mu}$, and $\hat{\sigma}^2$ are hence asymptotically unbiased and consistent. Additionally, the estimated asymptotic variance-covariance matrix of $\hat{a}$, $\hat{b}$, $\hat{\mu}$, and $\hat{\sigma}^2$ is determined by I^{-1} evaluated at $\hat{\boldsymbol{\theta}}$, represented by V . For these parameters, the corresponding $100(1 - \alpha)\%$ asymptotic confidence intervals are

$$\left(\hat{a} - z_{\frac{\alpha}{2}} \sqrt{V_{11}}, \hat{a} + z_{\frac{\alpha}{2}} \sqrt{V_{11}} \right), \quad \left(\hat{b} - z_{\frac{\alpha}{2}} \sqrt{V_{22}}, \hat{b} + z_{\frac{\alpha}{2}} \sqrt{V_{22}} \right)$$

$$\left(\hat{\mu} - z_{\frac{\alpha}{2}} \sqrt{V_{33}}, \hat{\mu} + z_{\frac{\alpha}{2}} \sqrt{V_{33}} \right), \quad \left(\hat{\sigma}^2 - z_{\frac{\alpha}{2}} \sqrt{V_{44}}, \hat{\sigma}^2 + z_{\frac{\alpha}{2}} \sqrt{V_{44}} \right)$$

where $z_{\frac{\alpha}{2}}$ is the upper $\frac{\alpha}{2}$ quantile of the standard normal distributions.

4 Simulation study

Employing the bias and mean squared error (MSE) criterion, a simulation analysis is carried out in this part to analyze the estimator's effectiveness for the process parameters a , b , and the distribution parameters μ and σ^2 . The simulation included the subsequent phases for the specified parameters, sample size n , and number of repetitions:

Step 1: Using the parameters μ and σ^2 , create a random sample $\{Y_1, \dots, Y_n\}$ from the log-normal distribution.

Step 2: Using the transformation $X_k = \left(\frac{Y_k}{a^{k-1}}\right)^{\frac{1}{(1+\log k)^b}}$ for $k = 1, \dots, n$, generate a realization of the DGP with parameters a and b .

Step 3: Based on the data set $\{X_1, \dots, X_n\}$, calculate the simulated bias and MSE of each parameter.

Step 4: Repeat the aforementioned steps for the specified number of iterations, and compute the average simulated bias and MSE of each estimator.

According to the above algorithm, we simulate the data for $n = 30, 50, 100$, $a = 0.95, 0.99, 1.05$, $b = -1.5, -0.5, 0.5, 1.5$, $\mu = 0$, and $\sigma^2 = 1$ with 1000 replications. The results of the simulation are presented in Tables 1 and 2.

For all of the estimated parameters, Tables 1 and 2 demonstrate that, generally speaking, the bias and MSE values decrease as n grows. Furthermore, it can be seen that the estimate of the parameters a and μ are closer to the actual values in all cases. For all values of the parameter b , the forecast of parameter b is less than its actual value in all cases. The simulation was also conducted using different values for the parameters μ and σ^2 , yielding the same results. Consequently, these are not included here.

5 Applications

Two data sets are examined in this part to show the value of the suggested process, DGPLN. These data sets, provided in [1] and [8], are the propulsion diesel engine failure data and the aircraft 3 data, respectively. The size of the first data set is 71 and is analyzed by Wu [24] and Pekalp et al. [18, 19]. Wu [24] finds that the DGP outperforms other models on the propulsion data set, and Pekalp et al. [19] demonstrate that this data set is well-fitted by the DGP with Weibull interarrival times. The Cox-Lewis model is used by Cox and Lewis [8] to examine thirteen data sets of the interarrival between subsequent air conditioning system failures in thirteen Boeing 720 aircraft. In order to examine the three biggest aircraft data sets (aircraft 3, 6, and 7), Lam [14] evaluates several models, including the geometric process, the renewal process, the Cox-Lewis model, and the Weibull process model. Lam comes to the conclusion that the geometric process model is the most appropriate. We examine the aircraft 3 data set, which includes 29 interarrival times, in this part.

The proposed model is compared to the following models:

- GP with Weibull interarrival times, GPW [2],
- GP with log-normal interarrival times, GPLN [15],

Table 1: The simulated biases and MSE's for the estimators of the parameters for $\mu = 0$ and $\sigma^2 = 1$.

$b = -1.5$									
a	n	$\hat{a}$		$\hat{b}$		$\hat{\mu}$		$\hat{\sigma}_2$	
		Bias	MSE	Bias	MSE	Bias	MSE	Bias	MSE
0.95	30	-0.0008	0.0017	-0.1921	0.5758	-0.0041	0.1438	0.117	2.0032
	50	-0.0009	0.0009	-0.1154	0.3321	-0.0074	0.0839	0.1266	1.4089
	100	-0.001	0.0005	-0.0567	0.1544	-0.0134	0.0515	0.1212	0.7428
0.99	30	0.0008	0.0006	-0.2161	0.6069	0.0131	0.1439	0.0478	1.2555
	50	-0.0004	0.0002	-0.1195	0.3316	-0.01	0.0784	0.1365	1.1885
	100	0.0003	0	-0.0906	0.1696	0.008	0.0478	0.0615	0.6813
1.05	30	-0.0001	0.0016	-0.2044	0.5578	0.0035	0.1309	0.0551	1.2247
	50	0.0005	0.001	-0.1301	0.329	-0.0101	0.0929	0.0779	0.9496
	100	0.0012	0.0005	-0.0604	0.1564	0.0067	0.0508	0.1085	0.725
$b = -0.5$									
a	n	$\hat{a}$		$\hat{b}$		$\hat{\mu}$		$\hat{\sigma}_2$	
		Bias	MSE	Bias	MSE	Bias	MSE	Bias	MSE
0.95	30	-0.0004	0.0017	-0.1892	0.5805	-0.0047	0.1447	0.1398	1.9052
	50	-0.0004	0.0008	-0.1223	0.3187	-0.0153	0.1054	0.1138	1.0667
	100	-0.0001	0.0004	-0.0733	0.1612	-0.0126	0.0455	0.0662	0.6475
0.99	30	0.0007	0.0005	-0.2017	0.5876	0.0163	0.1212	0.0865	1.4667
	50	-0.0001	0.0001	-0.1218	0.3077	0.0083	0.0903	0.0844	1.1206
	100	-0.0002	0	-0.0686	0.165	-0.0053	0.0401	0.0993	0.6756
1.05	30	0.0034	0.0022	-0.1638	0.5694	0.0061	0.1394	0.1416	1.7815
	50	0.0002	0.001	-0.1417	0.3153	0.0055	0.0871	0.0816	1.1134
	100	0.0006	0.0005	-0.0763	0.1616	0.012	0.0516	0.0751	0.6371

– DGP with Weibull interarrival times, DGPW [19].

To select the best model, we consider two information criteria, the corrected Akaike information criterion (AIC_c) and Bayesian information criterion (BIC), which are defined as

$$AIC_c = -2 \ln \hat{L} + 2k + \frac{2k(k+1)}{n-k-1},$$

$$BIC = -2 \ln \hat{L} + k \ln n,$$

where k is the model's parameter count, n is the sample size, and $\hat{L}$ is the model's maximum likelihood function value. It should be noted that the AIC_c is a modification

Table 2: The simulated biases and MSE's for the estimators of the parameters for $\mu = 0$ and $\sigma^2 = 1$.

$b = 0.5$									
a	n	$\hat{a}$		$\hat{b}$		$\hat{\mu}$		$\hat{\sigma}_2$	
		Bias	MSE	Bias	MSE	Bias	MSE	Bias	MSE
0.95	30	-0.0014	0.0016	-0.1818	0.5601	-0.0119	0.147	0.088	1.47
	50	-0.003	0.0009	-0.0663	0.2941	-0.0138	0.1015	0.1929	1.1897
	100	-0.0014	0.0005	-0.06	0.1747	-0.0142	0.0551	0.1309	0.7662
0.99	30	-0.0009	0.0006	-0.1884	0.5807	-0.0145	0.1429	0.1482	1.9069
	50	0	0.0001	-0.1018	0.3284	0.0052	0.086	0.1708	1.3835
	100	-0.0002	0	-0.0485	0.16	-0.0004	0.0444	0.1429	0.7823
1.05	30	0.0056	0.0026	-0.1664	0.5648	0.0361	0.144	0.1778	3.7664
	50	0.0015	0.001	-0.1141	0.3062	0.007	0.0837	0.1157	1.1342
	100	0.0008	0.0006	-0.071	0.1618	0.0073	0.0473	0.1029	0.7803
$b = 1.5$									
a	n	$\hat{a}$		$\hat{b}$		$\hat{\mu}$		$\hat{\sigma}_2$	
		Bias	MSE	Bias	MSE	Bias	MSE	Bias	MSE
0.95	30	0.0005	0.0014	-0.1981	0.5838	0.0105	0.1229	0.0956	1.5888
	50	0.0001	0.0008	-0.1266	0.3057	0.0037	0.0948	0.0702	0.9848
	100	-0.0002	0.0004	-0.0706	0.1574	-0.0147	0.0465	0.0816	0.6207
0.99	30	-0.0007	0.0007	-0.1662	0.5794	-0.0092	0.1415	0.1592	2.0632
	50	-0.0001	0.0001	-0.1461	0.3261	-0.0042	0.0839	0.0933	1.2909
	100	0	0	-0.0813	0.174	-0.0009	0.0412	0.0862	0.7116
1.05	30	0.006	0.0029	-0.1528	0.5605	0.02	0.1648	0.249	3.5479
	50	0.0025	0.0011	-0.1131	0.332	0.0144	0.0895	0.1422	1.3596
	100	0.001	0.0005	-0.0726	0.168	0.0196	0.0545	0.0835	0.6654

of the Akaike information criterion (AIC) used to measure model performance on small data sets. A model with the smallest AIC_c and BIC values provides the best fit to the data. The fitting results are shown in Table 3. According to this table, the DGPLN model has the best fit based on the AIC_c and BIC .

Furthermore, asymptotic confidence intervals for the model parameters may be obtained using the Fisher information matrix. The 95% asymptotic confidence intervals for the parameters for each data set are shown in Table 4.

Table 3: AIC_c and BIC criteria for the data sets.

Model	Parameters	Data set					
		Propulsion diesel engine failure data			Aircraft 3		
		Estimates	AIC_c	BIC	Estimates	AIC_c	BIC
GPW	$\hat{a}$	0.9853			0.9721		
	$\hat{\alpha}$	6.3168	1413.932	1420.362	1.3618	315.453	318.595
	$\hat{\beta}$	11829			59.852		
GPLN	$\hat{a}$	0.9821			0.9652		
	$\hat{\mu}$	9.1541	1452.79	1459.219	3.6022	312.617	315.759
	$\hat{\sigma}^2$	0.1306			0.6115		
DGPW	$\hat{a}$	1.0076			0.9321		
	$\hat{b}$	-0.3701	1228.848	1237.292	0.2846	317.369	321.172
	$\hat{\alpha}$	21.021			1.1202		
	$\hat{\beta}$	1575.7			91.344		
DGPLN	$\hat{a}$	1.0076			0.9096		
	$\hat{b}$	-0.3839	1213.062	1221.507	0.4134	310.954	314.757
	$\hat{\mu}$	7.2547			4.2029		
	$\hat{\sigma}^2$	0.0022			1.0468		

Appendix

Proof of Proposition 3.1 : Since $Y_k = a^{(k-1)}X_k^{(1+\log k)^b}$, $k = 1, \dots, n$, and Y_k 's are independent and log-normal random variables with the same parameters μ and σ^2 . For $k = 1, \dots, n$, we can easily obtain

$$E(\ln X_k) = (1 + \log k)^{-b}(E(\ln Y_k) - (k-1) \ln a) = (1 + \log k)^{-b}(\mu - (k-1) \ln a), \quad (5.1)$$

and

$$\begin{aligned} E(\ln X_k)^2 &= (1 + \log k)^{-2b} E(\ln Y_k - (k-1) \ln a)^2 \\ &= (1 + \log k)^{-2b} [\sigma^2 + (\mu - (k-1) \ln a)^2]. \end{aligned} \quad (5.2)$$

Table 4: 95% Asymptotic confidence interval for the model parameters.

Parameter	Propulsion diesel engine failure data		Aircraft 3	
	Lower bound	Upper bound	Lower bound	Upper bound
a	1.006796	1.008549	0.779493	1.039764
b	-0.397938	-0.369935	-0.350216	1.1770906
μ	7.189976	7.319935	2.808522	5.597453
σ^2	0.001477	0.002944	-0.205665	2.299267

Then, the elements of Fisher information matrix, $I(\boldsymbol{\theta}) = I_{ij}$, are computed as

$$I_{ij} = E \left[\left(\frac{\partial l(\boldsymbol{\theta})}{\partial \theta_i} \right) \left(\frac{\partial l(\boldsymbol{\theta})}{\partial \theta_j} \right) | \boldsymbol{\theta} \right], \quad (5.3)$$

and under certain regularity conditions, $I_{ij} = -E \left[\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_i \partial \theta_j} | \boldsymbol{\theta} \right]$. Using 5.1 and 5.2, we have

$$I_{11} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_1^2} \right) = \frac{n(n-1)(2n-1)}{6a^2\sigma^2},$$

$$I_{12} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_1 \partial \theta_2} \right) = \frac{1}{a\sigma^2} \sum_{k=1}^n (k-1)(\mu - (k-1)\ln a) \ln(1 + \log k),$$

$$I_{13} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_1 \partial \theta_3} \right) = \frac{-n(n-1)}{2a\sigma^2}, \quad I_{14} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_1 \partial \theta_4} \right) = 0,$$

$$I_{22} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_2^2} \right) = \frac{1}{\sigma^2} \sum_{k=1}^n (2\sigma^2 + \mu^2 + (k-1)^2(\ln a)^2 - 2(k-1)\mu \ln a) \\ \times (\ln(1 + \log k))^2,$$

$$I_{23} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_2 \partial \theta_3} \right) = \frac{1}{\sigma^2} \sum_{k=1}^n ((k-1)\ln a - \mu) \ln(1 + \log k),$$

$$I_{24} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_2 \partial \theta_4} \right) = \frac{-1}{\sigma^2} \sum_{k=1}^n \ln(1 + \log k),$$

$$I_{33} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_3^2} \right) = \frac{n}{\sigma^2}, \quad I_{34} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_3 \partial \theta_4} \right) = 0, \quad I_{44} = -E \left(\frac{\partial^2 l(\boldsymbol{\theta})}{\partial \theta_4^2} \right) = \frac{n}{2\sigma^4}.$$

References

- [1] Ascher, H. and Feingold, H. (1984). *Repairable systems reliability: Modelling, inference, misconceptions and their causes*. New York: Marcel Dekker.
- [2] Aydoğdu, H., Senoğlu, B. and Kara, M. (2010). Parameter estimation in geometric process with Weibull distribution. *Applied Mathematics and Computation*. **217**(6), 2657-2665.

- [3] Biçer, C. (2018). Statistical inference for geometric process with the power Lindley distribution. *Entropy*. **20**, 723-743.
- [4] Biçer, C., Biçer, H. D., Kara, M. and Aydoğdu, H. (2018). Statistical inference for geometric process with the Rayleigh distribution. *Communications Faculty of Science University of Ankara Series A1 Mathematics and Statistics*. **68**(1), 149-160.
- [5] Braun, W. J., Li, W. and Zhao, Y. Q. (2005) Properties of the geometric and related processes. *Naval Research Logistic*. **52**(7), 607-616.
- [6] Chan, S. K., Lam, Y. and Leung, Y. P. (2004). Statistical inference for geometric processes with gamma distribution. *Computational Statistics and Data Analysis*. **47**(3), 565-581.
- [7] Chan, J. S., Yu, P. L., Lam, Y. and Ho, A. P. (2006). Modelling SARS data using threshold geometric process. *Statistics in Medicine*. **25**(11), 1826-1839.
- [8] Cox, D. R. and Lewis, P. A. (1966). *The statistical Analysis of Series of Events*, Mathuen, London.
- [9] Kara, M., Aydoğdu, H. and Türkşen, Ö. (2014). Statistical inference for geometric process with the inverse Gaussian distribution. *Journal of Statistical Computation and Simulation*. **85**(16), 3206-3215.
- [10] Kara, M., Gven, G., enoğlu, B. and Aydoğdu, H. (2022). Estimation of the parameters of the gamma geometric process. *Journal of Statistical Computation and Simulation*. **92**(12), 2525-2535.
- [11] Lam, Y. (1988). A note on the optimal replacement problem. *Advances in Applied Probability*. **20**(2), 479-482.
- [12] Lam, Y. (1992). Nonparametric inference for geometric processes. *Communications in Statistics Theory and Methods*. **21**(7), 2083-2105.
- [13] Lam, Y. (1998). Geometric process and replacement problem. *Acta Mathematicae Applicatae Sinica*. **4**(4), 366-377.
- [14] Lam, Y. (2007). *The Geometric process and its applications*, Word Scientific Publishing Company.
- [15] Lam, Y. and Chan, S. K. (1998). Statistical inference for geometric processes with lognormal distribution. *Computational Statistics and Data Analysis*. **27**, 99-112.
- [16] Lam, Y., Zheng, Y. H. and Zhang, Y. L. (2003). Some limit theorems in geometric processes. *Acta Mathematicae Applicatae Sinica*. **19**(3), 405-416.
- [17] Lam, Y. and Zhu, L. X. (1997). Analysis of a series of events geometric process model. *Acta Mathematicae Applicatae Sinica*. **20**(2), 263-282.

- [18] Pekalp, M. H., İnan, G. E. and Aydoğdu, H. (2021). Statistical inference for doubly geometric process with exponential distribution. *Hacettepe Journal of Mathematics and Statistics*. **50**(5), 1560-1571.
- [19] Pekalp, M. H., Eroğlu İnan, G. and Aydoğdu, H. (2022). Statistical inference for doubly geometric process with Weibull interarrival times. *Communications in Statistics - Simulation and Computation*. **51**(6), 3428-3440.
- [20] Usta, I. (2019). Statistical Inference for geometric process with the inverse Rayleigh distribution. *Sigma Journal of Engineering and Natural Sciences*. **37**(3), 861-872.
- [21] Usta, I. (2022). Bayesian estimation for geometric process with the Weibull distribution. *Communications in Statistics - Simulation and Computation*. **53**(5), 2527-2553.
- [22] Yang A., Yu, H. and Yang, Z. (2006). The MLE of geometric parameter for a geometric process. *Communications in Statistics - Theory and Methods*. **35**(10), 1921-1930.
- [23] Wang, H., and Pham, H. (1996). A quasi renewal process and its applications in imperfect maintenance. *International Journal of Systems Science*. **27**, 1055-1062.
- [24] Wu, Sh. (2018). Doubly geometric processes and applications. *Journal of the Operational Research Society*. **69**(1), 66-77.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



A New Design of a Combined Control Chart Based on Exponentially Weighted Moving Average for Normal Quality Variables

Moradi, N.¹ Salehi, E.²

Department of Industrial Engineering, Birjand University of Technology, Birjand, Iran

Abstract

The production of high-quality products can be effectively achieved through the application of statistical quality control methods, particularly control charts. These tools facilitate precise monitoring of production processes and enable the rapid identification of any deviations. In this study, we focus on designing combined control charts by integrating Double Exponentially Weighted Moving Average (DEWMA) charts and Double Moving Average (DMA) charts. These combined charts are specifically developed to monitor quality characteristics that follow a normal distribution. The performance of these charts is evaluated using the Average Run Length (ARL) index under various parameters. The results demonstrate that the proposed charts significantly outperform traditional control charts in detecting small shifts in the process.

Keywords: Quality control, Control Chart, Double Exponential Weighted Moving Average (DEWMA), Average Run Length (ARL), Normal Distribution

¹n_moradi_stu@birjandut.ac.ir

²salehi@birjandut.ac.ir

1 Introduction

Nowadays, the use of control charts in Statistical Process Control (SPC) is considered an essential tool for quality management and identifying changes in processes. The control chart was first introduced by Shewhart[10]. This type of control chart, due to its focus on the data from the latest sample, only responds to large changes. In contrast, some charts, such as the Moving Average (MA) and exponentially weighted moving average (EWMA), which utilize all data points, have a higher capability to detect small changes. This characteristic allows these charts to perform better than Shewhart charts in monitoring and identifying small process shifts. In recent years, researchers' attention has increased towards control charts like MA, DMA, and EWMA. In a study, Khoo and Wang [6] introduced the DMA chart for quality characteristics with a normal distribution by calculating the MA twice. They demonstrated that the DMA chart performs better than the MA chart in detecting small to moderate changes, but shows similar performance in detecting large changes. Khan et al. [5] designed an EWMA control chart based on the MA statistic for quality characteristics with an exponential distribution. This control chart shows better performance compared to both the EWMA and MA charts in detecting small changes. The DMA control chart based on the EWMA statistic was proposed by Aslam et al. [3] for quality characteristics with an exponential distribution. This chart outperforms the charts developed by Khan et al. [5] and Khu and Wang [6] in detecting small changes. The Double Exponentially Weighted Moving Average (DEWMA) control chart for monitoring quality characteristics with an exponential distribution was designed by Adeoti [1]. This chart demonstrates better performance compared to EWMA and CUSUM charts in detecting process changes. In another study, Taburan and Sukparungsee [12] designed a hybrid EWMA-DMA control chart, combining the strengths of EWMA and DMA to effectively detect small to moderate changes in process parameters under various distributions. Using Monte Carlo simulations and real data, this research shows that the EWMA-DMA chart performs better than classical control charts in detecting small changes. Shanshool et al. [11] proposed a combined Shewhart and EWMA control chart based on the sign statistic for detecting changes in mean and variability in processes. The results indicate that this combined chart improves the ability to detect process changes by striking a balance between the performance of Shewhart and EWMA charts, demonstrating sensitivity to both small and large changes.

In this study, we have designed a novel hybrid control chart for situations where quality variables follow a normal distribution and have evaluated its performance. The paper is organized as follows: Section 2 is dedicated to the design of the hybrid control chart (a combination of DMA and EWMA). In Section 3, the performance of the control charts is assessed based on the Average Run Length. The application of the proposed control chart in the healthcare domain is examined in Section 4. Finally, the conclusions are presented in Section 5.

2 Design of the Control Chart

In this section, we design a new combined control chart, assuming that the quality characteristic of interest (X) follows a normal distribution with mean parameter μ_0 and variance

σ^2 . The combined control chart is presented as follows:

Step 1: In this step, a random sample of size n is selected from the i -th subgroup. The quality characteristic corresponding to these samples is denoted and measured by X_{ij} , where $j = 1, 2, \dots, n$. Then,

$$\bar{X}_i = \frac{1}{n} \sum_{j=1}^n X_{ij}.$$

Step 2: The MA statistic is calculated from the mean of the current and past subgroups with width w at time i , based on the following relation:

$$MA_i = \frac{\bar{X}_i + \bar{X}_{i-1} + \dots + \bar{X}_{i-w+1}}{w}.$$

Step 3: The $EWMA$ statistic at the i -th stage is calculated using the smoothing constant λ ($0 < \lambda \leq 1$), as follows:

$$M_i = \lambda MA_i + (1 - \lambda)M_{i-1} \quad (2.1)$$

Step 4: The DMA statistic is calculated from the mean of the current and past subgroups with width w at time i , based on the following relation:

$$DMA_i = \frac{M_i + M_{i-1} + \dots + M_{i-w+1}}{w}. \quad (2.2)$$

Remark 2.1. The desired statistic up to the third stage was designed by Khan and colleagues [5] for the case where the quality characteristic follows an exponential distribution. Furthermore, Aslam and colleagues [3] studied and investigated the DMA statistic up to the fourth stage for the case where the quality characteristic follows an exponential distribution.

In the following, we design a new control chart DME_i (a combination of the DMA and $EWMA$ charts). The statistic DME_i is derived from the DMA_i 's using the constant λ (where $0 < \lambda \leq 1$), and it is given by:

$$DME_i = \lambda DMA_i + (1 - \lambda)DME_{i-1} \quad (2.3)$$

Step 5: Decision-making is based on comparing DME_i with the control limits UCL and LCL . If DME_i falls between UCL and LCL at any time period, the process is considered in control. Otherwise, the process is out of control and requires investigation and corrective actions.

To calculate the control limits UCL and LCL , the approach similar to the Shewhart control chart can be used. For this, it is necessary to calculate the mean and variance of the DME_i statistic when the process is in control. It is assumed that when the process is in control, μ_0 is the mean of the normal data X_{ij} . In this case, the expected value and variance of the DME_i statistic are given by:

$$\begin{aligned} E(DME_i) &= \mu_0, \\ Var(DME_i) &= \frac{\sigma^2}{nw^2} \left(\frac{\lambda}{2 - \lambda} \right)^2. \end{aligned}$$

Therefore, the control limits UCL and LCL are obtained as follows:

$$UCL/LCL = \mu_0 \pm k \sqrt{\frac{\sigma^2}{nw^2}} \left(\frac{\lambda}{2 - \lambda} \right)^2.$$

where k is the control constant that needs to be determined. Typically, the control constant is set based on the allowable Type I error or the in-control ARL. The probability of detecting the process as in control at μ_0 is given by:

$$P_{in}^0 = P(LCL \leq DME_i \leq UCL | \mu_0).$$

If the statistic DME_i follows a normal distribution, by transforming equation (2.3) and after simplification, it becomes [2]:

$$P_{in}^0 = \Phi(k) - \Phi(-k) = 2\Phi(k) - 1.$$

where Φ is the cumulative distribution function of the standard normal distribution. The performance of each control chart is determined based on the ARL . The in-control Average Run Length, which is solely a function of the control constant k , is given by:

$$ARL_0 = \frac{1}{1 - P_{in}^0}. \quad (2.4)$$

Assume that the process changes from μ_0 to $\mu_1 = \mu_0 + c\sigma$. After calculating the expected value and variance of the DME_i statistic, the probability that the changed process is detected as in control (in other words, the probability of not detecting the change in the process) is given by: $P_{in}^1 = P(LCL \leq DME_i \leq UCL | \mu_1)$. This value can be calculated based on the cumulative distribution function of the standard normal distribution. The Average Run Length (ARL) when the process is out of control for the changed process is given by:

$$ARL_1 = \frac{1}{1 - P_{in}^1}$$

3 Performance Evaluation

The performance of the control chart is evaluated based on the ARL criterion, which is the most widely used measure for evaluating control charts [7]. The ARL is defined as the average number of subgroups until an out-of-control sample is observed. Table 1 shows the values of ARL_1 for different values of $ARL_0 = 300$, the constant $\lambda = (0.1, 0.4, 0.9)$, window width $w = (5, 4, 3)$, change magnitude $c = (0, 0.5)$, for a sample size of $n = 5$ in the proposed control chart. For $ARL_0 = 370$, the control constant $k = 2.9997$, and for $ARL_0 = 300$, $k = 2.9352$ is determined. From the data in these tables, it is observed that smaller values of λ detect changes in the process faster than larger values. Additionally, increasing the width w leads to faster detection of changes in the process. For example, when $ARL_0 = 300$, $\lambda = 0.1$, and $c = 0.01$, for $w = 3$, $ARL_1 = 20.7$, and for $w = 4$, $ARL_1 = 9.2$.

Next, we aim to evaluate the performance of the proposed control chart in comparison with two different control charts. These comparisons are based on simulation data, assuming the data follows a normal distribution with parameters $\mu_0 = 1$ and $\sigma^2 = 1$, as shown in Table ???. The comparisons are conducted for different values of $\lambda = 0.1$ and $n = 5$. Initially, we have proceeded to compare our proposed control chart with the control chart presented by Aslam

and colleagues [3]. Table 3 shows the values of ARL_1 for the proposed control charts and the DMA chart. The obtained results indicate that the proposed control chart detects small process changes significantly faster. As can be seen in Figure 1, the proposed control chart detects that the process is out of control at sample 10, whereas the DMA chart identifies this change at sample 19.

Next, the proposed control chart is compared with the chart presented by Khan et al. [5]. Table 4 shows the ARL_1 values for both charts, and Figure 2 also compares the performance of these two control charts. As can be observed, the proposed control chart goes out of control at sample 10. In contrast, the control limits of the chart by Khan and colleagues [5] are wider than those of the proposed chart and lack the ability to detect this change in the process.

Table 1: The ARL_1 of proposed chart when $ARL_0 = 300$

Change (c)	$w = 3$			$w = 4$			$w = 5$		
	λ			λ			λ		
	0.1	0.4	0.9	0.1	0.4	0.9	0.1	0.4	0.9
0	300.00	300.00	300.00	300.00	300.00	300.00	300.00	300.00	300.00
0.001	278.4	299.4	299.9	263.5	299.0	299.8	246.3	298.4	299.7
0.002	227.8	297.7	299.6	190.4	296.0	299.3	155.8	293.7	298.9
0.005	91.2	286.2	297.6	53.8	276.3	295.8	32.7	264.4	293.5
0.01	20.7	251.0	290.7	9.2	221.9	283.9	4.8	192.3	275.5
0.02	2.9	164.4	265.8	1.5	117.8	243.7	1.10	83.8	219.7
0.05	1.00	37.2	161.5	1.00	18.0	114.9	1.00	9.6	81.2
0.1	1.00	5.6	57.8	1.00	2.5	30.5	1.00	1.50	17.1
0.2	1.00	1.20	10.20	1.00	1.00	4.4	1.00	1.00	2.4
0.5	1.00	1.00	1.10	1.00	1.00	1.00	1.00	1.01	1.00

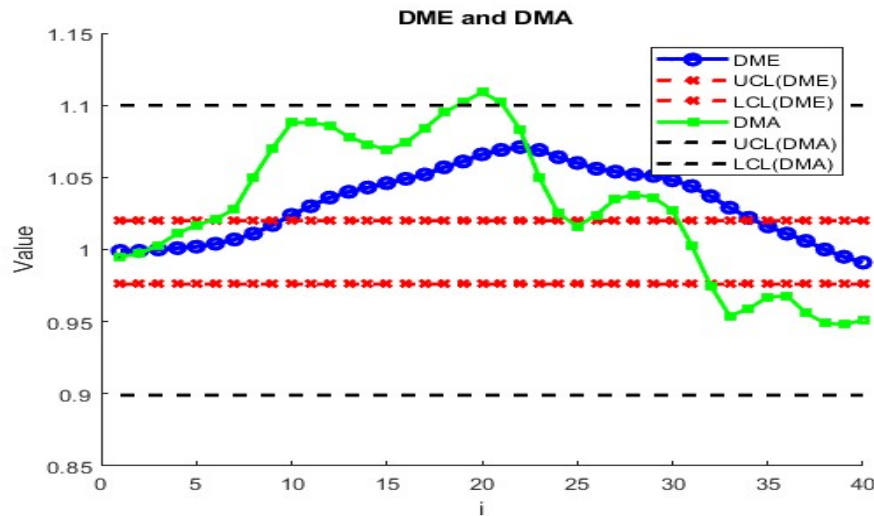


Figure 1: Comparison of the performance of the proposed control chart DME_i and DMA chart by Aslam et al.[3]

Table 2: Simulated data from with parameters $N(1,1)$

i	X_1	X_2	X_3	X_4	X_5	$\bar{X}$	MA	EWMA	DMA	DME_i
1	0.807	0.058	1.071	1.164	1.640	0.948	0.948	0.995	0.995	0.999
2	1.057	0.695	1.948	1.450	0.588	1.148	1.048	1.000	0.997	0.999
3	0.900	2.710	0.966	1.219	0.647	1.288	1.128	1.013	1.003	1.000
4	-0.323	0.595	2.051	0.612	1.209	0.829	1.088	1.020	1.011	1.001
5	1.541	0.816	0.163	1.179	0.504	0.840	0.986	1.017	1.017	1.002
6	0.518	1.104	2.378	1.906	2.446	1.670	1.113	1.027	1.021	1.004
7	0.191	1.855	1.697	0.313	1.012	1.014	1.175	1.041	1.028	1.007
8	1.404	2.070	2.047	1.451	1.132	1.621	1.435	1.081	1.050	1.011
9	0.965	0.351	1.188	-0.398	1.903	0.802	1.145	1.087	1.070	1.017
10	1.445	0.930	0.475	1.936	0.880	1.133	1.185	1.097	1.088	1.024
11	1.379	-0.439	1.154	0.726	1.378	0.840	0.925	1.080	1.088	1.030
12	0.975	1.713	0.883	2.395	0.415	1.276	1.083	1.080	1.086	1.036
13	1.294	-1.025	1.953	1.658	0.616	0.899	1.005	1.073	1.078	1.040
14	2.417	1.266	-1.474	1.110	0.755	0.815	0.997	1.065	1.073	1.043
15	2.151	1.655	2.075	1.851	0.487	1.644	1.119	1.071	1.069	1.046
16	0.638	0.055	2.493	0.800	2.039	1.205	1.221	1.086	1.074	1.049
17	0.143	-0.452	1.797	1.584	0.469	0.708	1.186	1.096	1.084	1.052
18	1.720	1.698	1.252	2.078	1.354	1.620	1.178	1.104	1.095	1.057
19	1.469	1.647	0.142	1.679	0.552	1.098	1.142	1.108	1.102	1.061
20	0.725	0.550	0.318	1.250	1.212	0.811	1.176	1.115	1.109	1.066
21	0.120	1.079	0.639	0.674	0.053	0.513	0.807	1.084	1.102	1.069
22	0.175	2.487	1.519	-0.607	1.191	0.953	0.759	1.051	1.083	1.071
23	-0.052	0.915	1.023	1.544	-0.461	0.594	0.687	1.015	1.050	1.069
24	0.170	4.102	1.644	1.223	-0.632	1.301	0.949	1.008	1.025	1.064
25	1.258	2.895	1.768	1.340	0.944	1.641	1.179	1.025	1.016	1.060
26	1.378	1.095	1.376	-1.044	-0.185	0.524	1.155	1.038	1.024	1.056
27	1.571	1.524	1.917	-0.124	0.185	1.015	1.060	1.040	1.035	1.054
28	1.520	2.679	1.387	1.554	-0.037	1.420	0.986	1.035	1.038	1.052
29	2.652	1.334	-0.509	-0.446	-0.102	0.586	1.007	1.032	1.036	1.051
30	-0.512	0.300	1.216	2.099	-0.566	0.507	0.838	1.013	1.027	1.048
31	1.325	0.282	1.362	-1.042	0.309	0.447	0.514	0.963	1.003	1.044
32	2.442	0.633	1.388	1.503	1.508	1.495	0.817	0.948	0.975	1.037
33	0.290	0.112	1.907	1.244	1.309	0.972	0.972	0.951	0.954	1.029
34	0.761	2.568	0.816	1.063	1.016	1.245	1.237	0.979	0.959	1.022
35	1.455	0.703	-1.116	0.882	0.363	0.457	0.892	0.971	0.967	1.016
36	0.572	-1.161	2.115	1.855	0.229	0.722	0.808	0.954	0.968	1.011
37	0.244	2.082	1.271	1.750	1.561	1.381	0.854	0.944	0.956	1.006
38	1.488	0.985	1.247	-0.377	1.051	0.879	0.994	0.949	0.949	1.000
39	0.134	-0.486	2.050	0.940	0.614	0.650	0.970	0.951	0.948	0.995
40	2.137	0.107	0.623	1.536	2.256	1.332	0.954	0.952	0.951	0.991

4 Application of the Proposed Control Chart

In this section, the performance of the proposed control chart is evaluated using real data from the healthcare sector. The data pertains to the duration (in days) between the admission and

Table 3: Comparison of ARL_1 between proposed control chart DME_i vs DMA Aslam et al. [3]

Change (c)	$w = 3$				$w = 4$			
	$ARL_0 = 300$		$ARL_0 = 370$		$ARL_0 = 300$		$ARL_0 = 370$	
	DME_i	DMA	DME_i	DMA	DME_i	DMA	DME_i	DMA
0	300	300	370	370	300	300	370	370
0.001	278.44	298.79	342.42	368.5	263.51	297.85	323.38	367.28
0.002	227.84	295.21	278.14	363.9	190.35	291.59	231.01	359.2
0.005	91.24	272.20	108.50	334.45	53.79	253.60	63.14	310.70
0.01	20.65	211.17	23.66	257.14	9.23	169.47	10.33	204.9
0.02	2.85	104.24	3.06	124.39	1.47	63.83	1.52	75.24
0.05	1.00	14.21	1.0	16.11	1.00	6.16	1.00	6.81
0.1	1.00	2.01	1.0	2.12	1.00	1.20	1.00	1.22
0.2	1.00	1.00	1.0	1.0	1.00	1.00	1.00	1.00
0.5	1.00	1.00	1.0	1.0	1.00	1.00	1.00	1.00

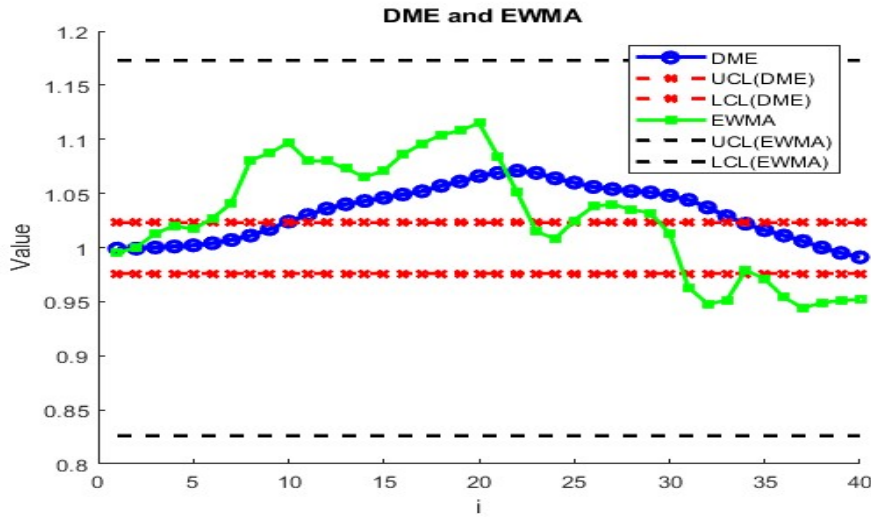


Figure 2: Comparison of the performance of the proposed DME_i control chart and the EWMA chart by Khan et al. [5]

discharge of male patients diagnosed with urinary tract infections (UTI) at a hospital. This dataset was first introduced by Santiago and Smith [9] in the design of the t -chart for the exponential distribution. Based on the study by Aslam et al. [3], the exponential data provided in this research was used and, after transformation to a normal distribution, is presented in Table 5. Considering $\lambda = 0.01$ and $w = 3$, the control limits were calculated based on an $ARL = 300$. The control limits for the proposed chart are: $UCL = 0.595$ and $LCL = 0.596$, while the control limits for the chart by Aslam et al. [3] are $UCL = 0.602$ and $LCL = 0.588$, as shown in Figure 3. As seen in this figure, the DMA control chart is in control, whereas the proposed control chart indicates that the process has gone out of control after sample 17.

Table 4: Comparison of ARL_1 between proposed control chart DME_i vs EWMA chart by Aslam et al. [5]

		$w = 3$				$w = 4$			
		$ARL_0 = 300$		$ARL_0 = 370$		$ARL_0 = 300$		$ARL_0 = 370$	
Change (c)		DME_i	$EWMA$	DME_i	$EWMA$	DME_i	$EWMA$	DME_i	$EWMA$
0		300	300	370	370	300	300	370	370
0.001		278.44	299.59	342.42	369.5	263.51	299.46	323.38	369.34
0.002		227.84	298.39	278.14	368.0	190.35	297.85	231.01	367.28
0.005		91.24	290.18	108.50	357.4	53.79	287.04	63.14	353.42
0.01		20.65	263.94	23.66	323.9	9.23	253.60	10.33	310.78
0.02		2.85	191.32	3.06	232.2	1.47	169.47	1.52	204.94
0.05		1.00	54.54	1.00	64.0	1.00	39.97	1.00	46.58
0.1		1.00	9.41	1.00	1.5	1.00	6.16	1.00	6.81
0.2		1.00	1.54	1.00	1.0	1.00	1.20	1.00	1.22
0.5		1.00	1.00	1.00	1.0	1.00	1.00	1.00	1.00

Table 5: Practical Data in Healthcare Domain (Aslam et al[3])

i	X_1	X_2	X_3	$\bar{X}$	MA	EWMA	DMA	DME_i
1	0.48574	0.54782	0.42293	0.48550	0.48550	0.59441	0.59441	0.59550
2	0.59341	0.20738	0.57385	0.45821	0.47185	0.59318	0.59380	0.59548
3	0.58578	0.55508	0.50718	0.54935	0.49768	0.59223	0.59327	0.59546
4	0.57790	0.43172	0.77715	0.59559	0.53438	0.59165	0.59235	0.59543
5	0.58963	0.36958	0.56296	0.50739	0.55077	0.59124	0.59171	0.59539
6	0.38876	0.73270	0.62489	0.58211	0.56170	0.59095	0.59128	0.59535
7	0.50718	0.88696	0.90865	0.76760	0.61903	0.59123	0.59114	0.59531
8	0.73913	0.58963	0.59712	0.64196	0.66389	0.59195	0.59138	0.59527
9	0.40374	0.30485	0.67776	0.46211	0.62389	0.59227	0.59182	0.59523
10	0.67776	0.40374	0.42293	0.50147	0.53518	0.59170	0.59198	0.59520
11	0.71250	0.80987	0.32433	0.61557	0.52639	0.59105	0.59168	0.59517
12	0.55418	0.65849	1.02394	0.74554	0.62086	0.59135	0.59137	0.59513
13	0.83425	0.71250	0.44007	0.66227	0.67446	0.59218	0.59153	0.59509
14	0.56123	0.84039	0.36958	0.59040	0.66607	0.59292	0.59215	0.59506
15	0.68040	0.59190	0.39319	0.55516	0.60261	0.59301	0.59270	0.59504
16	0.77566	0.83642	0.66968	0.76058	0.63538	0.59344	0.59312	0.59502
17	0.35891	0.49555	0.75276	0.53574	0.61716	0.59368	0.59338	0.59500
18	0.85549	0.55669	0.69569	0.70263	0.66632	0.59440	0.59384	0.59499

5 Conclusion

In this paper, a new control chart for quality characteristics that follow a normal distribution has been studied and evaluated. This chart is designed based on a combination of Double Exponential Weighted Moving Average (DEWMA) and Double Moving Average (DMA) statistics. The control limits of this chart are constructed using the well-known Shewhart-type methodology. Subsequently, the proposed control chart is compared with existing prior control charts,

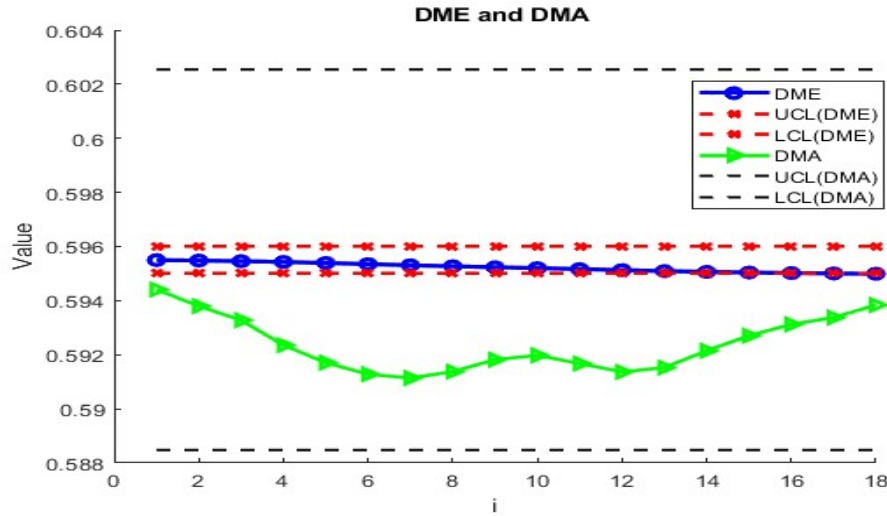


Figure 3: Comparison of the performance of the proposed control chart DME_i by the DMA control chart of Aslam et al. [3] for UTI data

and their performance is evaluated based on the common criterion of Average Run Length (ARL). The results indicate that the proposed control chart outperforms the existing prior control charts, particularly in detecting out-of-control conditions. It demonstrates higher accuracy and efficiency, especially in identifying small shifts in the process.

References

- [1] Adeoti, O. A. (2020), *On control chart for monitoring exponentially distributed quality characteristic*, Trans. Inst. Meas. Control, 42(2), 295–305.
- [2] Aslam, M., Khan, N., Azam, M. and Jun, C. H. (2014b), *Designing of a new monitoring T-chart using repetitive sampling*, Inf. Sci., 269, 210–216.
- [3] Aslam, M., Gui, W., Khan, N. and Jun, C. H. (2017), *Double moving average-EWMA control chart for exponentially distributed quality*, Commun. Stat. Simul. Comput., 46(9), 7351–7364.
- [4] Johnson, N. L. and Kotz, S. (1970), *Distributions in Statistics: Continuous Univariate Distributions, Vol. 2*, Houghton Mifflin.
- [5] Khan, N., Aslam, M. and Jun, C. H. (2016), *An EWMA control chart for exponentially distributed quality based on moving average statistics*, Qual. Reliab. Eng. Int., 32, 1179–1190.
- [6] Khoo, M. B. C. and Wong, V. H. (2008), *A double moving average control chart*, Commun. Stat. Simul. Comput., 37(8), 1696–1708.
- [7] Montgomery, D. C. (2009), *Introduction to Statistical Quality Control*, John Wiley & Sons, Inc.
- [8] Nelson, L. S. (1994), *A control chart for parts-per-million nonconforming items*, J. Qual. Technol., 26(3), 239–240.

- [9] Santiago, E. and Smith, J. (2013), *Control charts based on the exponential distribution: Adapting runs rules for the t chart*, Qual. Eng., 25(2), 85–96.
- [10] Shewhart, W. A. (1926), *Quality control charts*, Bell System Technical Journal, 5(4), 593–603.
- [11] Shanshool, M. K., Mahadik, S. B., Godase, D. G., and Khoo, M. B. C. (2024), *The combined ShewhartEWMA sign charts*, Quality and Reliability Engineering International, 40(7), 4111–4122.
- [12] 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.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Optimized Maintenance Planning Under Random Delayed Interventions: An RL-Based Approach

Rahimi, G. ¹ Karimi, M. ² Rafighi, M. ³ Vesali, N. ⁴ Safari, A. ⁵ Haghighi, F. ⁶

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

Abstract

In this article, we introduce a framework based on reinforcement learning (RL) for optimizing Condition-Based Maintenance (CBM) with inherent random delays. The degradation of the system is represented using a Gamma Process, and decision-making is executed through a Q-learning algorithm. A significant challenge in these systems is managing intervention delays, where an action taken does not produce immediate results. Unlike earlier research that treated the delay as a fixed amount, this study considers it a random variable, making the optimization process more realistic. The simulation study investigates the optimization of railway track maintenance strategies, focusing on the identification of cost-effective intervention policies at each inspection while accounting for both track degradation levels and stochastic intervention delays. The simulation results demonstrate that the proposed RL-based methodology achieves a significant reduction in long-term maintenance costs compared to a fixed-threshold strategy. Specifically, cost reductions of approximately 19% and 45% were observed for

¹ ghazale.rahimi@ut.ac.ir

² mahdiskarimi@ut.ac.ir

³ mohadese.rafighi@ut.ac.ir

⁴ niyayeshvesali@ut.ac.ir

⁵ a.safari@ut.ac.ir

⁶ fhaghighi@ut.ac.ir

scenarios with intervention delays shorter than and longer than the inspection interval, respectively.

Keywords: condition-based maintenance optimization, intervention delay, reinforcement learning, Q-learning

1 Introduction

Optimizing maintenance in systems that undergo gradual degradation over time, such as transportation infrastructure like railways, power plants, and heavy industries, is a critical challenge. These systems are subject to stochastic processes, such as the gamma degradation model, where degradation occurs progressively and unpredictably. The optimization of maintenance strategies for systems exhibiting gamma process to model degradation has been extensively studied over the years. Notably in 2002, Grall et al. explored the application of the gamma process in maintenance modeling [1]. In 2003, Brenguer et al. investigated a system with gamma-based degradation, incorporating continuous monitoring and intervention delays [2].

These traditional maintenance methods, however, often exhibit limitations in their ability to adapt to the dynamic nature of systems and their components. These methods typically rely on simplified assumptions and pre-determined thresholds for triggering specific interventions. This reliance can lead to inefficiencies in real-world environments characterized by various uncertainties. While techniques such as stochastic process modeling or renewal theory are frequently incorporated to enhance these models, they can introduce a degree of complexity. Consequently, with the increasing sophistication of machine learning techniques, particularly reinforcement learning (RL), there is a growing imperative to use these methods to enhance traditional approaches. In recent years, the application of RL in maintenance optimization has gained significant traction, demonstrating its potential to address the limitations of traditional approaches.

Numerous studies have demonstrated RLs crucial role in maintenance decision-making and long-term cost reduction. In infrastructure maintenance management, RL has been utilized to optimize maintenance policies for large-scale structures such as bridges, with models learning Q-values and refining strategies based on historical data [5]. Additionally, recent research has explored the integrated optimization of preventive maintenance and production scheduling in single-machine manufacturing systems. These studies employ Markov Decision Process (MDP) modeling and the R-learning algorithm to derive optimal policies, while novel RL approaches have been proposed to improve efficiency and adaptability in complex manufacturing environments [4].

Collectively, these studies highlight that Reinforcement Learning in maintenance and inspection offers smarter and more efficient decision-making compared to traditional methods, playing a key role in cost optimization and improving the performance of industrial systems. To that end, this study aims to leverage the power of RL in the field of maintenance by point out to one of the key challenges in finding an optimal policy for maintaining a system following a Gamma degradation process, which is intervention delays.

Intervention delays mean that after identifying the need for maintenance, a specific amount of time is required for its execution, which can lead to systems degradation and increase maintenance costs. This issue requires an optimal balance between timely interventions to prevent severe failures and controlling maintenance costs.

In 2009, Meier-Hirmer et al. proposed a maintenance optimization framework for systems undergoing gamma process degradation with fixed intervention delays [3]. Their study, which serves as the primary foundation for this work, employs Markov renewal techniques to compute long-term mean costs. The analysis considers two scenarios: one where the intervention delay is shorter than the inspection interval and another where it exceeds the inspection interval. In each of these scenarios, an optimal standard intervention initialization (SII) threshold is obtained with the purpose of minimizing the long-term mean costs. However, in both cases, the delay remains constant.

In contrast, this research aims to extend the optimization by leveraging the Q-learning algorithm within an RL framework to derive an optimal maintenance policy under conditions where intervention delays are stochastic rather than fixed. By embracing the uncertainty of the world, this makes the study to be more realistic. Our optimal policy has shown a significant mean cost reduction comparing to the optimal policy, obtained from Meier-Hirmer study. In addition, the trained agent learns to do a major intervention after the level of degradation surpasses 1.95. This introduces some sort of a preventive major intervention, rather than corrective, which results the optimal policy to be more conservative, but with a significant reduction in the average long-term cost.

The subsequent sections will first outline the underlying system assumptions. Next, in section 3, we will detail the types of interventions considered and introduce the primary challenge addressed in this work, which is the intervention delay. Following the maintenance model, we will provide a concise overview of the reinforcement learning framework and the principles of Q-learning in section 4. This foundation will then enable us to present the simulation study in section 5, and analyze the strengths and limitations of the proposed methodology in section 6.

2 System Assumption

In this system, degradation increases gradually and unpredictably, and the overall condition of the system is regularly assessed through inspections. As the degradation level increases, maintenance interventions become necessary a detailed explanation of the different type of interventions will be provided in section 3. For now, we focus on the system and its components, which include:

2.1 Degredation Process

The system's condition is represented through an age variable, which gradually increases over time. This age reflects the system's level of degradation, which changes continuously. We use Gamma process to model degradation, a random process that describes the progressive increase in degradation over time. If X_t represents the degradation level at time t , then:

- $X_0 = 0$: This implies that at the beginning of the model, the system has no degradation or damage and is in a healthy state.
- **X_t has independent increments**: This means that changes in the process at any time step are independent of the changes in prior time steps.

- **Increments follow a gamma distribution:** for $s > 0$ and $t \geq 0$, $X_{t+s} - X_t$ follows a Gamma distribution with parameters αs and β . The probability density function (PDF) is given by:

$$f(x; \alpha s, \beta) = \frac{\beta^\alpha}{\Gamma(\alpha s)} x^{\alpha-1} e^{-\beta x} \mathbb{I}(x > 0)$$

where $\alpha > 0$ is the shape parameter, $\beta > 0$ is the rate parameter, and $\Gamma(\alpha)$ is the gamma function.

2.2 Degredation Limit

Degradation Limit (DL) represents the critical boundary beyond which the system is considered to be in an undesirable condition, necessitating a corrective major intervention. Important thing is that the system still works after the level of degradation surpasses this limit, but the option for maintenance is only limited to a major intervention. This limit is usually prespecified by experts in the associated field, therefore finding its optimal value is not the objective of our study.

3 Maintenance Model

3.1 Inspections

The system is inspected at regular intervals δ to assess the degradation level and to evaluate the system's condition. The goal of these inspections is to accurately determine the system's status and whether a maintenance intervention is required. Depending on the context, any amount of time could determine a full interval, like weeks, months, or years.

3.2 Types of Interventions

When the systems condition reaches a level where degradation surpasses an acceptable threshold, maintenance actions begin. These interventions can be either standard intervention or major intervention, depending on the degree of degradation and the system's condition. Standard Intervention can be considered as some sort of preventive maintenance because its purpose is to make a repair at the time when the system is still in a healthy state. Major interventions, on the other hand, could be either preventive or corrective depending on whether the system has reached the state of being heavily deteriorated, which is explainable according to which side of the DL threshold the level of degradation at the corresponding inspection falls.

3.3 Intervention Thresholds

In this study, the maintenance process is governed by three critical thresholds, which dictate when an intervention should be initiated. The Standard Intervention Initiation (SII) threshold marks the point at which the degradation level exceeds a value that poses a potential risk to future performance. When the degradation level exceeds SII at an inspection time, a standard intervention is scheduled. But from some point on, standard intervention is not the best option anymore because system has deteriorated to the level which a more costly repair, like a major intervention is needed.

We introduce a third threshold called MII which stands for Major Intervention Initiation. When systems degradation surpasses this threshold, a preventive major intervention is scheduled. We call it preventive because it is upper bounded by the DL, meaning the system is still in a healthy condition, but needs a major intervention. In addition, after system exceeds the DL level, a corrective major intervention is initiated.

The objective of this study is to find the optimal values for SII and MII, since DL is a pre-specified threshold set out by experts that doesn't need to be estimated. Setting these thresholds involves optimizing the balance between intervention costs and the risks associated with system degradation, requiring careful consideration of the trade-offs involved. The maintenance actions based on the degradation level are summarized in Table 1.

Table 1: Maintenance actions based on degradation level

Degradation Level	Action
Degradation Level < SII	No Action
$SII \leq \text{Degradation Level} < MII$	Standard Intervention
$MI \leq \text{Degradation Level} < DL$	Preventive Major Intervention
Degradation Level $\geq DL$	Corrective Major Intervention

3.4 Intervention Delays

Intervention delays refer to the time gap between identifying the need for maintenance and its actual execution. Both of the interventions are executed with a delay. In this study, the delay is treated as a random variable rather than a fixed value.

- **Delay for Standard Interventions:** The delay for standard interventions is denoted as t_d , which can be either greater or smaller than the inspection interval δ . If the execution of a scheduled standard intervention is delayed more than an inspection interval ($t_d > \delta$), and the degradation level has exceeded DL, then the intervention type is switched to major, and a corrective major intervention is executed within a delay smaller than the inspection interval.
- **Delay for Major Interventions:** The delay for major interventions is denoted as T_D . The delay of a major intervention (preventive or corrective) is always smaller than the inspection interval δ ($T_D < \delta$). This assumption is due to the urgency and critical nature of major interventions, which require immediate attention to prevent further damage or system failure.

3.5 Gain

The gain quantifies the improvement in the systems condition after maintenance is performed and depends on the degradation level at the time of intervention and the type of intervention. The intervention gain comes from a normal distribution $g_{(\mu+\eta x),\sigma}$ with the mean $(\mu + \eta x)$ and standard deviation σ , where μ , η , and σ are given parameters, and x is the degradation level just before the intervention is executed. These parameters can be estimated based on historical data or through regression models that relate the gain to the degradation level before the intervention.

The random gain generated will be subtracted from the current level of degradation. Unlike the usual setup in maintenance models where we assume the gain will put the system in a healthier condition, this is not the case here. It means that the random numbers generated from the mentioned normal distribution are not necessarily positive. The dependence of the mean parameter on the level of degradation leads to the generation of negative numbers associated with smaller degrees of degradation. Ironically, this means that initiating any type of intervention too early not only wont improve the systems condition, but also will worsen it. This could be a bit strange since we always assume that any kind of repair will make the system healthier, but in reality, this is not always true. The random nature of gain and its means dependency on the level of degradation lead us to have a more realistic system that is not always affected in a good way as a result of an intervention.

3.6 Costs

The goal of the model is to minimize long-term costs, which include inspection and intervention costs.

where:

- C_i is the inspection cost,
- C_r is the standard intervention cost,
- C_{dl} is the major intervention cost.

4 Reinforcement Learning

Reinforcement Learning (RL) is a machine learning paradigm in which an agent interacts with an environment and learns optimal decision-making strategies through feedback. Unlike supervised learning, which relies on labeled datasets, RL enables an agent to learn directly from trial and error, using rewards as a guide. A defining characteristic of RL is the presence of non-immediate rewards, where the impact of an action may not be immediately apparent but influences future states and outcomes.

RL is typically modeled as a Markov Decision Process (MDP), which consists of a set of states, actions, transition probabilities, and rewards. The agents objective is to derive a policy that maximizes the expected cumulative reward over time. A key challenge in RL is the exploration-exploitation trade-off the agent must balance leveraging known high-reward actions

(exploitation) with trying new actions to discover potentially better strategies (exploration). Achieving this balance is critical for learning efficiency.

The effectiveness of RL largely depends on how the environments states are defined, as the agent makes decisions based on state information. In the context of maintenance optimization, system degradation serves as the primary state indicator. The degradation level, however, is a continuous variable, which poses computational challenges. To address this, discretization is applied to convert continuous degradation levels into manageable discrete states. This transformation ensures that RL algorithms operate efficiently without encountering infinite state spaces. The specific discretization approach and its impact on learning will be further detailed in the simulation study.

4.1 Q-Learning Algorithm

Q-learning is a model-free algorithm in reinforcement learning, offering an efficient approach to finding optimal policies in sequential decision-making problems. The term model-free highlights the fact that there is no presumed model for how the environment works, which is specified with the transition probabilities. This reduces the RL components to states, actions, and rewards. This algorithm operates based on the Q-function, which represents the expected value of cumulative rewards for each state-action pair. The primary goal of Q-learning is to approximate the Q-function to determine the best action in each state, such that the agent can learn a policy that maximizes the possible reward over time.

In this algorithm, the agent observes a state s , selects an action a , and, after executing that action, the environment returns a new state s' along with a reward value r to the agent. Then, the value of the Q-function for the state-action pair (s, a) is updated as follows:

$$Q(s, a) \leftarrow Q(s, a) + \tau \left(r + \gamma \max_{a'} Q(s', a') - Q(s, a) \right) \quad (4.1)$$

Where:

- $Q(s, a)$ is the value of taking action a in state s .
- τ is the learning rate, which determines how much new information affects previous values.
- r is the reward received after performing action a in state s .
- γ is the discount factor, which specifies the importance of future rewards.
- $\max_{a'} Q(s', a')$ is the maximum Q-value for the new state s' , indicating the expected reward for the best subsequent action.

In practice, an ϵ -greedy policy is commonly used for action selection, where the agent most often chooses the best action based on Q-values, but explores by taking random actions a percentage of the time (ϵ) to encourage further exploration.

The use of Q-learning in optimizing maintenance policies allows the agent to learn a policy that minimizes overall maintenance costs over time, without the computational complexity of the traditional methods. Unlike traditional models that operate based on pre-defined rules, this method learns empirically what actions yield the best results in the long term. In maintenance and repair problems, Q-learning can extract strategies for determining the appropriate time

to perform actions by receiving information from the system’s degradation level. This feature has made Q-learning a compelling alternative to static analytical methods and can be highly effective in complex and uncertain environments.

5 Simulation Study

5.1 Setup

This study investigates a reinforcement learning (RL) approach to optimizing track maintenance policies, leveraging the framework and parameterization established in Meier-Hirmer [3]. The system is checked at regular inspections. As mentioned in the reference study, the inspection interval is 6 weeks. In the referenced work, a Gamma process models the degradation of the track system and a degradation limit of 1.3 is defined, representing the point at which the system requires a major intervention prior to the subsequent inspection. The reference study aims to identify the optimal value for SII, which triggers preventive maintenance actions, by evaluating the asymptotic maintenance cost across a range of SII values under two scenarios: $t_d < \delta$ (scenario 1) and $t_d > \delta$ (scenario 2).

In our RL-based approach, we simulate the track maintenance problem by constructing an MDP environment mirroring the application described in Meier-Hirmer [3]. Consequently, the Gamma process, and gain parameters in the reference article are treated as pre-defined inputs to the RL environment. The formulation of the maintenance problem as an RL task requires the explicit definition of states, actions, and rewards, as detailed in the following sections. With the exception of the state space representation, the action space and reward function are implemented identically across both scenarios. The Python code used to implement the RL environment is available upon request for public use.

States

The state space encapsulates the information necessary for the RL agent to make informed decisions. A key component of the state representation is the level of degradation of the track. Given that degradation is a continuous variable, discretization is performed to enable the application of the Q-learning algorithm. This discretization process involves partitioning the continuous degradation range into a finite set of discrete states. Empirical evaluation, through experimentation with varying numbers of discrete states and threshold values, revealed that a state space comprising nine discrete states yielded the best performance in terms of achieving a near-optimal policy. Table 2 presents the mapping between degradation levels and their corresponding discrete states.

The degradation level at each inspection is determined by a Gamma process, parameterized as described in Meier-Hirmer study [3]. However, a modification was introduced to the Gamma process parameters to accelerate the learning process. Specifically, the rate of degradation was increased by a factor of 25. This adjustment was implemented to mitigate the slow learning convergence that resulted from the relatively slow degradation rate characteristic of real-world rail tracks. This parameter adjustment solely aims to reduce the computational time required for the agent to explore the state-action space and converge to an optimal policy and does not alter the fundamental interaction between the agent and the environment or the underlying cost

Table 2: Degradation level discretized to states

Level of Degradation	State
$0 < deg \leq 0.2$	0
$0.2 < deg \leq 0.4$	1
$0.4 < deg \leq 0.6$	2
$0.6 < deg \leq 0.8$	3
$0.8 < deg \leq 0.95$	4
$0.95 < deg \leq 1.195$	5
$1.195 < deg \leq 1.25$	6
$1.25 < deg < 1.3$	7
$1.3 \geq deg$	8

structure.

In the scenario where $t_d > \delta$, the state representation includes an additional component representing the delay incurred from a previously performed standard intervention. This component, denoted as $t_{d,before}$, quantifies the time elapsed since the last standard intervention. Therefore, in this scenario, the state space consists of both the discretized degradation level and the $t_{d,before}$ value, providing the agent with the necessary information to account for the delayed effects of maintenance actions.

Actions

At each inspection, the RL agent has access to a discrete action space comprising three possible actions:

1. Do Nothing (inspect)
2. Standard Intervention
3. Major Intervention

Each intervention action results in a stochastic reduction in the level of degradation, characterized by a gain sampled from a pre-defined distribution. The parameters of this distribution are directly derived from Meier-Hirmer study [3] to ensure consistency with the reference model.

Rewards

The reward function provides feedback to the RL agent, reflecting the consequences of its actions in each state. In this simulation, costs are used to define the reward structure. A fixed inspection cost of 50 was considered. A *Standard Intervention* incurs a cost of 1000, while a *Major Intervention* incurs a cost 100 times greater than a *Standard Intervention* (i.e., 100,000). These cost values are chosen to maintain consistency with the relative cost ratios established in Meier-Hirmer study [3]. Rewards are then calculated as the negative of these costs, incentivizing the agent to minimize maintenance expenditures. Table 3 represents the reward matrix for this study.

To ensure appropriate policy convergence, given the pre-defined degradation threshold of 1.3, the reward structure for the terminal state was designed to strongly incentivize the selection of a major intervention, reflecting the expert-defined necessity for such action at this degradation level, which has passed the DL.

Table 3: Rewards function matrix

	Do nothing (inspect)	Standard intervention	Major intervention
0	-50	-1050	-100050
1	-50	-1050	-100050
2	-50	-1050	-100050
3	-50	-1050	-100050
4	-50	-1050	-100050
5	-50	-1050	-100050
6	-50	-1050	-100050
7	-50	-1050	-100050
8	-2000050	-2001050	-400050

At each inspection, based on the action chosen, the agent is charged with a corresponding cost, which has already been introduced in Section 3. The objective is then to learn an optimal policy that minimizes the overall cost with the following objective function:

$$C(t) = C_i N_i(t) + C_r N_r(t) + C_{dl} N_{dl}(t) \quad (5.1)$$

where $N_i(t)$, $N_r(t)$, and $N_{dl}(t)$ represent the number of inspections, standard interventions, and major interventions at time t , respectively.

Parameter specification

There are two sets of parameters (and hyperparameters) that need to be specified. One set is entitled to the system and the maintenance model, while the other consists of hyperparameters used in Q-learning algorithms:

- Degradation process: The degradation is modeled using a Gamma process with the shape parameter $\alpha = 1.72$ and the rate parameter $\beta = 0.0238$, with a rate of degradation approximately equal to 0.04.
- Delays: In the first scenario where $t_d < \delta$, the delay for both interventions is randomly generated from a uniform distribution $U(0, 1)$. For the second scenario, however, the delay of the major intervention remains the same, but the standard intervention is executed with a delay that comes from a new uniform distribution $U(1, 2)$.
- Gain: As mentioned before, gain is a randomly generated number from a normal distribution $g_{(\mu+\eta x), \sigma}$. The estimation μ, η , and σ for the gain corresponding to a standard intervention are -0.39 , 0.8 , and 0.22 . On the other hand, the corresponding estimates for

the gain of a major intervention are -0.45 , 1.06 , and 0.08 , respectively. Figure 1 also provides a clear visualization of the mean gain for the two types of interventions, consistent with the explanation given in Section 3.

- Learning rate: This hyperparameter, denoted as τ , is first initialized at 1 and is decayed by 1% after each step.
- Epsilon in ϵ -greedy: This hyperparameter is also initialized and decayed similarly to the learning rate.
- Discount factor: $\gamma = 0.8$

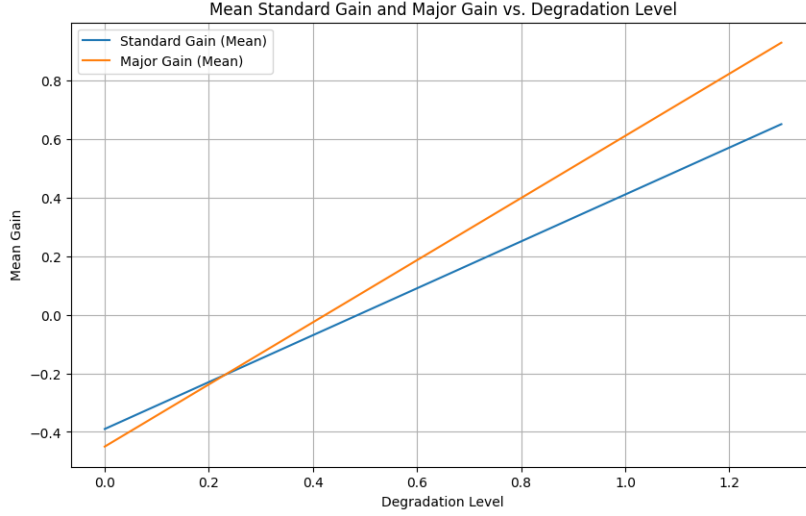


Figure 1: Mean standard and major gain vs. degradation level

- Delay: to change the constant nature of the delay proposed in the reference article, here we assume that delay is a random number, generated from a uniform distribution. For the first scenario the distribution for the delay of the standard intervention would be $U(0, 1)$ and for the second scenario this changes to $U(1, 2)$.

5.2 Results

Here we present the results obtained from the simulation which is implemented separately for each scenario. At each of them, an optimal value for two thresholds, SII and MII, have been computed. The optimal policy given by the Q-learning algorithm is then compared to the optimal policy derived in the Meier-Hirmer study. In other words, the mean cost of the optimal policy derived from our proposed method and the one proposed in the reference article, each corresponding to the founded optimal thresholds, have been compared.

Scenario 1: $t_d < \delta$

In this scenario any initiated intervention will be delivered at some point before the next inspection. With that in mind and incorporating the components defined earlier (states, actions and

rewards), the RL environment has been set up. Figure 2 shows the validity of the setup environment. At this point the agent is not trained yet and the purpose is to check if the environment is working in a legit, acceptable way.

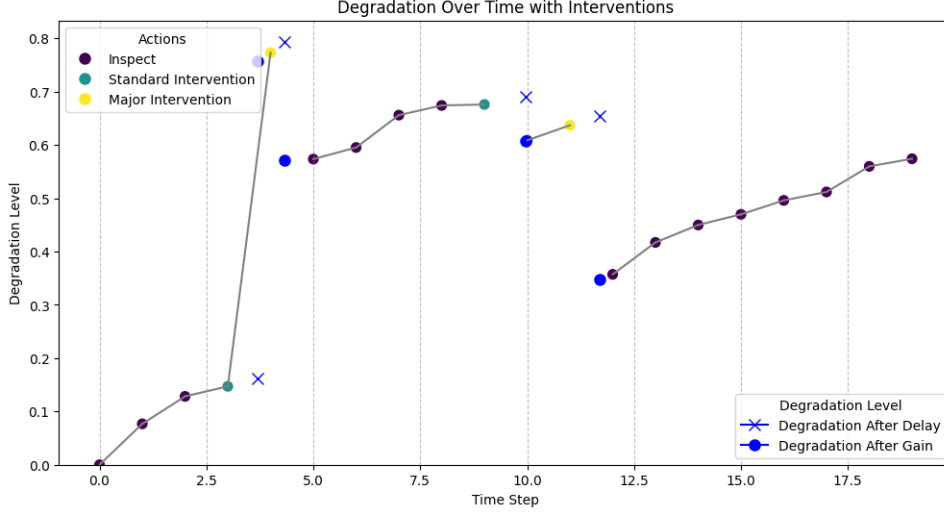


Figure 2: Degradation over time with interventions ($t_d < \delta$)

Following validation of the simulation environment, the Q-learning algorithm was employed to train the RL agent. The training process consisted of 2000 episodes, each comprising 5000 steps. Each step represents an inspection event, effectively discretizing time such that inspections occur at unit intervals (i.e., $\delta = 1$). This unit interval can be interpreted as a representative time period, such as six weeks. This temporal discretization and the interpretation of steps as inspection events hold true for both experimental scenarios.

After the agent is trained, an optimal policy was obtained from the results of the Q-table. The agent learns to defer intervention (choosing "Do Nothing") during the initial phases of degradation (states 0-4). This reflects a cost-effective strategy of allowing the system to operate without unnecessary intervention when its condition is good. When the degradation level reaches state 5, the agent initiates a "Standard Intervention," indicating that it recognizes the need for a preventative standard maintenance to mitigate further degradation. Finally, in states 6, 7, and 8, the agent consistently selects "Major Intervention," demonstrating its understanding of the necessity for preventive and corrective major interventions when the system reaches a critical degradation threshold. This policy specifies the optimal values of SII and MII, which are shown and compared to the optimal policy given by the reference article in Table 4. Additionally, Figure 3 demonstrates the trained agent going through an episode of 100 steps.

Table 4: Comparison between our optimal policy and Meier-Hirmer optimal policy ($t_d < \delta$).

Obtained thresholds	Our optimal policy	Obtained thresholds	Meier-Hirmer study optimal policy
$\text{deg} \leq 0.95$ (SII)	Do nothing	$\text{deg} \leq 0.98$ (SII)	Do nothing
$0.95 < \text{deg} \leq 1.95$ (MII)	Standard Intervention	$\text{deg} < 1.3$ (DL)	Standard Intervention
$1.95 < \text{deg} < 1.3$	Preventive Major Intervention	$\text{deg} \geq 1.3$	Corrective Major Intervention
$\text{deg} \geq 1.3$	Corrective Major Intervention		

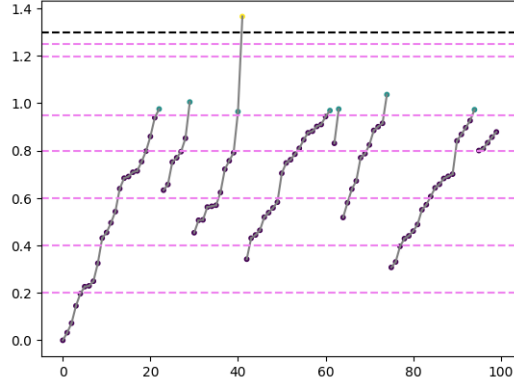


Figure 3: Trained agent on an episode of 100 steps ($t_d < \delta$)

To compare our optimal policy with the reference article, we let the trained agent go through 1000 episodes, each consisting of 500 steps. The mean cost derived from this procedure resulted in -238.564 , while the optimal policy given by the reference article has a mean cost of -294.478 . This shows a significant reduction in the mean cost by roughly 19%. By incorporating the use of confidence intervals, a mean cost distribution for both of these policies is demonstrated in Figure 4. The histogram shows a clear shift in the reward distribution towards higher values compared to the reference policy. This visually confirms that, on average, the RL agent achieves better outcomes, with a more favorable distribution of maintenance costs.

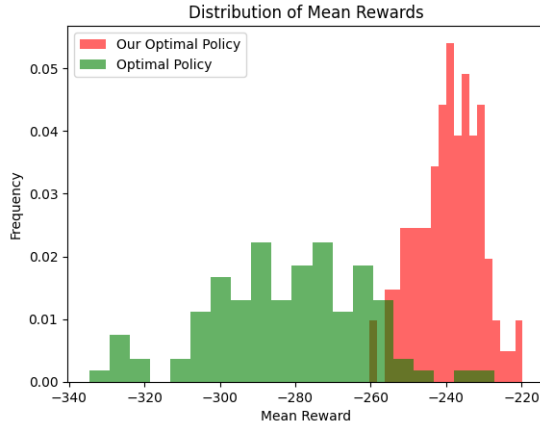


Figure 4: Comparison of distribution of mean costs between our optimal policy and Meier-Hirmer optimal policy ($t_d < \delta$)

Scenario 2: $t_d > \delta$

In this scenario, standard interventions are executed within a time delay longer than the inspection interval. As mentioned before, there is a slight change in the definition of states when we are in this scenario, which involves incorporating the use of $t_{d,\text{before}}$ denoting the time of delay after the last standard intervention call. This means that at each inspection, there is an additional component involved to explain the systems state, which depends on whether it is still in the delay period of the previously called standard intervention (`done = False`), with a randomly generated value for $t_{d,\text{before}}$ or there is no $t_{d,\text{before}}$ involved and the agent is ready to take a new

action (`done = True`). Figure 5 illustrates a sample trajectory of the agent’s actions, where interventions are followed by a delay and subsequent gain resulting from the delayed action.

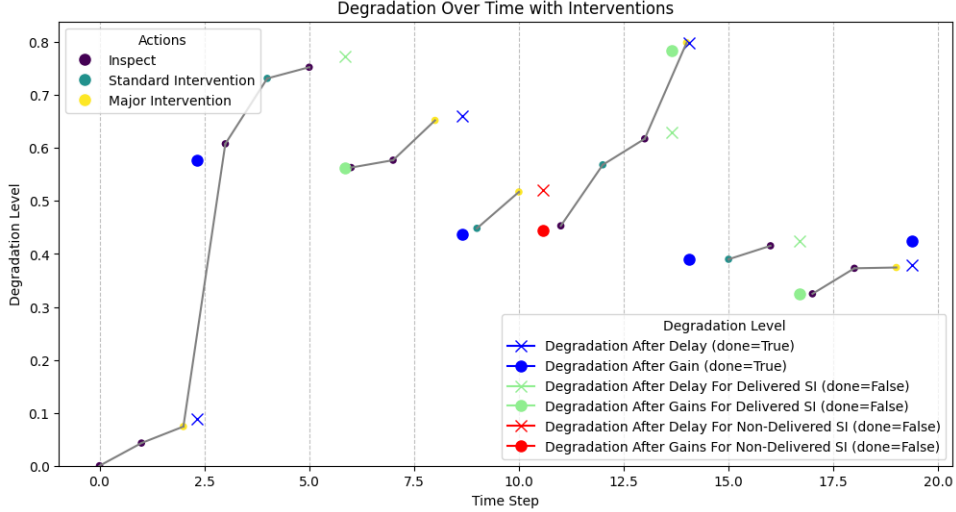


Figure 5: Degradation over time with interventions ($t_d > \delta$)

Since in this scenario, the environment becomes more complex, the learning process of the agent is affected. Therefore, we increased the number of steps to 10000 for this scenario. After running the Q-learning algorithm for 2000 episodes, an optimal policy is derived which is shown in table 5. In contrast to our expectation that this scenario would lead to a more conservative policy with smaller optimal value of SII, the optimal policy derived for both scenarios are the same. This may suggest that an optimal policy associated with less amount of mean cost might exist, but since our obtained policy has already shown a better performance comparing to the one from the reference article, the search for even a better policy was not our priority. Moreover, for a clearer visualization, figure 6 represents the trained agent, going through a 100- step episode.

Table 5: Comparison between our optimal policy and Meier-Hirmer optimal policy ($t_d > \delta$).

Obtained thresholds	Our optimal policy	Proposed thresholds	Meier-Hirmer study optimal policy
$\text{deg} \leq 0.95$ (SII)	Do nothing	$\text{deg} \leq 0.74$ (SII)	Do nothing
$0.95 < \text{deg} \leq 1.95$ (MII)	Standard Intervention	$0.74 < \text{deg} < 1.3$ (DL)	Standard Intervention
$1.95 < \text{deg} < 1.3$	Preventive Major Intervention	$\text{deg} \geq 1.3$	Corrective Major Intervention
$\text{deg} \geq 1.3$	Corrective Major Intervention		

Like the first scenario, After the agent is trained, we let it run for 1000 episodes, each consisting of 500 steps, to access an estimation of the mean costs. The mean cost for our obtained policy was -434.216, while the corresponding policy based on the reference has an average cost of -793.394. This indicates that the reduction of the mean of costs is even more significant in this scenario, being around 45%. Figure 7 shows the distribution of average costs for both policies which compared to the first scenario, visualizes a non-overlapping situation for both of the policies.

Overall, results suggest that finding an adjusted optimal value for SII, alongside with introducing a second threshold MII which lead us to incorporating not only the corrective major intervention, but also a preventive major intervention, lead to a significant cost reduction in

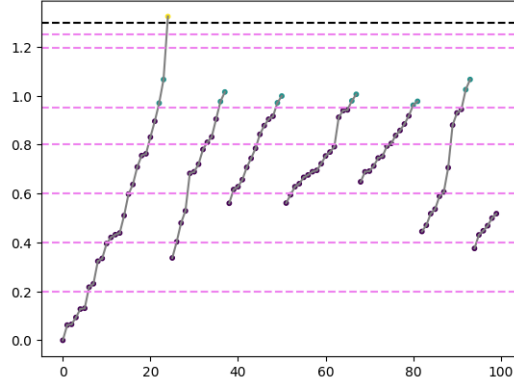


Figure 6: Trained agent on an episode of 100 steps ($t_d > \delta$)

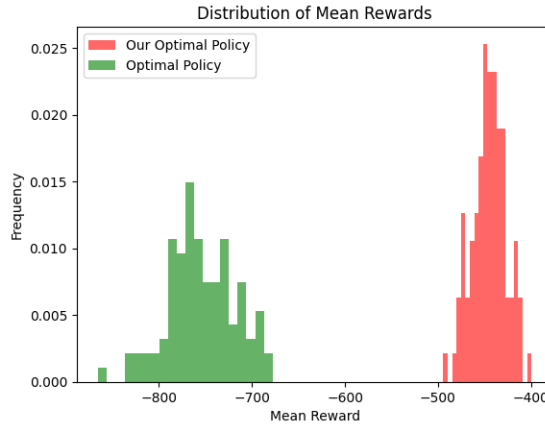


Figure 7: Comparison of distribution of mean costs between our optimal policy and Meier-Hirmer optimal policy ($t_d > \delta$)

both scenarios. Additionally, our study suggests by introducing a level of uncertainty to the delays, making them random instead of constant, the agent still manages to learn an optimal policy. Figure 8 validate this by demonstrating the convergence of rewards in both scenarios, after agent is trained for 2000 episodes.

6 Conclusion and Discussion

Our proposed optimal policies for both scenarios ($t_d < \delta$ and $t_d > \delta$) showed a significant decrease in long-term average cost compared to the policies presented in the Meier-Hirmer study. By adding uncertainty to the system through incorporating random delays, the Q-learning algorithm leads to an optimal policy without much time complexity. The use of the Q-learning algorithm provided a flexible and adaptable simulation environment, suitable for testing and trialing different sets of parameters and reward functions.

The proposed method showed some promising results; however, some initial restrictions were taken into account when implementing the simulation. For example, in the second scenario, if a standard intervention is initiated at an inspection, then at the next inspection when the system is still within the delay of the initiated intervention, the agent can initiate a second standard intervention with a delay that must be less than the inspection interval ($t_d < \delta$). In other words, our model is restricted to the assumption that a second call for a standard

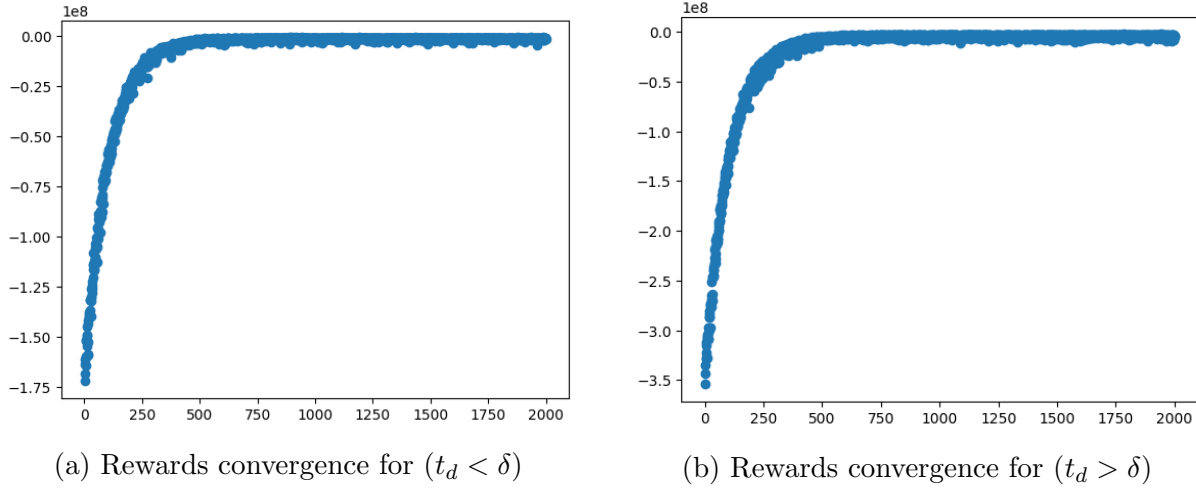


Figure 8: Comparison of Rewards Convergence for Different Scenarios

intervention, while still in the delay of the previous one, must lead to its execution prior to the next inspection. Therefore, future studies can focus on incorporating $t_d > \delta$ for all standard intervention initiations, even when the system is in delay (`done = False`).

Another aspect that can be improved is the introduction of a third scenario, where a random delay for standard intervention can be initiated, covering both scenarios ($t_d > \delta$ and $t_d < \delta$). Instead of investigating each scenario separately, both of these scenarios can be implemented within a single RL environment, with a random delay generated from a uniform distribution $U(0, 2)$.

With that being said, the proposed model has a predefined threshold for the time of delay. This means that it can only handle situations where the maximum delay time equals two inspection intervals. In the future, this restriction could also be removed by allowing the delay time to extend beyond two inspection intervals. While this could significantly impact the optimal policy and potentially worsen the system's condition due to the increased uncertainty, it also helps to create a more realistic environment.

References

- [1] Grall, A., Dieulle, L., Berenguer, C. and Roussignol, M. (2002). Continuous-time predictive-maintenance scheduling for a deteriorating system. *IEEE Transactions on Reliability*, **51**(2), 41150.
- [2] Brenguer, C., Grall, A., Dieulle, L. and Roussignol, M. (2003). Maintenance policy for a continuously monitored deteriorating system. *Probability in the Engineering and Informational Sciences*, **17**(2), 235250.
- [3] Meier-Hirmer, C., Riboulet, G., Sourget, F. and Roussignol, M. (2009). Maintenance optimization for a system with a gamma degradation process and intervention delay: Application to track maintenance. *Proceedings of the Institution of Mechanical Engineers, Part O: Journal of Risk and Reliability*, **223**(3), 189198.

- [4] Yang, H., Li, W. and Wang, B. (2021). Joint optimization of preventive maintenance and production scheduling for multi-state production systems based on reinforcement learning. *Reliability Engineering & System Safety*, **214**, 107713.
- [5] Wei, S., Bao, Y. and Li, H. (2020). Optimal policy for structure maintenance: A deep reinforcement learning framework. *Structural Safety*, **83**, 101906.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15. 2025



Reliability of Complex k -out-of- $n:d$ Systems

Razmkhah, M. ¹ Khatib Astaneh, B. ²

¹ Department of Statistics, Ferdowsi University of Mashhad, Mashhad, Iran

² Department of Mathematics and Statistics, University of Neyshabur, Neyshabur, Iran

Abstract

A system comprising n independent subsystems is considered, where each subsystem consists of d dependent components. The system is operational if and only if at least k out of the n subsystems are functioning. The Gumbel copula function is employed to model the dependency among the components within each subsystem. The system reliability is analytically evaluated, and graphical results are presented to illustrate the findings. The study demonstrates that ignoring dependency among components leads to significant bias in reliability estimation. Finally, key conclusions and implications are discussed.

Keywords: Coherent system, Dependency, Gumbel Copula, Reliability

1 Introduction

The reliability analysis of coherent systems is one of the most interesting topics in engineering and industries. Such systems have been described by Barlow and Proschan [1]. A k -out-of- n system is a kind of coherent systems which consists of n elements, and functions if and only if at least k of its elements function.

The independence assumption among the system elements has been considered in most previous research works. However, there are situations that all elements or a part of them are

¹razmkhah_m@um.ac.ir

²khatib_b@neyshabur.ac.ir

dependent to each other. Estimating reliability of such systems may be incorrect if dependencies are ignored. Navarro et al. [4] studied properties of coherent systems with dependent components. Bayramoglu [2] discussed the reliability and mean residual life of complex systems with two dependent components per element. Ota and Kimura [6] investigated coherent systems with dependent components due to common factors by using the concept of factor copula. Ghanbari et al. [3] investigated a stochastic comparisons of series and parallel systems with dependent log-logistic components. Reliability analysis for system with dependent components based on survival signature and copula theory was studied by Zheng and Zhang [10]. Zhang et al. [9] studied reliability and maintenance analysis of a system with dependent main and auxiliary components.

Here, we consider the systems with n independent subsystems each consisting of d dependent components, such that the system works properly if and only if at least k of its subsystems function. Moreover, the copula function is used to describe the dependence structure of the components. Such systems were called as complex k -out-of- $n:d$ systems by Saberzadeh et al [8]. They studied a Bayesian analysis for the special case of $d = 2$ when the components are subject to degradation. Similar research work has been done by Saberzadeh and Razmkhah [7].

The rest of this paper is organized as follows. In Section 2, the system used in the paper is described. Reliability of the proposed system is evaluated in Section 3. Some graphical results are presented in Section 4. Eventually, some conclusions are stated in Section 5.

2 Model description

Suppose that a system constitutes by n independent subsystems each containing d dependent components. The components within each subsystem may be arranged in a series, parallel or any other coherent structure. Assume that the system works if and only if at least k subsystems are functioning. Saberzadeh et al. [8] called such a system a complex k -out-of- $n:d$ series (parallel) system if the components of each subsystem are connected in a series (parallel) structure. They provided a Bayesian reliability analysis for such systems under degradation performance in the special case of $d = 2$. Here, we present some classical reliability analysis based on lifetime data. Assume that components of all subsystems have the same distribution. Let us denote the lifetime of the j th ($j = 1, 2, \dots, d$) component of each subsystem by T_j which obeys the marginal cumulative distribution function (cdf) $G_j(\cdot)$. Also, denote by X_i the lifetime of the i th ($i = 1, 2, \dots, n$) subsystem. To obtain the reliability function of X_i , we use a copula function to model dependency structure of the components. A copula function indeed connects individual marginal cdfs to obtain the joint cdf of $(T_1, \dots, T_d)$. According to Sklar's Theorem, there exists a unique copula $C(\cdot)$ such that for all real values $t_1, \dots, t_d$, the joint cdf of $(T_1, \dots, T_d)$ is

$$H(t_1, \dots, t_d) = C(G_1(t_1), \dots, G_d(t_d)). \quad (2.1)$$

Equation (2.1) demonstrates any multivariate distribution can be decomposed into a copula and its marginals. Thus, copula functions offer a much more flexible method to study multivariate distributions. For more details about the copula function and different kinds of it, the interested reader may refer to Nelsen [5]. In this paper, we use the Gumbel copula function as an upper tail dependence Archimedean copula. For $u_i \in [0, 1]$, $i = 1, \dots, d$, this copula is given by

$$C(u_1, \dots, u_d) = \exp \left\{ - \left(\sum_{i=1}^d (-\log u_i)^\theta \right)^{1/\theta} \right\}, \quad \theta \geq 1, \quad (2.2)$$

where θ is called the dependency parameter and shows the dependency between variables. The relationship between Kendall's correlation τ and the dependence parameter is $\tau = 1 - 1/\theta$. That is, this copula covers a wide range of dependency from independent to totally dependent variables.

3 Reliability evaluation

In this section, the reliability of a complex k -out-of- $n:d$ system is studied. It is obvious that the coherent structure among the components of each subsystem affects the reliability of the system. However, the series and parallel structures comprise the minimum and maximum reliabilities, respectively. So, we focus on complex k -out-of- $n:d$ series (parallel) system. Since the subsystems have a k -out-of- n structure, the system lifetime is the $(n - k + 1)$ th order statistic among $X_1, \dots, X_n$, denoted by $X_{n-k+1:n}$. Therefore, the system reliability is derived as

$$R(x) = \sum_{j=k}^n \binom{n}{j} (1 - F(x))^j (F(x))^{n-j}, \quad (3.1)$$

where $F(\cdot)$ stands for the cdf of X_i 's.

- When the subsystems have a parallel structure, their lifetime $X_1, \dots, X_n$ are some copies of

$$X = \max\{T_1, \dots, T_d\}$$

with the following cdf

$$\begin{aligned} F_p(x) &= P(T_1 \leq x, \dots, T_d \leq x) \\ &= C(G_1(x), \dots, G_d(x)), \end{aligned} \quad (3.2)$$

where the subscript p stands for the parallel subsystem. Using, the Gumbel copula (2.2), we get

$$F_p(x) = \exp \left\{ - \left(\sum_{i=1}^d (-\log G_i(x))^\theta \right)^{1/\theta} \right\}. \quad (3.3)$$

Substituting $F_p(x)$ instead of $F(x)$ in (3.1), the reliability of a complex k -out-of- $n:d$ parallel system is derived.

- Analogously, when the components within each subsystem are series, the subsystem lifetime are copies of

$$X = \min\{T_1, \dots, T_d\}$$

with the following cdf

$$\begin{aligned} F_s(x) &= 1 - P(T_1 > x, \dots, T_d > x) \\ &= 1 - C(\bar{G}_1(x), \dots, \bar{G}_d(x)) \\ &= 1 - \exp \left\{ - \left(\sum_{i=1}^d (-\log \bar{G}_i(x))^\theta \right)^{1/\theta} \right\}, \end{aligned} \quad (3.4)$$

where $\bar{G}_i(x) = 1 - G_i(x)$ and the subscript s stands for the series subsystem. Hence, the reliability of a complex k -out-of- $n:d$ series system is obtained by substituting $F_s(x)$ instead of $F(x)$ in (3.1).

Remark 3.1. For the special case of $\theta = 1$, the Gumbel copula associates with the independent random variables $T_1, \dots, T_d$. In this case, the equations (3.2) and (3.4) are simplified as

$$F_p(x) = \prod_{i=1}^d G_i(x) \quad \text{and} \quad F_s(x) = 1 - \prod_{i=1}^d \bar{G}_i(x),$$

respectively.

4 Graphical results

Here, we assume that the lifetime of the i th ($i = 1, \dots, d$) component of each subsystem follows Weibull distribution with the cdf

$$G(t) = 1 - \exp\{-(t/\alpha_i)^{\beta_i}\},$$

where α_i and β_i are the scale and shape parameters, respectively. In this section, we consider $(\alpha_1, \dots, \alpha_4) = (1, 1.2, 1.3, 1.4)$ and $(\beta_1, \dots, \beta_4) = (5, 6, 7, 8)$ to analyze systems with at most four components in each subsystem. The behavior of the reliability function is investigated from three perspectives. First of all, the effect of dependence parameter on the system reliability is shown in Figure 1 for complex 3-out-of-5:2 and 3-out-of-5:3 series and parallel systems. From this figure, it is observed that

1. The most (least) reliability corresponds to the systems with independent parallel (series) components. This is the case of $\theta = 1$.
2. When the components are dependent, the system reliability decreases (increases) for the subsystems with parallel (series) structure. Moreover, by increasing θ , the reliability decreases (increases) further when the structure of subsystems is parallel (series). Therefore, ignoring dependence among the components results in a bias in reliability.
3. Our computations show that increasing θ to values greater than 2 does not result in a significant difference in reliability.
4. As the number of dependent components d increases, the difference in reliability between systems with series and parallel components becomes greater.

Now, to study the effect of number of dependent components of each subsystem, the reliability functions of complex 3-out-of-5: d series and parallel systems are plotted in Figure 2 for $d = 2, 3, 4$ when $\theta = 1, 1.5$ is fixed. It is deduced that for a given θ , the reliability of complex 3-out-of-5: d parallel systems increase significantly when d increases. However, increasing the number of components results in a non-significant decrease in the reliability of complex 3-out-of-5: d series systems.

The final graphical result focuses on studying the effect of the number of subsystems on system reliability. Toward this end, the reliability functions of complex 3-out-of- n :3 series and parallel systems are plotted in Figure 3 for $n = 3, 5, 10$ when $\theta = 2$ is fixed. It is seen that

1. By increasing the number of subsystems n , the system reliability increases for both parallel and series structure of subsystems.
2. The reliability of a complex 3-out-of-10:3 series system is greater than that of a complex 3-out-of-3:3 parallel system for times before a specific time t_0 . Generally, reliability of a complex k -out-of- n_1 : d series system may be greater than that of a complex k -out-of- n_2 : d parallel system for some $n_1 > n_2$.

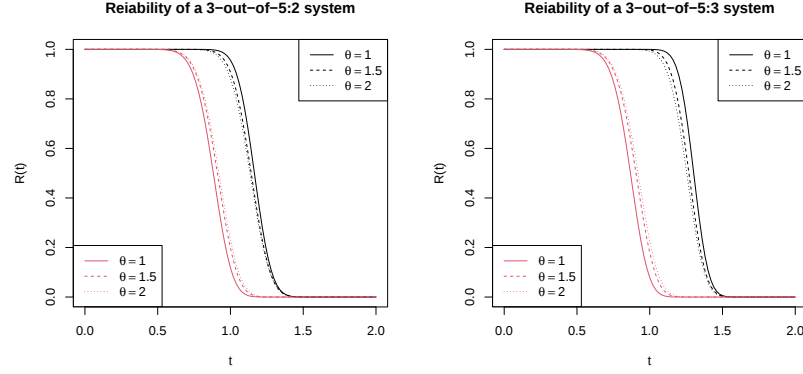


Figure 1: Reliability of complex 3-out-of-5: d parallel (plots in black) and series (plots in red) systems for $d = 2, 3$ and some choices of θ .

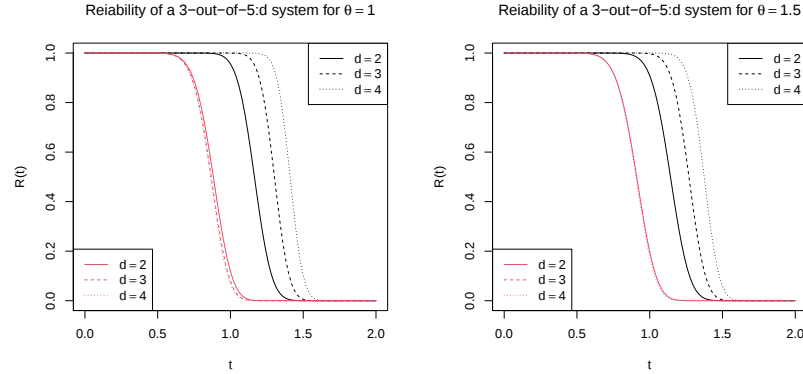


Figure 2: Reliability of complex 3-out-of-5: d parallel (plots in black) and series (plots in red) systems for $d = 2, 3, 4$ and $\theta = 1, 1.5$.

5 Conclusions

A k -out-of- $n:d$ system was considered and its reliability was studied. It was assumed that the components were put in a series or parallel arrangement. The dependency among the components was modeled by a Gumbel copula function with dependency parameter θ . By some graphical results, the behavior of system reliability with respect to varying n , d and θ was investigated. The following are the main results of the paper:

1. When the components are dependent, the system reliability decreases (increases) with respect to the case of independent components for the subsystems with parallel (series) structure.
2. As the dependency among components becomes stronger, the reliability decreases (increases) further when the structure of subsystems is parallel (series). Therefore, ignoring dependence among the components results in a bias in reliability.
3. When the structure of each subsystem is parallel, increasing the number of dependent components enhances the system reliability. On the other hand, for subsystems with series components, increasing the number of components slightly reduces the system reliability.
4. For a given k , increasing the number of subsystems increases the system reliability.

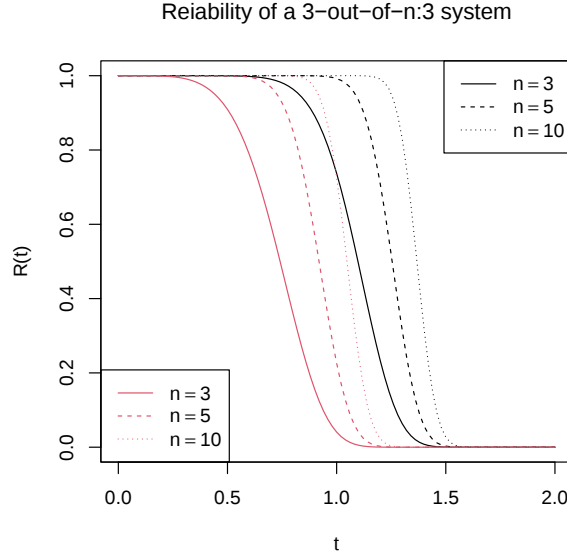


Figure 3: Reliability of complex 3-out-of- n :3 parallel (plots in black) and series (plots in red) systems for $n = 3, 5, 10$ when $\theta = 2$.

The results of the paper may extended to the following cases:

- Gumbel copula is a member of Archimedean class of copulas which was used in this paper. Other Archimedean copulas or even Elliptical or extreme-value copulas may be used to perform a comprehensive study for various types of tail dependence. For example, Gumbel and Joe copula functions have upper tail dependence, Clayton copula has lower tail dependence and Frank copula is symmetric with no tail dependence.
- In this paper, the special cases that subsystems have series or parallel components were studied. The results may be extended to other arrangement for the components or even other structures for subsystems. Moreover, subsystems with different sizes or different distributions may be investigated.

References

- [1] Barlow, R. E. and Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing*, Holt, Rinehart and Winston, New York.
- [2] Bayramoglu, I. (2013). Reliability and mean residual life of complex systems with two dependent components per element, *IEEE Transactions on Reliability*, **62**, 276-285.
- [3] Ghanbari, F., Barmalzan, G. and Hashemi, R. (2021). Stochastic comparisons of series and parallel systems with dependent log-logistic components, *Communications in Statistics–Theory and Methods*, **52**, 42594282.
- [4] Navarro, J., Ruiz, J. M. and Sandoval, C. J. (2007). Properties of Coherent Systems with Dependent Components, *Communications in Statistics–Theory and Methods*, **36**, 175-191.
- [5] Nelsen, R. B. (2007). *An introduction to copulas*, Springer, New York.

- [6] Ota, S. and Kimura, M. (2019). Reliability modeling of coherent systems with dependent components due to common factors, *IEEE 24th Pacific Rim International Symposium on Dependable Computing (PRDC)*, Kyoto, Japan, pp. 31-311.
- [7] Saberzadeh, Z. and Razmkhah, M. (2022). Reliability of degrading complex systems with two dependent components per element. *Reliability Engineering and System Safety*, **222**, 108398.
- [8] Saberzadeh, Z., Razmkhah, M. and Amini, M. (2023). Bayesian reliability analysis of complex k -out-of- $n:l$ systems under degradation performance. *Reliability Engineering and System Safety*, **231**, 109020.
- [9] Zhang, N., Zhang, J. and Shen, J. (2023). Reliability and maintenance analysis of a system with dependent main and auxiliary components, *Computers & Industrial Engineering*, **179**, 109188.
- [10] Zheng, Y. and Zhang, Y. (2023). Reliability analysis for system with dependent components based on survival signature and copula theory, *Reliability Engineering & System Safety*, **238**, 109402.



Eleventh Seminar on
Reliability Theory and its Applications
May 14-15, 2025



Analyzing Rayleigh-Poisson Model for Censored Data with Informative Random Removals

Sharafi, M.¹ Fallah Hasan, A.

Department of Statistics, Faculty of Science, Razi University, Kermanshah, Iran

Abstract

An important challenge when using progressive Type-II right censoring is determining a removal scheme. This can be fixed in advance or chosen randomly based on some discrete probability distributions. In this paper, we propose a new approach where the experiment itself influences the removal scheme, leading to the implementation of an informative censoring plan. For this purpose, we consider the Rayleigh Poisson distribution for lifetime data, along with two types of distributions for random removals: the truncated Poisson distribution and the truncated discrete Rayleigh distribution. Then, we obtain the maximum likelihood estimators and compare our new approaches with the patterns of removal obtained from the discrete binomial distributions via a Monte Carlo simulation study in terms of their biases, mean squared errors, and values of expected total time on the experiment.

Keywords: Binomial removals, Expected total time, Random removal, Weibull distribution.

1 Introduction

In clinical trials that involve time-to-event data, subjects are monitored until the event of interest occurs or until they are censored. Censoring for time-to-event outcomes can happen for

¹m.sharafi@razi.ac.ir

various reasons, including reaching the end of the follow-up period, losing contact with participants, or individuals discontinuing their participation in the study without further data being collected. Moreover, in life-testing experiments, various constraints such as time and cost may prevent researchers from observing the entire lifetime of every unit. There are different censoring schemes, with one of the most common being the progressive type-II censoring scheme. This scheme allows for the removal of units at points other than just at the end of the experiment. Units may also drop out randomly. This flexibility can help ensure the accuracy and validity of the experiment's results. For instance, the number of patients who drop out of a clinical trial at each stage is random and cannot be anticipated. In some life test studies, an experimenter may determine that it is inappropriate or too dangerous to continue testing certain units, even if those units have not yet failed. In such cases, the pattern of removal at each failure is random. [10] considered a situation where each experimental unit was dropped out from the life test independently of the others with the same removal probability p for the two-parameter Weibull distributed data. In this case, the number of test units removed at each failure time followed the binomial distribution. Several researchers have examined the inference of lifetime data with binomial removals, including [9], [7], [3]. Readers are also referred to [1] on pages 339-340.

The Binomial distribution presents challenges and may not accurately reflect the phenomena being studied. One key issue is the fixed removal probability p at each stage, which plays a significant role in determining expected test times and inferential results. Both small and large values of p can significantly impact statistical inference. The Binomial distribution presents challenges and may not accurately reflect the phenomena being studied. One key issue is the fixed removal probability p at each stage, which plays a significant role in determining expected test times and inferential results. Both small and large values of p can significantly impact statistical inference. When p is too small, there is a low chance of removing a surviving unit at each experimental stage, leading to most units being eliminated in the final step. This situation causes the progressive censoring scheme to resemble a Type-II censoring plan, thereby questioning the validity of implementing such a plan. Conversely, if p is large and approaches 1, remaining units tend to be removed early in the life test. As a result, the observed lifetimes are much closer to the tail of the failure time distribution. Therefore, the expected test time for progressive Type-II censoring with binomial removals is nearly equivalent to that of a complete sample. Additionally, assuming p to be random and following a specific probability distributionsuch as the beta distribution (refer to [8])can introduce further complications by adding more parameters to the model or requiring validation of the true probability distribution of p . As a result, analyzing binomial random removals is non-informative, as the distribution of life (survival) times provides no information about the distribution of random removals. Ghahramani et al. [4] and Sharafi et al. [9] proposed scenarios for dependent removals based on normalized spacings with random and fixed coefficients, consistent with progressively Type-II censored order statistics from the exponential and three-parameter Weibull distributions, respectively. They compared these new approaches with patterns of removal from discrete uniform and binomial distributions through Monte Carlo simulations and real data studies.

In this paper, we propose a new approach where the experiment itself influences the removal scheme, leading to the implementation of an informative censoring plan. For this purpose, we consider the Rayleigh Poisson distribution for lifetime data, along with two types of distributions for random removals: the truncated Poisson distribution and the discrete truncated Rayleigh distribution.

The structure of this article is as follows: Section 2 introduces the models of the Rayleigh Poisson distribution, the truncated Poisson distribution, and the truncated discrete Rayleigh distribution. We focus on three types of distributions for random removals: informative dis-

tributions, such as the truncated Poisson and the discrete truncated Rayleigh distributions, as well as a non-informative distribution, the binomial distribution. In Section 3, we conduct simulations to evaluate and compare the performance of estimators obtained through different methods.

2 Model

The probability density function (pdf) of Rayleigh-Poisson distribution is given by

$$f(x, \theta, \lambda) = \frac{2\theta x \lambda e^{-\lambda}}{1 - e^{-\lambda}} e^{-\theta x^2} e^{\lambda e^{-\theta x^2}}; \quad \theta > 0, \lambda > 0, x > 0, \quad (2.1)$$

The cumulative distribution function (CDF) is

$$F(x, \theta, \lambda) = \frac{e^{\lambda} - e^{\lambda e^{-\theta x^2}}}{e^{\lambda} - 1}; \quad \theta > 0, \lambda > 0, x > 0. \quad (2.2)$$

The Rayleigh-Poisson Model is one of the compounding of two most greeted probability distributions, i.e. Rayleigh and zero-truncated Poisson distribution. The Weibull Poisson distribution was pioneered by [6]. Therefore, we consider two informative distributions: the truncated Poisson distribution (DP) and the truncated discrete Rayleigh distribution (DR) for random removals.

The truncated Poisson truncated mass function (pdf) is given by

$$P(R_1 = r_1; \lambda) = \frac{e^{-\lambda} \lambda^{r_1} (n - m)!}{r_1! \Gamma(n - m + 1, \lambda)}, \quad (2.3)$$

where $\Gamma(s, y)$ is upper incomplete gamma distribution and $\Gamma(s, y) = \int_y^\infty t^{s-1} e^{-t} dt$. For $i = 2, 3, \dots, m - 1$, we get

$$P(R_i; \lambda) = P(R_i = r_i | R_{i-1} = r_{i-1}, \dots, R_1 = r_1) = \frac{e^{-\lambda} \lambda^{r_i} (n - m - \sum_{j=1}^{i-1} r_j)!}{r_i! \Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda)}. \quad (2.4)$$

and then,

$$P(\mathbf{R} = \mathbf{r}) = P(R_1 = r_1) \cdots P(R_{m-1} = r_{m-1} | R_{m-1} = r_{m-1}, \dots, R_1 = r_1). \quad (2.5)$$

Substituting from (2.3) and (2.4) into (2.5), we have

$$P(\mathbf{R} = \mathbf{r}; \lambda) = \frac{e^{-\lambda(m-1)} \lambda^{\sum_{i=1}^m r_i}}{\prod_{i=1}^{m-1} \Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda)} \prod_{i=1}^{m-1} \frac{(n - m - \sum_{j=1}^{i-1} r_j)!}{r_i!}. \quad (2.6)$$

The conditional likelihood function for fixed progressive Type-II censoring can be defined according to [2] as:

$$L_1(\theta, \lambda; \mathbf{x} | \mathbf{R} = \mathbf{r}) = c \prod_{i=1}^m f(x_i; \theta, \lambda) (1 - F(x_i; \theta, \lambda))^{r_i}, \quad -\infty < x_1 < \dots < x_m < \infty, \quad (2.7)$$

$n, m \in N$, $1 \leq i \leq m$ and $c = \prod_{i=1}^m \gamma_i$ where $\gamma_i = \sum_{j=1}^m (r_j + 1)$. Substituting (2.1) and (2.2) into (2.7), we have

$$L_1(\theta, \lambda; \mathbf{x} | \mathbf{R}) = c_1 \prod_{i=1}^m \frac{\lambda \theta x_i}{e^\lambda - 1} e^{-\theta x_i^2} e^{\lambda e^{-\theta x_i^2}} \left[\frac{e^{\lambda e^{-\theta x_i^2}} - 1}{e^\lambda - 1} \right]^{r_i}. \quad (2.8)$$

$$L(\theta, \lambda; \mathbf{x}, \mathbf{R}) = c_1 \prod_{i=1}^m \frac{\lambda \theta x_i}{e^\lambda - 1} e^{-\theta x_i^2} e^{\lambda e^{-\theta x_i^2}} \left[\frac{e^{\lambda e^{-\theta x_i^2}} - 1}{e^\lambda - 1} \right]^{r_i} \times \frac{e^{-\lambda(m-1)} \lambda^{\sum_{i=1}^m r_i}}{\prod_{i=1}^{m-1} \Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda)}. \quad (2.9)$$

The log-likelihood function may then be written as

$$\begin{aligned} \ell(\theta, \lambda; \mathbf{x}, \mathbf{R}) \equiv \ell = & \log A + m \log \lambda + m \log \theta + \sum_{i=1}^m \log x_i - \theta \sum_{i=1}^m x_i^2 + \lambda \sum_{i=1}^m e^{-\theta x_i^2} - \\ & n \log(e^\lambda - 1) + \sum_{i=1}^m r_i [\log(e^{\lambda e^{-\theta x_i^2}} - 1)] - (m-1)\lambda + \log \lambda \sum_{i=1}^{m-1} r_i \\ & - \sum_{i=1}^{m-1} \log \left(\Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda) \right). \end{aligned} \quad (2.10)$$

The normal equations are obtained by equating the first-order partial derivatives of the loglikelihood function w.r.t. the parameters to zero, i.e.,

$$\frac{\partial \ell}{\partial \theta} = \frac{m}{\theta} - \lambda \sum_{i=1}^m x_i^2 e^{-\theta x_i^2} [1 + r_i (1 - e^{-\lambda e^{-\theta x_i^2}})^{-1}] = 0. \quad (2.11)$$

$$\begin{aligned} \frac{\partial \ell}{\partial \lambda} = & \frac{n - r_m}{\lambda} + \sum_{i=1}^m e^{-\theta x_i^2} - n(1 - e^{-\lambda})^{-1} + \sum_{i=1}^m r_i e^{-\theta x_i^2} (1 - e^{-\lambda e^{-\theta x_i^2}})^{-1} - \\ & (m-1) + \sum_{i=1}^{m-1} \frac{\lambda^{n-m-\sum_{j=1}^{i-1} r_j} e^{-\lambda}}{\Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda)} = 0. \end{aligned} \quad (2.12)$$

$$\begin{aligned}
\frac{\partial^2 \ell}{\partial \theta^2} &= -\frac{m}{\theta^2} + \lambda \sum_{i=1}^m x_i^4 e^{-\theta x_i^2} [1 + r_i (1 - e^{-\lambda e^{-\theta x_i^2}})^{-1}] + \lambda^2 \sum_{i=1}^m r_i x_i^4 e^{-2\theta x_i^2} e^{-\lambda e^{-\theta x_i^2}} \times \\
&\quad (1 - e^{-\lambda e^{-\theta x_i^2}})^{-2}. \\
\frac{\partial^2 \ell}{\partial \lambda \partial \theta} &= -\sum_{i=1}^m x_i^2 e^{-\theta x_i^2} [1 + r_i (1 - e^{-\lambda e^{-\theta x_i^2}})^{-1}] + \lambda \sum_{i=1}^m r_i x_i^2 e^{-2\theta x_i^2} e^{-\lambda e^{-\theta x_i^2}} \times \\
&\quad (1 - e^{-\lambda e^{-\theta x_i^2}})^{-2}. \\
\frac{\partial^2 \ell}{\partial \lambda^2} &= \frac{-(n - r_m)}{\lambda^2} + n e^{-\lambda} (1 - e^{-\lambda})^{-2} - \sum_{i=1}^m r_i e^{-2\theta x_i^2} e^{-\lambda e^{-\theta x_i^2}} (1 - e^{-\lambda e^{-\theta x_i^2}})^{-2} + \\
&\quad \sum_{i=1}^{m-1} \frac{e^{-\lambda} \lambda^{n-m-\sum_{j=1}^{i-1} r_j - 1}}{\Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda)} \times \\
&\quad \left[(n - m - \sum_{j=1}^{i-1} r_j - \lambda) - \frac{e^{-\lambda} \lambda^{n-m-\sum_{j=1}^{i-1} r_j + 1}}{\Gamma(n - m - \sum_{j=1}^{i-1} r_j + 1, \lambda)} \right].
\end{aligned}$$

We also consider truncated discrete Rayleigh Distribution for random removals follow as:

$$P(R_i = r_i | R_{i-1} = r_{i-1}, \dots, R_1 = r_1) = \frac{e^{-\theta r_i^2} - e^{-\theta(r_i+1)^2}}{1 - \exp(-\theta(n - m - \sum_{j=1}^{i-1} r_j + 1)^2)}. \quad (2.13)$$

$$\begin{aligned}
L(\theta, \lambda; \mathbf{x}, \mathbf{R}) &= c_1 \prod_{i=1}^m \frac{\lambda \theta x_i}{e^\lambda - 1} e^{-\theta x_i^2} e^{\lambda e^{-\theta x_i^2}} \left[\frac{e^{\lambda e^{-\theta x_i^2}} - 1}{e^\lambda - 1} \right]^{r_i} \times \\
&\quad \prod_{i=1}^{m-1} \frac{e^{-\theta r_i^2} - e^{-\theta(r_i+1)^2}}{1 - \exp(-\theta(n - m - \sum_{j=1}^{i-1} r_j + 1)^2)}. \quad (2.14)
\end{aligned}$$

Suppose that an individual unit being removed from the experiment at the i th failure, $i = 1, 2, \dots, r - 1$ is independent of the others, but with same probability p . Then, the number R_i of units removed at the i th failure, $i = 1, 2, \dots, r - 1$ follows a binomial distribution with parameters $n - r - \sum_{\ell=1}^{i-1} r_\ell$ and p . Therefore,

$$P(R_1 = r_1) = \binom{n-r}{r_1} p^{r_1} (1-p)^{n-r-r_1}, \quad (2.15)$$

$$P(R_i = r_i | R_1 = r_1, \dots, R_{i-1} = r_{i-1}) = \binom{n-r-\sum_{j=1}^{i-1} r_j}{r_i} p^{r_i} (1-p)^{n-r-\sum_{j=1}^{i-1} r_j}, \quad (2.16)$$

where $0 \leq r_1 \leq n - r$ and $0 \leq r_i \leq n - r - \sum_{j=1}^{i-1} r_j$ for $i = 1, 2, \dots, r - 1$. Then

$$P(\mathbf{R} = \mathbf{r}) = \frac{(n-r)!}{\left(n-r-\sum_{j=1}^{r-1} r_j\right)! \prod_{j=1}^{r-1} (k_j)!} p^{\sum_{j=1}^{r-1} r_j} (1-p)^{(r-1)(n-r)-\sum_{j=1}^{r-1} (r-j)r_j}. \quad (2.17)$$

3 Simulation results

A Monte Carlo simulation study is performed to compare the performance of the proposed removal models, the discrete binomial random removals. The point estimation is evaluated by estimation value, bias, the mean square error (MSE) and the expected total time. For each set of simulated data, we have computed the estimators discussed in Section ???. The program codes of simulation results are written in R Version 3.5.2 and the maxLik function that is in the maxLik package is used for the constrained optimization (see [5]).

Based on the results obtained from the simulation, it was found that the absolute values of bias and RMSEs for all the estimators generally decrease as the proportion of effective samples, represented by r/n , increases. When n and r are fixed the DP and DR estimators provide the smallest bias and RMSE for the estimates in most cases. The aim of this paper is to compare different methods of generating random removals and analyze their impact on the termination point of an experiment. From Table 7 we have observed that the expected total time value (denoted by $E(X_{m:m:n})$) increases as the fixed sample size n and the effective sample size m increases, which is an expected result. When θ remains fixed, the value increases with an increase in λ . Furthermore, we have discovered that the values of $E(X_{m:m:n})$ under the Dp and DR approaches are significantly lower compared to the values based on the uniform and binomial models. The DR approach has the lowest expected test times compared to other approaches. This preference is maintained even for a small effective sample size m .

Table 1: Simulated results under different approaches of generating progressive censoring schemes with $\theta = 1$, and $\lambda = 0.5$

n	r	Removal plan	θ		λ		$Cov(\hat{\theta}, \hat{\lambda})$
			Bias	RMSE	Bias	RMSE	
10	4	DP	0.796670	1.08456	0.575843	0.648232	0.016190
		DR	-0.765323	0.767499	0.776061	0.778891	0.001685
		$Bin(p = 0.1)$	1.036848	1.310492	-0.456256	0.509548	0.083552
		$Bin(p = 0.5)$	0.890470	1.325872	-0.411014	0.508585	0.058377
		$Bin(p = 0.9)$	0.663518	1.271081	-0.309695	0.533567	-0.013407
10	8	DP	-0.096499	0.413476	1.074865	1.186895	-0.031768
		DR	0.414608	0.480656	-0.369837	0.386766	0.001440
		$Bin(p = 0.1)$	0.492979	0.739139	0.451601	0.489212	-0.000004
		$Bin(p = 0.5)$	0.323848	0.654739	-0.244050	0.581977	-0.071044
		$Bin(p = 0.9)$	0.250827	0.608951	-0.129341	0.659381	-0.102804
20	8	DP	0.296174	0.527188	0.463810	0.508576	0.0381324
		DR	-0.858559	0.859358	0.310477	0.339473	0.003943
		$Bin(p = 0.1)$	1.960943	2.062672	-0.498649	0.501384	-0.001849
		$Bin(p = 0.5)$	0.431538	0.714475	-0.359580	0.495358	-0.012145
		$Bin(p = 0.9)$	0.323980	0.649745	-0.239750	0.549935	-0.055366
20	12	DP	0.011615	0.336172	0.712770	0.745361	-0.000638
		DR	-0.444755	0.468910	0.161241	0.174318	0.002564
		$Bin(p = 0.1)$	0.464533	0.674128	-0.478708	0.486971	0.0126508
		$Bin(p = 0.5)$	0.264919	0.522251	-0.258315	0.532730	-0.052761
		$Bin(p = 0.9)$	0.194147	0.472740	-0.140793	0.624937	-0.088549
40	10	DP	0.060553	0.317891	0.543795	0.549955	0.017910
		DR	-0.959666	0.960091	0.499273	0.681056	-0.011599
		$Bin(p = 0.1)$	2.989949	3.013813	-0.4999365	0.4999422	-0.000028
		$Bin(p = 0.5)$	0.424362	0.644963	-0.383533	0.482636	-0.012328
		$Bin(p = 0.9)$	0.291363	0.573230	-0.243334	0.553691	-0.058365
40	20	DP	0.361036	0.499555	0.599302	0.629398	-0.017313
		DR	-0.807519	0.819948	0.509062	1.359281	0.153846
		$Bin(p = 0.1)$	0.612212	0.650481	-0.489968	0.492411	-0.001539
		$Bin(p = 0.5)$	0.174363	0.364727	-0.214446	0.562810	-0.065735
		$Bin(p = 0.9)$	0.130263	0.344292	-0.115335	0.617502	-0.085314

Table 2: Simulated results under different approaches of generating progressive censoring schemes with $\theta = 1$, and $\lambda = 1$

n	r	Removal plan	θ		λ		$Cov(\hat{\theta}, \hat{\lambda})$
			Bias	RMSE	Bias	RMSE	
10	4	DP	0.856504	1.148503	0.544941	0.673452	-0.016086
		DR	-0.759813	0.761948	0.356258	0.362575	0.001628
		$Bin(p = 0.1)$	1.280561	1.613403	-0.914484	0.967313	0.147068
		$Bin(p = 0.5)$	1.157727	1.618840	-0.873885	0.950463	0.108419
		$Bin(p = 0.9)$	0.947252	1.563040	-0.771216	0.912196	-0.012212
10	8	DP	-0.002809	0.430705	0.750119	0.922096	-0.029689
		DR	0.491223	0.543655	-0.840793	0.848419	0.001532
		$Bin(p = 0.1)$	0.783355	1.047515	-0.946689	0.969525	-0.004313
		$Bin(p = 0.5)$	0.5453434	0.906539	-0.689327	0.904756	-0.109108
		$Bin(p = 0.9)$	0.443818	0.834744	-0.549235	0.893292	-0.150765
20	8	DP	1.752811	1.903803	0.166317	0.404569	-0.060406
		DR	-0.857074	0.857931	-0.160627	0.216608	0.004359
		$Bin(p = 0.1)$	2.160376	2.216866	-0.997281	01.000281	0.000788
		$Bin(p = 0.5)$	0.681954	1.003857	-0.826088	0.913242	-0.031872
		$Bin(p = 0.9)$	0.519204	0.879217	-0.664569	0.875861	-0.096385
20	12	DP	0.171456	0.428815	0.444245	0.564980	0.004452
		DR	-0.414354	0.447677	-0.246904	0.252926	0.001125
		$Bin(p = 0.1)$	0.892921	0.988498	-0.974872	0.979085	0.011244
		$Bin(p = 0.5)$	0.454887	0.698979	-0.663312	0.868856	-0.094660
		$Bin(p = 0.9)$	0.380723	0.657622	-0.545551	0.865234	-0.137229
40	10	DP	0.217714	0.399259	0.256536	0.278283	0.0237559
		DR	-0.959171	0.959607	-0.014104	0.463999	-0.011874
		$Bin(p = 0.1)$	3.064853	3.072726	-0.999711	0.999738	0.000256
		$Bin(p = 0.5)$	0.656612	0.870465	-0.846094	0.914070	-0.024853
		$Bin(p = 0.9)$	0.488525	0.774250	-0.666767	0.865918	-0.087835
40	20	DP	0.537881	0.569743	0.551511	0.597849	0.014899
		DR	-0.781346	0.804431	0.150746	1.483783	0.244208
		$Bin(p = 0.1)$	0.747140	0.810147	-0.983222	0.984763	0.002487
		$Bin(p = 0.5)$	0.334267	0.521935	-0.606101	0.851210	-0.104655
		$Bin(p = 0.9)$	0.274564	0.4911501	-0.4673109	0.860157	-0.143801

Table 3: Simulated results under different approaches of generating progressive censoring schemes with $\theta = 1$, and $\lambda = 2$

n	r	Removal plan	θ		λ		$Cov(\hat{\theta}, \hat{\lambda})$
			Bias	RMSE	Bias	RMSE	
10	4	DP	0.949976	1.380469	0.426680	0.531513	-0.077311
		DR	-0.755098	0.757278	-0.475720	0.479894	0.001131
		$Bin(p = 0.1)$	2.075264	2.281572	-1.799881	1.857189	0.269367
		$Bin(p = 0.5)$	1.993314	2.424751	-1.821757	1.885634	0.164303
		$Bin(p = 0.9)$	1.619692	2.276713	-1.685438	1.788040	-0.019858
10	8	DP	0.277865	0.649238	-0.008074	0.636902	-0.035839
		DR	0.628548	0.668154	-1.743871	1.747737	0.000927
		$Bin(p = 0.1)$	1.635424	1.955222	-1.933171	1.949729	-0.027146
		$Bin(p = 0.5)$	1.114411	1.530757	-1.604613	1.738757	-0.226713
		$Bin(p = 0.9)$	0.959809	1.42016	-1.447492	1.625545	-0.307255
20	8	DP	2.193174	2.193733	2.808901	2.809196	-0.000139
		DR	-0.854980	0.855907	-1.114936	1.126141	0.005006
		$Bin(p = 0.1)$	2.455871	2.542069	-1.985654	1.990886	0.02217916
		$Bin(p = 0.5)$	1.365638	1.759053	-1.765731	1.829582	-0.085774
		$Bin(p = 0.9)$	1.114660	1.543881	-1.584415	1.709924	-0.199047
20	12	DP	0.488733	0.652888	-0.071874	0.497804	-0.063375
		DR	-0.373096	0.422384	-1.057920	1.058863	-0.000873
		$Bin(p = 0.1)$	2.173875	2.314614	-1.986946	1.989223	-0.008630
		$Bin(p = 0.5)$	0.918024	1.20099	-1.512092	1.643603	-0.182682
		$Bin(p = 0.9)$	0.825047	1.140509	-1.365074	1.57474	-0.267193
40	10	DP	0.493414	0.660410	-0.198542	0.247013	0.046098
		DR	-0.959899	0.9603153	1.005022	1.107683	-0.0111817
		$Bin(p = 0.1)$	3.147276	3.157539	-1.995608	1.996764	0.006408
		$Bin(p = 0.5)$	1.270180	1.535471	-1.772674	1.831776	-0.027184
		$Bin(p = 0.9)$	1.010975	1.341174	-1.543283	1.672645	-0.147383
40	20	DP	0.640712	0.673167	-0.233823	0.342422	-0.007910
		DR	-0.706468	0.773598	-0.527204	2.071396	0.556754
		$Bin(p = 0.1)$	1.395754	1.469774	-1.952170	1.957556	0.024468
		$Bin(p = 0.5)$	0.781479	0.993863	-1.417184	1.576036	-0.213008
		$Bin(p = 0.9)$	0.674218	0.902218	-1.25997	1.490964	-0.269267

Table 4: Simulated results under different approaches of generating progressive censoring schemes with $\theta = 2$, and $\lambda = 0.5$

n	r	Removal plan	θ		λ		$Cov(\hat{\theta}, \hat{\lambda})$
			Bias	RMSE	Bias	RMSE	
10	4	DP	0.675863	1.679309	0.719675	0.809675	-0.021040
		DR	-1.799003	1.799333	-0.150915	0.510314	0.015175
		$Bin(p = 0.1)$	4.816717	5.315164	-0.479416	0.532467	0.061400
		$Bin(p = 0.5)$	1.579545	2.148374	-0.353756	0.602593	0.183036
		$Bin(p = 0.9)$	1.100703	1.989597	-0.235106	0.638248	0.111937
10	8	DP	-0.228325	0.783441	1.071310	1.193528	0.008825
		DR	-0.320845	0.353834	-0.217049	0.242806	0.009126
		$Bin(p = 0.1)$	1.111747	1.536954	-0.446983	0.501986	0.006699
		$Bin(p = 0.5)$	0.653395	1.299587	-0.239579	0.587776	-0.125112
		$Bin(p = 0.9)$	0.514204	1.249572	-0.130703	0.642590	-0.181904
20	8	DP	2.984995	3.091608	0.355483	0.397226	0.004325
		DR	-1.890583	1.890650	-0.268063	0.359452	0.003654
		$Bin(p = 0.1)$	1.576723	1.769495	-0.4767762	0.5135182	0.03486098
		$Bin(p = 0.5)$	0.943429	1.462723	-0.357120	0.502406	-0.019248
		$Bin(p = 0.9)$	0.649551	1.291928	-0.229054	0.558881	-0.083841
20	12	DP	0.100474	0.540929	0.766207	0.826873	0.003020
		DR	-1.628941	1.631915	-0.263546	0.300236	0.011904
		$Bin(p = 0.1)$	1.992414	2.258767	-0.493598	0.496886	-0.004621
		$Bin(p = 0.5)$	0.494643	0.982561	-0.239213	0.520804	-0.067340
		$Bin(p = 0.9)$	0.397616	0.929589	-0.149203	0.593331	-0.140117
40	10	DP	0.0245974	0.550972	0.652785	0.656283	0.007489
		DR	-1.922945	1.923261	-0.089089	0.481273	-0.016496
		$Bin(p = 0.1)$	2.191582	2.216631	-0.491678	0.500235	0.012003
		$Bin(p = 0.5)$	0.905245	1.302205	-0.377524	0.511534	0.007244
		$Bin(p = 0.9)$	0.583422	1.123736	-0.242062	0.535099	-0.062602
40	20	DP	0.0635549	0.264257	0.800760	0.808675	-0.006519
		DR	-1.893807	1.894724	-0.322031	0.774752	0.774752
		$Bin(p = 0.1)$	1.103615	1.164188	-0.482623	0.488305	0.007109
		$Bin(p = 0.5)$	0.357851	0.696929	-0.219687	0.520507	-0.096259
		$Bin(p = 0.9)$	0.271554	0.694732	0.584041	0.814533	-0.156591

Table 5: Simulated results under different approaches of generating progressive censoring schemes with $\theta = 2$, and $\lambda = 1$

n	r	Removal plan	θ		λ		$Cov(\hat{\theta}, \hat{\lambda})$
			Bias	RMSE	Bias	RMSE	
10	4	DP	0.887888	1.932605	0.619140	0.812805	-0.115431
		DR	-0.765323	0.767499	0.776061	0.778891	0.001685
		$Bin(p = 0.1)$	5.399835	5.749903	-0.971697	1.010445	0.071944
		$Bin(p = 0.5)$	2.523632	3.137032	-0.845251	1.000454	0.167858
		$Bin(p = 0.9)$	1.588388	2.464263	-0.677242	0.957439	0.128468
10	8	DP	0.035469	0.898345	0.746134	0.977980	-0.015424
		DR	-0.185186	0.241864	-0.658783	0.665741	0.006928
		$Bin(p = 0.1)$	1.576623	2.013755	-0.9199345	0.970124	0.037686
		$Bin(p = 0.5)$	1.067422	1.723489	-0.672070	0.903500	-0.199353
		$Bin(p = 0.9)$	0.871449	1.575150	-0.530948	0.886075	-0.290203
20	8	DP	1.260455	1.274752	3.362319	3.364821	-0.015804
		DR	-1.890382	1.890450	-0.765577	0.803360	0.003755
		$Bin(p = 0.1)$	9.520243	10.067830	-1.998437	1.998720	-0.001882
		$Bin(p = 0.5)$	1.365309	1.853452	-0.786149	0.915009	-0.000991
		$Bin(p = 0.9)$	1.039254	1.664184	-0.651375	0.872822	-0.140124
20	12	DP	0.345842	0.711765	0.497694	0.6371248	-0.030404
		DR	-1.624862	1.628060	-0.756092	0.771487	0.013302
		$Bin(p = 0.1)$	2.490408	2.672304	-0.989658	0.993891	0.001767
		$Bin(p = 0.5)$	0.889570	1.383413	-0.655901	0.845870	-0.137371
		$Bin(p = 0.9)$	0.726762	1.266973	-0.536738	0.839662	-0.231424
40	10	DP	0.667162	0.795635	0.661251	0.679652	-0.038521
		DR	-1.922674	1.922988	-0.595126	0.756807	-0.016245
		$Bin(p = 0.1)$	2.333786	2.402780	-0.974864	0.990934	0.036694
		$Bin(p = 0.5)$	1.459103	1.886051	-0.842569	0.923266	-0.031124
		$Bin(p = 0.9)$	0.999954	1.542416	-0.666790	0.865725	-0.119296
40	20	DP	0.037282	0.411119	-0.8698366	0.8771684	-0.011605
		DR	-1.891440	1.892854	-0.797594	1.152194	0.054555
		$Bin(p = 0.1)$	2.011013	2.074511	-0.988127	0.990225	0.002315
		$Bin(p = 0.5)$	0.670510	0.988729	-0.612844	0.803629	-0.135901
		$Bin(p = 0.9)$	0.546187	0.936867	-0.472835	0.808630	-0.239367

Table 6: Simulated results under different approaches of generating progressive censoring schemes with $\theta = 2$, and $\lambda = 2$

n	r	Removal plan	θ		λ		$Cov(\hat{\theta}, \hat{\lambda})$
			Bias	RMSE	Bias	RMSE	
10	4	DP	1.22889	2.474895	0.417138	0.684719	-0.294316
		DR	-1.796328	1.796707	-1.593342	1.70100	0.019353
		$Bin(p = 0.1)$	6.251004	6.463149	-1.911211	1.970166	0.1488912
		$Bin(p = 0.5)$	4.684171	5.497140	-1.851269	1.921179	0.088352
		$Bin(p = 0.9)$	2.792927	3.696168	-1.569471	1.77193	0.117381
10	8	DP	0.595381	1.142995	0.031031	0.736262	0.019850
		DR	0.047980	0.169355	-1.573891	1.575164	0.003606
		$Bin(p = 0.1)$	2.886386	3.304118	-1.841664	1.893251	0.043101
		$Bin(p = 0.5)$	2.114786	2.796097	-1.538605	1.699319	-0.362834
		$Bin(p = 0.9)$	1.850795	2.650816	-1.391944	1.608242	-0.560924
20	8	DP	1.271592	1.329294	2.803088	2.808112	-0.061079
		DR	-1.889980	1.890050	-1.761345	1.778897	0.003877
		$Bin(p = 0.1)$	9.449705	9.981375	-1.998158	1.998495	-0.001276
		$Bin(p = 0.5)$	2.498893	3.044436	-1.67573	1.768481	-0.114531
		$Bin(p = 0.9)$	2.134053	2.822758	-1.555333	1.685981	-0.267690
20	12	DP	1.047993	1.292370	-0.111545	0.505654	-0.035245
		DR	-1.618089	1.621729	-1.737989	1.746052	0.015502
		$Bin(p = 0.1)$	3.227809	3.429729	-1.960547	1.969922	0.027320
		$Bin(p = 0.5)$	1.90027	2.417494	-1.527282	1.660721	-0.302862
		$Bin(p = 0.9)$	1.609975	2.189543	-1.364950	1.567577	-0.471689
40	10	DP	2.080985	2.080985	5.531152	5.531152	0
		DR	-1.922885	1.923200	-1.595031	1.661186	-0.016152
		$Bin(p = 0.1)$	3.225478	3.428345	-1.866812	1.927081	0.146805
		$Bin(p = 0.5)$	2.424775	2.806372	-1.682015	1.785097	-0.046206
		$Bin(p = 0.9)$	1.931227	2.475967	-1.508140	1.657105	-0.272037
40	20	DP	0.579208	0.720115	0.026422	0.270760	0.063556
		DR	-1.891946	1.893386	-1.799042	1.998606	0.058221
		$Bin(p = 0.1)$	3.578424	3.664448	-1.979444	1.981852	-0.009417
		$Bin(p = 0.5)$	1.531000	1.886428	-1.428797	1.555783	-0.282170
		$Bin(p = 0.9)$	1.304335	1.718655	-1.243484	1.455667	-0.451536

Table 7: Expected experiment time $\mathbf{E}(\mathbf{X}_{\mathbf{m:m:n}})$ under $PT - II$ CDRs

n	m	Removal plan	(θ, λ)					
			(1, 0.5)	(1, 1)	(1, 2)	(2, 0.5)	(2, 1)	(2, 2)
10	4	DP	0.773897	0.793424	0.815511	0.544894	0.557782	0.574961
		DR	0.633955	0.576761	0.479713	0.436182	0.633955	0.328340
		$Bin(p = 0.1)$	0.649096	0.590629	0.492089	0.458983	0.418376	0.349663
		$Bin(p = 0.5)$	0.995712	0.922755	0.791237	0.701100	0.652651	0.551753
		$Bin(p = 0.9)$	1.235588	1.161282	1.010470	0.872985	0.819633	0.716187
10	8	DP	1.524917	1.450465	1.284442	1.073981	1.024133	0.905458
		DR	1.288897	1.203590	1.052765	0.817908	0.761709	0.65362
		$Bin(p = 0.1)$	1.684310	1.086033	0.932754	0.825534	0.764151	0.656759
		$Bin(p = 0.5)$	1.432567	1.352806	1.190667	1.081586	0.959331	0.839135
		$Bin(p = 0.9)$	1.523949	1.446023	1.283438	.073169	1.020672	0.903571
20	8	DP	0.857386	0.959290	1.074334	0.605148	0.676639	0.764831
		DR	0.656978	0.596555	0.497055	0.448411	0.406779	0.338684
		$Bin(p = 0.1)$	0.745204	0.680822	0.567512	0.528238	0.401195	0.401557
		$Bin(p = 0.5)$	1.415599	1.335254	1.176457	1.000046	0.938964	0.824871
		$Bin(p = 0.9)$	1.498309	1.433021	1.254365	1.067404	1.007691	0.884732
20	12	DP	1.588240	1.527418	1.383272	1.122277	1.081502	0.979366
		DR	0.978579	0.902367	0.763766	0.617739	0.566083	0.475398
		$Bin(p = 0.1)$	1.071930	0.989992	0.849901	0.760409	0.702416	0.598310
		$Bin(p = 0.5)$	1.622987	1.546961	1.382916	1.146137	1.093763	0.970676
		$Bin(p = 0.9)$	1.650781	1.564857	1.397002	1.159969	1.113159	0.993210
40	10	DP	0.528883	0.499754	0.481889	0.373531	0.352848	0.340550
		DR	0.482843	0.435909	0.362097	0.336273	0.302451	0.251954
		$Bin(p = 0.1)$	0.606274	0.550474	0.458410	0.427647	0.390016	0.323233
		$Bin(p = 0.5)$	1.479524	1.390014	1.238049	1.040851	0.983938	0.8718177
		$Bin(p = 0.9)$	1.580225	1.495737	1.326816	1.108704	1.05438	0.939858
40	20	DP	1.611711	1.607371	1.497657	1.14485	1.143483	1.062358
		DR	0.819081	0.748535	0.624001	0.540879	0.490262	0.411089
		$Bin(p = 0.1)$	0.848973	1.110764	0.962994	1.143288	0.791200	0.678891
		$Bin(p = 0.5)$	1.780028	1.707421	1.544327	1.256243	1.206383	1.088869
		$Bin(p = 0.9)$	1.789945	1.720608	1.548423	1.266512	1.214980	1.098613

References

- [1] Balakrishnan N., and Cramer E. (2014). *The Art of Progressive Censoring*. Springer New York.
- [2] Cohen, A. C. (1963). Progressively Censored Samples in the Life Testing, *Technometrics*, **5**, 327-339.
- [3] Elshahhat, A., and Nassar, M. (2023). Analysis of adaptive Type-II progressively hybrid censoring with binomial removals. *Journal of Statistical Computation and Simulation*, **93** (7), 1077-1103.
- [4] Ghahramani M., Sharafi M., and Hashemi R. (2020). Analysis of the progressively Type-II right censored data with dependent random removals, *Journal of Statistical Computation and Simulation*, **90**(6), , 1001-1021.
- [5] Henningsen A and Toomet O. (2011). maxLik: A package for maximum likelihood estimation in R, *Computational Statistics*, **26**(3), 443-458.
- [6] Lu, W. and Shi, D. (2012). A new compounding life distribution: TheWeibull-Poisson distribution, *J. Appl.Stat.* , **39**, 21-38.
- [7] Pathak A, Kumar M, Singh SK, Singh U. (2020). Bayesian inference: Weibull Poisson model for censored data using the expectation-maximization algorithm and its application to bladder cancer data. *J Appl Stat.*, **49**(4), 926-948.
- [8] Singh S. K., Singh U. and Sharma, V. K. (2013). Expected total test time and Bayesian estimation for generalized Lindley distribution under progressively Type-II censored sample where removals follow the beta-binomial probability law, *Applied Mathematics and Computation*, **222**, 402-419.
- [9] Sharafi M. (2022), Inference of the Two-Parameter Lindley Distribution based on Progressive Type II Censored Data with Random Removals, *Communications in Statistics-Simulation and Computation*, **51**(4), 1967-1981.
- [10] Tse S. K., Yang C. and Yuen H. K. (2000). Statistical analysis of Weibull distributed lifetime data under Type II progressive censoring with binomial removals, *Journal of Applied Statistics*, **27**, 1033-1043.