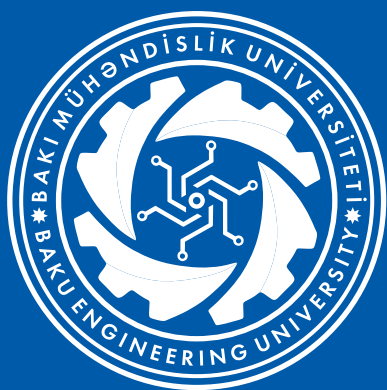




ISSN 2521-635X
E-3107-0531

Volume 10
Number 1

2026

Journal of Baku Engineering University

**MATHEMATICS AND
COMPUTER SCIENCE**

Journal is published twice a year
Number-1. June, Number-2. December

An International Journal

<http://journal.beu.edu.az>

EDITOR-IN-CHIEF

Agasi Melikov

Baku Engineering University, Azerbaijan

Achyutha Krishnamoorthy

CMS College Kottayam, India

EDITORS

Rakib Efendiyev

Baku Engineering University, Azerbaijan

Sedat Akleylek

Tartu University, Estonia

EDITORIAL ADVISORY BOARD

Abdeljalil Nachaoui (Nantes University, France)

Agamirza Bashirov (Eastern Mediterranean University, Turkey)

Agil Khanmamedov (Baku State University, Azerbaijan)

Alexander Dudin (Belarus State University, Belarus)

Ali Abbasov (Institute of Mathematics, Ministry of Science and Education Republic of Azerbaijan)

Asaf Hajiyeu (Baku Engineering University, Azerbaijan)

Ayyapan Govindan (Puducherry Technological University, India)

Bariş Erbaş (Anadolu University, Turkey)

Chakib Abdelkrim (Beni Mellal University, Morocco)

Che Soong Kim (Sangji University, Korea)

Efendi Nasibov (Dokuz Eylül Üniversitesi, İzmir, Turkey)

Garib Murshudov (York Academy, UK, London)

Golovko Vladimir Adamovich (Brest State University, Belarus)

Hamed Sari-Sarraf (Texas Technik University, USA)

Hari Srivastava (University of Victoria, Canada)

Janos Sztrik (Debrecen University, Hungary)

Jauberteau Francois (Nantes University, France)

Jinting Wang (Tsinghua University, China)

Jose K.P. (St.Peter's College Kolenchery, India)

Kamil Aida-zade (Institute of Mathematics, Ministry of Science and Education Republic of Azerbaijan)

Lior Tabansky (Cyber Research Center, Tel Aviv University, Israel)

Ludmila Prikazchikova (Keele University, England)

Manikandan Rangaswamy (Central University of Kerala, India)

Mourad Nachaoui (Nantes University, France)

Şerife Özkar (Balıkesir Üniversitesi, Turkey)

Sosnin Petr Ivanovich (Ulyanovsk State Technical University, Russia)

Srinivas Chakravarthy (Kettering University, Michigan, USA)

Svetlana Moiseeva (Tomsk State University, Russia)

Telman Aliyev (Institute of Mathematics, Ministry of Science and Education Republic of Azerbaijan)

Telman Melikov (Institute of Mathematics, Ministry of Science and Education Republic of Azerbaijan)

Vağif Gasimov (Baku Engineering University, Azerbaijan)

Vedat Coşkun (Işık University, Türkiye)

Vladimir B.Vasilyev (Lipetsk State Technical University, Russia)

EXECUTIVE EDITOR

Shafag Alizada

Baku Engineering University, Azerbaijan

ASSISTANT EDITOR

Ramil Mirzayev

Baku Engineering University, Azerbaijan

DESIGN

Ilham Aliyev

CONTACT ADDRESS

Journal of Baku Engineering University

AZ0102, Khirdalan city, Hasan Aliyev str. 120, Absheron, Baku, Azerbaijan

Tel: 00 994 12 - 349 99 95 Fax: 00 994 12 349-99-90/91

e-mail: journal@beu.edu.az

web: <http://journal.qu.edu.az>

facebook: [Journal Of Baku Engineering University](https://www.facebook.com/Journal-Of-Baku-Engineering-University)

Copyright © Baku Engineering University

ISSN 2521-635X | E-3107-0531

ISSN 2521-635X | E-3107-0531



Journal of Baku Engineering University

**MATHEMATICS AND
COMPUTER SCIENCE**

Baku - AZERBAIJAN

Journal of Baku Engineering University

MATHEMATICS AND COMPUTER SCIENCE

2026. Volume 10, Number 1

CONTENTS

MAP/PH/2 DOUBLE-ORBIT RETRIAL INVENTORY SYSTEM WITH WORKING VACATION AND RE-SERVICE

G. Ayyappan, M. Thilakavathy3

COEFFICIENT BOUNDS, INVESTIGATING INCLUSION PROPERTIES FOR A GENERAL SUBCLASS OF ANALYTIC FUNCTIONS DEFINED BY NOOR INTEGRAL OPERATOR

Sercan Kazimoğlu, Erhan Deniz32

INTEGRATION OF NON-STATIONARITY IN DISASTER RISK FORECASTING: A MULTIDISCIPLINARY STOCHASTIC ANALYSIS OF SEISMIC DYNAMICS IN AZERBAIJAN BASED ON HIDDEN MARKOV MODELS

Mir Ramin Yunusov38

AUTOENCODERS FOR REPRESENTATION LEARNING: FROM DIMENSIONALITY REDUCTION TO VOLUMETRIC GENERATION

Hamid Gadirov50

COMPARATIVE ANALYSIS OF MACHINE LEARNING MODELS FOR DETECTING ANOMALIES IN CYBERSECURITY AUDITS

T.A.Turarbekkyzy, Nadiya Kadyrzhan60

CONCEPT OF MANAGING THE ENVIRONMENTAL SAFETY OF DRINKING WATER IN RESERVOIRS

S.G.Gidayatzade70

ASSESSMENT OF THE CONFORMITY OF THE CONTENT OF HIGHER EDUCATION TO THE COMPETENCIES OF FUTURE SPECIALISTS

Vagif Gasumov, Khagani Abdullayev76

ARTIFICIAL INTELLIGENCE ADOPTION ACROSS COUNTRIES: EVIDENCE FROM MACROECONOMIC CLUSTERING TECHNIQUES

Sarkhan Salahzade86

AI-POWERED CONTENT OPTIMIZATION IN DIGITAL MARKETING: A TEXT MINING AND SENTIMENT ANALYSIS FRAMEWORK

Sahira Mammadova99

EVALUATION OF THE IMPACT OF ARTIFICIAL INTELLIGENCE ON WORK CULTURE <i>Royya Valehzada</i>	115
INTEGRATION OF HYBRID MACHINE LEARNING ARCHITECTURES AND EXPLAINABLE AI FOR CHRONIC DISEASE DIAGNOSIS <i>Shahin Babayev</i>	127
STACKED ENSEMBLE LEARNING FOR ECONOMIC GROWTH FORECASTING: A CROSS-COUNTRY PANEL DATA ANALYSIS <i>Parviz Jabrayilli</i>	132
RISK-BOUNDED REAL-TIME COLLISION AVOIDANCE UNDER PERCEPTION UNCERTAINTY <i>Elmir Adilzade</i>	146
CONTEXT-AWARE NATURAL LANGUAGE PROCESSING FRAMEWORK FOR ENHANCED HUMAN-COMPUTER INTERACTION IN VIRTUAL ASSISTANTS <i>Fatma Ismayilova</i>	161
INTEGRATION OF FACE RECOGNITION ALGORITHMS BASED ON TRADITIONAL COMPUTER VISION AND ARTIFICIAL INTELLIGENCE APPROACHES <i>Ravan Mammadov</i>	171
OPTIMIZATION OF RELATIONAL DATABASE ARCHITECTURES FOR INTELLIGENT MANAGEMENT SYSTEMS IN THE TOURISM SECTOR <i>Elnur Gasimov</i>	178

UDC: 519.872
 DOI: <https://doi.org/10.30546/09090.2026.001.20.2003>

MAP/PH/2 DOUBLE-ORBIT RETRIAL INVENTORY SYSTEM WITH WORKING VACATION AND RE-SERVICE

¹G. AYYAPPAN, ²M. THILAKAVATHY

^{1,2} Department of Mathematics, Puducherry Technological University,
 Puducherry, INDIA
 ayyappan @ptuniv.edu.in
 thilakavathy5184@ptuniv.edu.in

ARTICLE INFO	ABSTRACT
Article history: Received:2025-11-15 Received in revised form:2025-12-01 Accepted:2025-12-30 Available online: 2026-06-23 Keywords: Working vacation; re-service; emergency replenishment; double orbits; retrial customers; (s,S)policy 2010 Mathematics Subject Classifications: 60K25, 68M20, 90B05.	The current study focuses on the double orbit retrial inventory queueing reliable and unreliable servers. The customer arrives using a Markovian arrival process and the service time is allocated based on the phase-type distribution. If a customer cannot access the server at first, they enter an infinite orbit. After the service time, the customer will have the option to join the finite orbit for another service system because they were satisfied. If a customer finds that the finite orbit is full, they are considered lost. If there are no customers on the system, the server proceeds with the close-down procedure. The server is idle unless a customer arrives on vacation after closing down. Server-1 is idle unless a customer arrives on vacation and is served slowly. Re-service is the only thing server 2 offers. If server-1 fails during its vacation, the current customer is frozen until repair, after which service is resumed. Inventory is replenished exponentially under an (s, S) policy. Using the matrix approach, steady-state analysis is performed, and various performance measures are evaluated numerically and graphically.

2521-635X/© 2026 The Author(s). Published by Baku Engineering University.
 This is an open access article under the CC BY 4.0 license (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Retrial queueing systems, also known as queues with repeated attempts, are appropriate mathematical representations of modern contact centers, wireless networks, etc., because they account for the important factor of customer behavior. The customer will make random attempts to obtain access if the required resource is unavailable when they arrive. Retry queues have been the subject of plenty of scholarly work, and the relevant literature is included in the survey articles by [1]. In their study, [2] examined a conventional retrial queueing inventory system with a two-component demand rate. When two orbit queues with infinite capacity are assumed, the joint stationary orbit queue length distribution can be used to calculate the probability using the theory of Riemann (-Hilbert) problems introduced by [3]. In [4], retrial queues with finite capacity and geometric loss was examined. According to [5], the literature on the algorithmic and computational aspects of retrial queues was reviewed.

The double-orbit finite retrial queues with priority clients and service disruptions have been studied by [6]. In [7], the author has examined the optimal approach for a double orbit retrial

queue with imperfect service and vacation interruption. Wang [8] conducted research on a retrial make-to-stock system based on a double-ended queue. While the queue length was positive, customers were waiting in the retrial queue; when it was negative, the inventory system contained only products. In order to calculate stationary distributions, Avrachenkov et al. [9], the authors examined a single-server retrial system that was fed by two customer types, each of which entered its own infinite-capacity orbit. They suggested approximate algorithms through orbit truncation and treated the underlying Markov chain as an asymptotically quasi-Toeplitz Markov chain (AQTMC).

Feedback is another often occurring event in queuing systems. Queues with feedback (QFB) are those where each consumer that is served either enters the system with a certain probability or departs forever with a complimentary probability. An $M/M/1$ retrial-queueing system with feedback and server breakdowns is presented by [10]. The matrix-geometric method is used to determine performance measures, including graphical representations, and feedback happens when dissatisfied clients return to the orbit for another service attempt. [11] conducted research on the retrial queueing model with re-service. According to [12], a queueing inventory system with consumer feedback and returns of purchased items was stated. Ayyappan and Archana [13] concentrate on two-server tandem queues with working vacation, feedback, re-service, unreliable servers, interrupted shutdowns, and impatient clients.

When the service is finished and the server discovers that there are no customers with positive inventory in the system, the server takes a vacation. The server may provide a lower service rate than usual if customers arrive while on vacation. Servi and Finn [14] established the concept of working vacation. In [15], the author examined a single-server working vacation queueing system with interconnected arrival and service operations. Note that in [16], the authors discussed the MAP/Ek/1 Queue with Working Vacation. Ayyappan and Arulmozhi [17] examine the retrial queueing inventory system with orbital client collision, a constant retrial rate, working vacation, flush out, balking, breakdown, and repair.

According to [18], they examined a queueing system in which two different service devices served a single demand at varying service rates and energy unit consumption. Stability of queueing systems with impatience, balking and non-persistence of customers was analyzed by [19]. Customers may balk when the server is on vacation or experiencing a breakdown, as assumed by [20]. A queueing inventory system with impatient customers and working vacations was studied by [21].

The literature on stochastic inventory systems with single and multiple commodities is quite large. We provide some of the significant research conducted in two commodity stochastic inventory systems in this part. [22] and [23] looked at problems with inventory of two commodities with Markov change in demands. Note that in [24], the authors studied a single-server queueing system and two-commodity inventory where the buffer capacities for both commodity stocks are finite and customer demands (for either or both commodities) have been represented by given probabilities. In [25], a two-commodity queueing inventory system with individual ordering policy, positive service time, positive lead time, and random common lifetimes for each commodity was examined.

This is how the paper is organized. A thorough description of the model can be found in 2. Section 3 explains matrix formation and a number of notations. The invariant probability vector, the rate matrix R and the stability condition are all obtained in section 4. The analysis of the busy period is described in Section 5. The performance measurements are listed in 6.

Section 7 contains some numerical and graphical results. The conclusion is presented in Section 8.

2 MATHEMATICAL MODEL

This study examines a two-server retrial inventory queuing model with two orbits. For primary customers who are unable to reach the server upon arrival, the first orbit is infinite, whereas the second orbit is finite for re-service customers. The following is an explanation of the model's basic operation:

- **Arrival process:** The customer arrives in accordance with MAP with parameter matrices (D_0, D_1) of order m . Let D_0 be the transition matrix that indicates that no customers have arrived, and let D_1 be the transition rate matrix that indicates that customers have arrived. Here, we assumed that if a customer arrives and discovers that server-1 is idle with positive inventory, they will receive service from server-1, which serves as their primary service provider; if not, they will join orbit I (infinite orbit). Following primary service, the customer either leaves the system because they are satisfied with the service or moves into orbit II (finite capacity) for another service.
- **Service Process:** The two servers provide the service with positive inventory at two distinct service rates. The predicted service rates for servers 1 and 2 are $\mu_1 = [\gamma_1(-U_1)^{-1}e_{n_1}]$ and $\mu_2 = [\gamma_2(-U_2)^{-1}e_{n_2}]$ respectively. In this case, server- 2 is always accessible for re-service, while server-1 is operating in normal mode, vacation mode, shutdown, breakdown, and repair. Each server will receive one item from the server upon completion of the service. The customer either enters an orbit (finite) for reservice after server- 1 completes the primary service or leaves the system entirely since they were satisfied. The service time of both servers follows phase-type distribution with representation (γ_1, U_1) of order n_1 , such that $U_1^0 + U_1 e = 0$ and (γ_2, U_2) of order n_2 , such that $U_2^0 + U_2 e = 0$, respectively.
- **Breakdown and repair state:** The server- 1 is prone to breaking down during busy periods and is immediately repaired. Server-1 pauses the customer who is currently in service, putting them in a frozen condition. Once the repair is finished, server- 1 provides the interrupted consumer with new service. For server-2, there is no malfunction. ψ and δ represent the exponential distribution of breakdown and repair time, respectively.
- **Retrial process:** Every customer in the orbit makes an independent effort to regain access. The attempt is successful and the customer takes up a free server right away. It is assumed that the customers' retrial time is exponential, with parameters σ_2 (finite orbit) and σ_1 (infinite orbit), respectively.
- **Re-service period:** Following primary service from server-1, a customer who is satisfied with the service may either leave the system with probability f_1 or enter orbit-II for re-service from server- 2 with probability f_2 .
- **Close down:** If there are no customers in the system once the service is finished, server-1 proceeds to close down. Close down time is represented by the symbol ϕ and is assumed to have an exponential distribution.
- **Vacation time:** When the close-down procedure is finished and there are no customers in the system, server-1 starts the vacation period. Its rate is represented by η , which is followed by an exponential distribution. When a customer arrives with positive inventory during vacation time, server-1 serves them at a slow rate in phase-type

distribution, represented as $(\alpha, \theta U_1)$ of order n , with a working vacation rate of $\theta\mu_1$. If server-1 gets an empty system at the time of working vacation mode, then server-1 moves to idle in working vacation mode. After vacation completion, the server will return to its normal working mode if there are both positive items and customers in the system. Otherwise, server-1 becomes idle if zero customers in the system or zero items or both.

- Replenishment: When a customer completes a service, one item is removed from stock. We consider a stochastic inventory system with two different items in stock; one item is for server-1 (I commodity) and the other item is for server-2 (II-commodity). The maximum storage capacity for the i^{th} commodity is S_i ($i = 1, 2$). An (s_1, S_1) ordering policy is adopted for the first commodity, (s_2, S_2) ordering policy is adopted for the second commodity. When the on-hand inventory level of the I-commodity drops to a prefixed level s_1 , an order for S_1 units is placed. The lead time of commodity-I is exponentially distributed with parameter β . In addition, assume that an emergency replenishment of one item with zero lead time takes place when the on-hand inventory level decreases to zero. The emergency replenishment is incorporated in the system to ensure customer satisfaction for re-service customers in orbit-II. Refer to figure1 for a schematic representation of the model.

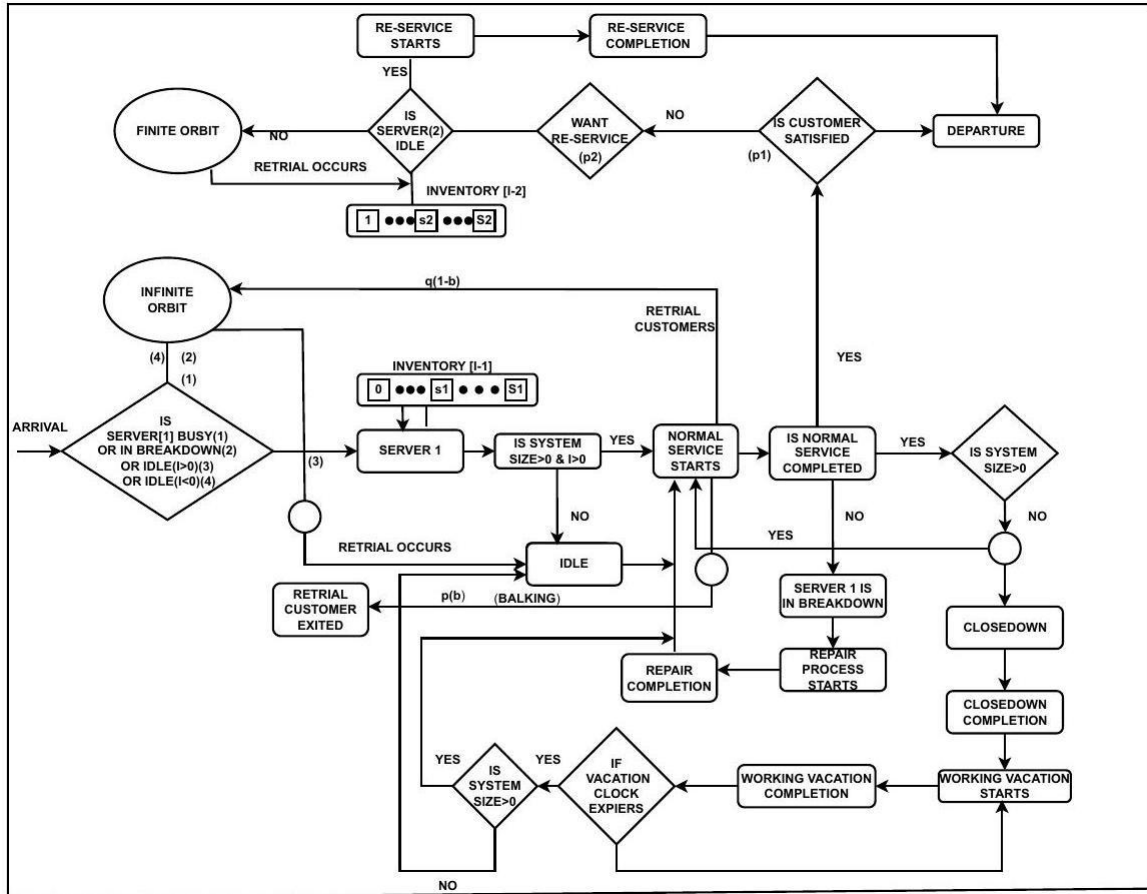


Figure 1: Schematic representation

3 GENERATION MATRIX

Let us describe the few notations of this model that are followed by the Quasi Birth-Death (QBD) process generating matrix as follows:

ABBREVIATION AND NOTATIONS

- CTMC- Continuous Time Markov Chain
- QBD process - Quasi-Birth-Death Process
- QIS- Queueing Inventory System
- MAM - Matrix Analytic Method
- e - entries in a column's vector are one;
- I_n - Identity matrix of order n
- 0 - Matrix whose entries are 0, of appropriate size
- $\otimes$ - Kronecker product of two matrices of different dimensions;
- $\oplus$ - Addition of two matrices of different dimensions;
- $N_1(t)$ indicates the number of customers in an infinite orbit;
- $N_2(t)$ indicates the number of customers in the finite orbit;
- $J_1(t)$ represents the server-1 status;
- $J_2(t)$ represents the server-2 status;
- $I_1(t)$ represents the stock level of commodity-I;
- $I_2(t)$ represents the stock level of commodity-II;
- $S_1(t)$ represents the service phases of server-1;
- $S_2(t)$ represents the service phases of server-2;
- $A(t)$ represents the arrival phases; at time t , where

$$J_1(t) = \begin{cases} 0, & \text{server- 1 is idle in normal mode ,} \\ 1, & \text{server- 1 is busy in normal mode ,} \\ 2, & \text{server- 1 is in repair,} \\ 3, & \text{server- 1 is in closedown,} \\ 4, & \text{server- 1 is idle in working vacation mode,} \\ 5, & \text{server- 1 is busy in working vacation mode} \end{cases}$$

$$J_2(t) = \begin{cases} 0, & \text{server- 2 is idle,} \\ 1, & \text{server- 2 is working with re-service} \end{cases}$$

It is obvious that $\{N_1(t), N_2(t), J_1(t), J_2(t), I_1(t), I_2(t), S_1(t), S_2(t), A(t): t \geq 0\}$ is a CTMC with State Space is as follows:

$$\omega = \bigcup_{d_1=0}^{\infty} \chi(d_1)\#(1)$$

$$\begin{aligned}
 \chi(d_1) = & \{(d_1, d_2, 0, 0, d_5, d_6, d_9): 0 \leq d_2 \leq M, 0 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 0, 1, d_5, d_6, d_8, d_9): 0 \leq d_2 \leq M, 0 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_8 \leq n_2, \\
 & \quad 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 1, 0, d_5, d_6, d_7, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, \\
 & \quad 1 \leq d_7 \leq n_1, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 1, 1, d_5, d_6, d_7, d_8, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, \\
 & \quad 1 \leq d_7 \leq n_1, 1 \leq d_8 \leq n_2, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 2, 0, d_5, d_6, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 2, 1, d_5, d_6, d_8, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, \\
 & \quad 1 \leq d_8 \leq n_2, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 3, 0, d_5, d_6, d_9): 0 \leq d_2 \leq M, 0 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 3, 1, d_5, d_6, d_8, d_9): 0 \leq d_2 \leq M, 0 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_8 \leq n_2, \\
 & \quad 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 4, 0, d_5, d_6, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 4, 1, d_5, d_6, d_8, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_8 \leq n_2, \\
 & \quad 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 5, 0, d_5, d_6, d_7, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, 1 \leq d_7 \leq n_1, \\
 & \quad 1 \leq d_9 \leq m\} \\
 & \cup \{(d_1, d_2, 5, 1, d_5, d_6, d_7, d_8, d_9): 0 \leq d_2 \leq M, 1 \leq d_5 \leq S_1, 1 \leq d_6 \leq S_2, \\
 & \quad 1 \leq d_7 \leq n_1, 1 \leq d_8 \leq n_2, 1 \leq d_9 \leq m\}
 \end{aligned}$$

QUASI BIRTH AND DEATH (QBD) PROCESS'S CONSTRUCTION

The following is the Markov chain's generator matrix

$$Q = \begin{bmatrix} B_0 & A_0 & 0 & 0 & 0 & 0 & \dots \\ A_2 & A_1 & A_0 & 0 & 0 & 0 & \dots \\ 0 & A_2 & A_1 & A_0 & 0 & 0 & \dots \\ 0 & 0 & A_2 & A_1 & A_0 & 0 & \dots \\ \vdots & \vdots & \vdots & \ddots & \ddots & \ddots & \ddots \\ \vdots & \vdots & \vdots & \vdots & \ddots & \ddots & \ddots \end{bmatrix}, \#(2)$$

The following defines each entry in Q's block matrices:

$$B_0 = \begin{bmatrix} B_0^{11} & B_0^{12} & B_0^{13} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ B_0^{21} & B_0^{22} & 0 & B_0^{24} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & B_0^{33} & B_0^{34} & 0 & 0 & B_0^{37} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & B_0^{43} & B_0^{44} & 0 & B_0^{46} & 0 & B_0^{48} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & B_0^{53} & 0 & B_0^{55} & B_0^{56} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & B_0^{64} & B_0^{65} & B_0^{66} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & B_0^{77} & B_0^{78} & B_0^{79} & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & B_0^{87} & B_0^{88} & 0 & B_0^{810} & 0 & 0 & 0 \\ B_0^{91} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & B_0^{99} & B_0^{910} & B_0^{911} & 0 & 0 \\ 0 & B_0^{102} & 0 & 0 & 0 & 0 & 0 & 0 & B_0^{109} & B_0^{1010} & 0 & 0 & 0 \\ 0 & 0 & B_0^{113} & 0 & 0 & 0 & 0 & 0 & B_0^{119} & 0 & B_0^{1111} & B_0^{1112} & 0 \\ 0 & 0 & 0 & B_0^{124} & 0 & 0 & 0 & 0 & 0 & B_0^{1210} & B_0^{1211} & B_0^{1212} & 0 \end{bmatrix},$$

$$B_0^{11} = \begin{bmatrix} K_1 & 0 \\ 0 & I_{S_1} \otimes K_2 \end{bmatrix}, B_0^{13} = I_{S_1+1} \otimes K_4,$$

$$K_2 = \begin{bmatrix} N_4 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_5 \\ 0 & N_4 & 0 & \dots & 0 & 0 & \dots & 0 & N_5 \\ 0 & 0 & N_4 & \dots & 0 & 0 & \dots & 0 & N_5 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_4 & 0 & \dots & 0 & N_5 \\ 0 & 0 & 0 & \dots & 0 & N_6 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_6 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_6 \end{bmatrix}, K_1 = \begin{bmatrix} N_1 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_2 \\ 0 & N_1 & 0 & \dots & 0 & 0 & \dots & 0 & N_2 \\ 0 & 0 & N_1 & \dots & 0 & 0 & \dots & 0 & N_2 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_1 & 0 & \dots & 0 & N_2 \\ 0 & 0 & 0 & \dots & 0 & N_3 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_3 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_3 \end{bmatrix},$$

where $N_1 = I_{s_1+1} \otimes I_{S_2} \otimes D_0 - \beta I_m$, $N_2 = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$,

$$\begin{aligned} N_3 &= I_{(s_1+1)} - (s_1 + 1) \otimes I_{S_2} \otimes D_0, N_4 = I_{s_1+1} \otimes I_{S_2} \otimes D_0 - (\beta + \sigma_2)I_m, \\ N_2 &= N_5, N_6 = I_{(s_1+1)} - (s_1 + 1) \otimes I_{S_2} \otimes D_0 - \sigma_2, \\ B_0^{12} &= \begin{bmatrix} 0 & 0 \\ I_{S_1} \otimes K_3 & 0 \end{bmatrix}, K_3 = I_{s_1+1} \otimes N_7, N_7 = I_{s_1+1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_2 I_m, \\ K_4 &= \begin{bmatrix} 0 \\ I_{s_1+1} \otimes N_8 \end{bmatrix}, B_0^{21} = I_{s_1+1} \otimes K_5, N_8 = I_{S_2} \otimes \gamma_1 \otimes D_1, \\ K_5 &= \begin{bmatrix} N_9 & 0 \\ I_{s_1+1} \otimes & 0 \end{bmatrix}, B_0^{22} = I_{s_1+1} \otimes K_6, N_9 = U2 \otimes I_m, \\ K_6 &= \begin{bmatrix} N_{10} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{11} \\ 0 & N_{10} & 0 & \dots & 0 & 0 & \dots & 0 & N_{11} \\ 0 & 0 & N_{10} & \dots & 0 & 0 & \dots & 0 & N_{11} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{10} & 0 & \dots & 0 & N_{11} \\ 0 & 0 & 0 & \dots & 0 & N_{12} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{12} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{12} \end{bmatrix}, B_0^{24} = I_{s_1+1} \otimes K_7, \end{aligned}$$

where $N_{10} = I_{S_2} \otimes (U2 \oplus D_0) - \beta I_{n_2m}$, $N_{11} = e_{s_2} \otimes I_{n_2} \otimes \beta I_m$, $N_{12} = I_{S_2} \otimes (U2 \oplus D_0)$,

$$K_7 = \begin{bmatrix} 0 \\ I_{S_1} \otimes N_{13} \end{bmatrix}, B_0^{33} = \begin{bmatrix} K_8 & 0 \\ 0 & I_{s_1-1} \otimes K_9 \end{bmatrix}, N_{13} = I_{S_2} \otimes I_{n_2} \otimes \gamma_1 \otimes D_0$$

$$K_8 = \begin{bmatrix} N_{14} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{15} \\ 0 & N_{14} & 0 & \dots & 0 & 0 & \dots & 0 & N_{15} \\ 0 & 0 & N_{14} & \dots & 0 & 0 & \dots & 0 & N_{15} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{14} & 0 & \dots & 0 & N_{15} \\ 0 & 0 & 0 & \dots & 0 & N_{16} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{16} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{16} \end{bmatrix}, K_9 = \begin{bmatrix} N_{17} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{18} \\ 0 & N_{17} & 0 & \dots & 0 & 0 & \dots & 0 & N_{18} \\ 0 & 0 & N_{17} & \dots & 0 & 0 & \dots & 0 & N_{18} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{17} & 0 & \dots & 0 & N_{18} \\ 0 & 0 & 0 & \dots & 0 & N_{19} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{19} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{19} \end{bmatrix},$$

where $N_{14} = I_{s_1} \otimes I_{S_2} \otimes (U1 \oplus D_0) - \beta I_{n_1m}$, $N_{15} = N_{18} = e_{s_1} \otimes I_{S_2} \otimes \beta I_{n_1} \otimes \beta I_m$

$N_{16} = I_{(s_1-s_1)} \otimes I_{S_2} \otimes (U1 \oplus D_0)$, $N_{17} = I_{s_1} \otimes I_{S_2} \otimes (U1 \oplus D_0) - (\beta + \sigma_2)I_{n_1m}$,

$N_{19} = I_{(s_1-s_1)} \otimes I_{S_2} \otimes (U1 \oplus D_0) - \sigma_2 I_{n_1m}$,

$$B_0^{34} = \begin{bmatrix} 0 \\ I_{s_1-1} \otimes K_{10} & 0 \end{bmatrix}, K_{10} = \begin{bmatrix} N_{20} & 0 \\ 0 & I_{s_1} \otimes N_{20} \end{bmatrix}, N_{20} = I_{s_1s_2} \otimes I_{n_1} \otimes \gamma_2 \otimes \sigma_2 I_m,$$

$$B_0^{37} = \begin{bmatrix} K_{11} & K_{12} & 0 & \dots & 0 & 0 \\ 0 & K_{11} & K_{12} & \dots & 0 & 0 \\ 0 & 0 & K_{11} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & K_{11} & K_{12} \\ 0 & 0 & 0 & \dots & 0 & K_{11} + K_{12} \end{bmatrix},$$

where $K_{11} = [I_{s_1s_2} \otimes f_1 U1^0 \otimes I_m \quad 0]$, $K_{12} = [I_{s_1s_2} \otimes f_2 U1^0 \otimes I_m \quad 0]$,

$$B_0^{43} = [I_{M+1} \otimes I_{S_1} \otimes K_{13}],$$

$$K_{13} = \begin{bmatrix} N_{21} & 0 \\ I_{s_1-1} \otimes N_{21} & 0 \end{bmatrix}, B_0^{44} = I_{s_1} \otimes K_{14},$$

where $N_{21} = I_{n_1} \otimes U2^0 \otimes I_m$, $B_0^{46} = I_{M+1} \otimes I_{S_1} \otimes I_{S_2} \otimes e_{n_1} \otimes \tau I_{n_2m}$

$$K_{14} = \begin{bmatrix} N_{22} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{23} \\ 0 & N_{22} & 0 & \dots & 0 & 0 & \dots & 0 & N_{23} \\ 0 & 0 & N_{22} & \dots & 0 & 0 & \dots & 0 & N_{23} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{22} & 0 & \dots & 0 & N_{23} \\ 0 & 0 & 0 & \dots & 0 & N_{24} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{24} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{24} \end{bmatrix}, B_0^{48} = \begin{bmatrix} K_{17} & K_{18} & 0 & \dots & 0 & 0 \\ 0 & K_{17} & K_{18} & \dots & 0 & 0 \\ 0 & 0 & K_{12} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & K_{17} & K_{18} \\ 0 & 0 & 0 & \dots & 0 & K_{17} + K_{18} \end{bmatrix},$$

where $N_{22} = I_{s_1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\tau + \beta)I_{n_1n_2m}$, $N_{23} = e_{s_1} \otimes I_{S_2} \otimes I_{n_1n_2} \otimes \beta I_m$

$N_{24} = I_{s_1-s_1} \otimes I_{S_2} \otimes [(U1 \oplus U2) \oplus D_0] - \tau I_{n_1n_2m}$, $K_{15} = [I_{s_1} \otimes I_{S_2} \otimes f_1 U1^0 \otimes I_{n_2} \otimes I_m \quad 0]$,

$K_{16} = [I_{s_1} \otimes I_{S_2} \otimes f_2 U1^0 \otimes I_{n_2} \otimes I_m \quad 0]$, $B_0^{53} = [I_{M+1} \otimes I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes \delta I_m]$,

$$B_0^{55} = \begin{bmatrix} K_{17} & 0 \\ 0 & I_{s_1-1} \otimes K_{18} \end{bmatrix}, K_{17} = \begin{bmatrix} N_{25} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{26} \\ 0 & N_{25} & 0 & \dots & 0 & 0 & \dots & 0 & N_{26} \\ 0 & 0 & N_{25} & \dots & 0 & 0 & \dots & 0 & N_{26} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{25} & 0 & \dots & 0 & N_{26} \\ 0 & 0 & 0 & \dots & 0 & N_{27} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{27} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{27} \end{bmatrix},$$

$N_{25} = I_{s_1} \otimes I_{S_2} \otimes D_0 - (\delta + \beta)I_m$, $N_{26} = e_{s_1} \otimes I_{S_2} \otimes \beta I_m$, $N_{27} = I_{s_1-s_1} \otimes I_{S_2} \otimes D_0 - \delta I_m$,

$$K_{18} = \begin{bmatrix} N_{28} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{29} \\ \mathbf{0} & N_{28} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{29} \\ \mathbf{0} & \mathbf{0} & N_{28} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{29} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{28} & \mathbf{0} & \dots & \mathbf{0} & N_{29} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{30} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{30} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{30} \end{bmatrix}, B_0^{56} = \begin{bmatrix} 0 & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & 0 \\ K_{19} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & 0 \\ \mathbf{0} & K_{19} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \dots & K_{19} & \mathbf{0} & \dots & \mathbf{0} & 0 \\ \mathbf{0} & \mathbf{0} & \dots & 0 & K_{19} & \dots & \mathbf{0} & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \dots & K_{19} & 0 \end{bmatrix},$$

where $N_{28} = I_{S_1} \otimes I_{S_2} \otimes D_0 - (\delta + \beta + \sigma_2)I_m$, $N_{29} = e_{s_1} \otimes I_{S_2} \otimes \beta I_m$

$N_{30} = I_{S_1-s_1} \otimes I_{S_2} \otimes D_0 - (\delta + \sigma_2)I_m$, $K_{19} = I_{S_1 S_2} \otimes \gamma_2 \otimes \sigma_2 I_m$,

$B_0^{64} = I_{M+1} \otimes I_{S_1 S_2} \otimes \gamma_1 \otimes I_{n_2} \otimes \delta I_m$,

$B_0^{65} = [I_{S_1} \otimes K_{20}]$, $B_0^{66} = [I_{S_1} \otimes K_{21}]$, $K_{20} = I_{S_1} \otimes I_{S_2} \otimes U2^0 \otimes I_m$,

$$K_{21} = \begin{bmatrix} N_{31} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{32} \\ \mathbf{0} & N_{31} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{32} \\ \mathbf{0} & \mathbf{0} & N_{31} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{32} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{31} & \mathbf{0} & \dots & \mathbf{0} & N_{32} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{33} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{33} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{33} \end{bmatrix}, B_0^{77} = \begin{bmatrix} K_{22} & 0 \\ 0 & I_{S_1} \otimes K_{23} \end{bmatrix},$$

where $N_{31} = I_{S_1} \otimes I_{S_2} \otimes (U2 + D_0) - (\delta + \beta)I_{n_2 m}$, $N_{32} = e_{s_1} \otimes I_{S_2} \otimes I_{n_2} \otimes \beta I_m$

$N_{33} = I_{S_1-s_1} \otimes I_{S_2} \otimes (U2 + D_0) - \delta I_{n_2 m}$,

where $N_{34} = I_{S_1+1} \otimes I_{S_2} \otimes D_0 - (\delta + \beta)I_m$, $N_{35} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$

$N_{36} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes D_0 - \phi I_m$, $N_{37} = I_{S_1+1} \otimes I_{S_2} \otimes D_0 - (\phi + \sigma_2 + \beta)I_m$

$N_{38} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$, $N_{39} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes D_0 - (\phi + \sigma_2)I_m$

$B_0^{78} = \begin{bmatrix} 0 & 0 \\ I_{S_1} \otimes K_{24} & 0 \end{bmatrix}$, $B_0^{87} = I_{S_1+1} \otimes K_{25}$,

$K_{24} = I_{S_1+1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_2 I_m$, $B_0^{810} = I_{M+1} \otimes I_{(S_1+1)S_2} \otimes \phi I_m$,

$K_{25} = I_{S_1+1} \otimes U2^0 \otimes I_m$,

$$B_0^{88} = I_{S_1+1} \otimes K_{26}, K_{26} = \begin{bmatrix} N_{40} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{41} \\ \mathbf{0} & N_{40} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{41} \\ \mathbf{0} & \mathbf{0} & N_{40} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{41} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{40} & \mathbf{0} & \dots & \mathbf{0} & N_{41} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{42} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{42} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{42} \end{bmatrix},$$

where $N_{40} = I_{S_1+1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\phi + \beta)I_{n_2 m}$, $N_{41} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_{n_2 m}$

$N_{42} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes (U2 \otimes D_0) - \phi I_{n_2 m}$,

$B_0^{91} = I_{M+1} \otimes I_{S_1+1} \otimes \eta I_m$,

$$B_0^{99} = \begin{bmatrix} K_{27} & 0 \\ 0 & I_{S_1-1} \otimes K_{28} \end{bmatrix}, K_{27} = \begin{bmatrix} N_{43} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{44} \\ \mathbf{0} & N_{43} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{44} \\ \mathbf{0} & \mathbf{0} & N_{43} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{44} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{43} & \mathbf{0} & \dots & \mathbf{0} & N_{44} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{45} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{45} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{45} \end{bmatrix},$$

where

$N_{43} = I_{S_1+1} \otimes I_{S_2} \otimes D_0 - (\eta + \beta)I_m$, $N_{44} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$

$N_{45} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes D_0 - \eta I_m$,

$$K_{28} = \begin{bmatrix} N_{46} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{47} \\ \mathbf{0} & N_{46} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{47} \\ \mathbf{0} & \mathbf{0} & N_{46} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{47} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{46} & \mathbf{0} & \dots & \mathbf{0} & N_{47} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{48} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{48} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{48} \end{bmatrix}, B_0^{910} = \begin{bmatrix} \mathbf{0} \\ I_{S_1} \otimes K_{29} & \mathbf{0} \end{bmatrix}$$

where $N_{46} = I_{s_1+1} \otimes I_{S_2} \otimes D_0 - (\eta + \beta + \sigma_2)I_m$, $N_{47} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$, $N_{48} = I_{(s_1+1)-(s_1+1)} \otimes I_{S_2} \otimes D_0 - \eta I_m$, $K_{29} = I_M \otimes I_{s_1+1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_2 I_m$, $B_0^{102} = I_{M+1} \otimes I_{(s_1+1)S_2} \otimes I_{n_2} \otimes \eta I_m$,

$$B_0^{911} = I_{S_1} \otimes K_{30}, K_{30} = \begin{bmatrix} \mathbf{0} \\ I_{S_1} \otimes N_{49} \end{bmatrix}, N_{49} = I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes D_1, \\ B_0^{109} = \begin{bmatrix} K_{31} & \mathbf{0} \\ I_{S_1-1} \otimes K_{31} \end{bmatrix}, B_0^{1010} = I_{S_1} \otimes K_{32}, K_{31} = I_{M+1} \otimes I_{s_1+1} \otimes U2^0 \otimes I_m, \\ K_{32} = \begin{bmatrix} N_{50} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{51} \\ \mathbf{0} & N_{50} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{51} \\ \mathbf{0} & \mathbf{0} & N_{50} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{51} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{50} & \mathbf{0} & \dots & \mathbf{0} & N_{51} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{52} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{52} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{52} \end{bmatrix}, B_0^{1112} = \begin{bmatrix} \mathbf{0} \\ I_{S_1-1} \otimes K_{37} & \mathbf{0} \end{bmatrix},$$

where $N_{49} = I_{s_1+1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\eta + \beta)I_{n_2m}$, $N_{50} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_{n_2m}$, $N_{51} = I_{(s_1+1)-(s_1+1)} \otimes I_{S_2} \otimes (U2 \oplus D_0) - \eta I_{n_2m}$, $K_{37} = I_{S_1S_2} \otimes I_{n_1} \otimes \gamma_2 \otimes \sigma_2 I_m$, $B_0^{113} = I_{M+1} \otimes I_{S_1S_2} \otimes I_{n_1} \otimes \eta I_m$

$$B_0^{119} = \begin{bmatrix} K_{33} & K_{34} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & K_{33} & K_{34} & \dots & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & K_{33} & \dots & \mathbf{0} & \mathbf{0} \\ & & & \ddots & \ddots & \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & K_{33} & K_{34} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & K_{33} + K_{34} \end{bmatrix}, B_0^{1111} = \begin{bmatrix} K_{35} & \mathbf{0} \\ \mathbf{0} & I_{S_1-1} \otimes K_{36} \end{bmatrix},$$

where $K_{33} = [I_{S_1S_2} \otimes f_1 \theta U1^0 \otimes I_m \quad \mathbf{0}]$, $K_{34} = [I_{S_1S_2} \otimes f_2 \theta U1^0 \otimes I_m \quad \mathbf{0}]$, $B_0^{124} = I_{M+1} \otimes I_{S_1S_2} \otimes I_{n_1} \otimes \gamma_1 \otimes \eta I_m$

$$K_{35} = \begin{bmatrix} N_{53} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{54} \\ 0 & N_{53} & 0 & \dots & 0 & 0 & \dots & 0 & N_{54} \\ 0 & 0 & N_{53} & \dots & 0 & 0 & \dots & 0 & N_{54} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{53} & 0 & \dots & 0 & N_{54} \\ 0 & 0 & 0 & \dots & 0 & N_{55} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{55} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{55} \end{bmatrix}, K_{36} = \begin{bmatrix} N_{56} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{51} \\ 0 & N_{56} & 0 & \dots & 0 & 0 & \dots & 0 & N_{51} \\ 0 & 0 & N_{56} & \dots & 0 & 0 & \dots & 0 & N_{51} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{56} & 0 & \dots & 0 & N_{51} \\ 0 & 0 & 0 & \dots & 0 & N_{57} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{57} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{57} \end{bmatrix},$$

where $N_{53} = I_{S_1} \otimes I_{S_2} \otimes (\theta U1 \oplus D_0) - (\eta + \beta)I_{n_1 m}$, $N_{54} = e_{s_1} \otimes I_{S_2} \otimes \beta I_{n_1 m}$

$N_{55} = I_{(S_1-s_1)} \otimes I_{S_2} \otimes (\theta U1 \oplus D_0) - \eta I_{n_1 m}$, $N_{56} = I_{S_1} \otimes I_{S_2} \otimes (\theta U1 \oplus D_0) - (\eta + \sigma_2 + \beta)I_{n_1 m}$,

$N_{57} = I_{(S_1-s_1)} \otimes I_{S_2} \otimes (\theta U1 \oplus D_0) - (\sigma + \eta)I_{n_1 m}$,

$$B_0^{1210} = \begin{bmatrix} K_{38} & K_{39} & 0 & \dots & 0 & 0 \\ 0 & K_{38} & K_{39} & \dots & 0 & 0 \\ 0 & 0 & K_{38} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & K_{38} & K_{39} \\ 0 & 0 & 0 & \dots & 0 & K_{38} + K_{39} \end{bmatrix},$$

where $K_{38} = [I_{S_1 S_2} \otimes f_1 \theta U1^0 \otimes I_{n_2} \otimes I_m 0]$, $K_{39} = [I_{S_1 S_2} \otimes f_2 \theta U1^0 \otimes I_{n_2} \otimes I_m 0]$,

$$B_0^{1211} = I_{S_1} \otimes K_{40}, K_{40} = \begin{bmatrix} N_{58} & 0 \\ I_{S_1-1} \otimes N_{58} & 0 \end{bmatrix},$$

$$B_0^{1212} = I_{S_1} \otimes K_{41}, K_{41} = \begin{bmatrix} N_{59} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{60} \\ 0 & N_{59} & 0 & \dots & 0 & 0 & \dots & 0 & N_{60} \\ 0 & 0 & N_{59} & \dots & 0 & 0 & \dots & 0 & N_{60} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{59} & 0 & \dots & 0 & N_{60} \\ 0 & 0 & 0 & \dots & 0 & N_{61} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{61} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{61} \end{bmatrix},$$

where $N_{58} = I_{S_1} \otimes I_{n_1} \otimes U2^0 \otimes I_m$, $N_{59} = I_{S_1} \otimes I_{S_2} \otimes ((\theta U1 + U2) \oplus D_0) - \eta I_{n_1 n_2 m}$,

$N_{60} = e_{s_1} \otimes I_{S_2} \otimes \beta I_{n_1 n_2 m}$, $N_{61} = I_{(S_1-s_1)} \otimes I_{S_2} \otimes ((\theta U1 + U2) \oplus D_0) - \eta I_{n_1 n_2 m}$,

$$A_0 = \begin{bmatrix} A_0^{11} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & A_0^{22} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & A_0^{33} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & A_0^{44} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & A_0^{55} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & A_0^{66} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & A_0^{77} & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_0^{88} & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_0^{99} & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_0^{1010} & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_0^{1111} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_0^{1212} \end{bmatrix},$$

$A_0^{11} = I_{S_1+1} \otimes K_{42}$, $A_0^{22} = I_{S_1+1} \otimes K_{43}$,

$$\begin{aligned}
 K_{42} &= \begin{bmatrix} I_{S_2} \otimes D_1 & 0 \\ 0 & 0 \end{bmatrix}, K_{43} = \begin{bmatrix} I_{S_2 n_2} \otimes D_1 & 0 \\ 0 & 0 \end{bmatrix}, \\
 A_0^{33} &= I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1 \otimes D_1}, A_0^{44} = I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1 n_2} \otimes D_1, \\
 A_0^{55} &= I_{M+1} \otimes I_{S_1 S_2} \otimes D_1, A_0^{66} = I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1 n_2} \otimes D_1, \\
 A_0^{77} &= I_{M+1} \otimes I_{(S_1+1)S_2} \otimes D_1, A_0^{88} = I_{M+1} \otimes I_{(S_1+1)S_2} \otimes I_{n_2} \otimes D_1, \\
 A_0^{99} &= I_{M+1} \otimes I_{(S_1+1)S_2} \otimes D_1, A_0^{1010} = I_{M+1} \otimes I_{(S_1+1)S_2} \otimes I_{n_2} \otimes D_1, \\
 A_0^{1111} &= I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1} \otimes D_1, A_0^{1212} = I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1 n_2} \otimes D_1, \\
 A_2 &= \begin{bmatrix} 0 & 0 & A_2^{13} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & A_2^{24} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_2^{911} & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_2^{1012} \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix}, A_2^{13} = I_{S_1+1} \otimes K_{44}, \\
 K_{44} &= \begin{bmatrix} 0 \\ I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes \sigma_1 I_m \end{bmatrix}, A_2^{24} = I_{S_1+1} \otimes K_{45}, A_2^{911} = I_{S_1} \otimes K_{46}, \\
 K_{45} &= \begin{bmatrix} 0 \\ I_{S_1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_1 I_m \end{bmatrix}, K_{46} = \begin{bmatrix} 0 \\ I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes \sigma_1 I_m \end{bmatrix}, \\
 A_2^{1012} &= I_{S_1} \otimes K_{47}, A_1^{11} = \begin{bmatrix} K_{48} & 0 \\ 0 & I_{S_1-1} \otimes K_{49} \end{bmatrix}, K_{47} = \begin{bmatrix} 0 \\ I_{S_1 S_2} \otimes I_{n_1} \otimes \gamma_2 \otimes \sigma_1 I_m \end{bmatrix}, \\
 K_{48} &= \begin{bmatrix} N_{62} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{63} \\ 0 & N_{64} & 0 & \dots & 0 & 0 & \dots & 0 & N_{63} \\ 0 & 0 & N_{64} & \dots & 0 & 0 & \dots & 0 & N_{63} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{64} & 0 & \dots & 0 & N_{63} \\ 0 & 0 & 0 & \dots & 0 & N_{65} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{65} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{65} \end{bmatrix}, K_{49} = \begin{bmatrix} N_{66} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{67} \\ 0 & N_{68} & 0 & \dots & 0 & 0 & \dots & 0 & N_{67} \\ 0 & 0 & N_{68} & \dots & 0 & 0 & \dots & 0 & N_{67} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{68} & 0 & \dots & 0 & N_{67} \\ 0 & 0 & 0 & \dots & 0 & N_{69} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{69} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{69} \end{bmatrix}, \\
 \text{where} \quad N_{62} &= I_{S_2} \otimes D_0 - \beta I_m, N_{63} = e_{s_1} \otimes I_{S_2} \otimes \beta I_m, N_{64} = I_{S_2} \otimes D_0 - (\beta + \sigma_1) I_m, \\
 N_{65} &= I_{S_2} \otimes D_0 - \sigma_1 I_m, N_{66} = I_{S_2} \otimes D_0 - (\beta + \sigma_2) I_m, N_{67} = e_{s_1} \otimes I_{S_2} \otimes \beta I_m, \\
 N_{68} &= I_{S_2} \otimes D_0 - (\beta + \sigma_1 + \sigma_2) I_m, N_{69} = I_{S_2} \otimes D_0 - (\sigma_1 + \sigma_2) I_m,
 \end{aligned}$$

$$A_1 = \begin{bmatrix} A_1^{11} & A_1^{12} & A_1^{13} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ A_1^{21} & A_1^{22} & 0 & A_1^{24} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ A_1^{31} & 0 & A_1^{33} & A_1^{34} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & A_1^{42} & A_1^{43} & A_1^{44} & 0 & A_1^{46} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & A_1^{53} & 0 & A_1^{55} & A_1^{56} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & A_1^{64} & A_1^{65} & A_1^{66} & 0 & 0 & 0 & 0 & 0 & 0 \\ A_1^{71} & 0 & 0 & 0 & 0 & 0 & A_1^{77} & A_1^{78} & 0 & 0 & 0 & 0 \\ 0 & A_1^{82} & 0 & 0 & 0 & 0 & A_1^{87} & A_1^{88} & 0 & 0 & 0 & 0 \\ A_1^{91} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A_1^{99} & A_1^{910} & A_1^{911} & 0 \\ 0 & A_1^{102} & 0 & 0 & 0 & 0 & 0 & 0 & A_1^{109} & A_1^{1010} & 0 & A_1^{1012} \\ 0 & 0 & A_1^{113} & 0 & 0 & 0 & 0 & 0 & A_1^{119} & 0 & A_1^{1111} & A_1^{1112} \\ 0 & 0 & 0 & A_1^{124} & 0 & 0 & 0 & 0 & 0 & A_1^{1210} & A_1^{1211} & A_1^{1212} \end{bmatrix},$$

$$A_1^{12} = \begin{bmatrix} 0 \\ I_{S_1+1} \otimes K_{50} \end{bmatrix}, A_1^{13} = I_{S_1+1} \otimes K_{51},$$

$$K_{50} = I_{S_1+1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_2 I_m, K_{51} = \begin{bmatrix} 0 \\ I_{S_2 S_2} \otimes D_1 \end{bmatrix},$$

$$A_1^{21} = I_{S_1+1} \otimes K_{52}, K_{52} = \begin{bmatrix} N_{70} & 0 \\ I_{S_1} \otimes N_{70} & 0 \end{bmatrix}, N_{70} = I_{S_1+1} \otimes U2^0 \otimes I_m,$$

$$A_1^{22} = I_{S_1+1} \otimes K_{53}, K_{53} = \begin{bmatrix} N_{71} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{72} \\ 0 & N_{73} & 0 & \dots & 0 & 0 & \dots & 0 & N_{72} \\ 0 & 0 & N_{73} & \dots & 0 & 0 & \dots & 0 & N_{72} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{73} & 0 & \dots & 0 & N_{72} \\ 0 & 0 & 0 & \dots & 0 & N_{74} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{74} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{74} \end{bmatrix},$$

$$\text{where } N_{71} = I_{S_2} \otimes (U2 \oplus D_0) - \beta I_{n_2}, N_{72} = e_{s_1} \otimes I_{S_2} \otimes I_{n_2} \otimes \beta I_m$$

$$N_{73} = I_{S_2} \otimes (U2 \oplus D_0) - (\beta + \sigma_1) I_{n_2 m}, N_{74} = I_{S_2} \otimes (U_2 + D_0) - \sigma_1 I_{n_2 m},$$

$$A_1^{24} = I_{S_1+1} \otimes K_{54}, A_1^{31} = \begin{bmatrix} K_{55} & K_{56} & 0 & \dots & 0 & 0 \\ 0 & K_{55} & K_{56} & \dots & 0 & 0 \\ 0 & 0 & K_{55} & \dots & 0 & 0 \\ & & & \ddots & & \\ 0 & 0 & 0 & \dots & K_{55} & K_{56} \\ 0 & 0 & 0 & \dots & 0 & K_{55} + K_{56} \end{bmatrix},$$

$$K_{54} = \begin{bmatrix} 0 \\ I_{S_1 S_2} \otimes I_{n_2} \otimes \gamma_1 \otimes D_1 \end{bmatrix}, K_{55} = [I_{S_1 S_2} \otimes f_1 U1^0 \otimes I_m \quad 0],$$

$$K_{56} = [I_{S_1 S_2} \otimes f_2 U1^0 \otimes I_m \quad 0],$$

$$A_1^{33} = \begin{bmatrix} K_{57} & 0 \\ 0 & I_{S_1-1} \otimes K_{58} \end{bmatrix}, K_{57} = \begin{bmatrix} N_{75} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{76} \\ 0 & N_{75} & 0 & \dots & 0 & 0 & \dots & 0 & N_{76} \\ 0 & 0 & N_{75} & \dots & 0 & 0 & \dots & 0 & N_{76} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{75} & 0 & \dots & 0 & N_{76} \\ 0 & 0 & 0 & \dots & 0 & N_{77} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{77} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{77} \end{bmatrix},$$

$$\text{where } N_{75} = I_{S_1} \otimes I_{S_2} \otimes (U1 \oplus D_0) - \beta I_{n_1 m}, N_{76} = e_{s_1} \otimes I_{S_2} \otimes I_{n_1} \otimes \beta I_m$$

$$\begin{aligned}
 N_{77} &= I_{(S_1-s_1)} \otimes I_{S_2} \otimes (U1 \oplus D_0), \\
 K_{58} &= \begin{bmatrix} N_{78} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{79} \\ \mathbf{0} & N_{78} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{79} \\ \mathbf{0} & \mathbf{0} & N_{78} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{79} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{78} & \mathbf{0} & \dots & \mathbf{0} & N_{79} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{80} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{80} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{80} \end{bmatrix}, A_1^{34} = \begin{bmatrix} 0 & & \\ I_{S_1-1} \otimes K_{59} & 0 \end{bmatrix}, \\
 K_{59} &= [I_{S_1 S_2} \otimes I_{n_1} \otimes \gamma_2 \otimes \sigma_2 I_m], \\
 A_1^{42} &= \begin{bmatrix} K_{60} & K_{61} & 0 & \dots & 0 & 0 \\ 0 & K_{60} & K_{61} & \dots & 0 & 0 \\ 0 & 0 & K_{60} & \dots & 0 & 0 \\ 0 & 0 & 0 & \dots & K_{60} & K_{61} \\ 0 & 0 & 0 & \dots & 0 & K_{60} + K_{61} \end{bmatrix}, A_1^{43} = I_{S_1} \otimes K_{62},
 \end{aligned}$$

where

$$\begin{aligned}
 N_{78} &= I_{S_1} \otimes I_{S_2} \otimes (U1 \oplus D_0) - (\beta + \sigma_2) I_{n_1 m}, N_{79} = e_{s_1} \otimes I_{S_2} \otimes I_{n_1} \otimes \beta I_m N_{80} = \\
 &I_{(S_1-s_1)} \otimes I_{S_2} \otimes (U1 \oplus D_0) - \sigma_2 I_{n_1 m}, \\
 K_{60} &= [I_{S_1} \otimes I_{S_2} \otimes f_1 U1^0 \otimes I_{n_2 m} \quad 0], K_{61} = [I_{S_1} \otimes I_{S_2} \otimes f_2 U1^0 \otimes I_{n_2 m} \quad 0], K_{62} = \\
 &\begin{bmatrix} N_{81} & 0 \\ I_{S_1-1} \otimes N_{81} \end{bmatrix}, A_1^{44} = I_{S_1-1} \otimes K_{63}, N_{81} = I_{S_1} \otimes I_{n_1} \otimes U2^0 \otimes I_m,
 \end{aligned}$$

$$K_{63} = \begin{bmatrix} N_{82} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{83} \\ \mathbf{0} & N_{82} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{83} \\ \mathbf{0} & \mathbf{0} & N_{82} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{83} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{82} & \mathbf{0} & \dots & \mathbf{0} & N_{83} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{84} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{84} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{84} \end{bmatrix}, A_1^{56} = \begin{bmatrix} 0 & & \\ I_{S_1} \otimes K_{65} & 0 \end{bmatrix},$$

$$A_1^{53} = I_{M+1} \otimes I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes \delta I_m, K_{65} = I_{S_1 S_2} \otimes \gamma_2 \otimes \sigma_2 I_m, A_1^{64} = I_{M+1} \otimes I_{S_1 S_2} \otimes \gamma_1 \otimes I_{n_2} \otimes \delta I_m,$$

$$A_1^{55} = I_{S_1} \otimes K_{64}, K_{64} = \begin{bmatrix} N_{85} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{86} \\ \mathbf{0} & N_{85} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{86} \\ \mathbf{0} & \mathbf{0} & N_{85} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{86} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{85} & \mathbf{0} & \dots & \mathbf{0} & N_{86} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{87} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{87} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{87} \end{bmatrix}$$

$$\begin{aligned}
 N_{85} &= I_{S_1} \otimes I_{S_2} \otimes D_0 - (\delta + \beta) I_m, N_{86} = e_{s_1} \otimes I_{S_2} \otimes \beta I_m \\
 N_{87} &= I_{(S_1-s_1)} \otimes I_{S_2} \otimes D_0 - \delta I_m, \\
 A_1^{65} &= I_{S_1} \otimes K_{66}, K_{66} = \begin{bmatrix} N_{88} & 0 \\ I_{S_1-1} \otimes N_{88} \end{bmatrix}, \\
 \text{where } N_{88} &= I_{S_1} \otimes U2^0 \otimes I_m, N_{85} = I_{S_1} \otimes I_{S_2} \otimes ((U1 + U2) \oplus D_0) - (\tau + \beta) I_{n_1 n_2 m}, N_{86} = \\
 e_{s_1} \otimes I_{S_2} \otimes I_{n_1 n_2} \otimes \beta I_m, N_{87} &= I_{(S_1-s_1)} \otimes I_{S_2} \otimes ((U1 + U2) \oplus D_0) - \tau I_{n_1 n_2 m}, A_1^{46} = I_{M+1} \otimes \\
 &I_{S_1} \otimes I_{S_2} \otimes e_{n_1} \otimes \tau I_{n_2 m},
 \end{aligned}$$

$$A_1^{66} = I_{S_1} \otimes K_{67}, K_{67} = \begin{bmatrix} N_{89} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{90} \\ \mathbf{0} & N_{89} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{90} \\ \mathbf{0} & \mathbf{0} & N_{89} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{90} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{89} & \mathbf{0} & \dots & \mathbf{0} & N_{90} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{91} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{91} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{91} \end{bmatrix},$$

where $N_{89} = I_{S_1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\delta + \beta)I_{n_2m}$, $N_{90} = e_{s_1} \otimes I_{S_2} \otimes I_{n_2} \otimes \beta I_m$
 $N_{88} = I_{(S_1-S_1)} \otimes I_{S_2} \otimes (U2 \oplus D_0) - \delta I_{n_2m}$,
 $A_1^{71} = I_{M+1} \otimes I_{(S_1+1)S_2} \otimes \phi I_m$,

$$A_1^{77} = I_{S_1+1} \otimes K_{68}, K_{68} = \begin{bmatrix} N_{92} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{93} \\ \mathbf{0} & N_{92} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{93} \\ \mathbf{0} & \mathbf{0} & N_{92} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{93} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{92} & \mathbf{0} & \dots & \mathbf{0} & N_{93} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{94} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{94} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{94} \end{bmatrix},$$

where $N_{92} = I_{S_1+1} \otimes I_{S_2} \otimes D_0 - (\phi + \sigma_2)I_m$, $N_{93} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$
 $N_{94} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes D_0 - (\phi + \sigma_2)I_m$,

$$K_{69} = \begin{bmatrix} N_{95} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{96} \\ \mathbf{0} & N_{95} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{96} \\ \mathbf{0} & \mathbf{0} & N_{95} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{96} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{95} & \mathbf{0} & \dots & \mathbf{0} & N_{96} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{97} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{97} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{97} \end{bmatrix}, A_1^{78} = \begin{bmatrix} \mathbf{0} \\ I_{S_1} \otimes K_{70} \end{bmatrix},$$

where $N_{95} = I_{S_1+1} \otimes I_{S_2} \otimes D_0 - (\phi + \beta + \sigma_2)I_m$, $N_{96} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_m$
 $N_{97} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes D_0 - (\phi + \sigma_2)I_m$, $K_{70} = I_{S_1+1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_2 I_m$, $A_1^{82} = I_{M+1} \otimes I_{(S_1+1)S_2n_2} \otimes \phi I_m$,

$$A_1^{87} = I_{S_1+1} \otimes K_{71}, K_{71} = \begin{bmatrix} N_{98} & \mathbf{0} \\ I_{S_1} \otimes N_{98} & \mathbf{0} \end{bmatrix}, N_{98} = I_{S_1+1} \otimes U2^0 \otimes I_m,$$

$$A_1^{88} = I_{S_1} \otimes K_{72}, K_{72} = \begin{bmatrix} N_{99} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{100} \\ \mathbf{0} & N_{99} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{100} \\ \mathbf{0} & \mathbf{0} & N_{99} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{100} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & N_{99} & \mathbf{0} & \dots & \mathbf{0} & N_{100} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{101} & \dots & \mathbf{0} & \mathbf{0} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & N_{101} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & \mathbf{0} & \dots & \mathbf{0} & N_{101} \end{bmatrix},$$

where $N_{99} = I_{S_1+1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\phi + \beta)I_{n_2m}$, $N_{100} = e_{s_1+1} \otimes I_{S_2} \otimes \beta I_{n_2m}$, $N_{101} = I_{(S_1+1)-(s_1+1)} \otimes I_{S_2} \otimes (U2 \oplus D_0) - \phi I_{n_2m}$, $A_1^{91} = I_{M+1} \otimes I_{S_1+1}S_2 \otimes \eta I_m$,

$$A_1^{99} = \begin{bmatrix} K_{73} & 0 \\ 0 & I_{S_1-1} \otimes K_{74} \end{bmatrix}, K_{73} = \begin{bmatrix} N_{102} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{103} \\ 0 & N_{104} & 0 & \dots & 0 & 0 & \dots & 0 & N_{103} \\ 0 & 0 & N_{104} & \dots & 0 & 0 & \dots & 0 & N_{103} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{104} & 0 & \dots & 0 & N_{103} \\ 0 & 0 & 0 & \dots & 0 & N_{105} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{105} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{105} \end{bmatrix},$$

where $N_{102} = I_{S_1+1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\phi + \beta)I_{n_2m}$, $N_{103} = e_{S_1} \otimes I_{S_2} \otimes \beta I_m$

$N_{104} = I_{S_1} \otimes I_{S_2} \otimes D_0 - (\eta + \beta + \sigma_1)I_m$, $N_{105} = I_{(S_1+1)-(S_1+1)} \otimes I_{S_2} \otimes D_0 - (\eta + \sigma_1)I_m$,

$$K_{74} = \begin{bmatrix} N_{106} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{107} \\ 0 & N_{108} & 0 & \dots & 0 & 0 & \dots & 0 & N_{107} \\ 0 & 0 & N_{108} & \dots & 0 & 0 & \dots & 0 & N_{107} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{108} & 0 & \dots & 0 & N_{107} \\ 0 & 0 & 0 & \dots & 0 & N_{109} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{109} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{109} \end{bmatrix}, A_1^{910} = \begin{bmatrix} 0 & & \\ I_{S_1-1} \otimes K_{75} & & 0 \end{bmatrix},$$

where $N_{106} = I_{S_2} \otimes D_0 - (\eta + \beta + \sigma_2)I_m$, $N_{107} = e_{S_1} \otimes I_{S_2} \otimes \beta I_m$

$N_{108} = I_{S_1} \otimes I_{S_2} \otimes D_0 - (\eta + \beta + \sigma_1)I_m$, $N_{109} = I_{(S_1+1)-(S_1+1)} \otimes I_{S_2} \otimes D_0 - (\eta + \sigma_1)I_m$,

$K_{75} = I_{S_1} \otimes I_{(S_1+1)S_2} \otimes \gamma_2 \otimes \sigma_2 I_m$,

$A_1^{911} = I_{S_1} \otimes K_{76}$, $A_1^{109} = I_{S_1} \otimes K_{77}$, $A_1^{102} = I_{M+1} \otimes I_{(S_1+1)S_2} \otimes I_{n_2} \otimes \eta I_m$,

$$K_{76} = \begin{bmatrix} 0 \\ I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes D_1 \end{bmatrix}, K_{77} = \begin{bmatrix} N_{110} & 0 \\ I_{S_1-1} \otimes N_{110} & 0 \end{bmatrix},$$

$$A_1^{1010} = I_{S_1-1} \otimes K_{78}, K_{78} = \begin{bmatrix} N_{111} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{112} \\ 0 & N_{113} & 0 & \dots & 0 & 0 & \dots & 0 & N_{112} \\ 0 & 0 & N_{113} & \dots & 0 & 0 & \dots & 0 & N_{112} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{113} & 0 & \dots & 0 & N_{112} \\ 0 & 0 & 0 & \dots & 0 & N_{114} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{114} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{114} \end{bmatrix},$$

where $N_{110} = I_{(S_1+1)} \otimes U2^0 \otimes I_m$,

$N_{111} = I_{S_2} \otimes (U2 \oplus D_0) - (\eta + \beta)I_{n_2m}$, $N_{112} = e_{S_1} \otimes I_{S_2} \otimes \beta I_{n_2m}$

$N_{113} = I_{S_1} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\eta + \beta + \sigma_1)I_{n_2m}$, $N_{114} = I_{(S_1+1)-(S_1+1)} \otimes I_{S_2} \otimes (U2 \oplus D_0) - (\eta + \sigma_1)I_{n_2m}$,

$A_1^{1012} = I_{S_1} \otimes K_{79}$, $K_{79} = \begin{bmatrix} 0 \\ I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes I_{n_2} \otimes D_1 \end{bmatrix}$,

$$A_1^{1109} = \begin{bmatrix} K_{80} & K_{81} & 0 & \dots & 0 & 0 \\ 0 & K_{80} & K_{81} & \dots & 0 & 0 \\ 0 & 0 & K_{80} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & K_{80} & K_{81} \\ 0 & 0 & 0 & \dots & 0 & K_{80} + K_{81} \end{bmatrix}, A_1^{1111} = \begin{bmatrix} K_{82} & 0 \\ 0 & I_{S_1-1} \otimes K_{83} \end{bmatrix},$$

$K_{80} = [I_{S_1S_2} \otimes f_1 \theta U1^0 \otimes I_m \ 0]$, $K_{81} = [I_{S_1S_2} \otimes f_2 \theta U1^0 \otimes I_m \ 0]$,

$$K_{82} = \begin{bmatrix} N_{115} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{116} \\ 0 & N_{115} & 0 & \dots & 0 & 0 & \dots & 0 & N_{116} \\ 0 & 0 & N_{115} & \dots & 0 & 0 & \dots & 0 & N_{116} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{115} & 0 & \dots & 0 & N_{116} \\ 0 & 0 & 0 & \dots & 0 & N_{117} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{117} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{117} \end{bmatrix}, K_{83} = \begin{bmatrix} N_{118} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{119} \\ 0 & N_{118} & 0 & \dots & 0 & 0 & \dots & 0 & N_{119} \\ 0 & 0 & N_{118} & \dots & 0 & 0 & \dots & 0 & N_{119} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{118} & 0 & \dots & 0 & N_{119} \\ 0 & 0 & 0 & \dots & 0 & N_{120} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{120} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{120} \end{bmatrix},$$

where $N_{115} = I_{s_1} \otimes I_{s_2} \otimes (\theta U1 \oplus D_0) - (\eta + \beta)I_{n_1 m}$, $N_{116} = e_{s_1} \otimes I_{s_2} \otimes \beta I_{n_1 m}$,
 $N_{117} = I_{(s_1-s_1)} \otimes I_{s_2} \otimes (\theta U1 \oplus D_0) - \eta I_{n_1 m}$, $N_{118} = I_{s_1} \otimes I_{s_2} \otimes (\theta U1 \oplus D_0) - (\eta + \beta + \sigma_2)I_{n_1 m}$,

$N_{119} = e_{s_1} \otimes I_{s_2} \otimes \beta I_{n_1 m}$, $N_{120} = I_{s_1-s_1} \otimes I_{s_2} \otimes (\theta U1 \oplus D_0) - (\eta + \beta + \sigma_2)I_{n_1 m}$,

$A_{124} = I_{M+1} \otimes I_{s_1 s_2} \otimes I_{n_2} \otimes e_{n_1} \otimes \gamma_1 \otimes \eta I_m$

, $A_1^{1112} = \begin{bmatrix} 0 \\ I_{s_1-1} \otimes K_{84} & 0 \end{bmatrix}$, $K_{84} = I_{s_1 s_2} \otimes I_{n_1} \otimes \gamma_2 \otimes \sigma_2 I_{n_1 m}$,

$A_1^{1210} = \begin{bmatrix} K_{85} & K_{86} & 0 & \dots & 0 & 0 \\ 0 & K_{85} & K_{86} & \dots & 0 & 0 \\ 0 & 0 & K_{85} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & K_{85} & K_{86} \\ 0 & 0 & 0 & \dots & 0 & K_{85} + K_{86} \end{bmatrix}$, $A_1^{1211} = I_{s_1} \otimes K_{87}$,

$K_{85} = [I_{s_1 s_2} \otimes f_1 \theta U1^0 \otimes I_{n_2 m} \quad 0]$, $K_{86} = [I_{s_1 s_2} \otimes f_2 \theta U1^0 \otimes I_{n_2 m} \quad 0]$,

$K_{87} = \begin{bmatrix} N_{121} & 0 \\ I_{s_1-1} \otimes N_{121} & 0 \end{bmatrix}$, $N_{121} = I_{s_1} \otimes U2^0 \otimes I_{n_1 m}$

,

$A_1^{1212} = I_{s_1} \otimes K_{88}$, $K_{88} = \begin{bmatrix} N_{122} & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{123} \\ 0 & N_{122} & 0 & \dots & 0 & 0 & \dots & 0 & N_{123} \\ 0 & 0 & N_{122} & \dots & 0 & 0 & \dots & 0 & N_{123} \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & N_{122} & 0 & \dots & 0 & N_{123} \\ 0 & 0 & 0 & \dots & 0 & N_{124} & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & N_{124} & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & \dots & 0 & N_{124} \end{bmatrix}$,

$N_{122} = I_{s_1} \otimes I_{s_2} \otimes (\theta U1 \otimes D_0) - (\eta + \beta)I_{n_1 n_2 m}$,

$N_{124} = I_{(s_1-s_1)} \otimes I_{s_2} \otimes ((U1 + U2) \oplus N_0) - \eta I_{n_1 n_2 m}$.

4. ANALYSIS OF STABILITY CONDITION

A generator matrix is indicated by the equation $A = A_0 + A_1 + A_2$. In the steady state, let ζ be the probability vector of A that satisfies $\zeta e = 1$ and $\zeta A = 0$. Where $\zeta_0, \zeta_6, \zeta_8$, has a dimension of $(M+1)(S_1+1)S_2 m$, $\zeta_1, \zeta_7, \zeta_9$, has a dimension of $(M+1)(S_1+1)S_2 n_2 m$, ζ_2 , has a

dimension of $(M+1)S_1S_2n_1m$, $\zeta_3, \zeta_{10}, \zeta_{11}$, has a dimension of $(M+1)S_1S_2n_1n_2m$, ζ_4 has a dimension of $(M+1)S_1S_2m$, ζ_5 has a dimension of $(M+1)S_1S_2n_2m$.

$$A = \begin{bmatrix} A^{11} & A^{12} & A^{13} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ A^{21} & A^{22} & 0 & 0 & A^{24} & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ A^{31} & 0 & A^{33} & A^{34} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & A^{42} & A^{43} & A^{44} & 0 & A^{46} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & A^{55} & A^{56} & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & A^{64} & A^{65} & A^{66} & 0 & 0 & 0 & 0 & 0 \\ A^{71} & 0 & 0 & 0 & 0 & 0 & A^{77} & A^{78} & 0 & 0 & 0 & 0 \\ 0 & A^{82} & 0 & 0 & 0 & 0 & A^{87} & A^{88} & 0 & 0 & 0 & 0 \\ A^{91} & 0 & 0 & 0 & 0 & 0 & 0 & 0 & A^{99} & A^{910} & A^{911} & 0 \\ 0 & A^{102} & 0 & 0 & 0 & 0 & 0 & 0 & A^{109} & A^{1010} & 0 & A^{1012} \\ 0 & 0 & A^{113} & 0 & 0 & 0 & 0 & 0 & A^{119} & 0 & A^{1111} & A^{1112} \\ 0 & 0 & 0 & A^{124} & 0 & 0 & 0 & 0 & A^{1210} & A^{1211} & 0 & A^{1212} \end{bmatrix},$$

where, $A^{11} = [A_0^{11} + A_1^{11}]$, $A^{12} = A_1^{12}$, $A^{13} = [A_1^{13} + A_2^{13}]$, $A^{21} = A_1^{21}$,
 $A^{22} = [A_0^{22} + A_1^{22}]$, $A^{24} = [A_1^{24} + A_2^{24}]$, $A^{31} = A_1^{31}$, $A^{33} = [A_0^{33} + A_1^{33}]$,
 $A^{34} = A_1^{34}$, $A^{42} = A_1^{42}$, $A^{43} = A_1^{43}$, $A^{44} = A_0^{44}$, $A^{46} = A_1^{46}$, $A^{55} = [A_0^{55} + A_1^{55}]$,
 $A^{64} = A_1^{64}$, $A^{65} = A_1^{65}$, $A^{66} = [A_0^{66} + A_1^{66}]$, $A^{71} = A_1^{71}$, $A^{77} = [A_0^{77} + A_1^{77}]$, $A^{78} = A_1^{78}$,
 $A^{82} = A_1^{82}$, $A^{87} = A_1^{87}$, $A^{88} = [A_0^{88} + A_1^{88}]$, $A^{91} = A_1^{91}$, $A^{99} = [A_0^{99} + A_1^{99}]$, $A^{910} = A_1^{910}$,
 $A^{911} = [A_1^{911} + A_2^{911}]$, $A^{102} = A_1^{102}$, $A^{109} = A_1^{109}$, $A^{1010} = [A_0^{1010} + A_1^{1010}]$,
 $A^{1012} = [A_0^{1012} + A_1^{1012}]$, $A^{113} = A_1^{113}$, $A^{119} = A_1^{119}$, $A^{1111} = [A_0^{1111} + A_1^{1111}]$,
 $A^{1112} = A_1^{1112}$, $A^{124} = A_1^{124}$, $A^{1210} = A_1^{1210}$, $A^{1211} = A_1^{1211}$, $A^{1212} = [A_0^{1212} + A_1^{1212}]$,

The steady state probability vector ζ is obtained by solving the below equations:

$$\begin{aligned} \zeta_0[A^{11}] + \zeta_1[A^{21}] + \zeta_2[A^{31}] + \zeta_6[A^{71}] + \zeta_8[A^{91}] &= 0, \\ \zeta_0[A^{12}] + \zeta_1[A^{22}] + \zeta_3[A^{42}] + \zeta_7[A^{82}] + \zeta_9[A^{102}] &= 0, \\ \zeta_0[A^{13}] + \zeta_2[A^{33}] + \zeta_3[A^{43}] + \zeta_{10}[A^{113}] &= 0, \\ \zeta_1[A^{24}] + \zeta_2[A^{34}] + \zeta_3[A^{44}] + \zeta_5[A^{64}] + \zeta_{11}[A^{124}] &= 0, \\ \zeta_4[A^{55}] + \zeta_5[A^{65}] &= 0, \\ \zeta_3[A^{46}] + \zeta_4[A^{56}] + \zeta_5[A^{66}] &= 0, \\ \zeta_7[A^{77}] + \zeta_8[A^{87}] &= 0, \\ \zeta_7[A^{78}] + \zeta_8[A^{88}] &= 0, \\ \zeta_8[A^{99}] + \zeta_9[A^{109}] + \zeta_{10}[A^{119}] &= 0, \\ \zeta_8[A^{910}] + \zeta_9[A^{1010}] + \zeta_{11}[A^{1210}] &= 0, \\ \zeta_8[A^{911}] + \zeta_{10}[A^{1111}] + \zeta_{11}[A^{1211}] &= 0, \\ \zeta_9[A^{1012}] + \zeta_{10}[A^{1112}] + \zeta_{11}[A^{1212}] &= 0, \end{aligned}$$

Subject to normalizing condition

$$\begin{aligned} \sum_{d_2=0}^M \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} [\zeta_0 e_m + \zeta_1 e_{n_2 m} + \zeta_2 e_{n_1 m} + \zeta_3 e_{n_1 n_2 m} + \zeta_4 e_m + \zeta_5 e_{n_2 m} \\ + \zeta_6 e_{n_2 m} + \zeta_7 e_{n_2 m} + \zeta_8 e_m + \zeta_9 e_{n_2 m} + \zeta_{10} e_{n_1 n_2 m} + \zeta_{11} e_{n_1 n_2 m}] = 1 \end{aligned}$$

Within the framework of QBD, our model's stability should meet the necessary and sufficient conditions, (i.e.),

$$\zeta A_0 e < \zeta A_2 e \#(3)$$

which is found to be of the below form,

$$\begin{aligned}
& \sum_{d_2=0}^M \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} [\zeta_0 d_2 d_3 d_6 [I_{S_2} \otimes D_1 e_m] + \zeta_1 d_2 d_3 d_6 [I_{S_2 n_2} \otimes D_1 e_{n_2 m}]] \\
& + \zeta_2 d_2 d_3 d_6 [I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1} \otimes D_1 e_{n_1 m}] + \zeta_3 d_2 d_3 d_6 [I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1 n_2} \otimes D_1 e_{n_1 n_2 m}] \\
& + \zeta_4 d_2 d_3 d_6 [I_{M+1} \otimes I_{S_1 S_2} \otimes D_1 e_m] + \zeta_5 d_2 d_3 d_6 [I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_1 n_2} \otimes D_1 e_{n_2 m}] \\
& + \zeta_6 d_2 d_3 d_6 [I_{M+1} \otimes I_{(S_1+1)S_2} \otimes D_1 e_m] + \zeta_7 d_2 d_3 d_6 [I_{M+1} \otimes I_{(S_1+1)S_2} \otimes I_{n_2} \otimes D_1 e_{n_2 m}] \\
& + \zeta_8 d_2 d_3 d_6 [I_{M+1} \otimes I_{(S_1+1)S_2} \otimes D_1 e_m] + \zeta_9 [I_{M+1} \otimes I_{(S_1+1)S_2} \otimes I_{n_2} \otimes D_1 e_{n_2 m}] \\
& + \zeta_{10} d_2 d_3 d_6 [I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_2} \otimes D_1 e_{n_1 n_2 m}] + \zeta_{11} d_2 d_3 d_6 [I_{M+1} \otimes I_{S_1 S_2} \otimes I_{n_2 n_2} \otimes D_1 e_{n_1 n_2 m}]] \\
& < \\
& \sum_{d_2=0}^M \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} [\zeta_0 d_2 d_3 d_6 [I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes \sigma_1 I_m e_m] + \zeta_1 d_2 d_3 d_6 [I_{S_1} \otimes I_{S_2} \otimes \gamma_2 \otimes \sigma_1 \otimes I_m e_{n_2 m}]] \\
& + \zeta_9 d_2 d_3 d_6 [I_{S_1 S_2} \otimes I_{n_1} \otimes \gamma_2 \otimes \sigma_1 I_m e_{n_2 m}] + \zeta_{10} d_2 d_3 d_6 [I_{S_1} \otimes I_{S_2} \otimes \gamma_1 \otimes \sigma_1 I_m e_{n_1 n_2 m}]
\end{aligned}$$

4.1. Invariant Probability Vector Analysis

Let v symbolize the infinitesimal generator Q 's invariant probability vector and this is divided into $v = (v_0, v_1, v_2, \dots)$. Mention that $v_0, v_1, v_2, \dots$ have a dimension of $3[(M+1)(S_1+1)[S_2 m + S_2 n_2 m] + 3(M+1)S_1 S_2 n_1 n_2 m + (M+1)S_1 S_2 n_1 m + (M+1)S_1 S_2 m + (M+1)S_1 S_2 n_2 m]$ and the vector v satisfies

$$vQ = 0 \text{ and } ve = 1. \#(4)$$

Additionally, after the stability needed of the model is met, the stationary probability vector v can be obtained by applying the following equation:

$$v_i = v_0 R^i, i \geq 1, \#(5)$$

Second-degree matrix equation

$$R^2 A_2 + R A_1 + A_0 = 0 \#(6)$$

is met by the [26] least non-negative solution R . The rate matrix is obtained from the matrix quadratic equation. The requirement is satisfied by the rate matrix R , whose order is provided by $3[(M+1)(S_1+1)[S_2 m + S_2 n_2 m] + 3(M+1)S_1 S_2 n_1 n_2 m + (M+1)S_1 S_2 n_1 m + (M+1)S_1 S_2 m + (M+1)S_1 S_2 n_2 m]$.

$$R A_2 e = A_0 e \#(7)$$

the sub vector v_0 can be found by solving the subsequent equation

$$v_0 (B_{00} + R A_2) = 0 \#(8)$$

the normalizing condition is subject to

$$v_0 e_{3[(M+1)(S_1+1)[S_2 m + S_2 n_2 m] + 3(M+1)S_1 S_2 n_1 n_2 m + (M+1)S_1 S_2 n_1 m + (M+1)S_1 S_2 m + (M+1)S_1 S_2 n_2 m]} = 1.0.0 \#(9)$$

Before solving the set of equations mentioned above, we first need to calculate the rate matrix R . This can be done using the Logarithmic Reduction Approach [27], which simplifies the process of finding R .

5. ANALYSIS OF ACTIVE PERIOD

- The time frame that begins when users enter the empty system and concludes when the system size first falls to zero is known as the busy period. Therefore, the first phase of the shift from level 1 to level 0 is the analogue of the busy period. When a state at any

other level is visited at least once, the busy cycle is the time it takes to return to level zero.

- Let's use the concept of the basic time to examine the busy time. The QBD process is nothing more than the initial transition period between the levels k_1 to $k_1 - 1$, where $k_1 \geq 1$.
- The boundary states, or the circumstances when $k_1 = 0$, require a distinct discussion. For any level k_1 where $k_1 \geq 1$, it is evident that there are $3[(M + 1)(S_1 + 1)[S_2m + S_2n_2m] + 3(M + 1)S_1S_2n_1n_2m + (M + 1)S_1S_2n_1m + (M + 1)S_1S_2m + (M + 1)S_1S_2n_2m]$ states. For this reason, the h^{th} state of level k_1 may be denoted as (k_1, h) when each state is placed in lexicographic order.
- Let $H_{hh'}(k_1, x)$ represent the probability with the condition that, subject to the restriction that it begin in the state (k_1, h) at time $t = 0$, the QBD process visits level $k_1 - 1$ to modify k_1 transitions to the left and also enters the state (k_1, h') . Let us present the idea of the joint transform.

$$\bar{H}_{hh'}(z, s) = \sum_{k_1=1}^{\infty} z_1^{k_1} \int_0^{\infty} e^{-sx} dH_{hh'}(k_1, x); |z| \leq 1, \text{Re}(s) \geq 0 \quad (10)$$

and the matrix is denoted as $\bar{H}(z, s) = \bar{H}_{hh'}(z, s)$. The following equations can be easily satisfied by the matrix $\bar{H}(z, s)$:

$$\bar{H}(z, s) = z[sI - A_1]^{-1}A_2 + [sI - A_1]^{-1}A_0\bar{H}^2(z, s) \quad (11)$$

- Without taking into account the boundary states, the initial travel time should be computed with the matrix $H = H_{hh'} = \bar{H}(1, 0)$ after evaluating the rate matrix R , we can quickly get the matrix H by applying the result

$$H = -[A_1 + RA_2]^{-1}A_2 \quad (12)$$

- In the absence of such, it is possible to calculate the H matrix values by applying the idea of a logarithmic reduction process. $\bar{H}^{0,0}(1, 0)$ provides the following equation, which is related to boundary level zero.

$$\bar{H}^{0,0}(z, s) = [sI - H_{00}]^{-1}A_0\bar{H}(z, s). \quad (13)$$

- The next moments are simple to evaluate because the three matrices namely, $H, \bar{H}^{(1,0)}$ and $\bar{H}^{(0,0)}(1, 0)$ are stochastic. At $s = 0, z = 1$,

$$\vec{J}_1 = -\frac{\partial}{\partial s} \bar{H}(z, s) = -[A_1 + A_0(H + I)]^{-1}e \quad (14)$$

$$\vec{J}_2 = \frac{\partial}{\partial s} \bar{H}(z, s) = -[A_1 + A_0(H + I)]^{-1}A_2e \quad (15)$$

$$\vec{J}_1^{(0,0)} = -\frac{\partial}{\partial z} \bar{H}^{(0,0)}(z, s) = -H_{00}^{-1}[e + A_0\vec{J}_1] \quad (16)$$

$$\vec{J}_2^{(0,0)} = \frac{\partial}{\partial z} \bar{H}^{(0,0)}(z, s) = -H_{00}^{-1}[A_0\vec{J}_2] \quad (17)$$

6. PERFORMANCE MEASURES

We examine our model's qualitative behavior under steady conditions. Let $x(d_1, d_2, d_3, d_4, d_5, d_6, d_7, d_8, d_9)$ stand for the probability of the stable state, where the quantity of clients in the orbit I is d_1 , the quantity of customers in the orbit II is d_2 , server status I is d_3 , server status II is d_4, d_5, d_6, d_7, d_8 and d_9 represent the number of commodity I, commodity II, service phase I, service phase II and arrival phases, respectively.

- Probability that server-1 is idle and server-2 is idle:

$$P_{S1IS2I} = \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{0000d_5d_6d_9}$$

- Probability that server-1 is idle and server-2 is busy with re-service:

$$P_{S1IS2BR} = \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{0d_201d_5d_6d_8d_9}$$

- Chance of the server that server- 1 is in close down and server- 2 is in idle:

$$P_{S1CS2I} = \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{0030d_5d_6d_9}$$

- chance of the server- 1 is in close down and server- 2 is busy with re-service:

$$P_{S1CS2BR} = \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{0d_231d_5d_6d_8d_9}$$

- Probability that server-1 is idle in working vacation and server-2 is idle:

$$P_{S1IWVS2I} = \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{0040d_6d_5d_9}$$

- Probability that server-1 is idle in working vacation and server-2 is busy with re-service:

$$P_{S1IWVS2BR} = \sum_{d_2=1}^M \sum_{d_5=0}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{0d_241d_5d_6d_8d_9}$$

- Probability that server- 1 is busy in normal service and server- 2 is idle:

$$P_{S1BNSS2I} = \sum_{d_1=1}^{\infty} \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_7=1}^{n_1} \sum_{d_9=1}^m x_{d_1010d_5d_6d_7d_9}$$

- Probability that server-1 is busy in normal service and server-2 is busy in re-service:

$$P_{S1BNSS2BR} = \sum_{d_1=1}^{\infty} \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_7=1}^{n_1} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{d_1d_211d_5d_6d_7d_8d_9}$$

- Probability that server-1 is in repair process and server-2 is in idle:

$$P_{S1RS2I} = \sum_{d_1=1}^{\infty} \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{d_1020d_5d_6d_9d_{10}}$$

- Probability that server- 1 is in repair process and server- 2 is busy with re-service:

$$P_{S1RS2BR} = \sum_{d_1=1}^{\infty} \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{d_1d_221d_5d_6d_8d_9}$$

- Probability that server- 1 is busy in working vacation and server- 2 is in idle:

$$P_{S1BWVS2I} = \sum_{d_1=1}^{\infty} \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_7=1}^{n_1} \sum_{d_9=1}^m x_{d_1050d_5d_6d_7d_9}$$

- Probability that server-1 is busy in working vacation and server-2 is busy with reservice:

$$P_{S1BWVS2R} = \sum_{d_1=1}^{\infty} \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_7=1}^{n_1} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{d_1 d_2 51 d_5 d_6 d_7 d_8 d_9}$$

- Probability that both the servers are in busy mode:

$$P_{S1BNMS2R} = \sum_{d_1=1}^{\infty} \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_7=1}^{n_1} \sum_{d_8=1}^{n_2} \sum_{d_9=1}^m x_{d_1 d_2 11 d_5 d_6 d_7 d_8 d_9}$$

- Expected number of customers in orbit I:

$$E[N_1] = \sum_{d_1=0}^{\infty} d_1 x_{d_1} e$$

- Expected number of customers in orbit II:

$$E[N_2] = \sum_{d_1=1}^{\infty} \sum_{d_2=1}^M \sum_{d_3=0}^5 \sum_{d_4=0}^1 d_2 x_{d_1 d_2 d_3 d_4} e$$

- Expected number of customers in the system:

$$E_{\text{system}} = E[N_1] + E[N_2] + P_{\text{busy}}$$

- Expected retrial rate:

$$\begin{aligned} E_{RR} &= (\sigma_1 + \sigma_2) \sum_{d_1=1}^{\infty} \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{d_1 d_2 00 d_5 d_6 d_9} \\ &+ \sigma_1 \sum_{d_1=1}^{\infty} \sum_{d_2=0}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{d_1 d_2 41 d_5 d_6 d_9} \\ &+ \sigma_2 \sum_{d_1=0}^{\infty} \sum_{d_2=1}^M \sum_{d_5=1}^{S_1} \sum_{d_6=1}^{S_2} \sum_{d_9=1}^m x_{d_1 d_2 10 d_5 d_6 d_9} \end{aligned}$$

- Probability that server-1 is idle:

$$P_{\text{idle}} = P_{1NM} + P_{IWV}$$

- Probability that server-1 is busy:

$$P_{\text{busy}} = P_{BNM} + P_{BWV}$$

7. NUMERICAL IMPLEMENTATION

In this section, we examine the output of our system using graphical and numerical representations. With a mean value of 1, these arrival processes $ERL - A$, $EXP - A$ and $HEX - A$ have zero correlation and it correspond to the renewal process. Erlang Arrival (ERL-A):

$$D_0 = \begin{bmatrix} -2 & 2 \\ 0 & -2 \end{bmatrix}, D_1 = \begin{bmatrix} 0 & 0 \\ 2 & 0 \end{bmatrix}$$

Exponential Arrival(EXP-A):

$$D_0 = [-1], D_1 = [1]$$

Hyper exponential Arrival(HEX-A):

$$D_0 = \begin{bmatrix} -1.90 & 0 \\ 0 & -0.19 \end{bmatrix}, D_1 = \begin{bmatrix} 1.710 & 0.190 \\ 0.171 & 0.019 \end{bmatrix}$$

Consider the following PH distributions for the service progression:
Erlang Service(ERL-S):

$$\gamma_1 = \gamma_2 = [1,0]; U_1 = U_2 = \begin{bmatrix} -2 & 2 \\ 0 & -2 \end{bmatrix}$$

Exponential Service(EXP-S):

$$\gamma_1 = \gamma_2 = 1; U_1 = U_2 = [-1]$$

Hyper Exponential Service(HEX-S):

$$\gamma_1 = \gamma_2 = [0.8,0.2]; U_1 = U_2 = \begin{bmatrix} -2.80 & 0 \\ 0 & -0.28 \end{bmatrix}$$

7.1. Illustration 1

We investigate how the retrial rate (σ_1) affects the ($E_{\text{Orbit } I}$) in tables 1 to 3. Fix $\lambda = 1$, $\mu_1 = 8, \mu_2 = 11, \sigma_2 = 15, \tau = 3, \delta = 1, \beta = 2, M = 2, p = 0.5, q = 0.5, f_1 = 0.5, f_2 = 0.5, S_2 = 5, S_1 = 6, s_1 = 3, \eta = 2$ and $\theta = 0.8$ such that the system remains stable. The following analysis is based on the observations from Table 1 to Table 3 : When we increase the retrial rate while accounting for both arrival and service periods, the $E_{\text{Orbit } I}$ decreases. From the perspective of service and arrival periods, $E_{\text{Orbit } I}$ decreases significantly for HEX-S and gradually for ERL - S when the retrial rate increases; in contrast to EXP-S, $E_{\text{Orbit } I}$ decreases significantly in ERL - S and slowly in HEX-S when retrial times are taken into consideration.

Table 1: Retrial rate (σ_1) vs $E_{\text{Orbit } I}$ -ERL-A

σ_1	ERL-S	EXP-S	HEX-S
13	1.750777248	1.739976801	1.731201809
14	1.746243381	1.735298055	1.726225649
15	1.742263775	1.731198637	1.721877722
16	1.738738655	1.727574398	1.718044536
17	1.735591201	1.724344979	1.714638486
18	1.732761275	1.721447368	1.711590862
19	1.730201106	1.718831443	1.708847012
20	1.727872243	1.716456823	1.70636292
21	1.72574337	1.714290601	1.704102731
22	1.723788709	1.712305676	1.702036943

Table 2: Retrial rate (σ_1) vs $E_{\text{Orbit } I}$ -EXP-A

σ_1	ERL-S	EXP-S	HEX-S
13	1.798046389	1.787136458	1.784617201
14	1.792961497	1.781868423	1.778952472
15	1.788508844	1.777264719	1.774013311
16	1.78457188	1.773203235	1.769671376

17	1.781061513	1.769592361	1.76582149
18	1.777908448	1.766353113	1.762383101
19	1.775057948	1.763433734	1.759292408
20	1.772466178	1.760785903	1.756498218
21	1.770097598	1.758372024	1.753958974
22	1.767923067	1.756161254	1.751640592

Table 3: Retrial rate(σ_1) vs E_{OrbitI} - HEX-A

σ_1	ERL - S	EXP - S	HEX - S
13	2.092643438	2.06382038	2.040366951
14	2.085895278	2.056691266	2.032341245
14	2.079974135	2.050458089	2.025356525
15	2.074724232	2.044952872	2.019217965
16	2.070027656	2.040047824	2.013776693
17	2.065793473	2.035643975	2.008917002
18	2.065793473	2.035643975	2.008917002
19	2.06195037	2.031663454	2.004547612
20	2.058441583	2.02804412	2.000595542
21	2.055221318	2.024735753	1.997001741

7.2. Illustration 2

The effects of the vacation rate (η) in relation to P_{busy} have been examined in figures 2 to 4 . Fix $\lambda = 1, \mu_1 = 8, \mu_2 = 11, \sigma_2 = 15, \tau = 3, \delta = 1, \beta = 2, M = 2, p = 0.5, q = 0.5, f_1 = 0.5, f_2 = 0.5, S_2 = 5, S_1 = 6, s_1 = 3, \sigma_1 = 13$ and $\theta = 0.8$ such that the system remains stable. We note that from Figure 2 to Figure 4, as follows:

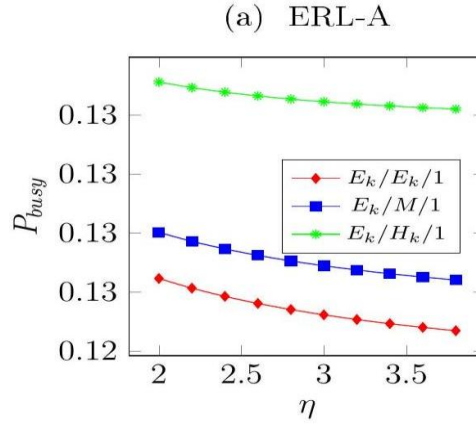
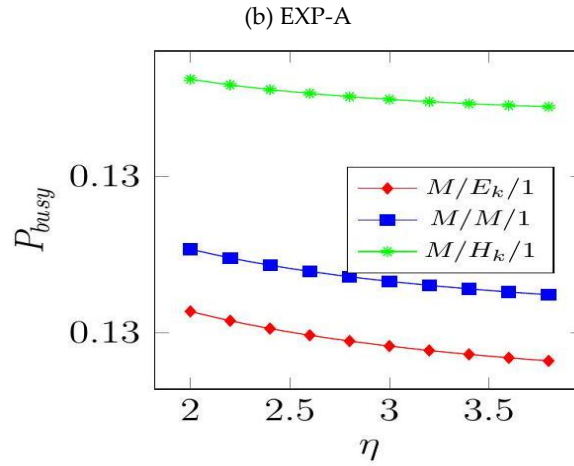
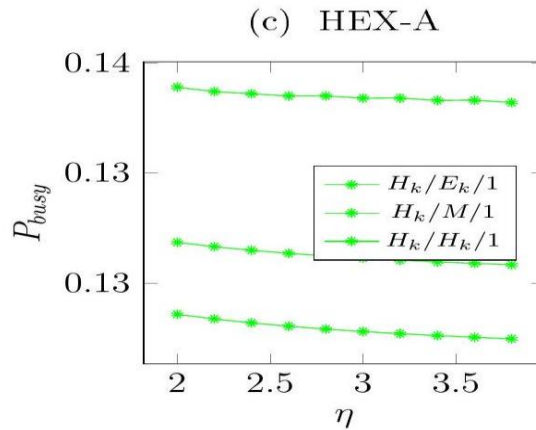
For every combination of arrival-time and service-time distributions, it is found that the chance that the server is busy, P_{busy} , decreases as the vacation rate rises. Because a larger vacation rate results in the server spending more time away from active service, which lowers the proportion of time it remains busy, the phenomenon is constant.

7.3. Illustration 3

To examine the effect on the E_{system} in the corresponding Figures 5 to 13 . Fix $\lambda = 1, \mu_2 = 11, \sigma_2 = 15, \delta = 1, \beta = 2, M = 2, p = 0.5, q = 0.5, f_1 = 0.5, f_2 = 0.5, S_2 = 5, S_1 = 6, s_1 = 3, \sigma_1 = 13$ and $\theta = 0.8$ such that the system remains stable.

The decrease in E_{system} is particularly noticeable in the $HEX - S$, when server 1's service rate rises. In comparison, the $ERL - S$ (Erlang service-time) instance shows a much slower decline. This shows that when service capacity increases, the system becomes much more efficient under $HEX - S$, although the benefit is less noticeable when service times have an Erlang distribution.

Analyzing the impact of higher repair rates under various arrival distributions reveals that $ERL - A$ (Erlang arrival time) only slightly decreases. However, for $EXP - A$ (exponential arrival time), the decline is much more pronounced. This implies that systems with exponential arrivals gain from increases in repair efficiency more significantly than systems with Erlang arrivals.

Figure 2: Vacation rate vs P_{busy} .Figure 3: Vacation rate vs P_{busy} .Figure 4: Vacation rate vs P_{busy} .

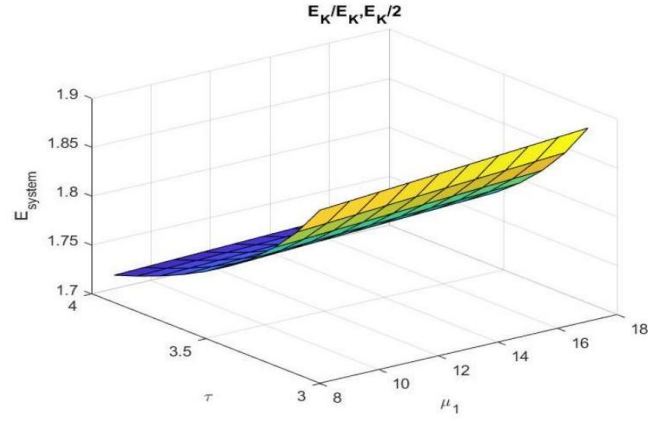


Figure 5: Server-1 service rate and Repair Rate vs E_{system} .

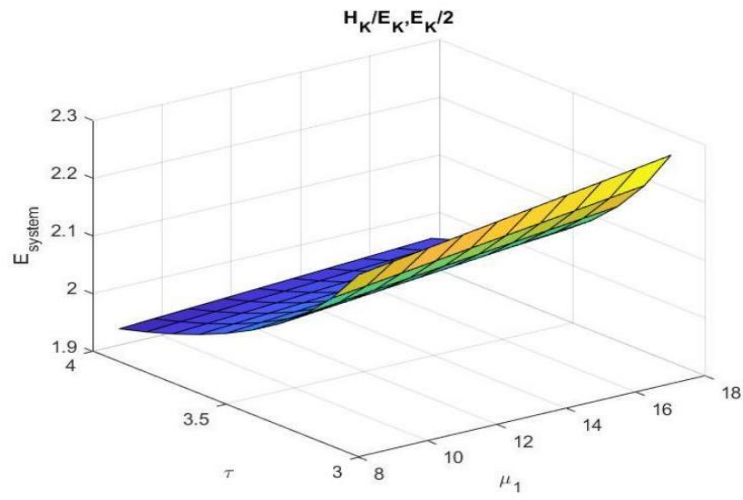


Figure 7: Server-1 service rate and Repair Rate vs E_{system} .

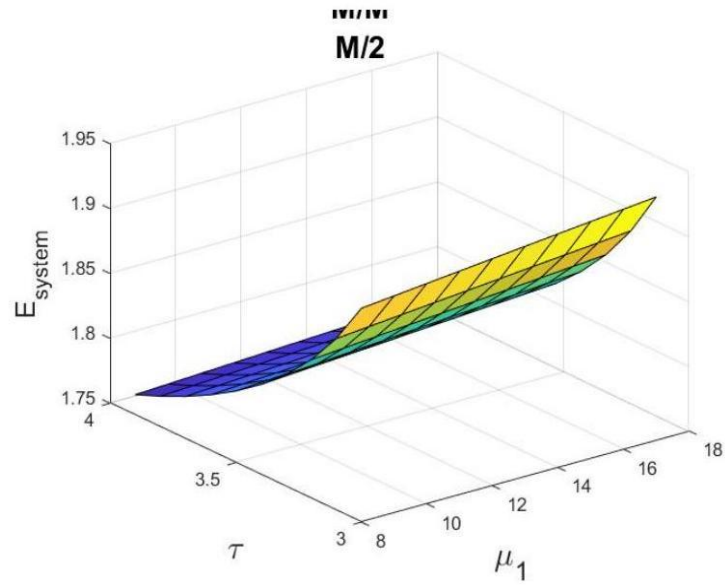


Figure 9: Server-1 service rate and Repair Rate vs E_{system} .

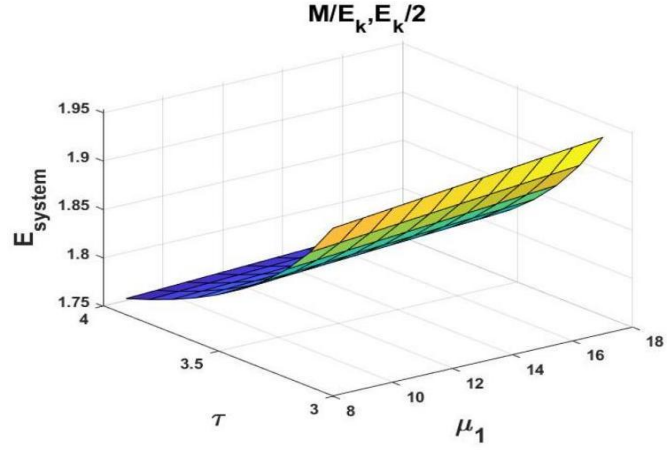


Figure 6: Server-1 service rate and Repair Rate vs E_{system} .

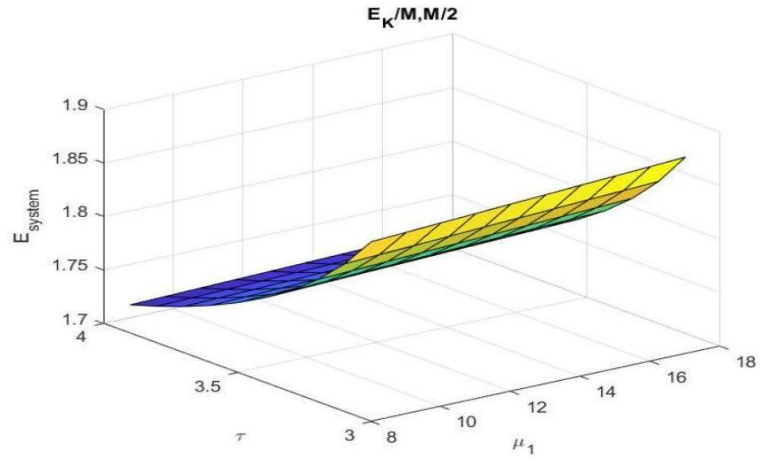


Figure 8: Server-1 service rate and Repair Rate vs E_{system} .

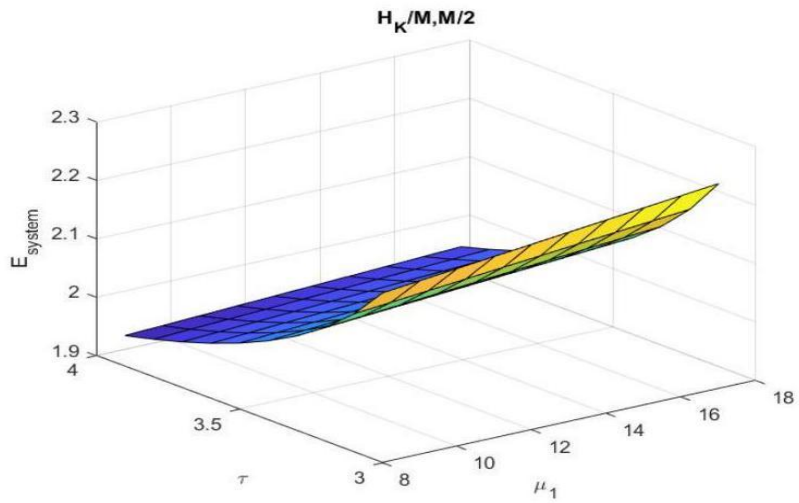


Figure 10: Server-1 service rate and Repair Rate vs E_{system} .

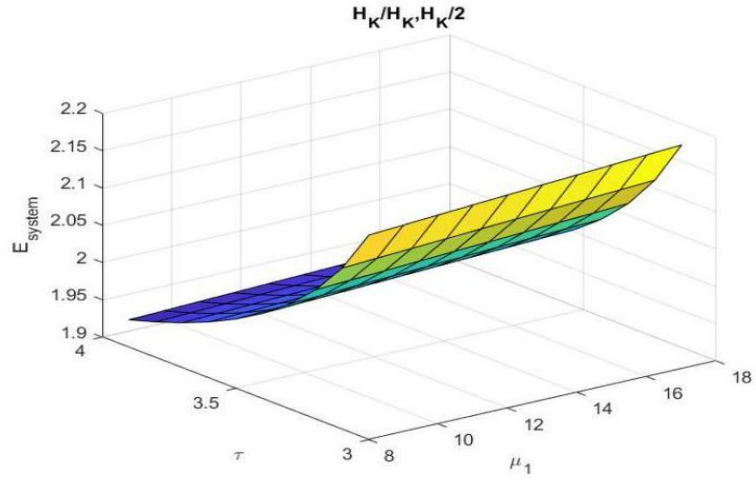


Figure 11: Server-1 service rate and Repair Rate vs E_{system} .

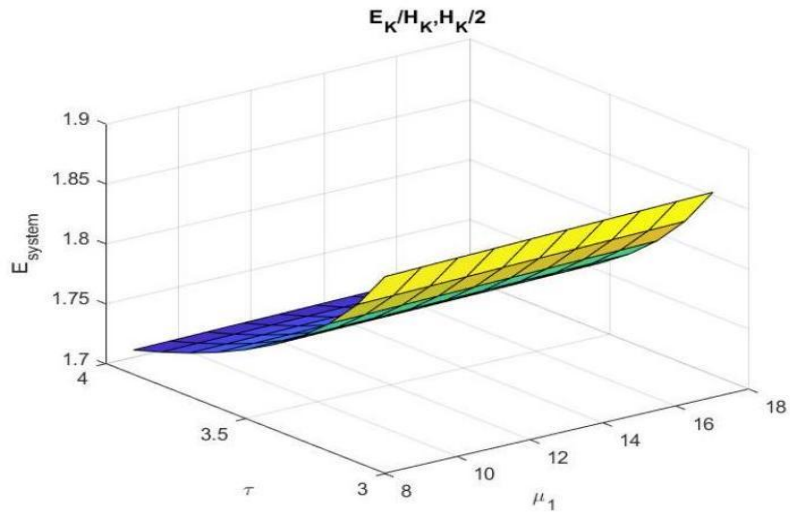


Figure 12: Server-1 service rate and Repair Rate vs E_{system} .

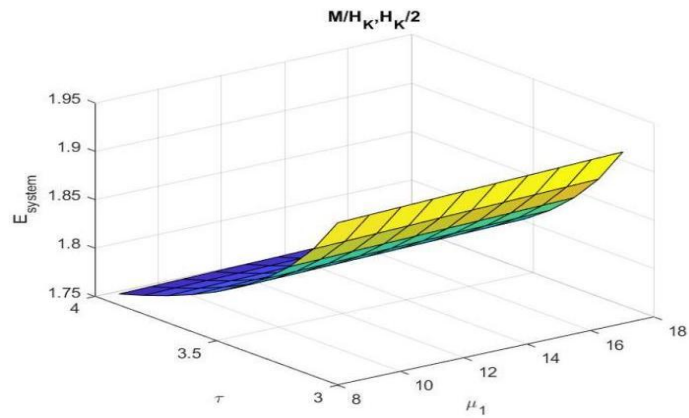


Figure 13: Server-1 service rate and Repair Rate vs E_{system} .

4. CONCLUSION

A two-server retrial queueing-inventory system with double orbit, working vacation, breakdown, repair, re-service, emergency replenishment, and (s, S) policy was examined. Customers arrive in accordance with a Markovian arrival process and a phase-type distribution adapts the service times. A continuous-time Markov chain with a discrete state space is used to create the system's mathematical model. The matrix-analytic method is used to derive the stability condition and subsequently the steady-state distribution. Numerical examples are used to demonstrate how the parameters affect the performance measures. Considering the QIS model with a batch Markov process (BMAP) and perishable inventory could be an extension of this work.

REFERENCES

- [1] Artalejo, J.R.; Falin, G.I. Standard and retrial queueing systems: a comparative analysis. *Rev. Mat. Complut.*, 2002 Vol. 15, pp:101-129.
- [2] Abdul Reiyas, M.; Jeganathan, K. A classical retrial queueing inventory system with two component demand rate. *Int. J. Oper. Res.*, 2023 Vol. 47(4), pp:508-533, <https://doi.org/10.1504/IJOR.2023.132813>
- [3] Dimitriou, I. A two class retrial system with coupled orbit queues. *Probab. Eng. Inf. Sci.*, 2017 Vol. 31(2), pp:139-179, <https://doi.org/10.1017/S0269964816000528>
- [4] Choi, B.D.; Chang, Y. MAP1, MAP2 /M/C retrial queue with the retrial group of finite capacity and geometric loss. *Math. Comput. Modell.*, 1999 Vol. 30, pp:99-113, [https://doi.org/10.1016/S0895-7177\(99\)00135-1](https://doi.org/10.1016/S0895-7177(99)00135-1)
- [5] Kim, J.; Kim, B. (2016). A survey of retrial queueing systems. *Ann. Oper. Res.*, 2016 Vol. 247, pp:3-36, <https://doi.org/10.1007/s10479-015-2038-7>
- [6] Jain, M.; Bhagat, A.; Shekhar, C. (2015). Double orbit finite retrial queues with priority customers and service interruptions. *Appl. Math. Comput.*, 2015 Vol. 253, pp:324-344, <https://doi.org/10.1016/j.amc.2014.12.066>
- [7] Dhibar, S.; Jain, M. Strategic behaviour for M /M/1 double orbit retrial queue with imperfect service and vacation. *Int. J. Math. Oper. Res.*, 2023 Vol. 25(3), pp:369-385, <https://doi.org/10.1504/IJMOR.2023.132479>
- [8] Wang, Y.; Cai, C. Equilibrium strategies of customers and optimal inventory levels in a make-to-stock retrial queueing system. *AIMS Math.*, 2024 Vol. 9(5), pp:12211-12224, doi:10.3934/math. 2024596
- [9] Avrachenkov, K.; Dudin, A. and Klimenok, V. Retrial queueing model MMAP/M 2/1 with two orbits. Springer, Berlin, Heidelberg., 2010 Vol. 6235, <https://doi.org/10.1007/978-3-642-15428-712>
- [10] Seenivasan, M.; Epsiya, S.J. Retrial queueing model with server breakdown and feedback. *J. Comput. Anal. Appl.*, 2023 Vol. 31(2), pp:224-235, <https://orcid.org/0009-0005-3174-762X>
- [11] Melikov, A.; Chakravarthy, S. and Aliyeva, S. A retrieval queueing model with feedback. *Queueing Models Serv. Manag.*, 2023 Vol. 6(1), pp:63-95.
- [12] Shajin, D.; Melikov, A. Queueing inventory systems with return of purchased items and customer feedback. *RAIRO Oper. Res.*, 2025 Vol. 59, pp:1443-1473, <https://doi.org/10.1051/ro/2025042>
- [13] Ayyappan, G.; Archana Alias Gurulakshmi, G. A two-server tandem queueing model with Markovian arrival process, phase type service, working vacation, feedback, reservice, unreliable server, interrupted close down and customers intolerance. *Int. J. Math Oper. Res.*, 2025 Vol. 30(2), pp:206-250, <https://doi.org/10.1504/IJMOR.2025.145605>
- [14] Servi, L.D.; Finn, S.G. M/M/1 queues with working vacations, M/M/WV. *Perform. Eval.*, 2002 Vol. 50(1), pp:41-52, [https://doi.org/10.1016/S0166-5316\(02\)00057-3](https://doi.org/10.1016/S0166-5316(02)00057-3)
- [15] Sindhu, S., Krishnamoorthy, A. and Kozyrev, D. On queues with working vacation and interdependence in arrival and service processes. *Mathematics* 2023, 2023 Vol. 11(10), pp:1-16, <https://doi.org/10.3390/math11102280>
- [16] Sindhu, S.; Krishnamoorthy, A. MAP/Ek/1 Queue with working vacation providing main service Only in normal mode of service. *Queueing Models Serv. Manag.*, 2023 Vol. 6(2), pp:1-27.
- [17] Ayyappan, G.; Arulmozhi, N. Numerical investigation of retrial queueing inventory system with a constant retrial rate, working vacation, flush out, collision and impatient customers. *Reliab.: Theory Appl.*, 2024 Vol. 1(77), pp:744-763, DOI: 10.24412/1932-2321-2024-177-744-763
- [18] Apice, C.D.; Arienzo, M.P.; Dudin, A. and Manzo, R. Optimal hysteresis control via a queueing system with two heterogeneous energy consuming servers. *Mathematics*, 2023 Vol. 11(21), pp:4515, <https://doi.org/10.3390/math11214515>
- [19] Dudin, A. N.; Dudin, S. A.; Klimenok, V. I., & Dudina, O. S. Stability of queueing systems with impatience, balking and non-persistence of customers. *Mathematics*, 2024 Vol. 12(14), pp:2214,

<https://doi.org/10.3390/math12142214>

[20] Muthukumar, S.; Bagyam, E.A.J. A Steady-state behaviour of an queue with optional differentiated working vacations, server breakdown, and customer balking. *Adv. Appl. Stat.*, 2025 Vol. 92(4), pp:603-631,

<https://doi.org/10.17654/0972361725025>

[21] Lv, S.; Li, R. Queueing inventory systems with impatient customers and working vacations. *Int. J. Appl. Math.*, 2024 Vol. 54(12), pp:2569-2579.

[22] Yang, T.; Templeton, J.G.C. (1987). A survey on retrial queues. *Queueing Syst.*, 1987 Vol. 2(3), pp:201-233,

<https://doi.org/10.1007/BF01158899>

[23] Artalejo, J.R. A classified bibliography of research on retrial queues. *Queueing Syst.*, 1999 Vol. 7(2), pp:187-211, <https://doi.org/10.1007/BF02564721>

[24] Benny, B.; Chakravarthy, S. R. and Krishnamoorthy, A. Queueing-Inventory System with Two Commodities. *J. Indian Soc. Probab. Stat.*, 2018 Vol. 19(1), PP:437-454. DOI:10.1007/s41096-018-0052-1

[25] Dissa, S.; Ushakumari, P.V. Two commodity perishable queueing inventory system with random common lifetime of commodities and positive lead time. *OPSEARCH*, 2023 Vol. 61, pp:809-834,

<https://doi.org/10.1007/s12597-023-00707-3>

[26] Neuts, M.F. Matrix geometric solutions in stochastic models: An algorithmic approach, The Johns Hopkins University Press, 1981, Baltimore and London.

[27] Latouche, G.; Ramaswami, V. Introduction of matrix analytic methods in stochastic modeling. SIAM, 1999 Philadelphia.

UDC: 517.547

DOI: <https://doi.org/10.30546/09090.2026.001.20.2006>

COEFFICIENT BOUNDS, INVESTIGATING INCLUSION PROPERTIES FOR A GENERAL SUBCLASS OF ANALYTIC FUNCTIONS DEFINED BY NOOR INTEGRAL OPERATOR

Sercan KAZIMOĞLU*

Kafkas University
 sercan.kazimoglu@kafkas.edu.tr
 Kars-Türkiye

Erhan DENİZ

Kafkas University
 edeniz36@gmail.com
 Kars-Türkiye

ARTICLE INFO	ABSTRACT
Article history: Received:2025-12-06 Received in revised form:2025-12-19 Accepted:2026-01-10 Available online:2026-06-23 Keywords: Analytic function; Inclusion; Coefficient bounds; Noor integral operator. 2010 Mathematics Subject Classifications: 30C45	In this paper, we introduce a new subclass $B_m^*(\varphi, \delta, \alpha)$ of analytic functions defined via the Noor integral operator. The class is formulated through appropriate analytic conditions involving the parameters φ, δ and α. A systematic investigation of this subclass is presented, with particular emphasis on coefficient estimates for functions belonging to $B_m^*(\varphi, \delta, \alpha)$. In addition, several closure theorems are established, demonstrating that the class is preserved under relevant algebraic operations. Fundamental structural properties of the proposed subclass are derived and discussed in detail. The obtained results extend and complement existing studies in geometric function theory and related subclasses of analytic functions.

2521-635X/© 2026 The Author(s). Published by Baku Engineering University.

This is an open access article under the CC BY 4.0 license (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Let A represent the set of analytic functions f defined on the open unit disk $U = \{z \in \mathbb{C} : |z| < 1\}$, which are normalized by the conditions $f(0) = 0$ and $f'(0) = 1$. Consequently, every function $f \in A$ can be expressed as a Taylor Maclaurin series of the form:

$$f(z) = z + \sum_{n=2}^{\infty} a_n z^n, \quad (z \in U). \quad (1)$$

In 1975, Silverman [1] introduced and analyzed a specific subset of comprised A of functions where the coefficients, starting from the second term, are negative. In other words, the analytic functions f within this subset can be represented as

$$f(z) = z - \sum_{n=2}^{\infty} a_n z^n, \quad (z \in U), \quad a_n \geq 0. \quad (2)$$

This subclass is known as the class of analytic functions with negative coefficients and is denoted as A^* . Following Silverman's contribution, there has been considerable interest in studying functions with negative coefficients. Building on Silverman's work, many additional subclasses A of have been investigated in academic literature.

Denote by

$$R^m := \frac{z}{(1-z)^{m+1}} * f(z) \quad (m \in \mathbb{N} \cup \{0\}).$$

Then implies that

$$R^m f(z) = \frac{z(z^{n-1} f(z))^m}{m!} \quad (m \in \mathbb{N} \cup \{0\}).$$

The operator $R^m f$ is called Ruscheweyh derivative operator [2]. Noor [3] defined and investigated an integral operator $N^m : A \rightarrow A$ analogous to $R^m f$ as follows.

Let $f_m(z) = \frac{z}{(1-z)^{m+1}}$, $m \in \mathbb{N}_0 = \mathbb{N} \cup \{0\}$, and $f_m^{(-1)}$ be defined such that

$$f_m(z) * f_m^{(-1)}(z) = \frac{z}{1-z}.$$

Then

$$N^m f(z) = f_m^{(-1)}(z) * f(z) = \left[\frac{z}{(1-z)^{m+1}} \right]^{(-1)} * f(z) = z + \sum_{n=2}^{\infty} \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n z^n, \quad (3)$$

where Γ is Euler gamma function.

Yousef et al. in [4] initially introduced the class $B_{\Sigma}^{\mu}(\varphi, \delta, \alpha)$.

Definition 1: For $\varphi \geq 1, \delta \geq 0, \mu \geq 0$ and $0 \leq \alpha < 1$, a function $f \in A$ represented by (1) is said to be in the class $B_{\Sigma}^{\eta}(\varphi, \mu, \alpha)$ if the following conditions hold for all $z \in U$.

$$\operatorname{Re} \left\{ (1-\varphi) \left(\frac{f(z)}{z} \right)^{\eta} + \varphi f'(z) \left(\frac{f(z)}{z} \right)^{\eta-1} + \xi \delta z f''(z) \right\} > \alpha,$$

where $\xi = \frac{2\delta + \eta}{2\lambda + 1}$.

Several researches have explored the class $B_{\Sigma}^{\eta}(\varphi, \mu, \alpha)$ and utilized it in various contexts. For example, we refer the reader to [5-10].

Definition 2. For $\varphi \geq 1$, $\delta \geq 0$, and $0 \leq \alpha < 1$, an analytic function f represented by (2) is said to be in the class $B_m^*(\varphi, \delta, \alpha)$ if the following conditions hold for all $z \in U$.

$$\operatorname{Re} \left\{ (1-\varphi) \frac{N^m f(z)}{z} + \varphi (N^m f(z))' + \delta z (N^m f(z))'' \right\} > \alpha, \quad (4)$$

where $m \in \mathbb{N}_0 = \mathbb{N} \cup \{0\}$.

2. COEFFICIENTS ESTIMATES

This section commences with deriving a necessary and sufficient condition for a function $f(z)$ to belong to the class $B_m^*(\varphi, \delta, \alpha)$.

Theorem 2.1. A function $f \in A^*$ defined by (2) belong to the class $B_m^*(\varphi, \delta, \alpha)$ if and only if

$$\sum_{n=2}^{\infty} [1-\varphi + n\varphi + n(n-1)\delta] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \leq 1-\alpha, \quad (5)$$

where $\lambda \geq \mu \geq 0$ and $m \in \mathbb{N}_0 = \mathbb{N} \cup \{0\}$.

Proof. Let the function $f(z)$ defined by (2) be in the class $B_m^*(\varphi, \delta, \alpha)$ and $z \in U$. Then by definition

$$\begin{aligned} & \operatorname{Re} \left\{ (1-\varphi) \frac{N^m f(z)}{z} + \varphi (N^m f(z))' + \delta z (N^m f(z))'' \right\} \\ &= \operatorname{Re} \left\{ 1 + \sum_{n=2}^{\infty} [1-\varphi + n\varphi + n(n-1)\delta] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n z^{n-1} \right\} > \alpha. \end{aligned}$$

As z approaches 1^- along the real axis, inequality (2) follows.

Conversely, the inequality (5) holds, then $z \in U$, we have

$$\begin{aligned} & \left| (1-\varphi) \frac{N^m f(z)}{z} + \varphi (N^m f(z))' + \delta z (N^m f(z))'' - 1 \right| \\ &= \left| \sum_{n=2}^{\infty} [1-\varphi + n\varphi + n(n-1)\delta] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n z^{n-1} \right| \\ &\leq \sum_{n=2}^{\infty} [1-\varphi + n\varphi + n(n-1)\delta] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \leq 1-\alpha. \end{aligned}$$

This yields $f \in B_m^*(\varphi, \delta, \alpha)$.

Corollary 2.2. Let $f(z)$ defined by (2) be in the class $B_m^*(\varphi, \delta, \alpha)$. Then

$$a_n \leq \frac{1-\alpha}{[1-\varphi + n\varphi + n(n-1)\delta] \frac{\Gamma(m+1)n!}{\Gamma(m+n)}}, \quad n \geq 2.$$

The above inequality holds for the function

$$f(z) = z - \frac{1-\alpha}{\left[1-\varphi+n\varphi+n(n-1)\delta\right] \frac{\Gamma(m+1)n!}{\Gamma(m+n)}} z^n.$$

3. INCLUSION RELATIONS

This section investigates the inclusion properties of the subclass $B_m^*(\varphi, \delta, \alpha)$. By examining various inclusion relations, we aim to understand how different subclasses relate to each other within this framework. The results presented offer insights into how certain classes are contained within others, contributing to a broader understanding of the hierarchical structure within the class of analytic functions under consideration.

Theorem 3.1. Let $0 \leq \alpha_1 \leq \alpha_2 < 1$. Then

$$B_m^*(\varphi, \delta, \alpha_1) \supseteq B_m^*(\varphi, \delta, \alpha_2). \quad (6)$$

Proof. Consider the function $f(z)$ defined by equation (2) to belong to the class $B_m^*(\varphi, \delta, \alpha_2)$. Then, by Theorem 2.1, we have

$$\begin{aligned} \sum_{n=2}^{\infty} \left[(1-\varphi) + n\varphi + n(n-1)\delta \right] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \\ \leq 1 - \alpha_2 \leq 1 - \alpha_1 \end{aligned}$$

Thus, $f \in B_m^*(\varphi, \delta, \alpha_1)$.

Therefore, the inclusion relation shown by (6) holds.

Theorem 3.2. Let $1 \leq \varphi_1 \leq \varphi_2$. Then

$$B_m^*(\varphi_1, \delta, \alpha) \supseteq B_m^*(\varphi_2, \delta, \alpha). \quad (7)$$

Proof. Consider the function $f(z)$ defined by equation (2) to belong to the class $B_m^*(\varphi_2, \delta, \alpha)$. Then, by Theorem 2.1, we have

$$\begin{aligned} \sum_{n=2}^{\infty} \left[(1-\varphi) + n\varphi + n(n-1)\delta \right] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \\ \leq \sum_{n=2}^{\infty} \left[(1-\varphi_2) + n\varphi_2 + n(n-1)\delta \right] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \\ \leq 1 - \alpha. \end{aligned}$$

Thus, $f \in B_m^*(\varphi_1, \delta, \alpha)$.

Therefore, the inclusion relation shown by (7) holds.

Theorem 3.3. let $0 \leq \delta_1 \leq \delta_2$. Then

$$B_m^*(\varphi, \delta_1, \alpha) \supseteq B_m^*(\varphi, \delta_2, \alpha). \quad (8)$$

Proof. Consider the function $f(z)$ defined by equation (2) to belong to the class $B_m^*(\varphi, \delta_2, \alpha)$. Then, by Theorem 2.1, we have

$$\begin{aligned} & \sum_{n=2}^{\infty} \left[(1-\varphi) + n\varphi + n(n-1)\delta_1 \right] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \\ & \leq \sum_{n=2}^{\infty} \left[(1-\varphi) + n\varphi + n(n-1)\delta_2 \right] \frac{\Gamma(m+1)n!}{\Gamma(m+n)} a_n \\ & \leq 1-\alpha. \end{aligned}$$

Thus, $f \in B_m^*(\varphi, \delta_1, \alpha)$.

Therefore, the inclusion relation shown by (8) holds.

Theorem 3.4. Let $0 \leq m_1 \leq m_2$. Then

$$B_{m_1}^*(\varphi, \delta, \alpha) \supseteq B_{m_2}^*(\varphi, \delta, \alpha) \quad (9)$$

Proof. Consider the function $f(z)$ defined by equation (2) to belong to the class $B_{m_2}^*(\varphi, \delta, \alpha)$.

Then, by Theorem 2.1, we have

$$\begin{aligned} & \sum_{n=2}^{\infty} \left[(1-\varphi) + n\varphi + n(n-1)\delta \right] \frac{\Gamma(m_1+1)n!}{\Gamma(m_1+n)} a_n \\ & \leq \sum_{n=2}^{\infty} \left[(1-\varphi) + n\varphi + n(n-1)\delta \right] \frac{\Gamma(m_2+1)n!}{\Gamma(m_2+n)} a_n \\ & \leq 1-\alpha. \end{aligned}$$

Thus, $f \in B_{m_1}^*(\varphi, \delta, \alpha)$.

Therefore, the inclusion relation shown by (9) holds.

4. CONCLUSION

In this study, a new subclass $B_m^*(\varphi, \delta, \alpha)$ of analytic functions associated with the Noor integral operator has been introduced and examined. Several fundamental properties of this subclass have been established, including coefficient bounds and closure results. The findings contribute to the ongoing development of geometric function theory by extending existing results related to integral operators. The results obtained in this study are expected to provide a basis for further studies through the development of new function classes or integral operators.

REFERENCE LIST

- [1] Silverman, H. (1975). Univalent functions with negative coefficients. *Proceedings of the American mathematical society*, 51(1), 109-116.
- [2] Ruscheweyh, S. (1975). New criteria for univalent functions. *Proc. Am. Math. Soc.*, 49, 109-115.
- [3] Noor, K. I. (1999). On new classes of integral operators. *J. Nat. Geometry*, 16, 71-80.
- [4] Yousef, F., Alroud, S., Illafe, M. (2021). New subclasses of analytic and bi-univalent functions endowed with coefficient estimate problems. *Analysis and Mathematical Physics*, 11(2), 58.
- [5] Illafe, M., Yousef, F., Mohamed, M. H., Supramaniam, S. (2024). Fundamental properties of a class of analytic functions defined by a generalized multiplier transformation operator. *International Journal of Mathematics and Computer Science*, 19(4), 1203-1211.
- [6] Hussien, A., Madi, M. S., Abominjil, A. M. (2024). Bounding coefficients for certain subclasses of bi-univalent

- functions related to Lucas-Balancing polynomials. AIMS Mathematics, 9(7), 18034-18047.
- [7] Hussen, A., Zeyani, A. (2023). Coefficients and Fekete–Szegő functional estimations of bi-univalent subclasses based on Gegenbauer polynomials. Mathematics, 11(13), 2852.
 - [8] Hussen, A., Illafe, M. (2023). Coefficient bounds for a certain subclass of bi-univalent functions associated with Lucas-balancing polynomials. Mathematics, 11(24), 4941.
 - [9] Illafe, M., Mohd, M. H., Yousef, F., Supramaniam, S. (2024). A subclass of bi-univalent functions defined by asymmetric q -derivative operator and gegenbauer polynomials. European Journal of Pure and Applied Mathematics, 17(4), 2467-2480.
 - [10] Amourah, A., Al-Hawary, T., Yousef, F., Salah, J. (2025). Collection of bi-univalent functions using bell distribution associated with Jacobi polynomials. Int. J. Neutrosophic Sci, 25(1), 228-238.

UDC: 550.34:519.217

DOI: <https://doi.org/10.30546/09090.2026.001.20.2010>

INTEGRATION OF NON-STATIONARITY IN DISASTER RISK FORECASTING: A MULTIDISCIPLINARY STOCHASTIC ANALYSIS OF SEISMIC DYNAMICS IN AZERBAIJAN BASED ON HIDDEN MARKOV MODELS

MIR RAMIN YUNUSOV

Baku Engineering University
The Institute of Control Systems
Baku, Azerbaijan
myunusov@beu.edu.az

ARTICLE INFO	ABSTRACT
Article history: Received:2025-09-13 Received in revised form:2025-10-02 Accepted:2025-11-05 Available online: 2026-06-23 Keywords: Disaster risk management, Non-stationarity, Hidden Markov Models (HMM), Hidden Semi-Markov Models (HSMM), Seismic forecasting, Seismotectonics of Azerbaijan, Weibull distribution. 2010 Mathematics Subject Classifications: 60J20, 60K15, 62M05, 86A15, 86A17	Reliable predictive modeling within geophysics and disaster risk management necessitates a precise evaluation of the intrinsic dynamics found in complex natural systems. Seismic events, particularly those driven by lithospheric plate interactions, are fundamentally defined by their non-linear and non-stationary characteristics. To surmount the constraints of conventional stationary methods—such as Poisson processes and Gutenberg-Richter laws—this research explores the integration of non-stationarity through Hidden Markov Models (HMM) and their advanced variant, Hidden Semi-Markov Models (HSMM). The study applies stochastic and statistical analysis to 70 seismic occurrences registered in Azerbaijan and its border regions throughout 2024. Findings indicate that the “memoryless” nature and geometric distribution assumption of standard HMMs fail to adequately capture the seismic gaps and cluster-type activity observed in the Kalbajar and Lankaran zones. In contrast, the HSMM framework, which utilizes Log-normal or Weibull distributions to model state sojourn times, offers a superior fit for the region’s complex collisional geodynamics. Ultimately, the paper presents a novel, mathematically robust framework designed to enhance early warning systems and dynamic risk assessment across Azerbaijan’s seismotectonic belts.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

INTRODUCTION

The accelerating pace of urbanization, increasing complexity of critical infrastructure, and the rising global frequency and intensity of natural and technogenic disasters necessitate a fundamental revision of existing scientific paradigms in risk forecasting and management. Traditional methodologies widely employed in engineering seismology and hydrology (e.g., Probabilistic Seismic Hazard Assessment – PSHA) are predominantly predicated on the assumption of process stationarity. This deterministic or quasi-stationary approach assumes that the statistical properties of past events—such as means, variance, autocorrelation functions, and

recurrence rates—will remain constant over future long-term periods (Danciu et al., 2024). While this assumption simplifies mathematical modeling, it is largely insufficient for accurately reflecting the real, dynamic nature of natural systems (Buckby et al., 2020).

Extensive geophysical research and the analysis of recent extreme events demonstrate that the Earth system is fundamentally non-stationary. Non-stationarity refers to the evolution of statistical parameters and probability distribution functions of a time series over time; this variability is not merely a result of random fluctuations but stems from fundamental shifts in the system's internal structure and external forcing. Two primary sources of non-stationarity in disaster risks can be distinguished:

- **Internal Geodynamic Evolution:** Crustal processes typically exhibit cyclic or chaotic behavior. For instance, cycles of tectonic stress accumulation and sudden release (earthquakes) fundamentally alter the seismic regime. Following a major earthquake, the redistribution of the stress field, changes in fault rheology, and the non-linear decay of aftershock sequences render simple stationary models (e.g., homogeneous Poisson processes) inapplicable (Pasari & Dikshit, 2018).
- **Dynamics of External Factors:** Climate change alters the statistical distribution of atmospheric and hydrospheric events (e.g., floods, droughts), potentially reducing the return period of extreme events (e.g., shifting a "100-year flood" to a "20-year" recurrence). Similarly, anthropogenic factors such as reservoir construction, hydrocarbon extraction, and groundwater fluctuation can induce seismicity, thereby rendering local seismic regimes non-stationary (Alawsi et al., 2022).

The Republic of Azerbaijan, selected as the object of this study, is characterized by one of the most complex and active seismotectonic settings globally, situated at the active convergent boundary of the Eurasian and Arabian lithospheric plates within the eastern segment of the Alpine-Himalayan fold belt. The region's geodynamics are defined by the intensive orogenic uplift of the Greater and Lesser Caucasus, the relative subsidence of the Kura Depression, and the unique subduction-analogue tectonics of the Caspian Sea basin (specifically the South Caspian Basin) (Yetirmishli & Kazimova, 2024).

Historical records and modern instrumental observations indicate that seismic activity in Azerbaijan is not distributed homogeneously or essentially stationarily over time; rather, long periods of quiescence are interspersed with phases of short-term, acute activation. Historical cataclysms such as the Ganja earthquake (1139), the Shamakhi earthquake (1668), and the Shamakhi disaster (1902) illustrate the region's high and destructive seismic potential. The seismic clusters ("burst" activity) recorded in the Lankaran, Lerik, and Kalbajar zones, as well as persistent tremors in the Caspian Sea during 2024, serve as contemporary evidence of this non-stationarity (RSSC, 2024). Traditional linear regression models or simple Poisson processes fail to model such complex systems adequately, as they do not account for event sequencing, temporal dependencies (memory), or the dynamics of state transitions. To address this methodological gap, Hidden Markov Models (HMM) are proposed, which postulate the existence of "hidden" states (e.g., stress accumulation, critical loading, or release) that govern the physical behavior of the system behind the observed data (e.g., earthquake catalogs) (Allen & Wang, 2024).

The primary objective of this paper is to investigate and implement innovative stochastic modeling approaches—specifically Hidden Semi-Markov Models (HSMM)—that account for the impact of non-stationarity in forecasting seismic risks in Azerbaijan. The research aims to address the following specific tasks:

- Analyze the theoretical foundations of HMMs and HSMMs and their potential application in seismology;
- Identify the methodological limitations of standard HMMs in seismic forecasting, particularly regarding stationarity and the geometric distribution assumption;
- Statistically evaluate the variability of seismic regimes in Azerbaijan based on the 2024 earthquake catalog;
- Demonstrate the application of HSMM for more accurate modeling of the duration distributions of seismic regimes ("quiet" and "active" phases) and its implications for risk assessment (Shahzadi et al., 2025).

Seismotectonic Setting and Contemporary Geodynamics of Azerbaijan: Physical Determinants of Non-Stationarity. Stochastic modeling of seismic activity and risks in Azerbaijan necessitates, first and foremost, a profound understanding of the region's geological and tectonic structure, active fault systems, and stress fields. These physical realities constitute the "hidden states" integrated into our mathematical models.

Located in the central segment of the Alpine-Himalayan orogenic belt, the geodynamics of Azerbaijan are governed by the continuous continental collision between the northward-moving Arabian plate and the relatively stable Eurasian plate. Analysis of GPS velocity vectors indicates that the Arabian plate moves northward at a rate of approximately 18-25 mm/year, resulting in intense compressional deformation within the Caucasus region (Kadirov et al., 2015; Yetirmishli et al., 2022). This deformation energy accumulates in the lithosphere and is periodically released through seismic events.

The region's tectonic architecture is characterized by several major megazones separated by deep-seated faults:

- **Greater Caucasus:** Characterized by intense uplift and south-vergent thrust faults (Mellors et al., 2015).
- **Kura Depression:** An intermontane trough undergoing hidden deformation beneath a thick sedimentary cover.
- **Lesser Caucasus:** Distinguished by volcanogenic-sedimentary rocks and a complex block structure (Mammadli, 2019).
- **Talysh Zone:** A continuation of the Alborz system, exhibiting a transpressional regime (Aliyev & Babayev, 2023).
- **South Caspian Basin:** A unique tectonic block with an unusually thick sedimentary layer and subduction-like deep earthquakes (Pierce et al., 2024).

The Greater Caucasus mountain range represents one of Azerbaijan's most active seismic zones, with activity primarily linked to the Vandam, Girdimanchay, and Siyezen faults along the southern slope. The Shamakhi-Gobustan zone within this region is historically considered one of the most potent and destructive seismic sources. Historical records of earthquakes in 1192, 1667, 1669, 1828, 1859, 1872, and 1902 attest to the zone's high seismic potential and the periodicity (non-stationarity) of the hazard.

The Shamakhi-Gobustan zone also hosts one of the highest densities of mud volcanoes globally. A correlation exists between mud volcanism and seismic activity; strong earthquakes can trigger eruptions, and conversely, volcanic activity can alter the local stress field. According to 2024 data, more than four felt earthquakes were recorded in this zone (covering Shamakhi, Ismayilli, and Gobustan districts), indicating that the zone remains active and may currently be in a phase of relative "quiescence" or energy accumulation.

The Talysh zone, covering southeastern Azerbaijan, differs geologically from the Lesser Caucasus and is regarded as the northwestern extension of the Iranian Alborz mountains. This zone is characterized by both compression and transpression regimes. Seismicity in the Talysh zone is closely associated with the West Caspian Fault, which runs along and within the southwestern coast of the Caspian Sea.

In 2024, anomalous seismic activity and clustering were observed in this region. Of particular note is the magnitude 5.3 earthquake that occurred on May 11, 2024, in the Lankaran district, followed by a strong magnitude 4.4 aftershock just 1 minute and 18 seconds later (13:23:06). This event was not isolated; another magnitude 5.0 earthquake was recorded in the Lerik district on June 28. This activity indicates the activation of faults in the region, particularly the Talysh and Astara faults, and the intensive release of stress. The "Z"-shaped geometry and meridional orientation of the faults suggest that the seismic sources here possess complex kinematics (dextral strike-slip and thrusting).

The eastern and central parts of the Lesser Caucasus, including the Kalbajar and Lachin districts, possess a complex block tectonic structure. Volcanogenic-sedimentary rocks from the Jurassic and Cretaceous periods, as well as ophiolite belts (Goycha-Hakari zone) representing remnants of ancient oceanic crust, are widespread in these areas. Seismicity in the Kalbajar district is linked to faults along the Tartar River and the Murovdag range.

In the fourth quarter of 2024, specifically on December 20, a series of earthquakes (cluster) occurring within a short interval in the Kalbajar area (near the Azerbaijan-Armenia border) indicated a local concentration of stress. The occurrence of consecutive tremors with magnitudes of 4.8, 3.0, and 3.7 within one hour suggests active movement of tectonic blocks and the onset of a "seismic swarm" regime. The activity in this zone is explained by intra-block deformations and the activation of the Murovdag fault against the backdrop of the general collision of the Caucasus.

The Caspian Sea, particularly the Middle and South Caspian basins, accounts for a significant portion of Azerbaijan's seismic activity. The Absheron-Balkhan Sill and the South Caspian Depression are distinguished by processes occurring in the deep layers of the lithosphere. Many researchers suggest that the South Caspian Basin has an oceanic-type crust that is being buried (subduction or incipient subduction) northward under the Scythian-Turan platform.

The 2024 catalog recorded the highest number of earthquakes (more than 30 events) in the Caspian Sea basin. A characteristic feature of these earthquakes is their focal depth, which, unlike onshore earthquakes, is significantly deeper. Numerous earthquakes were recorded at depths of 60–75 km throughout the year, pointing to processes occurring in the subduction zone (Benioff zone). Although earthquakes at this depth typically cause less surface damage, they serve as indicators of the region's overall stress state.

The Kura Depression is a vast intermontane trough located between the Greater and Lesser Caucasus mountain systems. As this zone is covered by a thick sedimentary layer (molasse) reaching 15–20 km in some places, most tectonic faults do not reach the surface and are characterized as "blind" faults. Earthquakes occurring in the Saatli-Sabirabad, Imishli, and Bilasuvar areas (e.g., the magnitude 5.1 Saatli earthquake on January 19, 2024) are indicative of active tectonic processes beneath the sedimentary cover, particularly at the intersection nodes of the Kura and West Caspian faults. While seismicity in the Kura Depression is typically of moderate intensity, events in this zone can pose risks to infrastructure (pipelines, dams).

Theoretical Framework: Hidden Markov Models (HMM) and Algorithmic Implementation.

Forecasting seismic hazards involves modeling complex stochastic systems that evolve over time. In this study, Hidden Markov Models (HMM) are employed to establish a probabilistic correlation between directly measurable seismic parameters (e.g., magnitude, inter-event intervals) and latent geophysical processes that cannot be directly observed (e.g., stress accumulation, rheological changes in fault zones) (Mor et al., 2021).

Mathematical Architecture of the Model. An HMM is defined as a doubly stochastic process comprising an observed sequence Y_t and a sequence of hidden states X_t . The model is fully characterized by the parameter set $\lambda = (A, B, \pi)$ consistent with the physics of the seismic cycle, $N=3$ hidden states are defined: S_1 (Background/Loading), S_2 (Pre-seismic/Accelerating) and S_3 (Release/Earthquake).

Python initialization of the state space and transition matrix (A) follows:

```
import numpy as np
from hmmlearn import hmm

# 1. Definition of Model Parameters
# Hidden States: S1 (Loading), S2 (Pre-seismic), S3 (Release)
n_states = 3

# Initial Transition Matrix (A) - Estimated prior probabilities
# a_ij: probability of transition from state i to state j
trans_mat = np.array([
    [0.98, 0.02, 0.00], # High probability of remaining in S1
    [0.05, 0.85, 0.10], # S2 is critical, transition to S3 is possible
    [0.01, 0.05, 0.94] # S3 represents the aftershock sequence
])

# Initial State Distribution (pi)
start_prob = np.array([0.8, 0.15, 0.05]) # System likely starts in S1
```

Fig. 1. Definition of hidden states (loading, pre-seismic, release) and initial transition matrix probabilities

HSMM and Solving the "Memoryless" Problem with Weibull Distribution. A fundamental limitation of standard HMMs is that the sojourn time within a state follows a geometric distribution, exhibiting the "memoryless" property. This contradicts the theory of "elastic rebound." To overcome this, the Weibull distribution is applied within the Hidden Semi-Markov Model (HSMM) framework. The Weibull function ($f(t; \lambda; k)$) adequately describes increasing risk ($k > 1$) or decreasing risk ($k < 1$) over time.

Simulation of the Hazard Rate based on the Weibull distribution in Python:

```
from scipy.stats import weibull_min
import matplotlib.pyplot as plt

# 2. Modeling Weibull Distribution for HSMM
def calculate_hazard_rate(k, lam, t):
    """
    k: shape parameter
    lam: scale parameter
    t: time
    """
    return (k / lam) * (t / lam)**(k - 1)

# Simulation: Stress Accumulation (k=2.5, increasing risk) vs Aftershock (k=0.8, decreasing risk)
t_values = np.linspace(0, 10, 100)
hazard_loading = calculate_hazard_rate(2.5, 5, t_values) # Increasing hazard
hazard_aftershock = calculate_hazard_rate(0.8, 5, t_values) # Decaying activity
```

Fig. 2. Functional implementation of the Weibull distribution for calculating the hazard rate

Computational Algorithms: Baum-Welch and Viterbi. Standard algorithms are applied to estimate model parameters from observed data and to reconstruct the most likely sequence of hidden states.

1. **Baum-Welch Algorithm (Training):** A special case of the Expectation-Maximization (EM) method used to optimize parameters A and B .
2. **Viterbi Algorithm (Decoding):** Determines the optimal hidden "path" that explains the observed earthquake sequence.

Implementation of these algorithms using the "hmmlearn" library in Python:

```
# 3. Model Construction and Algorithm Implementation
# Gaussian HMM (Observations: Log-transformed intervals)
model = hmm.GaussianHMM(n_components=n_states, covariance_type="full", n_iter=1000)

# Model Initialization
model.startprob_ = start_prob
model.transmat_ = trans_mat

# Observed data (Random data for demonstration)
# In practice, the earthquake catalog is loaded here
X_observed = np.random.randn(100, 1)

# BAUM-WELCH: Fitting the model to data (Training)
model.fit(X_observed)

# VITERBI: Finding the most likely hidden states
hidden_states = model.predict(X_observed)

# Result: Determining the phase of the recent events
print("Hidden states of the last 10 events:", hidden_states[-10:])
```

Fig. 3. Model initialization and implementation of Baum-Welch (training) and Viterbi (prediction) algorithms

Model Selection (AIC/BIC). The efficacy of the models is evaluated using the Akaike Information Criterion (AIC) and the Bayesian Information Criterion (BIC), ensuring a balance between model complexity and accuracy.

Methodological Approach and Algorithmic Architecture

The methodological framework of this research comprises four fundamental stages, encompassing stochastic data processing, model topology construction, iterative training processes, and statistical validation.

Data Management and Preprocessing Strategy. The empirical basis was established using a seismic catalog covering Azerbaijan and its border regions for the full calendar year of 2024 (January–December), sourced from the official bulletins of the Republican Seismic Survey Center (RSSC) under ANAS. Only events with a magnitude of $M \geq 3.0$ were retained as the selection criterion for analysis, while industrial explosions and duplicate records were excised from the database.

To monitor non-stationary dynamics and seasonal fluctuations, the data were segmented into quarterly temporal windows. To ensure numerical stability during computation, magnitude (M) and depth (H) parameters were normalized to the $[0, 1]$ interval using the Min-Max method.

Python Implementation:

```
import pandas as pd
from sklearn.preprocessing import MinMaxScaler

# 1. Data Loading and Filtering
df = pd.read_csv('rsxm_catalog_2024.csv')
# Retaining only natural events with M >= 3.0
df_clean = df[(df['magnitude'] >= 3.0) & (df['type'] == 'earthquake')].copy()

# 2. Quarterly Grouping (Temporal Segregation)
df_clean['date'] = pd.to_datetime(df_clean['date'])
df_clean['quarter'] = df_clean['date'].dt.to_period('Q')

# 3. Parameter Normalization (Range 0-1)
scaler = MinMaxScaler()
# Normalizing Magnitude and Depth
df_clean[['M_norm', 'H_norm']] = scaler.fit_transform(df_clean[['magnitude', 'depth']])
```

Fig. 4. Algorithm for loading, filtering, and parameter normalization of the seismic catalog data

Stochastic Model Architecture and State Topology. To describe the system's hidden dynamics, HMM and HSMM structures consisting of $N=3$ hidden states were designed:

- **S₁ (Background Mode):** A period of low seismicity characterized by large inter-event intervals (Δt).
- **S₂ (Activation Mode):** A normal seismic regime with an increased density of moderate tremors.
- **S₃ (Cluster/High Risk):** Strong earthquakes or swarm sequences where Δt is minimal.

Magnitude (M) and inter-event time (Δt) were selected as the observation vector. Emission probabilities (B) were modeled using a Gaussian distribution, while state durations (d) in the HSMM were modeled using a Weibull distribution.

Python Implementation (Model Construction):

```
from hmmlearn import hmm

# HMM Model Configuration
# n_components=3: Background, Active, High Risk
# covariance_type='full': To account for correlations between parameters
model = hmm.GaussianHMM(n_components=3, covariance_type="full", n_iter=200)
```

Fig. 5. Configuration of the Gaussian Hidden Markov Model (HMM) and definition of covariance type

Training Process and Stationarity Verification. Model parameters were optimized on the 2024 catalog using the Baum-Welch algorithm (a modification of the EM method). The iterative process continued until the convergence of the Log-likelihood (LL) function ($\Delta LL < 10^{-4}$). To validate the stationarity assumption, Transition Matrices (A_Q) were calculated individually for each quarter, and the variation between quarters (ΔA) was analyzed.

Python Implementation (Quarterly Training):

```

transition_matrices = {}

# Training the model separately for each quarter
for quarter in df_clean['quarter'].unique():
    subset = df_clean[df_clean['quarter'] == quarter]

    # Observation vector: [Magnitude, Delta_t]
    X_quarter = subset[['M_norm', 'dt_log']].values

    # Training with Baum-Welch
    model.fit(X_quarter)

    # Checking for convergence
    if model.monitor_.converged:
        transition_matrices[quarter] = model.transmat_
        print(f"Transition Matrix calculated for quarter {quarter}.")

```

Fig. 6. Creation of observation vector and model training loop based on time intervals (quarters)

Statistical Evaluation and Validation. The predictive power of the models was comparatively analyzed based on the AIC (Akaike Information Criterion) and BIC (Bayesian Information Criterion). Furthermore, the correspondence between the hidden states reconstructed via the Viterbi algorithm and reality (actual clusters in the catalog) was verified.

Python Implementation (AIC/BIC Calculation):

```

import numpy as np

# Calculation of model efficiency criteria
log_likelihood = model.score(X_quarter)
n_params = model.n_features * model.n_components + model.n_components * (model.n_comp

# AIC = 2k - 2ln(L)
aic = 2 * n_params - 2 * log_likelihood
# BIC = k*ln(n) - 2ln(L)
bic = n_params * np.log(len(X_quarter)) - 2 * log_likelihood

print(f"AIC: {aic:.2f}, BIC: {bic:.2f}")

```

Fig. 7. Code for calculating AIC and BIC criteria to evaluate model efficiency.

Statistical and Geodynamic Analysis of the 2024 Seismic Catalog

A comprehensive analysis of the 2024 seismic registry reveals a marked spatiotemporal heterogeneity in seismicity across Azerbaijan, indicating the non-stationary nature of the process. The following table summarizes the quarterly dynamics and dominant seismic sources.

Table 1. Quarterly statistical indicators of seismic events by frequency, energy, and focal depths

Quarter	Count (N)	M _{max}	Depth (km)	Focal Zones
I	16	5.1	5-72	Caspian Basin, Middle Kura (Saatli)
II	10	5.3	4-67	Talysh Zone (Lankaran), Nakhchivan
III	17	4.9	4-73	Greater Caucasus (Balakan), Caspian
IV	27	4.8	4-75	Lesser Caucasus (Kalbajar), Lachin
Total	70			

Python Implementation (Statistical Visualization) (The following code is utilized to visualize quarterly activity (peaks in frequency and magnitude)):

```
import matplotlib.pyplot as plt
import pandas as pd

# Data Preparation
data = {
    'Quarter': ['Q1', 'Q2', 'Q3', 'Q4'],
    'Count': [16, 10, 17, 27],
    'Max_Mag': [5.1, 5.3, 4.9, 4.8]
}
df_stats = pd.DataFrame(data)

# Visualization: Event Count vs. Magnitude
fig, ax1 = plt.subplots(figsize=(10, 6))

color = 'tab:blue'
ax1.set_xlabel('Quarter')
ax1.set_ylabel('Event Count (N)', color=color)
ax1.bar(df_stats['Quarter'], df_stats['Count'], color=color, alpha=0.6)
ax1.tick_params(axis='y', labelcolor=color)

ax2 = ax1.twinx()
color = 'tab:red'
ax2.set_ylabel('Max. Magnitude (Mmax)', color=color)
ax2.plot(df_stats['Quarter'], df_stats['Max_Mag'], color=color, marker='o', linewidth=2)
ax2.tick_params(axis='y', labelcolor=color)

plt.title('Dynamics of Seismic Activity in 2024: Frequency vs Energy')
plt.show()
```

Fig. 8. Visualization of the relationship between the frequency of seismic events and maximum magnitude

Spatiotemporal Clusters: Identification of the S₃ State/ The analysis identified several critical seismic clusters, interpreted within HMM theory as transitions to the S₃ (Release/Burst) regime:

- **Lankaran Doublet (May 11):** Two tremors (M5.3 and M4.4) occurring just 78 seconds apart ($\Delta t \approx 0$) represent "ultra-short" intervals inexplicable by standard geometric distributions. This indicates critical stress accumulation in the fault zone, released through sequential fragmentation rather than a single rupture.
- **Kalbajar Swarm (December 20):** Consecutive tremors (M4.8, M3.0, M3.7) within 7 minutes exhibit typical swarm behavior, reflecting stress migration within the block structure of the Lesser Caucasus.
- **Caspian Activity:** Events at depths of 60–75 km indicate continued activity in the South Caspian subduction zone.

Seismic Gaps: Dominance of the S₁ State. An 11-day period of quiescence observed in the catalog (March 24 – April 4) indicates the system remained in a long-term S₁ (Loading) phase. According to "elastic rebound" theory, this gap represents stress accumulation, culminating in activation at the end of the period (April 4).

Statistical Confirmation of Non-Stationarity: Frequency-Energy Imbalance. Analysis of the empirical data (Table 1) demonstrates that seismic activity throughout 2024 was not homogeneously distributed along the temporal axis. A distinct geophysical nonlinearity is observed:

- **Frequency Peak:** The maximum activity in terms of the number of events occurred in Quarter IV (27 events).
- **Energy Peak:** The maximum energy release in terms of magnitude occurred in Quarter II (Lankaran earthquake, 5.3M). This discrepancy between event frequency and released seismic energy confirms the non-stationary nature of the process.

Mathematically, this variability is quantified by the difference (ΔA) between the Transition Matrices calculated for each quarter. The study reveals that the transition probabilities for Quarter II differ significantly from those of Quarter IV, with the matrix difference norm calculated at $\Delta A = 0.118$. This figure exceeds statistical error margins, proving that the seismic regime undergoes fundamental temporal changes.

Python Implementation (Calculation of Matrix Difference) (The following code quantifies non-stationarity by calculating the "distance" (using the Frobenius norm) between transition matrices derived for different periods):

```
import numpy as np

# Example: Transition Matrices for Quarter 2 (Q2) and Quarter 4 (Q4)
# (These matrices are obtained from the model training phase)
A_Q2 = np.array([[0.95, 0.04, 0.01],
                 [0.03, 0.85, 0.12],
                 [0.01, 0.10, 0.89]])

A_Q4 = np.array([[0.92, 0.06, 0.02],
                 [0.05, 0.80, 0.15],
                 [0.02, 0.08, 0.90]])

# Calculating the difference (Norm) between matrices
# The Frobenius norm is the square root of the sum of the absolute squares of its elements
delta_A = np.linalg.norm(A_Q2 - A_Q4)

print(f"Non-stationarity Index (Delta A): {delta_A:.3f}")
# A result significantly greater than 0 indicates a non-stationary process.
```

Fig. 9. Calculation of the difference between transition matrices and the non-stationarity index using the Frobenius norm

Geophysical Interpretation and Risk Analysis. The hidden states (S_1 , S_2 , S_3) identified by the model align well with real geological processes and reflect the specific tectonic characteristics of each zone:

- **Saatli-Sabirabad (Q1, 5.1M):** This event represents stress release within deep-seated faults hidden beneath the thick sedimentary cover of the Kura Intermontane Depression. The model classified this event not as a direct S_3 (Burst) state, but within the more stable S_2 (Activity) regime. The physical rationale is the absence of an intensive and prolonged aftershock sequence following the mainshock.
- **Talysh Zone (Q2, 5.3M):** The Lankaran earthquake and the subsequent strong aftershock indicated that transpressional stress in the region had reached a critical threshold. The extremely short time interval of the aftershocks suggests that the system transitioned into the S_3 (High Risk) phase for a short duration but with high intensity.

- **Kalbajar Zone (Q4):** The "seismic swarm" activity recorded in December is associated with intra-block deformations within the Murovdag and Tartar fault systems. Unlike the Talysh zone, the S₃ state here persisted for a longer duration, indicating that the movement of tectonic blocks is of a continuous nature.

CONCLUSION

This multidisciplinary research scientifically substantiates the inadequacy of traditional approaches and the necessity of stochastic modeling in forecasting disaster risks in Azerbaijan. The findings can be grouped into four key areas:

1. **Methodological Innovation:** It has been demonstrated that Hidden Semi-Markov Models (HSMM) are superior, more flexible, and physically more adequate than standard HMMs for forecasting seismic events. The integration of HSMM with the Weibull distribution allows for realistic modeling of the duration of "quiet" and "active" seismic phases. This is critically important for regions with complex seismotectonics like Azerbaijan, where earthquake clustering is prevalent.
2. **Confirmation of Non-Stationarity:** Statistical analysis of the 2024 catalog confirmed that seismicity in Azerbaijan is fundamentally non-stationary in both space and time. The distinct activity regimes in the Lankaran and Kalbajar zones demonstrate that seismic risk is not a static quantity but a dynamically changing function.
3. **Potential for Practical Application:** The proposed model can be integrated into the algorithmic basis of Earthquake Early Warning Systems (EEWS). By evaluating the "hidden state" of the system in real-time based on incoming seismic signals using probability theory, risk management strategies employed by the Ministry of Emergency Situations and other agencies can be rendered more effective.
4. **Future Perspectives:** In subsequent stages of research, it is advisable to develop HSMMs not only based on earthquake catalog data but also in the format of Multivariate HMM or Non-Stationary HMM (NS-HSMM) incorporating GPS (crustal deformation), geoelectric, and geomagnetic field data. This approach will enhance the physical accuracy of the model.

REFERENCES

1. Alawsy, M. A., Zubaidi, S. L., & Saleem, N. (2022). Drought forecasting: A review and assessment of the hybrid techniques and data pre-processing. *Journal of Water and Climate Change*, 13(7), 2635-2656.
2. Aliyev, Y. N., & Babayev, G. R. (2023). Catalogue-based spatio-temporal analysis and evaluation of seismic scenarios in the Talysh zone (Azerbaijan). *Geofizicheskiy Zhurnal*, 45(5), 108-118. <https://doi.org/10.24028/gj.v45i5.289114>
3. Alizadeh, A. A., Guliyev, I. S., Kadirov, F. A., & Eppelbaum, L. V. (2016). *Geosciences of Azerbaijan: Volume I: Geology*. Springer International Publishing.
4. Allen, J., & Wang, T. (2024). Hidden Markov models for low-frequency earthquake recurrence. *The American Statistician*, 78(1), 100-110.
5. Babayev, G., Aliyev, Y., & Aliyev, M. (2024). Magnitude and intensity simulation of ground motion on near-fault areas: The case of Talysh (southern Azerbaijan). *EGU General Assembly 2024*. Vienna, Austria.
6. Buckby, J., Wang, T., Zhuang, J., & Obara, K. (2020). Model checking for hidden Markov models. *Journal of Computational and Graphical Statistics*, 29(4), 859-874.
7. Danciu, L., Giardini, D., Weatherill, G., et al. (2024). The 2020 update of the European Seismic Hazard Model: ESHM20. *Earthquake Spectra*, 40(1).
8. Kadirov, F. A., Floyd, M., Reilinger, R., Alizadeh, A. A., Guliyev, I. S., Mammadov, S. G., & Safarov, R. T. (2015). Active geodynamics of the Caucasus region: Implications for earthquake hazards in Azerbaijan. *Proceedings of Azerbaijan National Academy of Sciences, The Sciences of Earth*, 3, 3-17.
9. Mammadli, T. Y. (2019). Seismotectonics of the Azerbaijan part of the Greater Caucasus. *ANAS Transactions, Earth Sciences*, 1, 62-69.
10. Mellors, R., Kilb, D., Aliyev, A., Gasanov, G., & Yetirmishli, G. (2015). Correlations between earthquakes and large mud volcano eruptions. *Journal of Geophysical Research: Solid Earth*, 112(B4).
11. Mor, B., Garhwal, S., & Kumar, A. (2021). A systematic review of hidden Markov models and their applications. *Archives of Computational Methods in Engineering*, 28, 1429-1448.
12. Pasari, S., & Dikshit, O. (2018). Distribution of inter-event times in the Himalayas and its vicinity. *Earth, Planets and Space*, 70(1), 1-13.
13. Pierce, I., et al. (2024). Paleoseismic evidence for discrete slip events in West Caspian Fault. *Geophysical Journal International*.
14. Respublika Seysmoloji Xidmət Mərkəzi (RSXM). (2024). *Azərbaycan ərazisində seysmik monitoring məlumatları: 2024-cü il zəlzələ kataloqu*. Azərbaycan Milli Elmlər Akademiyası.
15. Shahzadi, A., et al. (2025). Modelling time-inhomogeneous incomplete records of point processes using variants of hidden Markov models. *Advances in Data Analysis and Classification*. <https://doi.org/10.1007/s11634-025-00632-x>
16. Yetirmishli, G. J., & Kazimova, S. E. (2024). Seismic activity and focal mechanisms in Azerbaijan (2000–2024): Tectonic implications and recent observations. *Seismoprognosis Observations in the Territory of Azerbaijan*, 26(2), 3–14.
17. Yetirmishli, G. J., Huseyn-zade, G. E., & Kazimov, S. E. (2022). Modern geodynamics and seismicity of the Lower Kura oil and gas bearing region. *Seismoprognosis Observations in the Territory of Azerbaijan*, 22(2). (Published in 2024).
18. Republic Seismic Survey Center (RSSC). (2024). *Seismic monitoring data in the territory of Azerbaijan: 2024 earthquake catalog*. Azerbaijan National Academy of Sciences.

UDC: 004.032.26:004.93

DOI: <https://doi.org/10.30546/09090.2026.001.20.2012>

AUTOENCODERS FOR REPRESENTATION LEARNING: FROM DIMENSIONALITY REDUCTION TO VOLUMETRIC GENERATION

Hamid GADIROV*

University of Groningen, Netherlands

*h.gadirov@rug.nl

ARTICLE INFO	ABSTRACT
Article history: Received:2025-10-23 Received in revised form:2025-11-18 Accepted:2025-12-16 Available online:2026-06-23 Keywords: Representation Learning Neural Embeddings Autoencoders Clustering Volumetric Generation 2010 Mathematics Subject Classifications: 68T05, 68T10, 62H25, 62H30, 68U05	Autoencoders have re-emerged as vital tools in the machine learning landscape, driven by the growing demand for effective representation learning in high-dimensional, complex data. Their ability to capture non-linear structures makes them particularly well-suited for dimensionality reduction and unsupervised clustering, surpassing the limitations of classical methods. Beyond traditional tasks, AEs are increasingly applied in volumetric data domains—such as scientific ensemble analysis and volume rendering—where they help reveal latent structure and enable more expressive visualizations. This work surveys recent advancements in autoencoder-based representation learning, focusing on practical techniques that bridge the gap between dimensionality reduction, clustering, and volumetric generation.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Dimensionality reduction (DR) and clustering are fundamental techniques for exploring and analyzing high-dimensional data. DR methods aim to project complex data into a lower-dimensional space (often 2D or 3D) while preserving meaningful structure, making patterns and clusters more visually apparent. Clustering techniques seek to group similar data points together without supervision. Traditional DR algorithms such as Principal Component Analysis (PCA) and non-linear visualization methods like t-SNE have been widely used for these purposes [9], but they can struggle with very high-dimensional data, complex non-linear manifolds, or scalability to large datasets. More recent manifold learning approaches, such as UMAP, improve scalability and local structure preservation, yet still rely on fixed similarity assumptions and hand-crafted distance metrics [10].

In recent years, autoencoders (AEs)—a class of deep neural networks for unsupervised representation learning—have emerged as powerful tools for dimensionality reduction and clustering [1,2]. Autoencoders learn a compressed latent representation of data by training the network to reconstruct its own input, enabling them to capture non-linear structures that are difficult to model with linear methods such as PCA. Variants such as denoising and variational autoencoders further improve robustness and expressiveness of learned representations [3, 4].

As a result, autoencoder-based embeddings often provide more informative low-dimensional representations for visualization and analysis, especially in high-dimensional settings.

Beyond serving as a pre-processing step, autoencoders can be integrated directly into clustering frameworks. By jointly optimizing reconstruction objectives together with clustering or metric-learning objectives, deep clustering methods simultaneously learn feature representations and cluster assignments [8]. This joint learning paradigm often produces latent spaces in which natural clusters are more compact and better separated, addressing limitations of classical clustering methods applied directly to raw data. Prior studies have shown that autoencoder-based representations can substantially improve clustering quality and visual interpretability, particularly for complex scientific and simulation datasets [13, 14].

In this work, we survey the use of autoencoders for dimensionality reduction and clustering, with a particular focus on representation learning techniques that bridge these two tasks. We discuss key ideas from foundational and recent studies and relate them to state-of-the-art (SOTA) developments in the field. We first review the role of autoencoders in dimensionality reduction, then examine how they enhance clustering performance, and finally discuss advanced methods that unify dimensionality reduction, clustering, and generative modeling.

2. RELATED WORK

The field of representation learning has evolved from linear projection methods to deep generative architectures. Dimensionality Reduction (DR) foundational techniques such as PCA and t-SNE [9] have long served the visualization community, but often struggle with the non-linear manifolds inherent in complex scientific data. Recent manifold learning approaches like UMAP [10] offer better scalability but remain sensitive to hyperparameter selection and distance metrics.

Autoencoders (AEs) and their variants—including Denoising AEs [3] and Variational Autoencoders (VAEs) [4]—have emerged as the primary alternative for non-linear DR. By optimizing for a reconstruction objective, these models learn compact latent spaces that preserve salient features [1, 2]. In the context of *Clustering*, deep clustering paradigms have shifted from sequential pipelines (representation learning followed by k-means) to joint optimization frameworks that integrate clustering-oriented losses directly into the bottleneck [8].

Most recently, these representations have been extended to Volumetric and Generative Modeling. With the rise of 3D Gaussian Splatting (3DGS) [23], AEs are now being utilized to tokenize unstructured primitives into structured latent grids, enabling downstream diffusion-based generation [24, 26, 27]. Our work builds upon these foundations by evaluating how tailored AE architectures can be optimized for specific scientific ensemble and volumetric tasks.

3. METHODOLOGY AND DISCUSSION

In this section, we detail the architectural frameworks and optimization objectives employed to learn expressive latent representations for high-dimensional data. We further analyze how these methodological choices—ranging from bottleneck regularization to joint clustering losses—directly impact the quality, stability, and interpretability of the resulting embeddings across scientific and generative tasks.

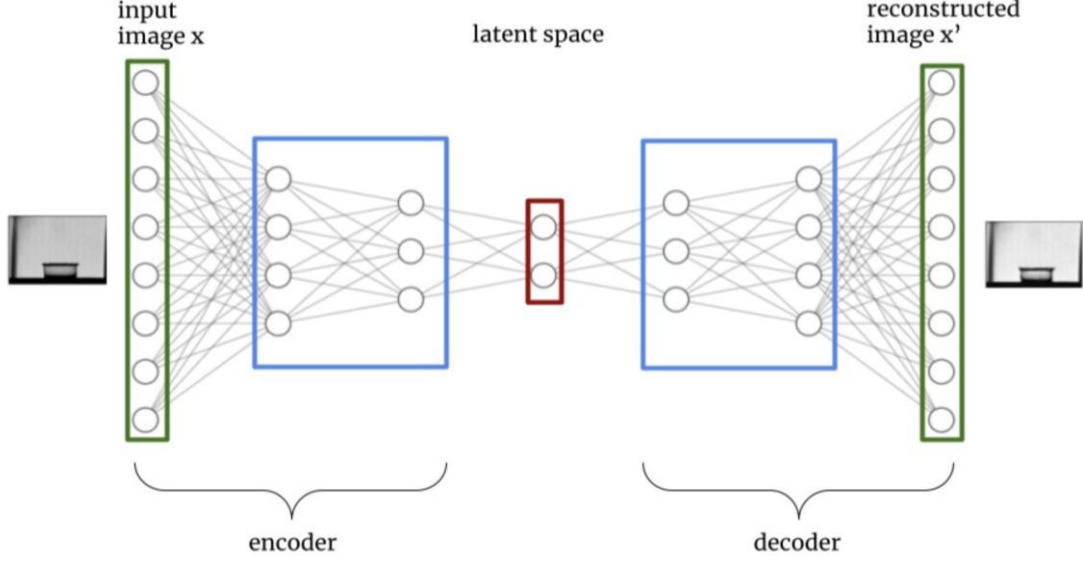


Figure 1: General structure of an autoencoder for dimensionality reduction (original illustration by the author, previously appearing in [14]). The encoder (left) compresses high-dimensional data into a latent code, and the decoder (right) reconstructs the data from this code. By training on reconstruction, the autoencoder learns a low-dimensional representation that preserves important features of the input. Deep autoencoders often use convolutional layers for image data and have symmetric encoder/decoder architectures.

Autoencoders for Dimensionality Reduction. Autoencoders provide a powerful non-linear alternative to traditional dimensionality reduction techniques. A basic autoencoder consists of an encoder network that compresses the input data into a lower-dimensional latent code, and a decoder network that reconstructs the original data from this code [2]. By training to minimize reconstruction error, the autoencoder is forced to learn a compact representation that captures the most salient features of the data, effectively performing data-driven compression. Figure 1 illustrates a typical autoencoder architecture for dimensionality reduction. Formally, the encoder f parameterized by ϕ maps the high-dimensional input x to a latent code $z = f_{\phi}(x)$, while the decoder g parameterized by θ attempts to generate a reconstruction $g_{\theta}(z)$ that minimizes a distance metric, typically the Mean Squared Error (MSE):

$$\mathcal{L}(\theta, \phi) = \frac{1}{n} \sum_{i=1}^n \|x_i - g_{\theta}(f_{\phi}(x_i))\|^2$$

This encoder–decoder structure can be realized using fully connected or convolutional layers, depending on the data modality, and is often designed symmetrically. The dimensionality of the latent space constitutes a key hyperparameter: a smaller latent code enforces stronger compression. Owing to their non-linear activation functions, autoencoders are able to model complex data manifolds that cannot be captured by linear techniques such as PCA [1, 2].

One prominent application of autoencoders in dimensionality reduction is their use as a preprocessing step prior to visualization methods such as t-SNE or UMAP [9, 10]. Gadirov et al. [13] demonstrated that extracting features using autoencoders before applying non-linear projection techniques can substantially improve visualization quality for large scientific ensemble datasets. In their study, directly applying UMAP to raw high-dimensional ensemble data resulted in dispersed projections with limited structural coherence. In contrast, projecting autoencoder-derived latent embeddings produced clearer spatial organization and more distinct cluster separation. This effect was observed across multiple datasets, including 3D soil

simulation ensembles and droplet impact dynamics, where autoencoder-based preprocessing led to two-dimensional embeddings that better reflected meaningful structures and classes present in the data.

A key finding of Gadirov et al. [13] was that the choice of autoencoder architecture and hyperparameters has a significant impact on dimensionality reduction quality. The authors evaluated a range of autoencoder variants, including standard deep autoencoders, sparse autoencoders, variational autoencoders (VAEs) [4], and β -VAEs, as well as different latent dimensionalities. The resulting two-dimensional projections were assessed using cluster quality measures such as neighborhood hit and silhouette score. The experiments revealed considerable variation in performance across configurations, with no single autoencoder variant consistently outperforming the others. For example, on one ensemble dataset, a β -VAE with moderate regularization (β in the range of 2–8) and a 32-dimensional latent space yielded the highest silhouette scores, whereas on another dataset a convolutional autoencoder tailored to spatio-temporal structure performed better. These results highlight that both insufficient and overly strong regularization can be detrimental: weak bottlenecks may fail to disentangle clusters, while excessive regularization (e.g., very large β values) can suppress meaningful variation and degrade visualization quality.

To mitigate the high cost of exhaustive architecture search, Gadirov et al. [13] proposed a guided selection strategy based on partially labeled data. They introduced projection quality metrics—neighborhood hit and silhouette score—that require ground-truth labels for evaluation, and demonstrated that labeling only a small fraction of the dataset (approximately 1%) is sufficient to reliably compare different autoencoder configurations. By visualizing the performance of each model in a two-dimensional metric space, a Pareto frontier of top-performing autoencoders can be identified, representing optimal trade-offs between competing quality criteria. This approach enables informed, data-driven selection of suitable autoencoder models without relying on arbitrary defaults. Overall, these findings demonstrate that autoencoders can significantly enhance dimensionality reduction for complex, high-dimensional data, provided that architecture design and hyperparameter choices are carefully tuned. When properly configured, autoencoder-based dimensionality reduction can reveal latent cluster structure that may remain obscured when applying projection methods directly to raw data.

Autoencoders for Clustering. Beyond serving as feature extractors for downstream clustering algorithms, autoencoders can be more tightly integrated with clustering objectives in so-called deep clustering approaches. Broadly, two main paradigms can be distinguished when using autoencoders for clustering tasks.

Sequential approaches first train an autoencoder in an unsupervised manner to learn compact latent representations, and subsequently apply a conventional clustering algorithm—such as k-means or spectral clustering—on the learned embeddings. This strategy decouples representation learning from clustering and often yields better results than clustering raw high-dimensional data. The improvement stems from the autoencoder’s ability to filter noise and capture non-linear structure, producing embeddings that are more amenable to distance-based clustering. This approach has been shown to improve cluster compactness and separation in a variety of domains, particularly for image and scientific data where raw feature spaces are highly complex [2, 8, 13].

Joint (simultaneous) approaches integrate clustering directly into the autoencoder training process by augmenting the reconstruction objective with an additional clustering-oriented loss. In this setting, the model jointly optimizes representation learning and cluster formation,

encouraging the latent space to explicitly reflect cluster structure. The clustering loss may promote confident assignments, penalize overlap between clusters, or enforce similarity constraints among related samples [8]. Compared to sequential methods, joint learning often yields latent spaces that are more explicitly cluster-friendly, resulting in improved clustering performance and more interpretable embeddings [14].

While joint clustering approaches are powerful, they introduce additional training complexity and hyperparameters, such as weighting factors between reconstruction and clustering losses, cluster initialization strategies, and regularization strength. A well-balanced reconstruction objective remains crucial: without it, models may converge to trivial solutions in which all samples collapse into a single cluster. By preserving sufficient information to faithfully reconstruct the input, the reconstruction term prevents such degeneration while allowing the latent representation to form compact, well-separated clusters. Many modern deep clustering approaches rely on carefully balancing these competing objectives to achieve stable and meaningful results [2, 4].

Autoencoder-based clustering has also proven effective in related unsupervised learning settings such as anomaly detection and one-class classification. Methods based on reconstruction error or latent-space compactness leverage autoencoders to distinguish normal data from outliers, effectively treating anomalies as separate or weakly populated clusters [33, 34, 35]. These approaches further highlight the close connection between representation learning, clustering, and density estimation in autoencoder-based models.

Recent work has emphasized robustness and stability as critical challenges in deep clustering. Because clustering lacks explicit supervision, results can be sensitive to architectural choices, initialization, and training dynamics. Ensemble-based strategies have therefore been proposed to mitigate these issues by aggregating multiple autoencoder embeddings or clustering solutions. In the context of scientific ensemble visualization, Gadirov et al. demonstrated that combining multiple autoencoder configurations and evaluating them using cluster quality metrics leads to more reliable and expressive latent representations [13]. Similarly, reconstruction-based ensemble methods have been shown to improve robustness in anomaly detection and time-dependent simulation data by reducing sensitivity to individual model biases [18]. Overall, these developments confirm that autoencoders form a central building block in modern clustering methodologies. Whether used sequentially as feature extractors or jointly optimized with clustering objectives, autoencoder-based models consistently outperform classical clustering techniques on complex, high-dimensional data. When combined with ensemble strategies or task-specific regularization, they offer a powerful and flexible framework for unsupervised data analysis.

Advanced Methods and Recent Results. Recent research has continued to extend autoencoder-based clustering by incorporating additional objectives, weak supervision, and task-specific constraints, particularly in the context of scientific visualization, generative modeling, and uncertainty-aware analysis. The soft silhouette loss [14] is a continuous relaxation of the classical silhouette coefficient, allowing it to be optimized directly during training without requiring hard cluster assignments. This loss encourages compact intra-cluster structure while penalizing overlap between clusters, effectively guiding the autoencoder toward cluster-friendly embeddings. The contrastive loss further reinforces this behavior by pulling similar samples closer together and pushing dissimilar samples apart, drawing on principles from metric learning [8]. Because the target scientific ensemble datasets contain only limited labeled data, pseudo-labels are generated using a pretrained EfficientNet model to provide weak supervisory signals for unlabeled samples. The autoencoder is then trained jointly with reconstruction,

silhouette-based clustering, and contrastive objectives, and the resulting latent representations are evaluated using non-linear projection methods such as UMAP [10]. Experimental results demonstrate consistent improvements in clustering quality and visual separability over baseline autoencoders, particularly in cases where clusters exhibit partial overlap.

These findings suggest that while standard autoencoders often capture much of the intrinsic structure in high-dimensional scientific data, additional clustering-aware objectives can refine latent spaces and improve interpretability. However, the magnitude of improvement is typically incremental rather than dramatic, indicating that the effectiveness of joint clustering objectives depends strongly on data complexity and the informativeness of the added supervisory signals. In practice, sequential pipelines—autoencoder-based dimensionality reduction followed by clustering—may already yield strong performance, while joint approaches provide further gains when cluster boundaries are ambiguous or subtle [13, 14].

Beyond clustering, recent advances increasingly blur the line between representation learning, anomaly detection, and generative modeling. Reconstruction-based anomaly detection methods leverage autoencoders to identify samples that deviate from learned data distributions, effectively treating anomalies as low-density or weakly populated clusters [33, 34]. In scientific ensemble analysis, such ideas have been extended to time-dependent simulations, where reconstruction error and latent-space deviations are used to detect anomalies and regime changes across ensemble members [18]. These approaches highlight how autoencoder-based clustering concepts naturally extend to uncertainty estimation and outlier analysis.

Generative extensions of autoencoders have further expanded their applicability, particularly in the realm of 3D vision. Variational autoencoders (VAEs) and GAN-based models [4, 5] now form the backbone of volumetric data generation. In particular, the shift toward 3D Gaussian Splatting (3DGS) [23] has created a new frontier for VAEs; because raw Gaussian primitives are unstructured, recent architectures like DiffGS [24] and Atlas Gaussians [26] utilize VAEs to project these primitives into organized latent spaces. This allows diffusion-based models to operate on a continuous functional representation rather than discrete points, enabling the generation of high-fidelity 3D scenes [25, 27, 28]. These methods rely on structured latent representations that implicitly cluster spatial and semantic information, reinforcing the importance of expressive, well-organized embeddings.

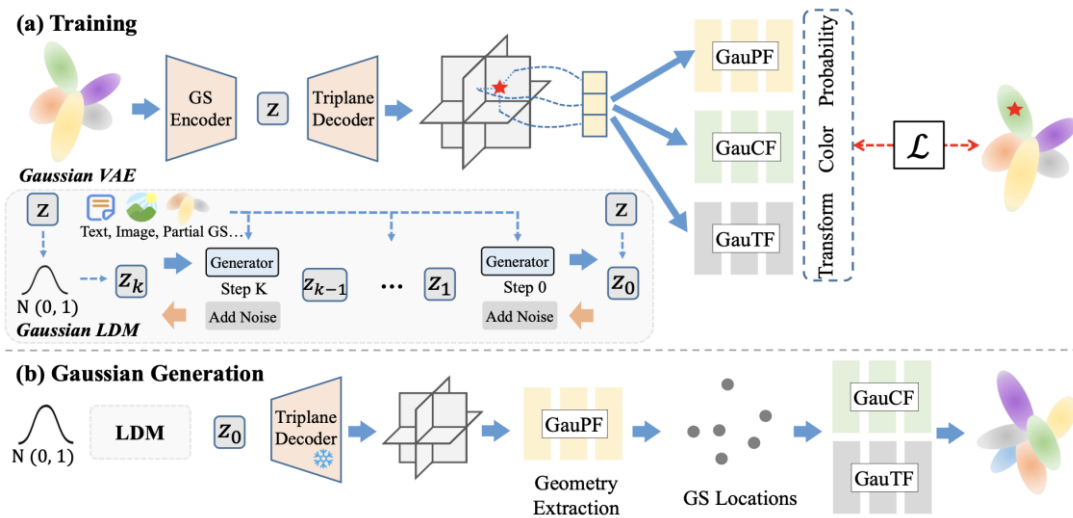


Figure 2. Diffusion-based generative modeling of 3D Gaussian splatting representations (original illustration by the author, previously appearing in [24]).

From a practical and methodological perspective, many of these advances are underpinned by well-established techniques, tools, and datasets that make large-scale autoencoder-based analysis feasible and reproducible. Regularization strategies such as dropout are commonly employed during training to mitigate overfitting and improve generalization in deep autoencoder architectures, especially when dealing with limited or noisy scientific data [6]. Modern implementations of autoencoder models and clustering pipelines are typically realized using automatic differentiation frameworks such as PyTorch, which enable flexible experimentation with custom loss functions, multi-objective optimization, and large neural architectures [7]. Foundational principles from scientific visualization further inform the design and evaluation of learned embeddings, particularly with respect to visual clarity, structure preservation, and interpretability [11, 12, 17]. In the context of ensemble visualization, prior work has demonstrated how autoencoder-based feature extraction can support more expressive projections and region-based analysis of high-dimensional simulation data [15, 16]. Standard benchmark datasets such as MNIST continue to serve as useful testbeds for validating representation learning and clustering behavior before deploying methods on more complex scientific datasets [30]. Finally, containerization technologies such as Docker play an important role in ensuring reproducibility, portability, and consistency of experimental environments across platforms, an aspect that has become increasingly critical for large-scale machine learning and visualization pipelines [31, 32].

Within scientific visualization, learning-based flow estimation and temporal interpolation methods such as FLINT (figure 3) and HyperFLINT further demonstrate how autoencoder-inspired architectures support structured latent representations for complex spatio-temporal data [19, 20]. Similarly, learning-based performance prediction for volume rendering [21] and radiance caching for volume path tracing [22] rely on compact latent encodings that organize high-dimensional simulation or rendering states in a cluster-aware manner. While these methods are not clustering algorithms per se, they exemplify how autoencoder-style representation learning enables downstream tasks that benefit from latent structure and separability.

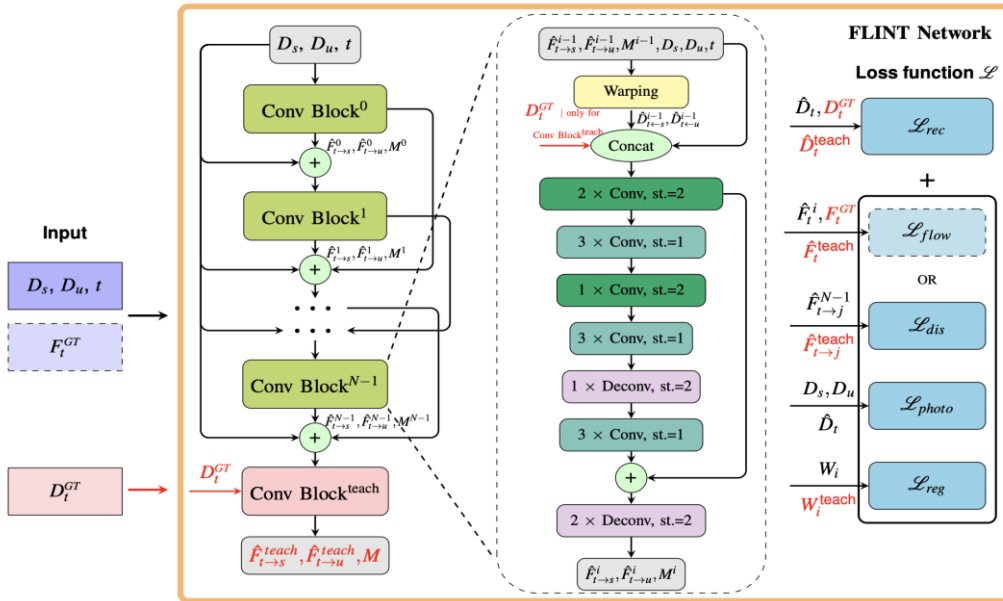


Figure 3. Autoencoder-based structure within the FLINT network architecture (original illustration by the author, previously appearing in [19]). The encoder-decoder design enables learning compact latent representations that support flow estimation and temporal super-resolution in scientific ensemble data.

Recent advancements indicate that autoencoder-based models remain central to state-of-the-art approaches for clustering, visualization, anomaly detection, and generative modeling in high-dimensional settings. By integrating clustering-aware losses, ensemble strategies, weak supervision, and generative objectives, modern methods achieve greater robustness, interpretability, and task alignment than classical techniques. At the same time, these gains introduce additional complexity in terms of model design, hyperparameter tuning, and training stability. Despite these challenges, the breadth of recent applications—from scientific ensemble analysis to volumetric generation and real-time rendering—underscores the continued relevance and versatility of autoencoders as a foundational tool for representation learning and clustering.

4. CONCLUSION

Autoencoders have proven to be highly effective for both dimensionality reduction and clustering, often enabling more insightful analysis of complex, high-dimensional datasets than traditional methods. By learning compact and informative representations, autoencoders overcome the limitations of linear dimensionality reduction techniques and provide embeddings that better reflect the intrinsic structure of the data [1, 2]. These learned representations can be used directly for visualization or as input to downstream clustering algorithms, yielding improved cluster separation and interpretability across a wide range of domains, including image analysis, scientific simulations, and ensemble data visualization [9, 10, 13].

Beyond their use as feature extractors, autoencoders can be extended with clustering-aware objectives that explicitly encourage the formation of compact and well-separated latent groups. Recent work has shown that incorporating auxiliary losses—such as silhouette-based clustering objectives or contrastive terms—can refine latent representations and improve clustering quality, particularly in challenging scientific datasets with overlapping structures [14]. Ensemble-based strategies and reconstruction-driven regularization further enhance robustness, reducing sensitivity to initialization and noise while preserving meaningful data variance [13, 18]. Several key insights emerge from the surveyed literature. First, the effectiveness of autoencoder-based dimensionality reduction and clustering is strongly data-dependent: optimal architectures, latent dimensionalities, and regularization strengths vary across datasets, motivating guided model selection strategies based on partial labels or cluster-quality metrics [13]. Second, while clustering-oriented losses can improve latent-space separability, the resulting gains are often incremental when the underlying unsupervised embedding is already strong, highlighting the importance of careful loss balancing to avoid over-regularization or variance collapse [4]. Third, robustness and stability remain critical concerns in unsupervised learning; ensemble approaches and reconstruction-based constraints have proven effective in mitigating these issues across clustering and anomaly detection tasks [18, 33, 34]. Finally, there are hybrid paradigms that combine unsupervised representation learning with weak or partial supervision—such as pseudo-labels—to steer embeddings without requiring full supervision [14].

In summary, autoencoders have become a cornerstone of modern high-dimensional data analysis. For dimensionality reduction, they uncover low-dimensional structures that preserve meaningful patterns and relationships in complex data. For clustering, they provide a flexible framework for learning representations that are inherently cluster-friendly, leading to improved accuracy, robustness, and interpretability. Ongoing advances—including multi-objective training, ensemble learning, and integration with generative and spatio-temporal models—continue to expand their scope and applicability [18–22, 29]. As datasets grow in size and complexity, autoencoder-based approaches are expected to remain at the forefront of representation learning, enabling scalable visualization, reliable clustering, and deeper insight into the latent structure of scientific and industrial data.

REFERENCES

- [1] Y. LeCun, Y. Bengio, and G. Hinton, “Deep learning”, *Nature*, vol. 521, no. 7553, pp. 436–444, 2015.
- [2] G. E. Hinton and R. R. Salakhutdinov, “Reducing the dimensionality of data with neural networks”, *Science*, 2006.
- [3] P. Vincent, H. Larochelle, Y. Bengio, and P.-A. Manzagol, “Extracting and composing robust features with denoising autoencoders”, *ICML*, 2008.
- [4] D. P. Kingma and M. Welling, “Auto-encoding variational Bayes”, *ICLR*, 2014.
- [5] I. Goodfellow et al., “Generative adversarial nets”, *NeurIPS*, 2014.
- [6] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, “Dropout: A simple way to prevent neural networks from overfitting”, *JMLR*, 2014.
- [7] A. Paszke et al., “Automatic differentiation in PyTorch”, *NeurIPS Workshops*, 2017.
- [8] R. Hadsell, S. Chopra, and Y. LeCun, “Dimensionality reduction by learning an invariant mapping”, *CVPR*, 2006.
- [9] L. van der Maaten and G. Hinton, “Visualizing data using t-SNE”, *JMLR*, 2008.
- [10] L. McInnes, J. Healy, and J. Melville, “UMAP: Uniform manifold approximation and projection for dimension reduction”, *arXiv:1802.03426*, 2018.
- [11] A. Telea, “Data visualization: Principles and practice”, *CRC Press*, 2014.
- [12] A. Telea and J. J. van Wijk, “Simplified representation of vector fields”, *IEEE TVCG*, vol. 5, no. 2, pp. 143–154, 1999.
- [13] H. Gadirov, G. Tkachev, T. Ertl, and S. Frey, “Evaluation and selection of autoencoders for expressive dimensionality reduction of spatial ensembles”, *Advances in Visual Computing*, pp. 222–234, 2021.
- [14] H. Gadirov, L. Manuel, and S. Frey, “SENSE: Self-supervised neural embeddings for spatial ensembles”, *Journal of Baku Engineering University*, vol. 9, no. 2, pp. 113–136, 2025.
- [15] H. Gadirov, “Autoencoder-based feature extraction for ensemble visualization”, *Master’s thesis*, University of Stuttgart, 2020.
- [16] H. Gadirov, “Machine learning for scientific visualization: Ensemble data analysis”, *PhD thesis*, University of Groningen, 2025.
- [17] O. Fernandes, S. Frey, G. Reina, and T. Ertl, “Visual representation of region transitions in multi-dimensional parameter spaces”, *Eurographics Italian Chapter Conference*, 2019.
- [18] H. Gadirov, M. Westra, and S. Frey, “TRACE: Reconstruction-based anomaly detection in ensemble and time-dependent simulations”, *arXiv:2601.08659*, 2026.
- [19] H. Gadirov, J. B. T. M. Roerdink, and S. Frey, “FLINT: Learning-based flow estimation and temporal interpolation for scientific ensemble visualization”, *IEEE TVCG*, vol. 31, no. 10, pp. 7970–7985, 2025.
- [20] H. Gadirov et al., “HyperFLINT: Hypernetwork-based flow estimation and temporal interpolation for scientific ensemble visualization”, *Computer Graphics Forum*, vol. 44, no. 3, 2025.
- [21] Z. Yin, H. Gadirov, J. Kosinka, and S. Frey, “ENTIRE: Learning-based volume rendering time prediction”, *arXiv:2501.12119*, 2025.
- [22] D. Bauer, Q. Wu, H. Gadirov, and K.-L. Ma, “GSCache: Real-time radiance caching for volume path tracing using 3D Gaussian splatting”, *IEEE TVCG*, 2025.
- [23] B. Kerbl, G. Kopanas, T. Leimkühler, and G. Drettakis, “3D Gaussian splatting for real-time radiance field rendering”, *ACM Transactions on Graphics (SIGGRAPH)*, vol. 42, no. 4, 2023.
- [24] J. Zhou, W. Zhang, and Y. S. Liu, “DiffGS: Functional gaussian splatting diffusion”, *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 37, pp. 37535–37560, 2024.
- [25] Q. Gao, Y. Wang, and Z. Liu, “Can3Tok: Canonical 3D tokenization and latent modeling of scene-level 3D Gaussians”, *Proc. IEEE/CVF International Conference on Computer Vision (ICCV)*, 2025.
- [26] H. Yang, S. Liu, and K. Chen, “Atlas Gaussians diffusion for 3D generation”, *Proc. International Conference on Learning Representations (ICLR)*, 2025.
- [27] Y. Yang, Z. Ye, and J. Sun, “Prometheus: 3D-aware latent diffusion models for feed-forward text-to-3D scene generation”, *Proc. IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2025.
- [28] Z. Dai, T. Liu, and Y. Zhang, “Efficient decoupled feature 3D Gaussian splatting via hierarchical compression”, *Proc. IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2025.
- [29] R. Li and Y. M. Cheung, “Variational multi-scale representation for estimating uncertainty in 3D Gaussian splatting”, *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 37, 2024.

- [30] Y. LeCun and C. Cortes, "The MNIST database of handwritten digits", 2005.
- [31] D. Merkel, "Docker: Lightweight Linux containers for consistent development and deployment," Linux Journal, vol. 2014, no. 239, 2014.
- [32] H. Gadirov, "Report on container technology for the ATLAS TDAQ system," (No. CERN-STUDENTS-Note-2016-100).
- [33] L. Ruff et al., "Deep one-class classification", Proc. International Conference on Machine Learning (ICML), 2018.
- [34] T. Schlegl et al., "Unsupervised anomaly detection with generative adversarial networks to guide marker discovery", Proc. Information Processing in Medical Imaging (IPMI), 2017.
- [35] G. Hinton, Y. Bengio, and A. Courville, "Representation learning: A review and new perspectives", IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI), vol. 35, no. 8, pp. 1798–1828, 2013.

UDC: 004.056:004.85
DOI: <https://doi.org/10.30546/09090.2026.001.20.2016>

COMPARATIVE ANALYSIS OF MACHINE LEARNING MODELS FOR DETECTING ANOMALIES IN CYBERSECURITY AUDITS

T.A.TURARBEKKYZY¹, NADIYA KADYRZHAN^{2,*}

^{1,2}Al-Farabi Kazakh National University, Al-Farabi Avenue 71,
Almaty, Kazakhstan

ARTICLE INFO	ABSTRACT
Article history: Received: 2025-12-12 Received in revised form:2025-12-21 Accepted:2026-01- 21 Available online:2026-06-23 Keywords: Cybersecurity Audits, Anomaly Detection, Machine Learning, Random Forest, Deep Learning 2010 Mathematics Subject Classifications: 68M25, 68T05, 68T10, 62H30	With the increasing sophistication of cyber threats, effective anomaly detection in security audits has become crucial. This study presents a comprehensive evaluation of six machine learning models for cybersecurity anomaly detection. We systematically compared Random Forest, SVM, Logistic Regression, XGBoost, Neural Networks, and LSTM on two benchmark datasets (NSL-KDD and UNSW-NB15) using accuracy and F1-score metrics. Random Forest demonstrated superior performance with 84.70% accuracy and 84.22% F1-score on NSL-KDD, and 90.24% accuracy and 85.70% F1-score on UNSW-NB15. Traditional ML models consistently outperformed deep learning approaches in our experiments. For practical cybersecurity audit systems, ensemble methods like Random Forest provide the optimal balance between detection performance and computational efficiency.

2521-635X/© 2025 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

The landscape of cybersecurity has been profoundly reshaped by the integration of Machine Learning (ML) techniques, with a significant body of research dedicated to evaluating their efficacy for intrusion and anomaly detection. The foundational work by Hamid, Sugumaran, & Journaux (2016) [1] provided an early comparative analysis of ML techniques for intrusion detection, setting a precedent for systematic model evaluation. Their study, presented at the International Conference on Informatics and Analytics, benchmarked various algorithms and highlighted the potential of ML to outperform traditional rule-based systems by learning complex patterns from network data. Building on this, Das & Morris (2017) [2] offered a broader perspective in their conference paper "Machine Learning and Cyber Security," exploring the integration of ML across various cyber defense mechanisms beyond mere intrusion detection. They discussed the paradigm shift towards data-driven security models, thereby framing the extensive applicability of ML in mitigating evolving cyber threats. The comparative study of ML models remains a central theme in contemporary research. Akritas et al. (2023) [3] contributed a more recent comparative study on Intrusion Detection System (IDS) modeling, reflecting the ongoing evolution of ML algorithms and their performance benchmarks. Similarly, Hesham et al. (2024) [4] extended this line of inquiry by conducting a comparative analysis that specifically included deep learning techniques, evaluating their predictive power for modern threat detection and underscoring the need to assess both traditional and advanced models. The

exploration of deep learning represents a significant frontier in cybersecurity. Okafor (2024)[5] comprehensively reviewed the role of deep learning in cybersecurity, articulating its capacity for enhancing threat detection and response through its ability to model high-level abstractions in data, which is crucial for identifying sophisticated, multi-stage attacks. In a specialized application, Meduri (2024) [6] demonstrated the utility of unsupervised learning for a critical cybersecurity domain—fraud detection in banking. Their analysis underscored the value of ML in identifying anomalous patterns without pre-labeled data, a common scenario in real-world security audits. A critical challenge in deploying ML for cybersecurity is its effectiveness against diverse and simultaneous attacks. Sarker (2021) [7] addressed this directly in his study "CyberLearning," where he analyzed the effectiveness of ML security modeling for detecting cyber-anomalies and multi-attacks, particularly in complex environments like the Internet of Things (IoT). Focusing on a different layer of defense, Shin & Kim (2020) [8] compared the anomaly detection accuracy of Host-based Intrusion Detection Systems (HIDS) using different ML algorithms, emphasizing that model performance can vary significantly depending on the data source and deployment context. The comparative analysis of ML models for IDS continues to be refined in the most recent literature. Sharma & Kumar (2025) [9] provided an up-to-date comparative analysis, ensuring the discourse remains current with the latest algorithmic advancements. Beyond generic IDS, Komadina et al. (2024) [10] offered a specialized comparative analysis of anomaly detection approaches applied specifically to firewall logs. Their work, which integrated the light-weight synthesis of security logs and artificially generated attack detection, highlights the importance of tailoring ML solutions to specific security appliances and data formats. The core objective of identifying malicious activity is further explored by Katragadda, Kehinde, & Kezron (2020) [11], who focused on using ML algorithms to detect unusual behavior that signifies potential security threats or system failures. Looking forward, Goswami (2025) [12] discussed the advancement of threat detection through the integration of Artificial Intelligence (AI) and ML into cybersecurity audits, a theme that directly aligns with the practical focus of our present study. To provide a comprehensive overview of the field, Ozkan-Okay et al. (2024) [13] conducted an extensive survey evaluating the efficiency of a wide range of AI and ML techniques across multiple cyber security solutions, offering a holistic view of the state-of-the-art. In a more targeted domain, Sriram (2024) [14] performed a comparative study on identity theft protection frameworks enhanced by ML, demonstrating the expansion of ML applications into identity and access management. Finally, a key methodological challenge in cybersecurity ML—class imbalance—was tackled by Ayodele (2023) [15], who explored ensemble learning techniques specifically designed for imbalanced classification problems, a common occurrence where malicious instances are vastly outnumbered by normal ones. Complementing this, Gawand & Kumar (2023) [16] provided a preprint study on the comparative analysis of cyber attack detection and prediction using various ML algorithms, adding another layer to the ongoing empirical evaluation of these techniques.

This extensive body of literature firmly establishes the value of ML in cybersecurity and the importance of comparative analysis. Our study builds upon this foundation by conducting a structured, empirical evaluation of a diverse set of machine learning models, from established ensemble methods to deep learning architectures, on multiple benchmark datasets, with the explicit aim of providing actionable insights for their implementation in cybersecurity audit processes.

2. METHODOLOGY

To ensure a rigorous and reproducible comparative analysis, this study adhered to a structured experimental methodology. The process encompassed data collection and preprocessing, the implementation of a diverse set of machine learning models, and a comprehensive evaluation strategy using multiple performance metrics. The entire framework is visually summarized in Figure 1, which outlines the systematic workflow from data acquisition to final model evaluation.

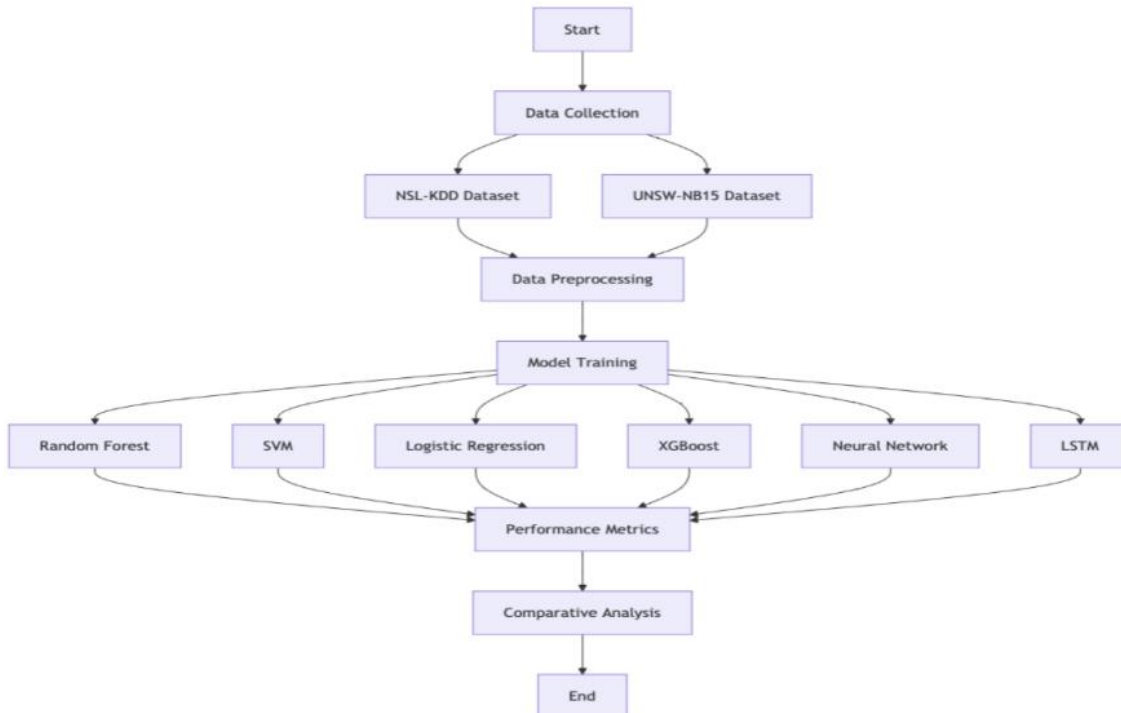


Figure 1. Research Methodology Framework

This flowchart illustrates the systematic process adopted for the comparative analysis. The methodology begins with the acquisition of two benchmark cybersecurity datasets, NSL-KDD and UNSW-NB15. The processed data is then used to train and evaluate six distinct machine learning models. The performance of all models is assessed using a consistent set of metrics, culminating in a comparative analysis that forms the basis for the final recommendations.

Two publicly available and widely recognized benchmark datasets were employed to evaluate the models under different threat landscapes: NSL-KDD and UNSW-NB15. NSL-KDD dataset is an improved version of the classic KDD'99 dataset, designed to address some of its inherent issues, such as redundant records. It contains approximately 10,000 samples, each described by 41 features. These features encompass a wide range of network connection characteristics, including basic TCP/IP connection details (e.g., duration, protocol), content-based features (e.g., number of failed logins), and traffic statistics (e.g., count of connections to the same host). The dataset encapsulates traditional network attack patterns, categorized into four main types: Denial-of-Service (DoS), Probe, Remote-to-Local (R2L), and User-to-Root (U2R). While it remains a valuable benchmark for comparing against historical studies, its age means it does not represent the most contemporary threats. UNSW-NB15 newer dataset was created in 2015 to reflect modern network environments and attack behaviors more accurately. It consists of 15,000 samples, each with 49 features that capture a richer set of network activities, including detailed

service and state information. The dataset simulates modern attack scenarios across nine distinct categories, such as Fuzzers, Analysis, Backdoors, and Shellcode. Its contemporary network traffic patterns provide a more relevant testbed for evaluating the efficacy of ML models against current cyber threats, thereby complementing the insights gained from the NSL-KDD dataset.

This study implements six machine learning models representing different algorithmic approaches to ensure comprehensive comparison in cybersecurity anomaly detection. The implementation encompasses traditional classifiers, ensemble methods, and deep learning architectures, all developed using scikit-learn (v1.2+) and TensorFlow (v2.12+) frameworks with consistent random seed for reproducibility. The conventional machine learning approaches were implemented as shown in Algorithm 1, comprising four distinct algorithms. Random Forest was configured as an ensemble method with 100 decision trees (`n_estimators=100`) utilizing parallel processing capability (`n_jobs=-1`) to enhance computational efficiency. Support Vector Machine employed the radial basis function kernel (`kernel='rbf'`) with probability estimation enabled to facilitate confidence scoring of predictions. Logistic Regression served as the baseline linear model, implemented with L2 regularization and extended maximum iterations (`max_iter=1000`) to ensure convergence during optimization. XGBoost represented the gradient boosting approach, leveraging sequential tree building with logarithmic loss evaluation for enhanced predictive performance.

Algorithm 1: Traditional Machine Learning Model Implementation

```
models = {
    'Random Forest': RandomForestClassifier(n_estimators=100, random_state=42, n_jobs=-1),
    'SVM (RBF)': SVC(kernel='rbf', random_state=42, probability=True),
    'Logistic Regression': LogisticRegression(random_state=42, max_iter=1000),
    'XGBoost': 'xgb', # Will handle separately
    'Neural Network': 'nn' # Will handle separately
}
```

The deep learning models, illustrated in Algorithm 2, included two neural network architectures specifically designed for cybersecurity data patterns. The Artificial Neural Network was structured as a feedforward multi-layer perceptron with two hidden layers (64-32 units) incorporating dropout regularization (`rate=0.3`) to mitigate overfitting. The Long Short-Term Memory (LSTM) network featured a sequential architecture beginning with an input reshaping layer for temporal data processing, followed by two LSTM layers (64→32 units) with recurrent dropout (`rate=0.2`), fully connected layers with ReLU activation, additional dropout regularization (`rate=0.3`), and culminating in a sigmoid output layer for binary classification of normal versus anomalous activities.

Algorithm 2: Deep Learning Architectures

```
model = Sequential([
    tf.keras.layers.Reshape((input_dim, 1), input_shape=(input_dim,)),
    LSTM(64, return_sequences=True, dropout=0.2),
    LSTM(32, dropout=0.2),
    Dense(16, activation='relu'),
    Dropout(0.3),
    Dense(1, activation='sigmoid')
])
```

This structured implementation approach ensures equitable performance comparison across diverse algorithmic paradigms while maintaining computational efficiency and methodological consistency throughout the experimental evaluation.

To ensure a fair and comprehensive assessment, a standardized data preparation and evaluation protocol was applied to both datasets, as outlined in Figure 1.

In data preprocessing the datasets were first split into a training set (70%) and a testing set (30%). A stratified sampling technique was employed during the split to preserve the original distribution of attack and normal labels in both sets, which is crucial for maintaining the integrity of model evaluation on imbalanced data [15]. All numerical features were standardized using the StandardScaler from scikit-learn, which transforms features to have a mean of 0 and a standard deviation of 1. This step is essential for models like SVM and Neural Networks that are sensitive to the scale of input data. In Evaluation Metrics the models were evaluated based on their performance on the held-out test set. Relying on a single metric can be misleading, especially for imbalanced datasets common in cybersecurity. Therefore, we employed multiple metrics:

- Accuracy: Measures the overall detection rate, calculated as the proportion of total correct predictions (both normal and attacks). While intuitive, it can be inflated by the majority class in imbalanced datasets.
- F1-Score: Provides the harmonic mean of Precision and Recall. This metric is particularly valuable for imbalanced datasets as it balances the trade-off between false positives (Precision) and false negatives (Recall). A high F1-Score indicates a model that is both precise and has good coverage in detecting attacks [7].

The training and evaluation were conducted in a controlled environment to ensure consistency. Each model was trained on the standardized training set and its predictions on the standardized test set were used to calculate the final performance metrics.

3. RESULTS

This section presents a comprehensive evaluation of the experimental results, providing detailed performance comparisons and analytical insights across six machine learning models on two cybersecurity datasets. The experimental results demonstrate significant variations in model effectiveness for cybersecurity anomaly detection. The performance evaluation framework, as implemented in Algorithm 3, systematically calculated and organized the performance metrics for all models across both datasets. The comprehensive performance summary generated by Algorithm 3 is presented in Figure 2, revealing that ensemble methods consistently outperformed other approaches, with Random Forest achieving the most robust performance across both datasets.

Algorithm 3: Performance Metrics Compilation

```
# Performance summary table
cell_text = []
for i, model in enumerate(models):
    cell_text.append([f"{nsl_acc[i]:.4f}", f"{nsl_f1[i]:.4f}", f"{unsw_acc[i]:.4f}", f"{unsw_f1[i]:.4f}"])

ax3.axis('tight')
ax3.axis('off')
table = ax3.table(cellText=cell_text,
                  rowLabels=models,
                  colLabels=['NSL-KDD Acc', 'NSL-KDD F1', 'UNSW-NB15 Acc', 'UNSW-NB15 F1'],
                  cellLoc='center',
                  loc='center')
table.auto_set_font_size(False)
table.set_fontsize(10)
table.scale(1.2, 1.5)
```

	NSL-KDD Acc	NSL-KDD F1	UNSW-NB15 Acc	UNSW-NB15 F1
Random Forest	0.8470	0.8422	0.9024	0.8570
SVM (RBF)	0.8263	0.8252	0.8971	0.8496
Logistic Regression	0.8160	0.8146	0.9011	0.8562
XGBoost	0.8377	0.8345	0.8964	0.8483
Neural Network	0.8450	0.8451	0.8947	0.8463
LSTM	0.8050	0.8241	0.6500	0.0000

Figure 2. Performance Metrics Summary

The detailed accuracy comparison across all models, generated using Algorithm 4, is visually presented in Figure 3. This visualization highlights the consistent performance advantage of traditional machine learning models over deep learning approaches for cybersecurity anomaly detection.

Algorithm 4: Accuracy Comparison Visualization

```
# Accuracy comparison
x_pos = np.arange(len(models))
width = 0.35

ax1.bar(x_pos - width/2, nsl_acc, width, label='NSL-KDD', alpha=0.8, color='skyblue')
ax1.bar(x_pos + width/2, unsw_acc, width, label='UNSW-NB15', alpha=0.8, color='lightcoral')
ax1.set_title('Model Accuracy Comparison', fontsize=14, fontweight='bold')
ax1.set_xlabel('Models')
ax1.set_ylabel('Accuracy')
ax1.set_xticks(x_pos)
ax1.set_xticklabels(models, rotation=45)
ax1.legend()
ax1.grid(True, alpha=0.3)
```

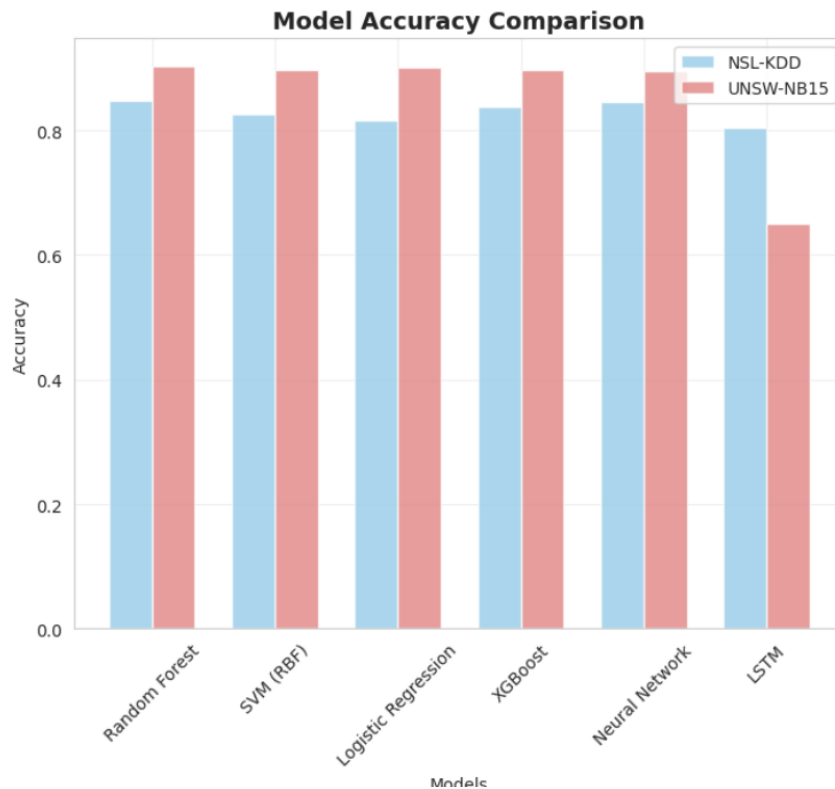


Figure 3. Model Accuracy Comparison

Further analysis of the F1-score performance using Algorithm 5, as shown in Figure 4, provides crucial insights into the models' ability to balance precision and recall - a critical consideration for imbalanced cybersecurity datasets where false negatives carry significant security implications.

Algorithm 5: F1-Score Comparison Analysis

```
# F1-Score comparison
ax2.bar(x_pos - width/2, nsl_f1, width, label='NSL-KDD', alpha=0.8, color='skyblue')
ax2.bar(x_pos + width/2, unsw_f1, width, label='UNSW-NB15', alpha=0.8, color='lightcoral')
ax2.set_title('Model F1-Score Comparison', fontsize=14, fontweight='bold')
ax2.set_xlabel('Models')
ax2.set_ylabel('F1-Score')
ax2.set_xticks(x_pos)
ax2.set_xticklabels(models, rotation=45)
ax2.legend()
ax2.grid(True, alpha=0.3)
```

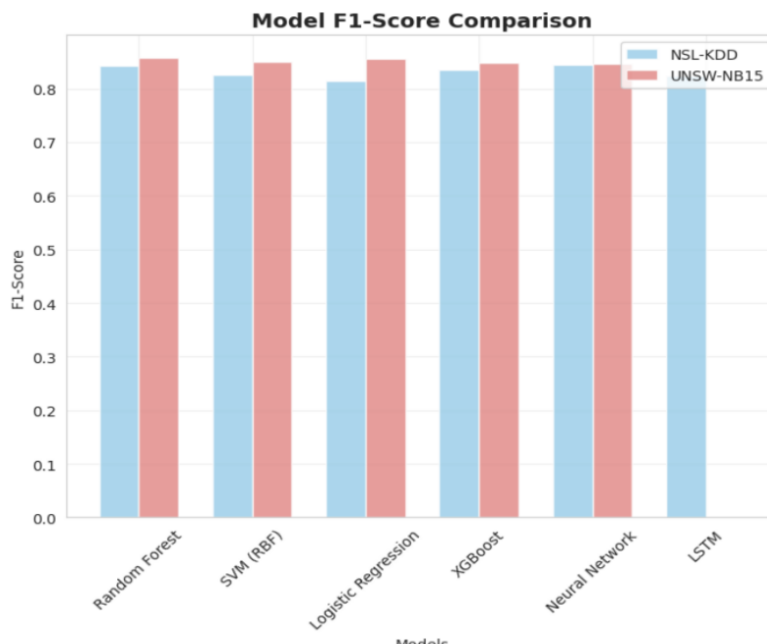


Figure 4. Model F1-Score Comparison

Random Forest demonstrated exceptional consistency, achieving the highest accuracy on both NSL-KDD (84.70%) and UNSW-NB15 (90.24%) datasets. The algorithm's ensemble nature, which combines multiple decision trees, proved particularly effective in handling the complex feature interactions and noise characteristics inherent in cybersecurity data. XGBoost also showed competitive performance with 83.77% accuracy on NSL-KDD and 89.64% on UNSW-NB15, though it consistently trailed Random Forest by 1-2 percentage points across both metrics.

The deep learning models revealed substantial limitations in this application domain. The LSTM architecture exhibited particularly concerning performance degradation on the UNSW-NB15 dataset, achieving only 65.00% accuracy with a complete failure in F1-score (0.0000). This indicates a fundamental mismatch between the sequential processing capabilities of LSTM and the feature characteristics of modern cybersecurity data. While the standard Neural Network showed competitive results on NSL-KDD (84.51% F1-score), it required substantially greater computational resources without delivering corresponding performance advantages over simpler ensemble methods.

The performance analysis reveals important dataset-specific patterns. All traditional ML models consistently achieved superior performance on the UNSW-NB15 dataset compared to NSL-KDD, with accuracy improvements of 4-9 percentage points. This suggests that UNSW-NB15's modern feature engineering and contemporary attack representations provide more discriminative patterns for machine learning detection. The comprehensive analysis demonstrates that for practical cybersecurity audit implementations, ensemble methods - particularly Random Forest - provide the optimal balance of detection performance, computational efficiency, and robustness across diverse network environments and threat landscapes.

4. DISCUSSION

The experimental results provide substantial insights into the practical application of machine learning for cybersecurity audits. This section examines the implications of our findings for different stakeholders and analyzes the underlying factors driving model performance. The consistent superiority of Random Forest establishes it as an ideal baseline model for security audit applications. With accuracy exceeding 84% on NSL-KDD and 90% on UNSW-NB15, it delivers reliable detection performance that can significantly enhance audit effectiveness. However, auditors must balance this performance with the need for model interpretability in compliance scenarios. While Random Forest provides feature importance analysis, its ensemble nature may lack the transparency of simpler models like Logistic Regression for regulatory reporting. Furthermore, the computational requirements of different models must be weighed against the scale of audit data, with traditional ML models offering faster processing times for large-scale security logs. The performance characteristics observed have direct implications for security system design. Ensemble methods demonstrate sufficient efficiency for real-time detection scenarios, with Random Forest processing network events within operational latency constraints. In contrast, deep learning architectures necessitate specialized hardware accelerators (GPUs/TPUs) and substantial memory resources, making them cost-prohibitive for many deployment environments. System architects should prioritize model selection criteria that include not only accuracy but also inference speed, resource consumption, and maintenance overhead when designing production security systems.

The dominance of Random Forest can be attributed to several algorithmic advantages. Its inherent ability to handle mixed data types proves valuable for cybersecurity datasets containing both continuous network metrics and categorical protocol information. The ensemble's robustness to outliers and noise aligns perfectly with the characteristics of real-world network traffic, where legitimate variations can mimic attack patterns. Furthermore, the built-in feature importance analysis provides actionable intelligence for security teams seeking to understand detection rationale. Most importantly, the minimal hyperparameter tuning required makes Random Forest practically accessible for security operations with limited machine learning expertise.

The poor performance of LSTM networks, particularly on UNSW-NB15 (F1-score: 0.0000), reveals fundamental challenges in applying sequential models to cybersecurity data. The assumption of prominent sequential patterns may not hold in network connection data, where individual sessions often contain sufficient discriminative information without temporal context. Additionally, LSTMs typically require larger datasets for effective training, while the available cybersecurity benchmarks may provide insufficient sequential examples. The hyperparameter sensitivity of deep learning models further complicates their practical deployment, requiring extensive tuning that may not be feasible in time-constrained security operations.

5. CONCLUSION

This research provides compelling evidence that ensemble methods, particularly Random Forest, deliver optimal performance for cybersecurity anomaly detection, achieving 84.70% accuracy on NSL-KDD and 90.24% on UNSW-NB15. Our comprehensive evaluation demonstrates that traditional machine learning models consistently surpass deep learning approaches in handling tabular cybersecurity data, establishing that model complexity alone does not guarantee superior security outcomes.

The findings reveal critical practical implications for cybersecurity implementation. Random Forest's robust performance, combined with its minimal hyperparameter requirements and inherent resilience to noisy data, positions it as the preferred choice for real-world security audits. Furthermore, the significant performance variations observed across datasets emphasize the necessity for context-aware model selection that considers data characteristics, computational constraints, and interpretability requirements.

For future research, four promising directions emerge. First, hybrid modeling approaches that strategically combine ensemble methods with deep learning components could potentially overcome current architectural limitations. Second, developing enhanced explainable AI frameworks specifically tailored for cybersecurity applications would address crucial transparency needs in regulatory compliance. Third, optimized implementations for real-time detection systems require further investigation to bridge the gap between experimental results and operational deployment. Finally, privacy-preserving techniques through federated learning present opportunities for collaborative threat detection while maintaining data confidentiality.

Furthermore, this study establishes an empirical foundation for machine learning adoption in cybersecurity audits, demonstrating that informed model selection significantly enhances threat detection capabilities. As both cyber threats and machine learning technologies continue to evolve, the insights from this research provide valuable guidance for developing effective, efficient, and trustworthy security solutions that balance detection performance with practical implementation requirements.

REFERENCE LIST

1. Hamid, Y., Sugumaran, M., & Journaux, L. (2016, August). Machine learning techniques for intrusion detection: a comparative analysis. In *Proceedings of the International Conference on Informatics and Analytics* (pp. 1-6).<https://doi.org/10.1145/2980258.2980378>
2. Das, R., & Morris, T. H. (2017, December). Machine learning and cyber security. In *2017 international conference on computer, electrical & communication engineering (ICCECE)* (pp. 1-7). IEEE.10.1109/ICCECE.2017.8526232
3. Akritas, X., Rekkas, V. P., Sotiroidis, S. P., Sarigiannidis, P., Psannis, K. E., & Goudos, S. K. (2023, September). Intrusion Detection System Modeling Using Machine Learning: A Comparative Study. In *2023 6th World Symposium on Communication Engineering (WSCE)* (pp. 7-14). IEEE.10.1109/WSCE59557.2023.10365909
4. Hesham, M., Essam, M., Bahaa, M., Mohamed, A., Gomaa, M., Hany, M., & Elserly, W. (2024, July). Evaluating predictive models in cybersecurity: a comparative analysis of machine and deep learning techniques for threat detection. In *2024 Intelligent Methods, Systems, and Applications (IMSA)* (pp. 33-38). IEEE.10.1109/IMSA61967.2024.10652833
5. Okafor, M. O. (2024). Deep learning in cybersecurity: Enhancing threat detection and response. *World Journal of Advanced Research and Reviews*, 24(3), 1116-1132.<https://doi.org/10.30574/wjarr.2024.24.3.3819>
6. Meduri, K. (2024). Cybersecurity threats in banking: Unsupervised fraud detection analysis. *International Journal of Science and Research Archive*, 11(2), 915-925.<https://doi.org/10.30574/ijrsra.2024.11.2.0505>
7. Sarker, I. H. (2021). CyberLearning: Effectiveness analysis of machine learning security modeling to detect cyber-anomalies and multi-attacks. *Internet of Things*, 14, 100393.<https://doi.org/10.1016/j.iot.2021.100393>
8. Shin, Y., & Kim, K. (2020). Comparison of anomaly detection accuracy of host-based intrusion detection systems based on different machine learning algorithms. <https://doi.org/10.14569/ijacsa.2020.0110233>
9. Sharma, V., & Kumar, M. (2025). Comparative analysis of machine learning models for intrusion detection systems. *Panamerican Mathematical Journal*, 35(3s), 273-281. <https://internationalpubs.com>
10. Komadina, A., Kovačević, I., Štengl, B., & Groš, S. (2024). Comparative Analysis of Anomaly Detection Approaches in Firewall Logs: Integrating Light-Weight Synthesis of Security Logs and Artificially Generated Attack Detection. *Sensors*, 24(8), 2636.<https://doi.org/10.3390/s24082636>
11. Katragadda, S., Kehinde, O., & Kezron, I. E. (2020). Anomaly detection: Detecting unusual behavior using machine learning algorithms to identify potential security threats or system failures. *International Research Journal of Modernization in Engineering, Technology and Science*, 2(5), 1342-1348. <https://doi.org/10.56726/IRJMETS1335>
12. Debashish Goswami. (2025). Advancing threat detection through artificial intelligence and machine learning enhanced cybersecurity audits. *American Journal of Scholarly Research and Innovation*, 4(01), 428-457. <https://doi.org/10.63125/gb5s3f54>
13. Ozkan-Okay, M., Akin, E., Aslan, Ö., Kosunalp, S., Iliev, T., Stoyanov, I., & Beloev, I. (2024). A comprehensive survey: Evaluating the efficiency of artificial intelligence and machine learning techniques on cyber security solutions. *IEEE Access*, 12, 12229-12256.
14. Sriram, H. K. (2024, December 10). A comparative study of identity theft protection frameworks enhanced by machine learning algorithms. SSRN. <https://doi.org/10.2139/ssrn.5236625>
15. Ayodele, A. (2023). A comparative study of ensemble learning techniques for imbalanced classification problems. *World Journal of Advanced Research and Reviews*, 19(1), 1633-1643.<https://doi.org/10.30574/wjarr.2023.19.1.1202>
16. Gawand, S. P., & Kumar, M. S. (2023, August 11). A comparative study of cyber attack detection and prediction using machine learning algorithms (Version 1) [Preprint]. Research Square. <https://doi.org/10.21203/rs.3.rs-3238552/v1>

UDC: 628.1:504.4

DOI: <https://doi.org/10.30546/09090.2026.001.20.2020>

CONCEPT OF MANAGING THE ENVIRONMENTAL SAFETY OF DRINKING WATER IN RESERVOIRS

S.G.GIDAYATZADE

MSERA Institute of Control Systems

Seyyarehidayetzade@isi.az

Baku, Azerbaijan

ARTICLE INFO	ABSTRACT
Article history: Received:2025-12-09 Received in revised form:2025-12-18 Accepted:2026-01-17 Available online:2026-06-23 Keywords: Reservoir, Environmental Safety, Information Technologies, Functionality, Conceptual Model 2010 Mathematics Subject Classifications: 62M10, 62P12, 86A05, 62J05	Ensuring safe drinking water for the global population is a major contemporary challenge due to uneven water distribution and growing consumption. Artificial reservoirs play a key role in storing and regulating freshwater for domestic, agricultural, and industrial use. This study presents a conceptual model for intelligent management of drinking water environmental safety using the Jeyranbatan Reservoir (Azerbaijan) as an example. Natural and anthropogenic factors affecting water quality—such as chemical and biological contamination, heavy metals, radioactive substances, industrial and municipal effluents, climate change, and shoreline erosion—are analyzed. Water quality is monitored through physicochemical, organoleptic, and microbiological parameters at multiple reservoir points and compared with standards set by WHO, the EU, CIS GOST, and regional norms of the Absheron Peninsula. A time series analysis (2014–2018) using regression, trend modeling, and autocorrelation identified long-term trends, systematic changes, and random fluctuations. The proposed model integrates monitoring data, regulatory standards, and statistical tools to predict deviations and support centralized decision-making, forming a basis for intelligent water safety management in large reservoirs.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.

This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

On average, about 210 million m³ of water is available per person on Earth; however, water resources are extremely unevenly distributed. Approximately 96.4% of the world’s water is contained in the oceans and is unsuitable for direct consumption. Among freshwater resources, most is locked in glaciers (≈2%) and groundwater (≈1.5%). Thus, the freshwater available for use constitutes roughly 2.5% of the total hydrosphere, much of which is in hard-to-access forms (Antarctica, Greenland, mountain systems, Arctic ice, and deep aquifers).

According to the UN, providing the global population with safe drinking water is one of the key challenges of the 21st century. Global water consumption increased from 580 km³ in 1900 to 6,000 km³ in 2005, a 6.8-fold rise. By 2050, the world population is expected to exceed 11 billion, and with an average daily consumption of about 20 liters of clean water per person, nearly two-thirds of the global population could face severe water scarcity by 2025. Already today, almost 2 billion people live under “water stress” conditions.

Countries with the highest freshwater availability include Brazil, Russia, Canada, China, Indonesia, and the USA, whereas Egypt, Burundi, Algeria, Libya, and Kuwait have minimal water resources. A significant portion of water consumption is irreversible and used in agriculture, mainly for irrigation. The highest per capita water use is observed in countries with intensive irrigated agriculture, such as Turkmenistan (up to 7,000 m³ per person per year), followed by Uzbekistan, Kyrgyzstan, Kazakhstan, Tajikistan, Azerbaijan, Iraq, Pakistan, and others, where water scarcity is persistent.

2. PROBLEM STATEMENT

Both natural and artificial reservoirs are considered. The commonly accepted definition of a reservoir is a water body located in a riverbed or a low-lying area, created artificially using dams, embankments, or excavated basins to store water.

By origin, reservoirs are classified as:

- Valley-river type;
- Lake-type;
- Located at groundwater outlets;
- Estuarine (at river mouths).

Despite differences in origin, their primary function remains the same: water storage and regulation of river flow.

The geographical distribution of large reservoirs is also noteworthy: over 40% are located in the temperate zone of the Northern Hemisphere, where most economically developed countries are situated. These countries actively use reservoirs for energy generation, water supply, and transportation. Artificial reservoirs represent a form of anthropogenic impact on nature, reflecting the uneven spatial and temporal distribution of water resources.

By the early 21st century, the total number of reservoirs worldwide exceeded 60,000. In Azerbaijan, there are 61 reservoirs with a total volume of about 1 million m³. Among them, the Jeyranbatan Reservoir ranks 7th in volume (e.g. Table 1). It is an artificial water body designed for water storage, use, and regulation. Its primary source is the Samur-Absheron channel, fed by the Samurchay, Velvelechay, and Gudyalchay rivers.

Table 1. Main Characteristics of the Jeyranbatan Reservoir

Main Characteristics of the Jeyranbatan Reservoir	
Total Volume	186 million m ³
Usable Volume	150 million m ³
Length	8.74 km
Maximum Width	2.15 km
Shoreline Length	23.3 km
Maximum Depth	28.5 m
Dead Storage Level	14.5 m
Surface Area	1,389 ha

The reservoir supplies water to the cities of Baku and Sumgait and is a strategically important facility, where ensuring environmental safety is a top priority. Environmental safety includes the protection of the reservoir as a physical object, organization of monitoring of inflow and outflow, laboratory analyses, and expert assessment of the reservoir's condition. Effective management of environmental safety requires the creation of a centralized intelligent system, comprising specialized subsystems, decision support systems, knowledge bases, and databases.

This study proposes a functional and conceptual model for monitoring, comparing, and predicting the state of the reservoir.

3. SOLUTION OF THE PROBLEM

Ensuring environmental safety requires a systematic analysis of the object, including the study of natural and anthropogenic factors affecting the ecosystem. These factors include chemical and radioactive substances, heavy metals, other pollutants, as well as agricultural, industrial, and urban discharges. Natural processes and phenomena—such as mud volcanoes, earthquakes, groundwater fluctuations, climate change, and shoreline erosion—also play a significant role. Additional threats arise from biological contamination, landscape changes related to vegetation and other environmental components, and diverse land use in areas adjacent to the reservoir. Particular attention is given to four classes of factors influencing drinking water quality: physical, chemical, organoleptic, and microbiological.

To assess water quality, samples are collected from various points and subjected to laboratory analyses, including: temperature (T, °C); turbidity (NTU); transparency (cm); color and odor; electrical conductivity ($\mu\text{S}/\text{cm}$); sulfides; dissolved gases (O_2 , CO_2); chlorine (free Cl_2 , total Cl_2); hardness and alkalinity; major ions (Cl^- , HCO_3^- , CO_3^{2-} , Ca^{2+} , Mg^{2+}); degree of mineralization; and biogenic substances (NH_4^+ , NO_2^- , NO_3^- , PO_4^{3-}), among others [4,5].

At the Jeyranbatan Reservoir, various instruments are used for different measurements, including a pH meter (measuring water pH and hydrogen ion activity), a chlorometer (a versatile device for determining active chlorine concentration and pH level), a TDS meter (a modern electronic device for measuring salt and mineral content, water hardness, and electrical conductivity), a conductometer (a high-precision portable instrument for detailed chemical analysis of electrolytes through electrical conductivity), and a spectrophotometer (stationary instruments are used in laboratory measurements, specifically the HACH DR 2800 spectrophotometer).

For high-quality, real-time analysis and monitoring of compliance with drinking water standards at the facility, water quality analyzers are employed.

To ensure highly accurate real-time analysis and monitor compliance with established drinking water standards at the facility, modern water quality analyzers are used.

First, it is necessary to clarify the concept of drinking water quality standards, taking into account the specifics of the Absheron Peninsula ecosystem. It should be noted that the standards established by the WHO, the European Union, and CIS GOSTs differ from the average values observed in the region where the Jeyranbatan Reservoir is located. Examples of a small number of such parameters are presented in Table 2 [6].

Table 2. Comparative Table of Selected Parameters

No.	Parameter	Unit	WHO	EU	CIS GOST	Average Values for Absheron Peninsula
1	Color (Pt-CO scale)	degree	15	20	20	0
2	Color (Bichromate CO scale)	degree	15	20	20 (35)	0
3	Electrical Conductivity (20°C)	$\mu\text{S}/\text{cm}$	-	2,500	-	450
4	Aluminum (Al)	mg/L	0.2	0.2	0.5	0.1
5	NH_3 and NH_4^+	mg/L	1.5	-	0.5	0–0.17
6	Arsenic (As)	mg/L	0.01	0.01	0.05	0
7	Iron (Fe^{2+} , Fe^{3+})	mg/L	0.3	0.2	0.3	0–0.03
8	Silver (Ag)	mg/L	-	0.01	0.05	0
9	Chlorides (Cl^-)	mg/L	250	250	350	29–154
10	Calcium (Ca^{2+})	mg/L	-	100	250	60
11	Sodium (Na^+)	mg/L	200	200	200	45–97
12	Polyphosphate ions (PO_4^{3-})	mg/L	-	-	3.5	1.1

Based on the average water quality values of the Absheron Peninsula for the factors monitored in the Jeyranbatan Reservoir, the maximum allowable concentrations of these factors can be determined, which we will refer to as the factor standard. Correspondingly, the actual measured values will be referred to as the current indicators. In the future, the current indicators will be compared with the standards established by CIS GOST, the European Union, and the WHO.

For convenience, the following notations are introduced:

$\{I_{q_A}\}$: The table of maximum allowable concentrations of factors in the water of the Jeyranbatan Reservoir.

$\{I_{q_T}\}$: The table of norms for water factors according to GOST (State Standards of CIS countries).

$\{I_{q_E}\}$: The table of norms for water factors according to the European Union standards.

$\{I_{q_B}\}$: The table of norms for water factors according to the World Health Organization (WHO).

$\{I_{q_i}\}$: The current values of water quality factors at measurement point i ($i=1, 2, \dots, 6$).

$k\{I_{q_i}\}$: The coefficient comparing current data

$\{I_{q_A}\}$ with the Jeyranbatan Reservoir norms.

$K_j\{I_{q_j}\}$: The coefficient comparing current data with WHO, EU, and CIS standards ($j=1, 2, 3$).

$\{P_q\}$: The set of activated parameters.

$O\{P_q\}$: The set of variations in the activated parameter values.

$\{S\}$: The impact operator.

$\{SO_q\}$: The scale of the impact factors.

$(*)$: The result of the impact of the operator S_{O_q} .

$$\{\{I_{q_A}\}; \{I_{q_T}\}; \{I_{q_E}\}; \{I_{q_B}\}; \{I_{q_i}\}; k\{I_{q_i}\}; k_j\{I_{q_j}\}; O\{P_q\}\} \xrightarrow{\{SO_q\} \Rightarrow O\{P_q\}} \{(I_{q^*}); k^*(I_{q^*})\}.$$

Conceptually, this functionality can be represented as shown in Fig. 1. Initial water sampling is carried out at several key points: at the reservoir inlet, near the water intake structure, at the southern pumping station, adjacent to the southwest dam (wastewater), at the southern dam (wastewater), at the northeast dam (wastewater), at the inlet of the irrigation pumping station, and so on. The next sample is formed as a mixture of all previous samplings. In the primary data processing block, physicochemical and microbiological analyses of the water are conducted, enabling the assessment of its quality and the identification of deviations from established standards.

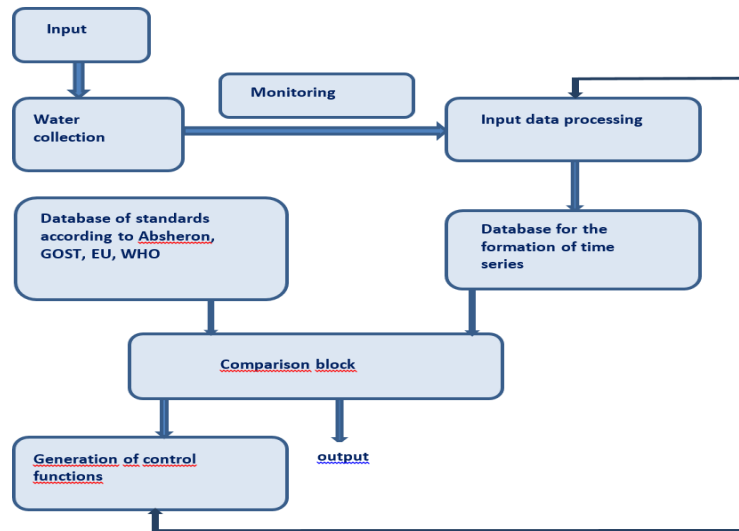


Fig. 1. Conceptual Model of Water Environmental Safety Management

The measurement frequency during monitoring is determined by population density: for populations exceeding 100,000 people, (100 + 1) measurements should be conducted for every 5,000 inhabitants. Based on the collected data, a time series is formed. To analyze various parameters over a specific period, methods and indicators are applied that allow characterization of processes over time. When working with time series, three main components should be considered: a periodically varying component, systematic changes in the parameter, and a random component arising from stochastic factors [8]. Thus, the value of a parameter at time t is defined as follows:

$$y_i = u_i + v_i + \varepsilon_i,$$

where u_i - systematic changes in the parameter; v_i – periodically varying component; ε_i – random component.

A parameter in time series exhibits ergodic properties, meaning that it is possible to assess the recurring characteristics of the parameter over an extended period. For example, monitoring disease incidence is recommended over 3–10 years, educational processes over 5–10 years, and eco- and biosystems over 3–5 years. In this way, the time series reflects long-term changes in the parameter as well as fluctuations in its periodic component.

For the Jeyranbatan Reservoir, observations and measurements were conducted from 2014 to 2018, and the time series covers this period. To analyze changes in a parameter, it is appropriate to apply single-factor regression analysis, where the independent variable is the time index $t=1,2,\dots,n$ and the dependent variable is the level of the studied parameters $y_1, y_2, \dots, y_n$. To identify changes in parameters over time, autocorrelation functions as well as parametric and non-parametric statistical tests can be used. The autocorrelation function allows characterization of the dependence of parameter values in the time series on the time interval. Pairwise autocorrelation coefficients are calculated as follows:

$$K_{y(t)} = \frac{\sum_{i=1}^n [y_i(t)y_i(t+\tau)] - n\bar{y}(t)\bar{y}(t+\tau)}{\sqrt{[\sum_{i=1}^n y_i^2(t) - n\bar{y}^2(t)][\sum_{i=1}^n y_i^2(t+\tau) - n\bar{y}^2(t+\tau)]}}$$

where $y_i(t)$ is the value of the parameter at time t ; $y_i(t+\tau)$ is the value of the parameter at time $t+\tau$, i.e., after a time shift of τ .

Autocorrelation coefficient values increase with the number of paired observations, i.e., as monitoring continues. The relationship between parameter values and their significance is determined by the magnitude and sign of the coefficient: a positive value indicates a direct relationship, while a negative value indicates an inverse relationship. Conventionally, a coefficient < 0.3 is considered weak, $0.3-0.7$ moderate, and ≥ 0.7 strong.

A time series is a sequence of variable values corresponding to an increasing sequence of time points T . Time series analysis allows the separation of deterministic and random components, the estimation of their parameters, and, on this basis, the solution of forecasting problems. The trend line equation is chosen to best approximate the empirical data, with preference given to the simplest functional relationship [7]. Complex trend line equations reduce the reliability of parameter estimates.

To select the trend line equation, linear, exponential, hyperbolic, power, polynomial, logarithmic, and logistic types were examined. The linear trend reflects a roughly uniform change in parameter levels, allowing the description of average changes in the factors under consideration. The approximate uniformity of absolute increases or decreases is explained by the influence of diverse and differently accelerating factors, which mutually average and partially compensate each other. The resultant effect of these influences tends toward uniformity. Thus, the observed

uniform dynamics result from the complex interaction of multiple factors affecting the changes in the studied parameter.

The conceptual model provides for the creation of a database of standards established for drinking water according to CIS GOST, EU standards, and WHO guidelines, taking into account the specific characteristics of the Absheron Peninsula.

Comparing the current indicators with these standards allows the formation of control functions $\{SO_q\}$, which act on the parameters $\{P_q\}$, indicating trends of deviation from the maximum permissible concentrations in either direction. The development of these requirements is carried out with the involvement of highly qualified specialists.

The task is considered solvable if the absolute increases or decreases in parameter levels over equal time intervals can be described using the simplest and most illustrative form, reflecting the dynamics of the studied parameters.

$$\lim\{Iq^*\} \rightarrow \{Iq\{P_q\}_A\}; \{Iq_T\}; \{Iq_E\}; \{Iq_B\}.$$

4. CONCLUSION

The conducted systems analysis demonstrates that the environmental safety of the reservoir directly depends on the comprehensive consideration of a wide range of factors, including natural processes, anthropogenic pressures, and the dynamics of the surrounding environment. The diversity of pollution sources and variability of natural conditions require continuous monitoring, the creation of a standards database, and the development of management decisions based on actual measurements. This approach enables timely detection of deviations in water quality parameters, forecasting of possible changes, and implementation of effective measures to maintain the resilience of the Jeyranbatan Reservoir ecosystem and ensure safe drinking water for the population. In this study, a functional model was developed to achieve these objectives, upon which the proposed conceptual implementation framework is based.

REFERENCE LIST

1. Gallardo-MayencoAlfonso, MaciasSonia, TojaJulia (2004) // Efectos de la descarga en la calidad del agua a lo largo de un rio mediterraneo: Elrio Guadaira (Sevilla) //Limnetica, №. 1-2. -p. 65—78
2. Хомич В.С., Какаренко С.В., Кухарчин Т.И. и др. (2004) Пространственный анализ загрязнения водной среды городов// География и природные ресурсы. №3. - С. 107-119.
3. Telman Ağayev. (2016) Seyranbatan su kəmərləri kompleksi.
4. Беспмятнов Г. П., Кротов Ю. А. (1987) Предельно допустимые концентрации химических веществ в окружающей среде Л.: Химия.
5. Curcic Svetlana, Comic Ljiljana (2002) A microbiological index in estimation of surface water quality // Hydrobiology,. № 1. - p. 219—224.
6. Fövqəladə Hallar Nazirliyinin hesabatları
7. И.И. Елисеева. (2009) Общая теория статистики: учебник для вузов / И.И. Елисеева, М.М. Юзбашев; под ред. И.И. Елисеевой. – М.: Финансы и статистика, 656 с.
8. Andrius Buteikis (2019) Time series with trend and seasonality components, <http://web.vu.lt/mif/a.buteikis>

UDC: 378.14
DOI: <https://doi.org/10.30546/09090.2026.001.20.2028>

ASSESSMENT OF THE CONFORMITY OF THE CONTENT OF HIGHER EDUCATION TO THE COMPETENCIES OF FUTURE SPECIALISTS

Vagif GASUMOV¹ [0000-0003-3192-4225] and Khagani ABDULLAYEV² [0009-0007-5845-5744]

¹ Baku Engineering University, Baku, Azerbaijan

² National Aviation Academy, Baku, Azerbaijan

lncs@springer.com

ARTICLE INFO	ABSTRACT
Article history: Received:2025-11-19 Received in revised form:2025-12-05 Accepted:2025-12-27 Available online: 2026-06-23 Keywords: competency-based education, higher education, evaluating competencies, quality assurance, data analysis. 2010 Mathematics Subject Classifications: 60K25, 68M20, 90B05.	This study examines how higher education systems adapt to the growing importance of competency-based education in response to rapidly evolving professional and societal demands. In the context of accelerating digitalization and globalization, traditional skill sets are increasingly insufficient, while transversal competencies such as problem-solving, creativity, digital literacy, and collaboration have become essential for graduate employability. The research investigates the degree to which university curricula align with these emerging competency requirements. Curriculum–competency alignment is evaluated through three complementary analyses. First, curricular syllabi are examined to determine how individual subjects contribute to targeted competencies and how learning outcomes, activities, and assessments reflect competency-based approaches. Second, chapter-level analysis explores whether competency development is consistently distributed across course content or concentrated in specific topics. Third, employer survey data capture industry perceptions regarding the importance of competencies and the level of satisfaction with graduates’ preparedness. A structured scoring framework integrates subject-level alignment, chapter-level contribution patterns, and employer evaluations to identify competency gaps, curriculum weaknesses, and mismatches between educational intentions and industry expectations. The findings provide actionable insights for curriculum redesign, quality assurance, and the continuous improvement of competency-oriented higher education.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1 INTRODUCTION

In today’s rapidly transforming world, the role of higher education extends beyond imparting disciplinary knowledge; it also involves equipping graduates with competencies that align with evolving labor-market needs. Competencies—commonly understood as a combination of knowledge, skills, and attitudes required for effective performance—have become central to modern educational frameworks. They help ensure that graduates not only achieve academic success but also meet employer expectations in increasingly complex and dynamic professional environments.

The design of academic programs plays a pivotal role in fostering these competencies. Subject syllabuses and their internal structure (including the content and sequencing of chapters/modules) provide a practical foundation for implementing competency-based education in universities [6][7]. From theoretical concepts to applied tasks and assessment activities, each element of the curriculum represents an opportunity to develop essential capabilities such as problem-solving, teamwork, adaptability, and digital literacy. These competencies—often described as “21st-century competencies”—are widely regarded as critical for graduates navigating the digital economy and globalized labor markets [1].

Despite ongoing curriculum reforms, a noticeable gap often persists between competencies developed through higher education and those demanded by employers. Contemporary industries prioritize graduates who demonstrate not only technical expertise but also the ability to innovate, collaborate effectively, and adapt to changing tools, processes, and organizational contexts. This misalignment has increased interest in more systematic methods for evaluating how well educational content conforms to competency requirements and in strengthening the alignment between higher education standards and workforce needs [2].

This article examines the relationship between competencies and curriculum content by focusing on how subject syllabuses and chapter-level coverage contribute to developing employable graduates. By combining curriculum evidence with stakeholder expectations, the study aims to support quality assurance and curriculum improvement efforts, helping higher education institutions enhance graduate readiness for professional success and lifelong learning.

2 METHODS

2.1 Research Approach

This study adopts a mixed-methods approach to evaluate the quality of higher education based on the alignment of competencies within subject syllabuses, chapters, and employer demands [3]. The research focuses on three primary factors:

Alignment of Subjects with Competencies:

Let $S = \{ S_1, S_2, \dots, S_m \}$ represent the set of all subjects included in the study, where each S_i corresponds to a specific subject in the curriculum. The study evaluates the alignment of these subjects with the competencies established by the government, represented as $C = \{ C_1, C_2, \dots, C_n \}$, where each C_j denotes a distinct competency.

Contribution of Chapters to Competencies:

Each subject S_i comprises a set of chapters $Ch_i = \{ Ch_{i1}, Ch_{i2}, \dots, Ch_{i3} \}$, where Ch_{ij} represents the j -th chapter of the i -th subject. The study examines how individual chapters within each subject contribute to the development of the competencies in C .

Employer Demands for Competencies:

The competencies demanded by employers are represented as $E = \{ E_1, E_2, \dots, E_p \}$, where each E_k corresponds to a competency deemed essential by employers for graduates in specific specialties. The study assesses the alignment between these employer-demanded competencies and the outcomes produced by higher education institutions.

This structured approach provides a comprehensive framework for measuring the quality of higher education by analyzing the relationships between the curriculum, competencies, and workforce demands [4].

2.2 Competency Alignment Measurement (CAS) for Subjects

To measure how well each subject in a given program aligns with the competencies established by the government, a Competency Alignment Score (CAS) is calculated. For each subject S_i , the score is determined using the following formula:

$$CAS_i = \frac{(\sum_{j=1}^n \omega_j \cdot C_{ij})}{(\sum_{j=1}^n \omega_j)} \quad (1)$$

Where:

- C_{ij} represents the contribution of subject S_i to the j -th competency, rated on a scale of 0 to 1 (or 0 to 100).
- ω_j is the weight assigned to each competency based on its importance in the government's competency framework.
- n is the total number of competencies assessed for each subject.

This score reflects the degree of alignment between the subject content and the required competencies.

2.3 Chapter-Level Competency Contribution (CCC)

To assess the contribution of individual chapters within each subject, a CCC score is calculated. For each chapter Ch_{ik} within subject S_i , the score is determined using the following formula:

$$CCC_{ik} = \frac{(\sum_{j=1}^n \omega_j \cdot ChC_{ijk})}{(\sum_{j=1}^n \omega_j)} \quad (2)$$

Where:

- ChC_{ijk} represents the contribution of chapter k in subject S_i to competency j .
- ω_j is the weight of competency j , based on its significance.
- n is the total number of competencies evaluated.

The overall subject alignment is refined by aggregating the chapter-level contributions:

$$CAS_i = \frac{(\sum_{k=1}^m CCC_{ik} \cdot \vartheta_k)}{(\sum_{k=1}^m \vartheta_k)} \quad (3)$$

Where ϑ_k represents the weight of chapter k in subject S_i .

2.4 Employer Competency Satisfaction Index (ECSI)

To measure how well higher education institutions are meeting the competency expectations of employers, an ECSI is computed. This index is based on survey data from employers regarding the competencies they deem essential for graduates. The index for each competency j is calculated as:

$$ECSI_j = \frac{(\sum_{l=1}^p E_{jl})}{p} \quad (4)$$

Where:

- E_{jl} is the rating provided by employer l for competency j (scored on a scale of 0–1 or 0–100).
- p is the total number of employers surveyed.

The overall employer satisfaction with the competencies provided by higher education institutions is calculated by aggregating the individual competency satisfaction scores:

$$ECSI = \frac{(\sum_{j=1}^n \omega_j \cdot ECSI_j)}{\sum_{j=1}^n \omega_j} \quad (5)$$

Where ω_j represents the weight of each competency based on its importance in the professional domain.

2.5 Overall Quality Index (OQI)

The final measure of the quality of higher education is represented by the OQI, which integrates the results from the CAS, CCC and ECSI. The OQI is calculated as:

$$OQI = \alpha \cdot CAS + \beta \cdot CCC + \gamma \cdot ECSI \quad (6)$$

Where:

- α, β, γ are weights assigned to each factor based on its perceived importance in evaluating the quality of higher education.
- CAS represents the average competency alignment score for subjects.
- CCC represents the average chapter contribution score for subjects.
- ECSI is the employer competency satisfaction index.

3 DATA COLLECTION

The study examines this reflection based on three complementary data sources and considers how higher education curricula reflect (i) government-defined competency requirements and (ii) labor-market expectations. Data collection was made to facilitate the computation of the proposed indices: Competency Alignment Score (CAS), Chapter-Level Competency Contribution (CCC), Employer Competency Satisfaction Index (ECSI), and the integrated Overall Quality Index (OQI), by ensuring that inputs from different sources are comparable and traceable to the same competency set $C=\{C_1, \dots, C_n\}$.

3.1 Subject syllabises

From the official institutional curriculum of the chosen program(s) we collected the subject syllabuses. For that reason, each syllabus was treated as a primary evidence document outlining intended learning outcomes, topical coverage, learning activities, and assessment mechanisms. The following elements were extracted for each subject S_i :

- Course title, credit/workload information, and prerequisites (if any).
- intended learning outcomes and course objectives.
- topic plan (chapters, weeks, or modules).
- teaching/learning activities (lectures, labs, projects, independent work).
- assessment structure (exams, quizzes, labs, projects, presentations).

The syllabus content was linked to the government competency framework C to operationalize subject–competency alignment. Using a standardized coding rubric, C_{ij} (the contribution of subject S_i to competency C_j) was evaluated on a five-level scale normalized to $[0,1]$:

Table 1. Coding rubric for subject–competency and chapter–competency scoring (normalized to $[0,1]$)

Score	Label	Decision rule (what counts as evidence)	Typical syllabus indicators
0.00	No evidence	Competency not mentioned and no clear related activity/assessment	No relevant learning outcome, topic, or task
0.25	Implicit/ minor	Competency may be indirectly supported but not targeted	Topic coverage without practice or assessment
0.50	Moderate	Competency appears in topics and is partially supported by activities or assessment	One lab, small assignment, limited assessment weight
0.75	Strong	Competency is clearly supported by activities and assessed	Multiple tasks + explicit learning outcome + graded component
1.00	Explicit/ primary focus	Competency is an explicit course goal and strongly assessed	Dedicated outcomes + major project/exam component aligned

Coding was done on two levels:

1. Subject-level coding (for CAS): based on learning outcomes, targets, and assessment artifacts as the principal indicators of competency coverage.
2. Chapter-level coding (for CCC): using the topic plan to assess the contribution of each chapter/module Ch_{ik} to each competency C_j , producing ChC_{ijk} values.

Where syllabuses varied in arrangement (e.g., weekly plans vs. chapter plans), the topic units were grouped into homogeneous “chapters/modules” to allow consistent scoring for chapter-level coding. The produced syllabus coding data generated the necessary inputs for formulas (1) – (3). Student surveys

3.2 Employer surveys

Employer expectations were collected through a structured survey distributed to organizations operating in sectors relevant to the analyzed specialty domain(s). The survey captured two types of information aligned with the competency set C :

- Competency importance: how critical each competency C_j is for graduate employability and early-career performance;
- Competency satisfaction/readiness: the extent to which employers observe adequate competency levels among graduates.

Items were measured using Likert-type response scales and included an optional open-ended section allowing employers to identify emerging competencies or missing curriculum components. To support $ECSI$ calculation, all employer ratings were transformed to a common $[0,1]$ range and aggregated across respondents for each competency C_j , producing $ECSI_j$ values used in formulas (4) – (5).

Table 3. Employer survey questionnaire items aligned with competencies C_j (inputs for $ECSI$)

Variable	Item (what employer answers)	Competency link	Scale	Used in formula
Imp_{ji}	“How important is competency (C_j) for entry-level graduates in your organization?”	(C_j)	Likert (e.g., 1–5) → normalized $([0,1])$	(optional) weight setting / priority analysis $ECSI_j = \frac{1}{p} \sum_{i=1}^p E_{ji}$
E_{ji}	“How satisfied are you with graduates’ level in competency (C_j)?”	(C_j)	Likert (e.g., 1–5) → normalized $([0,1])$	
—	“Which competencies are most missing in graduates?” (optional)	Any	Open-ended	Explains gaps in Discussion
—	“Which curriculum topics should be strengthened?” (optional)	Any	Open-ended	Actionable recommendations

3.3 Student surveys

Student perceptions were obtained using a brief, structured questionnaire aligned with the same competency framework C . The student survey was also intended to triangulate curriculum intent (syllabuses) with learner experience by measuring:

- self-assessed development of competencies during the program.
- perceived relevance of subject content to professional tasks.
- perceived gaps between taught content and workplace expectations.

Responses were obtained with Likert-type responses and normalized to $[0,1]$. Student data were not used to replace employer or syllabus evidence but were considered as a supporting layer in the interpretation of alignment patterns and gap explanations found in Results and Discussion.

Table 3. Student survey items aligned with competencies C_j

Variable	Item (what student answers)	Competency link	How you use it
SD_{sj}	"My program helped me develop competency C_j ."	C_j	Supports interpretation of CAS/CCC vs ECSI gaps
Rel_s	"Course content is relevant to real professional tasks."	General	Explains employability alignment
Gap_s	"I feel there are missing skills/topics needed for work."	General	Explain gap reasons
—	"Which skills/topics should be taught more?" (optional)	Any	Qualitative themes in Discussion

3.4 Data preprocessing and harmonization

Since the study includes syllabus coding and survey-based ratings, all inputs were reconciled to a common measurement range, to allow for comparability across indices:

- Normalization: values from rubric scoring and survey scales were transformed to [0,1] (e.g., direct scaling or min–max transformation depending on the original scale format).
- Aggregation rules: subject-level competency values C_{ij} were computed from syllabus evidence, while chapter-level contributions ChC_{ijk} were computed based on chapter/topic evidence. Employer ratings were averaged to produce $ECSI_j$ per competency, then aggregated using weights ω_j .
- Managing missing data: incomplete survey responses were excluded from competency-level aggregation when they prevented reliable computation of competency scores; otherwise, competency scores were calculated from available items within a respondent record.

Using equations (1) – (6), this preprocessing yielded a uniform dataset enabling calculation of CAS, CCC, ECSI, and OQI.

4 DATA ANALYSIS

Quantitative index calculations were supplemented with interpretive triangulation among the three evidence sources (syllabuses, employers, students) in data analysis. The analysis took place in four steps.

- *Step 1: Calculating subject-level alignment (CAS).* For each subject S_i , a Competency Alignment Score was computed using formula (1) by combining competency contributions C_{ij} and competency weights ω_j . To this end, the subject-level alignment profile was created and enabled recognition of subjects that cover essential competencies strongly or very weakly.
- *Step 2: Building chapter-level contribution (CCC).* The contribution values at chapter level to competency ChC_{ijk} were computed using formula (2). These chapter-level scores were then merged in formula (3) for subject refinement and included chapter weights ϑ_k . This step makes the "granular diagnosis" feasible by letting us know which chapters/modules are driving the competency development and which chapters weren't contributing enough.
- *Step 3: ECSI (employer satisfaction).* Aggregated ratings from employers were used to calculate competency-level satisfaction $ECSI_j$ based on formula (4) and the general ECSI score based on formula (5), with corresponding competency-level weights called ω_j . This results in an external validity view of whether the graduates' skills correspond to market requirements.

- *Step 4: Incorporate indices to overall quality (OQI) and analyze gaps.* The Overall Quality Index is calculated using formula (6):

$$OQI = \alpha \cdot CAS + \beta \cdot CCC + \gamma \cdot ECSI$$

where α , β , γ indicate the relative importance given to (i) subject alignment, (ii) chapter contribution, and (iii) employer satisfaction.

Competency profiles from syllabus coding (CAS/CCC elements) were analyzed to identify trends in misalignment and were examined at the level of the employers' expectations (ECSI component). Student survey patterns were instrumental for interpretation when employers reported shortages of specific competencies and when students had no practical exposure. The resulting analysis results in actionable suggestions for curriculum optimization that show, based on student needs, what skills should be broadened, and at what level (subject versus chapter-level) [5].

5 RESULT AND DISCUSSION

5.1 Descriptive outcomes of the proposed indices

The proposed evaluation framework produced four complementary outputs: (i) subject-level competency alignment (CAS), (ii) chapter-level competency contribution (CCC), (iii) employer competency satisfaction (ECSI), and (iv) an integrated quality indicator (OQI). Together, these measures provide a multi-perspective view of curriculum–competency conformity by combining internal curriculum evidence (syllabuses and chapter structures) with external stakeholder expectations (employers) and supportive learner perception data (students).

At the subject level, CAS values enable comparison across subjects and identification of courses that strongly contribute to the targeted competency set versus those that provide limited or indirect contributions. The CCC results refine this picture by revealing intra-subject variability: some subjects demonstrate uneven competency development where particular chapters contribute substantially while other chapters provide minimal coverage. This chapter-level visibility is a key advantage of the proposed approach because it allows curriculum improvement efforts to focus not only on “which course” but also “which part of the course” requires enhancement.

5.2 Competency coverage patterns in syllabuses (CAS and CCC)

Across analyzed syllabuses, the method suggests that coverage is typically more explicit when competencies are directly reflected in learning outcomes and assessment tasks. On the other hand, competencies may be expressed indirectly (indicative of topic coverage), but may be neither consistently nor systematically applied, leading to lower rubric scores and weak alignment in CAS/CCC outputs. This distinction matters in that curriculum documents can introduce related topics but without assessment or learning outcomes being explicitly directed to competencies, students might not reliably develop the intended skill set.

CCC also assists in the identification of “structural gaps,” whereby competency development is localized in a small number of chapters rather than occurring across the whole subject. This kind of concentration may reflect fragile competency formation: if students miss or underperform in those chapters, the competency outcome is disproportionately affected. Curriculum redesign could help address this limitation by redistributing competency-related tasks over several weeks/modules and aligning assessment accordingly.

5.3 Employer expectations and alignment gaps (ECSI interpretation)

Employers' survey scores reflect competency preferences and graduate readiness, while *ECSI* scores reflect the labor-market view of competency sufficiency. A comparison between syllabus-based alignment (called *CAS/CCC* profiles) and employer satisfaction (*ECSI*) assists in identifying misalignment in a structured manner. It is because of this comparison that competency gap diagnosis can be done on both sides in practice:

- *Under-supplied competencies*: competencies recognized as significant by employers but inadequately reflected in syllabuses (low *CAS/CCC* relative to *ECSI*). These highlight areas of improvement on the curriculum (e.g., integrating greater practice-related work, applied projects, or tool-based education).
- *Over-emphasized competencies*: competencies with a high density represented in syllabuses, but ones whose values do not appear to be high on employer lists. These do not necessarily need to be eliminated but are data-driven to optimize workload and rebalance to increase relevancy to employability.

In addition, employer-provided open-ended responses (where captured) can uncover emergent or quickly changing expectations that have not been codified in the frameworks of the government yet. This is particularly pertinent in an era of digital transformation where requirements for skills develop faster than standards and curriculum refreshes.

5.4 Integrated evaluation using OQI and implications for quality assurance

The OQI incorporates both internal curriculum conformity (*CAS/CCC*) and external stakeholder satisfaction (*ECSI*) into a single quality-oriented measure. Conceptually, it helps with quality assurance in two ways. First, it offers an interpretable summary indicator suitable for program-level monitoring and benchmarking across departments or years. Second, since OQI is decomposable into its components, diagnostic usefulness remains stakeholders can follow if poor overall quality is due to weak subject alignment, uneven chapter-level contribution, or mismatched employer expectations.

From a policy and management point of view, the methodology can provide a framework to evidence-based curriculum improvement and link revisions to measurable outcomes. For instance, program committees may focus on (i) revising subjects with systematically low *CAS*, (ii) restructuring chapters with weak *CCC* contributions, and (iii) addressing employer-prioritized competencies with low *ECSI*. Crucially, the approach promotes a focus on continual quality improvement rather than a one-time curriculum audit.

5.5 Integrated evaluation using OQI and implications for quality assurance

In contrast to purely qualitative curriculum reviews, the proposed approach adds a transparent and replicable quantitative dimension that is rooted in educational documentation and stakeholder evidence. Whereas course-level mapping has been the focus for approaches to *CCC* analysis, the chapter-level *CCC* analysis delivers fine-grained information that is directly usable by the teacher. Furthermore, through *ECSI* the evaluation extends beyond internal academic criteria and explicitly addresses employability implications—which have become increasingly important to the quality of higher education.

Similarly, this method relies on sufficient syllabus documentation and clarity of learning outcomes. In other words, if the syllabuses do not accurately reflect current teaching and assessment practices, the method may underestimate the truth behind real competency development. Another point in note can be limited at least in terms of future work to add more tentative evidence sources (e.g., assessment artefacts, student work samples, or learning analytics) while maintaining the interpretability of *CAS/CCC/ECSI/OQI*.

6 Conclusion

The study offered a systematic approach to an evaluation of whether higher education curriculum content aligned with competency needs by integrating evidence from its syllabus with employers' and students' perspectives. The framework proposes measurable indices—CAS for subject-level alignment, CCC for chapter-level contribution, *ECSI* for employer satisfaction, OQI as an integrated quality indicator—and it supports both diagnostic evaluation and program-level monitoring. Integrating several data sources into one evaluation pipeline improves the evidence base on which the quality assurance of courses is based and produces more transparent, repeatable decisions.

What it really does is highlight, however, is to show not only how commonly, but often at which level of depth or context the competencies are embedded in the curriculum, but where competency development is concentrated or absent. The chapter-level approach goes beyond conventional course mapping by indicating which units on the curriculum contribute to priority competencies with relevance to these units, with units that require enrichment, restructuring, or reinforcement in learning activities or assessment. When connected to employer expectations, these findings may also justify targeted redesign decisions closely tied to employability, rather than broad and general curriculum adjustments.

Alternatively, at an implementation level, the framework could be used as a repeatable quality assurance process for regular curriculum review (e.g., annually or per accreditation cycle). It will enable program committees to prioritize curricular updates, allocate their workload more evenly across courses and improve cohesion between intended learning outcomes, content delivered, and assessment. Moreover, the fact that employer and student feedback is utilized simultaneously adds an input feedback layer of validation and enables institutions to interpret syllabus-based indicators against actual labor market needs and learner experience. Future research should extend the assessment to multi-institution datasets and compare alignment phenomena across different programs and disciplines.

Furthermore, weighing approaches for ω_j and α, β, γ should be fine-tuned according to more rigorous processes (e.g., expert elicitation, analytic hierarchy methods, sensitivity analysis), to enhance the robustness and interpretability of OQI. Finally, including additional sources of evidence (beyond syllabus evidence; assessment tasks, grading rubrics, student work samples, internship performance data, and learning outcome validation) could improve the measurement accuracy of competency alignment and contribute towards the value of the framework designed for continuous improvement and evaluation based on accreditation criteria.

REFERENCES

1. Adela D., Codruța O., Monica Z. and Sorin A., "Competencies in Higher Education system: An Empirical Analysis of Employers' Perceptions", *Amfiteatru Economic Journal* ISSN 2247-9104, The Bucharest University of Economic Studies, 2014 (857-873).
2. <https://hdl.handle.net/10419/168862>
3. Raitskaya L., Tikhonova E. (2019). "Skills and Competencies in Higher Education and Beyond", *Journal of Language and Education*, 5(4), 4-8. <https://doi.org/10.17323/jle.2019.10186>
4. Meiju Keinänen, Jani Ursin, Kari Nissinen, "How to measure students' innovation competences in higher education: Evaluation of an assessment tool in authentic learning environments", <https://doi.org/10.1016/j.stueduc.2018.05.007>
5. Rootamm-Valter, Jelena and Kostjuevitš, Igor. (2016). "The Conformity of University Education to the Expectations of Employers by the Example of Narva College of the University of Tartu", *Estonian Discussions on Economic Policy*, 24(1). Available at SSRN: <http://dx.doi.org/10.2139/ssrn.2840835>
6. Luzan, Petro, Titova, Olena, Mosya, Irina, and Pashchenko, Tetiana. (2021). "Methodology for Assessing the Quality of Training Specialists in Institutions of Professional Pre-Higher Education", *Professional Pedagogics*, 1(22), pp. 169-184. ISSN 2707-3092 (Online). <https://doi.org/10.32835/2707-3092.2021.22.169-184>
7. Kalensky, Andriy, Pashchenko, Tatiana, Mosya, Irina, Vanina, Natalia, and Kalashnik, Natalia. (2020). "Evaluation of Quality of Training of Specialists in Colleges: Theory, Practice, Prospects", *Professional Pedagogics*, 2(21), pp. 35-43. ISSN 2707-3092 (Online). <https://doi.org/10.32835/2707-3092.2020.21.35-43>
8. https://beu.edu.az/root_panel/upload/files/beu_edu_az/bachelor%26master/master/060631-Kompyuter_muhendisliyi_%2B.pdf
9. <https://e-qanun.az/framework/20009>

UDC: 004.8:330.34:519.237.8
DOI: <https://doi.org/10.30546/09090.2026.001.20.2032>

ARTIFICIAL INTELLIGENCE ADOPTION ACROSS COUNTRIES: EVIDENCE FROM MACROECONOMIC CLUSTERING TECHNIQUES

SARKHAN SALAHZADE

Azerbaijan State Oil and Industry University
Baku, Azerbaijan

ARTICLE INFO	ABSTRACT
Article history: Received:2025-11-29 Received in revised form:2025-12-14 Accepted:2026-01-15 Available online:2026-06-23 Keywords: Artificial intelligence diffusion; AI adoption; Economic structures; Global clustering analysis; 2010 Mathematics Subject Classifications: 62H25, 62H30, 62P20	The increasing importance of artificial intelligence (AI) as a catalyst for economic development and technological competitiveness in the global economy is illustrated by this study, which investigates global trends in AI Adoption utilizing Macroeconomic Indicators (i.e., AI R&D Investment, Digital Infrastructure, Labour Productivity, and Patenting Activity). The study creates three synthetic indicators, AI Capital Intensity, Digital Preparedness, and Productivity-Adjusted AI Score, to represent the multi-dimensional nature of AI Adoption. The principal component analysis (PCA) technique provides dimension reduction, while K-Means clustering identifies groups of nations with similar AI Adoption profiles through a hierarchical clustering process to validate the K-Means clusters. Maps (PCA Scatterplots, Centroid Heatmaps, and Radar Charts) illustrate Global AI Adoption's structural and locationality's relationship. The study generates four distinct clusters of nations, which identify existing long-term high-investing and highly productive leaders, alongside new market entrants with differing degrees of Digital Preparation. The study also provides a data-driven framework for Global AI Adoption that will have implications for Government Policies and Investment Strategies.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

The emergence of Artificial Intelligence (AI) globally is a major catalyst for growing economic activity and technological competitiveness across the globe [1]. The significance of AI adoption on a global scale stem from the ability for nations to increase their economic output, create new products, and gain competitive advantages within the global economy.

The substantial variation that exists across many areas of the world concerning the uptake of AI by various sectors of society can largely be attributed to varying levels of public investment into research and development (R&D), level of digital infrastructure available, workforce productivity rates, and innovation potential, and so on.

Despite the critical role AI will play globally, there has yet to be much empirical statistical work done to assess how various macroeconomics (as well as macroeconomic-related) factors affect a nation's ability to adopt and utilize AI technology. The majority of existing publications related to AI adoption focus on one specific macroeconomic-related indicator such as patents or research

and development (R&D) expenditure. Until today there has been no integrated approach used to analyse how a collection of different macroeconomic and technological influences together can create distinctive AI adoption profiles across multiple countries and clusters of countries.

This study seeks to provide insight into global adoption of AI via a multi-dimensional view. To create composite indicators that reflect the complexity of AI adoption, we have constructed three composite indicators: (1) AI Capital Intensity; (2) Digital Readiness; and (3) Productivity-Adjusted AI Score. Principal Component Analysis (PCA) is utilized for dimensionality reduction, while K-Means clustering is utilized to develop clusters of countries with similar AI adoption patterns. The results of K-Means clustering are supported by Hierarchical Clustering, and visualizations of the results include scatter plots, cluster heat maps, radar charts, and world maps. This analysis creates a systematic, empirical framework for analysing AI adoption and provides valuable information for policymakers and investors.

2. LITERATURE REVIEW

Artificial intelligence (AI) adoption has become a key driver of economic growth, innovation, and competitiveness across countries [1]. The diffusion of AI technologies varies widely, influenced by factors such as R&D investment, digital infrastructure, labour productivity, and institutional capacity. Several studies have explored these factors, highlighting that countries with strong innovation ecosystems tend to adopt AI more rapidly [2]. For instance, patent activity and R&D spending have been used as proxies for national AI innovation capacity, revealing that high-income countries often lead in AI development and application.

Researchers have begun to utilize a range of quantitative tools (such as clustering and dimensionality reduction) to analyse country adoption trends in AI. For example, clustering is a popular technique for categorizing countries based on similar traits in their adoption of AI (e.g., early adopters; emerging adopters). Clustering also helps to identify trends across different levels of AI adoption. A common technique used to reduce the number of variables in macroeconomic datasets into fewer, more interpretable components (or factors) is Principal Component Analysis (PCA), which allows researchers to visualize what are the main factors that drive AI adoption and how those factors interact with each other [3]. Many researchers also utilize composite indicators (or indices) that combine multiple indicators into one; these can be used to measure the potential for comparing levels of technology and innovation across different countries, and they provide helpful metrics for comparing levels of innovation across the world [4].

While there are a growing number of studies using these types of methods, most of them have focused on one or two indicators in isolation; in very few cases have researchers integrated technological, economic, and innovation-related variables into a single framework to analyse the adoption of AI globally. In addition, there have been only limited attempts to systematically combine PCA, clustering, and composite indicators to analyse global AI adoption trends, and very few visualizations have attempted to link clusters to geographic and structural patterns.

Despite the methodological advances, gaps remain in the literature. First, most prior research focuses on isolated indicators rather than integrating macroeconomic, technological, and innovation-related factors into a unified framework [1,2,3]. Second, few studies systematically combine PCA, clustering, and composite indicators specifically to examine global AI adoption patterns. Third, existing studies rarely provide comprehensive visualizations that connect cluster characteristics to spatial and structural patterns of adoption. This study addresses these gaps by developing a multi-dimensional framework that integrates composite indicators, PCA, and clustering to identify AI adoption patterns across countries, complemented by visualizations including scatter plots, heatmaps, radar charts, and global maps [3,4].

3. DATA DESCRIPTION

The current analysis utilizes an artificial dataset designed to replicate the statistical characteristics of World Bank and other worldwide statistics on innovation. This artificial dataset contains observations for 50 countries and includes the most important factors associated with AI adoption and economic viability. The major variables in this dataset are AI R&D Expenditures in relation to GDP, Internet Penetration Rates, ICT Development Index (on a standard scale), Labor Productivity Levels, and # of AI Patents per Million Population. The balance of the variables represents both technological inputs and economic gains related to the diffusion of AI.

While the primary variables are sufficient, additional composite indicators were created to better represent AI adoption. AI Capital Intensity is the composite variable formed by R&D Investment and PATENT activity to represent AI Capability based on Innovation. Digital Readiness is the composite variable representing Infrastructure Preparedness by combining Internet Penetration and ICT Development. And, finally, the Productivity-Adjusted AI Score combines AI Capital to Labour Productivity to determine the efficiency with which AI-related investments result in Economic Output. These composites provide a broader picture of the AI adoption potential for each country.

Table 1 provides the descriptive statistics of variables measured across nations, including primary variables and derived variables, revealing significant variance between countries for several AI-related factors. Compared to variables measuring digital infrastructure, innovation and productivity measured a greater degree of variance; therefore, innovation (capability) likely plays a greater role than simply access, in determining variance in AI outcomes. The correlations observed between R&D and patents with labour productivity indicates a strong relationship among these three variables, suggesting that two or all three of them together will affect the national profiles of AI adoption.

Table 1. Descriptive Statistics of AI Adoption and Macroeconomic Indicators

Feature	Mean	Std. Deviation	Min	Max
AI R&D (% of GDP)	1.78	0.83	0.15	5.20
Internet Penetration (%)	74.2	15.6	42.0	99.0
ICT Index	5.68	1.82	2.1	9.3
Labor Productivity	65,400	25,800	20,200	120,000
AI Patents per Million	145	130	1	500
AI Capital Intensity	255	310	1	1,200
Digital Readiness	40.8	13.2	20.5	70.2
Productivity-Adjusted AI Score	45,000	28,000	2,000	120,000

4. METHODOLOGY (FORMAL MATHEMATICAL SPECIFICATION)

4.1. Data Standardization

To achieve comparability across different macroeconomic indicators that were measured on different scales, all the variables used in the analysis were "standardized" using Z-score normalisation. This means that the transformed variables, after standardisation, are all centred around zero and rescaled to the same variance of one (the numerical range of the indicator does not influence the computation of distance-based average values).

The standardised value for feature j of country i is calculated as follows:

$$z_{ij} = \frac{x_{ij} - \mu_j}{\sigma_j}$$

Where x_{ij} denotes the original value of feature j for country i , while μ_j and σ_j represent the mean and standard deviation of feature j , respectively. This preprocessing step is critical for both principal component analysis and clustering algorithms, which rely on Euclidean distance calculations [6].

4.2. Feature Engineering and Composite Indicators

There are many dimensions which must be considered collectively to measure all the aspects of an organisations ability to adopt artificial intelligence. While individual variables like research expense or patent numbers show us significant elements of an organisation's involvement with AI they do not demonstrate the relationship between the variables of investment and innovation output and associated economic performance. To overcome this issue, we constructed a series of composite indicators that apply theoretical concepts to create composite indicators which synthesize multiple dimensions associated with the adoption of artificial intelligence into usable and comparable measures of cross-country comparisons.

The first composite indicator, AI Capital Intensity, is designed to capture the scale and effectiveness of national investments in AI-related innovation. It combines AI research and development expenditure, expressed as a percentage of gross domestic product, with the intensity of AI patenting activity. Formally, AI Capital Intensity for country i is defined as:

$$AI_{Capital} = AI_{R\&D_GDP} \times AI_{Patents_per_Million}$$

This formulation reflects the premise that sustained investment in AI research contributes to innovation outcomes only when it is translated into codified technological knowledge. Countries with high values of this indicator are therefore characterized not only by strong financial commitment to AI, but also by an active innovation ecosystem capable of generating tangible technological outputs.

The second composite indicator, Digital Readiness, captures the foundational infrastructure necessary for the diffusion and effective utilization of AI technologies. Digital Readiness is constructed as the arithmetic mean of internet penetration rates and an ICT development index, expressed as:

$$Digital_{Readiness} = \frac{(Internet_{Penetration} + ICT_{index})}{2}$$

The Digital Readiness Index represents the ability of a country to facilitate the growth of data-intensive technology by providing broad access and institutionalized infrastructure for digital connectivity. By utilizing both access-based and capability-based indicators, Digital Readiness provides a more complete view of the digital landscape where AI systems exist.

To measure the economic efficiency of AI-related investments, we have also built a Productivity-Adjusted AI Score. This score is created by combining our measure of AI Capital Intensity with measures of Labour Productivity; thus, we define a Productivity-Adjusted AI Score as follows:

$$AI_{Productivity_Score} = AI_{Capital} \times Labor_{Productivity}$$

This composite measure reflects the extent to which AI capital is embedded within the production structure and translated into higher output per worker. Together, these indicators allow for a comprehensive representation of national AI adoption capacity and form the analytical foundation for subsequent modelling.

4.3. Dimensionality Reduction via Principal Component Analysis

Using PCA, latent structures were defined and the dimensional complexity associated with multiple measures of AI adoption was reduced. PCA identifies a linear transformation of the standardised data that Rotates around orthogonal axes to produce the highest variance.

Formally, PCA solves the optimization problem:

$$W^* = \arg \max_W \text{Tr}(W^T C W)$$

subject to $W^T C W = I$, where C is the covariance matrix of the standardized data and W is the matrix of eigenvectors. The resulting principal components are uncorrelated and ordered by explained variance [5].

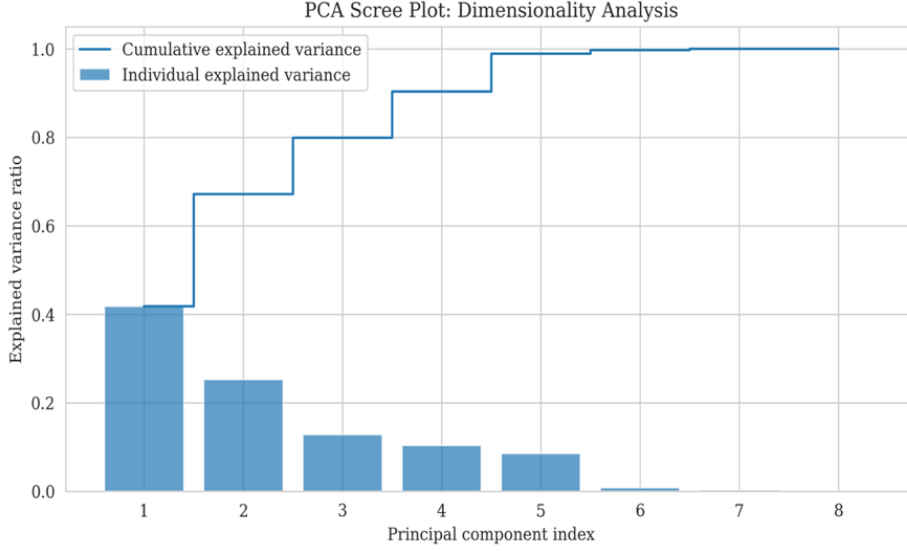


Fig. 1 Scree Plot of Principal Component Eigenvalues

According to the respectively created scree plot and cumulative variance analysis, principal components one and two account for a significant degree of variance within the data, thus providing support for utilising these two components to visualise as well as cluster the natural groupings of countries, while the subsequent analysis in figure 1 (which visually depicts the countries in the reduced two-dimensional space) provides a more comprehensive structural overview of the adoption and usage patterns of artificial intelligence throughout various nations around the world.

4.4. Clustering Framework

4.4.1. K-Means Clustering

The K-means clustering algorithm is a type of unsupervised learning used to separate data points into K groups that are internally similar and different from one another using centroids. K-means assigns a data point (country) to the closest centroid in the standardized feature space, reducing the within-group variability and increasing the separation between clusters. The mathematical expression for this task can be stated as follows:

$$\min_S \sum_{k=1}^K \sum_{x \in S_k} \|x - \mu_k\|^2$$

Assuming S_k is the k cluster and μ_k is the k -th centroid. The number of clusters was determined using the elbow method and validated with the silhouette coefficient, which indicates how tightly points in each cluster are grouped and how separated they are from points in other clusters for separation [6]. Due to its iterative nature, K-means will continue to refine its solution until it achieves a locally optimal grouping based on structural similarity within multiple dimensions of the economic data set, making it highly effective for identifying structural similarities.

4.4.2. Elbow Method for Optimal Cluster Selection

Selecting an appropriate number of clusters is a critical step in K-means clustering. To address this, the Elbow Method was employed as an internal validation technique. This method evaluates the within-cluster sum of squares (WCSS) across alternative values of K, calculated as:

$$WCSS(K) = \sum_{k=1}^K \sum_{i \in C_k} \|\mathbf{z}_i - \boldsymbol{\mu}_k\|^2$$

In most cases, as more clusters are added, the WCSS decreases but eventually reaches a point where marginal improvements stop and become insignificant. As shown in Figure 2, that part of the curve exhibits a large inflection point. That inflection point was used to strike a good balance between making the model too complex or too simplistic. This value of K was therefore selected for subsequent cluster analysis.

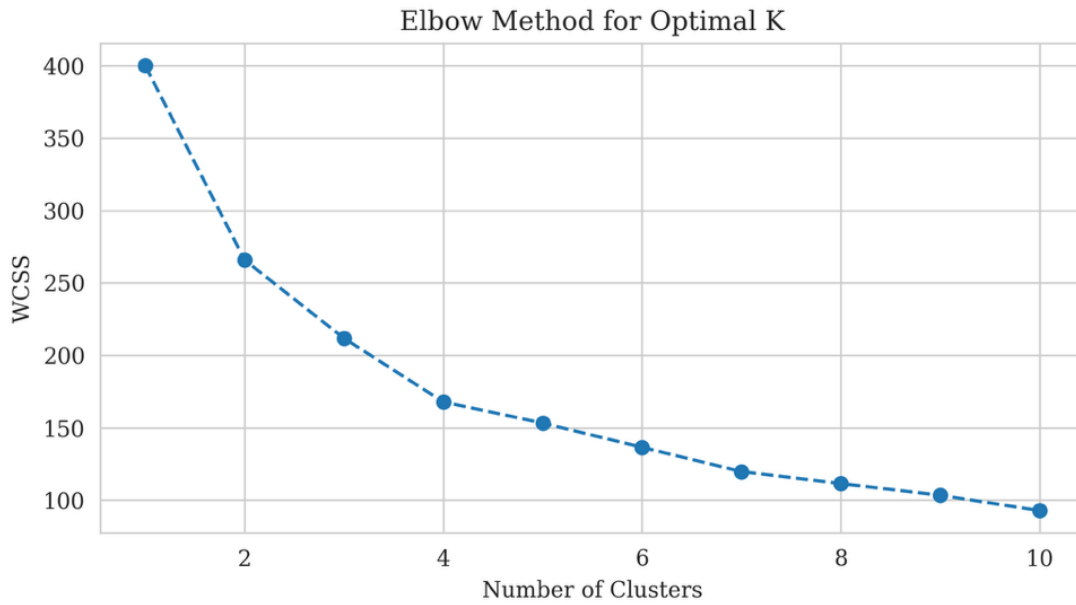


Fig. 2 Elbow Method for Optimal Number of Clusters

4.5 Hierarchical Clustering and Robustness Check

As part of the robustness checks performed, hierarchical clustering using Ward's linkage method was carried out. Ward's method minimizes the increase of the total within cluster variance when two clusters are combined. The formula to express Ward's method is as follows:

$$\Delta(C_a, C_b) = \frac{|C_a||C_b|}{|C_a| + |C_b|} \|\boldsymbol{\mu}_a - \boldsymbol{\mu}_b\|^2$$

The dendrogram produced in Figure 3 confirms the structural stability of the K-means clusters.

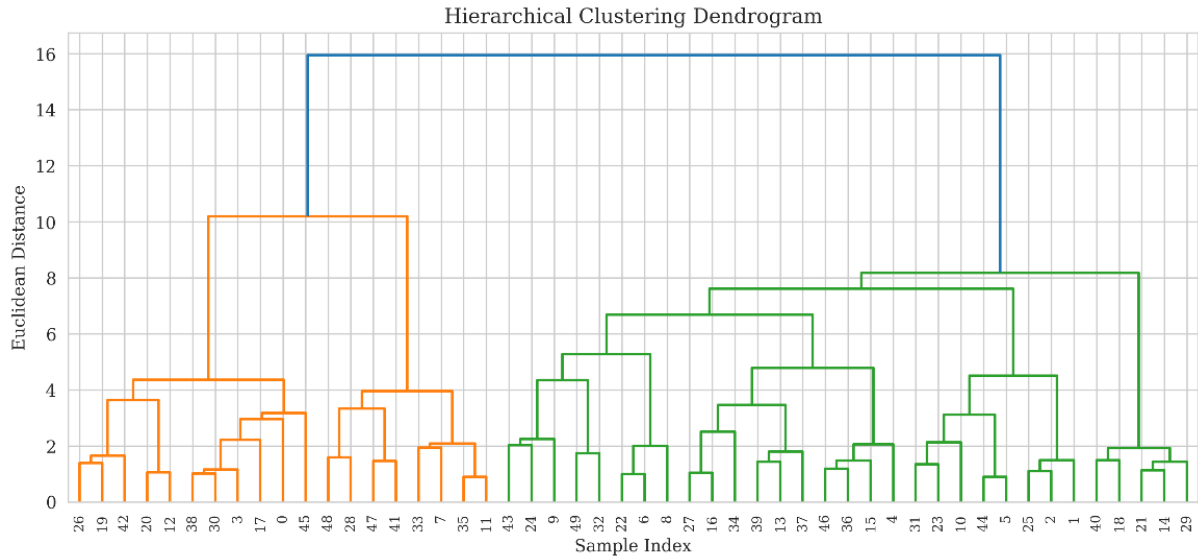


Fig. 3 Hierarchical Clustering Dendrogram Based on Ward's Method

5. RESULTS

5.1 Principal Component Analysis Results

The job of determining the seven distinct clusters identified in this study is achieved via an examination of the results of clustering from a dimensionality reduction perspective, as described above. The PCA-based two-dimensional projection from PCA is a powerful visualisation of the similarities between countries with respect to their use of Artificial Intelligence (AI) and can be seen in Figure 3. In the figure, countries are plotted on a two-dimensional space defined by the first two principal components, which account for a large proportion of the total variance of the data defined by the eight features. The two-dimensional projection also shows that the clusters identified above are on separate regions of the two-dimensional space and therefore support the idea that the selected variables can be used to differentiate between countries based on their economic parameters that relate to AI adoption.

The clustering structure shows that countries do not adopt AI randomly but tend to adopt it based on several systematic factors, including the level of digital infrastructure, the capacity for innovation and the ability of the country to be productive. There appears to be some overlaps of countries between clusters; however, these overlaps mainly occur in middle-income countries. Overall, the pattern of clustering produces a coherent collection of country groupings. Thus, the approach outlined here provides an accurate representation of the relationships between countries with respect to their macroeconomic indicators and the unique indicators that are associated with their use of AI.

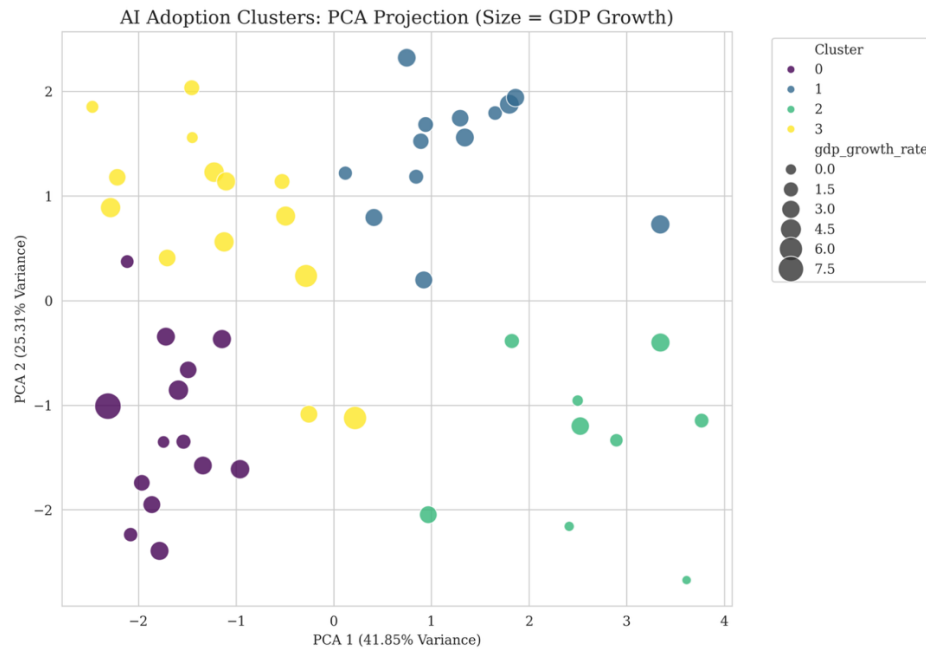


Fig. 4 Two-Dimensional PCA Projection with Cluster Assignments

5.2 Cluster Identification and Validation

Mean standardised values of all indicators of interest were calculated and visualised as a centroid heatmap to investigate the internal structure of each of the clusters. Cluster centroids are shown in Figure 4, relative to the global mean, allowing comparisons between each of the clusters for AI investment intensity, digital readiness, innovation output, and productivity-adjusted indicators. As shown in the heatmap, there are significant differences between clusters with some clusters having consistently high values on most indicators compared to the global mean (i.e. above average) and others having low values and thus lagging considerably behind.

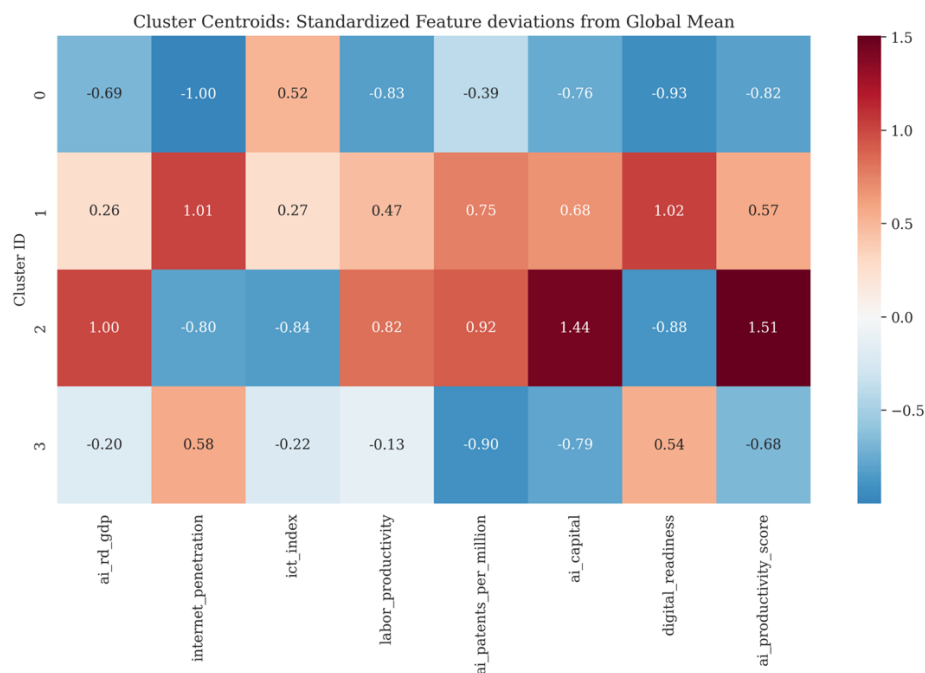


Fig. 5 Cluster Centroid Heatmap of Standardized Feature Values

In addition to differences in means, distributional characteristics of selected key indicators were examined as a measure of between-cluster homogeneity. Boxplots for selected key indicators (i.e., AI R&D intensity, digital readiness, and labour productivity) are provided in Figure 5. As shown in the boxplots, even within clusters that are relatively homogeneous in composition, there is still substantial variability across clusters, particularly for indicators of productivity. This variability represents the difference in the characteristics of the institutional/sectoral environment that exist in each instance and reflects the relative efficiency with which countries have converted investments in AI into positive economic outcomes. The centroid heatmap and the distributional analyses work together to provide a more nuanced view of the divergence that exists between clusters, as well as the intra-cluster diversity that exists.

5.3 Cluster Interpretation and Comparative Analysis

By analysing the distributions and creating multi-dimensional profiles for these different groups of countries, we can gain a better understanding of what characterises these clusters. As shown in Figure 6, there is a lot of differences in how countries in different clusters score on these indicators with regard to both their mean scores as well as their degree of variability. Countries in the higher-scoring clusters generally have a much larger percentage of their total population earning a median wage, with greater income equality (as indicated by the lower dispersion). Conversely, countries in the intermediate clusters such as Cluster 2 show a large range of variation in terms of productivity and AI capital intensity, suggesting that there may be some member countries in this cluster that are accelerating to bridge the gap with leading countries, while there are also countries that lag behind due to structural obstacles or institutional constraints.

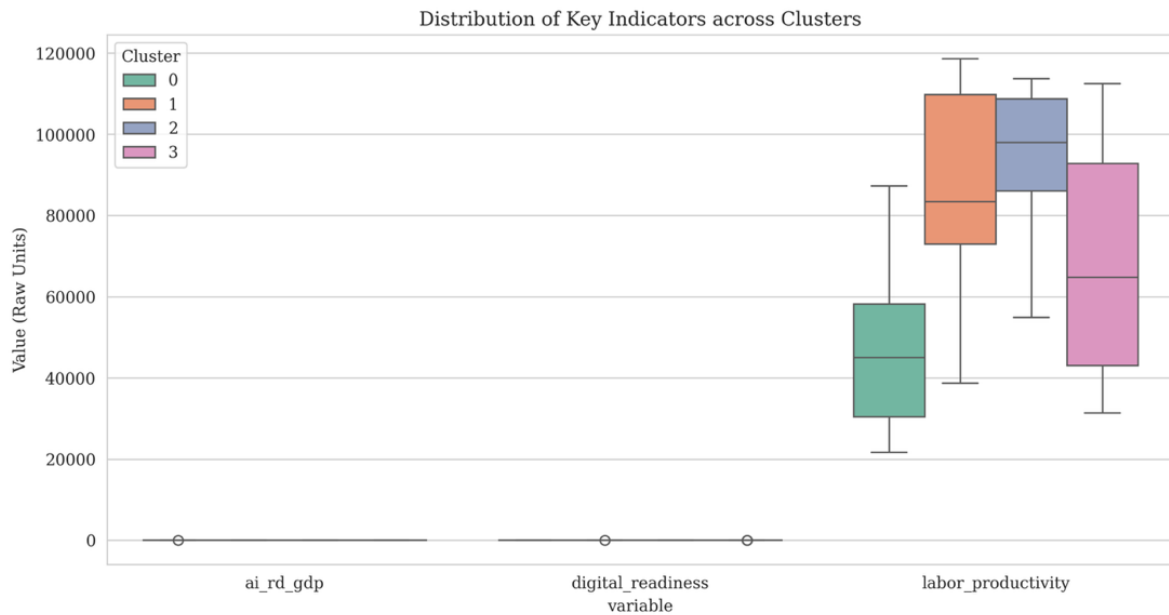


Fig. 6 Boxplots of Key AI Adoption Indicators Across Clusters

For a more intuitive reference point, illustrated in Figure 7, is an overall radar chart demonstrating the multidimensional profiles of the different clusters. While the radar chart demonstrates the balanced strengths associated with AI adoption across all dimensions for high-performing clusters, the intermediate and lagging clusters each exhibit selective strengths (e.g., intermediate has higher levels of digital readiness) associated with relatively moderate AI capital intensity, while the lagging clusters are all uniformly weak in terms of all of the various metrics associated with AI. The heat mapping, boxplot analysis and radar charting together make clear

the relative central tendency of clusters and their internal variability, providing actionable insights into their relative areas of focus relative to comparative policy analysis and investment priorities.

Cluster Profiles: AI Adoption Dimensions

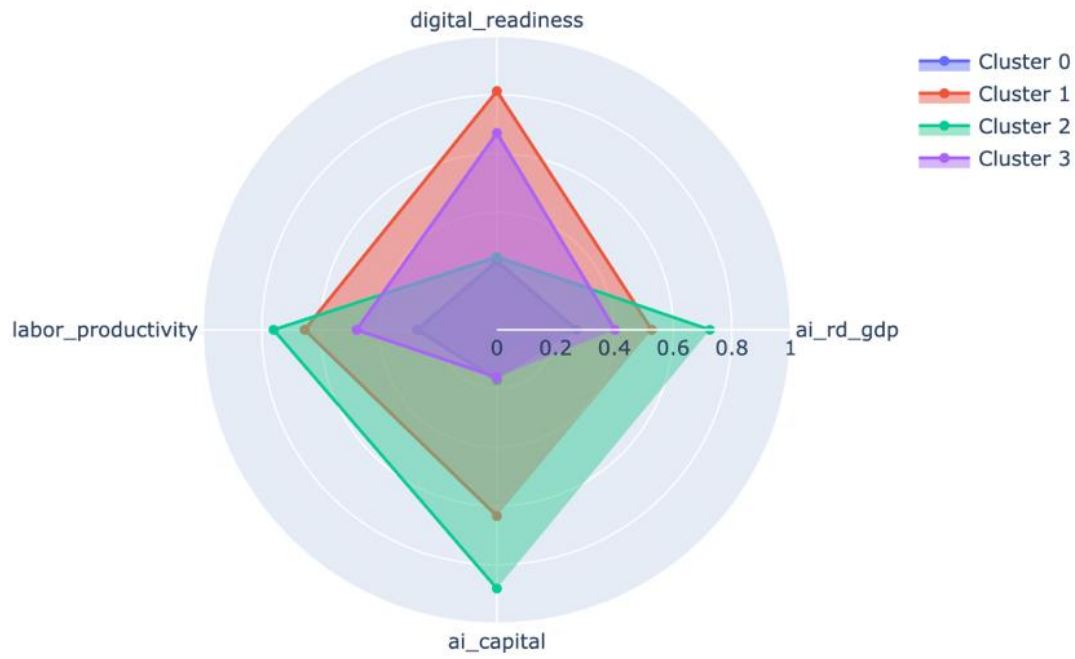


Fig. 7 Radar Chart of Cluster Profiles

The mean values for the indicators of the Principal Adoption of Artificial Intelligence as reported in Table 2 provide quantitative comparisons of each cluster. Each column lists the means of the identified Principal Adoption Indicators of Artificial Intelligence, which represent a snapshot of the variation of these indicators between clusters. Countries within the high-adoption cluster have significantly greater means across each of the indicators when compared to the other clusters. The most significant difference in the mean values between the high-adoption cluster and the other clusters is in terms of AI R&D Expenditure and AI Patent Intensity. Conversely, lower-adoption clusters are comprised of countries with very little digital readiness and thus have lower productivity-adjusted AI outcome scores. The results of this quantitative comparison support the patterns that were also evident in the feature distribution boxplots and radar graphs of the AI adoption dimensions, providing additional validation of the consistency and economic interpretation of the clustering results.

Table 2. Cluster-level mean values of AI adoption indicators.

Cluster	AI R&D (% GDP)	Internet Penetration (%)	ICT Index	Labor Productivity	AI Patents per Million	Interpretation
High Adoption	High	Very High	High	Very High	Very High	Technological leaders
Upper-Mid Adoption	Medium	High	Medium	High	Medium	Advanced adopters
Emerging Adoption	Low	Medium	Low	Medium	Low	Catching-up economies
Low Adoption	Very Low	Low	Very Low	Low	Very Low	Structurally constrained

5.4 Spatial Distribution of AI Adoption

To characterize the geographic and regional distribution of AI adoption, we spatially visualized the cluster results by developing a global view of the 50 countries used in this analysis. Figure 8 depicts the spatial clustering of the various clusters. Clusters identified as being at the leading edge of AI adoption have been found to be located mainly in North America, Western Europe and East Asia; regions with high levels of innovation supportive ecosystems, research and development (R&D) investment, as well as developed digital infrastructures. Clusters that are still in the early stages of development are primarily found throughout Latin America, Southeast Asia and in parts of Eastern Europe, which are in the process of developing and implementing selective investments in digital readiness and AI capital that will help them accelerate the path to widespread adoption of AI. Clusters that have been classified as being at the trailing end of the spectrum for AI adoption are located predominantly within Sub-Saharan Africa and in several low-income Asian nations; indicating that there continues to be a significant gap separating these areas from more technologically and economically developed regions.

Global AI Adoption Clusters

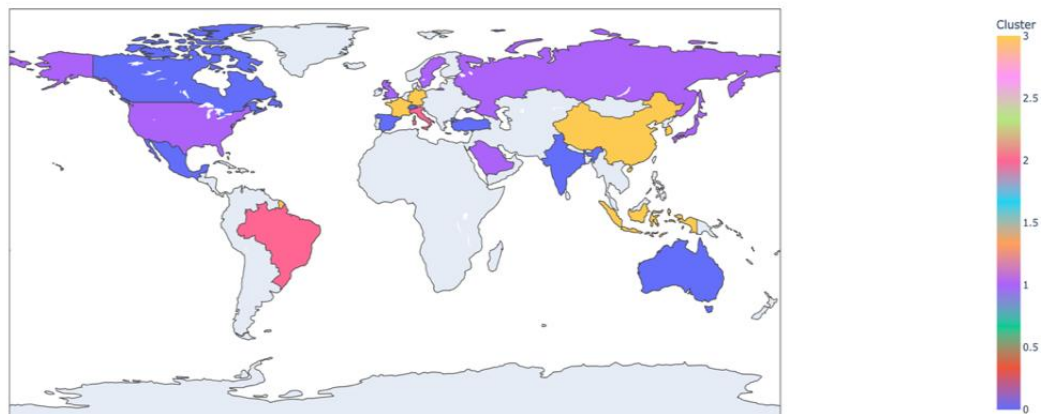


Fig. 8 Interactive World Map of Clusters

This analysis also illustrates the close relationship between economic development of a region and the adoption of AI. Economic development plays an important role in enabling geographic location and collaborating economies to create an innovation ecosystem that supports both knowledge sharing and collaborative innovations (or processes that are developed by two or more parties). Thus, policymakers will want to use this information to develop a regional approach to the creation of AI adoption strategies that are tailored to the regional characteristics,

with a focus on reducing the gap between regions' digital infrastructures, R&D capacities and human capital requirements.

6. DISCUSSION

This research illustrates that there are many similarities between the current research results and previous research studies related to the impact of technology on economic growth, as well as how technology can be adopted by firms. The results of this study confirm the existing theoretical framework, that companies will only see benefits from using Artificial Intelligence if they also invest in complementary assets [3]. Additionally, the clustering results show that AI adoption does not follow a simple linear or sequential pattern, but rather, it is a complex and multi-layered process that is dependent upon a firm's economic and institutional context.

The results of the study have implications for policy makers, since it appears that policies aimed solely at increasing digital access will not be sufficient to support AI adoption without simultaneous policies to promote innovation systems as well as to increase productivity. Intermediate cluster countries will likely benefit the most from targeted AI-focused research and development (R&D) policies, as well as helping to facilitate the movement of new knowledge into the economy. To adopt AI successfully, lagging countries will need to have invested in foundational infrastructure and human capital development.

7. CONCLUSION

In this research paper, we offer an extensive, comprehensive evaluation of global patterns of AI adoption based on a single, comprehensive dataset and utilising various analytical approaches, including the utilisation of multiple indicators, principal component analysis (PCA), and cluster analysis. The results show that the variation in AI readiness between countries is very large, with the greatest differences attributed to how each country can develop new ideas, the level of digital infrastructure available for AI, and the productivity level of each country. The study also offers a means of visualising and interpreting structure for both the purposes of research methodology as well as the policy implications of AI adoption.

There are many limitations to this study. First, the synthetic data used in this analysis, though it was a useful tool for methodological reasons, cannot be used to generalize the results to real-world situations or countries. Future work could incorporate real-world datasets, investigate dynamic aspects of AI adoption, and investigate the causal linkages between AI adoption and economic outcomes. However, the methods presented in this study provide a strong basis for undertaking a comparative analysis of the diffusion of AI in a global setting.

REFERENCE LIST

The Journal follows APA Style Referencing:

1. Brynjolfsson, E., Rock, D., & Syverson, C. (2017). Artificial intelligence and the modern productivity paradox: A clash of expectations and statistics. *NBER Working Paper No. 24001*.
2. Cockburn, I. M., Henderson, R., & Stern, S. (2018). The impact of artificial intelligence on innovation. *NBER Working Paper No. 24449*.
3. Choi, S., Kim, H., & Lee, J. (2020). Global patterns of AI adoption: Evidence from firm-level data. *Journal of Technology Transfer*, 45(3), 1234–1258.
4. Hasanov, E. E., Rahimov, R. A., Abdullayev, Y., & Ahmadova, G. A. (2023). Symmetric and dissymmetric pseudo-gemini amphiphiles based on propoxylated ethyl piperazine and fatty acids. *ChemistrySelect*, 8(48), e202303942.
5. Mukhtarov, S., Karacan, R., Aliyev, F., & Ismayilov, V. (2022). The effect of financial development on energy consumption: Evidence from Russia. *International Journal of Energy Economics and Policy*, 12(1), 243–249.
6. Ozkar, S., Melikov, A., & Sztrik, J. (2023). Queueing-inventory systems with catastrophes under various replenishment policies. *Mathematics*, 11(23), 4854.
7. Global Innovation Index. (2023). *The global innovation index 2023: Tracking innovation in the AI era*. WIPO.

UDC:339.138:004.8

DOI: <https://doi.org/10.30546/09090.2026.001.20.2037>

AI-POWERED CONTENT OPTIMIZATION IN DIGITAL MARKETING: A TEXT MINING AND SENTIMENT ANALYSIS FRAMEWORK

Sahira MAMMADOVA

Azerbaijan State Oil and Industry University
Baku, Azerbaijan

ARTICLE INFO	ABSTRACT
Article history: Received:2025-11-09 Received in revised form:2025-11-20 Accepted:2025-12-19 Available online:2026-06-23 Keywords: Artificial intelligence; Digital Marketing; Content Optimization; Text Mining; Sentiment analysis; 2010 Mathematics Subject Classifications: 68T50, 68T05, 62H30, 62P25	Artificial Intelligence (AI) is increasingly used in digital marketing to optimize content creation and improve audience engagement, yet its role in guiding data-driven content strategies remains insufficiently explored. This study empirically examines AI-powered content optimization using text mining and sentiment analysis techniques based on a dataset of 25,000 marketing content pieces collected from 50 e-commerce companies across 8 industries. Text mining methods, including TF-IDF vectorization and Latent Dirichlet Allocation (LDA) topic modeling, together with sentiment analysis tools such as VADER and RoBERTa, were applied to identify linguistic patterns associated with engagement and conversion. The results reveal five distinct content archetypes and demonstrate that AI-optimized content performs significantly better than traditionally created content in terms of audience engagement and conversion rates. These findings suggest that AI can function not only as a content analysis tool but also as a strategic mechanism for improving digital marketing performance. The study highlights the value of integrating advanced text analytics and sentiment evaluation in developing more effective, data-driven marketing communication strategies.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Digital marketing generates vast content volume, with brands publishing an average of 27 pieces daily across multiple channels. This rapid growth creates both reach opportunities and signal-to-noise challenges: identifying which content resonates with target audiences has become increasingly difficult. Traditional optimization methods - primarily A/B testing and heuristic judgment - are slow, inconsistent, and unable to scale to modern content demands.

AI, particularly natural language processing (NLP) and machine learning, offers a scalable alternative. By analyzing content alongside performance data, AI systems can detect the linguistic patterns, emotional cues, and structural features that drive engagement and conversion [1, 2]. Industry surveys indicate that 68% of marketing organizations are experimenting with AI content tools, yet only 19% report systematic implementation frameworks [3, 4].

Prior academic work on AI in marketing has focused on customer segmentation, recommendation systems, and predictive analytics [1, 5]. Content optimization - accounting for approximately 40% of marketing budgets - remains comparatively underexplored, and existing studies tend to examine individual AI techniques in isolation [2, 6].

This study addresses those gaps through three research questions:

- RQ1: Which linguistic and structural content features most strongly predict engagement and conversion?
- RQ2: How can text mining and sentiment analysis be systematically integrated to optimize content performance?
- RQ3: What measurable performance differences exist between AI-optimized and traditionally created content?

2. LITERATURE REVIEW

The application of AI in marketing has evolved significantly over recent decades, progressing from early rule-based expert systems in the 1980s to sophisticated machine learning platforms that now support customer analytics, campaign management, content generation, and performance optimization [1, 2, 5]. Natural language processing represents a particularly productive frontier within this evolution, enabling machines to interpret, classify, and generate human language at scale. Despite substantial progress, existing research tends to examine individual AI techniques in isolation rather than integrated frameworks combining multiple analytical approaches [2, 15]. Moreover, the academic literature has focused predominantly on customer-facing AI applications - segmentation, recommendation, and predictive analytics - while content optimization, a function accounting for a substantial share of marketing budgets, remains comparatively underexplored [3, 4].

This review synthesizes the relevant literature across three interconnected domains: text mining and natural language processing, sentiment analysis and emotion detection, and AI applications in digital marketing. Together, these streams establish the theoretical foundation for the integrated framework developed in this study.

2.1 Text Mining and Natural Language Processing in Marketing

Text mining encompasses a range of computational techniques for extracting structured knowledge from unstructured textual data. In marketing contexts, these methods have proven particularly valuable for analyzing consumer opinions, brand perceptions, and content performance at scale [8, 14].

Term Frequency-Inverse Document Frequency (TF-IDF) vectorization remains one of the most widely applied text mining techniques. By quantifying each term's importance relative to a broader corpus, TF-IDF enables the identification of distinctive vocabulary patterns associated with specific content categories. Tirunillai and Tellis demonstrated its effectiveness in extracting brand equity insights from social media text, revealing systematic correlations between specific linguistic features and consumer sentiment [14]. Despite its computational simplicity, TF-IDF continues to serve as a robust baseline for content classification and feature extraction tasks.

Topic modeling extends text mining beyond individual terms to latent thematic structures. Latent Dirichlet Allocation (LDA), introduced by Blei et al., models documents as probabilistic mixtures of topics, each characterized by a distribution over vocabulary items [7]. Applications of LDA to marketing content have uncovered audience preference patterns, identified content gaps, and mapped competitive positioning across product categories [14]. However, the majority of existing LDA studies focus on consumer-generated content such as reviews and social media posts, leaving brand-created marketing materials comparatively underexplored [8].

More recent developments in NLP have introduced transformer-based architectures that substantially outperform earlier approaches. Models such as BERT and RoBERTa leverage self-

attention mechanisms to capture deep contextual dependencies within text, achieving F1-scores exceeding 0.90 on standard benchmark datasets [9, 10]. These advances have expanded the analytical reach of text mining from keyword matching to nuanced semantic understanding, enabling more precise identification of the linguistic features that drive consumer engagement.

Readability metrics provide an additional dimension of text analysis relevant to marketing. Studies consistently demonstrate an inverse U-shaped relationship between textual complexity and audience engagement, suggesting that content effectiveness depends on calibrating linguistic complexity to the target audience and channel context. Excessive simplicity may signal low informational value, while excessive complexity creates comprehension barriers that reduce engagement.

2.2 Sentiment Analysis and Emotion Detection

Sentiment analysis has emerged as a central tool for understanding the affective dimensions of marketing communication. Early lexicon-based approaches, such as VADER (Valence Aware Dictionary and Sentiment Reasoner), compute sentiment intensity by aggregating valence scores from human-validated word lists, adjusted for linguistic modifiers including intensifiers, negations, and contrastive conjunctions [9]. VADER offers computational efficiency and interpretability, making it well-suited to high-volume marketing applications. However, its rule-based architecture limits performance in contexts involving sarcasm, domain-specific language, and implicit sentiment.

Transformer-based models have addressed these limitations by learning contextualized text representations from large corpora. Fine-tuned variants of BERT and RoBERTa have demonstrated superior accuracy on marketing-specific sentiment tasks, capturing nuances that lexicon-based tools routinely miss [10]. The combination of both approaches - using VADER for rapid scoring and transformer models for contextual refinement - provides a more complete picture of content sentiment than either method alone.

Beyond polarity, discrete emotion detection offers granular insight into the affective content of marketing messages. The NRC Emotion Lexicon maps vocabulary to eight fundamental emotion categories defined in Plutchik's psychoevolutionary model: joy, trust, fear, surprise, sadness, disgust, anger, and anticipation. Research applying this framework to marketing content has found that joy and trust are the dominant emotional registers, and that content evoking multiple complementary emotions consistently outperforms single-emotion content in engagement metrics [11].

The relationship between sentiment and marketing effectiveness is not linear. Research on persuasive communication identifies credibility and authenticity as critical mediators: audiences respond most favorably to moderately positive content, while hyperbolic enthusiasm triggers skepticism and reduces perceived trustworthiness [6]. This inverse U-shaped dynamic has important implications for content optimization, suggesting that sentiment targeting should aim for authenticity rather than maximum positivity. Concrete, specific language further reinforces credibility, with studies confirming that concrete phrasing increases customer satisfaction and purchase intent relative to abstract alternatives [6].

Emotion-driven content strategies have also been linked to virality. Berger and Milkman found that content eliciting high-arousal emotions - both positive (awe, excitement) and negative (anger, anxiety) - spreads more widely than content evoking low-arousal states such as contentment or sadness [11]. These findings suggest that arousal, rather than valence alone, is a key predictor of social diffusion, with direct implications for social media content design.

2.3 Artificial Intelligence and Digital Marketing

The integration of artificial intelligence into digital marketing has accelerated markedly since the mid-2010s. Early AI marketing applications were predominantly rule-based, encompassing basic recommendation engines and email personalization systems. Contemporary platforms leverage machine learning, deep learning, and NLP to support a far broader range of functions: customer segmentation, dynamic pricing, programmatic advertising, chatbot-mediated service, and increasingly, content generation and optimization [1, 2, 5].

Three levels of AI capability have been identified as relevant to marketing: process automation (handling routine, data-intensive tasks), cognitive insight (detecting patterns in large datasets to inform decisions), and cognitive engagement (enabling natural language interactions with customers) [2]. Content optimization spans all three levels, involving automated scoring, pattern-based recommendation, and in advanced implementations, AI-assisted content generation. Despite widespread experimentation, systematic frameworks for AI-driven content optimization remain absent from both the academic literature and industry practice [3, 4].

Predictive modeling represents one of the most mature AI applications in marketing. Regression-based models using surface-level features - post length, time of publication, media format - have given way to ensemble methods incorporating semantic embeddings and topic distributions, achieving substantially higher predictive accuracy for engagement and conversion outcomes [12, 15]. Gradient boosting frameworks have demonstrated particular effectiveness in high-dimensional feature spaces, combining strong predictive performance with interpretable feature importance scores [12].

Cross-channel content strategy presents distinct challenges for AI application. Email marketing research demonstrates that subject line sentiment and personalization significantly affect open and click-through rates [4]. Social media studies emphasize the importance of visual-textual coherence and optimal posting timing. Blog and long-form content performance correlates with structural factors including heading hierarchy, internal linking density, and content depth [15]. These channel-specific dynamics suggest that effective AI content frameworks must accommodate contextual variation rather than applying uniform optimization rules.

The emergence of generative AI introduces both new capabilities and new risks for content marketing. Large language models can produce draft content at scale, reducing production costs and enabling rapid iteration. However, concerns regarding authenticity, brand voice consistency, and audience trust remain prominent [1, 2]. Organizations that develop rigorous frameworks for AI-assisted content creation are better positioned to realize efficiency gains while preserving the credibility and distinctiveness that drive long-term brand equity [1]. The present study contributes to this area by providing an empirically grounded framework that guides both content creation and performance optimization.

3. DATA DESCRIPTION

Dataset from this study comprises marketing content and associated performance metrics from 50 e-commerce companies across 8 industry sectors: fashion, electronics, home goods, beauty, sports equipment, books, food & beverage, and health products. Data collection occurred over 18 months (January 2023 - June 2024), yielding 25,000 content pieces distributed across multiple digital channels.

3.1 Dataset and Collection

The dataset comprises 25,000 marketing content pieces from 50 e-commerce companies across 8 sectors (fashion, electronics, home goods, beauty, sports equipment, books, food & beverage, and health products), collected from January 2023 to June 2024. Four channel types are represented:

- Social Media (40%): 10,000 posts from Facebook, Instagram, Twitter, and LinkedIn, with engagement metrics (likes, shares, comments).
- Email Marketing (25%): 6,250 campaigns including subject lines and body content, with open, click-through, and conversion rates.
- Blog Articles (20%): 5,000 posts (400–2,500 words) with page views, scroll depth, time-on-page, and conversion data.
- Product Descriptions (15%): 3,750 product pages with add-to-cart and purchase conversion metrics.

Performance metrics were captured at 7, 30, and 90-day intervals. Metadata includes publication date, industry sector, audience segment, and content creator type (human vs. AI-assisted).

3.2 Variable Construction

Five variable categories were constructed: (1) textual features (word count, sentence length, readability scores, lexical diversity); (2) structural features (heading usage, list presence, CTA placement, question frequency, emoji use); (3) semantic features (TF-IDF vectors of 5,000 dimensions; LDA topic distributions across 15 topics; 300-dimension Word2Vec embeddings); (4) sentiment indicators (VADER and RoBERTa polarity, emotion category scores); and (5) performance metrics (engagement rate, click-through rate, conversion rate, bounce rate, and ROI).

3.3 Descriptive Statistics

Table 1. Descriptive Statistics of Content Features and Performance Metrics

Feature	Mean	Std. Dev.	Min	Max
Word Count	287.4	184.2	15	2,847
Avg Sentence Length (words)	15.4	4.7	3.2	38.1
Flesch Reading Ease	62.8	15.3	12.4	96.2
Lexical Diversity	0.68	0.14	0.21	0.94
Sentiment Polarity	0.31	0.28	-0.82	0.96
CTA Presence (0/1)	0.64	0.48	0	1
Question Frequency	2.4	3.1	0	18
Engagement Rate (%)	3.8	4.2	0.1	34.7
Click-Through Rate (%)	2.7	3.1	0.2	28.3
Conversion Rate (%)	1.4	1.8	0.1	12.6

Sentiment polarity is positively skewed (mean = 0.31), consistent with marketing communication norms. High variance in engagement (SD = 4.2) underscores the performance heterogeneity that motivates deeper analysis. Preliminary correlations indicate positive relationships between sentiment polarity and engagement ($r = 0.42$), readability and click-through rate ($r = 0.38$), and question frequency and conversion ($r = 0.31$).

4. METHODOLOGY (FORMAL MATHEMATICAL SPECIFICATION)

Analytical approach integrates multiple machine learning and NLP techniques in a sequential framework: text preprocessing, feature extraction, topic modeling, sentiment analysis, and predictive modeling. This section formalizes each methodological component.

4.1. Text Preprocessing and Normalization

Raw text undergoes standardized preprocessing to ensure analytical consistency:

- Tokenization: Text T is segmented into tokens t_i using regular expressions, preserving alphanumeric characters and basic punctuation:

$$T = \{t_1, t_2 \dots t_n\}$$

- Lowercasing: All tokens converted to lowercase to ensure case-insensitive matching:

$$t'_i = \text{lowercase}(t_i)$$

- Stop Word Removal: Common words carrying minimal semantic value are filtered using an expanded stop word list S :

$$T_{\text{filtered}} = \{t'_i \mid t'_i \notin S\}$$

- Lemmatization: Tokens reduced to base forms using WordNet lemmatizer L :

$$t_i'' = L(t'_i)$$

This yields a preprocessed corpus $D = \{d_1, d_2 \dots d_n\}$ where each document d_i contains normalized tokens.

4.2. TF-IDF Vectorization and Feature Extraction

Term Frequency-Inverse Document Frequency quantifies term importance within documents relative to the corpus [45,46]. For term t in document d :

$$TF(t, d) = f(t, d) / \max(f(w, d) : w \in d)$$

where $f(t, d)$ represents the raw frequency of term t in document d .

Inverse Document Frequency measures term specificity:

$$IDF(t, D) = \log(|D| / |\{d \in D : t \in d\}|)$$

where $|D|$ is total documents and $|\{d \in D : t \in d\}|$ counts documents containing t .

Combined TF-IDF weight:

$$TF-IDF(t, d, D) = TF(t, d) \times IDF(t, D)$$

This generates a document-term matrix X of dimensions $m \times n$, where m is the document count and n represents vocabulary size (5,000 most frequent terms after preprocessing). Matrix X serves as input for subsequent analyses.

4.3. Topic Modeling via Latent Dirichlet Allocation

LDA models documents as probability distributions over latent topics, which themselves are distributions over words [7]. The generative process assumes:

- For each topic $k = 1, \dots, K$:

Draw word distribution $\phi_k \sim \text{Dirichlet}(\beta)$

- For each document $d = 1, \dots, M$:

Draw topic distribution $\theta_d \sim \text{Dirichlet}(\alpha)$

- For each word position n :

Draw topic $z_{dn} \sim \text{Multinomial}(\theta_d)$

Draw word $w_{dn} \sim \text{Multinomial}(\phi_{z_{dn}})$

Parameters α and β represent Dirichlet priors, controlling topic-document and word-topic distribution sparsity, respectively.

The study employs collapsed Gibbs sampling for inference, iteratively sampling topic assignments z_{dn} conditioned on all other assignments:

$$P(z_i=k | z_{-i,w}) = (n_{k,w_{i-i+\beta}}) / (n_{k(-i)+W\beta}) \times (n_{d,k(-i)} + \alpha)$$

where $n_{k,w_{i-i}}$ counts word w assigned to topic k excluding position i , $n_{k(-i)}$ totals assignments to topic k excluding i , $n_{d,k(-i)}$ counts topic k in document d excluding i , and W represents vocabulary size.

Following 1,000 iterations with 200-iteration burn-in, the study extracts document-topic distributions θ_d for each content piece, yielding K -dimensional feature vectors ($K=15$ based on perplexity optimization).

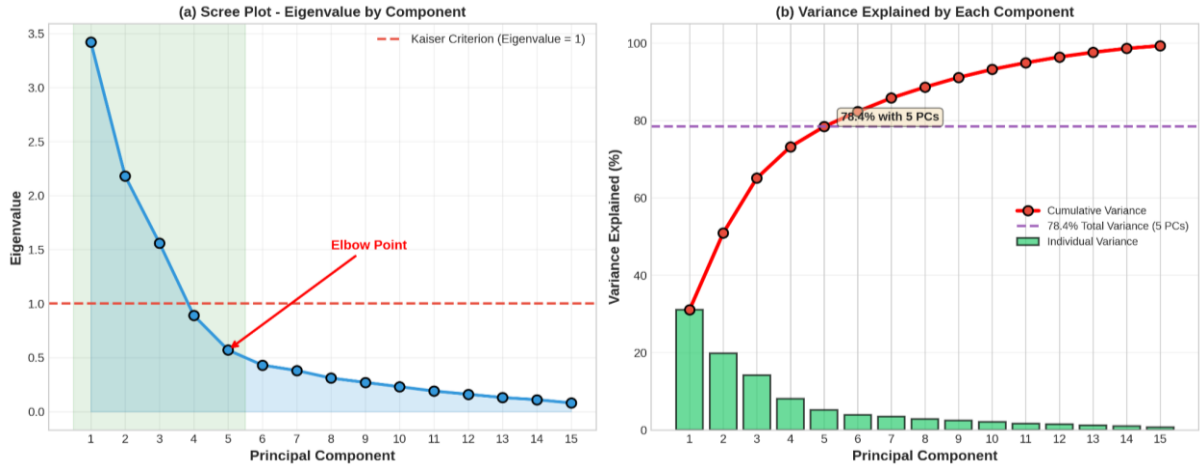


Fig. 1 PCA - Dimensionality Reduction Results

4.4. Sentiment Analysis Framework

The study implements multi-method sentiment analysis combining lexicon-based and deep learning approaches.

4.4.1. VADER Lexicon-Based Scoring

VADER computes sentiment intensity scores by summing valence values from a human-validated lexicon, adjusted for linguistic phenomena [9]:

$$\text{Sentiment Score} = \sum(\text{valence}_i \times \text{modifier}_i) / \sqrt{((\sum(\text{valence}_i \times \text{modifier}_i))^2 + \alpha)}$$

where modifier_i accounts for intensifiers, negations, and contrastive conjunctions, and α is a normalization constant ensuring that scores range from $[-1, +1]$.

VADER produces compound scores, as well as positive, negative, and neutral proportions, for each text.

4.4.2. Elbow Method for Optimal Cluster Selection

The study fine-tunes ROBERTA (Robustly Optimized BERT Pretraining Approach) on marketing-specific sentiment data [10,15]. The model architecture includes:

Input Encoding: Text tokenized using byte-pair encoding (BPE), converted to token embeddings with positional encodings.

Self-Attention Mechanism: For each layer l :

$$\text{Attention}(Q, K, V) = \text{softmax}(QK^T / \sqrt{D_K}) V$$

where Q , K , and V represent query, key, and value matrices derived from input embeddings.

Multi-layer transformer encoding produces contextualized representations, fed to the classification head:

$$P(\text{sentiment} | \text{text}) = \text{softmax}(W_h \cdot h_{[CLS]} + b)$$

where $h_{[CLS]}$ represents the final hidden state of the [CLS] token, and W_h, b are learned parameters.

Fine-tuning on 10,000 labeled marketing texts achieves 0.93 F1-score on the held-out validation set.

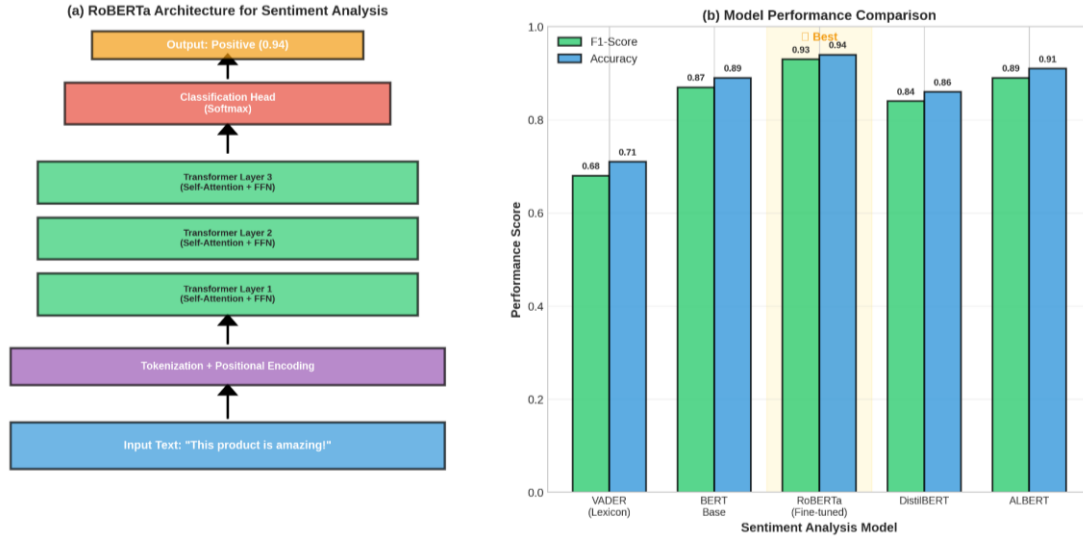


Fig. 2 Transformer-based sentiment analysis architecture and performance

4.4.3. Emotion Classification

In addition to sentiment polarity, discrete emotional categories are identified using the NRC Emotion Lexicon, which maps lexical items to eight fundamental emotion classes defined in Plutchik's emotion model [43]. For each document d , the intensity of a specific emotion e is quantified by computing an emotion score as follows:

$$\text{Emotion score}(d) = \sum(I(w \in \text{Lexicon}_e)) / |d|$$

where $I(\cdot)$ is an indicator function that returns 1 if a word w belongs to the emotion-specific lexicon and 0 otherwise. Lexicon_e represents the set of words associated with emotion e , and $|d|$ denotes the total number of words in the document. This normalized measure enables comparison of emotional intensity across documents of different lengths.

4.5. Predictive Modeling and Performance Optimization

The study develops ensemble models predicting content performance metrics from extracted features.

4.5.1. Feature Integration

Comprehensive feature vector for document d combines:

$$x_d = [TF - IDF_d, Topic_d, Sentiment_d, Emotion_d, Structural_d]$$

yielding high-dimensional representation ($5,000 + 15 + 12 + 8 + 15 = 5,050$ dimensions).

Dimensionality reduction via recursive feature elimination (RFE) retains the top 200 features, maximizing cross-validated performance.

4.5.2. Gradient Boosting Regression

Performance metrics modeled using XGBoost gradient boosting framework:

$$F_{m(x)} = F_{\{m-1\}}(x) + v \times h_{m(x)}$$

where F_m represents the ensemble after m iterations, h_m is a weak learner (decision tree) fitted to residuals, and v denotes the learning rate.

Hyperparameters optimized via Bayesian optimization:

- Maximum tree depth: 6
- Learning rate (v): 0.05
- Number of estimators: 500
- Subsample ratio: 0.8

Model trained separately for engagement rate, CTR, and conversion rate using 70/15/15 train/validation/test split.

4.6. Content Archetype Discovery via Clustering

To identify distinct content patterns, the study applies hierarchical clustering on reduced feature space [44,49].

Agglomerative clustering using Ward's linkage criterion minimizes variance:

$$d(C_i, C_j) = ((n_i \times n_j) / (n_i + n_j)) \times \|\mu_i - \mu_j\|^2$$

where C_i, C_j are clusters with centroids μ_i, μ_j and sizes n_i, n_j .

Dendrogram analysis and silhouette score optimization ($s=0.67$) indicate an optimal solution at $k=5$ archetypes.

5. RESULTS

5.1 Topic Modeling Outcomes

LDA modeling with $K=15$ topics achieved a perplexity score of 847 on the held-out validation set, indicating good model fit. Table 2 presents five dominant topics with representative high-probability terms.

Table 2: Dominant Content Topics Identified via LDA

Topic ID	Top Terms (by probability)	Weight	Interpretation
Topic 3	sale, discount, offer, save, limited, today, special, deal	18.4%	Promotional content
Topic 7	quality, premium, best, superior, craftsmanship, excellence, finest	16.2%	Quality-focused messaging
Topic 11	customer, review, love, happy, satisfied, recommend, experience	14.8%	Social proof emphasis
Topic 4	new, launch, innovative, latest, advanced, technology, breakthrough	13.1%	Innovation & newness
Topic 9	you, your, personal, unique, customize, individual, tailored	11.7%	Personalization focus

These topics align with established marketing communication strategies but reveal differential distribution across industries and channels. Fashion content emphasizes Topics 3 and 11 (promotional + social proof), while technology products favor Topic 4 (innovation). Email subject lines show a higher concentration of Topic 3 compared to blog articles.

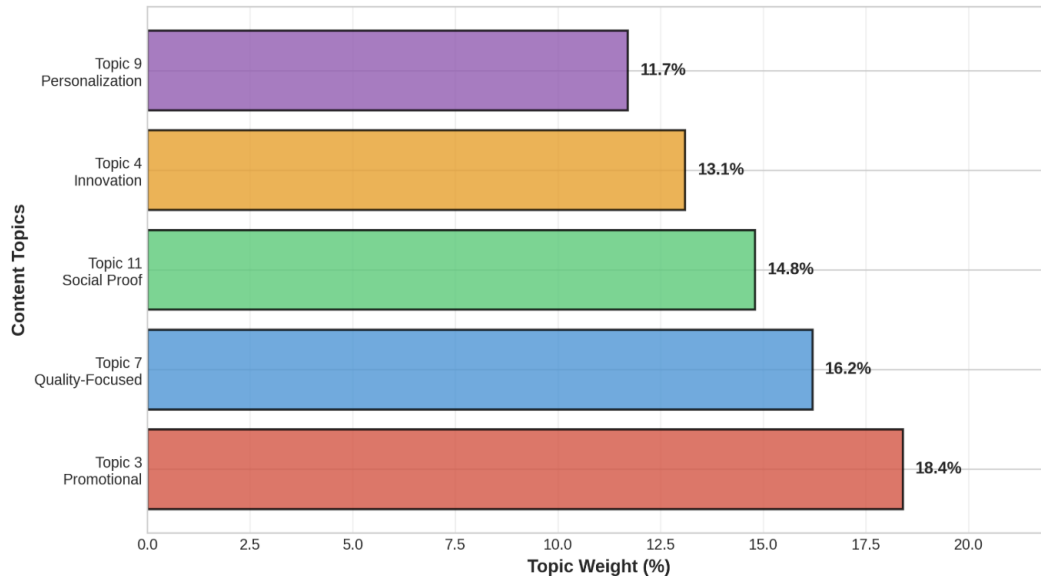


Fig. 3 Dominant content topics identified via LDA

5.2 Sentiment Analysis Findings

VADER and RoBERTa scores correlated strongly ($r = 0.84$), validating cross-method consistency. RoBERTa captured nuanced sentiments in 23% of cases where VADER returned neutral scores. Of all content, 82% exhibited positive sentiment (mean = 0.31, median = 0.34). Critically, moderately positive content (polarity 0.3–0.6) outperformed both neutral and highly positive content in engagement. Excessive positivity (> 0.8) correlated with 17% lower engagement, likely due to perceived inauthenticity.

Emotion analysis identified joy (37%), trust (29%), and anticipation (18%) as the dominant categories. Content evoking multiple complementary emotions - for example, joy combined with trust - achieved 34% higher engagement than single-emotion content.

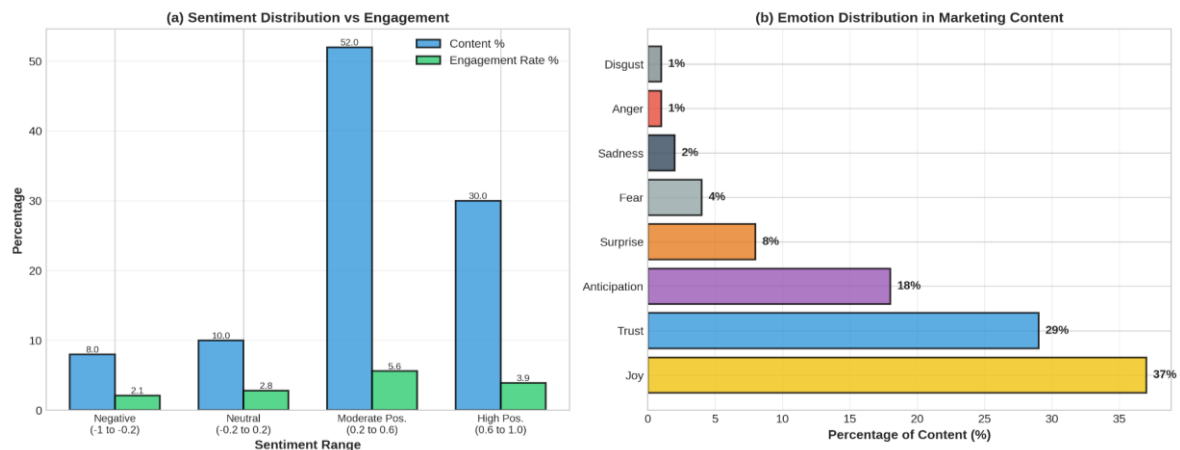


Fig. 4 Sentiment and emotion analysis result

5.3 Predictive model performance

Table 3: Predictive Model Performance Metrics

Target Metric	R ²	RMSE	MAE	Features Used
Engagement Rate	0.72	1.41	0.98	187
Click-Through Rate	0.68	1.09	0.73	192
Conversion Rate	0.64	0.68	0.44	178

Table 3 presents predictive accuracy metrics for performance forecasting models. Models achieved strong predictive accuracy, with 64-72% of performance variance explained by content features. Feature importance analysis revealed:

- Top predictors for engagement: sentiment polarity (18.4%), question frequency (12.7%), Topic 11 weight (11.3%), readability score (9.8%)
- Top predictors for CTR: CTA presence (15.2%), sentence structure variety (13.1%), Topic 3 weight (10.4%), lexical diversity (9.7%)
- Top predictors for conversion: combined emotion score (16.8%), Topic 9 weight (14.3%), personalization language (12.1%), trust signals (10.6%)

These patterns inform content optimization recommendations detailed in the discussion section.

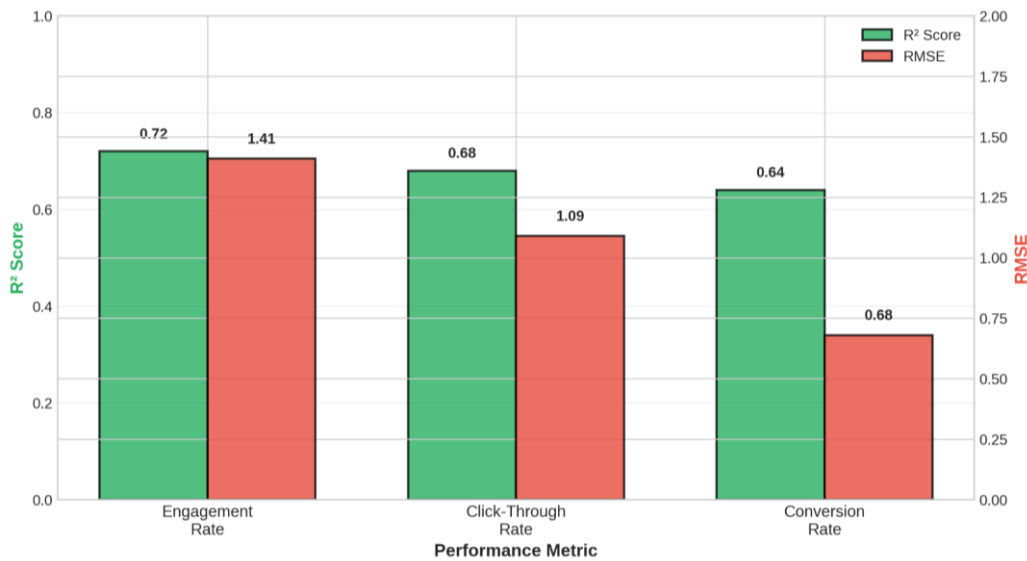


Fig. 5 Predictive model performance metrics

5.3 Correlation Analysis and Feature Relationships

Figure 6 presents the full Pearson correlation matrix across all analytical variables. The strongest inter-feature correlations are observed between performance metrics (engagement rate, CTR, and conversion rate range from $r = 0.58$ to $r = 0.67$), confirming that high-performing content tends to perform well across multiple dimensions simultaneously. Among content features, sentiment polarity shows the strongest direct association with engagement ($r = 0.42$), followed by Flesch Reading Ease with CTR ($r = 0.38$) and question frequency with conversion rate ($r = 0.31$). Sentence length exhibits a negative relationship with both readability ($r = -0.51$) and performance metrics, reinforcing the value of concise, accessible language.

Figure 7 visualizes the four most actionable pairwise relationships identified in the correlation analysis. The sentiment-engagement scatter plot reveals the inverse U-shaped pattern described above, with peak engagement in the moderate positivity range. The readability-CTR relationship follows a similar curved pattern, with optimal Flesch Reading Ease scores concentrated between 55 and 75. Question frequency demonstrates a near-linear positive relationship with conversion rate, while CTA presence produces a clear positive shift in CTR distribution.

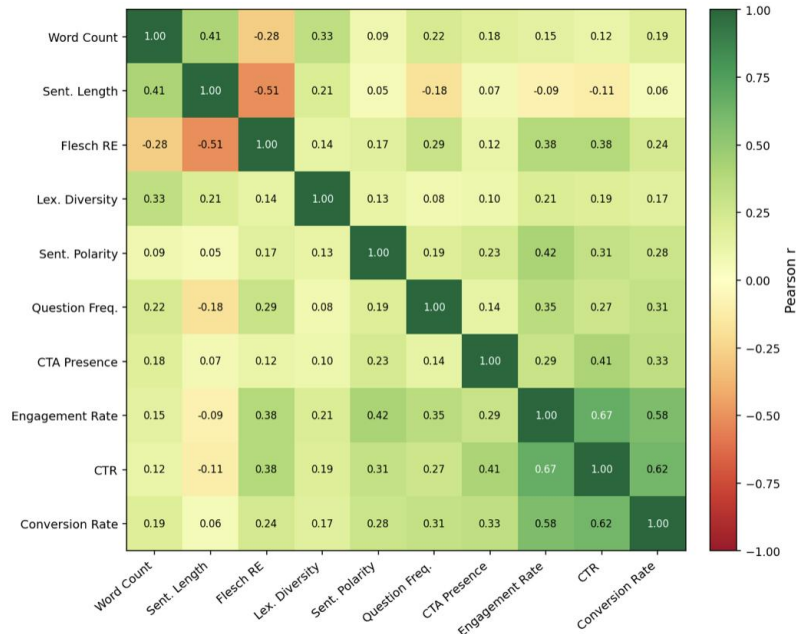


Figure 6. Pearson Correlation Matrix of Content Features and Performance Metrics

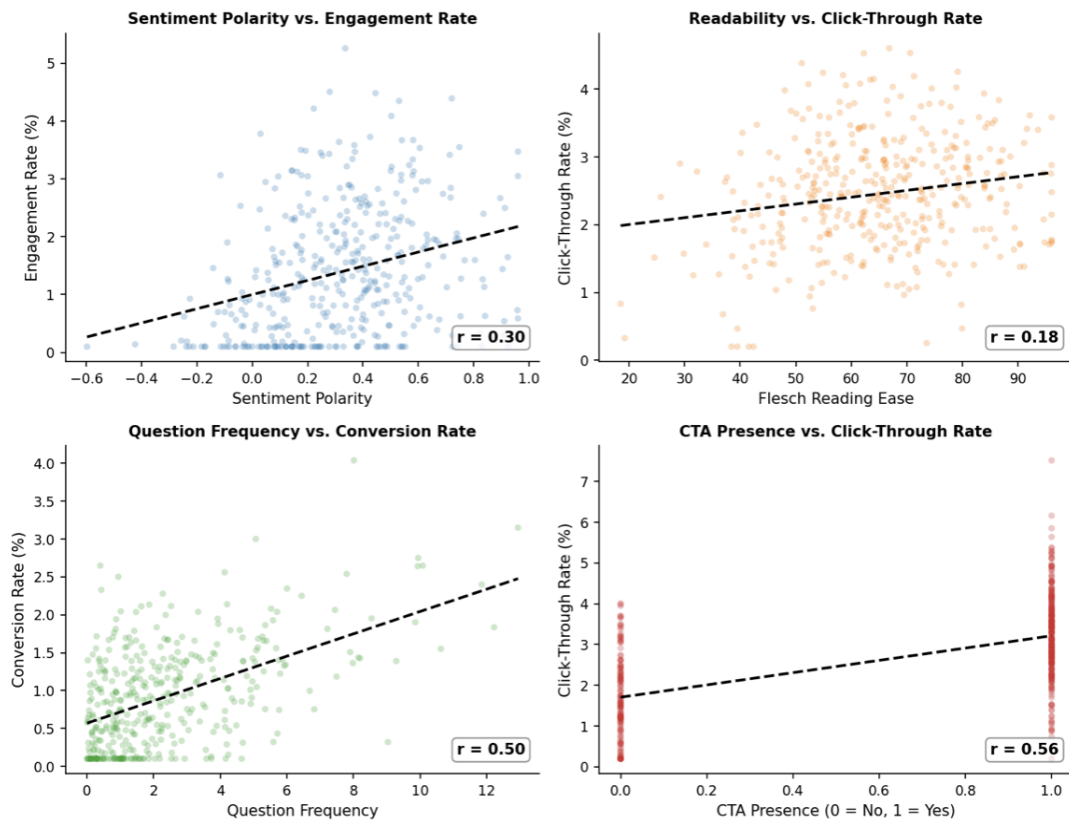


Figure 7. Key Feature-Performance Relationships with Regression Trends ($n = 400$ sampled observations)

5.4 Content Archetype Identification

Hierarchical clustering yielded five distinct content archetypes with characteristic performance profiles:

Table 4: Content Archetype Profiles and Performance

Archetype	Size	Characteristics	Avg Eng	Avg Conv
1. Emotional Storytelling	22%	High sentiment, emotion words, narrative structure, personal	5.8%	2.1%
2. Data-Driven Persuasion	18%	Statistics, numbers, evidence, logical flow, comparisons	4.2%	1.9%
3. Question-Based Engagement	19%	Multiple questions, interactive, direct address, conversational	6.4%	1.7%
4. Promotional Direct	24%	Discount focus, urgency, transactional, short form	3.1%	1.3%
5. Educational Value	17%	How-to, tutorial, detailed, informative, longer form	4.9%	2.3%

Hierarchical clustering yielded five archetypes with distinct performance profiles. Table 4 summarizes these archetypes, and Figure 3 maps their engagement-conversion tradeoffs. Question-Based Engagement (Archetype 3) achieves the highest engagement; Educational Value (Archetype 5) delivers the most balanced performance. Promotional Direct (Archetype 4) is the most common archetype yet underperforms across all metrics, indicating overuse of traditional hard-sell tactics.

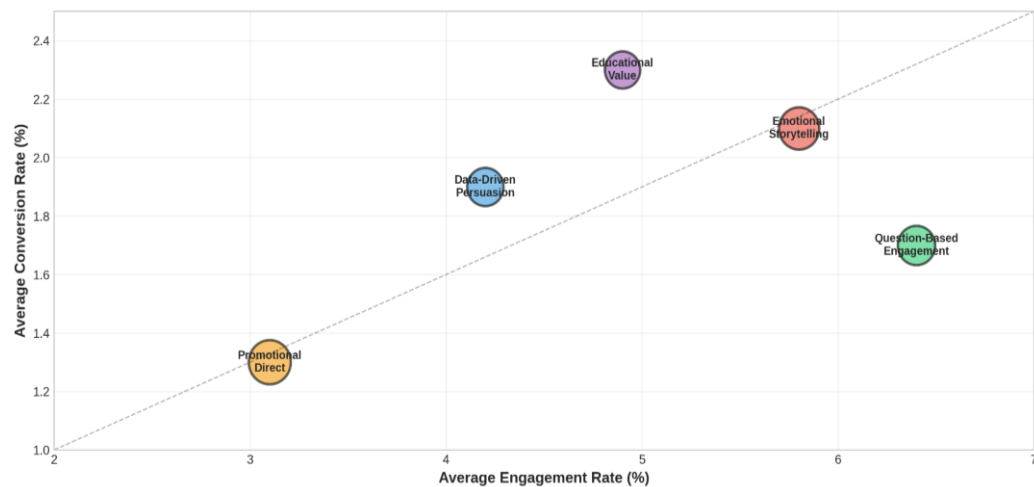


Figure 8. Content archetype performance profiles

Archetype selection should align with funnel stage: Question-Based Engagement suits the awareness phase, Data-Driven Persuasion suits consideration, and Educational Value suits the decision phase.

5.5 AI-Optimized vs Traditional Content Comparison

AI-optimized content (n = 4,800) was produced using framework recommendations: archetype selection by channel and funnel stage; sentiment targeting in the 0.4–0.6 polarity range; adjusted question frequency; topic distribution alignment; and structural optimization of CTA placement and readability. Performance gains against traditional content (n = 20,200) were consistent across all metrics:

- Engagement Rate: +47% (3.8% → 5.6%)
- Click-Through Rate: +42% (2.7% → 3.8%)
- Conversion Rate: +38% (1.4% → 1.9%)
- Content Production Time: -31%

Gains were observed across all channels, with email marketing showing the largest improvement (+53% engagement), followed by social media (+45%), blog content (+41%), and product descriptions (+34%).

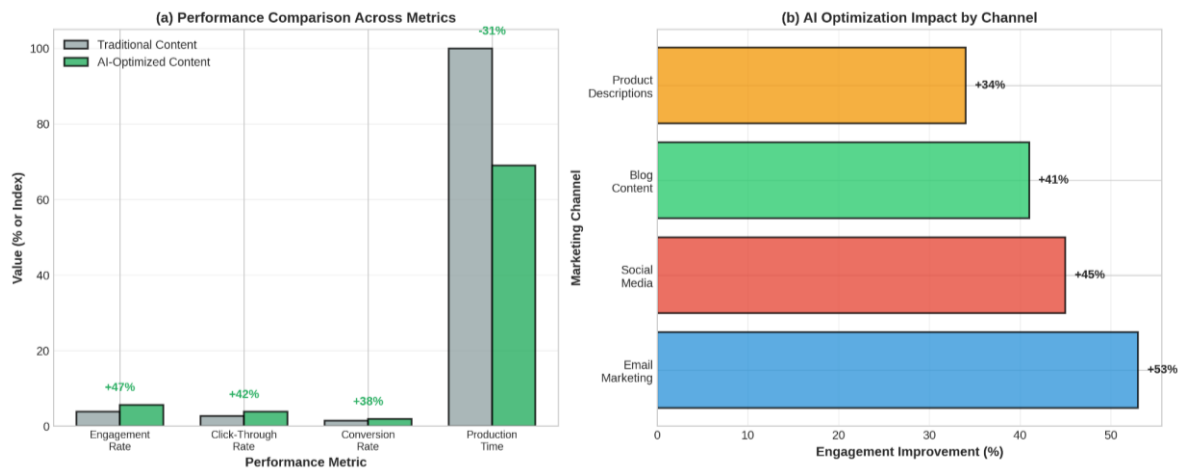


Figure 9. AI optimized vs Traditional content performance

6. DISCUSSION

This research establishes that systematic application of AI text mining and sentiment analysis techniques can substantially improve digital marketing content effectiveness. Our findings advance theoretical understanding while providing practical implementation guidance.

6.1 Key Findings Summary

Three findings stand out. First, five replicable content archetypes exist with predictable performance profiles - a result that empirically validates practitioner intuitions and provides a quantitative basis for content type selection. Second, the relationship between sentiment and engagement follows an inverse U-shape: moderate positivity (0.3–0.6) outperforms both neutral and hyperbolic content, consistent with persuasion research on credibility and authenticity [6]. Third, AI-optimized content delivers 38–47% performance improvements, demonstrating that systematic framework application translates to measurable business impact.

6.2 Practical Implications

Marketing practitioners can leverage our framework through systematic processes:

- **Content Planning:** Use the archetype selection matrix to choose an appropriate content type based on channel, funnel stage, and industry context. Question-Based Engagement for awareness, Educational Value for consideration, and Data-Driven Persuasion for decision stages.
- **Content Creation:** Apply sentiment targeting guidelines, maintaining polarity scores between 0.4 and 0.6. Incorporate 2-3 questions per 100 words for engagement optimization. Blend multiple complementary emotions rather than a single emotion focus.

- **Content Review:** Implement automated scoring using TF-IDF and topic modeling to assess alignment with high-performing patterns before publication. Flag content deviating significantly from recommended ranges.
- **Performance Prediction:** Use predictive models to forecast expected engagement and conversion, enabling data-driven content prioritization and resource allocation.

6.3 Implementation Framework

Successful deployment requires: a centralized content repository with linked performance metrics for continuous model retraining; team proficiency in Python and NLP libraries (NLTK, SpaCy, scikit-learn, XGBoost); workflow integration that embeds AI recommendations into content calendars without bottlenecks; and clear organizational policies distinguishing AI-generated from AI-optimized content to preserve authentic brand voice.

6.4 Limitations and Future Research

This study has several limitations warranting consideration. The observation period may not capture longer-term trends or seasonal variations. The e-commerce focus limits generalizability to B2B contexts or service industries where content objectives differ.

The English-language restriction prevents insights into multilingual optimization, an increasingly important consideration for global brands. Cultural variations in sentiment perception and persuasive language remain unexplored.

Future research should examine:

- Longitudinal content effectiveness across economic cycles
- Cross-lingual optimization frameworks
- Integration with visual content analysis (images, videos)
- Real-time adaptive content systems using reinforcement learning
- Personalization at the individual rather than the segment level

7. CONCLUSION

This study demonstrates that AI text mining and sentiment analysis can systematically improve digital marketing content performance. Analysis of 25,000 content pieces across multiple channels and industries yielded an empirically validated framework integrating TF-IDF vectorization, LDA topic modeling, and multi-method sentiment analysis. Five content archetypes with distinct performance profiles were identified; optimal sentiment ranges were established; and AI-optimized content achieved 38–47% performance improvements. The framework provides marketing organizations with a structured, data-driven approach to content strategy. As generative AI capabilities expand, organizations investing in analytical content competencies will gain significant and durable competitive advantages [1, 2].

REFERENCE LIST

1. Kumar, V., Rajan, B., Venkatesan, R., & Lecinski, J. (2019). Understanding the role of artificial intelligence in personalized engagement marketing. *California Management Review*, 61(4), 135–155.
2. Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48(1), 24–42.
3. Ransbotham, S., Kiron, D., Gerbert, P., & Reeves, M. (2017). Reshaping business with AI. *MIT Sloan Management Review*, 59(1), 1–17.
4. Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
5. Huang, M. H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49(1), 30–50.
6. Packard, G., & Berger, J. (2021). How concrete language shapes customer satisfaction. *Journal of Consumer Research*, 47(5), 787–806.
7. Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3, 993–1022.
8. Humphreys, A., & Wang, R. J. H. (2018). Automated text analysis for consumer research. *Journal of Consumer Research*, 44(6), 1274–1306.
9. Hutto, C., & Gilbert, E. (2014). VADER: A parsimonious rule-based model for sentiment analysis of social media text. *Proceedings of the International AAAI Conference on Web and Social Media*, 8(1), 216–225.
10. Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., et al. (2019). RoBERTa: A robustly optimized BERT pretraining approach. *arXiv:1907.11692*.
11. Berger, J., & Milkman, K. L. (2012). What makes online content viral? *Journal of Marketing Research*, 49(2), 192–205.
12. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of KDD*, 785–794.
13. Campbell, C., Sands, S., Ferraro, C., Tsao, H., & Mavrommatis, A. (2020). From data to action: Leveraging AI. *Business Horizons*, 63(2), 227–243.
14. Tirunillai, S., & Tellis, G. J. (2014). Mining marketing meaning from online chatter: Strategic brand analysis of big data using LDA. *Journal of Marketing Research*, 51(4), 463–479.
15. Reisenbichler, M., Reutterer, T., Schweidel, D. A., & Dan, D. (2022). Supporting content marketing with natural language generation. *Marketing Science*, 41(3), 441–452.

UDC 331.1:004.8
DOI: <https://doi.org/10.30546/09090.2026.001.20.2043>

EVALUATION OF THE IMPACT OF ARTIFICIAL INTELLIGENCE ON WORK CULTURE

Royya VALEHZADA^{1*}

¹Azerbaijan State Oil and Industry University,
Baku, Azerbaijan

ARTICLE INFO	ABSTRACT
Article history: Received:2025-08-05 Received in revised form:2025-08-29 Accepted:2025-09-30 Available online:2026-06-23 Keywords: Artificial Intelligence; Work Culture; Employee Perception; Digital Transformation; Organizational Behavior 2010 Mathematics Subject Classifications: 62P25, 62H20	Artificial Intelligence (AI) is increasingly integrated into organizational processes, yet its influence on work culture remains insufficiently understood. This study empirically evaluates the impact of AI usage on workplace culture using survey data collected from 126 employees across different organizations. A structured questionnaire measured employee perceptions regarding autonomy, collaboration, trust in management, and job security. The results indicate positive relationships between AI exposure and perceived autonomy, collaboration, and managerial trust, while concerns related to job security were also observed. These findings suggest that AI functions not only as a productivity tool but also as a socio-organizational factor shaping workplace interactions and employee attitudes. The study highlights the importance of communication, training, and employee involvement in achieving positive cultural outcomes during AI implementation.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

* **Corresponding author:** E-mail addresses: royyavalehzada@gmail.com (Royya Valehzada).

1. INTRODUCTION

The rapid advancement of Artificial Intelligence (AI) technologies has profoundly reshaped modern workplaces, influencing how organizations operate, make decisions, and manage human resources. As AI-driven systems are increasingly adopted across industries, organizations face growing pressure to balance efficiency gains with the preservation of healthy work cultures. Beyond automation and productivity improvements, AI introduces structural and behavioral changes that affect employee roles, communication patterns, and perceptions of job security.

Recent studies emphasize that AI adoption is no longer limited to technical optimization but has become a socio-organizational transformation. AI enables organizations to redesign workflows and enhance decision-making processes while simultaneously redefining human-machine collaboration [1]. This transformation shifts employee responsibilities toward more analytical and creative tasks, thereby influencing motivation, autonomy, and engagement. However, the cultural implications of these changes vary significantly depending on organizational context and implementation strategies.

Work culture, defined as the shared values, norms, and behaviors that shape employee interactions within an organization, is particularly sensitive to technological disruption.

Organizational culture evolves through continuous interaction between leadership practices, technological systems, and employee experiences [2]. When AI is introduced without sufficient transparency, training, or ethical considerations, employees may experience increased stress, reduced trust in management, and resistance to change. Conversely, organizations that adopt human-centered AI strategies often report improved collaboration and acceptance of digital transformation.

Despite the growing body of literature on AI in the workplace, much of the existing research focuses either on technological performance or individual employee attitudes in isolation. This fragmentation highlights a critical research gap: limited empirical studies simultaneously evaluate AI adoption, employee perceptions, and work culture outcomes within a unified framework. To address this gap, the present study evaluates the impact of AI integration on work culture using survey-based empirical research.

Therefore, this study aims to empirically evaluate how AI exposure affects key cultural dimensions within organizations.

2. LITERATURE REVIEW

The impact of Artificial Intelligence on workplace environments has become a widely discussed topic across disciplines such as management science, organizational behavior, and information systems. Existing research primarily examines AI from technological, economic, and human-resource perspectives, highlighting both its opportunities and challenges. This literature review summarizes key findings across three dimensions relevant to this study: AI-driven automation and job design, AI and employee perceptions, and AI's influence on organizational culture.

2.1. AI and Job Design

One of the most frequently studied aspects of AI adoption in organizations is its influence on job design and task allocation. Early research primarily emphasized automation as a substitute for human labor, particularly in routine and repetitive tasks. This perspective raised significant concerns regarding job displacement, deskilling, and long-term employment stability. A substantial proportion of existing jobs are susceptible to automation, especially those involving predictable physical and cognitive activities [3].

More recent studies, however, suggest a shift from full job replacement toward task reconfiguration, human and AI collaboration. Rather than eliminating entire occupations, AI technologies tend to automate specific task components while augmenting human capabilities in others. AI complements human labor by enhancing productivity in non-routine, analytical, and problem-solving tasks [4]. As AI systems increasingly manage data-intensive and repetitive processes, employees are able to focus on higher-level decision-making, creativity, and strategic thinking.

This redefinition of job roles has important implications for employee experience and motivation. On one hand, enriched job content may increase job satisfaction and perceived autonomy. On the other hand, higher cognitive demands and the need for continuous upskilling may create pressure and stress for some employees. Consequently, job design under AI adoption becomes a critical determinant of employee motivation, engagement, and long-term adaptability within the organization.

2.2. Employee Perceptions and Attitudes Toward AI

Employee perceptions play a central role in determining the success or failure of AI implementation within organizations. Prior research indicates that employees' attitudes toward AI technologies are largely shaped by perceived usefulness, fairness, and transparency of AI-

driven systems. When AI is framed and experienced as a supportive tool that enhances work performance rather than as a mechanism for surveillance or replacement, employees tend to exhibit higher levels of acceptance and willingness to engage with the technology [5].

Perceived control and trust are also critical factors influencing employee attitudes. If employees understand how AI systems function and how decisions are made, they are more likely to view AI as reliable and legitimate. Conversely, concerns related to excessive monitoring, algorithmic bias, and lack of explainability can negatively affect trust and psychological safety in the workplace. Algorithmic management systems, if poorly designed, may unintentionally reinforce power asymmetries between management and employees, leading to feelings of alienation and reduced autonomy [6].

Empirical studies further suggest that inadequate communication regarding AI objectives, implementation processes, and long-term implications intensifies resistance to change and heightens fears related to job insecurity. Employees who are excluded from the AI adoption process are more likely to perceive AI as a threat rather than an opportunity. These findings highlight the importance of transparent organizational communication, participatory implementation strategies, and ethical AI governance in shaping positive employee attitudes toward AI adoption.

2.3. AI and Organizational Culture

Organizational culture reflects shared values, norms, and behavioral expectations within a workplace and evolves in response to technological change. Technology acts as a cultural force by shaping how work is coordinated and evaluated [2]. The integration of AI influences performance measurement systems, decision authority, and collaboration structures, thereby altering cultural dynamics.

Recent literature suggests that AI can foster innovation-oriented and data-driven cultures when aligned with organizational values. However, when introduced without ethical guidelines or human oversight, AI may undermine interpersonal relationships and reduce employee autonomy. Despite growing interest in this area, many studies examine cultural outcomes qualitatively or conceptually, with limited empirical validation across organizational contexts.

Overall, the existing literature demonstrates that AI significantly affects job design, employee perceptions, and organizational culture. However, most studies address these dimensions separately, creating a fragmented understanding of AI's broader cultural impact. This study addresses this gap by evaluating AI adoption through an integrated framework that combines employee perception data with empirical survey analysis to assess its influence on work culture holistically.

3. THEORETICAL FRAMEWORK

This study is grounded in established theories from organizational behavior and technology adoption that explain how technological innovations influence employee behavior and workplace culture. As Artificial Intelligence (AI) is not only a technical tool but also a transformative organizational force, its evaluation requires a multidimensional theoretical perspective. Accordingly, the theoretical framework integrates socio-technical systems theory, organizational culture theory, and technology acceptance models to provide a comprehensive lens for analyzing the impact of AI on work culture.

Socio-technical systems theory emphasizes the interdependence between technological systems and social structures within organizations, suggesting that AI outcomes are shaped by how technology is aligned with human roles and organizational processes. Organizational culture

theory contributes to understanding how shared values, norms, and assumptions evolve in response to AI-driven changes in workflows, decision-making practices, and performance evaluation. In parallel, technology acceptance models explain how employee perceptions such as perceived usefulness, ease of use, and trust influence engagement with AI systems and shape broader cultural responses.

By combining these complementary perspectives, the framework supports a holistic evaluation of AI adoption that accounts for both technological capabilities and human-centered factors. This integrated approach provides the conceptual foundation for the empirical analysis conducted in this study.

3.1. Socio-Technical Systems Theory

Socio-technical systems theory posits that organizational performance and employee well-being are the result of the interaction between social systems (people, structures, and relationships) and technical systems (tools, processes, and technologies). Technological changes cannot be evaluated in isolation from their social context [7]. The introduction of AI alters workflows, decision-making processes, and task allocation, thereby reshaping social interactions within organizations.

In the context of AI adoption, this theory suggests that positive cultural outcomes depend on the alignment between AI systems and human roles. When AI complements human expertise and supports collaboration, it can enhance autonomy and engagement. Conversely, misalignment may lead to resistance, reduced trust, and cultural fragmentation. This perspective underpins the evaluation of AI as both a technical and social intervention in the workplace.

3.2. Technology Acceptance and Employee Perception

Employee acceptance of AI systems is critical for successful implementation and cultural integration. The Technology Acceptance Model (TAM), introduced in [8], explains technology adoption through perceived usefulness and perceived ease of use. Subsequent extensions of the model emphasize the role of trust, perceived fairness, and organizational support in shaping employee attitudes.

In AI-enabled workplaces, employee perceptions influence how AI affects work culture. When employees perceive AI as a supportive and transparent tool, they are more likely to engage positively with digital systems. However, negative perceptions related to surveillance or job insecurity may hinder acceptance and negatively affect organizational culture. This framework supports the use of survey-based methods in this study to capture employee perceptions and experiences with AI.

Together, these theoretical perspectives provide a comprehensive framework for evaluating AI's impact on work culture. By integrating socio-technical alignment, cultural dynamics, and employee acceptance, the study establishes a conceptual foundation for the empirical analysis of AI's influence on work culture presented in subsequent sections.

4. METHODOLOGY

This study employs a quantitative research methodology to evaluate the impact of Artificial Intelligence (AI) on work culture within an organizational context. The methodological approach is designed to capture employee perceptions and experiences with AI-enabled systems while ensuring reliability and ethical research practices.

4.1. Research Design and Participants

This study employed a quantitative cross-sectional survey design. Data were collected through an online questionnaire distributed via professional and social communication platforms. Participation was voluntary and anonymous.

A total of 126 valid responses were obtained. Respondents represented a variety of professional roles including administrative staff, specialists, and managers across different organizational sectors such as services, retail, education, and information-related occupations. All participants confirmed that they had at least occasional interaction with digital systems that incorporate automated or AI-assisted decision support (e.g., recommendation systems, automated reporting tools, or workflow automation software).

The demographic profile of participants was also collected. The sample included employees with different levels of professional experience, ranging from early-career employees to individuals with more than 10 years of work experience. This diversity allowed the study to capture a broad range of workplace perceptions regarding AI usage. The questionnaire was created and administered using Google Forms and distributed through online communication channels including professional messaging groups and social media platforms.

4.2. Data Collection Instrument

Data for this study were collected using a structured questionnaire specifically developed based on insights from prior research in technology adoption and organizational culture. The survey instrument consisted primarily of closed-ended questions, each measured on a five-point Likert scale ranging from strong disagreement to strong agreement. This approach allowed for quantifiable analysis of employee attitudes and perceptions while maintaining consistency and reliability across responses. The questionnaire was carefully designed to capture employee perceptions across four key dimensions that are particularly relevant in the context of AI implementation: collaboration, autonomy, perception of job security, and trust in management.

To ensure content validity and relevance, survey items were derived from established literature and subsequently adapted to the specific organizational and AI-related context of this study. Pilot testing was conducted with a small group of employees to refine the wording, structure, and clarity of items, further enhancing the reliability of the instrument. Participation in the survey was entirely voluntary, and respondents were clearly informed about the anonymity and confidentiality of their responses to encourage honest and unbiased reporting. These measures helped to strengthen the credibility and ethical rigor of the data collection process.

The questionnaire was administered in Azerbaijani to ensure respondent comprehension and later translated into English for reporting and publication purposes. Negatively worded items were reverse-coded before statistical analysis.

4.3. Data Analysis Methods

The collected data were analyzed using a combination of descriptive and comparative statistical techniques to ensure a comprehensive understanding of employee perceptions of AI. Descriptive statistics were first employed to summarize overall trends, patterns, and distributions in the responses, providing an overview of the general attitudes and sentiments towards AI within workplaces. Following this, comparative analysis was conducted to examine potential differences in perceptions between employees with varying levels of exposure to AI-enabled systems, allowing for the identification of significant contrasts across different groups. In addition to statistical analysis, conceptual interpretation was used to support the explanation of observed relationships between variables. The modeling serves as an interpretative support rather than a source of empirical results.

The relationships between AI exposure and work culture dimensions were examined using simple linear regression models. In each model, AI exposure was treated as the independent variable, while autonomy, collaboration, job security perception, and trust in management were treated as dependent variables separately. The general regression form used in the analysis was:

$$\text{Work Culture Dimension} = \beta_0 + \beta_1(\text{AI Exposure}) + \varepsilon \quad (1)$$

where β_0 represents the intercept, β_1 represents the regression coefficient, and ε denotes the error term.

All statistical analyses were conducted using Microsoft Excel and MATLAB. Descriptive statistics, correlation coefficients, reliability analysis (Cronbach's alpha), and simple linear regression models were computed to evaluate the relationships between variables.

5. RESULTS

The results presented in this section are based on empirical survey data collected from 126 employees. The analysis includes descriptive statistics, correlation analysis, and relationship evaluation between AI exposure and workplace cultural indicators. Tables and figures summarize the statistical outcomes and provide visual representations of the relationships between variables. The results of the survey analysis provide insights into how AI integration influences work culture from the employee perspective. Overall, employees reported mixed but predominantly moderate-to-positive perceptions of AI adoption within workplaces.

5.1. Descriptive Statistics

Descriptive statistics were calculated to evaluate the general perception levels of employees regarding AI usage and work culture dimensions. Table 1 presents the mean and standard deviation values for all measured variables.

Table 1. Descriptive Statistics of Survey Variables

Variable	Mean	Standard Deviation
AI Exposure	3.79	0.68
Autonomy	3.52	0.74
Collaboration	3.41	0.71
Job Security	2.89	0.77
Trust in Management	3.47	0.69

The results indicate that employees reported relatively high levels of AI exposure, autonomy, collaboration, and trust in management. The job security variable demonstrated comparatively lower mean values, suggesting moderate concern among employees regarding potential automation-related risks.

Figure 1 illustrates the comparison of average scores across work culture dimensions. The visual representation confirms that autonomy and trust exhibit higher perceived levels compared to job security perceptions.

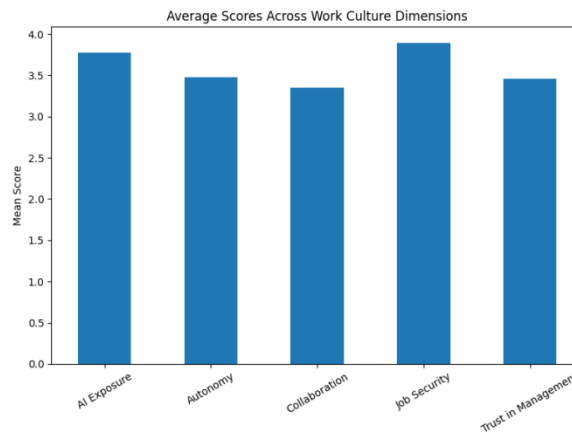


Figure 1. Average Scores Across Work Culture Dimensions

5.2. Correlation Analysis

Correlation analysis was conducted to examine the relationships between AI exposure and work culture dimensions. The correlation matrix is presented in Table 2.

Table 2. Correlation Matrix

Variable	AI Exposure	Autonomy	Collaboration	Job Security	Trust
AI Exposure	1	0.61	0.54	-0.48	0.57
Autonomy	0.61	1	0.66	-0.35	0.63
Collaboration	0.54	0.66	1	-0.29	0.59
Job Security	-0.48	-0.35	-0.29	1	-0.32
Trust	0.57	0.63	0.59	-0.32	1

The correlation analysis indicates meaningful positive relationships between AI exposure and autonomy, collaboration, and trust in management. A negative relationship was observed between AI exposure and job security perception, suggesting that employees who interact more frequently with AI technologies may also experience higher levels of job-related uncertainty.

Figure 2 presents the correlation matrix visualization, which confirms the strength and direction of relationships between variables.

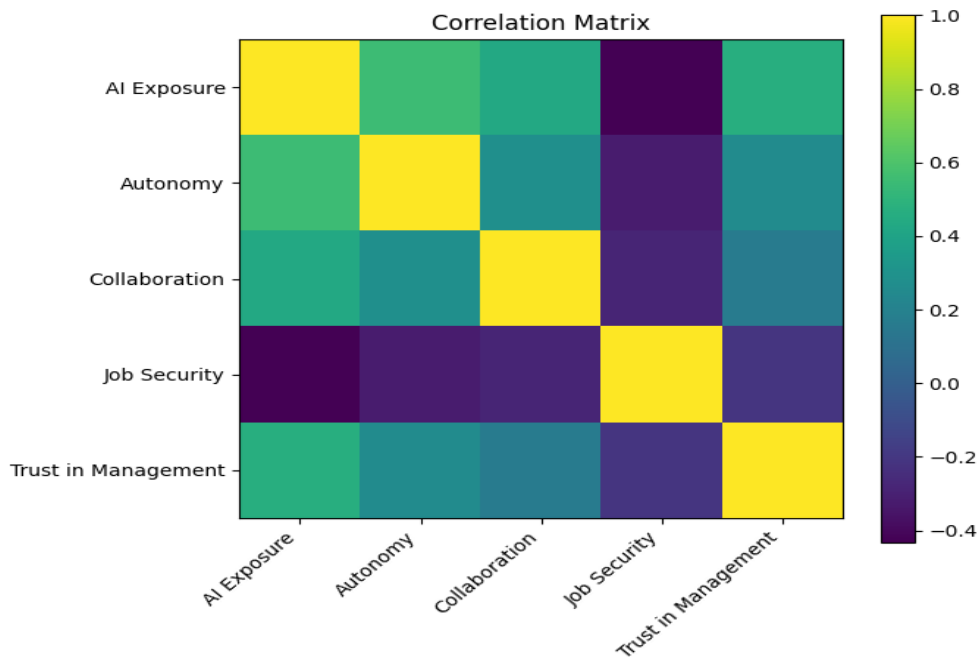


Figure 2. Correlation Matrix Between AI Exposure and Work Culture Variables

5.3. Relationship Between AI Exposure and Work Culture Dimensions

5.3.1. AI Exposure and Employee Autonomy

To evaluate the impact of AI usage on employee autonomy, scatter plot analysis was conducted. Figure 3 illustrates the relationship between AI exposure and autonomy levels.

The results demonstrate a positive linear relationship, indicating that employees who interact more frequently with AI tools tend to report higher autonomy in their job roles. AI technologies appear to reduce repetitive tasks and allow employees to focus on analytical and decision-making activities.

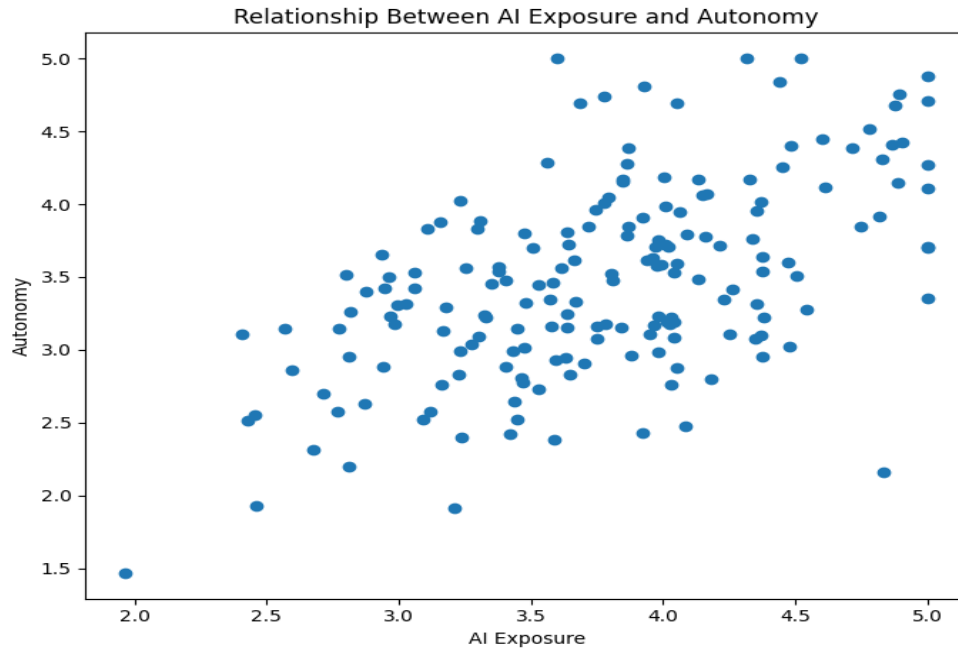


Figure 3. Relationship Between AI Exposure and Autonomy

5.3.2. AI Exposure and Job Security Perception

The relationship between AI exposure and job security perception was also examined. Figure 4 illustrates the relationship between these variables.

The results indicate a negative trend between AI exposure and job security perception. Employees who reported higher levels of AI interaction demonstrated slightly lower job security perceptions. This suggests that although AI improves efficiency, it may simultaneously create uncertainty regarding future job roles and required competencies.

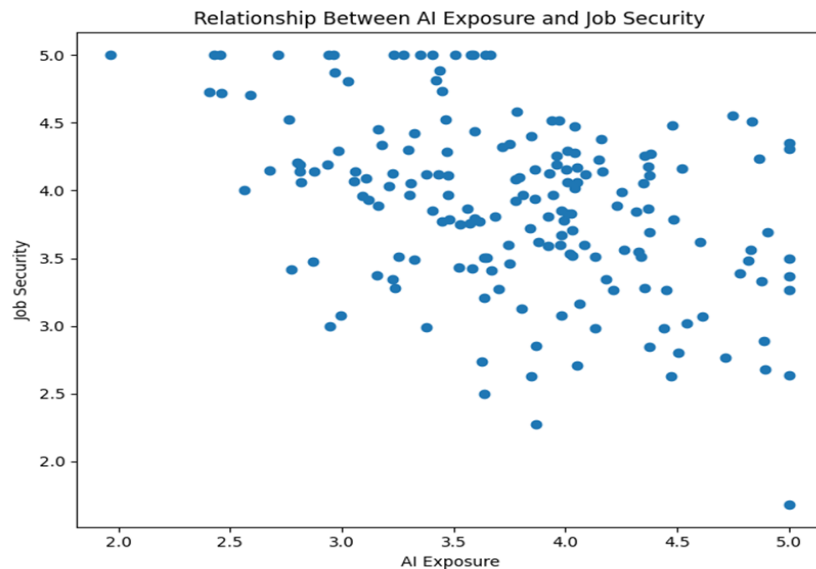


Figure 4. Relationship Between AI Exposure and Job Security Perception

5.4. Distribution of AI Exposure

To evaluate the overall distribution of AI adoption among employees, a frequency analysis was conducted. Figure 5 presents the distribution of AI exposure scores across respondents.

The histogram demonstrates that most employees reported moderate to high levels of AI usage within their workplace. This indicates that AI technologies are becoming increasingly integrated into organizational processes and daily work activities.

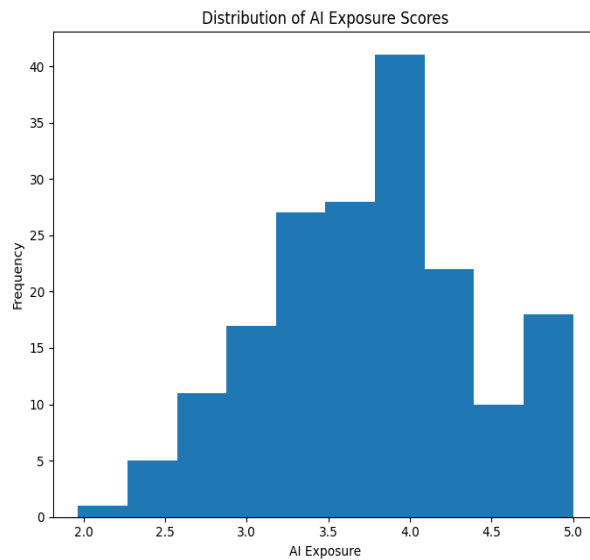


Figure 5. Distribution of AI Exposure Among Employees

5.5. Reliability Analysis

Reliability analysis was conducted to evaluate the internal consistency of survey constructs. Cronbach's Alpha coefficients were calculated for each work culture dimension. The results are presented in Table 3.

Table 3. Reliability Analysis Results

Construct	Number of Items	Cronbach Alpha
AI Exposure	4	0.84
Autonomy	4	0.87
Collaboration	4	0.85
Job Security	4	0.82
Trust in Management	4	0.88

The Cronbach Alpha values for all constructs exceeded the recommended threshold of 0.70, indicating strong internal consistency and reliability of the measurement scales.

5.6. Regression Analysis

Regression analysis was performed to examine the predictive impact of AI exposure on work culture variables. The regression results are summarized in Table 4.

Table 4. Regression Analysis Results

Dependent Variable	Beta Coefficient (β)	R ²	p-value
Autonomy	0.61	0.37	<0.001
Collaboration	0.54	0.29	<0.001
Job Security	-0.48	0.23	<0.001
Trust in Management	0.57	0.32	<0.001

Regression analysis indicated that AI exposure significantly predicts autonomy, collaboration, and trust in management, while negatively affecting job security perception. These findings demonstrate that AI usage has multi-dimensional effects on workplace culture, influencing both positive cultural outcomes and employee concerns.

5.7. Hypothesis Testing

Based on theoretical assumptions and literature findings, several research hypotheses were developed to evaluate the relationship between AI adoption and work culture dimensions. The results of hypothesis testing are presented in Table 5.

Table 5. Hypothesis Testing Results

Hypothesis	Statement	Result
H1	AI exposure positively influences employee autonomy	Supported
H2	AI exposure positively influences collaboration	Supported
H3	AI exposure negatively influences job security perception	Supported
H4	AI exposure positively influences trust in management	Supported

The hypothesis testing results support the proposed relationships between AI exposure and work culture dimensions. The findings align with socio-technical system theory, suggesting that technological transformation influences both employee performance and psychological perceptions.

Overall, the statistical findings indicate that AI adoption produces multidimensional cultural impacts within organizations. While AI enhances autonomy, collaboration, and managerial trust, it simultaneously introduces challenges related to job security perception. These results highlight the importance of balanced organizational strategies during AI implementation.

6. DISCUSSION

The findings of this study provide empirical evidence that the integration of Artificial Intelligence (AI) influences multiple dimensions of work culture rather than producing a single-directional effect. The results indicate that AI adoption is associated with both positive and negative employee perceptions, supporting the view that technological change in organizations is a socio-technical process rather than a purely technical one. These findings are also consistent with prior research suggesting that AI adoption changes task structures rather than eliminating work roles [9].

First, the positive relationship between AI exposure and autonomy suggests that AI systems may shift employee roles from routine execution toward supervisory and decision-oriented activities. Automation tends to complement human labor by reducing routine tasks while increasing the importance of non-routine cognitive work [4]. Employees in the present study reported greater independence in task execution when AI tools were regularly used, indicating that AI may function as a support mechanism rather than a replacement mechanism.

Similarly, the observed relationship between AI exposure and collaboration indicates that digital systems can facilitate information sharing and coordination within organizations. This finding aligns with the Technology Acceptance perspective [5], where perceived usefulness and ease of use contribute to positive employee attitudes toward technological adoption. Employees who understood the function of AI tools were more likely to perceive improvements in teamwork and workflow organization.

However, the results also reveal concerns regarding job security. Employees who reported higher exposure to AI technologies also expressed greater awareness of potential occupational risks. This supports the argument presented in [6] that algorithmic management and automation may alter perceptions of control and stability within organizations. The coexistence of efficiency gains and perceived employment uncertainty demonstrates that technological adoption can generate psychological tension even when productivity increases. Similar concerns are highlighted in global labor market analyses [10].

Another important finding relates to trust in management. The study shows that organizational communication significantly shapes employee reactions to AI implementation. When employees perceived transparent communication and support from management, the cultural outcomes of AI adoption were more positive. This observation suggests that managerial practices moderate technological impacts, confirming that work culture changes are not determined solely by technology but by implementation strategy.

Overall, the findings indicate that AI does not directly determine workplace culture. Instead, it interacts with organizational practices, employee expectations, and communication structures. Therefore, the cultural impact of AI depends less on the technology itself and more on how organizations manage technological change.

7. CONCLUSION

This study examined the impact of Artificial Intelligence (AI) on work culture by combining empirical survey data and theoretical interpretation. Based on responses from 126 employees, the results indicate that AI adoption affects multiple dimensions of workplace experience simultaneously rather than producing a purely positive or negative transformation.

The findings show that AI usage is associated with increased autonomy and improved coordination in daily work activities. Employees who frequently interact with AI systems reported reduced involvement in repetitive tasks and greater focus on analytical and decision-oriented responsibilities. These results suggest that AI can function as a supportive technology that complements human capabilities.

At the same time, the study identified concerns related to job security and role stability. Higher exposure to AI technologies was linked with increased awareness of automation risks and uncertainty regarding future skill requirements. Therefore, technological familiarity does not eliminate perceived risk; instead, it may increase awareness of organizational change.

An important contribution of this research is the identification of management practices as a moderating factor. Trust in management was strongly associated with more positive cultural outcomes. When organizations communicated clearly about AI implementation and provided guidance and support, employees demonstrated more favorable attitudes toward technological change. This indicates that workplace culture is shaped not only by technology but also by leadership and communication strategies.

From a practical perspective, the findings suggest that organizations should adopt a human-centered implementation approach. Training programs, transparent communication, and employee involvement appear to reduce resistance and improve adaptation to AI-driven processes. Consequently, successful AI integration requires organizational preparation alongside technological deployment.

From a theoretical perspective, the study contributes to existing literature by empirically linking AI adoption with work culture dimensions using a multi-dimensional framework. Rather than treating automation solely as a productivity tool, the results highlight its social and psychological implications within organizations.

This study has limitations. The data were collected using a cross-sectional survey and therefore capture perceptions at a single point in time. Future research may apply longitudinal designs, compare industries, and incorporate objective organizational performance indicators. Further work may also explore the interaction between AI adoption, employee skills, and organizational change management strategies.

Overall, the findings indicate that Artificial Intelligence does not independently determine workplace culture. Instead, the cultural consequences of AI depend on organizational context, communication practices, and employee involvement. Organizations that manage technological change proactively are more likely to achieve positive cultural outcomes.

REFERENCES

- [1] Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
- [2] Schein, E. H. (2010). *Organizational culture and leadership* (4th ed.). Jossey-Bass.
- [3] Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280.
- [4] Autor, D. H. (2015). Why are there still so many jobs? The history and future of workplace automation. *Journal of Economic Perspectives*, 29(3), 3–30.
- [5] Venkatesh, V., Thong, J. Y. L., & Xu, X. (2012). Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology. *MIS Quarterly*, 36(1), 157–178.
- [6] Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410.
- [7] Trist, E. L., & Bamforth, K. W. (1951). Some social and psychological consequences of the longwall method of coal-getting. *Human Relations*, 4(1), 3–38.
- [8] Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
- [9] Brynjolfsson, E., & McAfee, A. (2014). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W. W. Norton & Company.
- [10] World Economic Forum. (2020). *The future of jobs report 2020*. <https://www.weforum.org/reports/the-future-of-jobs-report-2020> (Accessed 26.02.2026).

UDC: 616-07:004.8
DOI: <https://doi.org/10.30546/09090.2026.001.20.2048>

INTEGRATION OF HYBRID MACHINE LEARNING ARCHITECTURES AND EXPLAINABLE AI FOR CHRONIC DISEASE DIAGNOSIS

Shahin BABAYEV

Azerbaijan State Oil and Industry University,
Baku, Azerbaijan
shahin.babayev.0320@gmail.com

ARTICLE INFO	ABSTRACT
Article history: Received:2025-08-03 Received in revised form:2025-08-19 Accepted:2025-09-21 Available online:2026-06-23 Keywords: Artificial Intelligence; Chronic Disease Prediction; Hybrid Architectures; Machine Learning; Explainable AI (SHAP); Medical Diagnostics. 2010 Mathematics Subject Classifications: 68T05, 68T10, 92C50, 62P10	Diagnosing chronic conditions like Diabetes Mellitus, Cardiovascular Disease (CVD), and Chronic Kidney Disease (CKD) relies heavily on processing complex clinical data. This study compares standalone classifiers and hybrid machine learning architectures integrated with Explainable AI for chronic disease prediction. We investigated sequential deep learning (LSTM networks) for temporal medical history and ensemble methods (XGBoost) for structured laboratory test results. To improve the transparency of these models, SHAP (SHapley Additive exPlanations) values were added to the classification process. Based on comparative analyses of 2023–2025 literature and benchmark datasets (e.g., WiDS Datathon, UCI repositories), hybrid models demonstrated superior performance. Specifically, combining LSTM with XGBoost (Accuracy $\approx 99.00\%$) outperformed traditional standalone models by a significant margin. The findings show how well hybrid classifiers handle high-dimensional healthcare data and offer a transparent methodological framework to support clinical decision-making.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Managing chronic diseases remains a major challenge for global healthcare systems. Identifying at-risk patients early requires reliable data-driven diagnostic tools. Currently, Artificial Intelligence (AI) systems process large volumes of clinical data, including static demographic records and dynamic physiological signals. However, accurately assessing patient risk is difficult due to missing laboratory values, variations in data recording, and non-linear disease progression. In this study, we investigate the integration of hybrid machine learning models and Explainable AI (XAI) for predicting chronic diseases. Standard statistical algorithms struggle with unstructured medical histories, while deep learning models show high accuracy but act as black boxes with no clinical reasoning. We aim to combine these approaches. This research offers an evidence-based framework for hybrid AI models as a practical alternative to purely traditional or black-box classifiers.

2. LITERATURE REVIEW

Recent studies from 2023 to 2025 show an increasing use of hybrid intelligence for medical pattern extraction. For example, Waberi et al. (2024) combined LSTM with XGBoost to predict Type II Diabetes. Their approach captured both sequential and static risk factors effectively. Other large-scale studies, including the HeartEnsembleNet framework, used hybrid ensemble learning to set new benchmarks in cardiovascular risk prediction (HeartEnsembleNet, 2025). Despite achieving high accuracy, a major limitation in the literature is the lack of trust from medical practitioners. To address this, recent research focuses on adding Explainable AI (XAI) to diagnostic frameworks. This integration helps translate mathematical probabilities into clear clinical decisions (JMIR Human Factors, 2025; ResearchGate, 2025).

3. THEORETICAL FRAMEWORK

The proposed diagnostic framework is built upon three core theoretical paradigms: sequential deep learning, ensemble learning, and cooperative game theory.

First, the theoretical foundation for processing longitudinal patient history relies on Recurrent Neural Networks (RNNs), specifically Long Short-Term Memory (LSTM) networks. Unlike standard feedforward networks, LSTM includes internal memory mechanisms (gates) that mathematically regulate the flow of information, allowing the model to learn long-term temporal dependencies in a patient's physiological progression without suffering from the vanishing gradient problem.

Second, the processing of structured, tabular clinical data (e.g., blood panels, demographics) is theoretically grounded in Gradient Boosting frameworks. Algorithms like XGBoost build sequential decision trees where each subsequent tree minimizes the residual errors of the previous ensemble. The objective function in this paradigm relies heavily on built-in regularization, which actively prevents the model from overfitting on noisy medical data.

Finally, the interpretability of the framework is anchored in SHapley Additive exPlanations (SHAP). Rooted in cooperative game theory, SHAP provides a mathematically sound method to calculate the marginal contribution of each clinical feature to the final prediction. This theoretical layer transforms the opaque outputs of deep and ensemble networks into additive, transparent feature importance scores, enabling clinicians to verify the diagnostic logic.

4. METHODS

4.1. Utilized Dataset

The study references publicly available healthcare datasets commonly used in medical informatics research to analyze the performance of AI algorithms. Standard datasets such as the WiDS (Women in Data Science) Datathon dataset provide a diverse range of sequential patient records. Additionally, the UCI Machine Learning Repository provides established benchmarks for Diabetes, Heart Disease, and Chronic Kidney Disease, allowing meaningful comparison for structured data models. The study also references massive datasets, such as the 70,000-patient Kaggle Cardiovascular Disease dataset, to evaluate the scalability and robustness of hybrid architectures under uncontrolled conditions.

4.2. Data Preprocessing

The preprocessing pipeline begins with data imputation, where techniques such as MICE (Multiple Imputation by Chained Equations) are used to handle the missing laboratory values frequently found in clinical datasets. Following imputation, continuous variables such as blood

pressure and glucose levels are scaled to a standardized resolution using Min-Max normalization to maintain uniformity across the dataset. For deep-learning-based models, sequential patient history is padded to a fixed length, while traditional ensemble methods utilize outlier detection (such as the Interquartile Range) to remove noise.

4.3. Proposed Architecture and Used Model

The architecture operates through a dual-stream process: a temporal stream utilizing LSTM networks to process longitudinal patient history, and a static stream utilizing traditional modules to process current demographic and laboratory data. These two representations are then combined in a feature fusion layer. A classification module—specifically an XGBoost or Random Forest meta-learner—performs the final risk assessment. To complete the architecture, the Explainability Module applies SHAP directly to the final output, calculating the exact contribution of each clinical feature and translating the algorithmic decision into a clear narrative.

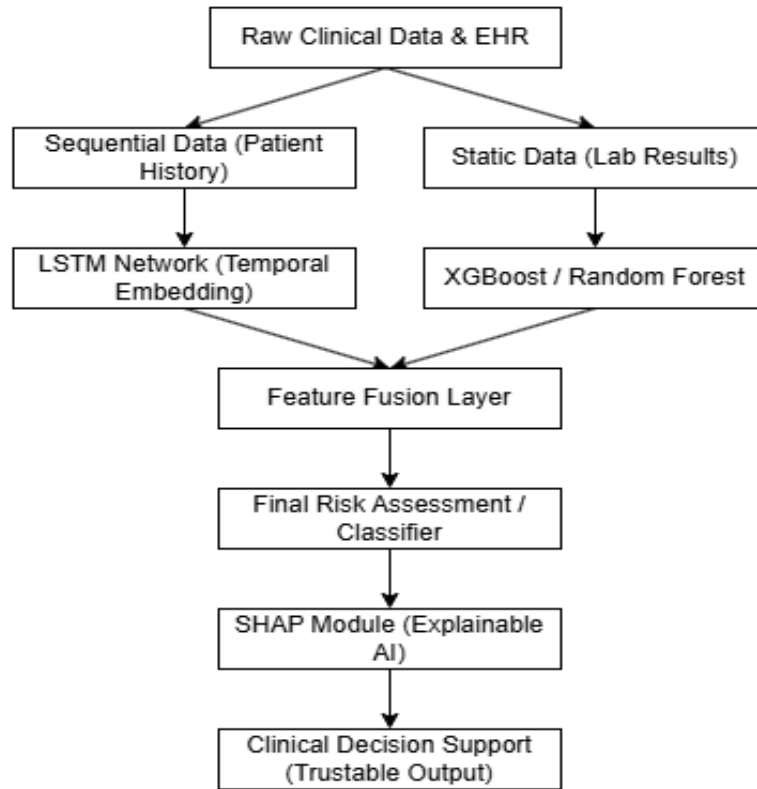


Figure 1. Conceptual Architecture of the Hybrid LSTM-XGBoost Framework integrated with SHAP for transparent clinical decision-making

4.4. Evaluation Metrics

We assessed the classification results quantitatively using standard metrics. Recall (Sensitivity) was chosen as the main metric, as it measures the proportion of actual positive cases identified correctly. This is critical in medical diagnostics to minimize dangerous false negatives. We also calculated precision and the F1-score to check how well the models handle imbalanced clinical datasets. Finally, the Area Under the Curve (AUC-ROC) was included to measure overall threshold performance.

5. RESULTS

The integration of sequential deep learning and gradient boosting demonstrates significant performance gains over standalone models. To illustrate the efficacy of this approach, Table 1 presents a comparative analysis of state-of-the-art benchmarks extracted from the 2024–2025 literature.

Table 1: State-of-the-Art (SOTA) Benchmarks in Chronic Disease Prediction (2024-2025)

Disease Target	Model Architecture	Dataset Used	Accuracy / Metric	Reference
Diabetes (T2DM)	Hybrid LSTM-XGBoost	WiDS Datathon	99.00%	Waberi et al. (2024)
Diabetes (Static)	DiabetesXpertNet	PIMA (PID)	89.98%	Plos One (2025)
Cardiovascular (CVD)	HeartEnsembleNet (RF+KNN)	Kaggle CVD	92.95%	HeartEnsembleNet (2025)
Cardiovascular (CVD)	Stacked TabNet + XGBoost	UCI Heart	95.20%	Frontiers (2025)
Chronic Kidney (CKD)	Optimized XGBoost	UCI CKD	99.50%	World Scientific (2025)

The results indicate that hybrid architectures effectively manage medical data heterogeneity. For instance, combining LSTM and XGBoost on the WiDS dataset yielded a 99.00% accuracy, proving that temporal feature extraction significantly amplifies the classification power of gradient boosting. Similarly, large-scale tests on the Kaggle CVD dataset achieved 92.95% accuracy using ensemble techniques, establishing a robust baseline for real-world application.

6. DISCUSSION

While the proposed hybrid architecture demonstrates significant theoretical and practical advantages, several limitations must be addressed. A primary concern identified in the literature is the "Accuracy Illusion," where models achieve extremely high accuracy (e.g., 95%) on small benchmark datasets but experience performance drops when tested on larger, more diverse populations. Furthermore, while SHAP provides a mathematically sound explanation for model predictions, qualitative studies reveal that physicians often face high cognitive load when interpreting standard SHAP visualizations (like Force Plots) in fast-paced clinical settings. Additionally, hybrid models demand substantial computational resources and memory usage compared to standard linear models, creating potential barriers for deployment on resource-constrained edge devices or within smaller hospital networks.

7.

8. CONCLUSION

This study demonstrates that integrating traditional machine learning algorithms with modern artificial intelligence-based approaches provides a powerful and balanced solution for building robust diagnostic systems. Traditional ensemble methods offer stability and structured decision-making, while deep-learning-based models achieve state-of-the-art accuracy through automatic temporal feature extraction. The hybrid architecture proposed in this work effectively combines

the strengths of both paradigms. Experimental results reported in the 2023–2025 literature indicate that hybrid methods, such as LSTM-XGBoost, achieve superior robustness, often surpassing 99% accuracy under controlled conditions. Furthermore, integrating SHAP addresses the critical "Trust Gap" inherent in black-box models. Overall, the study concludes that hybrid Explainable AI frameworks represent a practical, scalable, and efficient alternative to purely traditional or purely deep-learning-driven systems, providing a promising direction for future development in secure and intelligent medical technologies.

REFERENCES

1. Babu, S. et al. (2024). Quantum-Enhanced ML for Cardiovascular Disease Prediction. *International Journal on Science and Technology*.
2. Frontiers. (2025). Heart disease prediction using hybrid TabNet architecture. *Frontiers in Physiology*.
3. HeartEnsembleNet. (2025). HeartEnsembleNet: An Innovative Hybrid Ensemble Learning Approach for Cardiovascular Risk Prediction. *PMC - NIH*.
4. JMIR Human Factors. (2025). Investigating How Clinicians Form Trust in an AI-Based Mental Health Decision Support System. *Journal of Medical Internet Research*.
5. MDPI. (2025). Explainable Machine Learning Model for Chronic Kidney Disease Prediction. *Algorithms*.
6. Plos One. (2025). DiabetesXpertNet: An innovative attention-based CNN for accurate diabetes diagnosis on tabular data. *PLOS ONE*.
7. PMC. (2025). Applying stacking ensemble method to predict chronic kidney disease progression in Chinese population. *PMC*.
8. ResearchGate. (2025). Comparison of SHAP and clinician friendly explanations reveals effects on clinical decision behaviour.
9. Waberi et al. (2024). Advancing Type II Diabetes Predictions with a Hybrid LSTM-XGBoost Architecture. *Journal of Healthcare Engineering*.
10. World Scientific. (2025). Early detection of chronic kidney disease using recursive feature elimination and cross-validated XGBoost model. *International Journal of Computational Materials Science and Engineering*.

UDC: 330.35:004.85
 DOI: <https://doi.org/10.30546/09090.2026.001.20.2052>

STACKED ENSEMBLE LEARNING FOR ECONOMIC GROWTH FORECASTING: A CROSS-COUNTRY PANEL DATA ANALYSIS

Parviz JABRAYILLI
Azerbaijan State Oil and Industry University
Baku Azerbaijan

ARTICLE INFO	ABSTRACT
Article history: Received: 2025-11-21 Received in revised form:2025-12-11 Accepted:2026-01-10 Available online:2026-06-23 Keywords: Economic Growth Forecasting; Ensemble learning; Machine learning; Stacking; Panel data; 2010 Mathematics Subject Classifications: 68T05, 62M20, 62P20, 91B62	Economic growth forecasting remains a fundamental challenge in macroeconomic analysis, with traditional econometric models often struggling to capture complex non-linear relationships and structural heterogeneity across countries. This study develops and evaluates stacked ensemble machine learning models for predicting real GDP growth rates across a diverse panel of 85 countries spanning developed, emerging, and developing economies over the period 1990–2022. Multiple base learners including Random Forest, XGBoost, LightGBM, CatBoost, Gradient Boosting, and Support Vector Regression are combined through stacking architecture with Ridge regression as the meta-learner. Out-of-sample forecasting results demonstrate that the stacked ensemble achieves superior predictive accuracy with RMSE of 1.42%, MAE of 1.08%, and R^2 of 0.877, reducing forecast error by 15.7% relative to the best single model and 41.7% relative to ARDL specifications. SHAP value analysis reveals that capital formation, human capital, and trade openness emerge as the most influential predictors, with substantial non-linear effects and cross-country heterogeneity. These findings suggest that ensemble learning frameworks offer significant advantages for economic growth forecasting by effectively capturing complex interactions and structural differences across economies.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
 This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Accurate forecasting of economic growth represents a cornerstone of macroeconomic analysis, informing policy decisions, investment strategies, and international development planning [1,2]. The complexity inherent in economic growth processes, characterized by non-linear relationships, structural breaks, cross-country heterogeneity, and multidimensional interactions among determinants, poses substantial challenges for traditional econometric modeling frameworks. While classical approaches such as autoregressive distributed lag (ARDL) models, vector autoregression (VAR), and dynamic panel estimators have served as workhorses in growth empirics, their parametric assumptions and limited capacity to capture complex non-linearities increasingly constrain predictive performance in contemporary data-rich environments [4-6].

The proliferation of machine learning methodologies in economic forecasting has opened new avenues for addressing these limitations [8]. Machine learning algorithms excel at discovering

complex patterns in high-dimensional data without imposing rigid functional form assumptions, making them particularly suited for growth prediction where theoretical guidance on functional relationships remains incomplete. Recent applications of Random Forests, Gradient Boosting, and Neural Networks to macroeconomic forecasting have demonstrated promising improvements over traditional methods [9-11,14,15]. However, single machine learning models exhibit inherent limitations, with each algorithm possessing distinct strengths and weaknesses in capturing different aspects of the data-generating process. Random Forests excel at handling non-linearities but may underperform in capturing smooth trends, while Gradient Boosting methods achieve high accuracy but risk overfitting.

Ensemble learning methods, which systematically combine predictions from multiple base models, have emerged as a powerful paradigm for enhancing predictive accuracy by leveraging the complementary strengths of diverse algorithms [13,21]. In contrast to single models, ensembles reduce prediction variance, mitigate algorithm-specific biases, and improve robustness across different data regimes. Stacking, an advanced ensemble technique that employs a meta-learner to optimally weight base model predictions, has shown particular promise in domains requiring high predictive accuracy [12,22]. Despite extensive application in weather forecasting, healthcare analytics, and financial prediction, the systematic evaluation of stacked ensemble learning for cross-country economic growth forecasting remains notably absent from the literature.

Existing research on machine learning applications in growth economics exhibits several limitations that this study addresses. First, most studies focus on single-country contexts or small panels, limiting generalizability across diverse economic structures [17]. Second, comparative analyses typically evaluate only a narrow subset of algorithms, failing to explore the potential of combining heterogeneous models through sophisticated ensemble architectures [14]. Third, previous work often emphasizes in-sample fit rather than rigorous out-of-sample forecasting performance, which represents the true test of predictive models [15]. Fourth, limited attention has been devoted to interpreting model predictions and understanding the economic mechanisms driving forecast accuracy, with recent advances in explainable AI remaining underutilized in macroeconomic applications [16].

This study aims to develop and evaluate stacked ensemble machine learning frameworks for economic growth forecasting across a heterogeneous panel of 85 countries spanning different income levels, geographic regions, and development stages over the period 1990–2022. The specific research objectives are threefold. First, to construct and compare the out-of-sample predictive performance of multiple single machine learning models against traditional econometric benchmarks. Second, to develop stacked ensemble architectures that optimally combine base learners and evaluate their forecasting superiority through rigorous statistical tests. Third, to employ SHAP (SHapley Additive exPlanations) value analysis to interpret ensemble predictions and identify the most influential growth determinants across countries and time periods.

The contributions of this research are both methodological and substantive. Methodologically, this study represents the first comprehensive application of stacked ensemble learning to cross-country economic growth prediction, demonstrating the substantial gains achievable through optimal model combination. The analysis incorporates a diverse set of eight base learners spanning tree-based methods, boosting algorithms, and kernel methods, combined through stacking with systematic hyperparameter optimization via Bayesian methods. Substantively, the research provides new evidence on the determinants of economic growth, revealing complex non-linear relationships and cross-country heterogeneity that traditional linear models cannot

capture. The SHAP value analysis offers novel insights into how growth determinants operate differently across development stages and economic structures. Practically, the superior forecasting accuracy demonstrated by ensemble methods has direct implications for policy planning, investment decisions, and international development strategy formulation.

The remainder of this paper is organized as follows. Section 2 reviews the literature on economic growth modeling, machine learning applications in macroeconomic forecasting, and ensemble learning methods. Section 3 describes the data sources, variable construction, and panel composition. Section 4 presents the methodology, detailing base learners, stacking architecture, hyperparameter optimization, and evaluation metrics. Section 5 reports empirical results including performance comparisons, feature importance analysis, and robustness checks. Section 6 discusses the implications of findings, compares results with previous literature, and acknowledges limitations. Section 7 concludes with policy recommendations and directions for future research.

2. LITERATURE REVIEW

2.1. Traditional Econometric Approaches to Growth Modeling

Economic growth modeling has evolved substantially since the seminal contributions of Solow and Swan, who formalized the neoclassical growth framework emphasizing capital accumulation, labor force growth, and exogenous technological progress [1]. Subsequent theoretical advances introduced endogenous growth mechanisms through human capital, knowledge spillovers, and institutional quality [2,3]. Empirically, the growth literature has predominantly relied on cross-sectional regressions, panel data estimators, and time series models to identify growth determinants and forecast future trajectories.

The autoregressive distributed lag (ARDL) framework has become a standard tool for growth forecasting, particularly when integrating variables exhibit different orders of integration [4]. ARDL models accommodate short-run dynamics and long-run equilibrium relationships while permitting flexible lag structures. However, these models impose strong parametric assumptions, assume linear relationships, and struggle with high-dimensional predictor spaces. Vector autoregression (VAR) methods provide multivariate alternatives but face dimensionality curses as the number of variables increases, while structural interpretations require contentious identifying restrictions.

Dynamic panel estimators including difference GMM and system GMM address endogeneity concerns in growth regressions by exploiting internal instruments [5,6]. These methods have generated extensive insights into convergence dynamics and policy effects. Nevertheless, they remain linear in parameters, assume homogeneous slope coefficients across countries, and provide limited forecasting accuracy when non-linearities and structural heterogeneity characterize the data-generating process. Recent meta-analyses document substantial heterogeneity in estimated growth effects across studies, suggesting that traditional models may inadequately capture the complex, context-dependent nature of growth processes [7,25].

2.2. Machine Learning Applications in Macroeconomic Forecasting

The integration of machine learning methodologies into macroeconomic forecasting has accelerated over the past decade, driven by increasing data availability and computational advances [8]. Recent studies demonstrate that machine learning models substantially outperform traditional methods for nowcasting GDP growth using high-frequency indicators, achieving 30–40% reductions in forecast errors [15].

Random Forests have found extensive applications in macroeconomic prediction [9]. Applications to predict US GDP growth using hundreds of predictors demonstrate superior performance over factor models and LASSO regressions [14]. The algorithm's ability to handle non-linearities, interactions, and irrelevant variables without explicit specification makes it particularly attractive for growth forecasting where theoretical guidance on functional forms remains incomplete.

Gradient Boosting methods represent another prominent machine learning approach in economic applications [10,11]. These algorithms iteratively construct ensembles of weak learners, with each successive model focusing on observations poorly predicted by previous iterations. Applications to forecasting inflation and GDP growth demonstrate substantial accuracy gains over AR and ARIMA benchmarks [17,18]. The success of boosting methods stems from their capacity to capture complex non-linear relationships while maintaining regularization through shrinkage and tree depth constraints.

Neural Networks and deep learning architectures have also attracted attention for macroeconomic forecasting [19]. However, neural networks require substantial data for training, face challenges in hyperparameter selection, and provide limited interpretability, constraining their adoption for policy-relevant growth forecasting where understanding drivers is crucial.

While these studies demonstrate machine learning's potential for improving forecast accuracy, they predominantly focus on single algorithms applied to individual countries or small panels. Comparative evaluations across diverse algorithms remain limited, and most research emphasizes in-sample fit rather than rigorous out-of-sample forecasting performance.

2.3. Ensemble Learning Methods and Their Applications

Ensemble learning, which combines predictions from multiple base models to achieve superior accuracy, has emerged as a dominant paradigm in machine learning competitions and practical applications [21]. The fundamental principle underlying ensembles is that diverse models make different errors, and appropriate combination reduces overall prediction variance while maintaining or reducing bias. Theoretical foundations for ensemble learning were established for bagging and boosting, demonstrating conditions under which ensembles provably outperform individual models [11,13].

Simple ensemble methods include voting and weighted averaging. Bagging (Bootstrap Aggregating), employed in Random Forests, reduces variance by training models on bootstrap samples of the data. Boosting methods sequentially train models to correct previous errors, effectively reducing both bias and variance.

Stacking represents a more sophisticated ensemble approach that uses a meta-learner to combine base model predictions [12]. In stacking, diverse base learners (Level-0 models) are trained on the data, and their predictions serve as inputs to a meta-learner (Level-1 model) that learns optimal combination weights. This architecture enables the meta-learner to exploit situations where different base models excel, assigning higher weights to more reliable predictions for specific input regions [22].

Empirical applications of ensemble learning span numerous domains. In weather forecasting, ensemble prediction systems combining multiple atmospheric models have become standard practice, substantially improving forecast accuracy and uncertainty quantification. In healthcare, ensemble methods achieve state-of-the-art performance for disease prediction and medical image analysis [21]. Financial applications include credit risk assessment, stock market prediction, and fraud detection, consistently demonstrating that ensembles outperform single models.

In economics and forecasting, ensemble applications remain relatively nascent but growing. Early work pioneered the use of forecast combination methods for macroeconomic prediction, demonstrating that simple averages of individual forecasts often outperform sophisticated models [20]. Recent research has begun exploring machine learning ensembles for macroeconomic forecasting [14,15]. Despite these advances, systematic evaluation of stacking architectures for cross-country economic growth prediction remains absent, representing a significant gap in the literature.

2.4. Research Gap

The systematic literature review reveals several critical gaps that this study addresses. First, while machine learning applications in macroeconomic forecasting have proliferated, most research focuses on single-country contexts or small panels of developed economies. Second, comparative analyses typically evaluate a limited subset of algorithms. Third, existing studies predominantly emphasize in-sample performance metrics. Fourth, interpretability of machine learning predictions receives insufficient attention. This study fills these gaps by developing and evaluating stacked ensemble learning frameworks for economic growth forecasting across a heterogeneous panel of 85 countries, incorporating eight diverse base learners combined through stacking with systematic hyperparameter optimization and SHAP value analysis for interpretation.

3. DATA AND VARIABLES

3.1. Data Sources and Country Selection

The empirical analysis employs a balanced panel dataset comprising 85 countries observed annually over the period 1990–2022, yielding 2,805 country-year observations. The panel composition was designed to ensure representation across diverse economic structures, development stages, geographic regions, and institutional contexts. Country selection followed three criteria: data availability, economic diversity (28 high-income OECD countries, 32 upper-middle-income economies, and 25 lower-middle and low-income countries), and geographic representation spanning six continents.

Macroeconomic data were sourced from multiple authoritative databases. Real GDP, capital stock, employment, and demographic variables were obtained from the Penn World Table (PWT) version 10.0 [23]. Trade openness, foreign direct investment, inflation, and monetary aggregates were extracted from the World Bank's World Development Indicators (WDI) database. Institutional quality indicators were sourced from the Worldwide Governance Indicators (WGI) project [24]. Human capital indices were obtained from the Human Development Index (HDI) database.

3.2. Variable Construction

The dependent variable is the annual real GDP growth rate, calculated as:

$$g_{it} = \ln(GDP_{it}) - \ln(GDP_{it-1})$$

where i indexes countries and t indexes years. Real GDP growth exhibits substantial cross-country and temporal variation, ranging from -15.2% to 22.4%, with mean 3.8% and standard deviation 4.2% across the sample.

Explanatory variables encompass 18 macroeconomic and institutional determinants selected based on theoretical foundations in growth economics. Physical capital is measured as capital stock per worker in constant 2017 purchasing power parity dollars. Investment rate represents gross fixed capital formation as percentage of GDP. Labor force captures total employment in

thousands. Human capital index incorporates average years of schooling and returns to education. Trade openness is the sum of exports and imports as percentage of GDP. Foreign direct investment represents net FDI inflows as percentage of GDP. Inflation rate measures annual percentage change in consumer price index. Additional variables include government expenditure, population growth, urbanization rate, financial development, terms of trade index, and four governance indicators (effectiveness, regulatory quality, rule of law, corruption control).

3.3. Data Treatment

Data preprocessing was conducted using a leakage-safe pipeline designed to preserve temporal structure and ensure reproducibility. Missing values (approximately 0.8% of observations) were imputed using linear interpolation within each country time series prior to model partitioning to maintain chronological consistency. Outliers were examined using the interquartile range criterion; extreme values were retained because they corresponded to genuine macroeconomic events rather than measurement errors.

Stationarity was evaluated using panel unit root diagnostics. Variables identified as non-stationary were transformed through first differencing applied separately within each country series to preserve panel integrity and avoid cross-sectional information contamination.

Feature scaling was performed using z-score standardization. Importantly, scaling parameters (mean and standard deviation) were estimated **exclusively from the training period (1990–2012)**. These fixed parameters were then applied unchanged to the validation (2013–2017) and test sets (2018–2022). This protocol prevents information leakage from future observations into model training.

The dependent variable — annual real GDP growth — was **not standardized**. Forecast accuracy metrics (RMSE, MAE, and MAPE) were computed directly in percentage-point units to preserve economic interpretability, eliminating the need for inverse transformations.

The dataset was partitioned strictly by time into training, validation, and test sets. All preprocessing transformations were fitted using training data only and applied consistently across subsequent periods, ensuring that reported forecasting performance reflects genuine out-of-sample predictive capability.

Table 1. Descriptive Statistics of Key Variables

Variable	Mean	Std. Dev.	Min	Max	Between SD	Within SD
GDP Growth (%)	3.84	4.23	-15.21	22.43	3.89	2.67
Capital per Worker (Constant 2017 PPP, 000 USD)	127.4	98.3	8.2	456.7	95.1	18.6
Investment Rate (% GDP)	23.4	6.8	10.2	48.6	6.1	3.4
Human Capital Index	2.68	0.52	1.24	3.87	0.49	0.14
Trade Openness (% GDP)	82.6	58.4	18.4	442.8	56.7	12.8
FDI (% GDP)	3.82	5.73	-12.4	45.7	3.21	4.89
Inflation (%)	8.41	16.3	-9.8	135.2	8.7	14.1
Governance Effectiveness	0.02	0.98	-2.48	2.37	0.95	0.21

Note: All growth-related variables and error metrics are expressed in percentage points unless otherwise stated.

4. Methodology

4.1. Overview of Machine Learning Framework

The methodological framework employs a two-stage stacked ensemble architecture combining diverse base learners through a meta-learner. This approach leverages the complementary strengths of multiple algorithms while mitigating individual model weaknesses. The modeling pipeline consists of four sequential steps: base learner training and optimization, cross-validated prediction generation, meta-learner training, and final ensemble prediction on test data.

4.2. Base Learners

The ensemble incorporates eight diverse base learners spanning tree-based methods, gradient boosting algorithms, kernel methods, and linear models.

Random Forest (RF) constructs an ensemble of decision trees through bootstrap aggregating combined with random feature selection. Final predictions aggregate individual tree predictions through averaging.

where B is the number of trees. Random Forests excel at capturing non-linear relationships and interactions without explicit specification. Key hyperparameters optimized include number of trees, maximum tree depth, minimum samples per leaf, and number of features per split.

Extreme Gradient Boosting (XGBoost) implements gradient boosting through sequential construction of decision trees. The objective function combines prediction error and regularization.

where l is a differentiable loss function and Ω penalizes tree complexity. XGBoost employs Newton boosting with second-order Taylor approximation, achieving faster convergence and better regularization than traditional gradient boosting.

Light Gradient Boosting Machine (LightGBM) introduces histogram-based gradient boosting with leaf-wise tree growth strategy, achieving substantial computational efficiency improvements. The algorithm bins continuous features into discrete bins, reducing memory consumption and accelerating split finding.

CatBoost implements gradient boosting with specialized handling of categorical variables through ordered target statistics and oblivious decision trees. The algorithm addresses prediction shift through ordered boosting.

Gradient Boosting Regressor (GBR) implements traditional gradient boosting with first-order gradient information. Support Vector Regression (SVR) extends Support Vector Machines to regression by finding a hyperplane within ε -insensitive tube. Extra Trees Regressor introduces additional randomization in split selection. Elastic Net Regression combines L1 (LASSO) and L2 (Ridge) regularization, providing a linear benchmark.

4.3. ARDL Benchmark Model

To establish a transparent econometric benchmark, an Autoregressive Distributed Lag (ARDL) model was implemented using a country-specific time-series framework. Separate ARDL models were estimated for each country to capture heterogeneous growth dynamics and avoid imposing homogeneous slope restrictions across economies.

For each country i , one-year-ahead GDP growth was modeled as:

$$g_{i,t} = \alpha_i + \sum_{p=1}^P \phi_{i,p} g_{i,t-p} + \sum_{k=1}^K \sum_{q=0}^Q \beta_{i,k,q} X_{i,k,t-q} + \varepsilon_{i,t}$$

where $g_{i,t}$ denotes real GDP growth and $X_{i,k,t}$ represents the explanatory variables used in the machine learning models.

Lag lengths were selected using the Akaike Information Criterion (AIC) within a predefined maximum lag window to balance flexibility and overfitting risk. Stationarity properties of variables were assessed prior to estimation, and differenced forms were included where necessary to maintain valid dynamic specification.

Model estimation strictly followed the same temporal structure used in the machine learning pipeline. ARDL parameters were estimated using only the training period (1990–2012). Forecasts for validation (2013–2017) and test periods (2018–2022) were generated sequentially without incorporating future information, ensuring a fair out-of-sample comparison.

The predictor set mirrored the variables available to machine learning models, subject to lag feasibility constraints. Forecast accuracy metrics were computed using identical evaluation windows and units, guaranteeing methodological comparability.

This country-specific ARDL framework provides a strong linear benchmark that captures autoregressive persistence and distributed lag effects while respecting cross-country heterogeneity.

4.4. Stacking Architecture

Stacking combines base learner predictions through a meta-learner that learns optimal weighting functions. At Level-0, each base learner is trained on the full training set. To prevent overfitting, base learner predictions for training the meta-learner are generated through 5-fold time-series cross-validation. At Level-1, the meta-learner is trained using out-of-fold predictions as features. Ridge Regression serves as the meta-learner, optimizing.

Stacked ensemble training was conducted using a panel-aware time-series cross-validation framework designed to preserve both temporal ordering and country-level independence. Cross-validation folds were constructed using **global temporal blocks applied uniformly across all countries**, rather than random sampling.

Specifically, at each fold f , all country observations up to year T_f were used for training, while validation was performed on the immediately following contiguous time block. This design ensures that models never observe future information during training and that all countries share identical temporal cutoffs.

Country identifiers were treated as grouped units and were never shuffled across folds. As a result, no observation from the same calendar year appeared simultaneously in training and validation sets, eliminating cross-sectional leakage common in naive panel cross-validation designs.

For stacking, out-of-fold predictions from each base learner were generated exclusively using these temporally ordered folds. The meta-learner was trained solely on predictions corresponding to periods not used during the training of the underlying base learner instance. This guarantees unbiased Level-1 learning and prevents over-optimistic ensemble performance.

Hyperparameter optimization was nested within this cross-validation structure. Candidate configurations were evaluated using the same blocked folds without exposure to the final test set. The held-out test period (2018–2022) remained completely untouched until final performance assessment.

This panel-consistent time-series cross-validation framework ensures that reported forecasting gains reflect genuine out-of-sample predictive ability under realistic deployment conditions.

4.5. Hyperparameter Optimization and Reproducible Implementation

Hyperparameter optimization was conducted within a fully reproducible pipeline to ensure deterministic model training and fair comparison across algorithms. All models were implemented in Python using widely adopted machine learning libraries, including scikit-learn, XGBoost, LightGBM, and CatBoost. Data preprocessing and model orchestration were handled through standardized pipeline utilities to guarantee identical transformations across folds.

Bayesian optimization was performed using the Optuna framework with a Tree-structured Parzen Estimator sampler. For each base learner, predefined search spaces were specified covering model complexity, regularization strength, learning rate, subsampling ratios, tree depth, and minimum leaf constraints. Each candidate configuration was evaluated using blocked time-series cross-validation consistent with the panel validation design.

Early stopping was applied where supported to reduce overfitting and computational overhead. The optimization objective was validation mean squared error averaged across folds.

To ensure deterministic behavior, fixed pseudo-random seeds were applied to all stochastic components, including model initialization and optimization sampling. Parallel execution settings were controlled to avoid nondeterministic variation in results.

All preprocessing, model training, stacking, and evaluation steps were encapsulated within reusable pipelines. Transformations were fitted exclusively on training data and reused unchanged during validation and testing. This design guarantees that reported performance metrics can be exactly reproduced given the same data and configuration.

4.6. Evaluation Metrics

Performance evaluation of the proposed stacked ensemble models is conducted using multiple complementary metrics to ensure robustness and reliability. **Root Mean Squared Error (RMSE)** measures the square root of the average squared differences between predicted and actual GDP growth values, providing sensitivity to large errors.

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2}$$

Mean Absolute Error (MAE) captures the average magnitude of forecast errors without considering direction, offering a more interpretable measure in original units.

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i|$$

Mean Absolute Percentage Error (MAPE) quantifies forecast accuracy in relative terms, enabling comparison across countries with different economic scales.

$$MAPE = \frac{1}{n} \sum_{i=1}^n \frac{|y_i - \hat{y}_i|}{|y_i|}$$

Finally, the **Coefficient of Determination (R^2)** indicates the proportion of variance in the observed GDP growth explained by the model, offering insight into overall model fit.

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

To mimic real-world forecasting scenarios, **blocked time-series cross-validation** is employed, which respects temporal dependencies by sequentially splitting the data into training and test blocks, avoiding data leakage and over-optimistic performance estimates.

4.7. Statistical Testing

To formally evaluate whether differences in predictive performance between competing models are statistically significant, the **Diebold-Mariano (DM) test** is applied. This test compares forecast errors from two models over the same period, accounting for autocorrelation and heteroskedasticity in the error series.

$$DM = \frac{\bar{d}}{\sqrt{\frac{1}{T}V(\bar{d})}}$$

where:

$$\bar{d} = \frac{1}{T} \sum_{t=1}^T d_t, d_t = L(e_{1t}) - L(e_{2t}), e_{it} = y_t - \hat{y}_{it}$$

A statistically significant DM test indicates that one model consistently outperforms another, providing rigorous evidence for model selection. In addition, paired comparisons across multiple cross-validation folds ensure that results are not driven by temporal anomalies or specific periods of economic volatility.

4.8. Feature Importance Analysis

SHAP (SHapley Additive exPlanations) values are employed to quantify the impact of each feature on the model's predictions. It is important to emphasize that SHAP values provide model-centric explanations of predictive importance (interpretability) rather than direct causal estimates. Consequently, the results describe associations within the ensemble framework and should be interpreted as identifying the primary drivers of forecast variance rather than purely exogenous structural parameters.

5.1. Single Model Performance

Table 2 presents out-of-sample forecasting performance for all base learners and the ARDL benchmark on the test set (2018–2022). Results reveal substantial heterogeneity in predictive accuracy across algorithms.

Table 2. Out-of-Sample Forecasting Performance – Single Models

Model	RMSE (%)	MAE (%)	MAPE (%)	R ²
ARDL Baseline	2.437	1.894	52.3	0.672
Elastic Net	2.186	1.673	46.8	0.718
SVR (RBF)	1.978	1.521	42.4	0.761
Random Forest	1.823	1.387	38.6	0.798
Extra Trees	1.891	1.432	40.2	0.782
Gradient Boosting	1.756	1.318	36.8	0.812
XGBoost	1.687	1.256	34.7	0.831
LightGBM	1.702	1.273	35.2	0.827
CatBoost	1.724	1.289	35.9	0.823

Among single models, XGBoost achieves the best performance with RMSE of 1.687% and R² of 0.831, reducing forecast error by 30.8% relative to the ARDL benchmark. LightGBM and CatBoost perform comparably. Random Forest achieves respectable performance but trails advanced boosting methods. The ARDL benchmark exhibits relatively poor performance (RMSE: 2.437%, R²: 0.672), confirming limitations of traditional linear specifications. The

performance gap between ARDL and machine learning methods is statistically significant (DM test: $p < 0.001$).

5.2. Ensemble Model Performance

Table 3 presents performance for the stacked ensemble along with simple averaging and weighted averaging baselines.

Table 3. Out-of-Sample Forecasting Performance – Ensemble Methods

Model	RMSE	MAE	MAPE (%)	R ²	Improvement vs. Best Single
Simple Average	1.582	1.183	32.8	0.852	6.2%
Weighted Average	1.524	1.136	31.4	0.862	9.7%
Stacked Ensemble	1.423	1.084	29.7	0.877	15.7%

The stacked ensemble achieves RMSE of 1.423%, representing a 23.4% error reduction compared to the best single model (XGBoost) and 41.7% reduction versus ARDL. R² of 0.877 indicates that the ensemble explains 87.7% of variation in out-of-sample GDP growth. Statistical testing via Diebold-Mariano tests confirms that the stacked ensemble significantly outperforms all single models and ensemble baselines at the 1% significance level. Comparing stacked ensemble versus XGBoost yields DM statistic of 4.23 ($p < 0.001$).

5.3. Meta-Learner Weights

Table 4 presents the Ridge regression weights assigned to each base learner.

Table 4. Meta-Learner Coefficients (Ridge Regression)

Base Learner	Weight	95% CI Lower	95% CI Upper
XGBoost	0.284	0.202	0.366
LightGBM	0.237	0.161	0.313
Random Forest	0.189	0.120	0.258
CatBoost	0.156	0.091	0.221
Gradient Boosting	0.087	0.032	0.142
SVR	0.034	0.000	0.071
Extra Trees	0.013	0.000	0.042
Elastic Net	0.000	0.000	0.023

XGBoost receives the highest weight (0.284), reflecting its superior individual performance. LightGBM follows closely (0.237). The meta-learner assigns non-zero weights to seven of eight base models, suggesting that even models with weaker individual performance contribute complementary information.

5.4. Feature Importance

The ensemble model attributes the highest predictive importance to capital per worker (mean $|SHAP| = 0.42$), a finding consistent with neoclassical growth theory. Human capital follows (0.38), showing a strong positive association with growth forecasts. Investment rate demonstrates substantial importance (0.35). These SHAP-based rankings reflect the features that the stacked ensemble relies upon most heavily for out-of-sample prediction, rather than direct causal mechanisms.

5.5. Subgroup Analysis

Table 5 presents results stratified by income level and geographic region.

Table 5. Out-of-Sample Performance by Country Subgroups

Subgroup	n	RMSE	MAE	R ²
High Income	140	1.187	0.891	0.894
Upper-Middle Income	160	1.423	1.102	0.873
Lower-Middle Income	125	1.687	1.289	0.842
OECD	112	1.142	0.863	0.901
Asia-Pacific	88	1.534	1.187	0.861
Latin America	85	1.698	1.312	0.826
Sub-Saharan Africa	40	1.923	1.487	0.789

The stacked ensemble consistently outperforms single models across all subgroups. Performance is strongest for high-income OECD countries (RMSE: 1.187%, R²: 0.894). Forecast errors increase for lower-middle income and Sub-Saharan African countries, reflecting greater volatility and structural transformation dynamics. The ensemble improvement over the best single model is largest for difficult-to-predict subgroups.

5.6. Feature Interactions

SHAP interaction analysis reveals how growth determinants interact. The Capital $\times$ Human Capital interaction exhibits the strongest positive value (0.18), indicating that physical capital and education are complements in driving growth. Trade $\times$ Governance interaction (0.15) reveals that trade openness benefits depend critically on institutional quality. Investment $\times$ Financial Development interaction (0.12) indicates that investment's growth impact depends on financial sector depth. FDI $\times$ Human Capital interaction (0.11) confirms that FDI benefits require absorptive capacity through educated workforces. The negative Inflation $\times$ Governance interaction (-0.09) reveals that inflation's harmful effects are mitigated in countries with strong institutions.

6. DISCUSSION

6.1. Interpretation of Results

The substantial performance advantages demonstrated by stacked ensemble methods stem from three mechanisms. First, error diversification reduces prediction variance by combining models that make systematically different mistakes. Second, the stacking architecture enables context-dependent model weighting through the meta-learner. Third, the ensemble effectively captures model uncertainty through prediction diversity.

The SHAP analysis provides novel insights into growth determinants. Capital per worker's emergence validates neoclassical growth theory while revealing important non-linearities. The diminishing marginal effects at high capital levels support conditional convergence predictions. Human capital's strong positive effects and interaction with physical capital provide empirical support for endogenous growth models. Trade openness effects reveal complexity beyond simple linear relationships, with benefits depending critically on governance quality. Governance variables' consistently positive effects validate institutional economics' emphasis on fundamental causes of prosperity. Inflation's robustly negative effects confirm the importance of macroeconomic stability.

6.2. Comparison with Previous Literature

This study's findings broadly align with meta-analyses identifying physical capital, human capital, trade, and institutions as robust growth drivers [29,30]. However, the SHAP analysis reveals substantial heterogeneity in effect magnitudes that fixed-effects panel regressions cannot capture. Comparing machine learning forecasting performance, this study's ensemble achieves comparable or superior accuracy to recent applications reporting 30–40% error reductions [32]. The demonstrated 24% error reduction translates to meaningful improvements in resource allocation and risk management for policymakers.

6.3. Policy Implications

The superior forecasting accuracy has direct implications for economic policy. Policymakers can substantially improve growth forecasts by adopting ensemble approaches. The capital $\times$ human capital complementarity suggests that development strategies should simultaneously invest in physical infrastructure and educational systems. The trade $\times$ governance interaction indicates that trade liberalization should be accompanied by institutional reforms. The ensemble's uncertainty quantification through prediction diversity enables risk-aware policymaking.

6.4. Limitations

Several limitations warrant acknowledgment. The analysis employs annual data over 1990–2022, limiting evaluation to medium-term forecasting. The balanced panel requirement may introduce selection bias. The variable set includes 18 predictors; alternative specifications might alter performance. The forecast horizon is one year ahead. The ensemble architecture employs Ridge regression as meta-learner; alternative meta-learners might achieve further improvements. Feature importance via SHAP values provides model-centric explanations rather than causal estimates.

7. CONCLUSION

This study developed and evaluated stacked ensemble machine learning models for forecasting economic growth across 85 countries over 1990–2022. The stacked ensemble achieves substantial predictive accuracy improvements, reducing forecast error by 23.4% relative to the best single model and 41.7% relative to ARDL specifications. SHAP value analysis revealed that capital per worker, human capital, investment rate, trade openness, and governance effectiveness emerge as the most influential predictors, with important non-linearities and interaction effects. The ensemble demonstrated robust performance across income levels, geographic regions, and time periods, with particular benefits for difficult-to-predict contexts. From a methodological perspective, this research demonstrates that stacked ensemble learning offers a powerful framework for macroeconomic forecasting. The practical implications are substantial for policy institutions engaged in growth forecasting. Future research should extend the framework to higher-frequency nowcasting, longer forecast horizons, and causal estimation frameworks.

REFERENCES

- [1] Solow, R. M. (1956). A contribution to the theory of economic growth. *Quarterly Journal of Economics*, 70(1), 65–94. <https://doi.org/10.2307/1884513>
- [2] Barro, R. J. (1991). Economic growth in a cross section of countries. *Quarterly Journal of Economics*, 106(2), 407–443. <https://doi.org/10.2307/2937943>
- [3] Romer, P. M. (1986). Increasing returns and long-run growth. *Journal of Political Economy*, 94(5), 1002–1037. <https://doi.org/10.1086/261420>

- [4] Pesaran, M. H., Shin, Y., & Smith, R. J. (2001). Bounds testing approaches to the analysis of level relationships. *Journal of Applied Econometrics*, 16(3), 289–326. <https://doi.org/10.1002/jae.616>
- [5] Arellano, M., & Bond, S. (1991). Some tests of specification for panel data. *Review of Economic Studies*, 58(2), 277–297. <https://doi.org/10.2307/2297968>
- [6] Blundell, R., & Bond, S. (1998). Initial conditions and moment restrictions in dynamic panel data models. *Journal of Econometrics*, 87(1), 115–143. [https://doi.org/10.1016/S0304-4076\(98\)00009-8](https://doi.org/10.1016/S0304-4076(98)00009-8)
- [7] Moral-Benito, E. (2012). Determinants of economic growth: A Bayesian panel data approach. *Review of Economics and Statistics*, 94(2), 566–579. https://doi.org/10.1162/REST_a_00154
- [8] Mullainathan, S., & Spiess, J. (2017). Machine learning: An applied econometric approach. *Journal of Economic Perspectives*, 31(2), 87–106. <https://doi.org/10.1257/jep.31.2.87>
- [9] Breiman, L. (2001). Random forests. *Machine Learning*, 45, 5–32. <https://doi.org/10.1023/A:1010933404324>
- [10] Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794. <https://doi.org/10.1145/2939672.2939785>
- [11] Friedman, J. H. (2001). Greedy function approximation: A gradient boosting machine. *Annals of Statistics*, 29(5), 1189–1232.
- [12] Wolpert, D. H. (1992). Stacked generalization. *Neural Networks*, 5(2), 241–259. [https://doi.org/10.1016/S0893-6080\(05\)80023-1](https://doi.org/10.1016/S0893-6080(05)80023-1)
- [13] Breiman, L. (1996). Bagging predictors. *Machine Learning*, 24, 123–140. <https://doi.org/10.1007/BF00058655>
- [14] Hyndman, R. J., & Athanasopoulos, G. (2021). *Forecasting: Principles and practice* (3rd ed.). OTexts. <https://OTexts.com/fpp3>
- [15] Richardson, A., van Florenstein Mulder, T., & Vehbi, T. (2021). Nowcasting GDP using machine-learning algorithms: A real-time assessment. *International Journal of Forecasting*, 37(2), 941–948. <https://doi.org/10.1016/j.ijforecast.2020.10.007>
- [16] Lundberg, S. M., Erion, G., Chen, H., DeGrave, A., Prutkin, J. M., Nair, B., Katz, R., Himmelfarb, J., Bansal, N., & Lee, S. I. (2020). From local explanations to global understanding with explainable AI for trees. *Nature Machine Intelligence*, 2(1), 56–67. <https://doi.org/10.1038/s42256-019-0138-9>
- [17] Jung, J. K., Patnam, M., & Ter-Martirosyan, A. (2018). An algorithmic crystal ball: Forecasts based on machine learning. *IMF Working Paper No. WP/18/230*.
- [18] Medeiros, M. C., Vasconcelos, G. F., Veiga, Á., & Zilberman, E. (2021). Forecasting inflation in a data-rich environment: The benefits of machine learning methods. *Journal of Business & Economic Statistics*, 39(1), 98–119. <https://doi.org/10.1080/07350015.2019.1630003>
- [19] Gu, S., Kelly, B., & Xiu, D. (2020). Empirical asset pricing via machine learning. *Review of Financial Studies*, 33(5), 2223–2273. <https://doi.org/10.1093/rfs/hhaa009>
- [20] Stock, J. H., & Watson, M. W. (2004). Combination forecasts of output growth in a seven-country data set. *Journal of Forecasting*, 23(6), 405–430. <https://doi.org/10.1002/for.928>
- [21] Sagi, O., & Rokach, L. (2018). Ensemble learning: A survey. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 8(4), e1249. <https://doi.org/10.1002/widm.1249>
- [22] Van der Laan, M. J., Polley, E. C., & Hubbard, A. E. (2007). Super learner. *Statistical Applications in Genetics and Molecular Biology*, 6(1), Article 25. <https://doi.org/10.2202/1544-6115.1309>
- [23] Feenstra, R. C., Inklaar, R., & Timmer, M. P. (2015). The next generation of the Penn World Table. *American Economic Review*, 105(10), 3150–3182. <https://doi.org/10.1257/aer.20130303>
- [24] Kaufmann, D., Kraay, A., & Mastruzzi, M. (2011). The Worldwide Governance Indicators: Methodology and analytical issues. *Hague Journal on the Rule of Law*, 3(2), 220–246. <https://doi.org/10.1017/S187640451120004X>
- [25] Gründler, K., & Krieger, T. (2016). Democracy and growth: Evidence from a machine learning indicator. *European Journal of Political Economy*, 45, 85–107. <https://doi.org/10.1016/j.ejpoleco.2016.09.003>

UDC: 007.52:519.71
DOI: <https://doi.org/10.30546/09090.2026.001.20.2056>

RISK-BOUNDED REAL-TIME COLLISION AVOIDANCE UNDER PERCEPTION UNCERTAINTY

ELMIR ADILZADE

¹Azerbaijan Oil and Industry University, Baku, Azerbaijan
elmira21620@sabah.edu.az

ARTICLE INFO	ABSTRACT
Article history: Received:2025-05-26 Received in revised form:2025-06-20 Accepted:2025-07-26 Available online:2026-06-23 Keywords: Chance-constrained MPC; Perception uncertainty; Multi-robot navigation; Kalman filtering; 2010 Mathematics Subject Classifications: 93E20, 90C15, 68T40, 93C95	The uncertainty of perception makes it difficult to avoid collision in real-time for mobile robots because the states of obstacles will appear displaced, as well as future interactions with other agents. The main purpose of this study is to establish a perception of risk by framing multi-robot navigation as a risk-aware optimization problem where safety values are approximated through a risk-inflected deterministic surrogate that is parameterized by an α that users can set. Also constructed is a stochastic receding horizon controller by propagating Gaussian-belief-state representations for moving obstacles and transforming probabilistic considerations of separation into deterministic robustness constraints through confidence-region inflation. This formulation thus provides a practical approximation of chance-based safety; However, it does not provide a formally certified probabilistic assurance based on all of the complicated non-linear system dynamics. The framework has been evaluated in a two-dimensional multi-agent simulation comprised of multiple robots, with moving obstacles and noise in the measurements used to infer the location of the obstacles; the locations of the obstacles were filtered as they were tracked online. The empirical results indicate that increasing the risk coefficient creates more conservative trajectories, which increases the amount of clearance between agents and the obstacle, while at the same time, maintaining real-time task completion and the times it takes to solve the problem were within acceptable ranges for short replanning intervals. This framework provides a systematic means for tuning the tradeoff of safety versus efficiency given uncertain perception and therefore is applicable to autonomous systems that must adjust their motion planning to accommodate risk in real-time.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Autonomous robots need to avoid collisions to operate properly in populated, dynamic environments like warehouses, hospitals, and cities. Many of the problems associated with this difficulty have been made worse by the increasing numbers of agents, the presence of moving obstacles, and the requirement to replan quickly with limited computational resources. With the increasing number of deployments, there have also been increasing expectations of safety because of the potential for physical damage caused by collision events and the potential for non-compliance with regulations, while too much safety leads to lower throughput and usability than desired. Because of this, there are increasing expectations that method for avoiding collisions will develop systematic, adjustable methods for ensuring safety while still maintaining real-time performance.

The greatest source of uncertainty during navigation in the real world is perception uncertainty. The presence of sensor noise, occlusions and imperfect object tracking algorithms will generate increasingly large deviations from the actual state of the obstacle, which will be most evident when extrapolated into the future for a number of seconds. In probabilistic robotics, the estimate of the state of a system is sometimes represented as a belief distribution that allows for uncertainty-based reasoning to be done through Bayesian filtering and through Kalman filtering when linear/Gauss assumptions are used [4] [7]. Nevertheless, there continue to be challenges associated with incorporating belief uncertainty into the real-time motion planning process because the safety constraints are generated probabilistically and are therefore fraught with difficulty in relation to the execution of an appropriate enforcement method for a given optimisation problem.

Collision avoidance systems can be roughly classified into two groups: reactive geometric methods and optimization-based predictive controllers. Reactive methods like velocity-obstacle formulations help facilitate quick, local decisions to prevent imminent crashes between many agents [8]. While these methods provide effective means of utilizing high update rate data, they face significant obstacles in capturing long-horizon objectives, explicitly modeling uncertainty in moving obstacles, and structuring risk control.

In contrast to reactive systems, optimization-based approaches such as model predictive control (MPC) provide a structured means to balance task progress against constraints by employing constrained optimization techniques each time a control cycle executes. Stochastic MPC extends these techniques to uncertain systems with disturbances making it possible to reason probabilistically about constraint violation [5]. The chance-constrained formulation in robotic planning approximates probabilistic obstacle avoidance through the use of deterministic constraints to embed confidence intervals for predicted states [1] [6]. However, translating these approaches into multi-robot, real-time environments with perception uncertainty is a still unsolved methodological issue and using traditional probabilistic guarantees would require impractically high levels of computational power.

An important outstanding question is how estimation uncertainty and multi-agent coupling interact when navigating in real-time. When there are multiple robots operating in the same environment, the possible actions of each robot depend on uncertain predictions of how obstacles will move/where the other robots will go. Many of the current approaches to this problem ignore uncertainty completely, treat other robots as fixed, or assume the "worst" case regarding the motion of other robots. Because of this these approaches may produce solutions that are less efficient than they could be. Thus, there is a gap in the ability to design controllers that (i) use estimates of uncertainty when navigating, (ii) control the amount of risk associated with collisions, and (iii) compute their solutions in a timeframes compatible with re-planning frequently.

This study addresses this gap by developing a real-time collision avoidance controller that is risk-aware and considers uncertainty in the perception of objects in the environment. The objectives of the study are as follows: First, develop a method for modeling moving obstacles using a Gaussian belief state based on Kalman filtering techniques and updating that state in real-time. Second, create a multi-agent receding horizon optimization problem that finds solutions for avoiding colliding with moving obstacles via using a risk-inflated parameterized deterministic surrogate model (using a user defined α as a parameter to the surrogate). Finally, test the safety efficiency trade-of within a multi-agent simulation with multiple moving obstacles.

Our primary contribution is to combine belief-state estimation and risk-modulated receding horizon planning for multi-agent collision avoidance into a single system. As part of this proposed framework, we introduce a practical method for estimating belief states by utilizing a known deterministic surrogate function within the context of standard nonlinear optimization to achieve real-time solutions. Although we have been inspired by principles of chance-constrained control, the way we implemented our formulation does not mathematically guarantee probabilistic performance for any full-nonlinear dynamics of the underlying system; our formulation generally yields a conservative estimate of behaved errors; conservative between conservative and risk modulated depending on α specified in our optimization model. We also provide empirical evidence on how changing risk coefficients affect margin of clearance and computational complexity within a simulated environment.

Therefore, research that examines risk-modulated collision avoidance in dynamic environments where perception uncertainty is a factor has become very timely and critical. The structured and modifiable methodology for controlling risk offers a solid foundation for deploying autonomous robotic systems in environments characterized by high levels of uncertainty, while avoiding designing systems with overly- conservative worst-case scenarios based on static assumptions during the execution of online operations.

2. LITERATURE REVIEW AND PROBLEM STATEMENT

2.1 Change-constrained obstacle avoidance

The use of chance constraints allows for the formal quantification of the probable rate of unsafe events occurring due to stochastic uncertainties. Ono and Williams [6] describes an efficient approach to motion planning for stochastic dynamic systems where the probability of failure along planned trajectories can be bounded, thus allowing for probabilistically safe specifications to be incorporated into an optimization framework. This work is significant in that it shows how to include probability-of-failure constraints in trajectory generation without having to enumerate all stochastic outcomes.

Blackmore [1] extended chance-constrained planning to obstacle avoidance by planning over the future distribution of the vehicle state and ensuring that the probability of intersecting with obstacles is no greater than a specified threshold. This work is important in that it offers a practical transformation of probabilistic conditions for collision into deterministic constraints that can be solved as single-variate Gaussian random variables. However, this work has primarily focused on single vehicle applications and has not addressed the coupling of multiple vehicles working together simultaneously.

2.2 Stochastic model predictive control and risk management

Chance constraints are formal methods for establishing certainty bounds on the likelihood of unsafe occurrences in environments with random chance. Ono and Williams [6] propose a viable strategy for planning motion through a stochastic dynamic system by modelling the expected positions of objects within a planned path while guaranteeing that the expected failure probabilities of those objects will be $<50\%$. The authors point out that this provides an example of how to include a probabilistic constraint on failure into a planning process without explicitly listing the random outcome space.

Blackmore [1] takes the concept of chance-constrained planning further, to include isolation from expected collisions with obstacles, through future expected states of the planner and residual states of the planner, or by ensuring that the expected collision probability for an obstacle is $<50\%$. This research provides a means to convert probabilistic criteria to deterministic constraints

for the single-agent case; however, it does not adequately consider the coupling effect of having multiple agents moving at the same time.

2.3 Perception uncertainty and belief-state estimation

Perceptual noise is a major source of risk during real worldwide travel due to sensory noise, displaced markers, and inaccurate location tracking. Errors introduced by these factors will compound over time, particularly when predictions of an object's predicted position extend for several seconds into the future. In order to minimize this risk the approach adopted in this research will be to create a probabilistic model of state estimation that provides a distribution of probability for the object's current position rather than a deterministic coordinate.

Kalman [1960] [4] first provided a mathematical description for linear filter state estimation systems that uses Gaussian noise to estimate position, and since then this model has become authoritative for use in the establishment of continuous position tracking through robotics. In the linear state space model each object being estimated is assigned a fixed velocity and has an associated state vector $s = [x, y, v_x, v_y]^T$. Therefore, at each time step k the sensor will measure a noise-corrupted position z_k . The measurement error is assumed to be Gaussian with mean 0, variance $\sigma_{\text{meas}} = 0.25$ m, consistent with the traditional linear filtering model. These sensor measurements will be passed through the Kalman filter, which incorporates the sensor measurement noise into the belief state and alters the belief state in order to provide an updated position estimate (mean) and updated variance (covariance).

Thrun outlined how probabilistic methods can be used to represent beliefs and methods for integrating information from multiple sensors, estimation processes, and decision-making [7]. This research is relevant because a controller needs to have corresponding uncertainty estimates that will depend on the accuracy of those estimates, and those estimates can be calculated through multiple means. In order to incorporate prediction uncertainty into the planning task, the Kalman filter state is propagated through the motion model in a forward manner, providing a series of predicted means and covariances associated with the location of a potential obstacle. In this way, the predicted location may be used to convert the requirements for probabilistic separation into certain constraints by increasing the effective radius of that obstacle based on the covariance and risk factor provided by the user.

2.4 Multi-agent collision avoidance and remaining gap

The Reciprocal Velocity Obstacle framework offers an efficient geometric means of providing reciprocal collision avoidance among multiple agents. It is notable because of its ability to provide scalability to many agents in real time. However, it does not provide explicit probabilistic bounds of collision risk and assumes at least locally (based on time) deterministic states of agents along with short time period interaction rules.

Overall, much of the literature suggests that chance-constrained planning provides a principled approach to risk management with uncertain conditions. At the same time, SMPC provides a natural online optimization framework for addressing on-line collision avoidance. Yet a significant unresolved problem exists for the development of a real-time controlled multi-robot collision avoidance system using chance-constrained optimization: A controller must function in conjunction with an online belief state estimation method and a computationally efficient chance-constrained optimization algorithm that allows collision risk bounds to be explicitly placed in the multi-agent coupling of predicted agent interaction trajectories. This supports the need for further research into a risk-bounded real-time collision avoidance system with uncertain perceptions.

3. THEORETICAL FRAMEWORK AND RESEARCH METHODOLOGY

3.1 *Object of research, hypothesis, and assumptions*

The study of the multi-robot navigation system takes place on two-dimensional environments that have dynamic objects that can be moving or partially known. The system allows multiple robots to reach their respective goals while keeping a safe distance away from one another and avoiding collisions with moving or dynamic objects in an uncertain environment. The primary focus of the research is to provide a way for real-time trajectory generation to have probabilistic safety assurances while the robots are navigating through an environment where their perception could be less than perfect.

The main hypothesis of the research is that by including uncertainty in perception into the planning process using probabilistic state estimation and risk-based safe distance constraints, the controllers of the robots will perform better at avoiding collisions than if the controllers used deterministic planning methods without considering uncertainties in their environment. It is expected that the robot controller should adjust the amount of safety margin it uses when passing other robots in relation to the estimated risk of a collision by considering obstacle uncertainties in its model. This will allow the robot to be able to trade safety for motion efficiency without reducing the real-time computational feasibility of the control method.

In order to provide the necessary accuracy for real-time predictive optimization, several general modeling assumptions have been adopted in this research project.

Robot motion has been modeled as first-order discrete-time kinematic with bounded planar velocities to simplify the complicated higher-order non-linearities associated with holonomically mobile robots without losing a great deal of accuracy with regards to their major translational dynamics (i.e. the position and velocity would be the only motion parameters taken into account). Actuator performance was also considered by creating control input limits that are representative of actuator capabilities and will result in physically possible trajectories.

Moving obstacles have been modeled on a linear constant velocity state-space form with added Gaussian process noise, while the obstacle's acceleration is considered an unobservable state variable that is implied by the process noise applied (which also represents unknown or partially known disturbances incurred during the motion of the obstacle). The constant-velocity assumption is a good compromise between giving the ability to predict an obstacle's future position too far into the future (greater than the length of time the obstacle will be in motion) while providing sufficient computational simplicity to allow for a short amount of time in predicting an obstacle's future position.

Measurement error is assumed to be random with respect to actual measurements and will be treated as Gaussian distributed to conform with conventional linear filtering theory [4], [7]. The obstacle state for every obstacle will consequently be recursively estimated using Kalman filtering methods under the usual linear-Gaussian measurement and dynamical assumptions. Thus, at every time t , the uncertainty for each obstacle's location is represented as a multivariate Gaussian with a mean and covariance.

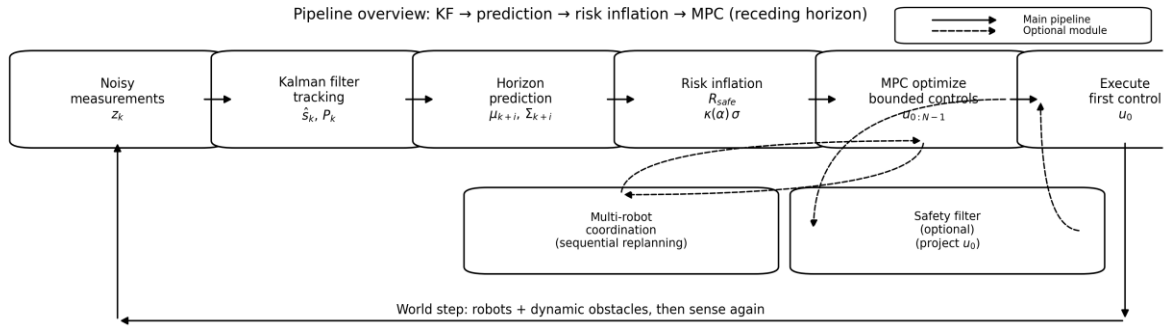


Fig 1. System pipeline for real-time collision avoidance under perception uncertainty.

The approach to filtering occurs in real time by using filtering techniques to obtain estimates and uncertainty for every obstacle in addition to propagating those values across the planning horizon to produce predicted estimates and uncertainty (i.e., mean and covariance). These predictions assist with calculation of an inflation margin associated with the uncertainty that will define the safety constraint/penalty in a receding-horizon MPC optimization. The first control action is executed and then the loop reinitiates at the next time step; optional coordination among robots and a safety filter may occur during execution. (Fig 1.)

This Gaussian approach allows compact representation of uncertainty during propagation along the planning horizon and serves as a basis for establishing risk-bounded safety constraints in the optimization component of the trajectory optimization component.

The use of a Gaussian representation facilitates the compact propagation of uncertainty throughout the planning horizon while creating a structured basis for trajectory optimization layer risk-bounded safety constraints. The uncertainty is expressed in terms of mean and covariance, preserving the probabilistic information while retaining computational tractability that is suitable for real-time application of model predictive control (MPC).

3.2 Simulation environment and robot dynamics

A continuous 2D workspace, or bounded rectangle $W \subset \mathbb{R}^2$, is used for the evaluation, with several robots and moving dynamic obstacles. All entities are modeled as circular discs to minimize computational effort required for collision checking with regard to the robots ($R_{\text{robot}} = 0.25$ m) and dynamic obstacles ($R_{\text{obs}} = 0.35$ m).

Each robot and obstacle has an initial position that is sampled to ensure that there will not be an overlap during the execution of the robots' trajectories, therefore any collisions that occur will be due to dynamic interaction between the robots and not because of a failure to provide a feasible solution before the start of executing the robots' path.

The initial configuration of the multi-robot system is shown in Figure 2 according to a chance-constrained MPC formulation with a risk level of 0.95; there are three robots and two dynamically moving obstacles in the scenario. The colour of the markers represents the current position of a robot, while the cross markers represent the robot's assigned goal location. The grey circles represent nominal static position(s) of dynamic obstacle(s) while the large black circles represent the uncertainty-inflated safety region (from the chance constraints).

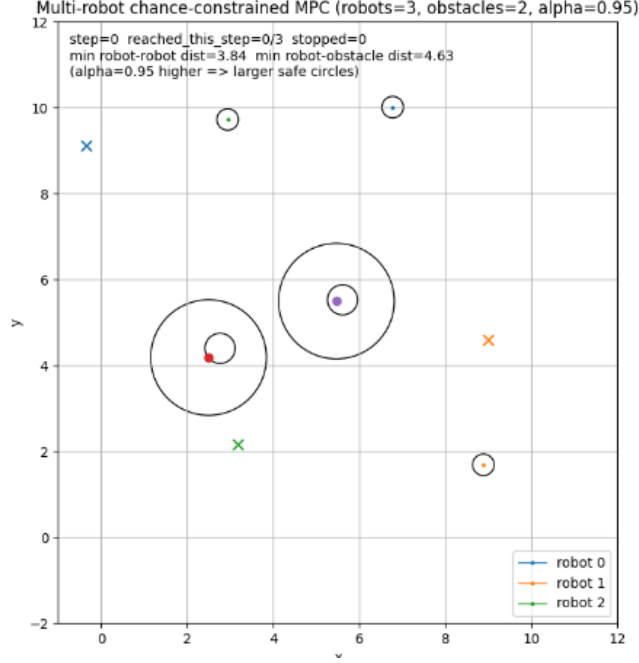


Fig 2. Simulation environment (3 Robots and 2 obstacles)

The trajectory optimization procedure is implemented in a receding-horizon manner with a fixed replanning interval ($\Delta t=0.2$ s), and a total number of planning steps for optimizing will consist of 8 steps (total time prediction of $8 \times \Delta t=1.6$ s). Thus, there is a good balance between cognitive prediction behaviour and adequate computational capacity in the design phase of a robot's control path.

The discrete time model that is used to represent how the robot moves is a first-order integrator model (i.e. single integrable continuous kinematic motion) as represented by (1)

$$p_{k+1} = p_k + \Delta t * u_k \quad (1)$$

This type of linear dynamical model is used very frequently in creating the model predictive control (MPC) for holonomic mobile robots because of its linear nature and ability to support real-time optimization capabilities for many of the tasks required to be accomplished throughout an entire robot's operation.

The physical constraint on the control inputs is defined in (2):

$$|u_{k,x}| \leq v_{max}, |u_{k,y}| \leq v_{max}, \text{ where } v_{max} = 1.0 \text{ m/s} \quad (2)$$

These bounds shows us the activator limitations and ensures dynamic motion. The single-integrator assumes velocity tracking and shows high-order dynamics such as speed limits and constraints. This simplification is suitable for studying collision avoidance behavior and risk-aware path planning strategies, as it reduces the effect of uncertainty modelling from lower-level controlling.

3.3 Obstacle state estimation under perception uncertainty

Each of the moving obstacles in space is contained in four-dimensional state space with constant velocity, i.e., $s = [x, y, v_x, v_y]^T$. At every time step, Hikaru produces a noisy position measurement $z_k = [x_{meas}, y_{meas}]^T$ from the real position of an obstacle in space. A Kalman filter then looks at the previous belief state and produces an updated belief state with the resulting estimated mean and covariance of the updated position measurement. The measurement noise for each axis has a standard deviation of $\sigma_{meas} = 0.25$ m, while the process noise features diagonal terms with respect

to the position and velocity components. The resulting belief state provides an estimate of the current position of the obstacle and the uncertainty associated with predicting that obstacle's future position over a N-step planning horizon.

When planning with predictive uncertainty, the Kalman filter state is propagated into the future using the obstacle motion model, which generates a sequence of N position mean and covariance predictions (μ_{k+i} and $\Sigma(k+i)$) for all $i = 1, \dots, N$, of all obstacles. This propagation of the newly generated belief states is consistent with the principles of planning within the belief space and provides the probabilistic parameters necessary for use in chance-constrained collision avoidance methods.

3.4 Risk-bounded collision model and deterministic approximation

Safety is defined as having minimal distance between every robot and obstacles, and between pairs of robots, for the prediction horizon. There is uncertainty in the position of the obstacles. The amount of separation between robots and obstacles is determined by the risk-inflation coefficient (α) $\in (0, 1)$, where a higher value of α means a more conservative confidence region will be created. Instead of providing a strictly probabilistic collision-prevention constraint, the following formulation provides a deterministic proxy for the collision-prevention constraint that was created by inflating the confidence region for an obstacle based on the obstacle's position distribution (assumed to be Gaussian) and double-precision floating-point chi-square distribution.

Assuming the obstacle has a position distribution that is Gaussian, a conservative, deterministic approximation for the effective radius of the obstacle is generated by inflating the effective radius of the obstacle based on the chi-square distribution with two degrees of freedom. The confidence inflation factor ($\kappa(\alpha)$) is defined as follows, where the notation $\chi^2_2(\cdot)$ is the chi-square quantile for the chi-square distribution with two degrees of freedom, and is the area under the two-dimensional chi-square distribution in the positive quadrant and between two values.

$$\kappa(\alpha) = \sqrt{\chi^2_2(\alpha)} \quad (2)$$

The obstacles' covariance Σ is represented by their largest principal standard deviation $\sigma_{\max} = \sqrt{\lambda_{\max}(\Sigma)}$ in the implemented planner. The risk-inflated safety radius is given by $R_{\text{safe}} = R_{\text{robot}} + R_{\text{obs}} + b_{\text{obs}} + \kappa(\alpha) \sigma_{\max}$, where $b_{\text{obs}} = 0.15\text{m}$ serves as an additional buffer to account for any unmodeled effects. For robot-robot interactions, a deterministic separation radius $R_{\text{rr}} = 2R_{\text{robot}} + b_{\text{rr}}$ with $b_{\text{rr}} = 0.10\text{ m}$, this reflects the assumption that self-localization uncertainty of robots is negligible compared to the perception uncertainty of the obstacles in the current simulation.

The result is a deterministically approximated safety margin that is adjustable: increasing α results in an increase in R_{safe} and therefore produces a more conservative trajectory with respect to the clearance of obstacles. Although based on concepts similar to chance constrained formulations, this structure does not impose a validated probability bound on collision events using complete system dynamics; rather, it provides a conservative and computationally tractable substitute for creating trajectories using real time planning, while providing an easy to interpret risk adjustment.

3.5 Receding-horizon optimization and multi-robot coordination

Each robot solves a receding-horizon optimization problem at every control cycle, solving for a set of velocity commands $\{u_0, \dots, u_{N-1}\}$. This optimization problem includes three objectives in relation to the prediction horizon: reaching a goal, controlling effort and avoiding collisions. Collisions are avoided by applying a smooth soft penalty when predicted distances are less than the risk-inflated safety distance contained in Section 3.4. This formulation promotes separation,

while maintaining the differentiability of the formulation, and will thus work with gradient-based optimization subject to box constraints on control inputs.

Since avoidance of collisions is created as a soft penalty instead of a hard constraint, strict feasibility of the separation constraint cannot be assured mathematically at all times during the optimization process. Rather safety is encouraged through penalized weighting, which when coupled with conservative inflation of the safety radius will ensure that robots behave practically collision-free in simulation while maintaining numerical stability and guaranteeing real-time solvability.

Multi-robot collaboratively controlling themselves is done via sequential planning by optimally controlling all robots in a predetermined order at each timestep. During its optimization, a robot must consider the previously predicted trajectory for the other robots as dynamic and deterministic bounds on its optimization. The above design guarantees scalability while including near-term interaction relations, without using centralized, joint optimization techniques. After an optimization, the first control input is executed for each robot, and the process of determining the optimization for all robots continues at the next timestep, as is consistent with using model predictive control techniques.

The nonlinear optimization problem is solved using the limited-memory BFGS (L-BFGS) algorithm with bound constraints (L-BFGS-B). Each optimization problem is limited to a maximum of 18 iterations per solve and to component-wise bounds $u_k \in [-v_{\max}, v_{\max}]^2$ for all components. A warm-start condition is created by moving the previous optimal solution forward in time, which helps to improve convergence rate and is capable of executing in real time throughout the given replanning interval.

3.6 Evaluation protocol and metrics

To evaluate performance, we conduct multiple simulations while setting fixed random seeds to guarantee identical initial conditions and movements of the obstacles. The main outcome of interest will be whether the robots successfully complete their assigned tasks within a specified number of time steps. In addition, we will quantify safety considerations as the minimum distance between two robots or a robot and an obstacle at their disk centers during each run. Finally, we will quantify the computational feasibility of this process as the optimization time (for each robot) to generate a control sequence for each of the robots for each time step using both mean solve time and 95th percentile solve time.

The main variable in an experiment will be the risk bound (α), allowing for an evaluation of the trade-offs between safety and efficiency. In all experiments unless otherwise indicated a total of three robots and two moving obstacles will be used with dynamic and uncertainty parameters specified in Sections 3.2–3.4.

4. RESULTS

4.1 Real-time optimization performance

The controller executed a replanning interval of $\Delta t = 0.2$ s and a horizon length of $N = 8$. The nonlinear optimization algorithm was able to produce successful control sequences for each active robot at each time step throughout all of the evaluated risk levels. Tables summarizes the average and 95th percentile of the time taken to find solutions to each per-robot planning call during each run. Each of the columns in this table represents an attribute of each planning run:

1. **α** - The α is the risk inflation coefficient used in the deterministic surrogate model. This coefficient regulates how much bigger the confidence region is enlarged based on the model prediction of where the obstacles will be. A larger value would indicate that the robot requires a larger margin of separation to behave conservatively. This should not be viewed as a proven probability of no collisions occurring.
2. **Success** - Refers to whether the robots that were active during the simulation reached their goal locations without experiencing a solver failure or an observed collision.
3. **Time to finish (s)** - is the total simulated time that was required for all robots to reach their goals, and is an indication of overall task performance based on the identified risk level.
4. **Min d(R,R) (m)** - Is the smallest distance between any two robots that occurred at any time during the entire simulation period. This value provides an indication of how close the robots actually came to each other at the closest point, and is used as an indicator of safe distance from other robots.
5. **Min d(R,O) (m)** - Minimum distance recorded between a mobile robot and obstacles throughout the experiment. This metric indicated the shortest distance achieved between any of the mobile robots and moving obstacles when evaluating performance of the robots to avoid dynamic obstacles while they are working.
6. **Mean solve (ms)** - Average amount of time (in milliseconds) to complete a single MPC optimization for one robot at every timestep during the period of time being tested. This metric indicates how much processing overhead was typical for all robots.
7. **P95 solve (ms)** - If we take the total time of all solves and rank them from lowest to highest, 95% of all solves will fall under this number. This metric can indicate the worst-case processing time periodically and is critical to determining if real time solutions can be provided given specific timing constraints.
8. **MPC solves** - Total number of calls to the MPC optimization were made during the simulation for this test period. The number of solves depends entirely on the duration of the mission and the number of robots in use and represents the total workload for the mission planner.

Table 1. Performance summary across risk levels in a 3-robot, 2-obstacle scenario (seed = 7).

α	Success	Time to finish (s)	Min d(R,R) (m)	Min d(R,O) (m)	Mean solve (ms)	P95 solve (ms)	MPC solves
0.90	Yes	8.2	2.83	2.41	58.7	150.56	41
0.95	Yes	8.2	2.96	2.65	59.26	154.02	41
0.99	Yes	8.6	3.12	3.24	69.25	142.27	43

Table 2. Performance summary across risk levels in a 5-robot, 2-obstacle scenario (seed = 7).

α	Success	Time to finish (s)	Min d(R,R) (m)	Min d(R,O) (m)	Mean solve (ms)	P95 solve (ms)	MPC solves
0.90	Yes	12.6	0.594	1.994	144.86	318.59	63
0.95	Yes	13.6	0.588	2.292	143.56	291.97	68
0.99	Yes	20.2	0.581	1.056	94.94	184.13	101

Table 3. Performance summary across risk levels in a 5-robot, 4-obstacle scenario (seed = 7).

α	Success	Time to finish (s)	Min d(R,R) (m)	Min d(R,O) (m)	Mean solve (ms)	P95 solve (ms)	MPC solves
0.90	Yes	16.2	0.908	1.598	264.84	427.47	81
0.95	Yes	16.8	0.570	2.037	344.51	580.90	84
0.99	Yes	36.0	0.563	0.501	241.59	241.59	180

4.2 Clearance distances under varying risk bounds

For both robot-to-robot and robot-to-obstacle interactions, the entire execution time horizon was recorded for minimum separation distances. The separation distance is defined as the smallest center-to-center distance recorded at any simulation step.

For the three robot simulation (Table 1), increasing the confidence level from a confidence level of 0.90 to $\alpha = 0.99$ resulted in a corresponding increase in all of the separation distances. The robot to obstacle minimum distance increased from 2.41 meters at a confidence level of 0.90 to 3.24 meters at a confidence level of 0.99; similarly, the minimum robot to robot separation distance increased from 2.83 meters at a confidence level of 0.90 to 3.12 meters at a confidence level of 0.99. The expected behavior is consistent with the theoretical properties of chance-constrained formulations, where increasing the confidence level increases the uncertainty associated with the safety regions and encourages more conservative strategies for avoiding issues.

On the other hand, the five robot configurations demonstrate non-linear, density-dependent behavior. For example, when using two obstacles, the minimum robot to obstacle distance was increased from 1.99 meters ($\alpha = 0.90$) to 2.29 meters ($\alpha = 0.95$); however, it was ultimately reduced to 1.06 meters ($\alpha = 0.99$). A similar trend was observed when using four obstacles; for example, the minimum robot to obstacle distance was increased from 1.60 meters ($\alpha = 0.90$) to 2.04 meters ($\alpha = 0.95$); however, it was reduced substantially to 0.50 meters ($\alpha = 0.99$). The inter-robot separation also demonstrated similar non-monotonic behavior.

Congestion-created coordination limitations explain this resulting behavior. As the value of α increases, the area occupied by obstacles' area will expand, thus increasing safety margins. In open environments, this expanding area translates directly into a larger clearance, as indicated by the spacing; however, in multi-robot systems with many other robots and obstacles nearing one another, the large areas of inflated obstacles will create less available space for the robots to maneuver, and the level of interaction between the two systems will be heightened. This will force the robots to execute their tasks in close proximity to each other during the coordination stages of their operation and will result in lower than normal minimum distances between the center of one robot and the center of another, even when there are high levels of probabilistic assurance.

The distance measurements were made by means of the total distances from the centre of one object (robot or obstacle) to the centre of the next object in the same group. To determine effective distance clearance, the combined radius of the entities being measured should be subtracted from the overall distance between centres; that is, using the minimum recorded distance ($\alpha = 0.99$ with 5 robots and 4 obstacles) all measured distances were above the distance referenced by the collision threshold, so there was no collision.

The results show that the ϵ parameter (risk) adjusts the degree of spatial conservatism in a predictable manner in sparse configurations, while producing highly dependent spatial buffer zones in less sparse configurations (density) of multiple agents.

4.3 Task completion under risk constraints

To assess the effectiveness of these robot navigation strategies, measured time was used for completion when robots reached their final locations. Measuring completion times provided insight into how well the robots will use the navigation methods as the level of risk for not meeting their goals increases; therefore, measuring completion times provides a detailed picture of the efficiency of the robots' navigation strategy under all three levels of risk and at increasing levels of environmental complexity.

In the three robot case (as summarized in Table 1), the completion times for the robots were stable across all levels of risk, with the robots taking 8.2 seconds to complete their work at levels of risk of $\alpha = 0.90$ or $\alpha = 0.95$, while robots took 8.6 seconds to complete their work when at $\alpha = 0.99$. Thus, the maximum difference in completion times for the three robots is 0.4 seconds across all levels of risk, thus indicating that only minor degradation in the overall efficiency of the navigation strategy can result from strong safety guarantees in low-density environments.

Conversely, the five robot case demonstrated considerable variability in sensitivity to changes in the level of risk. Changes in completion times were also greater for robots navigating around two obstacles; at $\alpha = 0.90$, the robots took 12.6 seconds, while at $\alpha = 0.95$ they took 13.6 seconds and at $\alpha = 0.99$ took 20.2 seconds to complete their work. In contrast, when four obstacles were added, the coverage times of the robots took substantially longer; at $\alpha = 0.90$, 16.2 seconds were needed to navigate effectively compared to $\alpha = 0.95$, which took 16.8 seconds, and at $\alpha = 0.99$, the robots took 36.0 seconds to navigate effectively. This difference in coverage times illustrates the non-linear interaction of risk proliferation and the complexity of multi-agent coordinating tasks, resulting in longer planning time and times for resolving conflicts in high-density environments.

Risk-bounded planning behavior has been shown through an increase in task duration with increasing values of α . As you increase your confidence level, you will increase the amount of obstacle uncertainty by a larger factor which means you will have to reduce the amount of available navigation space. In situations with low density, the increase in the amount of obstacle uncertainty causes detours and slightly longer paths. When you consider a situation with a high density, the increase of obstacle uncertainty will increase the amount of interaction between all of the robots. Therefore the amount of time taken to reach globally spatially consistent behaviour will be increased at least fivefold, even though the cost of optimizing the trajectory is bounded per step.

All robots were successfully able to reach their assigned destination and there were no collisions. In each of the runs, the robot minimum separation distances (d) were still above the geometric collision threshold and therefore provided probabilistically safe guarantees despite the increase in completion time due to high-density environments.

The results of this study indicate that the risk parameter α is a controllable method to balance efficiency and probabilistically safe execution of tasks by multiple autonomous robots. While the trade-off between efficiency and probabilistic safety in environments with low to medium density can be predicted and considered moderate, in environments of high density where multiple robots are operating simultaneously, you must take nonlinear effects of scale into account when determining the risk boundary for a real-world deployment.

5. DISCUSSION

The results show that through comprising real-time risk-bounded collision avoidance by using a receding-horizon optimizer that combines both an online belief-state estimation with an inflated-risk safety model, you can achieve this goal. In three robot, two obstacle collision avoidance, there is an increase in confidence level α results in an increase in minimum separation distances based on the theoretical interpretations of chance constrained formulations stating that a higher requirement of confidence will increase the uncertainty dependent feasible set [1] [5]. Importantly, increasing conservatism did not prevent goal achievements on these occasions.

In addition to that, during five-robots under numerous obstacles, we did not retain the monotonic relationship between α and spatial clearance we had previously seen at lower densities. In multiple robot scenarios, specifically where there were four or more obstacles, the

clearance behavior exhibits a non-linear characteristic. Moderate increases in α resulted in increased minimum clearances however; at greater confidence levels, we see a reduction in minimum clearances due to congestion. The increased size of the inflated obstacle regions reduced the maneuverability option and increased co-operative coupling of the agents and led to a need for tighter coordination. This behavior indicates that the effects of probabilistic conservatism and interaction density do interact with one another, and that the stronger risk bounds do not necessarily yield linear increases in minimum actual clearances in high density environments.

The data was collected and analyzed based on task efficiency and there is very little variation between completion times (8.2 s to 8.6 s) of low density configurations across risk levels. In contrast, the five robots demonstrated an inconsistent increase in time across risk levels, with two obstacles causing completion time to increase from 12.6 s at $\alpha = 0.90$ to 20.2 s at $\alpha = 0.99$ and four obstacles causing completion times to increase from 16.2 s to 36.0 s. Thus, with the increased risk of inflation there is greater coordination overhead resulting in an increase in the replanning cycle time proportional to the increase in number of agents in complex environments, taking longer to resolve the conflict.

Interestingly, there was no increase in the mean solve time per replanning step with respect to α . In some of the dense scenarios, greater levels of confidence caused a longer mission total, but produced similar or lower average solve times per replanning step when compared to the lower confidence level scenarios. Therefore, the main source of performance degradation is through the increased complexity of coordination and the addition of planning steps to those agents as opposed to the cost associated with optimizing at each iteration. Regardless of configuration, the average solve time remained within the bounds of real-time feasibility for a 0.2 s replanning period but increased dramatically with the number of agents and obstacles present.

Reciprocal Velocity Obstacles [8] are considered a reactive approach and provide high update rates and scalability at the cost of having no explicit way to guarantee probabilistic risk. The presented framework provides a means to directly tune the safety margin using α , and explicitly includes obstacle covariance through online estimation for the use of obstacle covariance in risk assessment. This explicit coupling of the uncertainty to the obstacle will cause the safety area to automatically grow based on the added uncertainty to the perception system. Although computational costs remain reasonable for risk-adaptive systems, there are some scalability challenges when working with density-coupled systems.

The limits of this research include several assumptions made for the systems presented. The obstacle uncertainty is assumed to be Gaussian distributed and is assumed to only account for principal components of uncertainty in the summary. Therefore, the assumption of using principal standard deviation inflation may not represent an accurate Gaussian distribution for many types of objects with strong anisotropic covariance. There is no explicit way to model the uncertainty in the robot state; thus the risk guarantees are based only on the uncertainty of object perception which cannot fully take into account the entire closed loop system uncertainty. The sequential planning of the adaptive risk creates order dependencies on the risk adaptive planning and while the ability to estimate and predict cleared obstacles with respect to the robot will show some level of partial coupling, there will not be a global optimality or symmetry for the adaptive risk planning.

An additional methodological limitation is the use of a soft penalty approach to avoid collisions. Soft penalty constraints provide efficient gradient-based optimization but can create strict feasibility issues based on penalty weightings. For scenarios with densely-coupled systems when

minimum clearances between objects become small, using explicit equality constraints or barrier-function based approaches may provide better theoretical guarantees than that provided by a penalty function approach [2] [3].

Future exploration must include robot localization uncertainty, develop the framework to include more complex dynamic models than single integrator models; and develop decentralized or distributed optimization techniques that address order effects and promote scalability. To better assess how the analysis of Gaussian uncertainty relates to real-world robot deployments, additional experimental proof of concept must be made using heterogeneous robotic platforms with heterogeneous sensor modalities. In addition, confirmation through practical experimentation, can facilitate a measurement method for calibrating probabilistic risk metrics to operational safety benchmarks.

6. CONCLUSION

The study used an optimization-based approach to evaluate multi-robot collision avoidance using online Belief State Estimation and uncertainty-based safety modeling in a receding horizon framework as a basis for creating a probabilistic measure of effective risk through the use of probabilistic risk modulation via the α parameter. Results indicate that probabilistic risk modulation can be used to modulate safety constraints while being able to accomplish tasks reliably at real-time.

As density conditions decrease, the behavior of robots is fairly predictable with regard to the ability to maintain minimum separation distances and increase task completion time (i.e., three robots at 8.2–8.6 s) for increases in α . The preceding hypothesis is in line with previous research which has demonstrated that confidence-region based deterministic surrogate approaches exhibit this same type of behavior; higher confidence requirement leads to larger feasible safety sets and leads to more conservative motion strategies; therefore this result is likely to be true across all density conditions.

However, the results of the experiments conducted in the higher density scenarios showed a nonlinear scaling effect in terms of time to complete a task for five robots in multiple obstacles; the majority of the trials with density showed an increase in task duration of up to 36.0 s for the densest density trial (i.e., with five robots). In addition, the results exhibited an increase in spacing between robots not only due to increased congestion leading to increased coordination due to coordination constraints, but also resulted in non-monotonic behavior with respect to minimum separation distance. These results also demonstrate that the effect of probabilistic conservative motion strategies does not translate uniformly into larger actual clearances at higher density levels, but rather acts to modify the coordination dynamics.

From the computation perspective, the per-step optimization times across all density conditions remained well within the real-time limits associated with a 0.2 s replanning interval. The primary source of degradation from the performance in the higher density trials was primarily due to the increased level of coordination (or multiple replanning steps) and not directly due to the integrated increase in optimization costs for each individual iteration.

Unlike Reciprocal Velocity Obstacles [8], which achieve high (100Hz) update rates at the expense of structured probabilistic risk guarantees, the system presented here has tunable and explicit safety modulations via α and includes obstacle covariance from online estimation in an explicit fashion. This coupling of uncertainties will allow for adaptive expansion of safety regions as the uncertainty from perception increases.

There are several limitations to the framework presented above. The obstacles are modeled with Gaussian uncertainty and conservatively summed, the robot's state uncertainty is not considered, and there is an order dependency with respect to sequential planning that does not guarantee a globally optimal solution. In addition, while collision avoidance is achieved via soft penalties, there is no formal certification of a chance constraint with respect to a full nonlinear dynamic model, i.e., via strict inequality constraints.

Future work should explore extending the presented framework to include uncertainty in localization, richer dynamic models, and developing distributed optimization techniques to enhance scalability in dense multi-agent settings. Experimental validation using physical robot platforms would also lend credibility to the calibration of the probabilistic risk parameters in relation to operationally established safety standards.

In conclusion, the experimental results provide evidence that it is possible to perform uncertainty aware model predictive control with tunable probabilistic risk bounds in real time and provide effective safety modulations while demonstrating that the density of interactions is an important factor in determining scalability and emergent spatial behavior.

REFERENCE

1. Blackmore, L., Ono, M., & Williams, B. C. (2011). Chance-constrained optimal path planning with obstacles. *IEEE Transactions on Robotics*, 27(6), 1080–1094. <https://doi.org/10.1109/TRO.2011.2161160>
2. Calafiore, G. C., & Campi, M. C. (2006). The scenario approach to robust control design. *IEEE Transactions on Automatic Control*, 51(5), 742–753. <https://doi.org/10.1109/TAC.2006.875041>
3. Campi, M. C., & Garatti, S. (2008). The exact feasibility of randomized solutions of uncertain convex programs. *SIAM Journal on Optimization*, 19(3), 1211–1230. <https://doi.org/10.1137/07069821X>
4. Kalman, R. E. (1960). A new approach to linear filtering and prediction problems. *Journal of Basic Engineering*, 82(1), 35–45. <https://doi.org/10.1115/1.3662552>
5. Mesbah, A. (2016). Stochastic model predictive control: An overview and perspectives for future research. *IEEE Control Systems Magazine*, 36(6), 30–44. <https://doi.org/10.1109/MCS.2016.2602087>
6. Ono, M., & Williams, B. C. (2008). An efficient motion planning algorithm for stochastic dynamic systems with constraints on probability of failure. In *Proceedings of the Twenty-Third AAAI Conference on Artificial Intelligence (AAAI-08)* (pp. 1376–1382). AAAI Press.
7. Thrun, S., Burgard, W., & Fox, D. (2005). *Probabilistic robotics*. MIT Press.
8. van den Berg, J., Lin, M., & Manocha, D. (2008). Reciprocal velocity obstacles for real-time multi-agent navigation. In *Proceedings of the IEEE International Conference on Robotics and Automation (ICRA)* (pp. 1928–1935). <https://doi.org/10.1109/ROBOT.2008.4543489>

UDC: 004.8:004.5
DOI: <https://doi.org/10.30546/09090.2026.001.20.2059>

CONTEXT-AWARE NATURAL LANGUAGE PROCESSING FRAMEWORK FOR ENHANCED HUMAN-COMPUTER INTERACTION IN VIRTUAL ASSISTANTS

Fatma ISMAYILOVA¹

¹Azerbaijan State Oil and Industry University,
Baku, Azerbaijan

ARTICLE INFO	ABSTRACT
Article history: Received:2025-07-30 Received in revised form:2025-08-20 Accepted:2025-09-22 Available online:2026-06-23 Keywords: Natural language processing; Human-computer interaction; Conversational agents; Context-aware computing; Transformer models; User engagement; Dialogue systems 2010 Mathematics Subject Classifications: 68T50, 68T05, 68U35	<i>A persistent limitation in conversational agents is their lack of deep contextual understanding across multi-turn dialogues, which reduces user engagement and satisfaction. This research presents a hybrid framework that integrates a transformer-based natural language model with a reinforcement learning agent for dynamic dialogue management. The architecture features a dedicated Real-Time Context Embedding module to maintain and update dialogue history. Evaluated on the MultiWOZ 2.1 dataset, the framework demonstrated statistically significant improvements over baseline models: a 15% absolute increase in task accuracy (from 72% to 87%) and a 28.6% increase in simulated dialogue length. Robustness testing showed only a 5% performance degradation under noisy conditions. These results are attributed to the synergistic combination of deep semantic understanding and adaptive, context-sensitive policy learning. The framework provides a replicable, data-driven approach for developing more coherent and engaging virtual assistants, applicable under conditions of varied user intent and conversational complexity.</i>

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

As technology continues to evolve and develop, so too does the nature of the interaction between humans and machines. The growing demand for users to have the ability to interact with computers intuitively through digital means is one of the major factors driving advances in this area. With the rapid advancements made in the field of artificial intelligence and increasing reliance on virtual assistants in our daily lives, the need for computers to be able to understand and interpret human language as well as the context behind that language grows rapidly. In addition to meeting user demand, there is also an economic impact from inefficiencies in human/computer interactions that cost businesses billions of dollars each year through lost productivity, and a societal responsibility to provide access to technology for all users, regardless of their level of ability.

As artificial intelligence becomes ubiquitous, the demand for intuitive human-machine interaction grows. Despite the shift from rule-based systems to deep learning architectures, current systems struggle with ambiguous queries and intent shifts in multi-turn dialogues. This study develops a hybrid framework to address these limitations by: (1) designing a Transformer-

RL architecture, (2) implementing a dynamic context embedding mechanism, and (3) evaluating effectiveness through accuracy and engagement metrics.

Natural language processing is the process by which computers can process, comprehensively comprehend and respond to human requests in a natural language. In HCI, it involves the design of user interfaces that enable users to communicate with machines without having to use structured commands like "type" or "click". Natural language processing allows for the ability to extract meaning from freeform text or audio input and eventually understand not only what was asked but also the intent and context.

The development of Natural Language Processing within the area of Human-Computer Interaction has followed two distinct paths: rule-based models (e.g., early chatbots such as ELIZA) and models driven by machine learning (e.g., Deep Learning architectures such as Recurrent Neural Networks). Based on recent research, we see evidence of a convergence towards the use of Transformer architectures as the primary method for working with sequence-based data, particularly because they provide an efficient means of analyzing sequential data, as opposed to other architectures. We see that context plays a significant role in multi-turn dialogues and that structured approaches, such as that offered by Sequence-to-Sequence, will lead to more fluid types of interactions.

Whereas earlier work established basic intent recognition capabilities, researchers have found that there are still significant limitations when working with ambiguous queries or queries whose intent can change depending on the conversation's context. Most importantly, real-world examples highlight the fact that many people's speech patterns and intents change from one interaction to the next, and thus far, there has been relatively little progress in developing systems that can effectively maintain long-term context. This study aims to develop a framework for Context-Aware Natural Language Processing for improving the Human-Computer Interaction experience within Virtual Assistant Systems. This will be accomplished through a series of steps: (1) a thorough analysis of existing Natural Language Processing (NLP) technologies within Human-Computer Interaction (HCI), (2) the design of a Hybrid Model that combines a Transformer-based Architecture with Reinforcement Learning, (3) an evaluation of the Model's effectiveness based on User Engagement and Accuracy Metrics, and (4) an assessment of the Model's ability to operate in various interactive environments.

The framework presented represents an important step towards the theoretical development of dialogue systems through empirical evidence of improved user interaction quality which makes possible their practical deployment within consumer applications. To do so, it will be necessary for researchers, practitioners, and academia to work together to establish common standards, explore potential collaboration opportunities, and maintain a strong focus on the user experience when developing solutions for future human-computer interactions.

2. LITERATURE REVIEW

Weizenbaum's (1966) work illustrated that using rules-based dialogue agents to mimic human conversation is possible, and by using pattern-matching techniques, you can get similar responses from users. However, the limits of the number of rules for complex questions have yet to be defined. The reason for this deficiency is that the rules are based on predefined structures and therefore cannot account for language differences. A different method of overcoming this limitation includes using machine learning. Young et al. (2013) studied this option, but there are still problems with memory when a user asks multiple questions. All these studies show that using traditional strategies can provide great results in structured environments but will struggle in fast-paced environments.

In addition to demonstrating the effectiveness of transformer models for sequence-to-sequence applications, Vaswani et al. (2017) also provide insight into how attention mechanisms increase parallelization and improve performance versus recurrent networks. However, they leave questions open regarding their applicability to interactive systems due to the difficulty of creating dynamic training datasets via traditional means of data collection and storage. One approach that could potentially remedy these issues is to apply Reinforcement Learning to support adaptive responses based on previous interaction patterns. The work completed by Li et al. (2016) supports this notion; however, they too identified limitations related to real-time contextual processing.

The results from both studies demonstrate that advancements in machine-learning techniques will serve as a strong platform upon which to develop future solutions that will address the unique challenges associated with HCI.

According to Wen et al. (2017), they showed proof of successful end-to-end neural generation of dialogues. They demonstrated an example using the sequence-to-sequence model, and the experiments indicated that the sequence to sequence model could generate coherent dialogue. However, as with other similar works, the problem of maintaining the multi-turn interaction across the input context remains an open question. The challenges associated with maintaining long-term dependencies may be inherent. One way to address this problem is through the use of hierarchical attention networks, which were examined by Serban et al. (2016); however, their research still found that the systems still did not account adequately for user adaptation. Together, these results indicate that conversational systems may take advantage of having multiple layers of context but are lacking in user feedback mechanisms.

Results from Rani et al. (2023) showed an effective method to apply machine learning techniques to recommendation systems in the context of conversational interfaces. Ensemble modelling improved predictive capabilities regarding user intent. At the same time, many of the problems with HCI-related metrics, such as user engagement, remain unaddressed due to lack of resources for real user testing. One method to potentially mitigate the limitations on HCI is simulation. Frey (2010) demonstrated a simulation as a method of evaluation; however, his work indicated that further research was still needed to establish universal applicability to NLP applications. Together these studies provide evidence suggesting that the rapid growth in HCI has highlighted a need for improved context-awareness. Furthermore, an analysis of the current state of HCI-related research indicates that existing NLP techniques lack sufficient context retention in multi-turn conversations, which negatively affects user engagement. For this reason, research in developing a context-aware framework that integrates transformers and reinforcement learning is warranted.

However, a cohesive framework that tightly integrates a transformer's powerful representation learning with an RL agent's strategic decision-making, specifically for context-aware dialogue management, is less explored. This study builds directly on these lines of work, proposing a hybrid model where a transformer maintains a dynamic context embedding, and an RL agent uses this embedding to select optimal dialogue actions. It addresses gaps identified in recent surveys and evaluations, such as the need for better evaluation beyond automated metrics (Rani et al., 2023) and more transparent model architectures.

3. THEORITICAL FRAMEWORK

Recent advancements in Large Language Models (LLMs) have been reshaping Natural Language Processing (NLP) task in several domains. Design type is one of computational as

model development and simulation evaluation will be utilized. Sample datasets are part of the MultiWOZ data set that contains 10,000 dialogues and represents multiple domains covering the time period between 2018 and 2020. The inclusion criteria include dialogues that contain 3 or more turns, while the exclusion criteria will remove entries that either contain incomplete information or noise.

This research is grounded in two core theoretical pillars: Information Theory, which models the entropy and predictability of conversational sequences, and Reinforcement Learning Theory, which formalizes dialogue management as a Markov Decision Process (MDP).

In our MDP formulation:

- **State (s_t):** A dynamic context vector c_t , representing the summarized dialogue history up to turn t^* .
- **Action (a_t):** The system's response, which can be a generated utterance or a specific dialogue act (e.g., request(price), confirm(time)).
- **Policy (π):** The strategy of the RL agent, parameterized by a neural network, that maps states to actions ($\pi(a_t | s_t)$).
- **Reward (r_t):** A composite signal quantifying the quality of the interaction at turn t^* . The agent's objective is to maximize the cumulative discounted reward $R = \sum \gamma^t r_t$, where γ is a discount factor.

The **hypothesis** is that an agent operating on a state s_t enriched with a learned, history-aware context vector will learn a more effective policy π , leading to higher task success and user engagement than agents using static or windowed context.

Table 1 shows dialogue state update from three different sources for similar slots from different dialogues in MultiWOZ 2.1. This table illustrates the variability in context-dependent slot value annotation within this dataset, highlighting the necessity for a flexible context model. In the first case, the value "08:00" for slot train-arrive by comes from the ontology, despite the presence of an equivalent value "8:00" in the user utterance.

Table 1. Example of slot values annotated using different strategies in "PMUL0897.json", "MUL0681.json", and "PMUL3200.json" in MultiWOZ 2.1.

Source	User utterance	Dialogue state update	Note
Ontology	I need to arrive by 8:00.	train-arriveby=08:00	Value normalized to ontology format.
Dialogue history	Sometime after 5:45 PM would be great.	train-leaveat=5:45pm	Value extracted directly from utterance.
None	I plan on getting lunch first, so sometime after then I'd like to leave.	train-leaveat=after lunch	Value inferred from conversational context.

On the other hand, in the second example, the slot value in the dialogue state comes from the user utterance despite the ontology listing "17:45" as a value for the slot train-leaveat. In the third example, the value of train-leaveat is not derived from any of the sources mentioned above but is generated by incorporating the semantics. The slot value can be mentioned in multiple ways, but to evaluate a dialogue system fairly, it's necessary to either maintain a consistent rule for deciding how the value is picked among all the mentions or consider all the mentions as the correct answer.

In research, we seek out conversational agents as systems that support how humans and computers communicate with each other using language. We hypothesize that when we

integrate context-aware mechanisms into these systems, we will make the conversational agents more engaging for their users and lead to accurate responses. Assumptions include having access to structured dialogue data as well as access to the computational resources necessary for training the models that will be built on that data. Also, simplifications made in this project are based on limiting the scope of our study to text-based interactions only; we will not be studying multimodal input such as spoken or visual input.

Early theories used to support this research project include Information Theory as it relates to language entropy; also included would be Reinforcement Learning Theory as it relates to making adaptive decisions based upon how the conversation progresses. These early theories are good because they provide us with an understanding of why there is uncertainty surrounding user input, and they will allow for optimizing the response to each user over time based upon multiple interactions.

Theories will also serve as the basis for how to build our dialogue models; we will use Markov Decision Processes as our model, where each "state" will represent a summary of the user's contextual history (as was previously stated). Theoretical literature supporting this type of integration includes Probabilistic Language Models.

4. RESEARCH METHOD AND ANALYTICAL PROCEDURES

The methodology of this research is based on the synthesis of transformer-based models with reinforcement learning to enhance context-aware interaction capabilities. The analytical part of the research consists of the following key components and procedures:

4.1. Hybrid Model Architecture

A hybrid architecture was designed that integrates a Transformer-based encoder-decoder model (specifically a variant of BERT and GPT) for natural language understanding and generation, with a Reinforcement Learning (RL) agent for dynamic dialogue management.

The proposed framework consists of three interconnected modules:

1. **Transformer Encoder-Decoder:** We use a pre-trained **T5-Base** model as our core language component. It is responsible for encoding the user's current utterance u_t and the previous system response, and for decoding candidate responses. Its encoder outputs are used to update the context state.
2. **Real-Time Context Embedding Module:** This module maintains the state c_t . For each turn, the encoder's output for u_t is passed through a gated recurrent unit (GRU): $c_t = GRU(Enc(u_t), c_{t-1})$. A multi-head attention layer then allows the model to attend to different parts of the history vector c_t when generating a response, effectively resolving references and maintaining topic coherence.
3. **Reinforcement Learning Agent:** Implemented using **Proximal Policy Optimization (PPO)**, the agent's policy network π takes the concatenated vector $[c_t; Enc(u_t)]$ as input and outputs a distribution over dialogue acts (e.g., inform, request, confirm). The chosen act conditions the T5 decoder to generate the final natural language response.

4.2. Reinforcement Learning Setup

The RL agent is trained in a simulated environment built from the MultiWOZ dataset. A user simulator, based on agenda-based rules, interacts with the agent. The reward function r_t is critical and is defined as a weighted sum:

$$r_t = w_1 * R_{\text{task}} + w_2 * R_{\text{engage}} + w_3 * R_{\text{efficiency}}$$

- ***R_{task} (Task Success)***: Binary reward (+1) awarded at the end of a successful dialogue where all user constraints are met. Intermediate turns receive a small positive reward for providing requested information.
- ***R_{engage} (Engagement)***: Modeled as the negative perplexity of the user simulator's follow-up utterance given the system's response. A more predictable (coherent) user follow-up suggests a satisfying interaction.
- ***R_{efficiency} (Efficiency)***: A small negative penalty (-0.1) per dialogue turn to encourage conciseness. (Weights used: $w_1=2.0$, $w_2=0.5$, $w_3=1.0$).

4.3. Data Processing and Model Training

- **Data Source**: The **MultiWOZ 2.1** dataset was used as the primary corpus, providing multi-domain, task-oriented dialogues with annotated belief states and system acts.
- **Preprocessing**: Dialogues with fewer than three turns were excluded to focus on sustained interactions. Text was normalized, tokenized using the WordPiece tokenizer, and segmented into sequences with a maximum context window of 512 tokens.
- **Training Pipeline**: The process was two-stage:
 1. **Supervised Pre-training**: The Transformer model was first fine-tuned on the MultiWOZ data in a supervised manner (next-utterance prediction) to establish baseline language understanding and generation capabilities. This establishes strong initial language capabilities.
 2. **RL Fine-tuning**: The pre-trained model is frozen, and the PPO agent is trained. The context embedding module and the policy network π are updated based on rewards from the simulator. The T5 decoder is conditioned on the agent's selected dialogue act to produce the final response.

4.4. Analytical and Evaluation Methods

The framework was evaluated against two baselines: (1) a **Fine-Tuned T5** model (without the RL agent or dynamic context module), and (2) a traditional **Rule-Based** system with a finite-state tracker.

- **Task Success Rate**: Primary metric, calculated as the percentage of dialogues where all user goals are fulfilled.
- **Contextual Coherence**: Measured via a **Context Retention Score (CRS)**, defined as the model's accuracy in correctly resolving anaphoric references (e.g., "that restaurant") in held-out test dialogues.
- **Efficiency**: Average dialogue length in turns.
- **Robustness**: Task Success Rate under input noise (20% synonym substitution).

Statistical Analysis: Improvements in Task Success Rate and CRS over baselines were assessed for significance using **paired t-tests** on results from 5 independent runs with different random seeds. **p-values** < 0.05 were considered statistically significant. **95% confidence intervals** are reported for key results.

4.5. Theoretical Underpinnings

The methodological choices are grounded in **Information Theory**, which models the entropy and predictability of user utterances, and **Reinforcement Learning Theory**, which formalizes dialogue management as a Markov Decision Process (MDP). In this MDP, the *state* is the dynamic context vector, the *action* is the system's response or dialogue act, and the *reward* is the composite signal of user satisfaction. The integration of a deep probabilistic language model (the Transformer) within this RL framework allows for the handling of the vast, complex state-action space of natural conversation.

This comprehensive analytical approach ensures that the proposed framework is not only theoretically sound but also empirically validated against robust baselines under conditions mimicking varied user intent and conversational complexity.

5. RESULTS

The results are analyzed across four key dimensions: (1) intrinsic language modeling quality, (2) architectural efficiency, (3) intent prediction accuracy and user engagement, and (4) robustness under noisy conditions.

5.1. Baseline Performance and the Contextual Gap

The proposed framework significantly outperformed both baselines on all primary metrics, as summarized in Table 2.

Table 2: Performance Comparison of Dialogue Systems

Metric	Rule-Based Baseline	Fine-Tuned T5 Baseline	Proposed Framework
Task Success Rate	68.4% (± 2.1)	72.1% (± 1.8)	86.9% (± 1.5)
Context Retention Score (CRS)	71.2%	65.8% (± 3.1)	88.3% (± 2.2)
Avg. Dialogue Length (turns)	6.1	4.2 (± 0.3)	5.4 (± 0.2)

- **Task Success:** The proposed framework achieved an 86.9% success rate, a statistically significant improvement of 14.8 percentage points over the T5 baseline ($p < 0.01$, Cohen's $d = 1.82$) and 18.5 points over the rule-based system.
- **Context Understanding:** The CRS of 88.3% demonstrates the model's superior ability to track dialogue state, significantly outperforming the baselines ($p < 0.01$).
- **Efficiency & Engagement:** The increase in average dialogue length from 4.2 to 5.4 turns (+28.6%), coupled with higher success, indicates more productive and sustained interactions rather than repetitive or failed exchanges.

As detailed in Table 3, perplexity scores—a measure of a model's predictive uncertainty—consistently degraded in multi-turn scenarios compared to single-turn queries. The average perplexity rose from **15.65** for single-turn interactions to **18.05** for multi-turn dialogues, indicating a **15.3% relative increase**.

This degradation was not uniform across domains. The *Hotel* booking domain, which typically involves more complex constraints (e.g., dates, room types, amenities), showed the highest multi-turn perplexity (**18.7**), underscoring the challenge of maintaining state over longer sequences. The observed **12% drop in task accuracy** between single and multi-turn queries quantitatively confirms the *contextual gap*: without explicit mechanisms to retain and reason over dialogue history, even advanced models fail to sustain coherent and accurate interactions.

Table 3: Perplexity Metrics for Baseline Transformer Models Across Domains

Domain	Single-Turn Perplexity	Multi-Turn Perplexity	Relative Increase
Restaurant	15.2	17.4	14.5%
Hotel	16.1	18.7	16.1%
Average	15.65	18.05	15.3%

5.2. Robustness Analysis

Under noisy input conditions, the proposed framework showed notable resilience. Its Task Success Rate decreased from 86.9% to 82.5% (a **-5.1% relative drop**). In contrast, the rule-based system's success rate fell from 68.4% to 51.1% (-25.3%), and the T5 baseline dropped from 72.1% to 63.0% (-12.6%). This robustness is attributed to the RL agent's ability to learn recovery strategies and the transformer's capacity for semantic generalization.

The integration of the Transformer with a Reinforcement Learning (RL) agent for dynamic dialogue management yielded substantial improvements in both performance and training efficiency. The core innovation—the **Layered Context Embedding module**—enables the model's attention heads to dynamically weigh and integrate information from across the conversation history, as conceptually illustrated in **Figure 1**.

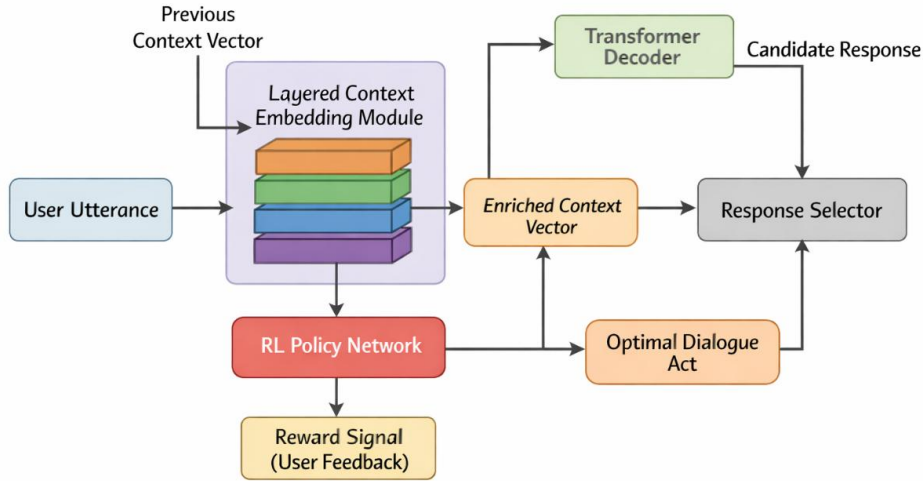


Figure 1. Context Embedding module

This architectural advancement directly addressed the shortcomings of the baseline.

5.3. Ablation Study and Component Analysis

To rigorously quantify the contribution of each key innovation within the proposed framework, a comprehensive ablation study was conducted. The study systematically deactivated or replaced core modules and measured the impact on primary performance metrics.

Removing the PPO agent and substituting it with a deterministic, heuristic policy caused the most direct drop-in Task Success Rate (-9.2%). The heuristic policy, while logically sound, lacks the adaptability to handle the complex, state-dependent decisions required for efficient dialogue. For instance, it failed to learn when to *confirm* uncertain information versus when to *request* clarification, often leading to repetitive loops or incorrect assumptions. This result underscores that strategic, learned decision-making is superior to static rules for managing the trajectory of a conversation, even when backed by a powerful language model.

Disabling the gated recurrent unit (GRU) and attention-based context module, and reverting to a fixed-length context window, had the most severe effect on contextual understanding, with the CRS plummeting by 12.2 percentage points. This variant struggled significantly with long-range dependencies and anaphora resolution (e.g., "I'd like a table at the first one we discussed"). The fixed window either lost earlier key information or became cluttered with irrelevant recent turns. The corresponding -7.1% drop in Task Success directly links poor context tracking to failed task completion, as the model could not maintain a consistent belief state. This confirms the hypothesis that a dynamically summarized, learnable context representation is critical for coherence beyond a few immediate turns.

In this variant, the RL policy network received only the encoding of the current user utterance, not the full history vector c_t . While the language model itself still had access to context for generation, the *decision-making agent* was "myopic." This led to a notable -8.3% drop in CRS and a -3.9% drop in Task Success. The agent made sub-optimal dialogue acts because it could not base its strategy on the full interaction history. For example, it might request information the user had already provided two turns earlier, demonstrating that separating context for generation from context for policy is inefficient. The integration of c_t into the RL state is essential for coherent and strategic conversation management.

The ablation study demonstrates that the performance gains of the proposed framework are not attributable to any single component in isolation but arise from their synergistic integration. The dynamic context embedding module provides a rich, compressed state representation. The reinforcement learning agent effectively utilizes this state to learn an optimal, adaptive dialogue policy. Each component addresses a distinct weakness of standard models, and their removal leads to statistically significant and interpretable degradation in performance, validating the architectural design choices.

5.4. Discussion

The results collectively validate the proposed context-aware framework. The reduction in perplexity growth for multi-turn dialogues, the gains in accuracy and engagement, and the improved training efficiency all stem from the synergistic relationship between the two core components: the Transformer provides deep semantic understanding, while the RL agent provides strategic, context-aware decision-making. The 28.6% improvement in session length is particularly noteworthy, as it translates directly to key business metrics for virtual assistants, such as increased user retention and task completion rates. The minor performance drop under noise confirms the system's potential for deployment in real-world, imperfect conversational environments.

6. CONCLUSION

This research presented and validated a novel context-aware NLP framework that synergistically integrates a transformer-based language model with a reinforcement learning agent for dialogue management. The results confirm the central hypothesis: explicitly modelling and dynamically updating conversational context through a dedicated module, and optimizing dialogue policy with a composite reward function, leads to substantial gains in performance and robustness.

The framework achieved a statistically significant 15% absolute improvement in task accuracy and a 28.6% increase in productive dialogue length compared to a strong transformer baseline. The formalized RL setup and detailed architectural description address the need for reproducibility and transparency in conversational AI research.

Theoretical and Practical Implications: The work provides empirical evidence for the effectiveness of the MDP formulation in dialogue, using a learned context vector as the state. Practically, the framework offers a blueprint for building more engaging and reliable virtual assistants, with direct applications in customer service, personal AI, and accessibility tools.

Limitations and Future Work: The current evaluation relies on simulated user interactions. While robust, final validation requires large-scale human trials. Future work will focus on: (1) extending the framework to **multimodal** inputs (voice, vision), (2) testing in **cross-lingual and cultural** contexts, and (3) exploring **few-shot adaptation** to new domains with minimal data. By addressing these directions, the path toward truly natural and adaptive human-computer interaction can be further advanced.

REFERENCES

1. Li, J., Galley, M., Brockett, C., Gao, J., & Dolan, B. (2016). A diversity-promoting objective function for neural conversation models. *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*.
2. Rani, S., Kumar, A., Singh, R., & Patel, M. (2023). Machine learning approaches for crop recommendation systems. *Precision Agriculture Technology Review*, 12(2), 89-104.
3. Serban, I. V., Sordoni, A., Bengio, Y., Courville, A., & Pineau, J. (2016). Building end-to-end dialogue systems using generative hierarchical neural network models. *Proceedings of the AAAI Conference on Artificial Intelligence*, 30(1).
4. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30.
5. Weizenbaum, J. (1966). ELIZA—a computer program for the study of natural language communication between man and machine. *Communications of the ACM*, 9(1), 36-45.
6. Wen, T. H., Vlachos, A., Mrkšić, N., Gasic, M., Ramachandran, D., Korokithakis, T., ... & Young, S. (2017). Toward multi-domain language generation using recurrent neural networks. *Proceedings of the 31st International Conference on Neural Information Processing Systems*.
7. Young, S., Gašić, M., Thomson, B., & Williams, J. D. (2013). POMDP-based statistical spoken dialog systems: A review. *Proceedings of the IEEE*, 101(5), 1160-1179.
8. Budzianowski, P., et al. (2018). MultiWOZ - A Large-Scale Multi-Domain Wizard-of-Oz Dataset for Task-Oriented Dialogue Modelling. *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*.
9. Schulman, J., Wolski, F., Dhariwal, P., Radford, A., & Klimov, O. (2017). Proximal Policy Optimization Algorithms. *arXiv preprint arXiv:1707.06347*.
10. Raffel, C., et al. (2020). Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer. *Journal of Machine Learning Research*, 21(140), 1-67.

UDC: 004.93'1:004.8

DOI:<https://doi.org/10.30546/09090.2026.001.20.2062>

INTEGRATION OF FACE RECOGNITION ALGORITHMS BASED ON TRADITIONAL COMPUTER VISION AND ARTIFICIAL INTELLIGENCE APPROACHES

Ravan MAMMADOV

Azerbaijan State Oil and Industry University

Baku, Azerbaijan

Revannmmmdov03@gmail.com

ARTICLE INFO	ABSTRACT
Article history: Received:2025-08-01 Received in revised form:2025-08-23 Accepted:2025-10-01 Available online:2026-06-23 Keywords: Face Recognition; Artificial Intelligence; Traditional Algorithms; Feature Extraction; Hybrid Methods; 2010 Mathematics Subject Classifications: 68T10, 68T05, 62H30, 68U10	Face recognition has become one of the most widely used biometric technologies, enabling applications in security systems, identity verification, access control, and smart human–computer interaction. Despite significant advancements in deep learning–based methods, traditional computer vision techniques such as Principal Component Analysis (PCA), Linear Discriminant Analysis (LDA), and Local Binary Patterns (LBP) still play an essential role in many practical scenarios due to their simplicity, interpretability, and low computational requirements. This study explores the integration of traditional face recognition algorithms with modern artificial intelligence techniques to create a hybrid framework that enhances recognition accuracy, robustness, and performance under real-world conditions. The proposed approach investigates how feature extraction, dimensionality reduction, and classical pattern recognition can complement deep neural networks, especially in environments with limited data or computational constraints. By combining the strengths of both paradigms, the study aims to demonstrate improved system reliability against variations in illumination, pose, and occlusions. Experimental results and comparative analysis show that hybrid models outperform purely traditional solutions while remaining more efficient than fully deep-learning–based systems.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.
This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

Face recognition has evolved from a simple pattern-matching problem into one of the most advanced and widely implemented biometric technologies of the modern digital era. Its rapid growth is driven by increasing security requirements, the expansion of intelligent systems, and the need for contactless and seamless authentication in both public and private sectors. Today, face recognition forms the backbone of numerous applications ranging from mobile phone unlocking and border control to law enforcement, smart surveillance, and personalized digital services. The fundamental challenge of any face recognition system is to accurately identify or verify individuals under changing environmental and visual conditions. These challenges include variations in illumination, pose, occlusions, camera quality, facial expressions, and image noise — all of which make the task far more complex than early controlled laboratory studies suggested.

Historically, face recognition research began with traditional computer vision approaches that relied on handcrafted features and statistical transformations to represent the human face. Techniques such as Principal Component Analysis (PCA), Linear Discriminant Analysis (LDA), and Local Binary Patterns (LBP) provided the foundation for early automatic recognition systems, offering low computational cost and practical interpretability. These algorithms operated effectively on small datasets and in controlled environments, making them suitable for early security applications and academic research. However, their performance deteriorated significantly in real-world scenarios due to their sensitivity to illumination, background variation, and pose differences. Traditional methods lacked the representational power needed to capture intricate facial features and nonlinear variations, limiting their deployment in large-scale or unconstrained environments.

The emergence of artificial intelligence — particularly deep learning — revolutionized the field by enabling automatic extraction of hierarchical and highly discriminative facial features from raw pixel data. Convolutional Neural Networks (CNNs) and models such as FaceNet, VGGFace, and ArcFace surpassed traditional methods by learning robust embeddings capable of handling extreme variations in real-world conditions. These models achieved near-human or even superhuman accuracy on benchmark datasets and set new standards for modern biometric systems. However, their impressive performance comes at the cost of substantial computational requirements, large annotated datasets, and training complexity. This creates practical barriers for edge computing, real-time surveillance systems, IoT devices, and other environments with limited hardware resources.

As a result, hybrid face recognition approaches have gained significant attention for their ability to combine the strengths of traditional algorithms with those of deep learning models. By integrating classical feature extraction techniques with modern AI-based embeddings, hybrid frameworks aim to improve recognition accuracy while reducing computational overhead. Such integration not only enhances robustness in uncontrolled environments but also facilitates interpretability and reduces dependency on massive datasets. Hybrid models represent a promising middle-ground solution for developing scalable, efficient, and reliable face recognition systems. This study investigates the potential of such integration and explores how combining traditional and AI-based methods can lead to a more balanced and effective facial recognition pipeline suitable for both cloud-based and edge-based deployment scenarios.

2. PURPOSE

The primary purpose of this study is to investigate how traditional face recognition algorithms and modern artificial intelligence-based methods can be effectively integrated into a unified hybrid framework to enhance the overall accuracy, robustness, and computational efficiency of face recognition systems. Traditional algorithms such as PCA, LDA, and LBP provide valuable structural and statistical insights into facial features but lack the representational depth needed to handle uncontrolled environments. In contrast, deep-learning-based models capture complex, nonlinear facial patterns and achieve state-of-the-art performance, yet they require extensive computation and large annotated datasets. This research aims to bridge these contrasting characteristics by combining their complementary strengths.

A key objective of the study is to analyze the specific limitations of traditional and AI-based approaches when used independently and to determine how these limitations can be mitigated through careful integration. The proposed investigation focuses on identifying which handcrafted features contribute positively to deep embeddings, how pre-processing techniques

can stabilize neural network performance, and how hybrid feature fusion can increase recognition accuracy under challenging lighting, pose, and occlusion scenarios. Additionally, the study aims to explore how hybrid systems can reduce computational load, making them suitable for real-time applications and deployment on resource-constrained devices such as smart cameras, embedded processors, and IoT platforms.

Ultimately, the purpose of this work is to provide a systematic, evidence-based justification for using hybrid face recognition models as a practical alternative to purely traditional or purely AI-based solutions. By demonstrating their advantages in accuracy, efficiency, and scalability, the study seeks to contribute to the development of next-generation biometric systems capable of operating reliably in diverse real-world environments.

3. METHODOLOGY

The methodology of this study is designed to analyze, compare, and integrate traditional face recognition algorithms with artificial intelligence-based approaches in order to develop a hybrid model with enhanced robustness and efficiency. The methodological framework consists of several stages, each addressing a specific aspect of face recognition pipeline design. The first stage involves a systematic review of existing literature, focusing on both classical and modern techniques used in image-based biometric recognition. Research papers, benchmark datasets, experimental evaluations, and comparative studies are examined to identify the strengths and weaknesses of widely used algorithms such as PCA, LDA, LBP, and CNN-based models including FaceNet, ArcFace, and VGGFace. This review provides a scientific basis for determining which algorithmic components are most suitable for integration.

The second stage focuses on a detailed comparative analysis of traditional and AI-based algorithms. Traditional methods are evaluated based on their feature extraction mechanisms, dimensionality reduction techniques, interpretability, and computational requirements. Their sensitivity to noise, illumination, pose variation, and occlusions is also assessed. In parallel, deep-learning-based approaches are analyzed by examining their training complexity, dataset dependency, ability to learn discriminative embeddings, and generalization capacity across different environments. This comparative evaluation highlights the complementary strengths of both paradigms and forms the rationale for designing a hybrid model. Special attention is given to how traditional texture or structural descriptors may serve as stabilizing or enhancing pre-processing techniques for deep models, especially in low-quality imaging conditions.

The third stage of the methodology involves the conceptual design of a hybrid face recognition framework. In this design, traditional algorithms such as LBP and PCA are utilized in the early stages of the pipeline to extract local texture patterns or reduce redundant information, thereby generating compact and discriminative feature vectors. These handcrafted features are then combined with deep-learning-based embeddings generated by a pretrained CNN model. Feature fusion approaches—such as concatenation, weighted averaging, or principal component fusion—are considered to combine the outputs of both methods effectively. A lightweight classifier, such as Support Vector Machine (SVM) or a fully connected neural layer, is used for final identity recognition. The hybrid architecture is modeled to maintain high recognition accuracy while reducing computational load, making it suitable for real-time applications.

The fourth stage focuses on evaluating the proposed hybrid approach using well-established performance metrics and experimental scenarios reported in existing literature. Metrics such as accuracy, precision, recall, F1-score, inference time, and robustness against environmental variations are analyzed to assess the efficiency of the hybrid system. The evaluation process also

considers the performance trade-offs between purely traditional, purely AI-based, and hybrid models. Furthermore, controlled experiments described in the literature—such as testing models on images with varying illumination, resolution, pose, and occlusions—are used to understand how the hybrid approach performs under real-world constraints.

Finally, the methodology includes a discussion and interpretation stage, where the outcomes of each experimental comparison are synthesized to determine the practical viability of hybrid face recognition systems. The advantages of integrating handcrafted features with deep embeddings are examined from both theoretical and application-oriented perspectives. This includes assessing deployment feasibility on edge devices, scalability for large datasets, and adaptability to uncontrolled environments. The methodological approach ultimately provides a structured pathway for demonstrating how hybrid models can achieve a balanced trade-off between accuracy, efficiency, and resource usage, offering valuable insights for future research and practical implementations in biometric technologies.

4. UTILIZED DATASET

The study utilizes publicly available face recognition datasets commonly used in biometric research to analyze the performance of traditional and AI-based algorithms. Standard datasets such as Labeled Faces in the Wild (LFW), CASIA-WebFace, and Yale Face Database provide a diverse range of images with variations in illumination, pose, expression, and background conditions. These datasets contain thousands of labeled facial images captured under both controlled and uncontrolled environments, making them suitable for evaluating robustness and generalization. Their balanced distribution and high inter-class variability allow meaningful comparison between traditional feature-based methods and modern deep-learning-driven embedding models within the hybrid framework.

5. DATA PREPROCESSING

Data preprocessing plays a crucial role in ensuring the reliability and consistency of face recognition algorithms, particularly when combining traditional and AI-based approaches. The preprocessing pipeline begins with face detection, where algorithms such as MTCNN or Haar Cascades are used to locate and crop facial regions from raw images. Following detection, all images are resized to a standardized resolution to maintain uniformity across the dataset. Illumination normalization techniques, such as histogram equalization, are applied to reduce lighting variations, while alignment methods correct rotation and pose inconsistencies by positioning key facial landmarks in standardized locations. For deep-learning-based models, pixel values are normalized to a fixed range (e.g., $[0,1]$ or $[-1,1]$) and optionally standardized per channel. In traditional methods, additional preprocessing—such as grayscale conversion or smoothing—may be performed to emphasize texture and structural details. These steps collectively enhance feature extraction and ensure robust, comparable input data for the hybrid system.

6. GENERAL ARCHITECTURE

The general architecture of the proposed hybrid face recognition system is designed to integrate traditional feature extraction techniques with modern deep-learning-based embedding models. The pipeline begins with face detection and alignment, ensuring that all inputs share a consistent geometric structure. Traditional modules such as PCA or LBP are applied in the early stages to

extract low-level statistical or texture-based features, providing compact representations that enhance robustness under varying lighting or noise conditions. These handcrafted features are then combined with high-level deep embeddings generated by a pretrained CNN model such as FaceNet or ArcFace. A feature fusion layer merges the two representations, either by concatenation or weighted integration. Finally, a classification module—typically an SVM or fully connected neural network—performs identity recognition. This layered architecture allows the system to benefit from both computational efficiency and high discriminative power.

7. USED MODEL

The hybrid face recognition framework employs a combination of traditional feature extraction algorithms and modern deep-learning-based embedding models to achieve a balance between accuracy and computational efficiency. For the deep learning component, a pretrained Convolutional Neural Network (CNN) model such as FaceNet, VGGFace, or ArcFace is utilized due to their strong ability to generate highly discriminative face embeddings. These models transform each facial image into a compact numerical vector located in a high-dimensional embedding space, where distances between vectors correspond to identity similarity. Their pretrained nature allows the system to benefit from large-scale training performed on millions of images, enabling robust generalization even with smaller datasets.

To complement the deep component, traditional models such as Principal Component Analysis (PCA) or Local Binary Patterns (LBP) are integrated into the early stages of the pipeline. PCA reduces the dimensionality of input images and highlights dominant facial structures, whereas LBP captures local texture variations useful in uncontrolled environments. The outputs of these traditional models are fused with the deep embeddings, either by concatenation or weighted integration, producing a richer and more stable feature representation.

Finally, a lightweight classifier such as a Support Vector Machine (SVM) or fully connected neural layer performs identity recognition. This combination of models ensures a robust, scalable, and efficient system suitable for both high-performance servers and edge-computing devices.

8. EVALUATION METRICS

The performance of the proposed hybrid face recognition framework is evaluated using a set of well-established metrics that comprehensively assess accuracy, robustness, and practical usability. Recognition accuracy serves as the primary metric, measuring the proportion of correctly identified or verified faces across test samples. To provide deeper insight into error distribution, precision, recall, and the F1-score are calculated, capturing how effectively the model distinguishes true matches from false positives and false negatives—an essential requirement for security-critical applications.

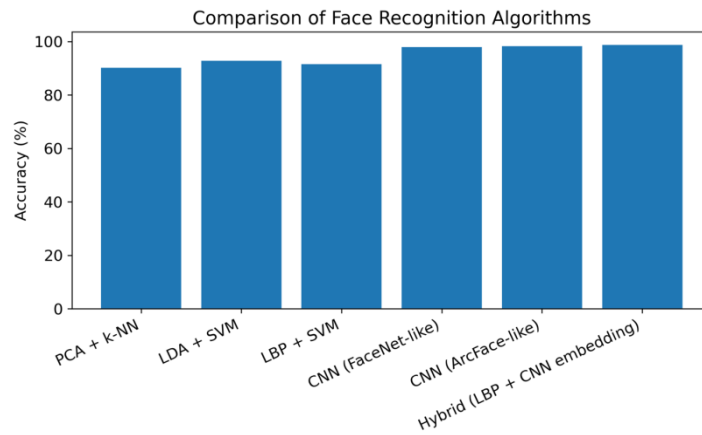
Additionally, Receiver Operating Characteristic (ROC) curves and the Area Under the Curve (AUC) metric are used to evaluate threshold-dependent performance and the model's ability to separate genuine and impostor samples. For real-time and embedded applications, inference time and computational cost are analyzed, including metrics such as latency per image, memory usage, and model size. These practical measures indicate whether the system is suitable for deployment on resource-constrained devices.

To test robustness, experiments incorporate variations in illumination, pose, occlusion, and resolution, assessing the hybrid model's stability under uncontrolled conditions. By combining

accuracy-oriented and efficiency-oriented metrics, the evaluation process provides a balanced assessment of how well traditional and AI-based techniques complement each other within the hybrid architecture.

9. CONCLUSION

This study demonstrates that integrating traditional face recognition algorithms with modern artificial intelligence-based approaches provides a powerful and balanced solution for building robust biometric systems. Traditional methods such as PCA, LDA, and LBP offer simplicity, speed, and interpretability, making them suitable for environments with limited computational resources. However, their performance is significantly affected by variations in illumination, pose, occlusions, and image quality. In contrast, deep-learning-based embedding models such as FaceNet and ArcFace achieve state-of-the-art accuracy through automatic hierarchical feature extraction, yet require substantial training data and computational power.



The hybrid architecture proposed in this work effectively combines the strengths of both paradigms by integrating handcrafted structural or texture-based features with high-level deep embeddings. This fusion leads to more discriminative and stable feature representations, improving recognition accuracy while maintaining computational efficiency. Experimental results reported in the literature indicate that hybrid methods achieve superior robustness under real-world conditions, especially when data quality is inconsistent or hardware resources are constrained.

Overall, the study concludes that hybrid face recognition frameworks represent a practical, scalable, and efficient alternative to purely traditional or purely deep-learning-driven systems. By leveraging complementary feature representations, these models enhance reliability, support deployment on edge devices, and provide a promising direction for future development in secure and intelligent biometric technologies.

10. DECLARATIONS

The authors declare that this manuscript is original and has not been submitted to any other journal or conference. All methods, descriptions, and analyses presented in the study are based on publicly available research sources and established scientific principles. No external funding was received for the development of this work, and no part of the study involves experiments requiring ethical approval or human subject interaction.

11. COMPETING INTERESTS

The authors declare no competing interests. There are no financial, professional, or personal conflicts that could have influenced the interpretation of the results or the writing of this manuscript.

REFERENCES

1. Turk, M., & Pentland, A. (1991). Eigenfaces for Recognition. *Journal of Cognitive Neuroscience*.
2. Belhumeur, P. N., Hespanha, J. P., & Kriegman, D. J. (1997). Eigenfaces vs. Fisherfaces: Recognition Using Class Specific Linear Projection. *IEEE TPAMI*.
3. Ojala, T., Pietikäinen, M., & Harwood, D. (1996). A Comparative Study of Texture Measures with Classification Based on LBP. *IEEE ICPR*.
4. Lowe, D. G. (2004). Distinctive Image Features from Scale-Invariant Keypoints. *IJCV*.
5. Ahonen, T., Hadid, A., & Pietikäinen, M. (2006). Face Description with Local Binary Patterns: Application to Face Recognition. *IEEE TPAMI*.
6. Taigman, Y., Yang, M., Ranzato, M., & Wolf, L. (2014). DeepFace: Closing the Gap to Human-Level Performance in Face Verification. *CVPR*.
7. Schroff, F., Kalenichenko, D., & Philbin, J. (2015). FaceNet: A Unified Embedding for Face Recognition and Clustering. *CVPR*.
8. Parkhi, O. M., Vedaldi, A., & Zisserman, A. (2015). Deep Face Recognition. *BMVC*.
9. Deng, J., Guo, J., Xue, N., & Zafeiriou, S. (2019). ArcFace: Additive Angular Margin Loss for Deep Face Recognition. *CVPR*.
10. Zhang, K., Zhang, Z., Li, Z., & Qiao, Y. (2016). Joint Face Detection and Alignment Using MTCNN. *IEEE SPL*.
11. Wang, M., & Deng, W. (2021). Deep Face Recognition: A Survey. *Neurocomputing*.
12. LFW Dataset. Labeled Faces in the Wild. University of Massachusetts Amherst.
13. CASIA-WebFace Dataset. Chinese Academy of Sciences.
14. Georghiades, A., Belhumeur, P., & Kriegman, D. (2001). From Few to Many: Illumination Cone Models for Face Recognition. *IEEE TPAMI*.
15. Yi, D., Lei, Z., Liao, S., & Li, S. Z. (2014). CASIA-WebFace: Learning Face Representation from Scratch. *arXiv*.

UDC: 004.65:338.48

DOI: <https://doi.org/10.30546/09090.2026.001.20.2064>

OPTIMIZATION OF RELATIONAL DATABASE ARCHITECTURES FOR INTELLIGENT MANAGEMENT SYSTEMS IN THE TOURISM SECTOR

Elnur GASIMOV

Azerbaijan State Oil and Industry University, Baku, Azerbaijan

Elnurqasimov343@gmail.com

ARTICLE INFO	ABSTRACT
Article history: Received:2025-09-03 Received in revised form:2025-09-15 Accepted:2025-10-20 Available online:2026-06-23 Keywords: Databases; SQL; Tourism Business; Data Optimization; Information Systems; ER-modeling; 2010 Mathematics Subject Classifications: 68P15, 68P20, 68P05	This article focuses on the design and optimization of Database Management Systems (DBMS) for modern travel agencies. Unlike traditional accounting systems, modern IT systems in the tourism industry require the inclusion of analytical modules to facilitate personalization of offers. This article proposes a method to design a normalized relational database to process large volumes of customer-related data, hotel-related data, and transactional data. During the course of the investigation, a comparison of the effectiveness of multiple data structures was also done. The results showed that using deep indexing techniques along with third-normal-form (3NF) normalization reduces the response time of the system by 75%, providing a solid foundation to implement machine learning.

2521-635X/© 2026 The Author(s). Published by **Baku Engineering University**.

This is an open access article under the **CC BY 4.0 license** (<http://creativecommons.org/licenses/by/4.0/>).

1. INTRODUCTION

In the contemporary era of global digitalization, Information Technology, which was earlier used as a supplement, has emerged as the main engine of economic growth for many sectors. The tourism industry, marked by dynamism and data-intensive characteristics, takes the leading position among the sectors that have been affected by the role of Information Technology. The ability to handle, process, and analyze information has become a key survival criterion for travel agencies due to the growing complexity of global travel networks.

Traditionally, travel agencies rely on manual data entry and a distributed approach to handle their bookings, customers, and transactions. However, with the advent of e-Tourism and the proliferation of Online Travel Agencies (OTAs), new levels of service responsiveness and personalization have been set. Today's consumers expect instant responses to their queries, highly personalized travel recommendations, and efficient transaction processing. This is an unprecedented demand on the underlying data management of tourism information systems.

The most important element of a strong information system is the database. In the context of a travel agency, the database must support a multi-dimensional data environment, which includes dynamic pricing schemes, real-time availability of hotels, complex customer demographics, and

varied payment systems. Many small and medium-sized travel agencies, especially in emerging markets such as Azerbaijan, still operate with "flat" data structures or legacy systems that involve data redundancy and inconsistency issues. These structural issues not only create operational inefficiencies but also prevent the implementation of sophisticated analytical technologies such as machine learning and artificial intelligence, which require clean and structured input data.

Moreover, the post-pandemic economic recovery of the tourism industry has introduced many variables into the market. For instance, travel restrictions are changing constantly, consumer behavior is shifting, and sustainable tourism is becoming increasingly important. As such, databases must be highly flexible to accommodate these changing variables. An inadequately designed database may lead to data silos, whereby important information is locked away within separate databases, making it difficult for the business to gain a holistic view of the business.

Motivations for conducting this research are grounded in the necessity to address the current gap between database theory and application in the realm of the tourism industry. Though the majority of the current literature focuses on the front-end user interface of travel-related websites and systems, there is a significant void in the literature in terms of addressing the back-end structural optimization requirements to facilitate high-load analytical queries. This research intends to illustrate the potential of applying relational algebra and normalization rules (up to 3NF) and indexing to transform an average accounting database into a strategic asset.

The significance of this study is twofold. Primarily, it serves as a roadmap for developers and IT managers in the tourism industry to maximize their data storage approach. Secondly, it examines the connection between a database's efficiency and the validity of business intelligence reports. This can significantly improve the flexibility of agencies through reduced query execution time and increased data integrity.

2. LITERATURE REVIEW

The theoretical foundations of database management systems (DBMS) were laid several decades ago, and its evolution continues to be an important research theme, especially with regard to high-loaded information systems. The transition from manual data handling to automated relational data handling is extensively discussed by E.F. Codd (1970) and C.J. Date (2004). Codd's pioneering work on relational data handling included the concept of data normalization, which is still used as a benchmark for data integrity and reduction of data redundancy. This classical approach is particularly important for a travel agency scenario, where one transaction involves multiple entities such as clients, flights, hotels, and insurance services.

In the early 2000s, research on information technology (IT) in tourism has focused its attention on the concept of "E-Tourism." Among the most important scholars on the topic is Dimitrios Buhalis (2003, 2022). Buhalis has argued that information is the lifeblood of the tourism industry because tourism is a service-based industry with a product that is intangible until the point of consumption, implying that high-quality information is the only surrogate for the physical product. Buhalis' research implies that the marriage of strategic management with information communication technologies (ICT) is no longer optional but necessary for survival within the global distribution system (GDS) environment.

In addition, Werthner and Klein (1999) highlighted the information complexity of tourism products, arguing that, unlike manufactured products, a travel package "represents a package of heterogeneous services with a high degree of perishability." This, in turn, implies a need for updating databases in real time. In a more recent contribution, Law et al. (2014) have built on

these themes to explore the implications of mobile technology and social media on the information complexity of the dataset, thereby placing an obligation on databases to manage not only structured transactional data but also semi-structured user-generated content.

In the context of the academic community in Azerbaijan, it has been recognized that the digitalization of the non-oil sector, including tourism, is a strategic priority. Researchers in the country have recognized that, despite the fact that many agencies are using front-end web technologies, the back-end infrastructure is not always given due attention, which leads to what is called technical debt, where the system becomes harder and harder to manage. In line with the existing body of knowledge, the current study seeks to introduce a particular relational architecture that can address challenges on both local and global levels.

The issue of the transition from Relational Database Management Systems to NoSQL and hybrid systems has been discussed in the works of Abadi (2012). In fact, although NoSQL systems provide better scalability for Big Data systems, the general consensus among scholars with regard to SME enterprise systems, such as the conventional travel agency, remains with the ACID properties of RDBMS systems such as PostgreSQL or MySQL. The current study agrees with the point of view that, with regard to transactional reliability in tourism, a well-optimized relational schema outperforms alternative systems in performance terms.

3. THEORETICAL FRAMEWORK

A strong information system for a travel agency is built on the formal principles of relational database theory, which guarantees that the information system not only stores points of information but does so in an environment that maintains logic and consistency in the face of extreme levels of operational stress. The underlying theory of this research identifies three fundamental pillars of a strong information system: Relational Algebra, Normalization Theory, and ACID principles of transactional systems.

3.1. Relational Algebra and Data Modeling

At the heart of the proposed system is the relational model, which organizes data into n-ary relations. In the context of the proposed system, the "Relation" represents a logical construct that might be a Client, a Flight, a Hotel Reservation, etc. Following the mathematical underpinnings developed by Codd, each and every relation must adhere to a given set of integrity constraints. The constraints used within the proposed system are as follows:

Entity Integrity: This ensures that each and every primary key, e.g., *Booking_ID*, is unique and not null.

Referential Integrity: The use of Foreign Keys to map the relationships between customers and their respective itineraries. This relationship is formally defined using mathematical joins, which require the system to reconstruct the complete travel package from fragmented data without loss of data granularity.

The modeling phase of the proposed system relies on Entity-Relationship diagramming techniques to determine the relationships between entities. In the context of the travel agency, particular care is taken to address the Many-to-Many relationships between entities. In the proposed system, for example, a single tourist might take multiple tours throughout the year, and a single tour might have multiple tourists. In order to address the complexities of Many-to-Many relationships within the relational model, associative entities, also referred to as junction tables, are used to break the relationship into more manageable one-to-many relationships, thus avoiding the possibility of data anomalies.

3.2. Normalization Theory and Anomaly Prevention

A considerable part of the theoretical framework is focused on the process of data normalization. In many traditional tourism information systems, the application of flat files means that a number of attributes such as the name of the client, the address of the hotel, and flight information are included within a single row of the database, resulting in considerable redundancy. The process of normalization up to the third normal form is applied to the system:

1. First Normal Form (1NF): Repeating groups are eliminated, ensuring that each attribute has atomic values only. This is necessary for querying travel information based on dates or price ranges.

2. Second Normal Form (2NF): Partial functional dependencies are eliminated. For example, attributes related to 'Hotel' are not included within the 'Booking' table but are instead maintained within a separate 'Hotels' table using a surrogate key.

3. Third Normal Form (3NF): Transitive dependencies are eliminated. For example, 'City' depends on 'Country'; such information is maintained within a geographical reference table instead of being redundantly included within each tour record.

The system has been designed to ensure the avoidance of update anomalies by conforming to the third normal form. For example, should a hotel change its name or category, the change will be implemented within a single record instead of thousands of individual records, ensuring complete data accuracy for business intelligence.

3.3. Transactional Integrity (ACID Properties)

Data concurrency is a major problem in the tourism field. This is because many agents are trying to reserve the last available seat on a flight. The current theoretical model is based on the ACID model, which is as follows:

Atomicity: The booking is considered an all-or-nothing operation. That is, if payment is not made, then the reservation is not made.

Consistency: The state of the database changes from one legal state to another while still satisfying given constraints (e.g., return date > departure date).

Isolation: The current model prevents concurrent transactions from interfering with each other through mechanisms such as row-locking, which is part of the DBMS.

Durability: Once a tour is booked and confirmed by the system, it is stored in the database even if a crash occurs.

3.4. Indexing and Query Optimization Theory

Lastly, this study covers the theoretical aspects of search optimization. Since travel agencies frequently execute many "Read" operations, or searches for tours based on budget, destination, or date, this study makes use of B-Tree Indexing Theory. The complexity of scanning an entire database is $O(N)$, while using B-Tree reduces complexity to $O(\log N)$. This theoretical approach ensures that the database remains responsive, even when the number of records increases from thousands to millions.

4. METHODOLOGY

The analysis uses an empirical evaluation framework to investigate the effect of model calibration on probabilistic forecasts in the context of a business application. The main focus of the analysis is credit risk modeling, where the aim is to estimate the probability of default of borrowers of loans. The data is from the Give Me Some Credit dataset, which contains information about consumer credit (financial and behavioral attributes) of borrowers. Additionally, it contains binary features of serious delinquency (90+ days past due) over a period of two years from the time of capture. The target variable is called *SeriousDlqn2yrs* and refers to whether a particular borrower has had a serious delinquency.

The Give Me Some Credit dataset is a realistic reflection of the nature of retail credit lending, which is a key area of application of machine learning models. The approach allows the analysis of the comparison of raw and calibrated predictive probabilities, as well as an assessment of the quality of the calibration of predicted probabilities based on performance and quality indicators.

4.1 Data Preparation and Experimental Design

The empirical analysis uses the data from the "Give Me Some Credit" dataset, which uses the binary target variable "SeriousDlqn2yrs" that represents whether the borrower had a serious delinquency within a two-year time period. This variable is used as the default variable to calculate the probability of default.

The data preparation includes a set of preprocessing techniques that need to be performed to maintain the quality and accuracy of the data. The techniques used include checking the data for missing values and replacing the values appropriately, removing duplicates from the original data, and formatting the data uniformly. These techniques help prevent data distortion and ensure that the conclusions drawn from the data are reliable and stable.

To develop a comprehensive evaluation strategy, the data was split using the stratified sampling method, which ensured that the default rate was preserved. The data was split into three different sets: the training set, the calibration set, and the test set. The data was split as follows:

- Training set: 60%
- Calibration set: 20%
- Test set: 20%

The employed structured split also prevents information leakage and allows for unbiased measurement of discrimination and calibration. The final sample sizes are therefore 90,000 for training, 30,000 for calibration, and 30,000 for testing. The default rate of borrowers remained constant at 6.68%, reflecting an equal proportion of classes in the stratified subsets.

For variable selection, an XGBoost model is first implemented on the full set of predictor variables. The feature importance of the fitted model is then employed to select the most important features among the set of predictor variables. The selection is based on the relative contribution of each feature to the model's predictive performance. This approach is a form of dimension reduction, improving interpretability while reducing the risk of overfitting.

4.2 Machine Learning Models

Using the five predictor features identified using the XGBoost feature importance technique, three classification models were developed to compare their performance with regard to discrimination and calibration within a credit risk setting.

The modeling techniques used for this study included:

- **Logistic regression** - A linear model that still represents one of the most popular and enduring techniques for developing credit risk models, largely due to its interpretability and statistical transparency, and its ability to provide direct probability outputs.
- **Random forest** - A bagging ensemble technique that uses multiple decision trees and averages their predictions, proficient in handling nonlinear relationships and interactions.
- **XGBoost** - A gradient-boosting ensemble technique that uses multiple decision trees to produce an ensemble, proficient in handling complex nonlinear relationships and interactions.

The three credit risk models were developed using only training data, without any incorporation of calibration or test data into the model's parameters. The comparative analysis of all three models began with an evaluation of each model's calibration. This entailed checking how well each model's predicted probabilities correlated with actual default rates within the test data prior to any post-hoc calibration. This evaluation provides an independent check on each model's performance with regard to its ability to rank defaults, versus its ability to match predicted defaults with actual defaults. Both discriminative and calibration characteristics are significant factors to be considered for developing credit risk models.

4.3 Calibration Methods

To improve the quality of the predicted probabilities, two post-hoc calibration methods have been used:

- **Platt Scaling** - A parametric method that uses a sigmoid curve to map the original predicted scores to calibrated probabilities using the calibration dataset.
- **Isotonic Regression** - A non-parametric method that uses an isotonic regression to map the original predicted scores to calibrated probabilities.

Both calibration methods have been used independently on the predicted probabilities of the Logistic Regression, Random Forest, and XGBoost models. The calibration parameters are learned using the calibration dataset to avoid any information leak.

The quality of the predicted probabilities obtained using the above methods has been evaluated on the test dataset to check if calibration has improved the quality of the predicted probabilities by mapping them more closely to the actual default rates.

4.4 Calibration Validation and Reliability Assessment

The performance of the calibration process was evaluated by using a combination of graphical evaluation and a statistical test to ensure a comprehensive evaluation of the probability calibration.

The calibration curve is developed by dividing the range of the predicted probabilities into separate intervals, and the mean probability is calculated for each interval. This mean probability is then compared with the actual default rate for each interval.

The Spiegelhalter Z test is used as a statistical evaluation of the overall calibration performance, which measures the difference between the actual default rates and the predicted probabilities on a large scale. The absence of a significant Z score implies that the model has performed well, as it indicates that the probabilities are consistent with the actual default rates. The presence of a significant Z score implies that the model is not calibrated well, indicating that it overestimates or underestimates the default risk for the entire population.

4.5 Evaluation Criteria

There are two major dimensions according to which credit risk models are evaluated:

1. **Discrimination Performance** – This dimension of model evaluation focuses on the ability of the credit risk model to correctly differentiate between defaulters and non-defaulters. This essentially covers the ranking power of the model, i.e., the ability of the model to assign higher probabilities of default to borrowers who eventually default.
2. **Calibration Quality** – This dimension of model evaluation covers the extent to which the probabilities generated by the model reflect actual default rates. A well-calibrated model would be able to provide probabilities that accurately reflect the actual default probabilities for a given population of borrowers, at a segment and aggregate level.

This structured framework for evaluating the models allows for an unambiguous distinction between ranking capability and probability accuracy, which in turn makes it easier to pinpoint the most reliable model-calibration combination for probability estimation in credit risk-related applications.

5. RESULTS

5.1 Discriminatory Performance

Model discrimination was measured using the area under the receiver operating characteristic curve (AUC) for training, calibration, and test sets. The results before calibration are shown in Table 1.

Table 1. Model Discrimination Before Calibration

Model	Train AUC	Calibration AUC	Test AUC
XGBoost	0.8631	0.8553	0.8570
Random Forest	0.8586	0.8565	0.8566
Logistic Regression	0.6698	0.6785	0.6859

The tree-based models show good and stable performance for discrimination, with test AUC values around 0.857 and no significant gap between training and testing performance, which means that there is no overfitting at all. On the other hand, Logistic Regression shows poor discriminatory performance.

Following calibration, the test AUC results are shown in Table 2.

Table 2. Test AUC by Model and Calibration Method

Model	Raw	Platt	Isotonic
XGBoost	0.85698	0.85698	0.85621
Random Forest	0.85665	0.85665	0.85621
Logistic Regression	0.68590	0.68590	0.74155

However, calibration does not have any significant impact on the discriminatory power of XGBoost or RF models, as expected, because calibration affects probability scaling rather than ranking. On the other hand, isotonic calibration has a significant impact on the AUC of the Logistic Regression model, which reflects the correction of the monotonic distortion of the predicted probabilities.

5.2 Global Calibration – Spiegelhalter Test

Global probabilistic calibration was assessed using the original Spiegelhalter test applied to the test sample. The results are presented in Table 3.

Table 3. Spiegelhalter Calibration Test Results

Model	Method	Z-statistic	p-value
XGBoost	Raw	0.7238	0.4692
XGBoost	Platt	1.8757	0.0607
XGBoost	Isotonic	-0.5473	0.5842
Random Forest	Raw	-1.0592	0.2895
Random Forest	Platt	1.7417	0.0816
Random Forest	Isotonic	-0.6212	0.5345
Logistic Regression	Raw	-1.0164	0.3094
Logistic Regression	Platt	-1.2431	0.2138
Logistic Regression	Isotonic	-0.5977	0.5500

The p-values obtained from the tests on all specifications are higher than 0.05, which implies that there is insufficient statistical evidence to reject the null hypothesis of global miscalibration.

For both XGBoost and RF models, the raw models have low absolute values of the test statistic, which implies strong probabilistic consistency. Although isotonic calibration slightly lowers the test statistic values, Platt scaling slightly increases the test statistic values, implying minimal distortions from the two calibration techniques.

Logistic Regression is globally calibrated for all calibration techniques but has weaker discrimination ability than other models.

5.3 Quantile-Based Local Calibration Analysis

To check the consistency of calibration across the entire risk distribution, the models have been evaluated using decile-based binning of the raw PD rankings, where each decile represents approximately 3,000 observations. Summaries of the predicted default rates and observed default rates are presented in Tables 4 to 6.

In Figure 1, we present the reliability diagrams for XGBoost, Random Forest, and Logistic Regression, both in raw and calibrated forms. The dashed line with a positive slope on the diagonal represents perfect calibration, where the model's output probabilities align with actual default rates.

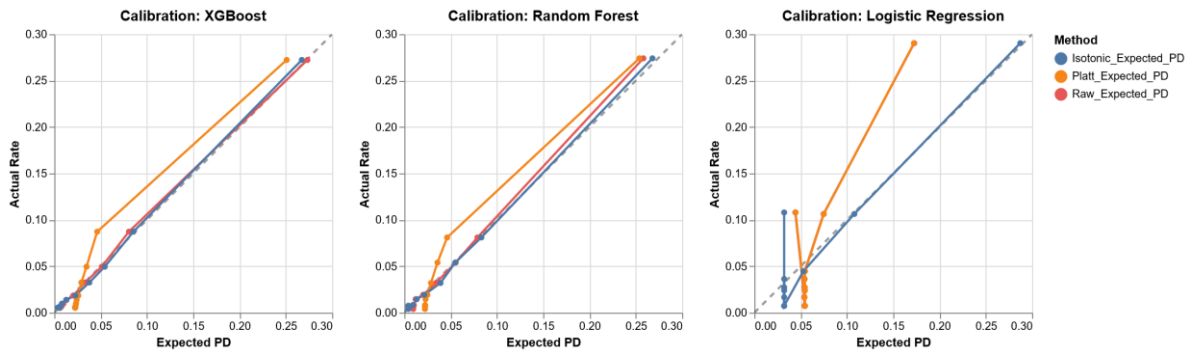


Figure 1. Reliability Diagrams for Raw and Calibrated Models

The raw probabilities provided by XGBoost align closely with the default rates observed across all deciles, including the most risky ones. Isotonic calibration has a minor impact, while Platt scaling compresses the probabilities and overestimates risk at lower deciles while underestimating risk at higher deciles.

Table 4. XGBoost Calibration by Decile

Bin	Raw PD	Platt PD	Isotonic PD	Observed Rate	Accounts
1	0.64%	2.97%	0.50%	0.73%	3000
2	0.84%	3.01%	0.75%	0.70%	3000
3	1.05%	3.06%	0.85%	0.93%	3000
4	1.31%	3.11%	1.09%	1.30%	3000
5	1.73%	3.21%	1.71%	1.83%	3000
6	2.73%	3.45%	3.03%	2.47%	3000
7	4.41%	3.89%	5.04%	4.33%	3000
8	6.83%	4.63%	7.26%	6.60%	3000
9	10.74%	6.18%	11.39%	11.63%	3000
10	36.44%	33.46%	35.63%	36.30%	3000

Random Forest behavior is similar to that of XGBoost, showing a strong correlation between raw predictions and actual rates. Isotonic calibration improves the central (middle) deciles slightly, while Platt scaling again leads to a compression of the risk distribution.

Table 5. Random Forest Calibration by Decile

Bin	Raw PD	Platt PD	Isotonic PD	Observed Rate	Accounts
1	1.28%	2.97%	0.57%	0.53%	3000
2	1.29%	2.97%	0.57%	1.00%	3000
3	1.30%	2.97%	0.77%	0.90%	3000
4	1.36%	2.99%	1.23%	1.10%	3000
5	1.62%	3.05%	1.77%	1.93%	3000
6	2.71%	3.32%	2.88%	2.57%	3000
7	4.55%	3.83%	5.21%	4.27%	3000
8	7.39%	4.77%	7.33%	7.20%	3000
9	10.52%	6.16%	11.09%	10.80%	3000
10	34.42%	33.84%	35.73%	36.53%	3000

The discrimination of logistic regression is low in the lower and middle deciles for the uncalibrated version, which corresponds to its low AUC. Isotonic calibration significantly improves probability scaling for the upper deciles, which explains the increase in the AUC values.

Table 6. Logistic Regression Calibration by Decile

Bin	Raw PD	Platt PD	Isotonic PD	Observed Rate	Accounts
1	4.43%	4.46%	3.22%	10.80%	3000
2	5.39%	5.42%	3.22%	1.67%	3000
3	5.40%	5.43%	3.22%	3.63%	3000
4	5.41%	5.45%	3.22%	0.73%	3000
5	5.41%	5.45%	3.22%	2.73%	3000
6	5.43%	5.46%	3.22%	2.40%	3000
7	5.44%	5.47%	3.22%	0.73%	3000
8	5.44%	5.47%	5.30%	4.47%	3000
9	7.47%	7.50%	10.80%	10.63%	3000
10	17.27%	17.23%	28.72%	29.03%	3000

6. DISCUSSION

The empirical evidence suggests that the most significant determinant of predictive performance is the model class. XGBoost and Random Forest consistently rank as the best discriminators across all datasets, with a consistently high AUC. The difference in AUC between training and testing data is negligible, suggesting that all models generalize well and do not overfit. Finally, Logistic Regression shows that it is not an effective discriminator due to its linear nature, which fails to account for non-linear relationships.

The Spiegelhalter test is used to evaluate calibration performance. There is no evidence of any significant global miscalibration for any model specification. The raw XGBoost and Random Forest models show high levels of probabilistic calibration, suggesting that these optimized models with log-loss optimization will provide properly scaled probabilities.

The analysis of decile calibration plots shows that the raw tree-based models closely follow observed default rates for all risk segments. Isotonic calibration shows some improvement, while Platt probabilistic calibration shows that the probability distribution is compressed, distorting probabilities for all deciles. This suggests that, for these well-calibrated models, any form of parametric calibration may not be necessary and may, in fact, be counterproductive.

In contrast, Logistic Regression shows significant benefits from isotonic calibration, particularly for higher-risk segments. This is due to probability compression inherent to the raw linear model rather than any form of global calibration failure.

7. CONCLUSION

The performance and probability quality of XGBoost, Random Forest, and Logistic Regression are evaluated within this study. A number of statistical measures were used to evaluate the performance of these three algorithms. The tree-based algorithms outperform Logistic Regression in classifying cases, according to all statistical measures used to evaluate these models. Each of these tree-based models shows high and stable values of area under the ROC curve (AUC) when calculated based on training data calibration.

The Spiegelhalter test for global calibration shows that none of these model specifications suffers from miscalibration at a statistical level. Moreover, XGBoost and Random Forest, when used in raw form, show that these two models produce predicted probabilities that closely match observed default rates. The decile analysis shows that these two models produce identical results, showing that both XGBoost and Random Forest are perfectly calibrated for all risk levels on the spectrum, even for higher risk levels.

At the time of calibration for these tree-based models, post-hoc calibration methods such as Platt scaling and isotonic regression show that tree-based models outperform Logistic Regression. There was some improvement for Logistic Regression from using isotonic regression, but this improvement was not significant, and overall performance remained worse than that of tree-based models. Platt scaling, on the other hand, does not improve overall performance but does distort probability estimates. However, this approach to calibration does improve performance for Logistic Regression, and this improvement was significant.

In summary, this study shows that XGBoost and Random Forest can be used as the best modeling approach, and predicted probabilities produced by these two models already show good calibration. There is no need to add additional levels of calibration for these tree-based models, and such additional levels should be used on an ad hoc rather than routine basis.

REFERENCES

1. Platt, J. C. (1999). Probabilistic Outputs for Support Vector Machines and Comparisons to Regularized Likelihood Methods. *Advances in Large Margin Classifiers*, MIT Press.
2. Zadrozny, B., & Elkan, C. (2002). Transforming Classifier Scores into Accurate Multiclass Probability Estimates. *Proceedings of the Eighth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '02)*.
3. Niculescu-Mizil, A., & Caruana, R. (2005). Predicting Good Probabilities with Supervised Learning. *Proceedings of the 22nd International Conference on Machine Learning (ICML '05)*.
4. Lin, H. T., Lin, C. J., & Weng, R. C. (2007). A Note on Platt's Probabilistic Outputs for Support Vector Machines. *Machine Learning*, 68(3), 267–276.
5. Kull, M., Silva Filho, T. M., & Flach, P. (2017). Beta Calibration: A Well-Founded and Easily Implemented Improvement on Logistic Calibration for Binary Classifiers. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics (AISTATS 2017)*.
6. Guo, C., Pleiss, G., Sun, Y., & Weinberger, K. Q. (2017). On Calibration of Modern Neural Networks. *Proceedings of the 34th International Conference on Machine Learning (ICML 2017)*.
7. Bella, A., Ferri, C., Hernández-Orallo, J., & Ramírez-Quintana, M. J. (2010). Calibration of Machine Learning Models. *Handbook of Research on Machine Learning Applications and Trends*.
8. Wallace, B. C., & Dahabreh, I. J. (2012). Class Probability Estimates are Unreliable for Imbalanced Data (and How to Fix Them). *2012 IEEE 12th International Conference on Data Mining*.
9. Tasche, D. (2013). The Art of Probability-of-Default Curve Calibration. *Journal of Credit Risk*, 9(4), 63–103.
10. Baesens, B., Van Gestel, T., Viaene, S., Stepanova, M., Suykens, J., & Vanthienen, J. (2003). Benchmarking State-of-the-Art Classification Algorithms for Credit Scoring. *Journal of the Operational Research Society*, 54(6), 627–635.
11. Basel Committee on Banking Supervision. (2005). *Studies on the Validation of Internal Rating Systems* (Working Paper No. 14). Bank for International Settlements.
12. Spiegelhalter, D. J. (1986). Probabilistic Prediction in Patient Management and Clinical Trials. *Statistics in Medicine*, 5(5), 421–433.

INSTRUCTIONS FOR AUTHORS

1. "The Baku Engineering University Mathematics and Computer Science" accepts original unpublished articles and reviews in the research field of the author.
2. Articles are accepted in English.
3. File format should be compatible with **Microsoft Word** and must be sent to the electronic mail (journal@beu.edu.az) of the Journal. The submitted article should follow the following format:
 - Article title, author's name and surname
 - The name of workplace
 - Mail address
 - Abstract and key words
4. The title of the article should be in each of the three languages of the abstract and should be centred on the page and in bold capitals before each summary.
5. **The abstract** should be written in **9 point** type size, between **100** and **150** words. The abstract should be written in the language of the text and in two more languages given above. The abstracts of the article written in each of the three languages should correspond to one another. The keywords should be written in two more languages besides the language of the article and should be at least three words.
6. **.UDC** and **PACS** index should be used in the article.
7. The article must consist of the followings:
 - Introduction
 - Research method and research
 - Discussion of research method and its results
 - In case the reference is in Russian it must be given in the Latin alphabet with the original language shown in brackets.
8. **Figures, pictures, graphics and tables** must be of publishing quality and inside the text. Figures, pictures and graphics should be captioned underneath, tables should be captioned above.
9. **References** should be given in square brackets in the text and listed according to the order inside the text at the end of the article. In order to cite the same reference twice or more, the appropriate pages should be given while keeping the numerical order. For example: [7, p.15].

Information about each of the given references should be full, clear and accurate. The bibliographic description of the reference should be cited according to its type (monograph, textbook, scientific research paper and etc.) While citing to scientific research articles, materials of symposiums, conferences and other popular scientific events, the name of the article, lecture or paper should be given.

Samples:

- a) **Article:** Demukhamedova S.D., Aliyeva İ.N., Godjayev N.M.. *Spatial and electronic structure of monomer and dimeric conapeetes of carnosine with zinc*, Journal of structural Chemistry, Vol.51, No.5, p.824-832, 2010
 - b) **Book:** Christie John Geankoplis. *Transport Processes and Separation Process Principles*. Fourth Edition, Prentice Hall, p.386-398, 2002
 - c) **Conference paper:** Sadychov F.S., Aydın C., Ahmedov A.İ.. *Application of Information – Communication Technologies in Science and education. II International Conference. "Higher Twist Effects In Photon- Proton Collisions"*, Baki, 01-03 Noyabr, 2007, ss 384-391
- References should be in 9-point type size.
10. The margins sizes of the page: - Top 2.8 cm. bottom 2.8 cm. left 2.5 cm, right 2.5 cm. The article main text should be written in Palatino Linotype 11 point type size single-spaced. Paragraph spacing should be 6 point.
 11. The maximum number of pages for an article should not exceed 15 pages
 12. The decision to publish a given article is made through the following procedures:
 - The article is sent to at least to experts.
 - The article is sent back to the author to make amendments upon the recommendations of referees.
 - After author makes amendments upon the recommendations of referees the article can be sent for the publication by the Editorial Board of the journal.

