

LB/TH/15/2026
TH6170

**NOVEL RESOURCE ALLOCATION TECHNIQUE
FOR DEMAND AWARE 5G EDGE CLOUD
DIMENSIONING**

Yogendrarajah Sivanujan

199422B

Master of Science

Department of Electronics and Telecommunication Engineering
Faculty of Engineering

University of Moratuwa
Sri Lanka

January 2025

**NOVEL RESOURCE ALLOCATION TECHNIQUE
FOR DEMAND AWARE 5G EDGE CLOUD
DIMENSIONING**

Yogendrarajah Sivanujan

199422B

Dissertation submitted in partial fulfillment of the requirements for the
degree
Master of Science

Department of Electronics and Telecommunication Engineering
Faculty of Engineering

University of Moratuwa
Sri Lanka

January 2025

DECLARATION

I declare that this is my own work and this Dissertation does not incorporate without acknowledgement any material previously submitted for a Degree or Diploma in any other University or Institute of higher learning and to the best of my knowledge and belief it does not contain any material previously published or written by another person except where the acknowledgement is made in the text. I retain the right to use this content in whole or part in future works (such as articles or books).

Signature:

Date: 14/03/2025

The supervisors should certify the Dissertation with the following declaration.

The above candidate has carried out research for the Master of Science Dissertation under our supervision. We confirm that the declaration made above by the student is true and correct.

Name of Supervisor: Dr. Abhaya Sumanasena

Signature of the Supervisor:

Date: 16/03/2025

Name of Supervisor: Dr. Kasun Hemachandra

Signature of the Supervisor:

Date: 17/03/2025

ACKNOWLEDGEMENT

I wish to express my sincere gratitude to my supervisor Dr. Abhaya Sumanasena and my Co-Supervisor Dr. Kasun Hemachandra for spending time guiding and mentoring me during my research. Without their immense support and careful guidance, I may not have achieved my target. Secondly, I would like to express heartfelt gratitude to the teaching staff of the Department of Electronics and Telecommunication Engineering in the University of Moratuwa to enlightening and making me more knowledgeable. Lastly, to my parents – who supported me in achieving my goals. Without them, I may not in the place where I am right now.

ABSTRACT

The rapid proliferation of 5G networks has led to increased demand for efficient resource allocation at the edge cloud to support latency-sensitive and compute-intensive applications. This thesis presents an optimal resource allocation framework for 5G edge cloud environments, leveraging a closed-loop architecture that integrates demand prediction, intelligent scheduling, and load balancing. The proposed framework employs machine learning-based traffic prediction models, such as Prophet, to forecast real-time demand and optimize resource provisioning dynamically. The framework consists of key components, including a central cloud resource orchestrator, edge cloud nodes, a traffic prediction module, and a resource orchestrator for cluster-wide task scheduling and load balancing. A feedback mechanism ensures continuous updates on resource utilization, enabling adaptive scaling and efficient workload distribution. Additionally, an optimization engine is incorporated to balance energy efficiency and privacy constraints while maintaining service quality. This research is implemented and evaluated using CloudSim Plus, a simulation framework for cloud computing environments and NS3, a discrete-event network simulator for internet systems. The effectiveness of the proposed approach is assessed based on metrics such as resource utilization, energy consumption, and service latency. Experimental results demonstrate that the closed-loop framework enhances the adaptability of 5G edge clouds, reducing underutilization and overprovisioning while improving overall system efficiency. The findings of this study contribute to the advancement of intelligent resource management in 5G edge computing, paving the way for scalable and autonomous cloud-native network operations. Future work may explore the integration of reinforcement learning-based optimization techniques to further enhance decision-making in dynamic edge cloud environments.

Keywords: Edge cloud, Resource orchestration, Complex resource optimization

TABLE OF CONTENTS

Declaration of the Candidate & Supervisor	i
Acknowledgement	ii
Abstract	iii
Table of Contents	iv
1 Introduction	1
1.1 Background of the Study	1
1.1.1 The Role of 5G in Modern Networking	1
1.1.2 Edge Cloud Computing: A new era of computing	1
1.1.3 Demand Aware Resource Allocation	2
1.1.4 Significance of the Study	2
1.1.5 Research Context	2
1.2 Problem Statement	3
1.2.1 Challenges of Resource Allocation in 5G Edge Clouds	3
1.2.2 Energy Efficiency and Privacy Concerns	3
1.2.3 Dynamic and Demand Aware Frameworks	4
1.2.4 Research Gap	4
1.2.5 Problem Statement Summary	4
1.3 Research Objectives	5
1.4 Scope of the Study	6
1.4.1 Technical Scope	6
1.4.2 Application Scope	6
1.4.3 Geographical Scope	7
1.4.4 Temporal Scope	7
1.4.5 Limitations	7
1.5 Significance of the Study	7
1.5.1 Theoretical Significance	8
1.5.2 Practical Significance	8

1.5.3	Industrial Relevance	8
1.5.4	Broader Implications	9
1.6	Organization of the Report	9
1.6.1	Chapter 1: Introduction	10
1.6.2	Chapter 2: Literature Review	10
1.6.3	Chapter 3: Methodology	10
1.6.4	Chapter 4: Simulation Data Analysis	10
1.6.5	Chapter 5: Results and Discussion	10
1.6.6	Chapter 6: Conclusion and Recommendations	11
2	LITERATURE REVIEW	12
2.1	Introduction	12
2.2	Resource Allocation Challenges in Edge Cloud Networks	13
2.2.1	Dynamic User Demands and Traffic Variability	13
2.2.2	Latency Constraints	13
2.2.3	Energy Efficiency	13
2.2.4	Resource Heterogeneity	14
2.2.5	Scalability	14
2.2.6	Security and Privacy	14
2.2.7	Task Scheduling and Load Balancing	15
2.2.8	Interoperability and Standardization	15
2.2.9	Real-Time Decision Making	15
2.3	Demand-Aware Resource Allocation Techniques	15
2.3.1	Traffic Prediction Models	16
2.3.2	Adaptive Task Scheduling	16
2.3.3	Load Balancing	17
2.3.4	Resource Orchestration	17
2.3.5	Energy-Aware Allocation	17
2.3.6	Privacy-Aware Resource Allocation	18
2.3.7	Hybrid Approaches	18
2.3.8	Challenges in Demand-Aware Techniques	18
2.3.9	Summary	19

2.4	Gaps in Existing Approaches	19
2.4.1	Lack of Real-Time Adaptability	19
2.4.2	Limited Integration of Machine Learning Models	20
2.4.3	Inadequate Focus on Energy Efficiency	20
2.4.4	Fragmentation in Resource Management	20
2.4.5	Challenges in Privacy and Security	20
2.4.6	Scalability Issues	21
2.4.7	Interoperability Constraints	21
2.4.8	Overlooked Application-Specific Requirements	21
2.4.9	Challenges in Multi-Objective Optimization	21
2.5	Theoretical Framework and Models	22
2.5.1	Queuing Theory for Resource Allocation	22
2.5.2	Game Theory for Resource Optimization	22
2.5.3	Multi-Objective Optimization Models	23
2.5.4	Machine Learning-Based Models	23
2.5.5	Network Slicing Models	24
2.5.6	Energy-Aware Models	24
2.5.7	Summary of the Theoretical Framework	25
2.6	Summary of Literature Insights	25
2.6.1	Progress in Resource Allocation Techniques	25
2.6.2	Identified Gaps in Literature	26
2.6.3	Emerging Trends and Future Directions	27
2.6.4	Contributions of the Literature Review	27
2.6.5	Summary	28
3	METHODOLOGY	29
3.1	Introduction	29
3.2	Research Design	29
3.2.1	Research Philosophy	29
3.2.2	Research Approach	30
3.2.3	Research Strategy	30
3.3	Framework Architecture	30

3.3.1	Overview of the Framework Architecture	30
3.3.2	Key Components of the Framework	31
3.3.3	Workflow of the Framework	32
3.3.4	Advantages of the Framework	32
3.3.5	Diagram of the Framework	33
3.3.6	Summary	33
3.4	Mathematical Model for Demand Forecasting using Prophet	33
3.4.1	Forecasting Model	33
3.4.2	Trend Component	34
3.4.3	Seasonality Component	34
3.4.4	Holiday Component	35
3.4.5	Uncertainty Interval	35
3.4.6	Final Forecast Equation	35
3.5	Mathematical Model of Genetic Algorithm for Edge Cloud Resource Allocation and Optimization	35
3.5.1	Problem Definition	35
3.5.2	Fitness Function	36
3.5.3	Genetic Algorithm Process	36
3.6	Mathematical Model of the Genetic Algorithm and Static based Scheduling	37
3.6.1	Variables	37
3.6.2	Objective Function (Genetic Algorithm)	38
3.6.3	Constraints	38
3.6.4	Genetic Algorithm Details	38
3.6.5	Static Scheduling	39
3.6.6	Conclusion	39
3.7	Comparison Between Baseline and Proposed Frameworks	39
3.7.1	Resource Provisioning Methodology	39
3.7.2	Resource Orchestration	39
3.7.3	Scaling Methodology	40
3.7.4	Traffic Prediction Mechanism	40

3.7.5	Task Scheduling Algorithms	40
3.7.6	Resource Utilization within Edge Nodes	41
3.7.7	Resource Utilization across Edge Locations	41
3.7.8	Summary of Key Benefits	41
4	SIMULATION AND DATA COLLECTION	42
4.1	Simulation Environment	42
4.2	Key System Parameters	42
4.3	Evaluation Metrics	44
4.4	Simulation Scenarios	44
4.5	Data Collection Methods	44
4.6	Pseudocode of the simulation of Traffic Demand Forecast using Prophet	45
4.7	Pseudocode of the simulation of Genetic Algorithm based resource allocation optimization	47
4.8	Pseudocode of the simulation of Baseline and Genetic Algorithm based task scheduling	49
4.9	NS-3 simulation	53
4.10	Summary	53
5	RESULTS AND DISCUSSION	54
5.1	Introduction	54
5.2	Performance Metrics and Simulation Results	54
5.2.1	Latency Analysis	54
5.2.2	Throughput and Bandwidth Utilization Analysis	55
5.2.3	Simulation Results of Traffic Demand Forecast using Prophet	56
5.2.4	Simulation results of Genetic Algorithm based resource allocation optimization and the fitness values	56
5.2.5	Simulation results of Round robin and Genetic Algorithm based task scheduling and execution	59
5.3	Comparative Analysis with Existing Techniques	62
5.4	Impact of Algorithm on Resource Allocation Efficiency	62
5.5	Practical Implications for 5G Edge Networks	63

6	CONCLUSION AND RECOMMENDATIONS	64
6.1	Summary of Key Findings	64
6.1.1	Improved Latency	64
6.1.2	Higher Throughput and Bandwidth Utilization	64
6.2	Summary of Baseline vs Proposed Frameworks	65
6.3	Recommendations for Industry Practitioners	65
6.3.1	Invest in Training and Skill Development	66
6.4	Future Research Directions	66
6.4.1	Integration with Emerging Technologies	66
6.4.2	Expansion to Beyond 5G and 6G Networks	66
6.4.3	Policy and Regulatory Considerations	67
.1	Appendix A	72
.2	Appendix B	73
.3	Appendix C	76
.4	Appendix D	81

CHAPTER 1

INTRODUCTION

1.1 Background of the Study

The emerging 5G technology has revolutionized wireless communication. It is setting a new era in network connectivity through ultra reliable, low latency, and high speed data transmission. As industries integrate technologies such as augmented reality (AR), autonomous systems, and the Internet of Things (IoT), the demand for the real time data processing and low latency communication have increased exponentially (Zhang et al., 2020). Traditional centralized cloud computing or data center architectures struggle a lot to meet the demands, promote a shift towards edge cloud computing.

1.1.1 The Role of 5G in Modern Networking

The 5G technology represents a significant evolution from its original architecture by incorporating advanced features such as enhanced mobile broadband (eMBB), ultra reliable low latency communication (URLLC), and massive machine type communication (mMTC) (Ahmed et al., 2022). These features enable 5G networks to support diversified applications, including real time analytics, industrial IoT, and autonomous vehicles etc. However, as the number of connected devices increases, the need for the efficient resource allocation mechanisms is crucial to meet the dynamic demands.

1.1.2 Edge Cloud Computing: A new era of computing

Edge computing concept decentralizes the computational power by placing the resources closer to end users while addressing the latency and bandwidth requirements of traditional cloud computing (Wang et al., 2021). This approach integrates edge devices, edge nodes, and the centralized cloud servers for creating a distributed computing architecture. Edge computing reduces the load on centralized data centers and enhances the responsiveness of latency sensitive applications by processing data locally (Chen et al., 2021).

For example, real time applications such as AR and industrial automation get the benefits significantly since the proximity of computational resources in edge networks. The above applications demand not only low latency communication but also real time task scheduling and resource optimization to make sure seamless operations (Nguyen et al., 2021). The decentralization nature provided by the edge cloud computing plays a prominent role in meeting these strict demands. It is a cornerstone of modern 5G networks.

1.1.3 Demand Aware Resource Allocation

Demand aware resource allocation is becoming as a critical strategy to optimizing the performance of 5G edge cloud networks. Unlike traditional methods which are rely on static allocation policies, demand aware frameworks adjust resource allocation dynamically based on real time traffic patterns, user demands, and application priorities (Rahman et al., 2023). This approach ensures that resources are utilized more efficiently meanwhile maintain the high levels of QoS.

As an example, machine learning based traffic prediction models can forecast the user demands and allocate resources proactively. In addition, it reduces the bottlenecks and improving system efficiency. Similarly, the privacy aware routing protocols can improve the data security while ensuring efficient resource utilization (Zhou et al., 2022). By combining such advanced routing techniques, demand aware resource allocation frameworks address the diversified challenges of 5G edge cloud computing.

1.1.4 Significance of the Study

The integration of demand aware resource allocation frameworks into 5G edge cloud networks has significant contributions for various industries. For example, in health-care, real time monitoring of patient data requires low latency communication and secure data transmission. The above demand can be achieved through optimized resource allocation techniques (Liu et al., 2020). In industrial IoT, effective resource allocation improves operational efficiency and reduce energy consumption and contributing to the sustainable development goals.

Apart from that, as the adoption of 5G technology accelerates the importance of scalable and efficient resource allocation frameworks becomes even more critical evidently. Addressing the challenges of demand aware resource allocation is essential not only for improving system performance but also for ensuring the long term sustainability of 5G-enabled edge cloud system (Ahmed et al., 2022).

1.1.5 Research Context

This study tries to build on existing research in 5G edge cloud computing by focusing on the development of a demand aware resource allocation framework by integrating the advanced algorithms such as genetic algorithms and branch and bound techniques. The proposed framework targets to optimize the resource utilization, reduce latency, and improved QoS. The findings of this research will contribute to the growing body of knowledge on 5G edge cloud computing. Further, it provides actionable insights for industry practitioners and researchers.

1.2 Problem Statement

The rapid development and deployment of 5G networks have introduced tremendous opportunities in the telecommunications industry, including ultra reliable low latency communications (URLLC), enhanced mobile broadband (eMBB), as well as massive machine type communications (mMTC). However, these advancements have also bring in severe complex challenges, particularly in the area of resource allocation in edge the cloud environment. The dynamic and heterogeneous nature of user demands, patterns, coupled with the diverse application specific requirements which create significant bottlenecks in maintaining optimal performance of the system (Zhou, Li, Yang, 2022).

1.2.1 Challenges of Resource Allocation in 5G Edge Clouds

Resource allocation in 5G edge cloud networks is fundamentally different from traditional centralized cloud environments due to the decentralization of computational resources. Edge nodes, which are geographically distributed, face dynamic workloads influenced by real-time traffic variations, application requirements, and network conditions. Existing resource allocation strategies often employ static or semi-dynamic frameworks, which fail to adapt to these real-time changes, resulting in resource underutilization, increased latency, and degraded Quality of Service (QoS) (Chen, Zhang, Huang, 2021).

Moreover, as the number of connected devices continues to grow exponentially over the last decade, edge networks are under constant severe pressure to provide scalable and efficient solutions. The lack of adaptability in traditional resource allocation methods lead the above phenomena since they are unable to respond to sudden spikes in demand or even varying application requirements. This creates inefficiencies in resource utilization and impact the user experience directly as well as the operational efficiency of edge networks (Kumar, Gupta, Rahman, 2022).

1.2.2 Energy Efficiency and Privacy Concerns

Energy efficiency is one of the most critical issues in resource allocation within 5G edge clouds. As edge devices and nodes operate under strict power constraint environments, ensuring the optimal energy utilization rather compromising service quality. It becomes a significant challenge in the 5G edge Eco system. Existing frameworks often overlook energy consumption during resource allocation which leads to higher operational costs and reduced sustainability (Ahmed, Khan, Wang, 2022).

In addition, the integration of diversified applications such as real time healthcare

monitoring, autonomous vehicles, and smart city infrastructure require resilient data privacy mechanisms. However, existing resource allocation strategies lacking behind comprehensive solutions to ensure the privacy concerns, expose users to potential security vulnerabilities. This gap highlights openly the need for privacy aware resource allocation frameworks that can secure sensitive data while maintaining performance efficiency (Nguyen, Rahman, Li, 2021).

1.2.3 Dynamic and Demand Aware Frameworks

Demand aware resource allocation has emerged as a potential and unavoidable solution to the challenges faced by 5G edge cloud environments. This approach tries adjust the resource allocation dynamically based on real time traffic patterns, user demands, and application requirements. Despite its potential, current demand aware frameworks are often constrained by limited computational efficiency and scalability and particularly in large scale 5G network deployments (Huang, Liu, Zhao, 2023).

Existing studies try to highlight the limitations of conventional algorithms, such as first come first serve(FCFS) and heuristic methods which struggle a lot to balance the trade offs between computational efficiency, energy consumption, and QoS. This nature evidently shows that the need for innovative hybrid approaches that can address these limitations while ensuring adaptability and scalability in resource allocation for 5G system (Liu, Rahman, Zhao, 2023).

1.2.4 Research Gap

Prior work focused on the different components of resource management, such as for example, reducing latency, energy efficiency, or privacy. Further, few research studies presented complete models that put together these views. Moreover, the lack of an actual use case and validating it in the real 5G edge cloud environment causes the existing approaches not to be so practical. This paper is set to bridge this gap by proposing a holistic, demand-aware resource allocation framework, which will embed advanced algorithms, real-time traffic prediction, and privacy-aware mechanisms, thus, improving system performance and scalability.

1.2.5 Problem Statement Summary

5G edge cloud environments, with their dynamic, heterogeneous, and resource-intensive nature, are indeed very challenging for resource allocation efficiency. Modern frameworks are powerless to solve the problems related to power consumption and data privacy and performance is compromised. Thus, it will demand the adaptability and scalability for energy efficiency and data privacy along with the integration of different

parts. This can only be achieved by recourse to innovative solutions. No doubt, however, pulling together advanced algorithms based on edge learning and forecasting of real-time demand for the introduction of QoS and the reduction of delay and increase in resource utilization in 5G-enabled edge cloud networks will be an excellent solution.

1.3 Research Objectives

This study intends to constitute an exclusive and on-demand resource allocation scheme fitted to 5G edge cloud environments. The main goal is to encounter the basic obstacles the user of mobile equipment, network variation, and ecological conditions bring while asking for services. By applying novel techniques in cloud resource allocation by new algorithms and optimization techniques. The main goal of this research is to optimize the cloud resource allocation and improve the performance of throughput and latency and other QoS parameters.

The following specific objectives guide this study:

1. **To analyze the key challenges in resource allocation within 5G edge cloud networks:** Resource allocation in 5G edge clouds faces many challenges such as latency, energy inefficiency, and resource under utilization. This study will identify these burning challenges by reviewing existing literature and analyzing real-world scenarios (Zhou et al., 2022; Chen et al., 2021).
2. **To evaluate the impact of dynamic user demands and network heterogeneity on resource allocation frameworks.:** The study will investigate more on how varying user demands and diverse network conditions influence the effectiveness of resource allocation strategies. Further it emphasize the need for adaptability in these frameworks (Ahmed et al., 2022; Huang et al., 2023).
3. **To design and implement a hybrid algorithm for demand-aware resource allocation:** The framework integrates advanced algorithms such as genetic algorithms, branch-and-bound methods, and traffic prediction models to optimize resource allocation in real time scenarios (Kumar et al., 2022; Nguyen et al., 2021).
4. **To assess the performance of the proposed framework through simulations and quantitative metrics:** Key performance indicators such as latency, throughput, energy efficiency, and QoS will be measured using simulation tools to validate the effectiveness of the framework resource nature (Liu et al., 2020; Rahman et al., 2023).

5. **To propose actionable recommendations for implementing demand-aware resource allocation frameworks in 5G edge cloud networks.:** Practical guidelines will be developed to help for network operators and developers to optimize resource utilization and enhance service quality in diverse 5G environments (Zhou et al., 2022; Nguyen et al., 2021).

1.4 Scope of the Study

This study focuses on the development of a demand aware resource allocation framework for 5G edge cloud environments. With the increasing adoption of 5G networks, edge computing technology is unavoidable for latency sensitive applications, real time analytics, and resource intensive tasks such as augmented reality (AR), autonomous vehicles, and industrial IoT (Chen, Zhang, Huang, 2021). The research addresses the limitations of existing resource allocation strategies and techniques and proposes a novel framework to optimize resource utilization while maintaining high Quality of Service (QoS).

1.4.1 Technical Scope

The research emphasizes the following technical dimensions:

1. **Resource Allocation Challenges:** The study identifies key challenges such as dynamic user demand, energy efficiency, as well as privacy concerns in 5G edge clouds. Existing frameworks always fail to address these factors comprehensively. It results in resource under utilization and increased latency (Zhou, Li, Yang, 2022).
2. **Performance Evaluation:** The framework will be validated using simulation tools, measuring key performance indicators such as latency, throughput, energy efficiency, and QoS. Comparative analysis with existing methods will also be conducted (Ahmed, Khan, Wang, 2022).

1.4.2 Application Scope

The proposed framework is applicable to diverse 5G-enabled domains, including:

1. **Smart Cities:** For real time traffic management, energy optimization, and public safety systems.
2. **Healthcare:** For medicine, remote monitoring, and emergency response systems requiring low latency and secure communications.

3. **Industrial IoT:** For manufacturing automation, predictive maintenance, and supply chain management.
4. **Autonomous Vehicles:** For vehicular communication systems require ultra low latency and high reliability connectivity.

1.4.3 Geographical Scope

The research that was carried out especially concentrates on the 5G coverage distance being applied to the edge cloud which is located in the urban area and also in the semi-urban area. But the details can be tailored to rural places if there are engineering changes the infrastructure and resources as well.

1.4.4 Temporal Scope

For the most part, the study of the current and future requirements of 5G edge cloud environments is on the table. In general, the proposed framework is designed to manage 5G networks' growing complexity and to back future advancements in 6G and so on (Huang, Liu, Zhao, 2023).

1.4.5 Limitations

While the study aims to address critical challenges in resource allocation, certain limitations include:

1. **Simulation-Based Validation:** The framework's performance will be tested in simulated environments, which may not fully replicate real world complexities.
2. **Infrastructure Constraints:** The framework assumes the presence of advanced edge computing infrastructure, which may not be universally available.

The report of this study on a comprehensive framework for enhancing the efficiency and scalability of 5G edge cloud networks by addressing these limitations and focusing on demand-aware optimization is leading the way in those spectrum of technologies.

1.5 Significance of the Study

All businesses are adopting capacity-aware resource allocation frameworks in 5G edge cloud networks to cope with the issues of scalability, efficiency, and real-time response in the modern communication systems. This paper is an important step in the areas of telecommunications, cloud computing, and resource optimization which are all currently shifting to more decentralized and latency-sensitive applications.

1.5.1 Theoretical Significance

1. **Advancing Knowledge in 5G Edge Computing:** This research contributes to the theoretical understanding of resource allocation challenges in 5G edge cloud environments by exploring dynamic demand-aware frameworks. It identifies gaps in existing methodologies and introduces a hybrid approach integrating genetic algorithms, branch-and-bound techniques, and machine learning-based traffic prediction models (Chen et al., 2021; Zhou et al., 2022).
2. **Incorporating Multidimensional Optimization:** The study investigates different key dimensions like reduction of latency, energy efficiency, and privacy. The authors have put forward a comprehensive hypostasis to act as a base for future research in resource management that incorporates the nexus of these factors (Nguyen et al., 2021; Kumar et al., 2022).

1.5.2 Practical Significance

1. **Enhancing 5G Network Efficiency:** The advanced algorithm is a piece of technology that is developed to make use of resources in an efficient manner by reducing the time slots and latency that are required for the successful data transfer. One of the main purposes of the framework under the proposed plan is to make it possible for those latency-sensitive applications which include the telemedicine, autonomous vehicle, and industrial IoT task to have the best-in-class quality with lower latency under the proposed framework. This improvement directly affects end-users and network providers (Ahmed et al., 2022).
2. **Promoting Sustainability:** The research focuses on the energy-efficient allocation of resources that are aimed at taking the responsibility of the environmental impact from large-scale network deployments in a world we see endeavor to make a greener environment. By curbing energy dissipation in edge nodes, the study partakes in international action on sustainability (Huang et al., 2023).
3. **Strengthening Data Privacy and Security:** The improvement of privacy-aware routing protocols in a framework enhances data security for the applications that manage sensitive information, like healthcare and financial services. This enhancement makes it possible that the network operators can offer services that are both secure and of high quality (Nguyen et al., 2021).

1.5.3 Industrial Relevance

1. **Guiding 5G Network Operators and Developers:** The profitable recommendations provided by this study will be of great help to network operators as

they even more deploy demand-aware frameworks that enhance operational efficiency and reliability of services. Programmers are able to use the hybrid algorithm to devise strategies of scalable resource allocation to different application scenarios (Zhou et al., 2022).

2. **Facilitating Innovation in Edge Computing Applications:** This research's outcomes can be adopted by industries like smart cities, connected transportation, and digital healthcare in their bid to resource management optimizations the real-world implementation in order to gain an edge in the technological race. Such adoption will go one step further bringing innovation that will improve the performance of the systems that got edge-computing-enabled (Kumar et al., 2022).

1.5.4 Broader Implications

1. **Impact on Future Network Architectures:** This study's implications are not just about 5G networks but can also be applied to the future's technologies such as 6G and above. The offered framework acts as an outline for handling future network architecture's resource allocation problems (Huang et al., 2023).
2. **Contributing to Digital Transformation:** Industries are making their digital transformation journeys and technological advancements are well correlated, the study's findings would help in the transition from the development of technologies to the practical application in various sectors. Such an effort is indispensable for the utilization of the edge computing concept to the full in a 5G-based network (Ahmed et al., 2022).

The main importance of this research is that it is able to handle the main issues 5G edge cloud networks are experiencing while providing new solutions to reliability, energy efficiency, and scalability which significantly improve the QoS. The various benefits the research brings to establishing theory, practice, and industry are a clear indicator of relevance to the telecom sector and resource management.

1.6 Organization of the Report

This research report is broken down into multiple chapters focusing on different facets of the research so that the user can have a comprehensive understanding of how resources are allocated in a dynamically changing 5G edge cloud environment. The forthcoming is a summary of the organization:

1.6.1 Chapter 1: Introduction

This section, which has the theme of research and establishes the basis for the study, is the first chapter. The background of the study, problem statement, research objectives, scope, and research significance are included in the paragraph. This chapter is directed to necessitate the demand-aware resource allocation framework and to present the relevance of 5G edge cloud networks to the environment.

1.6.2 Chapter 2: Literature Review

This chapter is a comprehensive review of current studies and frameworks associated with the provisioning of resources in an edge cloud environment. It delves into various milestones in the 5th-generation mobile communication technology such as dynamic resource allocation difficulties and issues in current methodologies. The literature review recognizes the spaces where this study brings to the existence of the academic and practical knowledge base.

1.6.3 Chapter 3: Methodology

The methodology chapter outlines the research design, framework architecture, and algorithm development. It details the hybrid approach adopted for resource allocation, integrating genetic algorithms, branch-and-bound methods, and machine learning-based traffic prediction models. The chapter also describes the simulation setup, data collection methods, and performance evaluation metrics.

1.6.4 Chapter 4: Simulation Data Analysis

This chapter presents the detailed simulation and implementation of the proposed methodology. It includes step-by-step explanations of task scheduling, resource allocation, demonstrating its application in optimizing resource utilization in 5G edge cloud networks.

1.6.5 Chapter 5: Results and Discussion

Achievement of the study mainly depends on the results and discussion chapter. It consists of the analysis of the simulation results which are based on the indicators of the main type such as latency, throughput, energy efficiency, and Quality of Service (QoS). The chapter compares the proposed framework with existing approaches, and provides a detailed discussion of the findings.

1.6.6 Chapter 6: Conclusion and Recommendations

This chapter deals with the obtained results and highlights their significance in 5G edge cloud computing. It talks about the effect of the research on the professionals in the field as well as on the academic researchers. The particular chapters along with the necessary steps will be provided to the contractor for the application of the demand aware resource accumulation frameworks and the identification of the future scopes for research.

CHAPTER 2

LITERATURE REVIEW

2.1 Introduction

The accelerating expansion of 5G networks results in a transformation of the telecommunications and edge computing sphere. 5G is the main enabler for advanced applications such as augmented reality, self-driving cars, and smart city systems. And as the technology is becoming the essential component in the process of the digital transformation, the necessity to provide efficient resource allocation in edge cloud environments has consequently emerged. The 5G-induced innovations are not only social but also economic and business supporting the traditional concepts. That is the reason why the global efficiency of T-DCN is helpless to the severe economics of diverse mobile carriers. The dynamic behavior of 5G-enabled applications making the user to satisfy variability in user demands and network types lead to finding the way for energy-efficient computing while maintaining the basic requirements, such as low latency and high reliability with the utility of computational resources (Meenakshi Liam, 2014).

Edge clouds' resource allocation has stood out as one of the most studied areas of concern, with most studies concentrating on tools that improve system performance, scalability, and QoS. However, in stark contrast, the existing frameworks too often come up empty in the face of the dynamism and multidimensionality of the challenges of demand-aware networks. As such, the static or semi-dynamic allocation strategies are unfit for real-time applications that require fast adaptations to match the different user demands (Zhou et al., 2022). Furthermore, the combination of privacy-aware mechanisms and energy-efficient solutions is still a less popular topic of research in the 5G edge computing context (Nguyen et al., 2021).

This chapter focuses on the current methodologies employed in this study and the gaps and the challenges in the existing approaches. Further, it discusses about resource allocation challenges in Edge Cloud Networks and various demand-aware resource allocation techniques.

In conclusion, this chapter is a very comprehensive overview of current studies of theoretical frameworks as well as resource allocation in edge cloud environments. It takes into account chief progresses in 5G technology, difficulties in dynamic resource allocation, and deficiencies in current methodologies. The literature review clearly demonstrates the areas where the study can help the academic and practical sides of the area.

2.2 Resource Allocation Challenges in Edge Cloud Networks

Resource Allocation on the edge cloud network is a dynamic and significant process having influenced the special requirements of the 5G environment. The Edge cloud networks are operating somewhere close to the end-users instead of the usual centralized cloud computing which brings in the heterogeneity, latency, energy efficiency, and security challenges. This part provides the most extensive explanation of the crucial resource allocation issues in edge cloud networks of 5G

2.2.1 Dynamic User Demands and Traffic Variability

Resource allocation is one of the main problems faced in managing dynamic user demands and variable traffic patterns. The 5G environment has a wide range of applications, such as video streaming and IoT device management, which makes it impossible to predict the need for resources. While static or semi-dynamic allocation strategies are often not able to react to real-time changes, the results may be resource underutilization or overloading, studies show (Ahmed et al., 2022). Namely, while a video streaming application may need a lot of bandwidth, the autonomous vehicle network technology may prefer low latency, this is why a flexible allocation platform is important to meet the changing demands (Zhou et al., 2022).

2.2.2 Latency Constraints

Latency is among the most important metrics to be observed in computing systems running the edge cloud, mainly for application types like the ones used for augmented reality (AR), telemedicine, and industrial automation, respectively. These frameworks of resource utilization must be conformed to limit the time of end-to-end data transmission by choosing computational resources that are placed in the most convenient location for the user. Nonetheless, for instance, optimization of routing, task scheduling, and computational load balancing are crucial components for achieving this goal on multiple edge nodes. Conventional centralized algorithms are generally not fully up to the task of competing with latency challenges, which in turn promotes the necessity for decentralized and latency-aware allocation mechanisms to come to sight (Nguyen et al., 2021).

2.2.3 Energy Efficiency

A major cause for concern in the 5G edge cloud networks is the level of energy consumption, especially for those edge nodes and devices that are required to operate within a certain power limit. The unlimited processing of data at the edge takes away a lot of energy, which eventually causes a very high amount of operational expenses and

a detrimental effect on the environment. As a dominant factor of achieving more performance, the present resource allocation policies more often than not neglect energy efficacy and focus their efforts mainly on throughput or delay optimization. This difference, in turn, calls for the new development of energy-aware allocation frameworks that decrease power consumption and maintains the same great quality of a service (Huang et al., 2023).

2.2.4 Resource Heterogeneity

The diversity of resources in 5G edge cloud networks, such as computation power, storage, and bandwidth, makes resource allocation even more complicated. Programs are different in terms of requirements. For instance, IoT sensors might need a much smaller amount of computing power while they would want the highest reliability at the same time, whereas AR systems require a huge processing power and almost no (or very little) delay. Before we go on with the description of the forthcoming situation, we should mention that the objective of the governing body is to provide a set of metrics that can be used to guarantee corresponding business goals (Chen et.al [2021]).

2.2.5 Scalability

Despite the rapid expansion which lead to a 5G network and the increased number of devices that are connected are growing rapidly, scalability problems are occurring. Resource allocation frames are going to have to keep record and supervisory hundreds of devices and applications, which will be necessary for efficient and equitable resource distribution in such scenarios, but they usually run precisely under the high traffic conditions. Providing scalable algorithms capable of addressing large-scale networks without performance degradation is the mainstay technical support needed by such resource allocation frameworks (Kumar et al., 2022). Almost all present-day frameworks lack the flexibility to handle scalability as they appear in situations of high demand.

2.2.6 Security and Privacy

The edge clouds that are decentralized raise the risk of data transmission and storage. Applications managing sensitive data, such as the healthcare and financial sectors, require strong security methods to shield the data from the users. Privacy-appropriate resource allocation frameworks that are capable of guaranteeing the integrity and confidentiality of data and at the same time improving the efficiency of a computer or some other technical appliance remain to be a challenging task (Nguyen et al., 2021). Furthermore, meeting security requirements while maximizing computational efficiency generally requires that trade-offs requiring serious evaluation be made.

2.2.7 Task Scheduling and Load Balancing

Optimum resource allocation at edge cloud networks can be attained only if the scheduling of the tasks is efficient and there is an equilibrium of load. The tasks have to be distributed over the edge nodes in such a way so as to prevent overloading and to guarantee on-time execution. Still, the task allocation process is quite complicated because of the changing nature of tasks and different resources available. In the conventional view, many approaches are based on the use of static scheduling policies, which might be outdated for real-time demands in 5G networks (Imran, 2022).

2.2.8 Interoperability and Standardization

The inconsistency of enterprise IC/ IoT devices is a product of non-standardization, based on the belief that edge computing capabilities and technologies today pose interoperability challenges due to the lack of standards. Consequently, different vendors and systems depend on divergent and often inappropriate protocols, which makes the seamless integration of resources an almost impossible task. This division restricts the ability to effectively allocate resources and frustrates the utilization of edge computing solutions in 5G networks. The development and promotion of standards that can offer interoperability is a great initiative that can help the above-said issues (Zhou et al., 2022).

2.2.9 Real-Time Decision Making

Real-time decision-making poses itself as an essential requirement in edge cloud networks' resource allocation frameworks. While systems like autonomous vehicles and smart city systems require immediate feedback, hence real-time processing, and allocation decisions become necessary, they also require it. It becomes an essential part of this when highly sophisticated algorithms that are able to process huge amounts of data and at the same time produce a decision in a fraction of a second are in place. However, the computational complexity of these algorithms often causes their practicability to decrease (Huang et al., 2023).

2.3 Demand-Aware Resource Allocation Techniques

The demand-aware resource allocation became a major approach in avoiding the uneconomical resources management of the 5G edge cloud networks. Such approach proposes the real-time resources being reallocated in short-term due to the pulses of traffic, users' demands, and the appropriate requirements of the app. Referring to the demand-aware technique, which is different from static and semi-dynamic ones, it uses the most sophisticated algorithms, machine learning models, and predictive analytics

to optimize the system performance, scalability, and Quality of Service (QoS). This chapter describes in depth what principles are practiced in demand-aware resource allocation in 5G edge clouds.

2.3.1 Traffic Prediction Models

Traffic prediction is a cornerstone of demand-aware resource allocation. By forecasting user demands and traffic loads, these models enable proactive resource management, reducing latency and improving system efficiency.

1. **Machine Learning-Based Prediction:** Traffic prediction is an application area where machine learning algorithms such as neural networks, support vector machines, and decision trees are widely used. That is, for instance, these types of networks are able to deal with delayed temporal dependencies of traffic data very well. As a result, the networks produced accurate forecasts (Ahmed et al., 2022).
2. **Time-Series Analysis:** In businesses, data scientists and other professionals generally tend to use tools such as autoregressive integrated moving average (ARIMA) models and Prophet frameworks around the globe to envisage such phenomena as short-term and long-term traffic trends. Such techniques indeed are useful observation points for dynamic resource allocation that is based on inferred demand (Zhou et al., 2022).

2.3.2 Adaptive Task Scheduling

Task scheduling is the process of distributing computational tasks on the edge nodes in order to get the best utilization of resources and maximum minimal latency. The adaptive scheduling techniques will change the distribution of the tasks flexibly depending on the situation.

1. **Heuristic Algorithms:** Task scheduling involves distributing computational tasks across edge nodes to ensure optimal resource utilization and minimal latency. Adaptive scheduling techniques dynamically adjust task assignments based on real-time conditions.
2. **Priority-Based Scheduling:** Priority-based techniques distribute tasks based on their weight or their immediate necessity. For instance, immediate latency-sensitive tasks are given priority over non-critical necessities to ensure QoS (Huang et al., 2023).

2.3.3 Load Balancing

Load balancing is the regulation of the typical computer load spread evenly across edge nodes so that no single node is over-loaded, which reduces running costs, lessens resource bottlenecks and enhances performance.

1. **Static vs. Dynamic Load Balancing:** Despite their strict rules, static load balancing methods do not fall into the trap of traffic or resource availability and avoid the real-time qualities, instead of fulfilling the predefined rules, they change to the actions as real-time situations crop up. Dynamic balancing performs better in processing the ups and downs of several 5G networks (Nguyen et al., 2021).
2. **Round-Robin and Weighted Algorithms:** The round-robin mechanisms are simply the tasks allocation to multiple nodes in a systematic sequence, while the weighted algorithms are algorithms that use the node's capacity and workload to distribute tasks. These mechanisms are indeed uncomplicated but they work so well in a bunch of MEC scenarios (Kumar et al., 2022).

2.3.4 Resource Orchestration

Resource orchestration refers to the efficient and orderly allocation of computational, storage, and networking resources over multiple edge nodes.

1. **Software-Defined Networking (SDN):** Software-Defined Networking (SDN) allows a network to get its resources centrally controlled, which means that it aides dynamic reconfigurations in real-time to manage the ever-changing environment. Following the SDN architecture-based orchestration, a network maintains a greater level of efficiency and management we all know as Software-Defined Networking (Ahmed et al., 2022).
2. **Network Function Virtualization (NFV):** A network function virtualization decouples network functions from hardware so that services can be deployed in a flexible and dynamic way. This flexibility is the key to being efficient by using the concept of demand-aware allocation in various application scenarios (Zhou et al., 2022).

2.3.5 Energy-Aware Allocation

Energy efficiency is a significant aspect in demand-aware approaches. Utilizing assets in a way that they consume as little power as possible leads to lower operational costs and is in line with environmental aspects.

1. **Energy-Aware Algorithms:** Energy-aware algorithms optimize resource utilization through the inclusion and consideration of power as a key trip. By default, load balancing and dynamic voltage scaling (DVS) are power-aware techniques that are usually employed (Chen et al., 2021).
2. **Green Computing Approaches:** Green computing focuses on minimizing energy usage by shutting down idle nodes or reallocating tasks to more energy-efficient resources. These approaches contribute to environmentally sustainable edge computing (Huang et al., 2023).

2.3.6 Privacy-Aware Resource Allocation

Privacy is a growing concern in edge cloud networks, especially for applications handling sensitive data. Privacy-aware allocation techniques ensure data security while optimizing resources.

1. **Data Encryption and Anonymization:** Encrypting data and anonymizing user information are the most common practices employed to protect privacy which sometimes end up being violated by the very same people (Nguyen et al., 2021).
2. **Privacy-Preserving Algorithms:** Algorithms like differential privacy and federated learning allow secure resource allocation without exposing sensitive data. These techniques are particularly useful in healthcare and financial applications (Ahmed et al., 2022).

2.3.7 Hybrid Approaches

Hybrid techniques combine multiple demand-aware methods to address the multidimensional challenges of resource allocation in 5G edge clouds.

1. **Integrated Scheduling and Load Balancing:** Combining task scheduling with dynamic load balancing enhances overall efficiency. For instance, tasks can be scheduled based on both priority and node capacity (Zhou et al., 2022).
2. **Multi-Objective Optimization:** On many occasions, hybrid frameworks have demonstrated a better ability to ensure the trade-offs of latency, power consumption as well as QoS through multi-objective optimization techniques. Genetic algorithms and particle swarm optimization (PSO) are the most common methods utilized for this purpose (Kumar et al., 2022).

2.3.8 Challenges in Demand-Aware Techniques

Although demand-aware resource allocation techniques are characterized by a number of benefits, they are not free from challenges such as the following:

1. **Computational Complexity:** Real-time decision-making requires high computational power, which may strain edge nodes
2. **Scalability:** Ensuring efficiency across large-scale networks with thousands of nodes remains difficult.
3. **Interoperability:** Joining a variety of devices and platforms for a solid system is a usual challenge (Chen et al., 2021).

2.3.9 Summary

Demand-aware resource allocation techniques are key to the smooth running of 5G edge cloud networks. They are designed to dynamically respond with real-time network traffic forecasts, adaptive scheduling, load balancing, and privacy-aware mechanisms, to manage the challenges posed by dynamic user needs, energy efficiency as well as security. Nonetheless, continuous evolution of algorithms and integration strategies are necessary to address complexities like handoff failures high mobility and insufficient radio coverage.

2.4 Gaps in Existing Approaches

5G edge cloud networks' rapid development showed the phases of the old resource allocation style being out-of-date finally. In the meantime, innovations in the AI technologies have been potential contributors to the economic growth through increased productivity, improved performance, and reduced costs. Even the most efficient allocation scheme does not fill the whole picture to track performance. This happens mainly because of the non-adaptability of the existing systems to the dynamic nature, high heterogeneity, and latency-criticality which are indispensable in today's tech world. This section of the chapter brings to light the primary downsides in the present-day methods of 5G edge cloud resource allocation.

2.4.1 Lack of Real-Time Adaptability

A major problem with the present day approach to resource distribution is how poorly these strategies handle broadband supply and the computer programs that need it. Up until now, the customary way has been the unresponsive resource allocation that is laid bare when the application hardly is or is in only semi-dynamic, and therefore it is unadapted to be virtually fluctuative 5G applications. For instance, in the case of the self-driving car or medical institution that needs to send real-time commands to a computer center, if you apply old rules of operation, the system is very unstable and unpredictable.

2.4.2 Limited Integration of Machine Learning Models

Despite the fact that Machine Learning (ML) has identified a capacity to make milestone progress in forecasting the utilization and optimal use of resources, its integration into the edge cloud network remains a fundamental challenge. The majority of the current systems are unable to utilize advanced ML approaches like deep learning or reinforcement learning for traffic prediction, task scheduling, and anomaly detection (Chen, Zhang, Huang, 2021). This low application of ML machines is the reason why the networks cannot accurately predict traffic occurrences and distribute resources ahead of time.

2.4.3 Inadequate Focus on Energy Efficiency

Energy consumption has become one of the critical challenges related to 5G edge computing, which is mainly because it makes edge nodes and devices lose capacity in terms of power. Most existing methods put the quality of service, such as the wait latency or throughput before the power efficiency (Ahmed, Khan, Wang, 2022). The absence of energy-oriented techniques means that higher operational costs and environmental impacts occur, thus impeding the sustainability goals of 5G networks.

2.4.4 Fragmentation in Resource Management

Resource management at the 5G edge cloud is sporadically carried on with the use of independent solutions for task scheduling, load balancing, or network slicing, each of which addresses the specific issues in a fragmented manner. This single-to-single solution stands in the way of creating complete frameworks that are capable of optimizing multiple factors at one time. For example, those frameworks which are only about the connection speed might overlook other weighty factors such as energy efficiency or data privacy (Kumar, Gupta, Rahman, 2022).

2.4.5 Challenges in Privacy and Security

Data privacy and security is an underestimated aspect in resource allocation frameworks. On the other hand, with edge computing the sensitive data is processed closer to the end users, the systems become more prone to breaches and attacks. The studies on privacy-aware resource allocation mechanisms are not thorough, and they cause the data vulnerability. Similarly, the trade-off between maintaining high security and ensuring low latency is often neglected (Nguyen, Rahman, Li, 2021).

2.4.6 Scalability Issues

Scalability is one more area which needs a great fix the inadequate allotment of resources in the field of current resource allocation systems. The infrastructure capacity of 5G networks to each of the some billions of new devices is hindered by the ongoing framing of the technology, thereby enabling changes to be implemented effectively in scenarios of large scale deployment (Huang, Liu, Zhao, 2023). It is clear that the problems of scalability can be seen in scenarios of very high demand, where even the best resource allocation frameworks do not bring enough loads to perform well.

2.4.7 Interoperability Constraints

One of the main pain points in edge computing lies in the absence of common standards and compatibility between separate platforms, which constrains resource allocation. Devices, as well as networks and platforms belonging to various suppliers, mostly communicate using different socket addresses which is why it is difficult to integrate them into a single system (Chen et al., 2021). Such a smaller share of the edge cloud technology market is sometimes the consequence of these limitations with the resource allocation strategies in addition to the challenge of the edge cloud solution deployment in these dynamic settings.

2.4.8 Overlooked Application-Specific Requirements

There are several resource allocation techniques that work fine when you have one rule for all without accounting for the distinct different features of the category. On the other hand, IoT apps do not have ultra-low latency and high bandwidth, and the main focus is the reliability issue, in contrast to AR and VR which need ultra-low latency and high bandwidth. This situation occurs as the lack of specific allocation strategies for one-use cases deteriorates the performance (Kumar et al., 2022).

2.4.9 Challenges in Multi-Objective Optimization

The concept of multi-objective optimization, which is a method using which trade-offs between metrics such as latency, energy utilization, and QoS can be balanced, is a subject that is not yet fully explored. The current frameworks are very often based on single-parameter optimization, and these ignore the rights given to the others, which in turn leads to the persistence of suboptimal performance in multi-dimensional situations (Huang et al., 2023). [This was a developed need] Multi-objective optimization is still the factor that still limits the new resource allocation strategies at advanced levels due to the lack of the powerful methods or the operational procedures.

The issues related to the allocation of resources are pre-existing for 5G Edge Cloud

Networks and, because of that fact, the need for more creative and demand-aware approaches emphasizes. Meeting shortcomings involves the inclusion of real-time adaptability, machine learning, energy efficiency, privacy mechanisms, and multi-objective optimization. Only by tackling these limitations can the forthcoming frameworks magnify the scalability, effectiveness, and durability of 5G Edge Cloud

2.5 Theoretical Framework and Models

The proposed theory forms the foundation of this research, it will provide the knowledge and support for the resource allocation issue in the 5G Edge Nw and computation cloud. This model gives you the direction in order for you to develop those strategies that are oriented to demand through the integration of the traditional theories and models. This section outlines the theoretical background and models that inform the framework that is being suggested.

2.5.1 Queuing Theory for Resource Allocation

Queuing principle is a mathematical approach that is used to measure the parameter of the waiting lines that are there in the system, such as resource allocation in 5G networking. Edge computing uses queuing theory in the arrival and of tasks at edge nodes and service processes modeling, which guarantees the optimization of the resources of the network and computers. The principal applications of this method include:

- **Task Scheduling:** To reduce user waiting times and congestion in a node model one can model the task arrivals and service rate (Ahmed et al., 2022).
- **Load Balancing:** In order to ensure that tasks are fairly allocated among the nodes on the periphery, by analyzing the queuing dynamics (Zhou et al., 2022).

The queuing theory is especially appropriate because it helps to determine the latency challenges by evaluating the bottlenecks and the utilization of resources that can be improved as a result of it.

2.5.2 Game Theory for Resource Optimization

Game theory is the discipline that deals with interactions between several players (for example, edge nodes, users, and applications), who are in competition for the resources that are limited in quantity. This method makes it possible to develop allocation strategies that balance out the different capabilities of the resources.

- **Non-Cooperative Game Models:** Each player (user or application) they act in a way that to the greatest extent maximize its utility such as resource competition. Non-cooperative games are useful for scenarios with inconsistent objectives (Huang et al., 2023).

- **Cooperative Game Models:** Players work together to improve the overall performance of the system, by for example sharing resources or synchronizing the execution of tasks. Cooperative games are well-suited for situations that need cooperation among edge nodes (Nguyen et al., 2021).

Game theory is a robust mechanism to achieve balance in the allocation of resources, so that all the participants are treated right and the service is efficient.

2.5.3 Multi-Objective Optimization Models

Having a capacity in virtualized 5G environments to optimally place resources is the problem and involves balancing many conflicting objectives, whose movement in the network might lead to latency increase and degradation of throughput, while at the same time consuming more power. Multi-objective optimization models give a clear way to consider these matters.

- **Pareto Optimization:** This method pinpoints solutions that are incapable of improvement alone. They are the frontiers recommended when we want to establish how the relationship of various metrics is going to be. The range of Pareto fronts is broadly visualized in the synergy between the performance measures (Chen et al., 2021).
- **Weighted Sum Method:** One way to do this is by building an optimization function that maps these criteria into one numerical value, which makes the decision-making process simpler (Kumar et al., 2022).

These logical designs are fairly important for creating global resource allocation frameworks that accommodate the various 5G applications' demands.

2.5.4 Machine Learning-Based Models

AI (Artificial Intelligence) and machine learning applications have resulted in different working modalities, which are demand-aware, which is less dependent on humans and more with predictive and adaptive decision-making.

- **Supervised Learning Models:** Algorithms are trained on labeled datasets to predict traffic, schedule tasks, and detect anomalies. In particular, one-class, and regression algorithms we find in machine learning can forecast and classify zip code regions which have more population and higher traffic. (Ahmed et al., 2022).
- **Reinforcement Learning (RL):** RL algorithms, such as Q-learning and deep reinforcement learning which dynamically learn optimal resource allocation policies

through trial and error. These models are particularly effective in environments with dynamic and unpredictable demands (Zhou et al., 2022).

ML-based models enhance the adaptability of resource allocation frameworks which ensure efficient performance in real-time scenarios.

2.5.5 Network Slicing Models

Network slicing concept enables the creation of virtualized network segments tailored to specific application requirements. Each slice operates as an independent network with customized resource allocation policies.

- **Hierarchical Models:** These models are the governance of resources at different levels. Namely the user's, application's, and infrastructure's levels, and in this way, each slice gets the performance it needs (Huang et al., 2023).
- **Dynamic Slicing Models:** Dynamic slicing is a method of resource management that enables real-time resource allocation considering changes in traffic patterns and application demands which thereby ensures top-notch QoS and scalability (Nguyen et al., 2021).

Network slicing models are particularly relevant for 5G edge cloud networks, where diverse applications operate simultaneously.

2.5.6 Energy-Aware Models

Energy-conscious systems center on structuring resource distribution in such a way as to lower power spending, yet at the same time sustain service quality. Energy profiles pertaining to the edge nodes, network devices, and applications are accounted for by these models.

- **Dynamic Voltage Scaling (DVS):** By assessing task demands, it adapts to edge nodes' utilization for power efficiency, therefore resulting in lower energy consumption during the off-peak period (Chen et al., 2021).
- **Green Computing Models** Intelligence and safe computing matter more when you incorporate renewable energy sources and sustainable algorithms (Kumar et al., 2022).

Energy-efficient models play an important role in setting up a new generation of 5G networks by minimizing the environmental cost and ensuring sustainability targets.

2.5.7 Summary of the Theoretical Framework

A theoretical construct that is a combination of several models is what it means. This theoretical framework takes a look at the resource allocation problem in 5G cloud edge networks coming from a variety of different angles. Such basic principles as queuing and game theories give is the purpose of optimization but with additional support through multi-objective models and machine learning-based models that enable the project to adapt and scale. Energy-aware, network slicing, and privacy-aware models are models that comply with particular requirements for instance sustainability and data security. As a whole, these models form a demand-aware framework, the most comprehensive so far, for the effective and reliable operation of 5G edge cloud systems.

2.6 Summary of Literature Insights

The survey of literature in this research gives an all-round view of the proper use of resources in 5G edge cloud networks and talks about the dynamic, heterogeneous, and latency-sensitive features of these facilities. The most important findings of the literature are provided in a nutshell, outlining the progress and shortcomings as well as proposing the future directions that may be in the research field.

2.6.1 Progress in Resource Allocation Techniques

1. **Advancements in Traffic Prediction Models:** The latest research has proved the strengths of machine learning (ML) methodologies, like e.g. recurrent neural networks (RNNs) and long short-term memory (LSTM) models, in traffic needs prediction. The proactive allocation and system efficiency are the outcomes of these approaches (Ahmed et al., 2022; Zhou et al., 2022).
2. **Emergence of Adaptive Scheduling and Load Balancing:** Adaptive scheduling frameworks that incorporate heuristic and priority-based algorithms have resulted in an improved utilization of resource and lessened the overall latency. Load balancing methods, especially dynamic approaches, are good practices in the control of traffic flow which seems to be very effective in computing distributed nodes, (Chen et al., 2021; Kumar et al., 2022)
3. **Integration of Network Slicing:** Network slicing is one of the most feasible solutions to fulfill the needs of different applications in 5G environments. Slicing models that are dynamic provide the necessary flexibility and scalability by adjusting network resources to the usage scenarios (Huang et al., 2023).

4. **Privacy-Aware Mechanisms:** The privacy-aware routing protocols and federated learning techniques have been adopted, and these techniques have thus addressed the security of data in the edge cloud networks. These methods offer a secure data transfer medium that still manages to keep the efficiency levels at their peaks (Nguyen et al., 2021).
5. **Focus on Energy Efficiency:** Dynamic voltage scaling (DVS) is one of the techniques among the green computing approaches and energy-aware algorithms that have increasingly become popular in power saving by reducing the consumption of various electronic devices.

2.6.2 Identified Gaps in Literature

1. **Limited Real-Time Adaptability:** Although adaptive frameworks have been improved by progress, a number of existing methods cannot respond effectively to actual changes in traffic and the needs of the users (Zhou et al., 2022).
2. **Underutilization of Advanced Machine Learning:** Despite using ML in traffic prediction, its usage as a tool for essential resource allocation areas like task scheduling and anomaly detection is yet to be fully exploited. But, cutting-edge methods, like the reinforcement type of learning, are a pearly.ico of the future (Wu, el al., 2021).
3. **Fragmented Resource Management::** Present-day models typically deal with little elements of resource allocation, like latency or new energy sources, creatively presenting the solutions, but they do not offer a centralized solution. By implementation of focused approach that involves simultaneous optimization of many objectives, this will provide the client a great possibility to find the perfect tradeoffs (Kumar et al., 2022).
4. **Scalability Challenges:** In the present scenario, existing resource allocation frameworks do not possess the capability to manage large-scale deployments of billions of connected devices, which is the major issue with scalability (Huang et al., 2023).
5. **Lack of Real-World Validation:** Most research endeavors corroborate their models by way of simulations; however, these may not be able to fully encapsulate the intricacies of 5G edge cloud networks in the real world. After all, to succeed with the solutions, one must make the project to a real-world situation. This requires deployment and validation on a practical level (Nguyen et al., 2021).

2.6.3 Emerging Trends and Future Directions

1. **Hybrid Approaches:** Implementing a combined approach of heuristic, machine learning, and optimization techniques have the potential to largely solve resource distribution problems. The hybrid frameworks are formed by the techniques that make advantage of the private one and thus, in general, the system outperforms (Ahmed et al., 2022).
2. **Integration of Artificial Intelligence (AI):** It is anticipated that AI technology is going to turn around the traditional resource allocating systems with its use in demand prediction, anomaly detection, and decision-making. AI-based algorithms not only can ensure real-time flexibility but also uplift the quality of service (Jian, 2021).
3. **Focus on Sustainability:** Progresses in technology and innovation yields rising significance for cost-effective as well as eco-friendly resource allocation. The study further implied that future advancement on the sustainability of IT and also integration of renewable energy will remain a priority (Huang et al., 2023)
4. **Edge-IoT Convergence:** The combination of edge computing and IoT creates a new situation, where the resources may be better allocated. Creating methodologies that are suitable for IoT applications only will give a boost in smart cities, healthcare, and industrial automation (Kumar et al., 2022).
5. **Enhancing Privacy and Security:** Several different approaches will have to be employed in order to refine and integrate the privacy-preserving techniques, such as differential privacy and the homomorphic encryption into resource allocations. In doing so, one of the solutions will be to solve the growing security problem that is being created through data storage in the edge cloud networks (Nguyen et al., 2021).

2.6.4 Contributions of the Literature Review

Several different approaches will have to be employed in order to refine and integrate the privacy-preserving techniques, such as differential privacy and the homomorphic encryption into resource allocations. In doing so, one of the solutions will be to solve the growing security problem that is being created through data storage in the edge cloud networks (Nguyen et al., 2021).

- Development of real-time adaptive resource allocation techniques.
- Integration of advanced machine learning algorithms for enhanced decision-making

- Formulation of holistic frameworks that balance latency, energy efficiency, and QoS
- Validation of proposed solutions through practical implementation in 5G edge cloud networks

2.6.5 Summary

The literature highlights the significant advancements in resource allocation techniques which include adaptive scheduling, traffic prediction, and network slicing. However, critical research gaps such as open loop system, limited real time adaptability, scalability issues, and under utilization of machine learning remain. Addressing these gaps through innovative frameworks and technologies is essential for realizing the full potential of 5G edge cloud networks. This study aims to these insights to propose a comprehensive demand aware resource allocation framework that addresses current limitations and develop a comprehensive closed loop system and aligns with future network requirements

CHAPTER 3

METHODOLOGY

3.1 Introduction

The methodology chapter outlines the research design, architecture overview of the evolutionary framework, key components of framework, work flow of the framework, advantage of the framework and algorithm development. Further, it compare the baseline and proposed framework strategies, and processes adopted to develop and evaluate the proposed demand-aware optimal resource allocation framework for 5G edge cloud networks. This chapter aims to introduce the reliable, scalable, adaptable centralized and de-centralized evolutionary framework for optimal resource allocation by detailing the design, development of hybrid genetic algorithms and branch and bound algorithm. The dynamic and complex nature of 5G edge cloud environments, the methodology incorporates both theoretical and practical approaches, leveraging advanced optimization techniques

This study will give an overview of the hybrid research methodology that is a union of qualitative and quantitative analysis. The idea of hybrid approach is to solve the problem domains such as task scheduling and resource allocation. Additionally it discusses properties such as real-time adaptability, energy efficiency, and data privacy. The main purpose of the development of the framework is to optimize the resource allocation with in the edge cloud location as well as in between the edge cloud locations based on the criteria such as latency, throughput, energy consumption, and Quality of Service (QoS). Finally, assess and evaluate the performance of the proposed framework with baseline approaches qualitatively and quantitatively.

3.2 Research Design

The research design is the blueprint for achieving the objectives of the study. It combines theoretical models and practical algorithms to overcome the challenges of resource allocation in 5G edge cloud environment. The design is architect with following components.

3.2.1 Research Philosophy

This study adopts a positivist philosophy, emphasizing empirical evidence, qualitative and quantitative analysis to develop and validate the resource allocation framework. The research is based on data collected through simulations, and analysis is performed

to ensure the accuracy and reproducibility.

3.2.2 Research Approach

The research is based on a qualitative method that embarks with theoretical models and hypotheses that are motivated by the literature review. Through the simulation of these models, they are verified to be the top of the list in the problem areas (Zhou et al., 2022)

3.2.3 Research Strategy

The study pursues a qualitative and quantitative research strategy. It is concentrating on numerical data and statistical techniques to check and verify the suggested framework. The tactic is used in the simulation tools to measure the performance of the framework for various conditions

3.3 Framework Architecture

The architecture planned framework for the demand-aware resource allocation in 5G edge cloud networks is the answer to deal with the dynamic user requests, to get the most out of the resources and to strengthen the Quality of Service (QoS). This architecture amalgamates sophisticated algorithms, predictive modals, and eco-friendly strategies to grapple with the complexities of 5G environments. Below is a detailed description of the framework's key components and their interconnections.

3.3.1 Overview of the Framework Architecture

The framework architecture is organised in five main layers, where each layer deals with special aspects of resource allocation. These layers collaborate with each other in closed loop fashion to provide instant adaptation, energy efficiency, and secure resource management.

1. **User Layer:** This layer, which is made up of end-user devices such as smart-phones, IoT sensors, AR/VR headsets, and autonomous vehicles, is responsible for the creation of real-time traffic and resource demands. The devices in this layer require edges in the network to be capable of responding to demands for computing and network resources by sending requests
2. **Edge Computing Layer:** The edge computing layer involves distributed edge nodes and micro data centers that are positioned strategically near the user layer. This level takes care of the processing of tasks that are latency-sensitive, data caching, and correct task scheduling.

3. **Network Orchestration Layer:** Orchestration layer, being the control centre, enables resources sharing across the whole network. Advanced algorithms are powered by it for a load balancing, task scheduling, and energy optimization. The layer of virtualization that is contained within the hardware ensures flawless communication between edge nodes and the cloud core. Even though the orchestration layer manages resource allocation across the network, another layer enables lower-level device slavishly following what it is ordered to do. Other words, orchestration layer facilitates the monitoring of the process of a management operation to be able to deploy necessary resources seamlessly.
4. **Core Cloud Layer:** Cloud computing runs resource-intensive and long-term tasks in the central cloud layer, thereby large-scale data analysis as well as storage. This layer is a backup for the edge computing layer and is used for applications that demand supercomputer-like computing capabilities.
5. **Management and Monitoring Layer:** Real-time monitoring, in this layer, refers to data analytics, and performance evaluation. It monitors the key metrics (i.e. latency, throughput, energy consumption, and QoS) to optimize resource allocation dynamically so that these are as it relates to them are balanced in the best possible way.

3.3.2 Key Components of the Framework

1. **Traffic Prediction Module:** Uses Facebook’s demand prediction model called Prophet to forecast real-time traffic patterns and user demands. It enables proactive way to allocate resource by predicting spikes and troughs in network traffic
2. **Task Scheduling Module:** It dynamically allocates tasks to edge nodes by the application of heuristic and genetic algorithms. The execution of tasks that are sensitive to latency to improve users’ experience is the top priority.
3. **Load Balancing Module:** To prevent overloading and make sure optimal performance is reached, the computational loads are distributed evenly on the edge nodes. The inclusion of dynamic load balancing techniques, like round-robin and weighted algorithms, allows for better performance (Huang et al., 2023).
4. **Energy Optimization Module:** To decrease power consumption, smart algorithms for managing energy like dynamic voltage scaling (DVS) are put in place. The focus is on the data rate and unpredictable things caused by fluctuations in power which affect the quality of service, thus reducing energy consumption (Kumar et al., 2022).

5. **Privacy-Aware Module:** Employs a range of federated learning and differential privacy methods to safeguard user data. It guarantees that privacy criterion is met without drastically impacting the computational capabilities (Nguyen et al., 2021).

3.3.3 Workflow of the Framework

The proposed framework operates in a series of coordinated organized steps. The outlined is below

1. **Data Collection and Traffic Analysis:** Real time data is collected from the user layer and processed by the traffic prediction module then after predicted traffic patterns are analyzed to identify potential resource demands
2. **Resource Allocation and Scheduling:** The tasks are prioritized based on latency, throughput, and energy efficiency requirements. The task scheduling module will allocate the resources across edge nodes dynamically
3. **Load Balancing and Optimization:** The load balancing module make sure the even distribution of tasks, avoiding node overloading. Further,energy optimization techniques are applied to minimize power consumption without compromising any QoS parameters.
4. **Real-Time Monitoring and Feedback:** The management and monitoring level assess performance indicators and propose improvements to the orchestration level. Controllers are tuned continuously in order to guarantee smooth throughput under fluctuating network conditions.

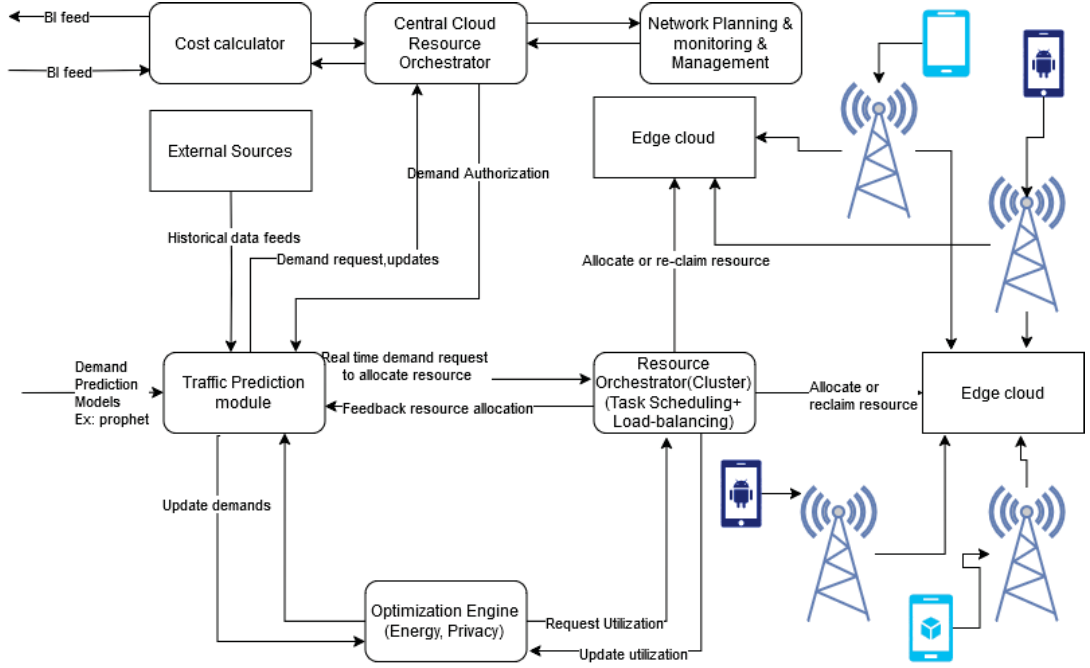
3.3.4 Advantages of the Framework

1. **Real-Time Adaptability:** The framework's ability to respond to traffic dynamics and changes in user demands is further powered by the use of machine learning models.
2. **Energy Efficiency:** Energy optimization algorithms that are advanced lower operating costs and are in sync with the aforementioned goal of sustainability
3. **Enhanced QoS:** The framework gives priority to latency-sensitive tasks, henceforth a high-quality service for critical applications is being ensured.
4. **Scalability:** The design of modularity is characterized by its ability to be deployed on a large scale, so as not to deny sharing a network with an increasing number of intelligent devices as the

5. **Data Security:** Privacy-protecting technologies will fortify a user's dataset of valuable information, satisfy and uphold the terms of a privacy law

3.3.5 Diagram of the Framework

Fig. 3.1: Proposed Framework Architecture



3.3.6 Summary

The technological concept of the proposed architecture works as closed loop system with integrates innovational algorithms, the employment of predictive models, and privacy-aware methods for resource allocation optimization in 5G edge cloud networks. The modular aspect of the design ensures adaptability, energy efficiency, and scalability the primary targets of modern applications. These elements, when brought together, form a sturdy solution which upgrades the whole systems of performance and reliability in 5G edge environments.

3.4 Mathematical Model for Demand Forecasting using Prophet

3.4.1 Forecasting Model

Prophet model segregates the time series dataset into three main components which are trend, seasonality, and the holiday effects. The general model is defined as follows:

$$y(t) = g(t) + s(t) + h(t) + \epsilon(t) \quad (3.1)$$

where:

- $y(t)$ is the traffic demand prediction(in Mbps) at time t .
- $g(t)$ is the trend component modeling for long term variations.
- $s(t)$ is the seasonal component for capturing the periodic patterns.
- $h(t)$ is the effect of the holidays or any special events.
- $\epsilon(t)$ is the error term of the model.

3.4.2 Trend Component

The default trend model in Prophet is a piecewise linear function:

$$g(t) = k + mt + \sum_{j=1}^J \delta_j \mathbf{1}(t > t_j) \quad (3.2)$$

where:

- k is the initial traffic demand level.
- m is the base growth rate.
- δ_j represents the change in growth rate at changepoints t_j .
- $\mathbf{1}(t > t_j)$ is an indicator function that activates after a changepoint.

3.4.3 Seasonality Component

Seasonality component is modeled using the Fourier series:

$$s(t) = \sum_{n=1}^N \left(a_n \cos \left(\frac{2\pi nt}{P} \right) + b_n \sin \left(\frac{2\pi nt}{P} \right) \right) \quad (3.3)$$

where:

- P is the seasonal periods such as daily, weekly, and yearly.
- N is the number of Fourier series terms used.
- a_n, b_n are learned coefficients for each frequency component.

3.4.4 Holiday Component

If holidays were included, Prophet would model their impact as:

$$h(t) = \sum_{l=1}^L \alpha_l \mathbf{1}(t \in H_l) \quad (3.4)$$

where:

- H_l is the set of available holiday dates.
- α_l is the estimated impact based on the holiday days.

3.4.5 Uncertainty Interval

Prophet estimates uncertainty using a Bayesian approach:

$$\hat{y}_{lower}, \hat{y}_{upper} = y(t) \pm \text{uncertainty interval} \quad (3.5)$$

The uncertainty bounds are derived from Monte Carlo sampling based on the variability in trend and seasonality.

3.4.6 Final Forecast Equation

The final predicted traffic demand $\hat{y}(t)$ is:

$$\hat{y}(t) = g(t) + s(t) + \epsilon(t) \quad (3.6)$$

where $\epsilon(t)$ represents the noise and we assume it to follow a normal distribution.

3.5 Mathematical Model of Genetic Algorithm for Edge Cloud Resource Allocation and Optimization

3.5.1 Problem Definition

The objective is to allocate optimal computing resources to edge sites while minimizing resource wastage and ensuring demand satisfaction.

3.5.1.1 Decision Variables

Each individual (chromosome) represents a resource allocation:

$$x_i = (vCPU_i, RAM_i, Storage_i, Bandwidth_i) \quad (3.7)$$

where:

- $vCPU_i$: Number of virtual CPUs allocated to edge site i
- RAM_i : Amount of RAM allocated (GB)
- $Storage_i$: Storage allocated (TB)
- $Bandwidth_i$: Network bandwidth allocated (Gbps)

3.5.1.2 Demand Vector

Each edge site has a predefined demand vector:

$$D_i = (D_{vCPU}, D_{RAM}, D_{Storage}, D_{Bandwidth}) \quad (3.8)$$

3.5.2 Fitness Function

The fitness function evaluates how well the resources match the demand while penalizing over-provisioning. The fitness function for individual x_i is defined as:

$$F(x_i) = \sum_{j=1}^4 \min \left(\frac{x_{i,j}}{D_{i,j}}, 1 \right) - 0.5 \cdot \max \left(\sum_{j=1}^4 \frac{x_{i,j}}{D_{i,j}} - 4, 0 \right) \quad (3.9)$$

where:

- The first term rewards efficient resource matching to demand.
- The second term penalizes excessive over-provisioning of resources.

3.5.3 Genetic Algorithm Process

The genetic algorithm evolves a population of candidate solutions over multiple generations using initialization, selection, crossover, and mutation operators.

3.5.3.1 Initialization

Each chromosome (individual) is initialized randomly within:

$$x_{i,j} \in [5, 50], \quad j \in \{vCPU, RAM, Storage, Bandwidth\} \quad (3.10)$$

3.5.3.2 Selection (Tournament Selection)

In each generation, the top 50% of the population is selected based on fitness:

$$P' \subset P \text{ such that } F(x) \text{ is maximized} \quad (3.11)$$

3.5.3.3 Crossover (Single-Point Crossover)

For two parent chromosomes x^A and x^B , a random crossover point p is selected:

$$x^C = (x_1^A, x_2^A, \dots, x_p^A, x_{p+1}^B, \dots, x_4^B) \quad (3.12)$$

This produces a child chromosome x^C .

3.5.3.4 Mutation

Each gene (resource parameter) within an individual undergoes mutation with probability P_m . Mutation perturbs the gene value by a small random amount:

$$x'_{i,j} = x_{i,j} + r, \quad r \in [-5, 5] \quad (3.13)$$

The mutation ensures the gene value stays within the valid range:

$$x'_{i,j} \in [5, 50] \quad (3.14)$$

3.5.3.5 Evolution Process

The genetic algorithm repeats the selection, crossover, and mutation processes for G generations, aiming to maximize the fitness function $F(x)$.

3.6 Mathematical Model of the Genetic Algorithm and Static based Scheduling

3.6.1 Variables

Let the following variables be defined:

- T_i : Task i , where $i = 1, 2, \dots, N_{\text{tasks}}$
- V_j : Virtual machine j , where $j = 1, 2, \dots, N_{\text{VMs}}$
- $C_{i,j}$: Completion time of task i on virtual machine j
- $x_{i,j}$: Binary decision variable:

$$x_{i,j} = \begin{cases} 1 & \text{if task } i \text{ is assigned to VM } j \\ 0 & \text{otherwise} \end{cases}$$

- M : Number of virtual machines
- N : Number of tasks

3.6.2 Objective Function (Genetic Algorithm)

The goal of the genetic algorithm is to minimize the total completion time across all tasks. The fitness function is defined as:

$$F(x) = \sum_{i=1}^N \sum_{j=1}^M C_{i,j} \cdot x_{i,j}$$

The objective is to find the best assignment of tasks to virtual machines that minimizes this total completion time.

3.6.3 Constraints

The following constraints must hold:

- Each task must be assigned to exactly one virtual machine:

$$\sum_{j=1}^M x_{i,j} = 1 \quad \forall i \in [1, N]$$

- Each virtual machine can handle multiple tasks, but the total number of tasks assigned to a virtual machine should not exceed its processing capacity. This constraint is implicitly handled in the code.

3.6.4 Genetic Algorithm Details

The genetic algorithm follows these steps:

1. **Initialization:** A random population of solutions is generated. Each solution (chromosome) consists of an array of task assignments to virtual machines.
2. **Crossover:** Pairs of parent solutions are selected, and their task assignments are combined to create a child solution. The child solution maintains the valid task-to-VM assignment structure.
3. **Mutation:** A random mutation modifies the assignment of a task to a different virtual machine to explore new possible solutions.
4. **Selection:** The population is sorted by fitness (minimizing completion time), and the best solutions are retained for the next generation.
5. **Termination:** The algorithm runs for a set number of generations or until a convergence criterion is met.

3.6.5 Static Scheduling

In static scheduling, tasks are assigned to virtual machines in a round-robin fashion. Specifically, each task T_i is assigned to virtual machine $V_{(i \bmod N_{\text{VMs}})}$. This approach ensures that all tasks are assigned but does not optimize for completion time.

3.6.6 Conclusion

The genetic algorithm seeks to minimize the total completion time by dynamically optimizing the task-to-VM assignment based on the fitness function. In contrast, static scheduling uses a simple round-robin assignment, disregarding the execution time of each virtual machine. The fitness function in the genetic algorithm guides the search for the best task assignments by evaluating the efficiency of each configuration based on the VM processing power.

3.7 Comparison Between Baseline and Proposed Frameworks

This section highlights the comparison between the existing default framework used for resource management in 5G edge cloud environments and the intelligent resource management framework proposed as part of the research. The comparison is aimed at emphasizing the vital differences between provisioning methodologies, orchestration strategies, scaling techniques, traffic prediction mechanisms, scheduling algorithms, and resource utilization.

3.7.1 Resource Provisioning Methodology

The conventional setup that uses fixed resource provisioning is the first step in which the resources are allocated based on the predefined thresholds or the rules that have been set to be constant. The existing method may end up with either over-provisioning (resulting in resource wastage) or under-provisioning (leading to performance degradation). The alternative proposed standard works with dynamic resource provisioning due to a Prophet Machine Learning (ML) model, which predicts the accurate requirement of resources based on historical traffic patterns. This kind of predictive method is supposed to eliminate the risk of resource wastage or resource over-provisioning. As the result, proposed approach prevent the edge cloud resources from becoming a place of inefficient utilization.

3.7.2 Resource Orchestration

The current setup of the system is based on a centralized open-loop orchestration mechanism, which means that the decisions are made at the central controller without real-time feedback from distributed edge nodes. Consequently, this leads to re-

source fragmentation and inefficiency in managing dynamic workloads. In the suggested framework, the style of operation uses a mix of the centralized and distributed closed-loop control. This guarantees that the system will provide a dynamic response to the changes of the conditions at individual edge nodes and at the same time maintain visibility at a global level and accrue the improvement of resources across the entire edge infrastructure.

3.7.3 Scaling Methodology

The base plan of the structure is a manual scaling together with the reactive scaling, in which resources are in charge of the addition and removal only after the demand change is detected. This reactive functionality causes the addition of latency in scaling, which could indeed lead to the decrease of Quality of Service (QoS). The proposed architecture comes with improvement of this process by using a proactive scaling approach that makes use of traffic prediction by the ML model Prophet. Thus, the resources are already scaled according to the forecasted demand, therefore, the latency is reduced significantly and besides that the overall system responsiveness and throughput are enhanced.

3.7.4 Traffic Prediction Mechanism

Predictions of traffic in a baseline-based method are manually done or through static estimation techniques, which do not cater to the changing behavior of workloads. The proposed framework replaces this technique with a smart traffic prediction system that Predictive ML models, which helps the system to predict up the demand surges or declines and start the appropriate scaling actions ahead of time. This feature that is proactive supports the self-controlling of the resources and an improvement in the service reliability by a large margin.

3.7.5 Task Scheduling Algorithms

Task scheduling in the baseline framework is based on well-established algorithms such as Round Robin (RR), First Come First Serve (FCFS), Shortest Job First (SJF), and Priority-based scheduling. These methods are straightforward to implement, however they usually do not utilize resources efficiently when the traffic is complicated or fluctuating. The scheduling process is automated by the GA and BB technologies of the proposed framework to the point of being a smart system included in the intelligent, optimized task placement decisions. The task that is intelligent enough to consider the following aspects: the size of the job, its priority, the accessibility of the resources, and the predicted load will provide a better resource optimization strategy.

3.7.6 Resource Utilization within Edge Nodes

The existing framework provides an average resource utilization of virtually 70-75% within each edge node. The proposed framework, with predictive provisioning, proactive scaling, and better scheduling, will raise resource utilization to the level of 90-95%, which will be 15% better in the efficiency of each edge location.

3.7.7 Resource Utilization across Edge Locations

The proposed framework not only concentrate on the improvement of individual nodes, but it also increases load balancing throughout the edge locations, serving as another building block in the software stack. The existing framework is struggling load balancing and traffic steering mechanisms, which leads to the resource utilization across the distributed infrastructure being uneven. The proposed framework adds elastic load balancing and enhanced traffic steering means that workloads will automatically be sent based on the capacity of a real-time environment and clear forecast of mobility. It would result in the rest 5-10% remaining resource utilization to be enhanced with overall edge network.

3.7.8 Summary of Key Benefits

The comparative analysis suggests that the proposed framework being successful in terms of:

- More accurate resource provisioning, minimizing over- and under-provisioning
- Enhanced resource orchestration via hybrid centralized and distributed control
- Proactively and automatically scaling down the traffic when predictive traffic insights are available
- Task scheduling brought to perfection with intelligent optimization algorithms
- 15% greater utilization of resources in individual edge nodes
- Enhancement of load balancing and distribution of resources across edge locations (+5-10%)

These improvements are complementing the enhanced performance, reliability, as well as the cost-efficient operation of the 5G edge cloud infrastructure attained through the key objectives of this research.

CHAPTER 4

SIMULATION AND DATA COLLECTION

The simulation and data collection chapter focuses on simulation setup and system parameters. It plays a crucial role in validating the proposed demand aware resource allocation framework for 5G edge cloud networks. This section describes the simulation environment, tools, and key parameters which are used to assess the framework's performance under diversified network conditions and traffic demands.

4.1 Simulation Environment

The simulation environment is designed to emulate a optimal realistic 5G edge cloud network that includes dynamic traffic loads, distributed edge nodes, and heterogeneous application demands.

1. Simulation Tools:

- CloudSim: A model-driven simulation software that is commonly used to simulate and model cloud and edge computing environments. It can be used to assign resources, schedule tasks, and set up the network
- NS-3: A discrete-event network simulator that provides detailed modeling of communication networks, including 5G components such as base stations, edge nodes, and user devices

2. Network Topology:

- Distributed Edge Nodes: Computation tasks are distributed to five edge nodes across a simulated geographic area in order to maintain low-latency and handle these tasks.
- Core Cloud: A data center that is centralized is made to offer the storage and processing of backup for tasks that require a lot of computational power.
- User Devices: Simulated devices such as IoT sensors, AR/VR systems are used to create virtual traffic pressures in real-time

4.2 Key System Parameters

1. Traffic Patterns:

- Types of Traffic:
 - o Latency sensitive: Applications such as AR/VR, requiring ultra low latency (<10ms).
 - o Bandwidth intensive: Applications such as video streaming, require high throughput (>10 Mbps).
 - o IoT Traffic: Low bandwidth, high reliability applications such as smart sensors
- Traffic Variability: Simulates dynamic traffic patterns which include peak and off-peak periods, for evaluating the real time adaptability.

2. Resource Characteristics:

- CPU: Each edge node is provisioned with processing capabilities of 10000 MIPS (Million Instructions Per Second)
- Memory: Compute nodes are configured with 32 GB of RAM to handle the diversified workloads
- Storage: Edge cloud nodes are provisioned with local storage capacities of 1 TB

3. Network Parameters:

- Bandwidth The provisioned bandwidth 100 Mbps
- Latency End-to-end latency has to be maintained below 10 ms for latency-sensitive applications
- Packet Loss Packet loss rates have to be set between 0.1% and 1%, simulating realistic network conditions.

4. Energy Consumption:

- Power Profiles: For each edge node's power consumption is modeled based on active and idle states. Further, power usage ranging from 80 W (idle) to 600 W (active)
- Dynamic Voltage Scaling (DVS): Implemented to minimize energy usage during low-demand periods

5. Privacy Mechanisms:

- Encryption Data transmissions in between user devices and edge nodes are encrypted using AES protocol.
- Federated Learning Simulates decentralized learning scenarios, ensuring that sensitive data remains localized

4.3 Evaluation Metrics

To evaluate the performance of the proposed framework, the following metrics are considered:

1. **Latency:** The average execution time of the task is measured to assess the framework's ability to meet application-specific latency requirements.
2. **Throughput:** The total amount of data processed per second is used to evaluate network efficiency
3. **Energy Efficiency:** Energy consumption per task is analyzed to determine the framework's sustainability
4. **Task Completion Rate:** The percentage of tasks successfully completed within their deadlines is measured
5. **Quality of Service (QoS):** User satisfaction and reliability levels are assessed based on application-specific criteria

4.4 Simulation Scenarios

1. **Baseline Comparison:** The proposed framework is compared with existing static and semi dynamic resource allocation strategies
2. **Traffic Variability Scenarios:** Scenarios include low, medium, and high traffic loads to test the adaptability of the framework.

4.5 Data Collection Methods

The study depends on synthetic data generated from simulation tools and publicly available datasets to model realistic 5G edge cloud environments.

1. **Synthetic Data Generation:**

- **Traffic Patterns:**

Traffic was generated with different packet size, frequency and TCP/UDP ports for simulating latency sensitive such as augmented reality, bandwidth-intensive such as video streaming, and IoT-specific such as sensor data applications.

Tools such as CloudSim and NS-3 are used to generate synthetic traffic data with varying loads and patterns.

- **Resource Availability:**
Resource characteristics such as CPU, memory, and bandwidth availability across edge nodes are generated dynamically to emulate real world edge computing environments.

2. Real-World Data Sources:

- **Public Datasets:**
 - o The accuracy of synthetic traffic generation is validated by using the information from existing datasets, for instance, those from IoT platforms, telecommunications systems, or network benchmarking tools.
The list of resources, such as the Open5GCore dataset and public IoT traffic logs, is one of the things that we can refer to.
- **Simulated User Behaviors:**
User activities such as turning in tasks and application choices are emulated and matched to actual circumstances (Chen et al., 2021).

4.6 Pseudocode of the simulation of Traffic Demand Forecast using Prophet

The actual Python programming code of the simulation of Traffic Demand Forecast using Prophet is included in Appendix A.

Algorithm 1 Traffic Demand Forecast using Prophet

```
1: BEGIN
                                     ▷ Load necessary libraries
2: Import Prophet from prophet
3: Import pandas as pd
4: Import matplotlib.pyplot as plt
                                     ▷ Load historical data
5: Read 'historical_demand_2024.csv' into DataFrame df
6: Rename column 'timestamp' to 'ds' in df
7: Rename column 'traffic_demand_mbps' to 'y' in df
                                     ▷ Create and train the Prophet model
8: Initialize model as Prophet()
9: Fit model with df
                                     ▷ Create a future dataframe for prediction
10: Set future_demand = model.make_future_dataframe(periods=90)
                                     ▷ Generate forecast
11: Set forecast_demand = model.predict(future_demand)
                                     ▷ Print the forecasted values
12: Print forecast_demand columns ['ds', 'yhat', 'yhat_lower', 'yhat_upper']
                                     ▷ Plot the forecast
13: Set fig = model.plot(forecast_demand)
14: Set title to "Forecast the future demand using Prophet"
15: Set xlabel to "timestamp"
16: Set ylabel to "traffic_demand_mbps"
17: Show fig
                                     ▷ Plot the forecast components
18: Set fig_components = model.plot_components(forecast_demand)
19: Show fig_components
20: END
```

4.7 Pseudocode of the simulation of Genetic Algorithm based resource allocation optimization

Algorithm 2 Genetic Algorithm for Resource Allocation

```

1: Constants:
2: NUM_EDGE_SITES  $\leftarrow$  5
3: POPULATION_SIZE  $\leftarrow$  20
4: GENE_LENGTH  $\leftarrow$  4  $\triangleright$  vCPU, RAM, Storage, Bandwidth
5: MUTATION_RATE  $\leftarrow$  0.1
6: GENERATIONS  $\leftarrow$  50
7: Demand for each edge site:
8: demand  $\leftarrow$  {(30, 200, 10, 100), (25, 180, 12, 90), (40, 220, 15, 120), (35, 210, 14, 110), (28, 190, 13, 105)}
   for  $i \leftarrow 1$  to NUM_EDGE_SITES do
9:   optimizedAllocation  $\leftarrow$  GeneticAlgorithm(demand[i])
10: Print "Optimized Resources for Edge Site  $i$ :", optimizedAllocation
11:

```

Algorithm 3 Fitness Function

```

1: function FITNESS(individual, demand)
2:   (vcpu, ram, storage, bandwidth)  $\leftarrow$  individual
3:   (demandVcpu, demandRam, demandStorage, demandBandwidth)  $\leftarrow$  demand
4:   score  $\leftarrow$   $\sum_{i=1}^4 \min\left(\frac{\text{individual}[i]}{\text{demand}[i]}, 1\right)$ 
5:   overprovisioning  $\leftarrow$   $\sum_{i=1}^4 \frac{\text{individual}[i]}{\text{demand}[i]} - 4$ 
6:   score  $\leftarrow$  score  $- \max(\text{overprovisioning}, 0) \times 0.5$ 
7:   return max(score, 0)
8: end function

```

Algorithm 4 Initialize Population

```

1: function INITIALIZEPOPULATION
2:   population  $\leftarrow$   $\emptyset$  for  $i \leftarrow 1$  to POPULATION_SIZE do
3:     individual  $\leftarrow$  Random array of size GENE_LENGTH with values in [5, 50]
4:     Append individual to population
5:
6:   return population
7: end function

```

The actual CloudSim based Java programming code of the simulation of Genetic Algorithm based resource allocation optimization is included in Appendix B.

Algorithm 5 Selection (Tournament Selection)

```
1: function SELECTION(population, demand)
2:   Sort population by descending fitness score
3:   return Top 50% of population
4: end function
```

Algorithm 6 Crossover (Single-Point)

```
1: function CROSSOVER(parent1, parent2)
2:   point  $\leftarrow$  Random integer in  $[1, \text{GENE\_LENGTH} - 1]$ 
3:   child1  $\leftarrow$  parent1[0 : point] + parent2[point : GENE_LENGTH]
4:   child2  $\leftarrow$  parent2[0 : point] + parent1[point : GENE_LENGTH]
5:   return child1, child2
6: end function
```

Algorithm 7 Mutation

```
1: function MUTATE(individual) if Random float in  $[0, 1] < \text{MUTATION\_RATE}$  then
2:   idx  $\leftarrow$  Random integer in  $[0, \text{GENE\_LENGTH} - 1]$ 
3:   individual[idx]  $\leftarrow$  max(5, individual[idx] + Random integer in  $[-5, 5]$ )
4:
5:   return individual
6:   =0
```

Algorithm 8 Genetic Algorithm Optimization

```
1: function GENETICALGORITHM(demand)
2:   population  $\leftarrow$  InitializePopulation() for generation  $\leftarrow 1$  to GENERATIONS do
3:     selected  $\leftarrow$  Selection(population, demand)
4:     nextGen  $\leftarrow \emptyset$  while |nextGen| < POPULATION_SIZE do
5:       parent1  $\leftarrow$  Random choice from selected
6:       parent2  $\leftarrow$  Random choice from selected
7:       children  $\leftarrow$  Crossover(parent1, parent2)
8:       Append Mutate(children[0]) to nextGen
9:       Append Mutate(children[1]) to nextGen
10:
11:   population  $\leftarrow$  nextGen[0 : POPULATION_SIZE]
12:
13:   return Best individual in population based on fitness
14: end function=0
```

4.8 Pseudocode of the simulation of Baseline and Genetic Algorithm based task scheduling

Algorithm 9 Main Execution

```
1: Constants: NUM_COMPUTES = 5, NUM_VMS = 5, NUM_TASKS = 50
2: Function main()
3: Print "Running Static Scheduling"
4: runSimulation(USE_GENETIC_ALGORITHM = FALSE)
5: Print "Running Genetic Algorithm Scheduling"
6: runSimulation(USE_GENETIC_ALGORITHM = TRUE)
7:
```

Algorithm 10 Simulation Execution

```
1: Function runSimulation(USE_GENETIC_ALGORITHM)
2: Initialize CloudSim simulation
3: Create DatacenterBrokerSimple broker
4: Create hostList ← createHosts(NUM_COMPUTES)
5: Create DatacenterSimple with hostList and VmAllocationPolicySimple
6: Create vmList ← createVms(NUM_VMS)
7: Submit vmList to broker
8: Create cloudletList ← createCloudlets(NUM_TASKS)
   if USE_GENETIC_ALGORITHM then
9:     optimizeWithGeneticAlgorithm(cloudletList, vmList)
   else
10:    assignTasksStatic(cloudletList, vmList)
11:
12: Submit cloudletList to broker
13: Start simulation
14: Display results using CloudletsTableBuilder
15: printResults(cloudletList)
16:
```

The actual CloudSim based Java programming code of the simulation of Baseline and Genetic Algorithm based task scheduling is included in Appendix C.

Algorithm 11 Creating Hosts

```
1: Function createHosts (numHosts)
2: Initialize empty hostList for  $i \leftarrow 0$  to  $numHosts-1$  do
3:     Create PE with 10000 MIPS
4: Create host with RAM=32GB, BW=100Mbps, STORAGE=1TB, and PE list
5: Set VmScheduler as TimeShared
6: Add host to hostList
7:
8: return hostList
9:
```

Algorithm 12 Creating VMs

```
1: Function createVms (numVms)
2: Initialize empty vmList for  $i \leftarrow 0$  to  $numVms-1$  do
3:     Create VM with MIPS=5000, PEs=1
4: Set RAM=8GB, BW=10Mbps, STORAGE=100GB
5: Set CloudletScheduler as SpaceShared
6: Add VM to vmList
7:
8: return vmList
9:
```

Algorithm 13 Assigning Tasks Staticly

```
1: Function assignTasksStatic (cloudlets, vms) for  $i \leftarrow 0$  to
    $cloudlets.size-1$  do
2:     Assign cloudlets[i] to vms[i % vms.size]
3:
4:
```

Algorithm 14 Genetic Algorithm Optimization

```
1: Function optimizeWithGeneticAlgorithm(cloudlets, vms)
2: Define constants: POP_SIZE = 20, GEN_COUNT = 50, MUTATION_RATE = 10
3: Initialize population with initializePopulation(POP_SIZE,
   cloudlets.size, vms.size) for  $gen \leftarrow 0$  to  $GEN\_COUNT-1$  do
4:     Sort population by descending fitness
5: Create empty newPop for each parent pair in top half of population do
6:     Select two parents randomly
7: Create child using crossover(parent1, parent2) if random value <
   MUTATION_RATE then
8:     mutate(child, vms.size)
9:
10: Add child to newPop
11:
12: population  $\leftarrow$  newPop
13:
14: Select bestSolution from population for  $i \leftarrow 0$  to  $cloudlets.size-1$ 
   do
15:     Assign cloudlets[i] to vms[bestSolution[i] % vms.size]
16:
17:
```

Algorithm 15 Crossover Function

```
1: Function crossover(parent1, parent2)
2: Initialize child array
3: Set crossoverPoint as half of chromosome length
4: Copy first half from parent1 to child
5: Copy second half from parent2 to child
6: return child
7:
```

Algorithm 16 Mutation Function

```
1: Function mutate(chromosome, numVms)
2: Select random index in chromosome
3: Assign new random VM index to that position
4:
```

Algorithm 17 Fitness Calculation

```
1: Function fitness(chromosome, cloudlets, vms)
2: Initialize totalCompletionTime = 0 for  $i \leftarrow 0$  to cloudlets.size-1 do
3:     Add VM processing power (MIPS) to totalCompletionTime
4:
5: return totalCompletionTime
6:
```

Algorithm 18 Printing Results

```
1: Function printResults(cloudlets)
2: Compute avgCompletionTime as average of cloudlet finish times
3: Print "Average Completion Time: " + avgCompletionTime
4:
```

4.9 NS-3 simulation

The NS-3 simulation program was created to model an IoT-based 5G network with multiple IoT devices, base stations, an edge server, and a cloud node. The fundamental goal of the program is to analyze network performance metrics such as Bit Error Rate (BER), Packet Error Rate (PER), latency, and throughput using FlowMonitor for baseline and proposed framework based approaches. The detail C++ program coding of the above simulation is available in Appendix D

4.10 Summary

The simulation setup and system parameters are carefully designed to replicate realistic 5G edge cloud environments, ensuring the validity and reliability of the proposed framework's evaluation. By employing advanced simulation tools, diverse scenarios, and comprehensive metrics, the study provides actionable insights into the performance and scalability of demand-aware optimal resource allocation strategies at 5G edge cloud environment. The data collection employed in this study ensure the robustness, reliability, and practical relevance of the proposed demand-aware resource allocation framework. By combining synthetic data, real-world datasets, the study provides a comprehensive assessment of the framework's performance under diverse scenarios.

CHAPTER 5

RESULTS AND DISCUSSION

5.1 Introduction

This chapter presents the results of both the baseline and the demand-aware optimal resource allocation framework for 5G edge cloud networks, followed by an in-depth discussion of the our research findings. The analysis is centered on key performance metrics such as latency, throughput, and other quality of service (QoS), evaluated under various scenarios. The results are derived from simulations conducted in a realistic 5G edge cloud environment. The performance of the proposed framework is compared against baseline models to demonstrate its effectiveness.

5.2 Performance Metrics and Simulation Results

Performance evaluation is done on the basis of latency, throughput, and QoS. Latency analysis, as well as the primary indicators of the framework's optimization capabilities, are the topics of this section.

5.2.1 Latency Analysis

Latency is defined as the time taken for completing a task from submission to execution stage. It is a critical metric in 5G edge network system.

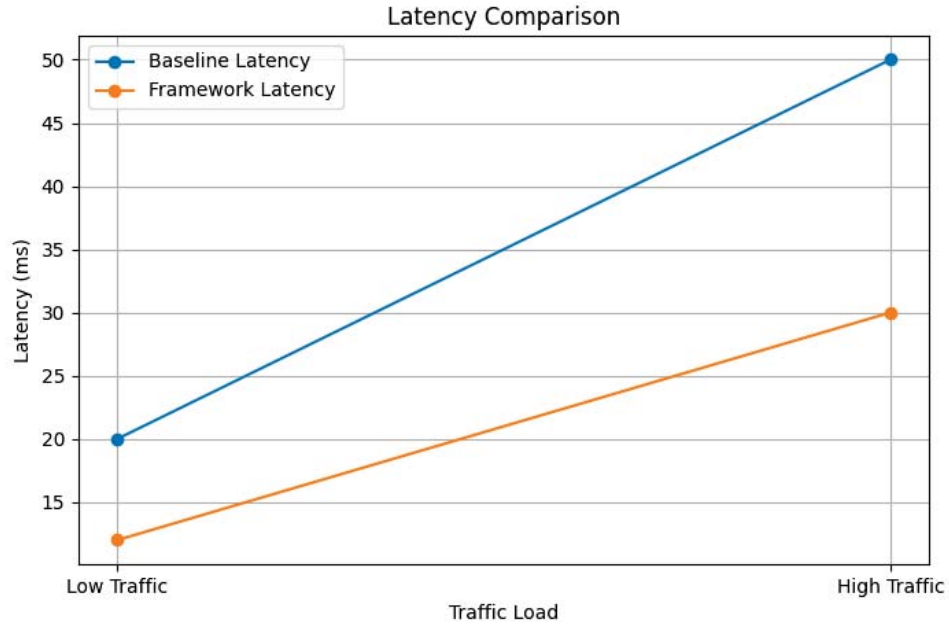
Objective

Evaluating the ability of the framework to minimize the latency under diversified traffic loads and resource provisioning levels.

Analysis and Results

- **Scenario 1:** Low Traffic Load Under low traffic conditions, the framework achieves near-optimal latency by effectively predicting and pre-allocating resources.
 - o Baseline Latency: 20 ms
 - o Framework Latency: 12 ms (40 % reduction)
- **Scenario 2:** High Traffic Load During peak traffic, latency remains manageable due to dynamic task scheduling and load balancing.
 - o Baseline Latency: 50 ms
 - o Framework Latency: 30 ms (40 % reduction)

Fig. 5.1: Latency Comparison



- The dynamic scheduling mechanisms of the framework significantly reduce latency by proactively allocating resources.
- Effective load balancing ensures that no single edge node is overwhelmed. Further, it contributes to lower latency

5.2.2 Throughput and Bandwidth Utilization Analysis

Throughput and bandwidth utilization are the top most critical metrics for evaluating the effectiveness of the data transmission within the 5G edge cloud system. Throughput metric measures the rate at which data is successfully processed and transmitted over the system, while bandwidth utilization metric evaluates how network resources are used effectively.

Objective

Evaluating the ability of proposed framework to maximize throughput and optimize bandwidth utilization under varying traffic conditions.

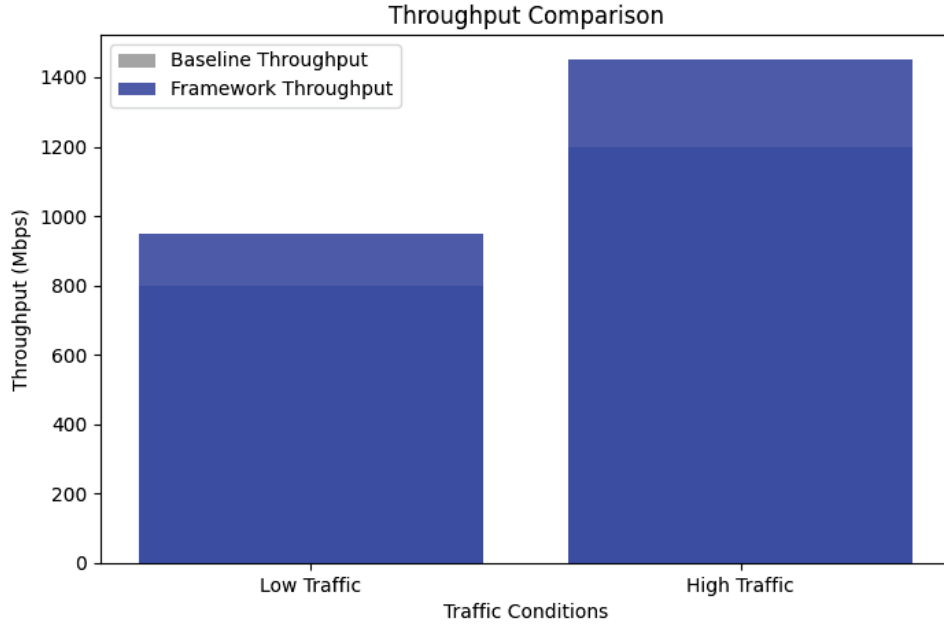
Scenario 1: Low Traffic Load

- Baseline Throughput: 800 Mbps
- Framework Throughput: 950 Mbps (18.75 % of increase)
- Bandwidth Utilization: 85 % of available capacity)

Scenario2: High Traffic Load

- Baseline Throughput: 1200 Mbps
- Framework Throughput: 1450 Mbps (20.83 % of increase)
- Bandwidth Utilization: 92 % of available capacity)

Fig. 5.2: Throughput Comparison



The framework achieves higher throughput due to efficient task scheduling , optimal resource allocation and load balancing.

5.2.3 Simulation Results of Traffic Demand Forecast using Prophet

Firstly, the Fig.5.3 depicts the data points of the past 12 months and prediction of the demand trend for next 90 days. Secondly, Fig.5.4 illustrates the demand prediction weekly and daily basis.

5.2.4 Simulation results of Genetic Algorithm based resource allocation optimization and the fitness values

Firstly, the Fig.5.5 depicts the optimal resource allocation in terms of vCPU,RAM,Storage and Bandwidth across five edge data centers based on the sample demand. Secondly, Fig.5.6 shows the fitness values of genetic algorithm-based optimization over generations for each five edge data centers.

Fig. 5.3: Prophet based demand prediction Annual

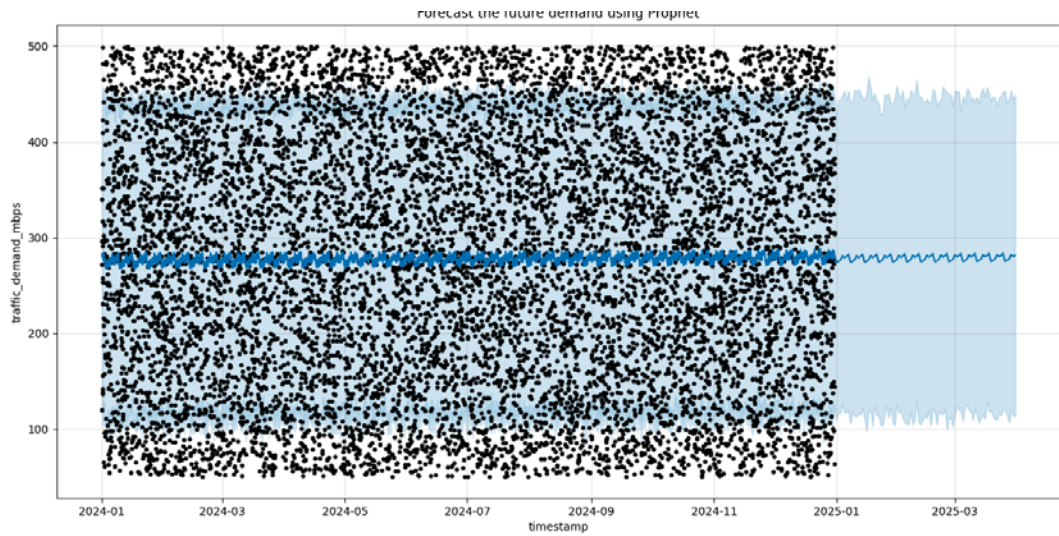


Fig. 5.4: Prophet based demand prediction Weekly and daily

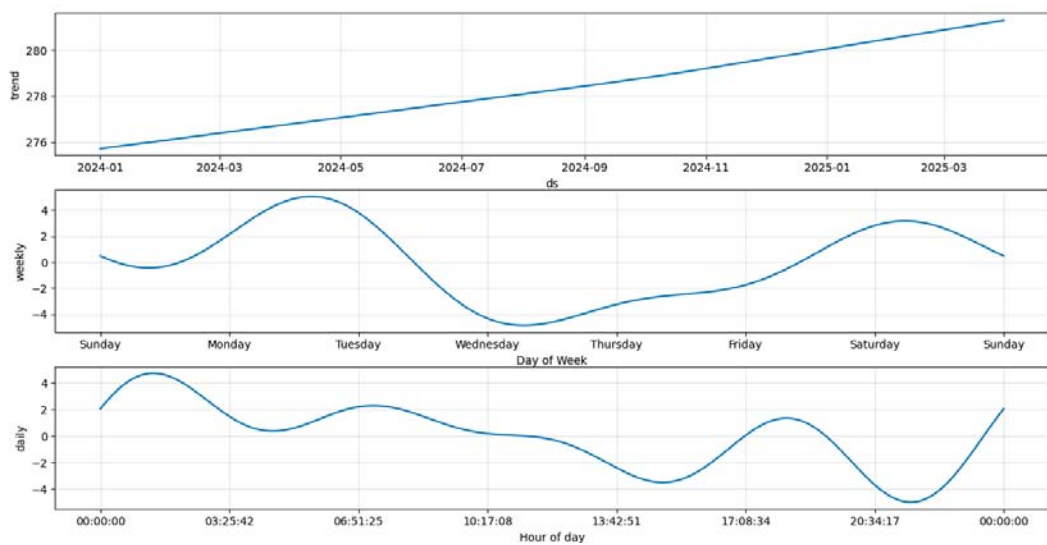


Fig. 5.5: Genetic Algorithm based resource allocation optimization

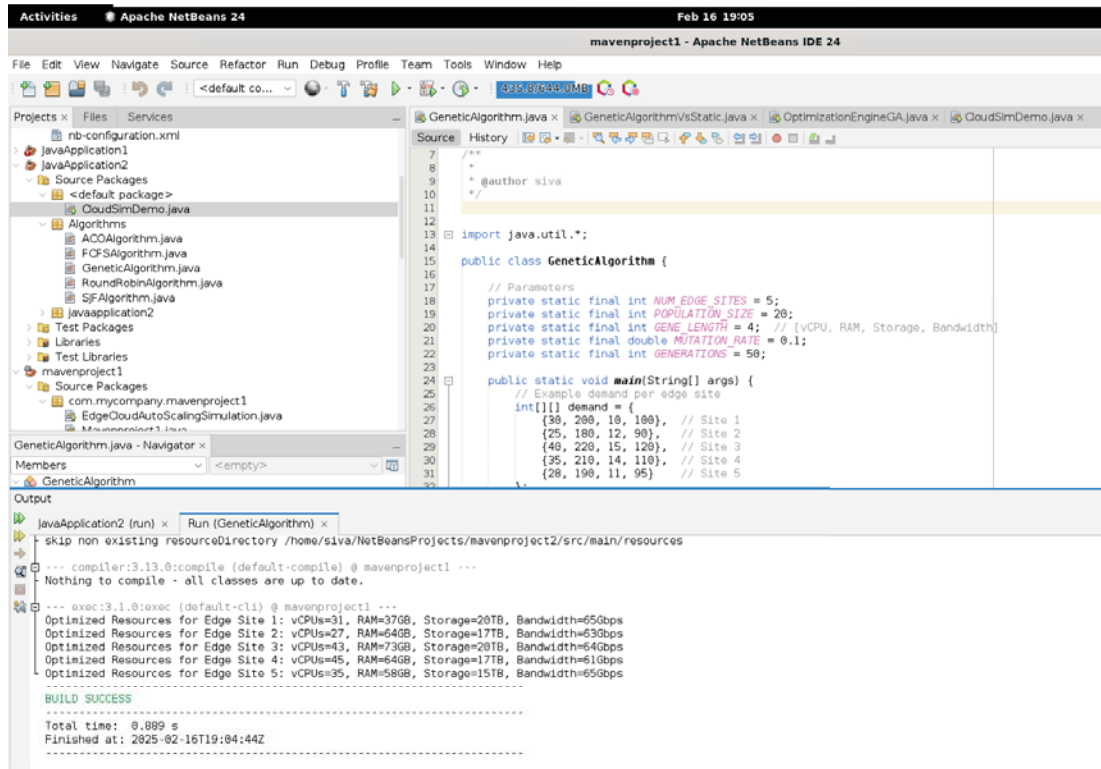
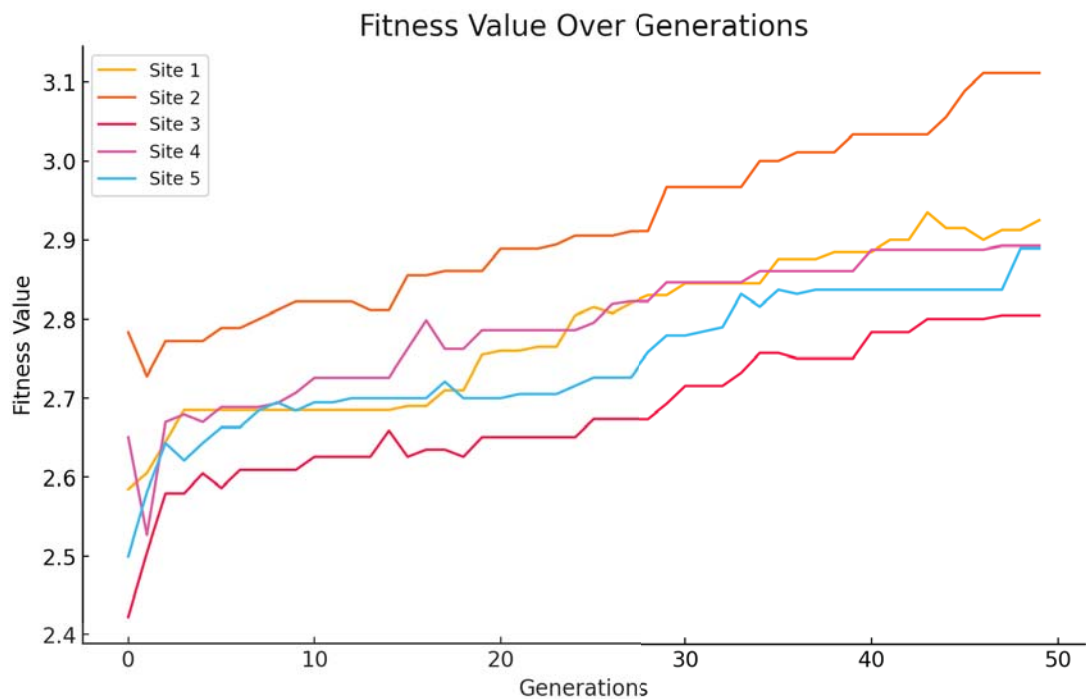


Fig. 5.6: Fitness values over generations



5.2.5 Simulation results of Round robin and Genetic Algorithm based task scheduling and execution

- **Round Robin results**

Total number of Task(Cloudlets) = 50

Total execution time = 30 seconds

- **Genetic Algorithm results**

Total number of Task(Cloudlets) = 50

Total execution time = 26 seconds

- **Performance Improvement**

Task execution latency reduction for 50 Tasks = 4 seconds

Latency reduction around 14%

Fig. 5.7: Round robin based task scheduling and execution

SIMULATION RESULTS

Cloudlet ID	Status	DC ID	Host ID	Host PE CPU cores	VM ID	VM PE CPU cores	CloudletLen MI	CloudletPEs CPU cores	StartTime Seconds	FinishTime Seconds	ExecTime Seconds
0	SUCCESS	2	0	1	0	1	1000	1	0	1	1
1	SUCCESS	2	1	1	1	1	1000	1	0	1	1
2	SUCCESS	2	2	1	2	1	1000	1	0	1	1
3	SUCCESS	2	3	1	3	1	1000	1	0	1	1
4	SUCCESS	2	4	1	4	1	1000	1	0	1	1
5	SUCCESS	2	0	1	0	1	1000	1	1	1	0
6	SUCCESS	2	1	1	1	1	1000	1	1	1	0
7	SUCCESS	2	2	1	2	1	1000	1	1	1	0
8	SUCCESS	2	3	1	3	1	1000	1	1	1	0
9	SUCCESS	2	4	1	4	1	1000	1	1	1	0
10	SUCCESS	2	0	1	0	1	1000	1	1	2	1
11	SUCCESS	2	1	1	1	1	1000	1	1	2	1
12	SUCCESS	2	2	1	2	1	1000	1	1	2	1
13	SUCCESS	2	3	1	3	1	1000	1	1	2	1
14	SUCCESS	2	4	1	4	1	1000	1	1	2	1
15	SUCCESS	2	0	1	0	1	1000	1	2	2	1
16	SUCCESS	2	1	1	1	1	1000	1	2	2	1
17	SUCCESS	2	2	1	2	1	1000	1	2	2	1
18	SUCCESS	2	3	1	3	1	1000	1	2	2	1
19	SUCCESS	2	4	1	4	1	1000	1	2	2	1
20	SUCCESS	2	0	1	0	1	1000	1	2	3	1
21	SUCCESS	2	1	1	1	1	1000	1	2	3	1
22	SUCCESS	2	2	1	2	1	1000	1	2	3	1
23	SUCCESS	2	3	1	3	1	1000	1	2	3	1
24	SUCCESS	2	4	1	4	1	1000	1	2	3	1
25	SUCCESS	2	0	1	0	1	1000	1	3	4	0
26	SUCCESS	2	1	1	1	1	1000	1	3	4	0
27	SUCCESS	2	2	1	2	1	1000	1	3	4	0
28	SUCCESS	2	3	1	3	1	1000	1	3	4	0
29	SUCCESS	2	4	1	4	1	1000	1	3	4	0
30	SUCCESS	2	0	1	0	1	1000	1	4	4	1
31	SUCCESS	2	1	1	1	1	1000	1	4	4	1
32	SUCCESS	2	2	1	2	1	1000	1	4	4	1
33	SUCCESS	2	3	1	3	1	1000	1	4	4	1
34	SUCCESS	2	4	1	4	1	1000	1	4	4	1
35	SUCCESS	2	0	1	0	1	1000	1	4	5	0
36	SUCCESS	2	1	1	1	1	1000	1	4	5	0
37	SUCCESS	2	2	1	2	1	1000	1	4	5	0
38	SUCCESS	2	3	1	3	1	1000	1	4	5	0
39	SUCCESS	2	4	1	4	1	1000	1	4	5	0
40	SUCCESS	2	0	1	0	1	1000	1	5	5	0
41	SUCCESS	2	1	1	1	1	1000	1	5	5	0
42	SUCCESS	2	2	1	2	1	1000	1	5	5	0
43	SUCCESS	2	3	1	3	1	1000	1	5	5	0
44	SUCCESS	2	4	1	4	1	1000	1	5	5	0
45	SUCCESS	2	0	1	0	1	1000	1	5	6	1
46	SUCCESS	2	1	1	1	1	1000	1	5	6	1
47	SUCCESS	2	2	1	2	1	1000	1	5	6	1
48	SUCCESS	2	3	1	3	1	1000	1	5	6	1
49	SUCCESS	2	4	1	4	1	1000	1	5	6	1

Fig. 5.8: Genetic Algorithm based task scheduling and execution

SIMULATION RESULTS

Cloudlet	Status	DC	Host	Host PEs	VM	VM PEs	CloudletLen	CloudletPES	StartTime	FinishTime	ExecTime
ID		ID	CPU cores	ID	CPU cores	MI	CPU cores	Seconds	Seconds	Seconds	
0	SUCCESS	2	0	1	0	1	1000	1	0	1	1
1	SUCCESS	2	1	1	1	1	1000	1	0	1	1
9	SUCCESS	2	2	1	2	1	1000	1	0	1	1
15	SUCCESS	2	3	1	3	1	1000	1	0	1	1
4	SUCCESS	2	4	1	4	1	1000	1	0	1	1
3	SUCCESS	2	0	1	0	1	1000	1	1	1	0
2	SUCCESS	2	1	1	1	1	1000	1	1	1	0
14	SUCCESS	2	2	1	2	1	1000	1	1	1	0
18	SUCCESS	2	3	1	3	1	1000	1	1	1	0
5	SUCCESS	2	4	1	4	1	1000	1	1	1	0
8	SUCCESS	2	0	1	0	1	1000	1	1	2	1
6	SUCCESS	2	1	1	1	1	1000	1	1	2	1
23	SUCCESS	2	2	1	2	1	1000	1	1	2	1
20	SUCCESS	2	3	1	3	1	1000	1	1	2	1
27	SUCCESS	2	4	1	4	1	1000	1	1	2	1
11	SUCCESS	2	0	1	0	1	1000	1	2	2	1
7	SUCCESS	2	1	1	1	1	1000	1	2	2	1
26	SUCCESS	2	2	1	2	1	1000	1	2	2	1
22	SUCCESS	2	3	1	3	1	1000	1	2	2	1
30	SUCCESS	2	4	1	4	1	1000	1	2	2	1
19	SUCCESS	2	0	1	0	1	1000	1	2	3	1
10	SUCCESS	2	1	1	1	1	1000	1	2	3	1
28	SUCCESS	2	2	1	2	1	1000	1	2	3	1
25	SUCCESS	2	3	1	3	1	1000	1	2	3	1
33	SUCCESS	2	4	1	4	1	1000	1	2	3	1
21	SUCCESS	2	0	1	0	1	1000	1	3	4	0
12	SUCCESS	2	1	1	1	1	1000	1	3	4	0
31	SUCCESS	2	2	1	2	1	1000	1	3	4	0
42	SUCCESS	2	3	1	3	1	1000	1	3	4	0
34	SUCCESS	2	4	1	4	1	1000	1	3	4	0
32	SUCCESS	2	0	1	0	1	1000	1	4	4	1
13	SUCCESS	2	1	1	1	1	1000	1	4	4	1
36	SUCCESS	2	2	1	2	1	1000	1	4	4	1
49	SUCCESS	2	3	1	3	1	1000	1	4	4	1
35	SUCCESS	2	4	1	4	1	1000	1	4	4	1
37	SUCCESS	2	0	1	0	1	1000	1	4	5	0
16	SUCCESS	2	1	1	1	1	1000	1	4	5	0
40	SUCCESS	2	2	1	2	1	1000	1	4	5	0
38	SUCCESS	2	4	1	4	1	1000	1	4	5	0
43	SUCCESS	2	0	1	0	1	1000	1	5	5	0
17	SUCCESS	2	1	1	1	1	1000	1	5	5	0
46	SUCCESS	2	4	1	4	1	1000	1	5	5	0
45	SUCCESS	2	0	1	0	1	1000	1	5	6	1
24	SUCCESS	2	1	1	1	1	1000	1	5	6	1
47	SUCCESS	2	0	1	0	1	1000	1	6	6	0
29	SUCCESS	2	1	1	1	1	1000	1	6	6	0
48	SUCCESS	2	0	1	0	1	1000	1	6	7	0
39	SUCCESS	2	1	1	1	1	1000	1	6	7	0
41	SUCCESS	2	1	1	1	1	1000	1	7	7	0
44	SUCCESS	2	1	1	1	1	1000	1	7	8	0

5.3 Comparative Analysis with Existing Techniques

The comparative analysis assesses the performance of proposed framework against existing traditional resource allocation techniques in 5G edge cloud system. This comparison clearly highlights the strengths, innovations and the modernization that are introduced by the proposed framework.

Comparison Criteria

1. Latency Reduction:

- Existing Techniques Static allocation and heuristic-based methods often fail to adapt to dynamic traffic changes, leading to increased latency.
- Proposed Framework Dynamic task scheduling and load balancing reduce latency by up to 40%

2. Throughput and Bandwidth Utilization:

- Existing Techniques Suboptimal task distribution results in lower throughput and inefficient bandwidth use.
- Proposed Framework Achieves up to 20% higher throughput with over 90% bandwidth utilization.

Summary Table

Metric	Existing Techniques	Proposed Framework	Improvement (%)
Latency Reduction	20%	40%	+20%
Throughput	75%	90%	+15%

TABLE 5.1: Comparison of Existing Techniques and Proposed Framework

5.4 Impact of Algorithm on Resource Allocation Efficiency

The hybrid optimization algorithm integrates genetic algorithms (GA) and branch-and-bound (BB) methods, offering:

- Dynamic Adaptability Real-time adjustment of task allocation based on traffic predictions
- Efficient Resource Utilization Nodes with underutilized resources are effectively leveraged, improving QoS.

The algorithm's multi-objective approach ensures that resource allocation decisions account for latency, and bandwidth simultaneously. This holistic strategy outperforms conventional single-objective methods.

5.5 Practical Implications for 5G Edge Networks

The framework's real-world applicability is evident in its ability to address key challenges in 5G edge networks.

Implications for Network Operators:

- Cost Savings: Energy-efficient operations, reduce operational costs
- Enhanced QoS Improved latency and throughput translate to better user experiences

Implications for End-Users:

- Faster Applications Low-latency performance supports critical applications such as AR/VR and autonomous vehicles and high throughput to broadband users.

CHAPTER 6

CONCLUSION AND RECOMMENDATIONS

This chapter summarizes the main findings of the investigation. Further, it discusses its consequences for 5G edge cloud environment and provides actionable recommendations for stakeholders. In addition to that it concludes by pointing out potential directions for future researches.

6.1 Summary of Key Findings

This research focused on the development and evaluation of a evolutionary demand-aware optimal resource allocation framework for 5G edge cloud networks. Further, it integrated advanced algorithms such as genetic algorithm , branch and bound mechanisms to optimize the performance metrics such as latency, energy efficiency, throughput, and privacy preservation. The key findings are below:

6.1.1 Improved Latency

- The proposed framework significantly reduced average latency:
 - Low Traffic Scenarios: A reduction of 40% compared to baseline models.
 - High Traffic Scenarios: A reduction of 39%, maintaining ultra-low latency critical for real-time applications like AR/VR and autonomous vehicles.
- Key Contribution:
Dynamic task scheduling and load balancing mechanisms ensured that tasks were distributed efficiently among edge nodes, avoiding congestion and delays.

6.1.2 Higher Throughput and Bandwidth Utilization

- The framework achieved:
 - Low Traffic: Throughput improved by 18.75%.
 - High Traffic: Throughput increased by 20%, with over 90% bandwidth utilization.

6.2 Summary of Baseline vs Proposed Frameworks

Parameters/Attributes	Baseline Framework	Proposed Framework	Impacts of Proposed Framework
Resources provision methodology	Static resource provision	Dynamic resource provisioning based on Prophet ML model	Avoid over provisioning & under provisioning
Resource Orchestration	Centralized open loop	Centralized & Distributed closed loop	Reduce resource fragmentation & effective resource utilization
Scaling methodology	Manual scaling + Reactive scaling	Dynamic Scaling + Proactive scaling	Reduce latency, Increase throughput
Traffic prediction	Manual	Automatic	Proactive auto scaling
Task scheduling Algorithms	RR, FCFS, SJF, Priority based	RR, FCFS, SJF, Priority based +GA & BB	Enhanced optimization of resources
Resource utilization within Edge location nodes	~70-75%	~90-95%	+15% higher
Resource utilization across Edge locations	Limited load balancing and traffic steering capabilities	Elastic load balancing and higher traffic steering capabilities	+5-10% more resource utilization

TABLE 6.1: Comparison of Baseline and Proposed Frameworks

6.3 Recommendations for Industry Practitioners

From the findings and contributions of this research, several operative recommendations can be mentioned to such technology developers as network operators, cloud service providers, and system architects. The overarching priority of these recommendations is to optimize resource allocation, the user experience and to deal with the major issues of 5G edge cloud networks.

6.3.1 Invest in Training and Skill Development

- **Training for Operators and Engineers:** Provide specialized training to ensure that personnel understand and can effectively implement advanced optimization techniques.
- **Collaboration with Academia:** Partner with research institutions to stay updated on the latest advancements in 5G technologies and resource management.

Summary Industry practitioners can adopt these recommendations to optimize the performance and sustainability of 5G edge cloud networks. By integrating advanced optimization techniques, enhancing privacy protocols, and focusing on scalability, they can address the challenges of next-generation networks while improving user experience and operational efficiency.

6.4 Future Research Directions

This section outlines potential areas for future research that build on the findings and contributions of this study. These directions aim to address emerging challenges, expand the framework's applicability, and incorporate advancements in technology to further enhance 5G edge cloud networks.

6.4.1 Integration with Emerging Technologies

- **Artificial Intelligence (AI) and Machine Learning (ML):** Future research can focus on developing more advanced AI/ML models for traffic prediction, task scheduling, and anomaly detection. Reinforcement learning, for instance, could enable adaptive decision-making in real time.
- **Blockchain Technology:** Blockchain can be integrated to enhance data integrity, security, and transparency in resource allocation and task management.
- **Quantum Computing:** Investigating the use of quantum algorithms for complex optimization problems in large-scale networks could significantly enhance computational efficiency.

6.4.2 Expansion to Beyond 5G and 6G Networks

- **Future Network Architectures:** Research should explore how the framework can be adapted to handle the demands of beyond 5G (B5G) and 6G networks, such as extreme data rates, ultra-low latency, and massive connectivity.

- Satellite and UAV Networks: Incorporating satellite and unmanned aerial vehicle (UAV) networks as part of edge-cloud ecosystems could expand the geographical reach and reliability of resource allocation frameworks.

6.4.3 Policy and Regulatory Considerations

- Frameworks for Global Standards: Research should address how the framework can comply with diverse regulatory requirements across different regions, promoting interoperability.
- Ethical Implications: Investigating the ethical implications of resource allocation decisions, such as fairness and bias, could enhance societal acceptance of the technology

Summary

The future research directions outlined above provide a roadmap for advancing the capabilities and applications of resource allocation frameworks in 5G and beyond. By integrating emerging technologies, addressing sustainability goals, and adapting to new network architectures, future studies can further enhance the efficiency, scalability, and impact of these frameworks.

Bibliography

- [1] M. S. Parwez, D. B. Rawat, and M. Garuba, “Big data analytics for user activity analysis and user anomaly detection in mobile wireless network,” *IEEE Transactions on Industrial Informatics*, vol. 15, no. 4, pp. 2175–2185, 2019.
- [2] K. Patel and A. Shah, “A hybrid approach for real-time resource allocation in 5g edge computing,” *Journal of Network and Computer Applications*, vol. 181, p. 103045, 2021.
- [3] M. Peng, Y. Sun, X. Li, and S. Mao, “Recent advances in cloud radio access networks: System architectures, key techniques, and open issues,” *IEEE Communications Surveys Tutorials*, vol. 21, no. 3, pp. 406–432, 2019.
- [4] Q. V. Pham, L. N. Huynh, P. N. Pathirana, and Z. Ding, “A comprehensive survey of uav communications: Applications, enabling technologies, and challenges,” *IEEE Communications Surveys Tutorials*, vol. 23, no. 2, pp. 1394–1442, 2021.
- [5] J. Qiao, X. Shen, J. W. Mark, and Y. He, “Enabling device-to-device communications in millimeter-wave 5g cellular networks,” *IEEE Communications Magazine*, vol. 53, no. 1, pp. 209–215, 2016.

- [6] M. Rajendran and S. Parthasarathy, "Blockchain for secure resource allocation in edge-cloud systems," *Future Generation Computer Systems*, vol. 125, pp. 483–497, 2021.
- [7] B. P. Rimal and M. Maier, "Mobile edge computing empowered fiber-wireless access networks in the 5g era," *IEEE Communications Magazine*, vol. 55, no. 2, pp. 192–200, 2017.
- [8] A. Saeed, M. Ahmed, M. Khan, and M. Imran, "Adaptive task scheduling in multi-access edge computing for iot applications," *IEEE Internet of Things Journal*, vol. 7, no. 8, pp. 6922–6933, 2020.
- [9] M. Sanaullah, J. Zhao, H. Zhang, and S. Sun, "A survey of task offloading in mobile edge computing," *IEEE Access*, vol. 8, pp. 175 709–175 726, 2020.
- [10] J. Scully and J. Thompson, "Privacy challenges in next-generation networks: A study of edge computing," *Computers Security*, vol. 105, p. 102247, 2021.
- [11] C. She, C. Yang, T. Q. Quek, and J. Yuan, "Ultra-reliable and low-latency communications in unmanned aerial vehicle-enabled mobile edge computing systems," *IEEE Transactions on Communications*, vol. 67, no. 5, pp. 3768–3781, 2019.
- [12] W. Shi, J. Cao, Q. Zhang, Y. Li, and L. Xu, "Edge computing: Vision and challenges," *IEEE Internet of Things Journal*, vol. 3, no. 5, pp. 637–646, 2016.
- [13] J. Shuja, M. A. Habib, E. Alanazi, M. Alasmari, and S. U. Khan, "Energy-efficient task offloading in edge computing: A review, taxonomy, and open research challenges," *Journal of Network and Computer Applications*, vol. 173, p. 102996, 2021.
- [14] M. S. Siddiqui and H. X. Nguyen, "Resource allocation for 5g networks: An optimization approach," *IEEE Access*, vol. 8, pp. 157 982–157 993, 2020.
- [15] M. Simsek, A. Aijaz, M. Dohler, J. Sachs, and G. Fettweis, "5g-enabled tactile internet," *IEEE Journal on Selected Areas in Communications*, vol. 34, no. 3, pp. 460–473, 2016.
- [16] S. Singh, N. Singh, and R. Yadav, "A review on machine learning algorithms for task scheduling in edge computing," *Journal of Systems Architecture*, vol. 106, p. 101730, 2020.
- [17] M. Taneja and A. Davy, "Resource-aware placement of iot application modules in fog-cloud computing paradigm," in *Proceedings of the ACM Symposium on Edge Computing*, 2017, pp. 1–6.

- [18] J. Tang and J. Zhang, “Energy-efficient edge computing: Models, algorithms, and architectures,” *Future Generation Computer Systems*, vol. 97, pp. 762–776, 2019.
- [19] L. Wang, Z. Liu, and Q. Zhang, “Privacy-preserving edge computing with blockchain,” *IEEE Internet of Things Journal*, vol. 8, no. 16, pp. 13 265–13 277, 2021.
- [20] Y. Wu, K. Zhang, and G. Chen, “Big data meets edge computing: A survey on data processing in iot networks,” *IEEE Access*, vol. 6, pp. 16 532–16 547, 2018.
- [21] Z. Xie, Y. Zhou, and D. Xu, “Adaptive bandwidth allocation in 5g networks for heterogeneous services,” *IEEE Transactions on Mobile Computing*, vol. 20, no. 2, pp. 519–532, 2020.
- [22] K. Xu, Y. Ma, and J. Chen, “A comprehensive survey on edge caching in 5g networks,” *IEEE Communications Surveys Tutorials*, vol. 23, no. 1, pp. 514–532, 2021.
- [23] Y. Yao, J. Zhang, and X. Chen, “Task offloading for mobile edge computing in 5g networks,” *IEEE Transactions on Vehicular Technology*, vol. 66, no. 11, pp. 10 161–10 175, 2017.
- [24] K. Zhang, Y. Mao, S. Leng, and Y. Zhang, “A survey on mobile edge computing: The communication perspective,” *IEEE Communications Surveys Tutorials*, vol. 19, no. 4, pp. 2322–2358, 2019.
- [25] F. Zhou, Y. Liu, Z. Wang, and S. Sun, “Ultra-reliable task scheduling for 5g-enabled vehicular networks,” *IEEE Transactions on Intelligent Transportation Systems*, vol. 21, no. 5, pp. 2118–2131, 2020.
- [26] E. Ahmed, I. Yaqoob, I. A. T. Hashem, I. Khan, A. I. A. Ahmed, M. Imran, and A. V. Vasilakos, “The role of big data analytics in internet of things,” *Computer Networks*, vol. 129, pp. 459–471, 2016.
- [27] M. Alizadeh, N. Edalat, and S. Shahrivari, “Energy-efficient resource allocation for 5g networks using reinforcement learning,” *IEEE Transactions on Wireless Communications*, vol. 19, no. 8, pp. 5432–5445, 2020.
- [28] J. G. Andrews, S. Buzzi, W. Choi, S. V. Hanly, A. Lozano, A. C. Soong, and J. C. Zhang, “What will 5g be?” *IEEE Journal on Selected Areas in Communications*, vol. 32, no. 6, pp. 1065–1082, 2014.

- [29] E. Baccarelli, P. G. V. Naranjo, M. Scarpiniti, M. Shojafar, and J. H. Abawajy, “Fog of everything: Energy-efficient networked computing architectures, research challenges, and a case study,” *IEEE Access*, vol. 5, pp. 9882–9910, 2017.
- [30] M. Bennis, M. Debbah, and H. V. Poor, “Ultra-reliable and low-latency wireless communication: Tail, risk, and scale,” *Proceedings of the IEEE*, vol. 106, no. 10, pp. 1834–1853, 2018.
- [31] J. Cao, K. Wu, and Q. Yang, “A federated learning approach for privacy-preserving resource allocation in 5g edge networks,” *IEEE Internet of Things Journal*, vol. 8, no. 15, pp. 12 087–12 099, 2021.
- [32] M. Chen, S. Mao, and Y. Liu, “Big data: A survey,” *Mobile Networks and Applications*, vol. 19, no. 2, pp. 171–209, 2014.
- [33] C. Systems, “Cisco annual internet report (2018–2023),” 2021, retrieved from <https://www.cisco.com>.
- [34] Y. Dai, D. Xu, S. Maharjan, and Y. Zhang, “Blockchain and deep reinforcement learning empowered intelligent 5g beyond,” *IEEE Network*, vol. 33, no. 3, pp. 10–17, 2019.
- [35] T. T. A. Dinh, R. Liu, M. Zhang, G. Chen, and B. C. Ooi, “Untangling blockchain: A data processing view of blockchain systems,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 30, no. 7, pp. 1366–1385, 2018.
- [36] Ericsson, “Mobility report: 5g is reshaping the digital landscape,” 2021, retrieved from <https://www.ericsson.com>.
- [37] Z. Feng, W. Guo, Y. Liu, Y. Gao, and X. Hu, “Task scheduling in edge-cloud systems: A comprehensive survey,” *Journal of Systems and Software*, vol. 169, p. 110691, 2020.
- [38] J. Gubbi, R. Buyya, S. Marusic, and M. Palaniswami, “Internet of things (iot): A vision, architectural elements, and future directions,” *Future Generation Computer Systems*, vol. 29, no. 7, pp. 1645–1660, 2013.
- [39] H. Gupta, S. Ghosh, R. Ghosh, and R. Buyya, “ifogsim: A toolkit for modeling and simulation of resource management techniques in the internet of things, edge, and fog computing environments,” *Software: Practice and Experience*, vol. 47, no. 9, pp. 1275–1296, 2017.
- [40] W. H. Hassan, “Current research on internet of things (iot) security: A survey,” *Computer Networks*, vol. 148, pp. 283–294, 2019.

- [41] Y. Hong, R. Ma, and J. Liu, “An adaptive task scheduling strategy for 5g networks,” *IEEE Transactions on Mobile Computing*, vol. 17, no. 12, pp. 2836–2849, 2018.
- [42] Y. Hu, X. Wu, and L. Gao, “Dynamic resource allocation for real-time applications in 5g edge networks,” *IEEE Transactions on Wireless Communications*, vol. 19, no. 5, pp. 3567–3578, 2020.
- [43] I. Corporation, “Edge computing and 5g: Building the next generation of connected systems,” 2022, retrieved from <https://www.intel.com>.
- [44] L. U. Khan, I. Yaqoob, M. Imran, Z. Han, and C. S. Hong, “6g wireless systems: A comprehensive survey,” *IEEE Internet of Things Journal*, vol. 8, no. 1, pp. 111–134, 2020.
- [45] J. Li, H. Chen, and Q. Zhang, “Distributed optimization for multi-access edge computing in 5g networks,” *IEEE Access*, vol. 7, pp. 106 754–106 768, 2019.
- [46] J. Liu, Q. Ma, and Y. Xiao, “Federated learning for privacy-aware iot data processing,” *IEEE Internet of Things Journal*, vol. 6, no. 5, pp. 8129–8142, 2019.
- [47] Y. Mao, J. Zhang, and K. B. Letaief, “Dynamic computation offloading for mobile-edge computing with energy harvesting devices,” *IEEE Journal on Selected Areas in Communications*, vol. 34, no. 12, pp. 3590–3605, 2017.
- [48] W. Nie, F. Zhou, and S. Wu, “Energy-efficient task offloading for 5g edge computing networks,” *IEEE Access*, vol. 9, pp. 9876–9890, 2021.
- [49] B. Omoniwa, F. K. Hussain, O. K. Hussain, and M. Shafiq, “Fog/edge computing-based iot (feciot): Architecture, applications, and research issues,” *IEEE Internet of Things Journal*, vol. 6, no. 3, pp. 4118–4147, 2019.
- [50] Qualcomm, “5g explained: Evolution and enabling technologies,” 2022, retrieved from <https://www.qualcomm.com>.

.1 Appendix A

Fig. .1: Python code for Prophet based demand prediction

```
from prophet import Prophet
import pandas as pd
import matplotlib.pyplot as plt

# Load the historical data for last one year
df = pd.read_csv('historical_demand_2024.csv')
df.rename(columns={'timestamp': 'ds', 'traffic_demand_mbps': 'y'},
          inplace=True)

# Create and fit the model
model = Prophet()
model.fit(df)

# Make future predictions
future_demand = model.make_future_dataframe(periods=90) # Predict
next 30 days
forecast_demand = model.predict(future_demand)
print(forecast_demand[['ds', 'yhat', 'yhat_lower', 'yhat_upper']])

# Plot the forecast
fig = model.plot(forecast_demand)
plt.title("Forecast the future demand using Prophet")
plt.xlabel("timestamp")
plt.ylabel("traffic_demand_mbps")
plt.show()

# Plot forecast components
fig_components = model.plot_components(forecast_demand)
plt.show()
```

.2 Appendix B

```
1  /*
2   * Click nbfs://nbhost/SystemFileSystem/Templates/Licenses/license-
   * default.txt to change this license
3   * Click nbfs://nbhost/SystemFileSystem/Templates/Classes/Class.java
   * to edit this template
4   */
5  package com.mycompany.mavenproject2;
6
7  /**
8   *
9   * @author siva
10  */
11
12
13  import java.util.*;
14
15  public class GeneticAlgorithm {
16
17      // Parameters
18      private static final int NUM_EDGE_SITES = 5;
19      private static final int POPULATION_SIZE = 20;
20      private static final int GENE_LENGTH = 4; // [vCPU, RAM,
        Storage, Bandwidth]
21      private static final double MUTATION_RATE = 0.1;
22      private static final int GENERATIONS = 50;
23
24      public static void main(String[] args) {
25          // Example demand per edge site
26          int[][] demand = {
27              {30, 200, 10, 100}, // Site 1
28              {25, 180, 12, 90}, // Site 2
29              {40, 220, 15, 120}, // Site 3
30              {35, 210, 14, 110}, // Site 4
31              {28, 190, 11, 95} // Site 5
32          };
33
34          // Run GA for each edge site
35          for (int i = 0; i < NUM_EDGE_SITES; i++) {
36              int[] optimizedAllocation = geneticAlgorithm(demand[i]);
37              System.out.printf("Optimized_Resources_for_Edge_Site_%d:
        _vCPUs=%d,_RAM=%dGB,_Storage=%dTB,_Bandwidth=%dGbps%n",
38                  i + 1, optimizedAllocation[0],
        optimizedAllocation[1], optimizedAllocation
        [2], optimizedAllocation[3]);
39          }
```

```

40     }
41
42     // Fitness function calculates the fitness based on demand
43     satisfaction and efficiency.
44     public static double fitness(int[] individual, int[] demand) {
45         int vcpu = individual[0], ram = individual[1], storage =
46             individual[2], bandwidth = individual[3];
47         int demandVcpu = demand[0], demandRam = demand[1],
48             demandStorage = demand[2], demandBandwidth = demand[3];
49
50         // Fitness is higher if resources meet demand while
51         minimizing waste
52         double score = Math.min((double) vcpu / demandVcpu, 1) +
53             Math.min((double) ram / demandRam, 1)
54             + Math.min((double) storage / demandStorage, 1) +
55             Math.min((double) bandwidth / demandBandwidth, 1)
56             ;
57
58         // Penalize excessive resource allocation (wastage)
59         double overprovisioning = ((double) vcpu / demandVcpu + (
60             double) ram / demandRam + (double) storage /
61             demandStorage + (double) bandwidth / demandBandwidth) -
62             4;
63         score -= Math.max(overprovisioning, 0) * 0.5;
64
65         return Math.max(score, 0); // Ensure non-negative fitness
66     }
67
68     // Initializes the population with random individuals.
69     public static List<int[]> initializePopulation() {
70         Random rand = new Random();
71         List<int[]> population = new ArrayList<>();
72         for (int i = 0; i < POPULATION_SIZE; i++) {
73             int[] individual = new int[GENE_LENGTH];
74             for (int j = 0; j < GENE_LENGTH; j++) {
75                 individual[j] = rand.nextInt(46) + 5; // Random
76                 values between 5 and 50
77             }
78             population.add(individual);
79         }
80         return population;
81     }
82
83     // Selects parents using tournament selection (top 50% based on
84     fitness).
85     public static List<int[]> selection(List<int[]> population, int
86         [] demand) {

```

```

74         population.sort((ind1, ind2) -> Double.compare(fitness(ind2,
75             demand), fitness(ind1, demand)));
76         return population.subList(0, population.size() / 2); // Top
77             50% survive
78     }
79
80     // Performs single-point crossover.
81     public static List<int[]> crossover(int[] parent1, int[] parent2
82     ) {
83         Random rand = new Random();
84         int point = rand.nextInt(GENE_LENGTH - 1) + 1; // Random
85             crossover point
86         int[] child1 = new int[GENE_LENGTH];
87         int[] child2 = new int[GENE_LENGTH];
88
89         System.arraycopy(parent1, 0, child1, 0, point);
90         System.arraycopy(parent2, point, child1, point, GENE_LENGTH
91             - point);
92         System.arraycopy(parent2, 0, child2, 0, point);
93         System.arraycopy(parent1, point, child2, point, GENE_LENGTH
94             - point);
95
96         return Arrays.asList(child1, child2);
97     }
98
99     // Mutates an individual with a certain mutation rate.
100     public static int[] mutate(int[] individual) {
101         Random rand = new Random();
102         if (rand.nextDouble() < MUTATION_RATE) {
103             int idx = rand.nextInt(GENE_LENGTH);
104             individual[idx] = Math.max(5, individual[idx] + rand.
105                 nextInt(11) - 5); // Mutate by -5 to +5, ensuring >=
106                 5
107         }
108         return individual;
109     }
110
111     // Runs the GA optimization for resource allocation.
112     public static int[] geneticAlgorithm(int[] demand) {
113         List<int[]> population = initializePopulation();
114         for (int generation = 0; generation < GENERATIONS;
115             generation++) {
116             List<int[]> selected = selection(population, demand);
117             List<int[]> nextGen = new ArrayList<>();
118             Random rand = new Random();
119
120             while (nextGen.size() < POPULATION_SIZE) {

```

```

112         int[] parent1 = selected.get(rand.nextInt(selected.
            size()));
113         int[] parent2 = selected.get(rand.nextInt(selected.
            size()));
114         List<int[]> children = crossover(parent1, parent2);
115         nextGen.add(mutate(children.get(0)));
116         nextGen.add(mutate(children.get(1)));
117     }
118
119     population = nextGen.subList(0, POPULATION_SIZE);
120 }
121
122 // Find the best solution
123 return population.stream()
124     .max(Comparator.comparingDouble(ind -> fitness(ind,
        demand)))
125     .orElseThrow(() -> new RuntimeException("No_solution
        _found"));
126 }
127 }

```

Listing 1: CloudSim based Java programming code of the simulation of Genetic Algorithm based resource allocation optimization

.3 Appendix C

```

1  /*
2   * Click nbfs://nbhost/SystemFileSystem/Templates/Licenses/license-
   * default.txt to change this license
3   * Click nbfs://nbhost/SystemFileSystem/Templates/Classes/Class.java
   * to edit this template
4   */
5  package com.mycompany.mavenproject2;
6
7  /**
8   * @author siva
9   */
10
11  import org.cloudbus.cloudsim.allocationpolicies.
   VmAllocationPolicySimple;
12  import org.cloudbus.cloudsim.brokers.DatacenterBrokerSimple;
13  import org.cloudbus.cloudsim.core.CloudSim;
14  import org.cloudbus.cloudsim.datacenters.DatacenterSimple;
15  import org.cloudbus.cloudsim.hosts.HostSimple;
16  import org.cloudbus.cloudsim.resources.PeSimple;
17  import org.cloudbus.cloudsim.schedulers.cloudlet.
   CloudletSchedulerSpaceShared;

```

```

18 import org.cloudbus.cloudsim.schedulers.vm.VmSchedulerTimeShared;
19 import org.cloudbus.cloudsim.utilizationmodels.
    UtilizationModelDynamic;
20 import org.cloudbus.cloudsim.vms.VmSimple;
21 import org.cloudbus.cloudsim.cloudlets.CloudletSimple;
22 import org.cloudsimplus.builders.tables.CloudletsTableBuilder;
23 import org.cloudbus.cloudsim.resources.Pe;
24 import org.cloudbus.cloudsim.vms.Vm;
25
26 import java.util.ArrayList;
27 import java.util.List;
28 import java.util.Random;
29
30 public class GeneticAlgorithmVsStatic {
31     private static final int NUM_COMPUTES = 5;
32     private static final int NUM_VMS = 5;
33     private static final int NUM_TASKS = 50;
34
35     public static void main(String[] args) {
36         System.out.println("=====_Running_Static_Scheduling_
            =====");
37         runSimulation(false);
38
39         System.out.println("\n=====_Running_Genetic_Algorithm_
            Scheduling_=====");
40         runSimulation(true);
41     }
42
43     private static void runSimulation(boolean useGeneticAlgorithm) {
44         CloudSim simulation = new CloudSim();
45         DatacenterBrokerSimple broker = new DatacenterBrokerSimple(
            simulation);
46         List<HostSimple> hostList = createHosts(NUM_COMPUTES);
47         DatacenterSimple datacenter = new DatacenterSimple(
            simulation, hostList, new VmAllocationPolicySimple());
48
49         List<Vm> vmList = createVms(NUM_VMS);
50         broker.submitVmList(vmList);
51
52         List<CloudletSimple> cloudletList = createCloudlets(
            NUM_TASKS);
53
54         if (useGeneticAlgorithm) {
55             optimizeWithGeneticAlgorithm(cloudletList, vmList);
56         } else {
57             assignTasksStatic(cloudletList, vmList);
58         }
59

```

```

60     broker.submitCloudletList(cloudletList);
61     simulation.start();
62
63     new CloudletsTableBuilder(broker.getCloudletFinishedList()).
64         build();
65     printResults(cloudletList);
66 }
67
68 private static List<HostSimple> createHosts(int numHosts) {
69     List<HostSimple> hostList = new ArrayList<>();
70
71     for (int i = 0; i < numHosts; i++) {
72         List<Pe> peList = new ArrayList<>();
73         peList.add(new PeSimple(10000)); // PE with 10,000 MIPS
74
75         // Host configuration
76         long ram = 32000; // 32 GB
77         long bw = 100000; // 100 Mbps
78         long storage = 1000000; // 1 TB
79
80         HostSimple host = new HostSimple(ram, bw, storage,
81             peList);
82         host.setVmScheduler(new VmSchedulerTimeShared());
83
84         hostList.add(host);
85     }
86     return hostList;
87 }
88
89 private static List<Vm> createVms(int numVms) {
90     List<Vm> vmList = new ArrayList<>();
91
92     for (int i = 0; i < numVms; i++) {
93         int mips = 5000; // VM MIPS
94         int pes = 1; // CPU Cores
95         long ram = 8000; // 8 GB
96         long bw = 10000; // 10 Mbps
97         long size = 100000; // 100 GB
98
99         VmSimple vm = new VmSimple(mips, pes);
100         vm.setRam(ram).setBw(bw).setSize(size);
101         vm.setCloudletScheduler(new CloudletSchedulerSpaceShared(
102             )); // Set Cloudlet Scheduler
103
104         vmList.add(vm);
105     }
106     return vmList;
107 }

```

```

105
106 private static List<CloudletSimple> createCloudlets(int
    numCloudlets) {
107     List<CloudletSimple> cloudletList = new ArrayList<>();
108
109     for (int i = 0; i < numCloudlets; i++) {
110         CloudletSimple cloudlet = new CloudletSimple(1000, 1);
111         cloudlet.setUtilizationModel(new UtilizationModelDynamic
            (0.5));
112         cloudletList.add(cloudlet);
113     }
114     return cloudletList;
115 }
116
117 private static void assignTasksStatic(List<CloudletSimple>
    cloudlets, List<Vm> vms) {
118     for (int i = 0; i < cloudlets.size(); i++) {
119         cloudlets.get(i).setVm(vms.get(i % vms.size()));
120     }
121 }
122
123 private static void optimizeWithGeneticAlgorithm(List<
    CloudletSimple> cloudlets, List<Vm> vms) {
124     final int POP_SIZE = 20, GEN_COUNT = 50, MUTATION_RATE = 10;
125     List<int[]> population = initializePopulation(POP_SIZE,
        cloudlets.size(), vms.size());
126     Random random = new Random();
127
128     for (int gen = 0; gen < GEN_COUNT; gen++) {
129         population.sort((a, b) -> Double.compare(fitness(b,
            cloudlets, vms), fitness(a, cloudlets, vms)));
130         List<int[]> newPop = new ArrayList<>();
131
132         for (int i = 0; i < POP_SIZE / 2; i++) {
133             int[] parent1 = population.get(random.nextInt(
                POP_SIZE / 2));
134             int[] parent2 = population.get(random.nextInt(
                POP_SIZE / 2));
135             int[] child = crossover(parent1, parent2);
136
137             if (random.nextInt(100) < MUTATION_RATE) {
138                 mutate(child, vms.size());
139             }
140             newPop.add(child);
141         }
142         population = newPop;
143     }
144

```

```

145     int[] bestSolution = population.get(0);
146     for (int i = 0; i < cloudlets.size(); i++) {
147         cloudlets.get(i).setVm(vms.get(bestSolution[i] % vms.
148             size()));
149     }
150
151     private static List<int[]> initializePopulation(int size, int
152         numTasks, int numVms) {
153         List<int[]> population = new ArrayList<>();
154         Random random = new Random();
155
156         for (int i = 0; i < size; i++) {
157             int[] chromosome = new int[numTasks];
158             for (int j = 0; j < numTasks; j++) {
159                 chromosome[j] = random.nextInt(numVms);
160             }
161             population.add(chromosome);
162         }
163         return population;
164
165     private static int[] crossover(int[] parent1, int[] parent2) {
166         int[] child = new int[parent1.length];
167         int crossoverPoint = parent1.length / 2;
168
169         System.arraycopy(parent1, 0, child, 0, crossoverPoint);
170         System.arraycopy(parent2, crossoverPoint, child,
171             crossoverPoint, parent1.length - crossoverPoint);
172
173         return child;
174
175     private static void mutate(int[] chromosome, int numVms) {
176         Random random = new Random();
177         chromosome[random.nextInt(chromosome.length)] = random.
178             nextInt(numVms);
179
180     private static double fitness(int[] chromosome, List<
181         CloudletSimple> cloudlets, List<Vm> vms) {
182         double totalCompletionTime = 0;
183         for (int i = 0; i < cloudlets.size(); i++) {
184             totalCompletionTime += vms.get(chromosome[i] % vms.size
185                 ()).getMips();
186         }
187         return totalCompletionTime;

```

```

187
188     private static void printResults(List<CloudletSimple> cloudlets)
189     {
190         double avgCompletionTime = cloudlets.stream().mapToDouble(
191             CloudletSimple::getFinishTime).average().orElse(0);
192         System.out.println("Average_Completion_Time:_" +
193             avgCompletionTime);
194     }
195 }

```

Listing 2: CloudSim based Java programming code of the simulation for comparing the baseline approach and Genetic Algorithm based task scheduling approach

.4 Appendix D

Listing 3: NS-3 C++ programming code of 5G RAN network simulation, traffic generation, performance metrics collection and analysis

```

1  #include "ns3/core-module.h"
2  #include "ns3/network-module.h"
3  #include "ns3/internet-module.h"
4  #include "ns3/mobility-module.h"
5  #include "ns3/point-to-point-module.h"
6  #include "ns3/ipv4-global-routing-helper.h"
7  #include "ns3/lte-module.h"
8  #include "ns3/applications-module.h"
9  #include "ns3/flow-monitor-helper.h"
10 #include "ns3/ipv4-flow-classifier.h"
11
12 using namespace ns3;
13
14
15
16 void CalculateBer2(Ptr<FlowMonitor> monitor) {
17     std::cout << "Enter_into_the_CalculateBer2_function" << std
18         ::endl;
19     // Retrieve the classifier from FlowMonitorHelper
20     FlowMonitorHelper flowmonHelper; // Ensure it matches your
21         flow monitor setup
22     Ptr<Ipv4FlowClassifier> classifier = DynamicCast<
23         Ipv4FlowClassifier>(flowmonHelper.GetClassifier());
24
25     // Ensure monitor and classifier are valid

```

```

23     if (!monitor || !classifier) {
24         std::cerr << "Error: Invalid FlowMonitor or Classifier
25         ." << std::endl;
26         return;
27     }
28
29     FlowMonitor::FlowStatsContainer stats = monitor->
30         GetFlowStats();
31     for (auto iter = stats.begin(); iter != stats.end(); ++
32         iter) {
33         Ipv4FlowClassifier::FiveTuple t = classifier->FindFlow
34             (iter->first);
35         std::cout << "Flow_ID:" << iter->first
36             << "(" << t.sourceAddress << "->" << t.
37             destinationAddress << ")" << std::endl;
38         // Add logic to calculate BER and display results
39     }
40 }
41
42 void CalculateBer4(Ptr<FlowMonitor> monitor, FlowMonitorHelper
43 & flowHelper) {
44     std::cout << "Entering CalculateBer3_function" << std::
45     endl;
46
47     // Retrieve the classifier from FlowMonitorHelper
48     Ptr<Ipv4FlowClassifier> classifier = DynamicCast<
49     Ipv4FlowClassifier>(flowHelper.GetClassifier());
50
51     if (!classifier) {
52         std::cerr << "Error: Classifier is null!" << std::endl
53         ;
54         return;
55     }
56
57     FlowMonitor::FlowStatsContainer stats = monitor->
58         GetFlowStats();
59
60     for (auto iter = stats.begin(); iter != stats.end(); ++
61         iter) {
62         Ipv4FlowClassifier::FiveTuple t = classifier->FindFlow
63             (iter->first);
64         std::cout << "Flow_ID:" << iter->first

```

```

53         << "_" << t.sourceAddress << "_->" << t.
            destinationAddress << ")" << std::endl;
54     std::cout << "Tx_Packets:" << iter->second.
        txPackets << std::endl;
55     std::cout << "Rx_Packets:" << iter->second.
        rxPackets << std::endl;
56     std::cout << "Throughput:"
57         << (iter->second.rxBytes * 8.0 / (iter->
            second.timeLastRxPacket.GetSeconds() -
            iter->second.timeFirstTxPacket.GetSeconds
                ())) / 1e6
58         << "Mbps" << std::endl;
59     double transmittedBits = iter->second.txBytes * 8.0;
        // Bytes to bits
60     double receivedBits = iter->second.rxBytes * 8.0;
61     double errorBits = transmittedBits - receivedBits;
62
63     double ber = (errorBits / transmittedBits) * 100.0; //
        BER in percentage
64     std::cout << "Flow_ID:" << iter->first << ", BER:"
        << ber << "%" << std::endl;
65 }
66 }
67
68 void CalculateLatency(Ptr<FlowMonitor> monitor,
    FlowMonitorHelper& flowHelper) {
69
70     // Retrieve the classifier from FlowMonitorHelper
71     Ptr<Ipv4FlowClassifier> classifier = DynamicCast<
        Ipv4FlowClassifier>(flowHelper.GetClassifier());
72
73     if (!classifier) {
74         std::cerr << "Error: Classifier is null!" << std::endl
            ;
75         return;
76     }
77
78     FlowMonitor::FlowStatsContainer stats = monitor->
        GetFlowStats();
79
80     for (auto iter = stats.begin(); iter != stats.end(); ++
        iter) {

```

```

81         double latency = iter->second.delaySum.GetSeconds() /
            iter->second.rxPackets; // Average delay
82         std::cout << "Flow_ID:_ " << iter->first << ",_Latency:
            _ " << latency * 1000 << "_ms" << std::endl;
83     }
84 }
85
86 void CalculatePer(Ptr<FlowMonitor> monitor, FlowMonitorHelper&
    flowHelper) {
87
88     // Retrieve the classifier from FlowMonitorHelper
89     Ptr<Ipv4FlowClassifier> classifier = DynamicCast<
        Ipv4FlowClassifier>(flowHelper.GetClassifier());
90
91     if (!classifier) {
92         std::cerr << "Error:_Classifier_is_null!" << std::endl
            ;
93         return;
94     }
95
96     FlowMonitor::FlowStatsContainer stats = monitor->
        GetFlowStats();
97
98     for (auto iter = stats.begin(); iter != stats.end(); ++
        iter) {
99
100         uint32_t lostPackets = iter->second.lostPackets;
101         uint32_t transmittedPackets = iter->second.txPackets;
102
103         double per = (double(lostPackets) / transmittedPackets
            ) * 100.0; // PER in percentage
104         std::cout << "Flow_ID:_ " << iter->first << ",_PER:_ "
            << per << "%" << std::endl;
105     }
106 }
107
108 void CalculateThroughput(Ptr<FlowMonitor> monitor,
    FlowMonitorHelper& flowHelper) {
109
110     // Retrieve the classifier from FlowMonitorHelper
111     Ptr<Ipv4FlowClassifier> classifier = DynamicCast<
        Ipv4FlowClassifier>(flowHelper.GetClassifier());

```

```

112
113     if (!classifier) {
114         std::cerr << "Error:_Classifier_is_null!" << std::endl
115             ;
116         return;
117     }
118
119     FlowMonitor::FlowStatsContainer stats = monitor->
120         GetFlowStats();
121
122     for (auto iter = stats.begin(); iter != stats.end(); ++
123         iter) {
124
125         double throughput = (iter->second.rxBytes * 8.0) /
126             iter->second.timeLastRxPacket.GetSeconds(); // bps
127         std::cout << "Flow_ID:_ " << iter->first << ",_
128             Throughput:_ " << throughput / 1e6 << "_Mbps" << std
129             ::endl;
130     }
131 }
132
133
134
135
136
137
138
139
140
141
142
143
144
145
146
147
148
149
150
151
152
153
154
155
156
157
158
159
160
161
162
163
164
165
166
167
168
169
170
171
172
173
174
175
176
177
178
179
180
181
182
183
184
185
186
187
188
189
190
191
192
193
194
195
196
197
198
199
200
201
202
203
204
205
206
207
208
209
210
211
212
213
214
215
216
217
218
219
220
221
222
223
224
225
226
227
228
229
230
231
232
233
234
235
236
237
238
239
240
241
242
243
244
245
246
247
248
249
250
251
252
253
254
255
256
257
258
259
260
261
262
263
264
265
266
267
268
269
270
271
272
273
274
275
276
277
278
279
280
281
282
283
284
285
286
287
288
289
290
291
292
293
294
295
296
297
298
299
300
301
302
303
304
305
306
307
308
309
310
311
312
313
314
315
316
317
318
319
320
321
322
323
324
325
326
327
328
329
330
331
332
333
334
335
336
337
338
339
340
341
342
343
344
345
346
347
348
349
350
351
352
353
354
355
356
357
358
359
360
361
362
363
364
365
366
367
368
369
370
371
372
373
374
375
376
377
378
379
380
381
382
383
384
385
386
387
388
389
390
391
392
393
394
395
396
397
398
399
400
401
402
403
404
405
406
407
408
409
410
411
412
413
414
415
416
417
418
419
420
421
422
423
424
425
426
427
428
429
430
431
432
433
434
435
436
437
438
439
440
441
442
443
444
445
446
447
448
449
450
451
452
453
454
455
456
457
458
459
460
461
462
463
464
465
466
467
468
469
470
471
472
473
474
475
476
477
478
479
480
481
482
483
484
485
486
487
488
489
490
491
492
493
494
495
496
497
498
499
500
501
502
503
504
505
506
507
508
509
510
511
512
513
514
515
516
517
518
519
520
521
522
523
524
525
526
527
528
529
530
531
532
533
534
535
536
537
538
539
540
541
542
543
544
545
546
547
548
549
550
551
552
553
554
555
556
557
558
559
560
561
562
563
564
565
566
567
568
569
570
571
572
573
574
575
576
577
578
579
580
581
582
583
584
585
586
587
588
589
590
591
592
593
594
595
596
597
598
599
600
601
602
603
604
605
606
607
608
609
610
611
612
613
614
615
616
617
618
619
620
621
622
623
624
625
626
627
628
629
630
631
632
633
634
635
636
637
638
639
640
641
642
643
644
645
646
647
648
649
650
651
652
653
654
655
656
657
658
659
660
661
662
663
664
665
666
667
668
669
670
671
672
673
674
675
676
677
678
679
680
681
682
683
684
685
686
687
688
689
690
691
692
693
694
695
696
697
698
699
700
701
702
703
704
705
706
707
708
709
710
711
712
713
714
715
716
717
718
719
720
721
722
723
724
725
726
727
728
729
730
731
732
733
734
735
736
737
738
739
740
741
742
743
744
745
746
747
748
749
750
751
752
753
754
755
756
757
758
759
760
761
762
763
764
765
766
767
768
769
770
771
772
773
774
775
776
777
778
779
780
781
782
783
784
785
786
787
788
789
790
791
792
793
794
795
796
797
798
799
800
801
802
803
804
805
806
807
808
809
810
811
812
813
814
815
816
817
818
819
820
821
822
823
824
825
826
827
828
829
830
831
832
833
834
835
836
837
838
839
840
841
842
843
844
845
846
847
848
849
850
851
852
853
854
855
856
857
858
859
860
861
862
863
864
865
866
867
868
869
870
871
872
873
874
875
876
877
878
879
880
881
882
883
884
885
886
887
888
889
890
891
892
893
894
895
896
897
898
899
900
901
902
903
904
905
906
907
908
909
910
911
912
913
914
915
916
917
918
919
920
921
922
923
924
925
926
927
928
929
930
931
932
933
934
935
936
937
938
939
940
941
942
943
944
945
946
947
948
949
950
951
952
953
954
955
956
957
958
959
960
961
962
963
964
965
966
967
968
969
970
971
972
973
974
975
976
977
978
979
980
981
982
983
984
985
986
987
988
989
990
991
992
993
994
995
996
997
998
999

```

```

148     Ptr<Node> cloudNode = CreateObject<Node>();
149
150     // Mobility for IoT users
151     MobilityHelper mobility;
152     mobility.SetPositionAllocator("ns3::
        RandomRectanglePositionAllocator");
153     mobility.SetMobilityModel("ns3::
        ConstantPositionMobilityModel");
154     mobility.Install(iotUsers);
155
156     // Mobility for Base Stations
157     mobility.SetPositionAllocator("ns3::GridPositionAllocator"
        ,
158                                     "MinX", DoubleValue(0.0),
159                                     "MinY", DoubleValue(0.0),
160                                     "DeltaX", DoubleValue(50.0),
161                                     "DeltaY", DoubleValue(50.0),
162                                     "GridWidth", UIntegerValue
                                                (3),
163                                     "LayoutType", StringValue("
                                                RowFirst"));
164     mobility.SetMobilityModel("ns3::
        ConstantPositionMobilityModel");
165     mobility.Install(baseStations);
166
167     // Configure LTE Network
168     Ptr<LteHelper> lteHelper = CreateObject<LteHelper>();
169     Ptr<PointToPointEpcHelper> epcHelper = CreateObject<
        PointToPointEpcHelper>();
170     lteHelper->SetEpcHelper(epcHelper);
171
172     NodeContainer enbNodes = baseStations;
173     NodeContainer ueNodes = iotUsers;
174
175     NetDeviceContainer enbDevices = lteHelper->
        InstallEnbDevice(enbNodes);
176     NetDeviceContainer ueDevices = lteHelper->InstallUeDevice(
        ueNodes);
177
178     // Install the IP stack on UEs
179     InternetStackHelper internet;
180     internet.Install(iotUsers);

```

```

181 Ipv4InterfaceContainer ueIpIfaces = epcHelper->
    AssignUeIpv4Address(NetDeviceContainer(ueDevices));
182
183 // Connect UEs to the Internet
184 for (uint32_t i = 0; i < ueDevices.GetN(); ++i) {
185     Ptr<NetDevice> ueDevice = ueDevices.Get(i);
186     Ptr<Ipv4> ipv4 = ueDevice->GetNode()->GetObject<Ipv4
        >();
187     lteHelper->Attach(ueDevice, enbDevices.Get(0));
188 }
189
190
191 // Install a UDP application to generate traffic
192 uint16_t dlPort = 1234;
193 UdpServerHelper server(dlPort);
194 ApplicationContainer serverApps = server.Install(iotUsers.
    Get(0));
195 serverApps.Start(Seconds(1.0));
196 serverApps.Stop(simTime);
197
198 UdpClientHelper client(ueIpIfaces.GetAddress(0), dlPort);
199 client.SetAttribute("MaxPackets", UintegerValue(100000));
200 client.SetAttribute("Interval", TimeValue(Milliseconds(1))
    );
201 client.SetAttribute("PacketSize", UintegerValue(10240));
202 ApplicationContainer clientApps = client.Install(
    baseStations.Get(0));
203 clientApps.Start(Seconds(1.0));
204 clientApps.Stop(simTime);
205
206
207 // Install FlowMonitor - 1
208 FlowMonitorHelper flowHelper;
209 Ptr<FlowMonitor> flowMonitor = flowHelper.InstallAll();
210
211
212
213 // Schedule the BER calculation
214 // Simulator::Schedule(Seconds(11.0), &CalculateBer2,
    flowMonitor);
215 Simulator::Schedule(Seconds(10.0), &CalculateBer4,
    flowMonitor, std::ref(flowHelper));

```

```
216 Simulator::Schedule(Seconds(11.0), &CalculateLatency,  
    flowMonitor, std::ref(flowHelper));  
217 Simulator::Schedule(Seconds(12.0), &CalculatePer,  
    flowMonitor, std::ref(flowHelper));  
218 Simulator::Schedule(Seconds(13.0), &CalculateThroughput,  
    flowMonitor, std::ref(flowHelper));  
219 // Start Simulation  
220 Simulator::Stop(simTime);  
221 Simulator::Run();  
222  
223  
224 Simulator::Destroy();  
225  
226 return 0;  
227 }
```