



ADDIS ABABA UNIVERSITY
SCHOOL OF ELECTRICAL AND COMPUTER
ENGINEERING
TELECOMMUNICATION ENGINEERING GRADUATE PROGRAM

Thesis on:

**Performance Comparison of Machine Learning-Based Dynamic
Resource Allocation Methods for LTE-A Using Realistic Data**

By
Tiblets Kinfu

Advisor
Dr. Eng Yihenew Wondie

September 30, 2023
Addis Ababa, Ethiopia

ADDIS ABABA UNIVERSITY
SCHOOL OF ELECTRICAL AND COMPUTER
ENGINEERING
TELECOMMUNICATION ENGINEERING GRADUATE PROGRAM

Thesis on

Performance Comparison of Machine Learning-Based Dynamic
Resource Allocation Methods for LTE-A Using Realistic Data

By:

Tiblets Kinf

Dr. Eng Yihenew Wondie

Advisor

Signature

Date

Evaluator

Signature

Date

Evaluator

Signature

Date

Declaration

I, the undersigned, declare that this thesis is my original work and does not incorporate without recognition of any material previously submitted for a degree or diploma in any other university and all sources of materials used for the thesis have been fully acknowledged. To the best of my knowledge, it does not contain any material previously published or written by another person unless mentioned in the text.

Tiblets Kinf

Name

Signature

Addis Ababa

Place

_____/_____/_____
Date of Submission

This thesis has been submitted for examination with my approval as a university advisor.

Abstract

In the last few years, fourth generation (4G) mobile networks have become much more complex due to deployments of macro, micro, pico, and femto-cells to boost network capacity and quality of Services(QoS). Long-Term Evolution Advanced (LTE-A) system is a 4G wireless communication technology that offers mobile devices high-speed data transmission, increased capacity, and improved network performance. While LTE-A offers improved capabilities, it also poses unique resource allocation problems. The following are some difficulties in allocating LTE-A resources: heterogeneous network deployment, diverse QoS requirements and dynamic traffic patterns. Traditional resource management approaches usually use static strategies or predefined rules and are insensitive to the current demand on the network. To allocate resources effectively based on the current network conditions, resource allocation must be adaptable and dynamic. Adaptive resource allocation methods that can handle the dynamic nature of LTE-A networks can be developed with the help of machine learning and artificial intelligence. This thesis evaluates the performance of machine learning-based resource allocation strategies in LTE-A networks using realistic data. Four different machine learning-based resource allocation strategies are compared: Long Short Term Memory (LSTM), Conventional Neural Network and Long Short Term Memory (CNN-LSTM), Random Forest (RF) and K Nearest Neighbor(KNN) algorithm. Performance measures include the accuracy of resource allocation, Root Mean Square Error (RMSE), Mean Absolute Error(MAE) and speed of convergence(running time). The results show that LSTM is the model with higher accuracy score of resource allocation that is 98.56% and in terms RMSE and MAE random forest has lowest value of 0.294 and 0.08. In the case of running time KNN takes shortest time of 1.58 second.

Keywords: *LTE, LTE-Advanced, Machine Learning, Radio Resource Allocation, Channel quality indicator*

Acknowledgment

First and foremost, I thank the almighty God for giving me all the help, courage, and strength to accomplish the thesis.

Second, I would like to express my truly grateful to my advisor Dr. Eng Yihenew Wondie for his continuous support of my thesis, for his patience, motivation, interest to guide me.

Likewise, I would like to express my sincere gratitude to Dr. Beneyam Berehanu and Dr. Yalemzewd Negash, for his constructive comments and helpful suggestions during thesis progress report presentations.

Finally, I'd like to express my heartfelt gratitude to my beloved husband and my children's. I am able to complete this research because of your unwavering support, unconditional love, and your precious time.

Contents

Abstract	ii
Acknowledgment	iii
List of Figures	vii
List of Tables	ix
List of Abbreviations	x
1 Introduction	1
1.1 Background	1
1.2 Statement of the Problem	5
1.3 Objectives of the Thesis	6
1.3.1 General Objective	7
1.3.2 Specific Objectives	7
1.4 Literature Review	7
1.4.1 Dynamic Resource Allocation Strategies	8
1.4.2 Traffic Prediction Models	9
1.4.3 Machine Learning Based Dynamic Resource Allocation	12
1.5 Methodology	14
1.6 Scope and Limitations	15
1.7 Contribution of the Thesis	15
1.8 Thesis Organization	16
2 Machine Learning-Based Resource Allocation in LTE-A Network	17
2.1 Overview of Rradio Resource Allocation in LTE/LTE-A . . .	17
2.1.1 Overview of LTE/LTE-Advanced	17
2.1.2 Conventional Resource Allocation Strategies	20
2.2 Machine Learning Basics	23

2.3	Deep Learning (DL)	24
2.4	Deep Learning-Based Resource Allocation	25
2.4.1	Supervised Deep Learning	25
2.4.2	Artificial Neural Network(ANN)	26
2.4.3	Conventional Neural Network(CNN)	27
2.4.4	Recurrent Neural Network(RNN)	28
2.4.5	Deep Reinforcement Learning	28
3	System Model and Data Preparation	31
3.1	System Model	31
3.1.1	Resource Allocation	31
3.2	Data Preparation and Feature Extraction	33
3.2.1	Data Collection	33
3.2.2	Data Preprocessing and Feature Extraction	34
3.3	Machine Learning Algorithms	34
3.3.1	Long Short Term Memory(LSTM)	35
3.3.2	Conventional Neural Network and Long Short Term Memory(CNN-LSTM)	38
3.3.3	Random Forest(RF)	39
3.3.4	K-Nearest Neighbor(KNN)	41
3.4	Performance Metrics	42
3.4.1	Accuracy	42
3.4.2	Root Mean Square Error(RMSE)	42
3.4.3	Mean Absolute Error(MAE)	42
3.4.4	Mean Absolute Percentage Error(MAPE)	42
4	Results and Discussions	44
4.1	Results	44
4.1.1	Long Short Term Memory	44
4.1.2	Conventional Neural Network-Long Short Term Mem- ory (CNN-LSTM)	49

4.1.3 Random Forest	52
4.1.4 K-Nearest Neighbor(KNN)	54
4.2 Performance Comparison	55
5 Conclusions and Recommendations	58
5.1 Conclusions	58
5.2 Recommendations	59
Reference	60
Appendix: Manuscript	67

List of Figures

Figure 1: Frame of long term evolution(LTE)	3
Figure 2: Physical model of a downlink resource Allocation . .	3
Figure 3: Intelligent module for better decision-making in managing network resources to serve users	5
Figure 4: The overall methodology of the research	15
Figure 5: LTE architecture	18
Figure 6: Scheduling block in an LTE downlink	19
Figure 7: A deep neural network with three hidden layers . . .	25
Figure 8: Train stage	26
Figure 9: Test trained model	26
Figure 10: Artificial neural network	27
Figure 11: Structure of LSTM network	28
Figure 12: Deep Reinforcement Learning	29
Figure 13: Enhanced proportional fair scheduling algorithm . . .	32
Figure 14: Overall system model	32
Figure 15: Traffic Distribution	33
Figure 16: (a) Average active user (b) Average power used in dBm (c) DL PRB Utilization (d) The pattern of the combination traffic, power and PRB utilized	34
Figure 17: Types of machine learning algorithm	35
Figure 18: LSTM cell structure	37
Figure 19: The structure of CNN-LSTM model	39
Figure 20: The structure of random forest algorithm	41
Figure 21: LSTM model output of each epoch train accuracy, test accuracy, train loss and test loss	47
Figure 22: (a) LSTM accuracy at 10 epochs (b) LSTM accuracy at 20 epochs (c) LSTM train test loss at 10 epochs (d) LSTM train test loss at 20 epochs	48

Figure 23: (a)Predicted using LSTM model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized	49
Figure 24: (a) CNN-LSTM accuracy at 10 epochs (b) CNN-LSTM accuracy at 20 epochs (c) CNN-LSTM train test loss at 10 epochs (d) CNN-LSTM train test loss at 20 epochs	51
Figure 25: (a)Predicted using CNN-LSTM model Vs observed value of power and PRB utilized (b)Average power used (c) PRB utilized	52
Figure 26: (a)Predicted using RF model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized	53
Figure 27: (a)Predicted using KNN model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized	55
Figure 28: Predicted Vs observed value of average power used for each model	57
Figure 29: Predicted Vs observed value of PRB utilized for each model	57

List of Tables

Table 1: Summary of literature review	13
Table 2: Summary of difference between LTE and LTE-A . . .	17
Table 3: Summary of conventional resource allocation algorithms	23
Table 4: Hyperparameter and its value	46
Table 5: LSTM model performance metrics value	46
Table 6: CNN-LSTM model hyperparameter possible value . . .	50
Table 7: CNN-LSTM model performance metrics value	50
Table 8: Random forest model hyperparameter possible value .	52
Table 9: RF model evaluation result of possible value	53
Table 10:KNN model performance metrics value	54
Table 11:Summary of performance evaluation of each model . .	55

List of Abbreviations

3GPP	3rd Generation Partnership Project
4G	Fourth Generation
AI	Artificial Intelligence
ANN	Artificial Neural Network
AMC	Adaptive Modulation and Coding
AR	Auto-Regressive
ARIMA	Auto-Regressive Integrated Moving Average
BCQI	Best Channel Quality Indicator
BS	Base Stations
BSR	Buffer Status Report
CAPEX	Capital Expenditures
CNN	Conventional Neural Network
CoMP	Cooperative Multi-Point
CQI	Channel Quality Information
DCI	Downlink Control Information
DDPG	Deep Deterministic Policy Gradient
DNN	Deep Neural Network
DQN	Deep Q Network
DQL	Deep Q Learning
DRX	Discontinuous Reception
eNB	Evolved Node B
EPC	Evolved Packet Core
EPF	Enhanced Proportional Fair
EPS	Evolved Packet System
HSS	Home Subscriber Server
IMAC	Interfering Multiple-Access Channel
ITU	International Telecommunication Union
KNN	K Nearest Neighbor

LSTM	Long Short-Term Memory
LTE-A	Long-Term Evolution Advanced
MAE	Mean Absolute Error
MAC	Medium Access Control
MAPE	Mean Absolute Percentage Error
MCS	Modulation and Coding Scheme
ML	Machine Learning
MME	Mobility Management Entity
MS	Mobile Station
NRT	Non-Real-Time
OFDM	Orthogonal Frequency Division Multiplexing
OPEX	Operational Expenditures
PBCH	Physical Broadcast Channel
PCFICH	Physical Control Format Indicator Channel
PDCCH	Physical Downlink Control Channel
PDSCH	Physical Downlink Shared Channel
PF	Proportional Fair
P-GW	Packet Gateway
PUCCH	Physical Uplink Control Channel
QoE	Quality of Experience
QoS	Quality of Service
RAB	Radio Access Bearer
RAN	Radio Access Network
RB	Resource Block
RF	Random Forest
RMSE	Root Mean Square Error
RNC	Radio Network Controllers
RNN	Recurrent Neural Network
RR	Round-Robin
RRM	Radio Resource Management

RT	Real-Time
SARIMA	Seasonal Auto Regressive Integrated Moving Average
SB	Scheduling Block
SC-FDMA	Single Carrier Frequency Division Multiple Access
S-GW	Serving Gateway
SVR	Support Vector Regression
TP	Transmission Point
TTI	Transmission Time Interval
UE	User Equipment
UMTS	Universal Mobile Telecommunication System
UPN	Unsupervised Pre-trained Networks
UTRAN	Universal Terrestrial Radio Access Network
WLAN	Wide Area Network

1 Introduction

1.1 Background

The scale and complexity of fourth generation (4G) mobile networks have grown significantly during the past few years. By adding more cells of various types to the present deployments, which include macro-, micro-, pico-, and femtocells, operators are constantly looking to increase network capacity and Quality of Service (QoS). The Operating and Capital Expenditures (OPEX/CAPEX) of the mobile network providers significantly increase as a result of this increase in complexity. To permanently decrease these costs, operators are looking for network solutions that would automatically configure, manage, and optimize networks while reducing the need for human interaction based on the traffic demand. Radio Resource Management (RRM) making use of the available adaptation techniques to serve users depending on their QoS parameters to ensure the efficient use of available radio resources [1, 2, 4]. To enhance the capacity of the network and absorb the mobile data traffic increment, ethiotelecom has been conducting a lot of expansion projects, base station densifications on a greenfield or rooftop, load balancing and optimizations.

In order to inform the User Equipment (UE) about its allotted time/frequency resources and transmission formats to be employed by UE, eNodeB performs scheduling on the available radio resources. UE capability, QoS, and measurement reports from the UE are used to inform Radio Resource (RR) scheduling [3]. Each eNodeB must distribute transmission powers throughout each resource block while taking into account the impact on UE throughput and interference on other UEs connected to other base stations. A resource allocation framework in the time and frequency domains is defined by the Long-Term Evolution Advanced (LTE-A) standard. The frequency dimension is broken into sub-carriers that are separated by 15 KHz. The time dimension is split into ten 1 ms sub-frames and then di-

vided into 10 radio frames of 1ms each. It divides each sub-frame into two 0.5 ms slots. The lowest unit of resource, known as a resource element, is one sub carrier lasting one Orthogonal Frequency Division Multiplexing (OFDM) signal. The length of a resource block is one millisecond, with 12 constant sub-carriers. A centralized resource block scheduling method chooses which UE will be assigned to each resource block. Only one user may be given access to a resource block [4, 37].

A resource block (RB) is the smallest unit of resource allocation for any one user. Twelve sub-carriers and six or seven OFDM symbols are used by one resource block. Consequently, a resource block may contain a total of 72 or 84 resource elements. Radio resource elements are separated into time and frequency resource elements in LTE. As in Figure 1 resource elements are separated into a certain number of sub-carriers on the frequency axis and a certain number of OFDM symbols on the time axis. The smallest resource unit, known as a resource element (RE), uses both an OFDM symbol and a single sub-carrier. Flexible bandwidths supported by the technology include 1.4MHz, 3MHz, 5MHz, 10MHz, 15MHz, and 20MHz. Each bandwidth can have a variable number of resource blocks, namely 6, 15, 25, 50, 75, and 100 RBs. The five physical channels Physical Downlink Control Channel (PDCCH), Physical Downlink Shared Channel(PDSCH), Physical Broadcast Channel(PBCH), Physical Control Format Indicator Channel(PCFICH), and Physical Hybrid ARQ Indicator Channel(PHICH) suggested by the LTE standard are for use in the downlink. The shared data traffic channel is called PDSCH, and the control channel is called PDCCH. In a 1ms transmit time interval (TTI) sub-frame, a number of resource blocks are allotted to these channels. 10 sub-frames, each ranging 1ms, make up a single 10ms LTE radio frame as in Figure 1. Two resource blocks (RBs) of 0.5ms each are included in each sub-frame [5, 17, 35].

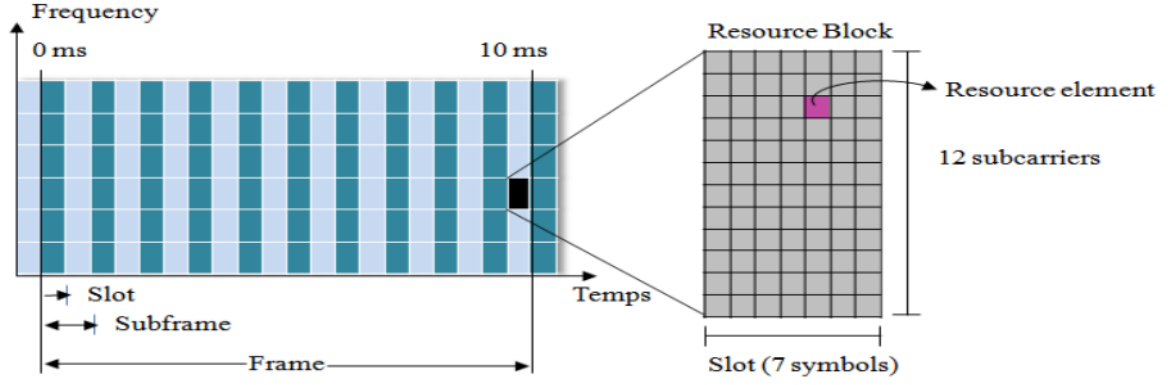


Figure 1: Frame of long term evolution(LTE) [7]

The user measurement entity reports the Channel Quality Indicator (CQI) to the base station (BS) for every TTI to give time and frequency channel quality information for improved spectral efficiency and resource allocation. Users can alert the BS about the channel quality for downlink RBs by using the Physical Uplink Control Channel (PUCCH). By employing the PDCCH as shown in Figure 2, BS distributes downlink RB allocations and Modulation and Coding scheme(MCS) assignments to all users [6].

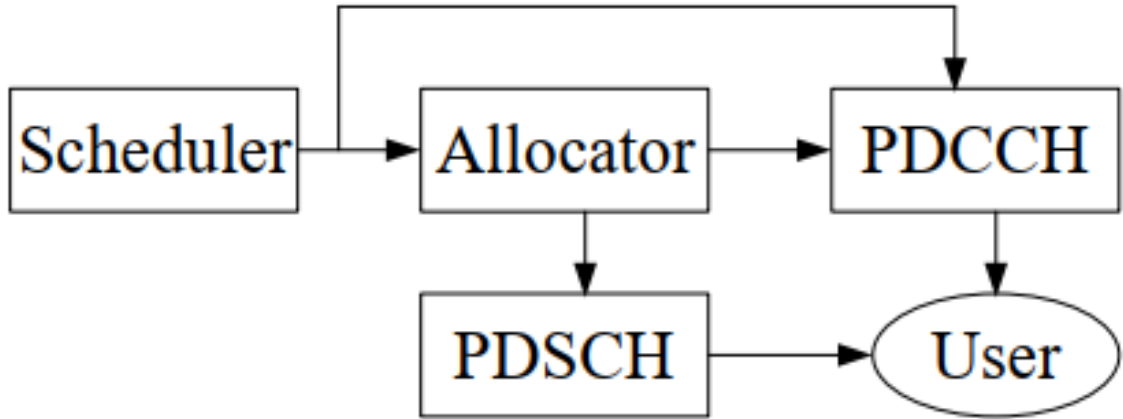


Figure 2: Physical model of a downlink resource Allocation [8]

The International Telecommunication Union's (ITU) guideline is not met by the LTE standard. Since, it is a 3.9G rather than a full 4G technology. As a result, LTE was further developed in accordance with ITU

rules to become LTE-A, which is a true 4G technology that offers up to 500 MBPS in the uplink and up to 1 GBPS in the downlink. It enables to have such increased data rates and low latency using higher order Multiple Input Multiple Output(MIMO), relay technology, Cooperative Multipoint(CoMP), carrier aggregation, and other technologies [5, 1].

To obtain high performance, it is vital to utilize limited network resources as effectively as possible in the most advantageous way possible while also maintaining the best network utilization, to model and predict exactly cellular networks will load [9]. Besides, to assign the available network resources to the devices requesting service while controlling the level of network resources, such as the quantity of radio resources, the Buffer Status Report (BSR) of users is typically used in traditional approaches to network resource management to provision and allocate resources as shown in Figure 3. However, in this case, take a proactive approach and look for chances to implement resource provisioning and allocation. Then, based on the present and past states of user traffic behavior, channel condition, required quality of service, tries to decide how to serve a group of users who coexist. It must do a through analysis of each user's traffic behavior [10]. This thesis compares the performance of different distinct machine learning based dynamic resource allocating strategies, to jointly allocate downlink Physical Resource block(PRB) utilization and average power used by each user using realistic data collected from Addis Ababa LTE-A network.

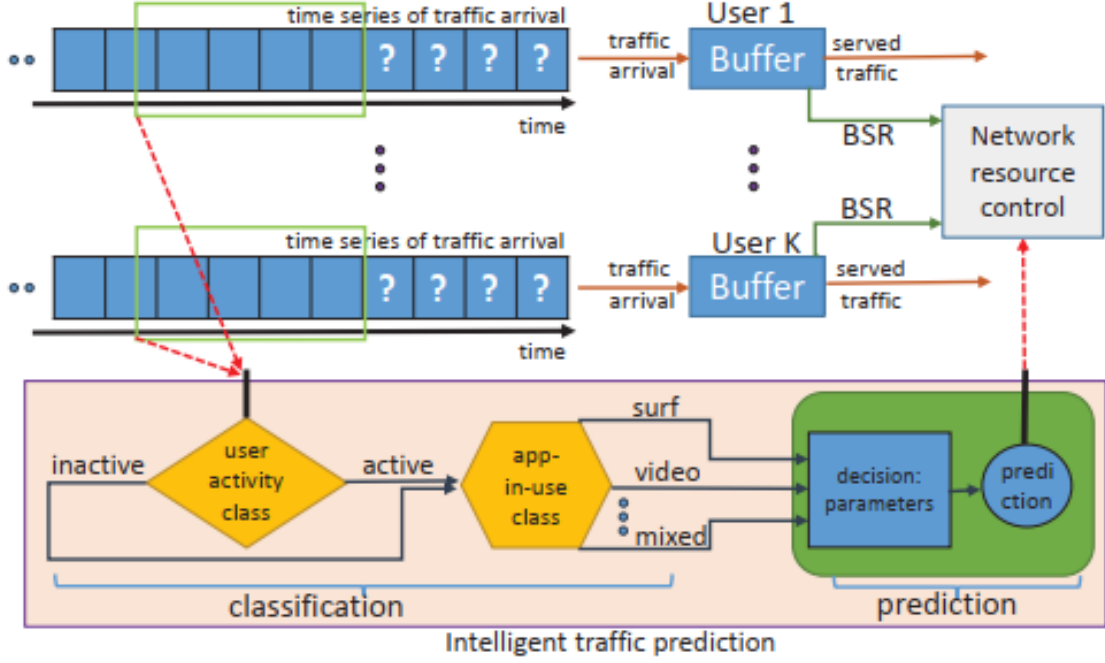


Figure 3: Intelligent module for better decision-making in managing network resources to serve users [10]

To utilize network resources effectively, considering traffic volume, active user traffic, available resources is necessary. Accurate network traffic prediction at access points provides effective resource allocation to guarantee acceptable service quality [11].

1.2 Statement of the Problem

In the last few years, 4G mobile networks have become much more complex as operators are always looking to expand their deployments of macro, micro, pico, and femtocells as well as other cell types to boost network capacity and QoS. Traditional resource management approaches usually use static strategies and are insensitive to the current demand on the networks. These static methods can cause certain cells to waste those few resources while also adding to congestion in a select few cells. As a result, such static solutions are ineffective at resolving the capacity issue facing the network of the future. Planning and management of the network

are deployed to address this increasing demand for data traffic, and they depend on how well the network resource is used and the volume of traffic at the time.

In LTE-A, resource allocation refers to the distribution of various network resources, such as frequency, time, and power, among users and various services. While LTE-A offers improved capabilities, it also poses unique resource allocation problems. The following are some difficulties in allocating LTE-A resources: Spectrum scarcity, interference management, Heterogeneous network deployment, diverse QoS requirements and Dynamic traffic patterns. To allocate resources effectively based on the current network conditions, resource allocation must be adaptable and dynamic. Advanced resource allocation algorithms, clever network design, and optimization strategies are needed to overcome these obstacles. Adaptive resource allocation methods that can handle the dynamic nature of LTE-A networks can be developed with the help of machine learning and artificial intelligence. The goal of this thesis is to provide a machine learning-based resource allocation framework that enhances network performance across all bands in LTE-A systems while overcoming the drawbacks of conventional approaches and comparing the performance of machine learning based resource allocation strategies. The aim of this work intends to develop a more intelligent, flexible, and efficient resource allocation solution for LTE-A networks by utilizing actual data and sophisticated machine learning algorithms.

1.3 Objectives of the Thesis

The general and specific objectives of the research are as follows.

1.3.1 General Objective

The general objective of this thesis is to compare the performance of machine learning-based dynamic resource allocation methods by utilizing actual data of ethiotelecom LTE-A network in Addis Ababa.

1.3.2 Specific Objectives

The specific objectives of this research are as follows

- To review machine learning-based resource allocation-related literature for LTE-A
- To Collect and preprocess the collected data from ethiotelecom LTE-A network
- To make dynamic resource allocation decisions in real-time.
- To develop machine learning-based resource allocation model using historical data of ethiotelecom LTE-A network for performing joint resource allocation (i.e., both power and physical resource block) to maximize the overall network throughput..
- To evaluate the performance of each algorithm in terms of accuracy of resource allocation and speed of convergence
- To compare the performance of different models and identify the best-performing model.

1.4 Literature Review

Researchers are addressing the challenge of resource allocation by utilizing innovative traffic forecasting techniques, by developing intelligent systems that considers channel and user conditions to optimize resource allocation following the traffic data flow during network operation. Some of the related literature's are presented as follows.

1.4.1 Dynamic Resource Allocation Strategies

A fair and effective resource allocation system is crucial for enhancing network performance [12]. To enhance network capacity, QoS, and bandwidth usage, it introduces a dynamic resource allocation technique. To maximize the total network throughput, the resource controller distributes transmission bandwidth to subnetworks in a heterogeneous network cooperatively. We examine a heterogeneous LTE/WLAN network with a spectrally efficient distribution of transmission resources across its member networks [13].

The authors in [14] present that using RRM, the LTE Advanced system distributes radio resources for various services. Different service kinds, including real-time (RT), non-real-time (NRT), and control signaling, are dynamically active in networks, necessitating dynamic resource allocation for the transmission of pertinent data packets. When findings from the suggested technique are compared to those from existing methods, it becomes clear that the current solution yields superior outcomes. There are descriptions of several service types and their schedule. It demonstrates service prioritizing optimization strategies in which the weights of service types are optimized. For resource allocation, it does not, however, take into account the volume of traffic, user count, channel status, etc.

Since it is responsible for ensuring user QoS in LTE-A networks, multi-user scheduling is one of the key components. Typically, scheduling algorithms strive for maximum throughput while retaining a certain level of fairness. The smallest resource unit in LTE-A is the RB, which has a frequency domain size of 180 KHz and a time domain division into two slots with a duration of 0.5 ms each. The eNodeB makes a scheduling decision every 1 ms, which is the size of the Transmission Time Interval (TTI). Physical resources are distributed to users in the frequency or time grid throughout time, and multi-user diversity is controlled in both of these time and

frequency domains. Signaling limitations prevent subcarriers from being allocated separately, so they should be combined on an RB basis [15, 16]. The efficiency of Universal Terrestrial Radio Access Network (EUTRAN) when using various resource allocation algorithms, including round robin, proportional fair, best channel quality indicator (CQI), maximum transmission point (TP), and resource fair, was evaluated using MATLAB. Fair resource distribution among UE is provided via round robin, proportional fair, resource fair, and best CQI. High cell throughput with selective resource allocation is provided by MAX TP. The best CQI and resource fair strategy offers a network with degraded resource occupancy high throughput and fair resource distribution. Although the UE throughput is relatively smaller with the round robin technique, resource allocation and occupancy are quite fair. The authors recommend adapting mixed strategies in addition to channel imperfection in which the performance of resource allocation depends [17, 18]. The paper [19] compares how well two fundamental techniques proportional fairness (PF) and round robin perform. According to this, the PF approach offers great fairness but just average throughput. However, from a scheduling standpoint, it does not distinguish between different service types; instead, a technique for raising QoS called proportional fairness is offered. Similar to this, [20] also describes adaptive QoS for PF methods. These methods, however, cannot have the capability to determine the current service type prior to resource scheduling. Dynamic resource allocation is necessary for dynamic service behavior. This research focuses on performance comparison of machine learning-based dynamic resource allocation methods for LTE-A using realistic data.

1.4.2 Traffic Prediction Models

Time-series analysis techniques have typically been used to study the prediction of mobile traffic patterns. Prediction issues are divided into tem-

poral and spatiotemporal prediction issues based on how cellular traffic is predicted. The statistical, machine learning, and deep learning models of artificial intelligence are subcategories of prediction models [21]. To capture the trends of the temporal evolution of the mobile traffic, the majority of previous related works use statistical approaches like Auto-Regressive (AR), Auto Regressive Integrated Moving Average (ARIMA) and Seasonal Auto-Regressive Integrated Moving Average (SARIMA). Since the prediction tends to over-reproduce the average of the previously observed values, one of the known disadvantages of such algorithms is their low resistance to the time-series' rapid fluctuations. These techniques also function with homogeneous time-series, in which the input and the forecast both fall into the same range of values [22].

It is necessary to replace the legacy design approach with user-centric predictive analysis of mobile network traffic and proactive network resource management since resource provisioning and operation control based on cell-aggregated traffic scenarios are not as energy and cost-efficient. By creating traffic prediction tools that combine per-user data with common machine learning (ML) techniques like LSTM and ARIMA. In this investigation, the effects of several parameters, including the temporal granularity, the length of future forecasts, and feature selection are examined. It gives a thorough empirical evaluation of the provided solutions over a real network traffic dataset. The presentation of ML models for dynamic Discontinuous reception (DRX) parameter adaptation to user traffic. The findings demonstrate that DRX parameter adaptation using online traffic prediction offers significantly more energy savings at low latency cost than the more traditional cell-wide DRX parameter adaptation [23]. The author in [24] Base Stations (BSs) can conserve energy by applying deep learning to predict the light-traffic data periods for them and sleeping during those times. To control the on-demand allocation of network resources and to lower network operation costs, traffic-exchange prediction

is typically based on an hourly granularity [9].

A hybrid deep learning model for spatiotemporal prediction that uses LSTM units for temporal modeling and an unique autoencoder-based deep model for spatial modeling. The findings demonstrate that, when compared to two widely used baseline approaches, ARIMA and Support Vector Regression (SVR), the proposed deep learning model greatly increases prediction accuracy [25].

Recently, a few research has been undertaken to predict the Universal Mobile Telecommunication System (UMTS) and LTE network traffic in Addis Ababa city using neural network models. In reference [26] use ConvLSTM to predict spatiotemporal UMTS mobile traffic data considering space and time variation of the UMTS traffic data by collecting real time dataset from 739 base stations for three months. It recommends to examining user-level data traffic use patterns to provide better information for spatiotemporal traffic data prediction and improve resource management. Besides in [27] to predict time series traffic data for UMTS network using LSTM by collecting from five Radio Network Controllers (RNCs) of the UMTS for ten months and it recommends to improving forecast accuracy use a combining approach, different dissimilar models. The authors in paper [28] use hybrid Convolutional Neural Network and Long Short Term Memory (CNN-LSTM) to predict temporal LTE traffic data by collecting data from 690 eNodeBs for four months on an hourly granularity and by clustering the base stations using k-means then predicted on cluster basis. It recommends to consider more specific multivariate features such as amount of spectrum utilized and Radio Access Bearer(RAB) attributes like maximum source data, traffic type, and maximum bit rate to improve the model performance.

1.4.3 Machine Learning Based Dynamic Resource Allocation

In [29], the author presented supervised deep learning model for allocation of sub-band and power for multi-cell wireless network to maximize the total network throughput. The proposed model uses CQI values of each user of the multicell network as input and allocates the required radio resources. The proposed model needs only one-time offline training and after training, it can allocate radio resources in an online fashion. The optimality of this work is no guarantee because it uses Genetic algorithm to generate training and testing data. In [30], the authors propose a method for distributed power regulation based on a data-driven Convolutional Neural Network (CNN) that uses spatial data in channel gain to estimate the transmit power. Maximizing throughput through power allocation is the major goal. There is no assurance that the suggested approach would produce the centrally optimum result. The authors in [31] propose two-step training framework that is a Deep Q Network (DQN) is trained with Deep Q Learning (DQL) algorithm with the off-line learning environment and DQN is fine-tuned in on-line training with real time data. They considered Interfering Multiple-Access Channel (IMAC) cellular network with the problem of power allocation.

Table 1: Summary of literature review

Ref	Objective	Method	Contribution	Limitations
[32]	Energy efficiency	A two-step neural network (NN)	By allocating RB and Power using unsupervised and supervised neural networks to reduce the computational cost achieve, better energy efficiency ,throughput and fairness	The dataset is generated by a genetic algorithm not from real deployment.
[33]	Power efficiency	Deep Reinforcement Learning	DRL-based framework for dynamic resource allocation in a cloud RAN for power efficiency, user demand satisfaction and the capability of handling highly dynamic cases.	Historical data used to train the DNN offline is not known
[30]	Power efficiency	CNN	It propose a Deep power control strategy using CNN to allocate transmit power based on the channel condition to maximize spectral efficiency or power efficiency	It considers only CSI
[29]	Power and throughput maximizing	Deep neural-network	To solve the sub band and power allocation problem in multi cell networks using supervised Deep neural network for maximizing through put	The dataset is generated by a genetic algorithm not represent real deployment
[43]	Maximized throughput	Dynamic neural Qlearning	How to intelligently select the best scheduling rule based on the estimated values of throughput and cell fairness index	Uplink and downlink in the network are considered to be one
[34]	To increase data rate	LSTM	Allocate bandwidth and power to each UE based on data rate	The only data rate is considered

This thesis focuses on performance comparison of machine learning-based dynamic resource allocation using the realistic data of Addis Ababa LTE-A network. It considers the actual active user traffic, traffic volume, channel quality indicator and service types rather than predicting traffic first. Besides, traditional resource allocation strategies follow predefined rules so that cannot adapt the dynamic nature of the network.

1.5 Methodology

In this study, a multivariate time series data approach are employed to construct accurate machine learning based radio resource allocation model to allocate average power used and physical resource block based on channel state information and quality of service requirement of each user. Related literature's are reviewed to select most fit machine learning based resource allocation algorithms for regression problem . Historical data is collected from the real deployment of ethiotelecom LTE-A network in Addis Ababa to build the machine learning based resource allocation model for future online allocation of radio resources. It is collected two months hourly basis historical data of 450 cells and finally one month 213 cells data is used to build the resource allocation models by removing the cells that have down hours. The collected data is preprocessed for the way that is suitable for the machine learning models. Data preparation, which involves locating and removing outliers and noisy values as well as missing information, is carried out. Correlation analysis is used to choose the pertinent features from multivariate time series data. The best machine learning based resource allocation model that enhance the network performance in terms of throughput and QoS is selected by evaluating and comparing the performance of each model. The performance of each model is evaluated in terms of accuracy of resource allocation and speed of convergence.

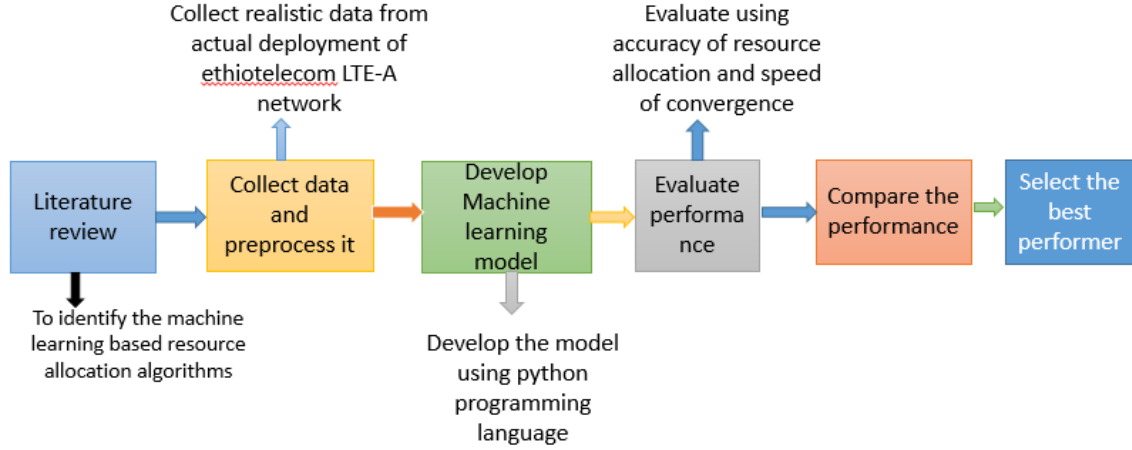


Figure 4: The overall methodology of the research

1.6 Scope and Limitations

This thesis mainly focused on performance comparison of ML-based resource allocation for LTE-A systems using realistic data. In this thesis the evaluation of different machine learning (ML) methods for resource allocation like LSTM, CNN-LSTM, Random forest and K nearest neighbor is investigated using the real data collected form ethiotelecom. However, the implementation of the models for examining the realistic systems is not considered.

1.7 Contribution of the Thesis

This thesis focuses on performance comparison of machine learning based dynamic resource allocation methods for LTE-A networks using realistic data for jointly allocating radio resources such as, average power used and physical resource block utilized. This allocates based on channel conditions, quality of service requirements, service type, average traffic volume and active user traffic to adapt the dynamic nature of the traffic and heterogeneity of the network.

By assessing LSTM, CNN-LSTM, RF, and KNN, the research offers useful insights into the advantages and limitations of each technique, empowering

network operators to choose the best course of action for their unique network requirements and conditions.

A significant contribution comes from using realistic data in the experiments. The research guarantees that the results are more useful in real-world settings by using operational LTE-A network data. By giving a more precise picture of LTE-A network activity and traffic patterns, this feature increases the results' validity and credibility.

1.8 Thesis Organization

The rest of the paper is organized as follows: Chapter two focuses on an overview of machine learning based resource allocation in LTE-A networks. Chapter three presented system model and data preparation. This describes about the overall processes of the study and data collection. In chapter four the results and discussions of the research are presented. The results and performance evaluations of each model is discussed in this chapter. Finally, the conclusions and the recommendation of the research are presented in chapter five.

2 Machine Learning-Based Resource Allocation in LTE-A Network

One of the major difficulties in designing and operating cellular wireless networks is allocating radio resources optimally. For particular network scenarios, these issues must be resolved while taking into account the rapidly changing wireless channels and the users' demands for QoS [29]. LTE-A, which provides great spectral efficiency, low latency, and high peak data rates, is now the leading Orthogonal Frequency Division Multiple Access (OFDMA) wireless mobile broadband technology [35]. It distributes radio resources across user equipment (UEs) using OFDMA [32].

2.1 Overview of Radio Resource Allocation in LTE/LTE-A

2.1.1 Overview of LTE/LTE-Advanced

In order to meet the International Mobile Telecommunications Advanced (IMT-Advanced) specifications established by the ITU for 4G networks, 3GPP created Release 10, also known as LTE Advanced. More than a new version, LTE Advanced is a development of Release 8 and should seem like conventional LTE networks to legacy LTE terminals because it is backward compatible with LTE technology [35, 4]. Table 2 below summarizes the main differences between the two releases.

Table 2: Summary of difference between LTE and LTE-A

	LTE(Rel.8)	LTE-A(Rel.10)
DL peak data rate	300 Mbps	1 Gbps
UL peak data rate	75 Mbps	500 Mbps
Maximum bandwidth	20MHZ	100MHZ
DL MIMO Support	up to 4	Up to 8

The Evolved Packet System(EPS), which consists of the Evolved Packet Core(EPC) and the radio portion of the core network, is the major component of LTE's overall design. Mobility Management Entity(MME), Home

Subscriber Server (HSS), Serving Gateway, and Packet-data (S-GW) and (P-GW) are the control elements that make up the EPC. Connecting with other 3rd Generation Partnership Project(3GPP) and non-3GPP networks is the responsibility of the EPC. The radio part of the network is composed of eNodeB and UE as shown in Figure 5 [4].

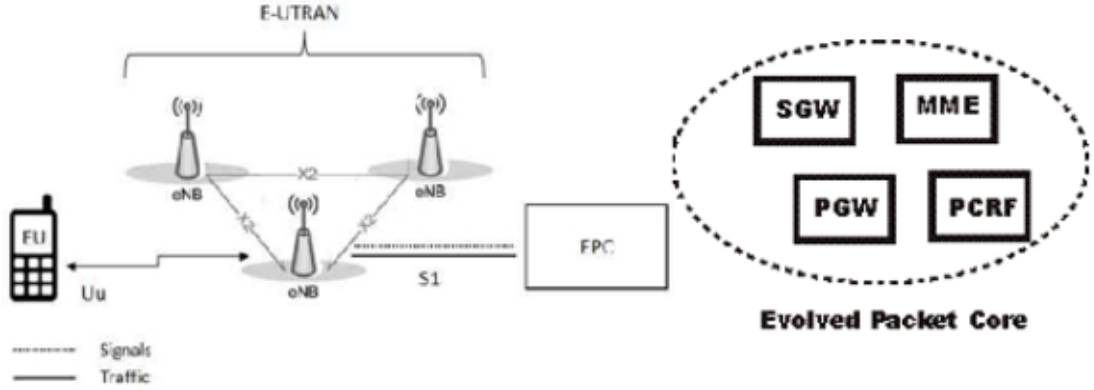


Figure 5: LTE architecture [4]

Sub-carriers are frequently distributed in groups due to the limited signaling resources; specifically, sub-carriers are grouped into Resource Blocks (RBs) of 12 contiguous sub-carriers with a 15 kHz inter-subcarrier spacing. One RB represents six or seven OFDM symbols and 0.5 ms (one time slot) in the time domain. A Scheduling Block (SB), which is made up of two consecutive RBs, is the lowest resource unit that may be assigned to a user and represents a sub-frame time length of 1 millisecond [6] as shown in Figure 6.

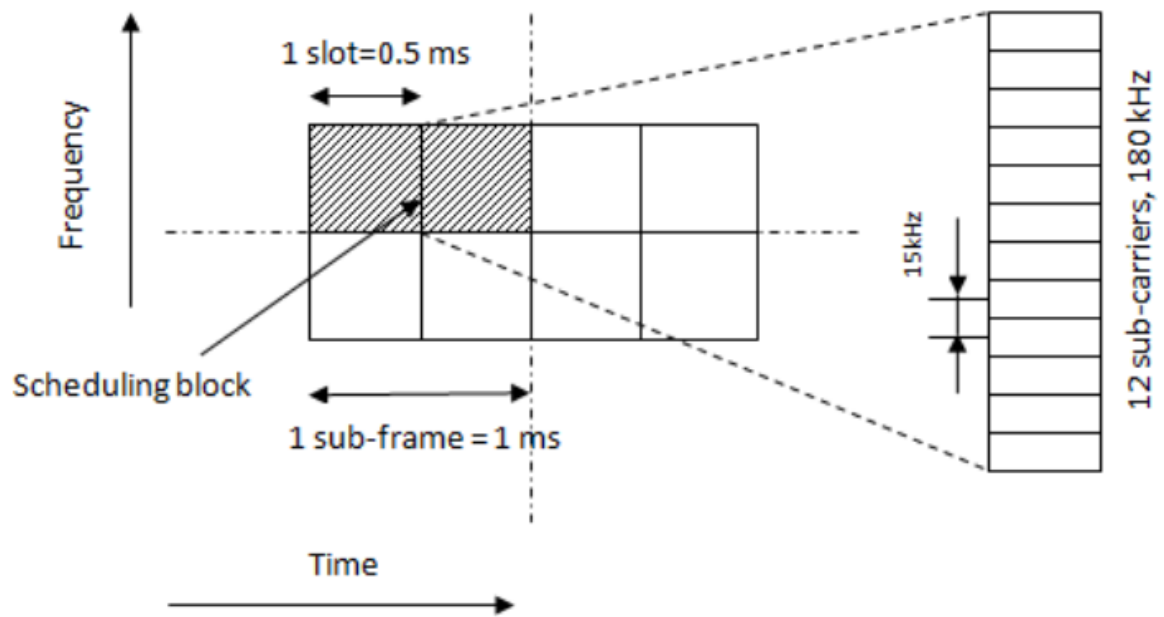


Figure 6: Scheduling block in an LTE downlink [6]

Multi-user scheduling is one of the key components of LTE systems because it is responsible for ensuring that the QoS of all active users is maintained.

In wireless connections, scheduling algorithms are crucial for ensuring QoS metrics as throughput, latency, jitter, fairness, and packet loss rate. Due to the fact that it allocates RB to various users for transmission, the scheduler is crucial to both current and future cellular communications. The connection state, the number of sessions, the reserved rates, and the state of the session queues are all taken into account by the scheduling algorithm when making a decision. The downlink transmission makes it simple to get the data needed by a scheduler set up in the base station. According to the demand of the cellular customers, the schedulers of the Medium Access Control (MAC) Layer in Evolved Node B (eNodeB) are responsible for distributing resource blocks (RB) among them. The schedulers assign both uplink and downlink resources, with eNodeB serving as their primary location of operation. At each TTI, it distributes the shared resources to

each cellular user according to a scheduling algorithm's technique.

A reference signal is sent by eNodeB to each piece of user equipment during downlink scheduling. User equipment decodes this signal, computes the CQI indicator, and then transmits it back to the base station. Through channel quality management, the eNodeB obtains CQI information. The resource block's modulation and coding mechanism (QAM, 16QAM, or 64QAM) is chosen by the MAC scheduler utilizing AMC and link adaptation. The channel is better the higher this CQI indicator rises [12].

2.1.2 Conventional Resource Allocation Strategies

Today's wireless communications are expected to ensure better QoS, as the majority of multimedia applications demand higher data rates. This necessitates the creation of strong resource management plans that not only guarantee high throughput but also efficient use of the resources already on hand. These algorithms are divided into two groups: distributed algorithms and centralized algorithms. Downlink transmission often uses centralized algorithms to effectively meet the fairness and QoS criteria. Where the decision on the distribution of resources among users is made by the eNodeB. In the case of distributed algorithms, these are frequently used in uplink transmission, where users can choose how to allocate resources [32]. One of the specific problems with LTE-A networks is resource allocation, where inefficient RB allocation could lead to higher energy usage and lower throughput [35]. To ensure QoS metrics including throughput, latency, jitter, fairness, and packet loss rate, energy and spectral efficiency in wireless links, scheduling algorithms distribute RB to multiple users [36]. The following are some of the conventional strategies to distribute radio resources to cellular users in LTE-A network.

Round Robin(RR)

A round robin scheduler is a non-aware scheduling method that recursively allots resources to users one at a time, starting with the first user. It doesn't account for the current channel situation. Because of this, it sacrifices performance to give greater user fairness. As users might be given fading channels, it reduces the system throughput as a whole. Additionally, because the CQI element is not taken into account, the overall quantity of resources are not managed effectively. Due to its simplicity of use, this scheduler is utilized by numerous systems [36].

Proportional Fair(PF)

By utilizing the channel variations, PF increases spectral efficiency and raises the level of fairness in the system. By combining CQI and level of fairness, this scheduler distributes the resource blocks to the mobile users with the optimal link quality. Because it maintains a balance between fairness and maximum cell throughput, it also enables mobile users to attain a maximum QoS. The throughput at a particular TTI and the typical throughput of the cellular customer are taken into account while allocating the RB.

Best Channel Quality Indicator(BCQI)

For a specific RB at every TTI, the best CQI scheduler allocates resource blocks to the user with the best radio connection circumstances or channel quality. For scheduling purposes, each cellular user sends an eNodeB a CQI. In the downlink channel, the eNodeB sends the cellular user a reference signal (also known as a downlink pilot). Cellular users receive this reference signal, which is used to calculate their unique CQI. The better the channel condition, the higher the CQI value. Although this scheduler can boost individual cell throughput, it has a problem with

cellular consumers' lack of fairness. The majority of cellular customers who are far from the base station (often near the cell's edge) are unlikely to ever be scheduled.

Enhanced Proportional Fair

PF chooses an MS based on the current instantaneous data rate, which is normalized by its average throughput over a given appropriate window of time. Each MS will eventually receive an equal amount of airtime because to the PF's goal of airtime fairness. As a result, each MS's throughput only depends on its own data rate and is unrelated to the data rates used by other MSs. This implies that PF results in comparatively low throughput of users on faulty channels [37] PF is an improved PF-based utility function that is anticipated to support QoS and offer a simple solution for multi-channel OFDMA networks.

Table 3: Summary of conventional resource allocation algorithms

Scheduling Policy	Effect Factor	Scheduling Priority	Usage Scenario
BCQI	Channel condition	The UE with the best channel quality is most likely to be scheduled, whereas cell edge users are unable to be scheduled at all	Increase the throughput of individual cells.
Round Robin	None	Each UE has an equal chance of being scheduled.	To confirm the highest limit of scheduling proportionality
Proportional fair	Service rate and channel condition	The UE with a low ratio of service rate to channel quality is given higher priority in scheduling.	To ensure the system's fairness and throughput.
Enhanced proportional fair	channel condition, Service rate, and QoS requirement		In operating networks Used in Ethio telecom LTE-A network

2.2 Machine Learning Basics

Machine learning is a branch of AI that enables a computer to learn from history in order to carry out a particular activity. Recently, RRM for wireless networks is an area where ML is receiving a lot of interest. In particular, during the past ten years there has been a considerable rise in the application of ML to enhance the performance of wireless networks. Supervised learning, unsupervised learning, semi-supervised learning, and reinforcement learning are the four basic categories that machine learning methods can be divided into [38, 32].

In supervised learning, certain labeled data are used to accomplish the first training. The labeled data represents a set of known inputs and outputs with matching inputs. As a result, applications using historical data are ideally suited for supervised learning techniques. The whole process is under supervision. It can be classification or regression type. Machines are

trained using unlabeled and unclassified datasets in unsupervised learning. Algorithms for unsupervised learning seek to extract features from the data, hence inferring the implied structure. There is no any supervision it extracts patterns from the data. In case of Semi-supervised learning was built taking into account both the advantages and disadvantages of supervised and unsupervised learning. The data used by semi-supervised learning systems might be labeled or unlabeled. Instead of using previous data, reinforcement learning makes use of data from the implementation. Using feedback from the environment, reinforcement learning aims to increase an agent’s performance in a certain task. As a result, the agent’s objective is to foresee the next move that will result in the highest final reward. Although reinforcement learning is unsupervised, it differs from other unsupervised learning methods in the way that it facilitates learning. Reinforcement learning aims to discover the best behaviors inside the operating medium rather than studying the structure of some data. As a result, reinforcement learning is excellent for problems involving a succession of decisions because it can capture the environment through feedback and take actions, i.e., follow a policy of actions according to the observed environment’s state [38, 44].

2.3 Deep Learning (DL)

Deep learning (DL), a subclass of machine learning (ML), uses numerous layers of nonlinear processing units to automatically extract high-level characteristics and correlations from input data. DL is a particular method for constructing and training neural networks. Unsupervised pre-trained networks (UPNs), convolutional neural networks (CNNs), recurrent neural networks (RNNs), and recursive neural networks are the four main neural network architectures used in deep learning (DL). Three learning models: supervised, unsupervised, and reinforcement learning can be utilized in conjunction with DL structures [29, 40, 30].

2.4 Deep Learning-Based Resource Allocation

DL is a particular method for constructing and training neural networks. A multi-layered feed-forward neural network that permits feature extraction and manipulation of incoming data is the foundational DL model. There are several hidden layers, an input layer, and an output layer that make up a Deep neural network (DNN)(Figure 7). Several units known as neurons are present in each layer. On the inputs, each neuron applies a non-linear action.

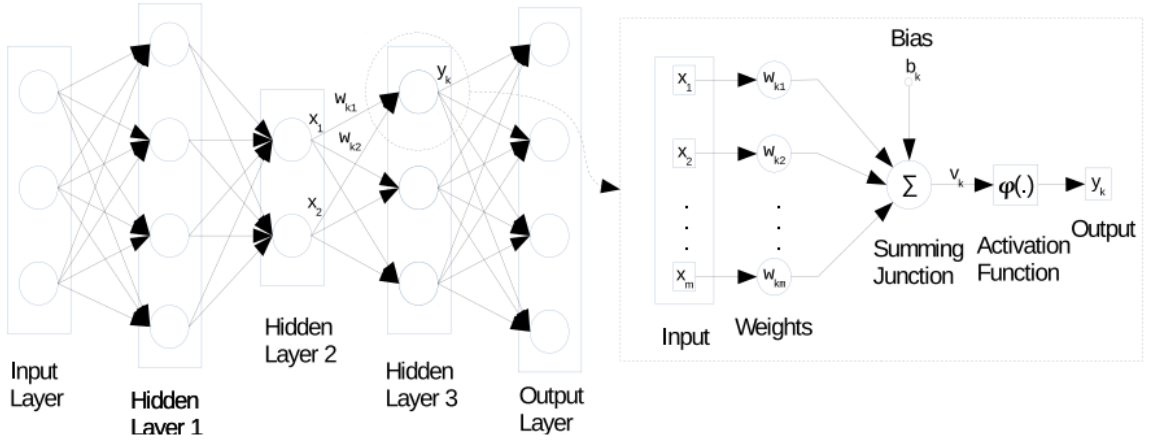


Figure 7: A deep neural network with three hidden layers [29]

Recently, DL methods have been applied to numerous wireless resource management issues (e.g. Allocation of channels and power, throughput maximization, and spectrum sharing). We divide the significant efforts into two categories: supervised DL and DRL.

2.4.1 Supervised Deep Learning

The authors in [30] Convolutional neural networks (CNNs) that are data-driven and rely on spatial features in channel gain to predict transmit power are proposed as a distributed power control approach. Maximizing throughput through power allocation is the major goal. In [39] the authors Create a trained DNN that predicts the transmit power using the communication channels as input(Figure 9). This DNN will serve as the

basis for an energy-efficient power control strategy. Maximizing energy efficiency is the main goal. There are three types of supervised deep learning algorithms.

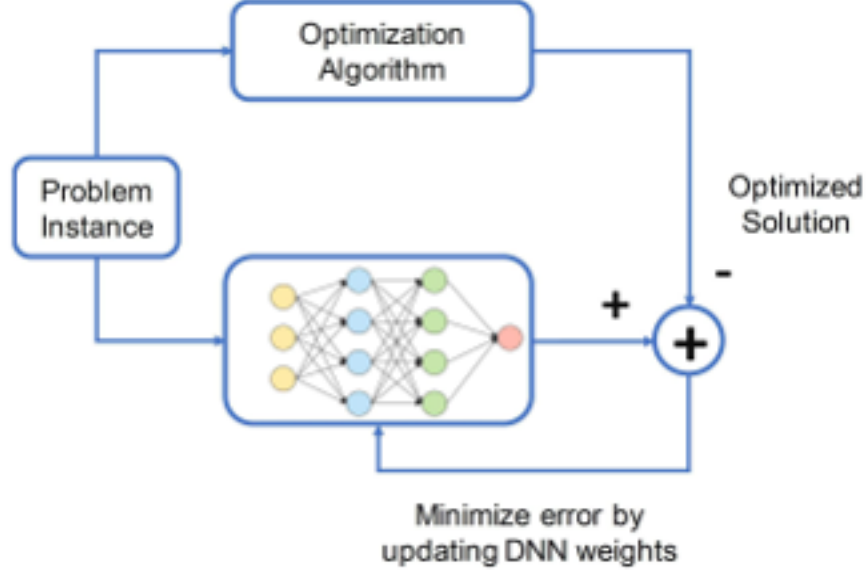


Figure 8: Train stage [40]

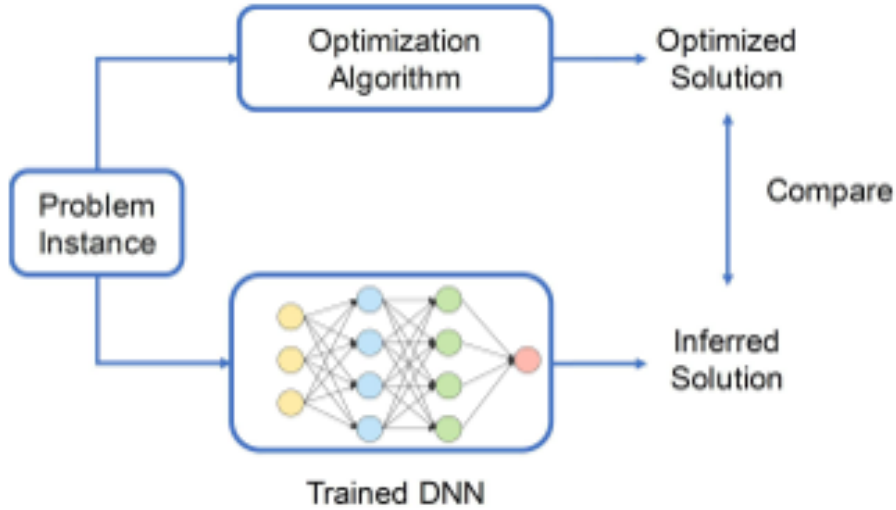


Figure 9: Test trained model [40]

2.4.2 Artificial Neural Network(ANN)

An artificial neural network is a component of a computing system created to mimic how the human brain thinks, processes information, and makes

decisions. ANN, a fundamental component of deep learning, finds solutions to issues that people might find impossible or extremely challenging. It functions Like a human brain composed of many billions of neurons, and each one is made up of a cell body that computes information by sending it to hidden neurons for ultimate output [41].

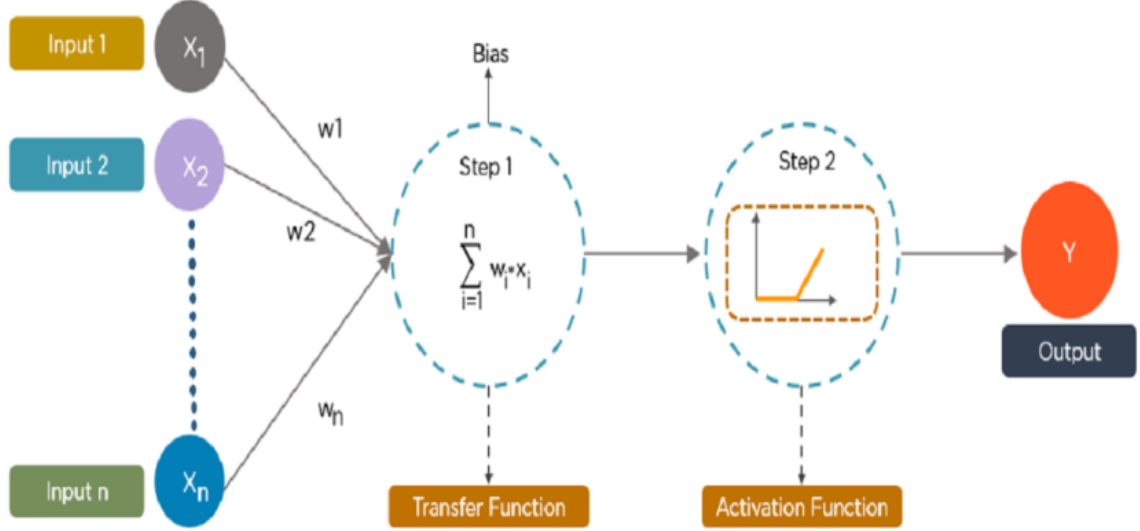


Figure 10: Artificial neural network [41]

An artificial neural network (ANN) first learns to recognize patterns using inputs provided to the input layer during the training phase. This stage involves comparing ANN's output with the real output, and any discrepancy between the two is referred to as an error. As in Figure 10 the weight and bias of the interconnection aim to minimize this error which is known as backpropagation.

2.4.3 Conventional Neural Network(CNN)

CNN is a supervised form of deep learning that is most frequently employed in computer vision and image recognition. The image is processed and key features are extracted by CNN using many layers.

2.4.4 Recurrent Neural Network(RNN)

RNN is a sort of supervised deep learning that works best with sequential data. RNN is most frequently employed in picture captioning, time-series analysis, natural language processing, handwriting recognition, and machine translation. A specific kind of RNN network called LSTM was developed to address issues with exploding and vanishing gradients as well as short-term dependency [22].

The LSTM is a particular kind of recurrent neural network (RNN) that can learn long-term dependencies and recall information for prolonged periods of time [34]. Three gates in the LSTM network Figure 11 determine which data should be added to or subtracted from the cell state, and the desired data is stored in the cell state memory. These are the input gate accomplished basic operations, forget gate selects whether information should be maintained or throw away and cell state represents the memory unit.

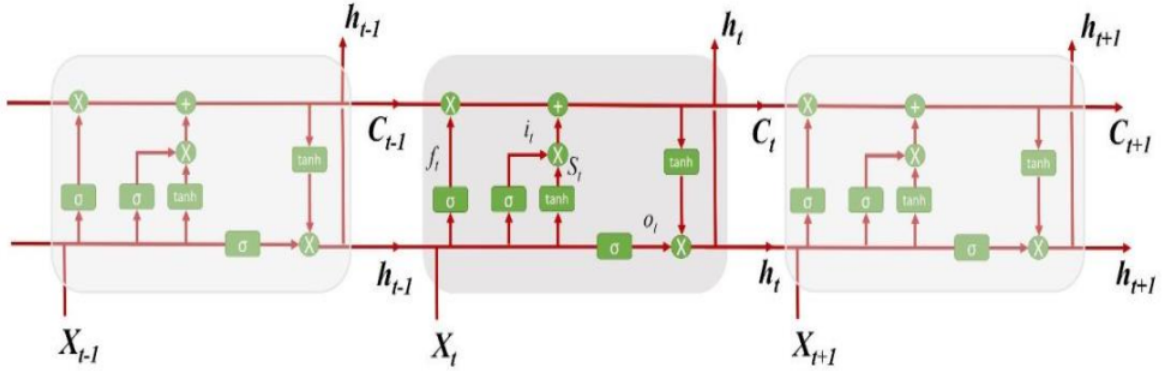


Figure 11: Structure of LSTM network [26]

2.4.5 Deep Reinforcement Learning

In reinforcement learning, a software agent interacts with the environment to learn. The agent assesses its own state as well as the status of the surrounding environment before deciding on a course of action. There are consequences for every action it takes. When an agent makes a wise decision, they are either rewarded or penalized. The agent's main respon-

sibility is to take a set of acts that will maximize the cumulative reward. A DNN is known as a DRL when it serves as an agent. Hence, DRL is the outcome of merging DNN with RL as shown in Figure 12.

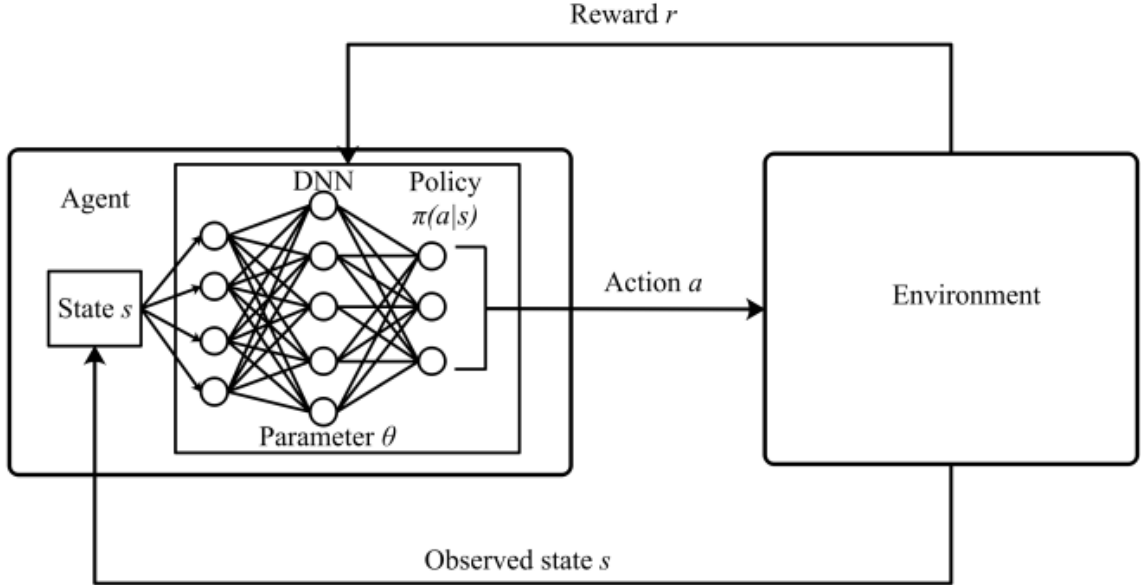


Figure 12: Deep Reinforcement Learning [29]

The three types of deep reinforcement learning (DRL) algorithms include value-based, policy-based, and actor-critical techniques. In value-based DRL, the agent learns the best course of action via the action-state value function. The most widely used value-based DRL techniques are Deep Q-learning (DQL) and SARSA. With a policy-based approach like REINFORCE learning, the agent directly optimizes the policy. In actor-critic approaches, the actor updates the policy while the critic modifies the action-value function. The actor-critic based algorithm is Deep Deterministic Policy Gradient (DDPG) [29, 40].

Q-Learning As a member of the Machine Learning (ML) family, deep reinforcement learning (DRL) has been hailed as a promising technique that can interact with the wireless environment to learn the best way to allocate resources. To reduce overall power usage and ensure that each user receives the power they need, the authors in [33] use a DQL for power

allocation in a cloud-RAN. The suggested model uses user demand as an input before allocating power to the distant radio heads. The authors in [42], in order to address bandwidth utilization and energy efficiency in an IIoT system, we suggest a DQN based approach. It creates a DQN model, which is made up of two deep neural networks (DNN) and a Q-learning model, in detail. The DNN network obtains the estimated Q-function for the Q-learning model by abstracting the features from the highly dimensional inputs. The Q-learning model may produce the Q-table and reward function based on the Q-function. Following training, the DQN model can choose the agent's suitable actions.

The authors in [43] propose dynamic neural Q learning radio resource allocation algorithm for downlink transmission in LTE-A network to make a good trade-off between fairness and throughput. The main concept behind this algorithm is how to intelligently select the best scheduling rule based on the estimated values of cell throughput and cell fairness index.

3 System Model and Data Preparation

3.1 System Model

The time and frequency domain resource allocation framework is defined by the LTE standard. The subcarriers that make up the frequency dimension are separated by 15 kHz. The time dimension is first split into ten 1 ms subframes, followed by ten 10 ms radio frames. Two 0.5 ms slots separate each subframe. One subcarrier for one OFDM symbol makes up a resource element, which is the lowest unit of resource. For the duration of one slot, a resource block has 12 constant subcarriers [4].

3.1.1 Resource Allocation

OFDMA is the foundation of the LTE downlink. The OFDMA technology offers a flexible multiple-access scheme that enables the frequency domain allocation of transmission resources of variable bandwidth to various users and the utilization of various system bandwidths without changing the core system parameters or equipment design. The LTE physical layer standards state that radio resources are organized into frequency domain RBs, which are collections of 12 subcarriers with a 15 kHz subcarrier spacing [44]. LTE-A networks in mobile broadband technology support a wide range of communication services with different requirements. One of the most significant and distinctive problems in wireless communication networks is resource allocation for each UE, wherein inefficient RB distribution could significantly reduce performance in LTE-A networks [35]. PDCCH carries the Resource Allocation information over Downlink Control Information (DCI). One of the key fields of the DCI is resource allocation.

The scheduling algorithm used for LTE-A in ethiotelecom network is enhanced proportional fair and considers channel quality indicator(CQI), QoS requirement, buffer status report and UE capability.To inform the

eNB about the channel quality in the LTE-A system, the UE uses the CQI. It ranges from 0 to 15 for the CQI reported value. It shows the degree of coding and modulation that the UE is capable of.

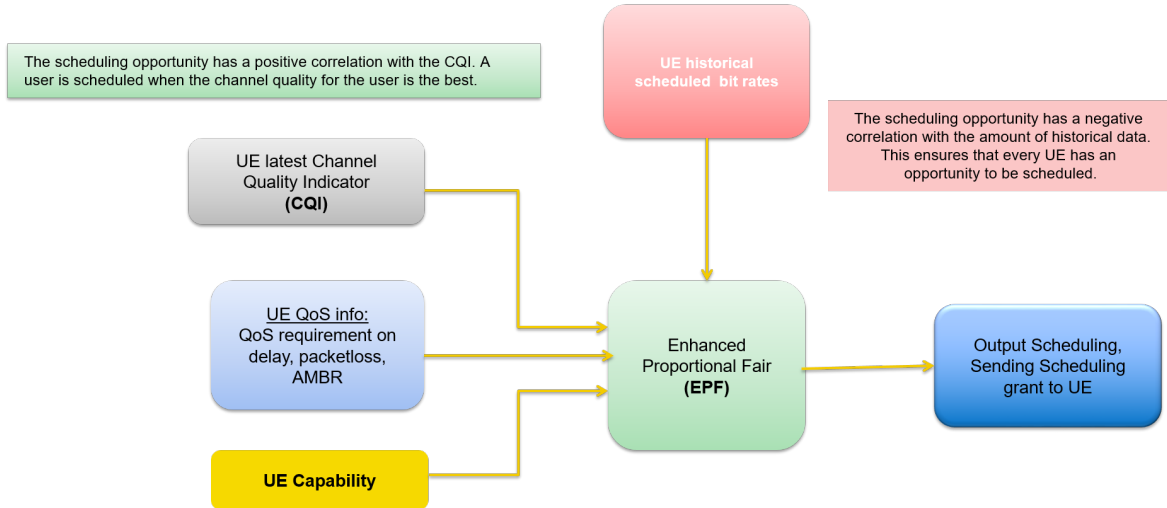


Figure 13: Enhanced proportional fair scheduling algorithm

QoS requirement is represented in terms of QoS class Identifier(QCI) that defines priority, service type, packet loss rate and delay.

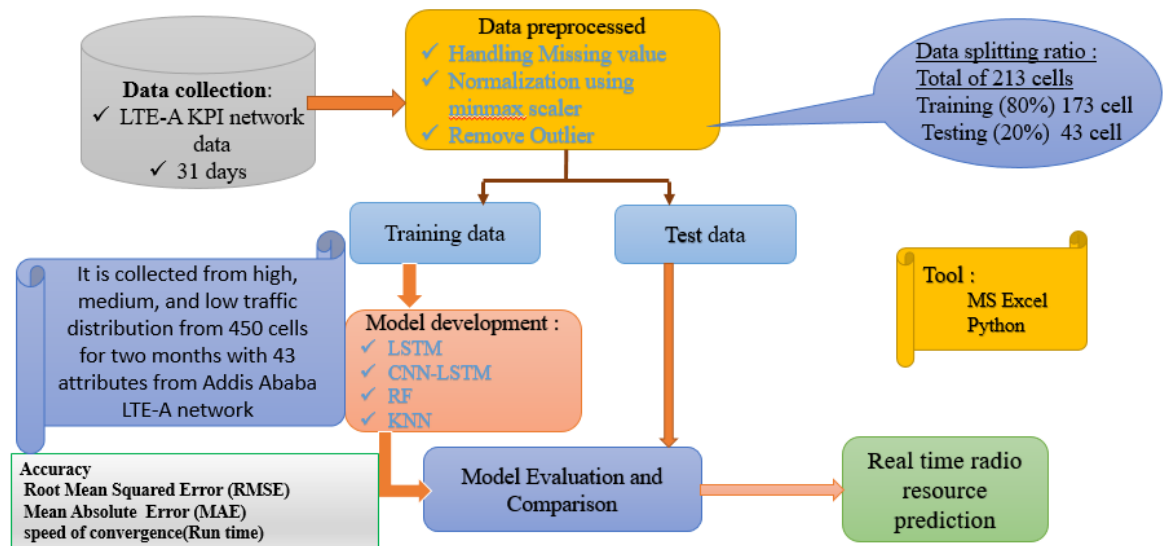


Figure 14: Overall system model

3.2 Data Preparation and Feature Extraction

3.2.1 Data Collection

In this study machine learning based resource allocation model is developed for accurate jointly allocating PRB and power for each user using historical data of ethiotelecom LTE-A network. Hourly data is collected from the Key Performance Indicator(KPI) that indicates the performance of real deployment of the LTE-A network in Addis Ababa from 450 cells for two months. The collected data is taken from high, medium and low traffic areas in Addis ababa (Figure 15)

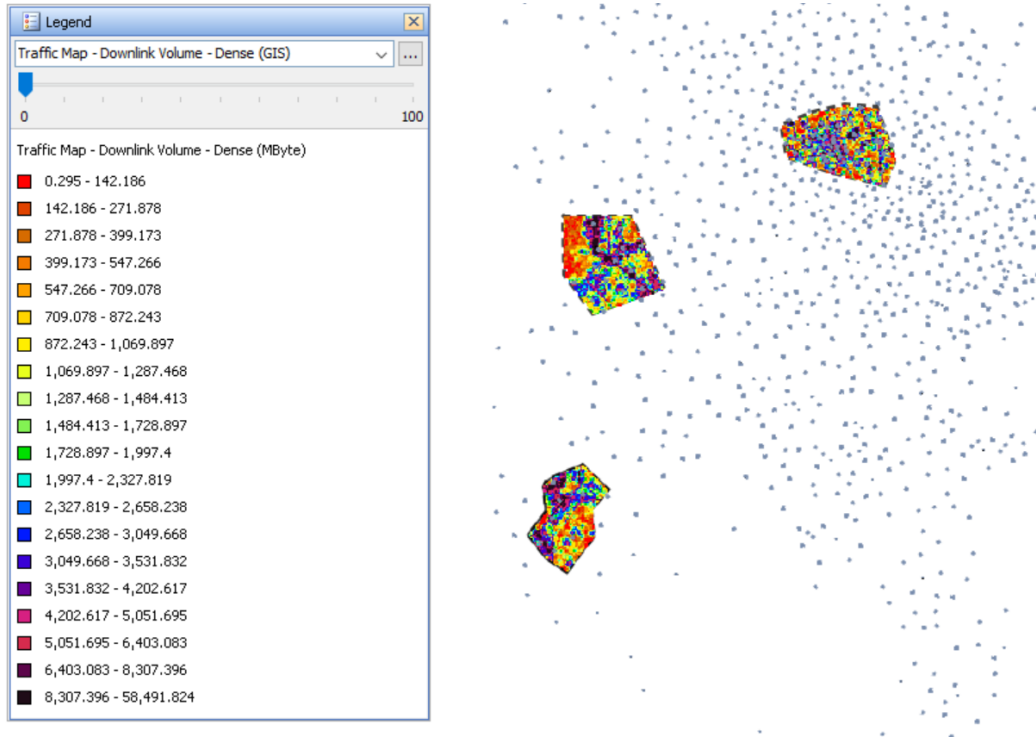


Figure 15: Traffic Distribution

This data includes channel state information(CQI), quality of service class indicator(QCI), MCs, traffic volume, average active user traffic, average user traffic, available PRB and maximum power per cell with 16 features 14 input and 2 output features. Finally, by removing the sites with missing hours one month data of 213 cell is used.

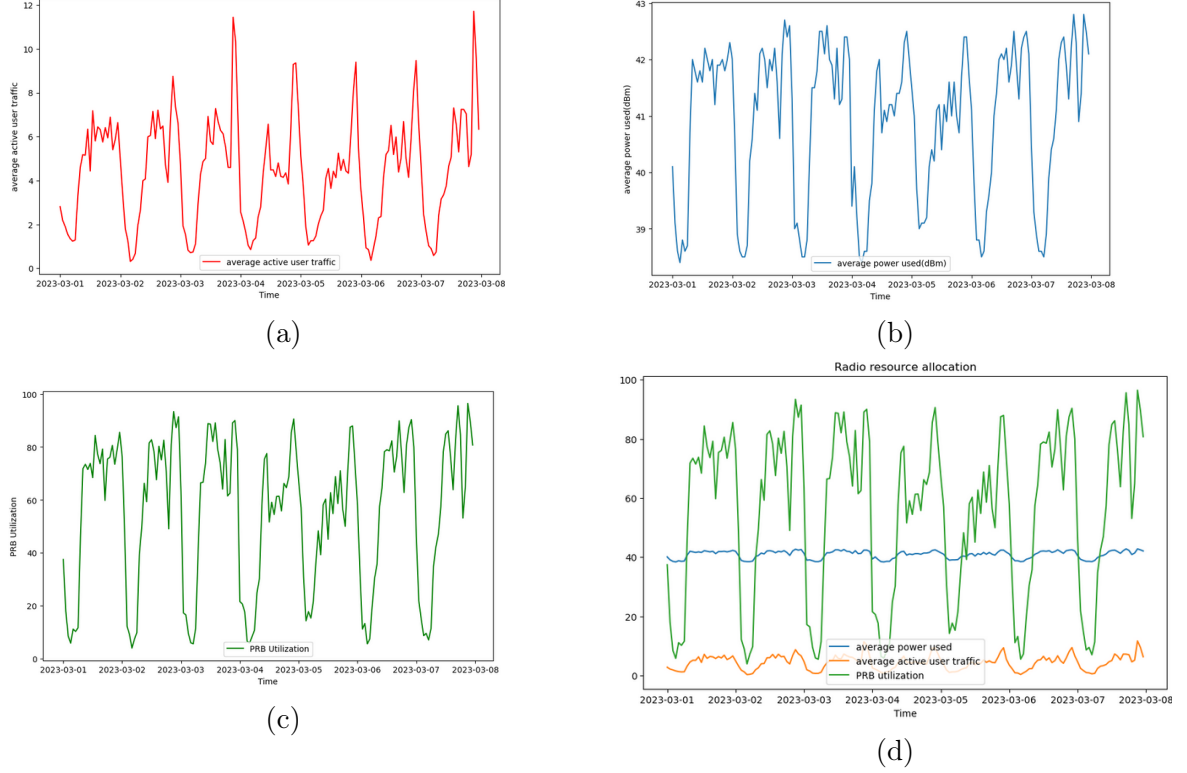


Figure 16: (a) Average active user (b) Average power used in dBm (c) DL PRB Utilization (d) The pattern of the combination traffic, power and PRB utilized

3.2.2 Data Preprocessing and Feature Extraction

Data preprocessing perform steps like data cleaning, filtering, and normalizing the dataset. Data cleaning is the process of clean the dataset by removing any missing or irrelevant data points. Normalize the data to ensure uniformity and prevent biases toward specific attributes. In case of feature extraction, extract relevant features from the dataset that can serve as input to the machine learning models and affect the resource allocation decisions. This study extracts 16 features with 14 input and 2 target features from 43 total collected features.

3.3 Machine Learning Algorithms

On the basis of prior data, machine learning models are employed to forecast UEs' resource needs. The LTE-A network's realistic data is used to

train the machine learning algorithms.

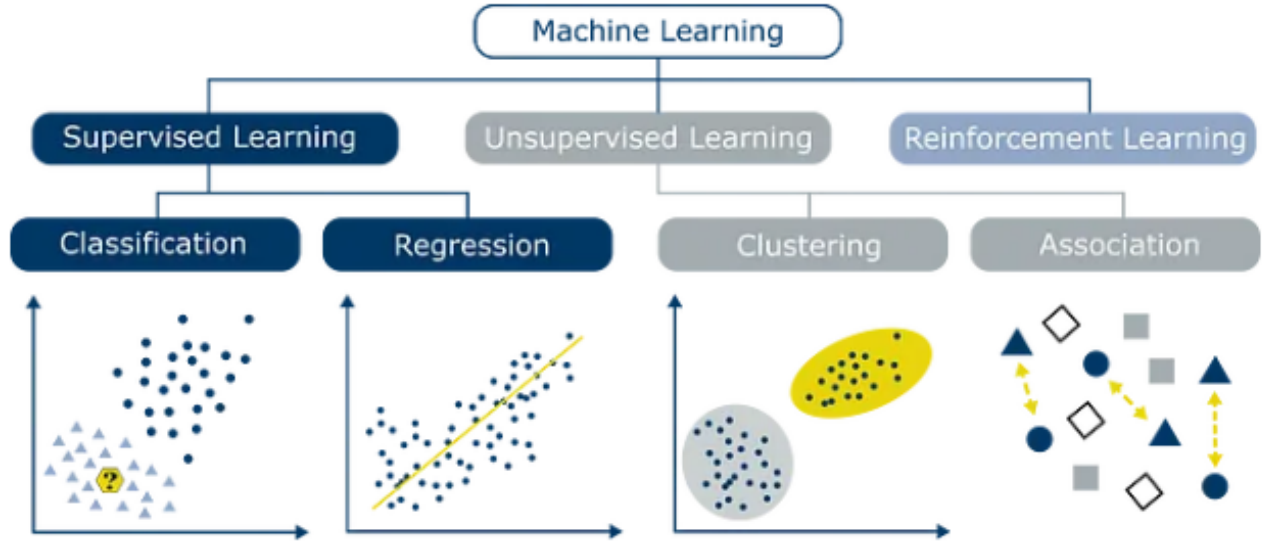


Figure 17: Types of machine learning algorithm [38]

Regression and Classification are subcategories that simply differ in terms of output value. While classification categorizes the dataset, regression produces continuous results. This thesis uses regression supervised machine learning algorithm for jointly allocating resource. The algorithms used in this thesis are Long short term memory(LSTM), the hybrid model of Conventional Neural Network and Long Short Term Memory(CNN-LSTM), Random Forest(RF) and K-Nearest Neighbor(KNN) algorithm.

3.3.1 Long Short Term Memory(LSTM)

Long Short-Term Memory (LSTM) is type of Recurrent neural networks (RNNs) which are especially made to handle sequential data, including time series, speech, and text. To address the vanishing gradient problem, that typical RNNs have, the LSTM kind of Recurrent Neural Network (RNN) architecture was developed. The goal of the LSTM approach is to develop a model that has a long-term memory while also being able to forget about unimportant data from the training set. To do this, three

significant variations from conventional RNNs are introduced [45].

- Two types of activation functions; tanh, sigmoid function
- A cell state which serves as a neuron's long-term memory;
- The neuron, also known as a cell, has a complicated structure that includes four gates that control information flow. By integrating LSTMs with a deep learning framework like TensorFlow, Keras, PyTorch, or Theano, they can be used in code.

In this thesis keras library is used to built LSTM for joint resource allocation from TensorFlow back end.

Activation Functions

Tanh is the most commonly used activation function in LSTM networks. This aids in controlling data flow through the network and prevents the phenomenon known as "exploding gradients" that ranges $<-1,1>$. The LSTM networks continue to use tanh, but they also incorporate the sigmoid activation function that ranges $<0,1>$ and this allows to discard irrelevant information.

Cell State and Hidden State

The hidden state serves as both the network's memory and the hidden layer's output in the traditional RNN architecture. A cell state is also implemented by the LSTM networks. Their hidden state acts as a working memory during the short term. The cell state, on the other hand, serves as a long-term memory for essential information from the past.

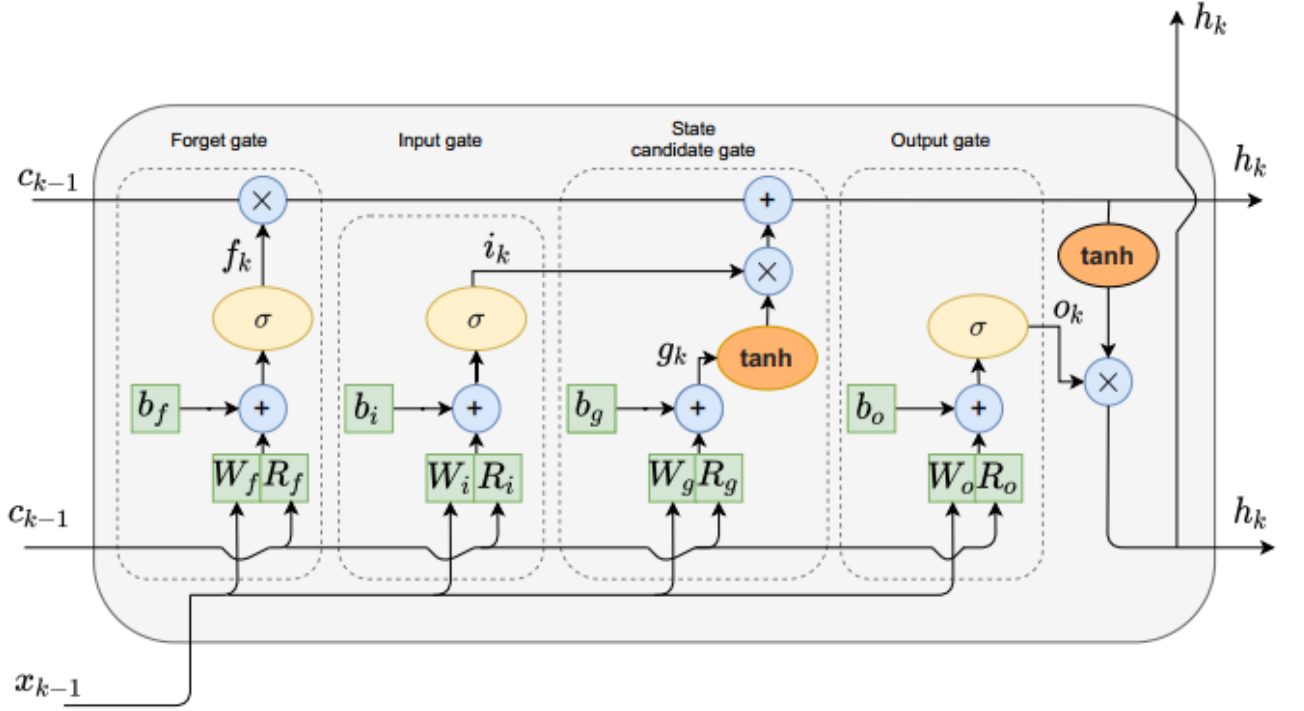


Figure 18: LSTM cell structure[46]

Gates

A technique known as gates allows the LSTM network to change the value of the cell state. Figure 18 consists of four gates:

- **Forget gate f_k** :-Used to identify and remove the unwanted data
- **Cell state**:-utilized to store information for a long time.
- **Input gate i_k** :-Updates values by passing them through the sigmoid function with values from the current input and the prior hidden state as inputs.
- **Output gate o_k** :-it determines the value of the next hidden layer.

$$f_k = \sigma(W_f [h_{k-1}, X_k] + b_f) \quad (1)$$

$$i_k = \sigma(W_i [h_{k-1}, X_k] + b_i) \quad (2)$$

$$o_k = \sigma(W_o [h_{k-1}, X_k] + b_o) \quad (3)$$

$$c_k = \tanh(W_c [h_{k-1}, X_k] + b_c) \quad (4)$$

$$c_k = i_k \times c_k + f_k \times c_{k-1} \quad (5)$$

$$h_k = o_k \times \tanh c_k \quad (6)$$

where

- W_f, W_i, W_o, W_c represents the weight matrix associated with the forget gate, input gate output gate and cell state.
- $[h_{k-1}, X_k]$ denotes the concatenation of the current input and the previous hidden state.
- b_f, b_i, b_o, b_c is the bias with the forget gate, input gate, output gate and cell state.
- σ is the sigmoid activation function.

This thesis uses Rectified linear unit(ReLU) activation function in order to remove completely the drawbacks like vanishing gradient problems of TanH and sigmoid functions. Besides, Adam optimizer is used.

3.3.2 Conventional Neural Network and Long Short Term Memory(CNN-LSTM)

In hybrid model CNN model combined with an LSTM backend can be used to interpret the input subsequences supplied as a sequence to the LSTM model.

The architecture of CNN-LSTM was developed with CNN layers on the front end as shown in Figure 19. Its goal is to extract the input dataset features. The outputs of CNN layers were then passed to the LSTM and a dense layer added at the output to aid with sequence prediction [47].

For reading throughout the subsequence, the CNN model first uses a convolutional layer that calls for a number of filters and a kernel size to be supplied. The quantity of filters corresponds to how many times the input

sequence has been read or interpreted. Each 'read' operation on the input sequence includes a certain number of time steps, which is known as the kernel size.

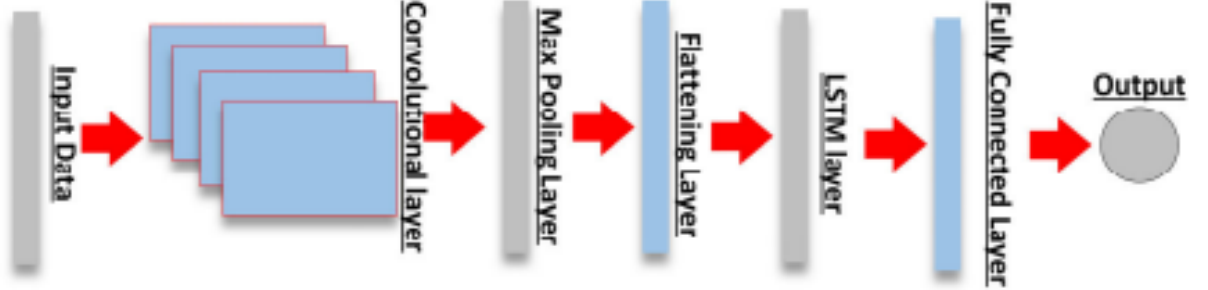


Figure 19: The structure of CNN-LSTM model [47]

The max pooling layer, which reduces the filter maps to half their size and incorporates the most important features, is added after the convolution layer. These structures are then reduced to a single one-dimensional vector and used as the LSTM layer's only input time step.

3.3.3 Random Forest(RF)

The bootstrap aggregation-based bagging approach is the foundation of the random forest algorithm , which consists of a collection of decision trees as shown in Figure 20. The goal is to prevent overfitting by taking into account numerous trees and using each one to make predictions. Each tree in the RF model is trained to fit to a random sample of the original data's rows and columns. The average of each tree response is used for regression tasks. A prediction on a regression problem is the mean of the predictions made by all of the ensemble's trees. RF is resilient to noisy data since it is nonlinear. Additionally, the approach was created to decrease error variance while just slightly increasing bias [48, 49]. Our random forest regression model will be trained using the RandomForestRegressor function of the sklearn package. It chooses from a wide range of parameters for this model, according to the RandomForestRegressor

documentation. The following is a list of some of the crucial variables:

- **N-estimators**:- define how many decision trees planned to use in the model
- **Criterion**:- define loss function to evaluate the model such as Mean Absolute Error (MAE) and Mean Squared Error (MSE). The default value is MSE.
- **Max-depth** :- Sets the maximum possible depth of each decision tree
- **Max-features**:- The important features that the model will take into account when deciding on a split
- **Bootstrap**:- True is its default value, meaning the model follows bootstrapping principles that is technique of randomly selecting dataset subsets over a predetermined number of rounds.
- **Max-samples**:- In order for this parameter to be applicable, bootstrapping must be set to True. The biggest size of each sample for each tree is set by this value when the value is True.

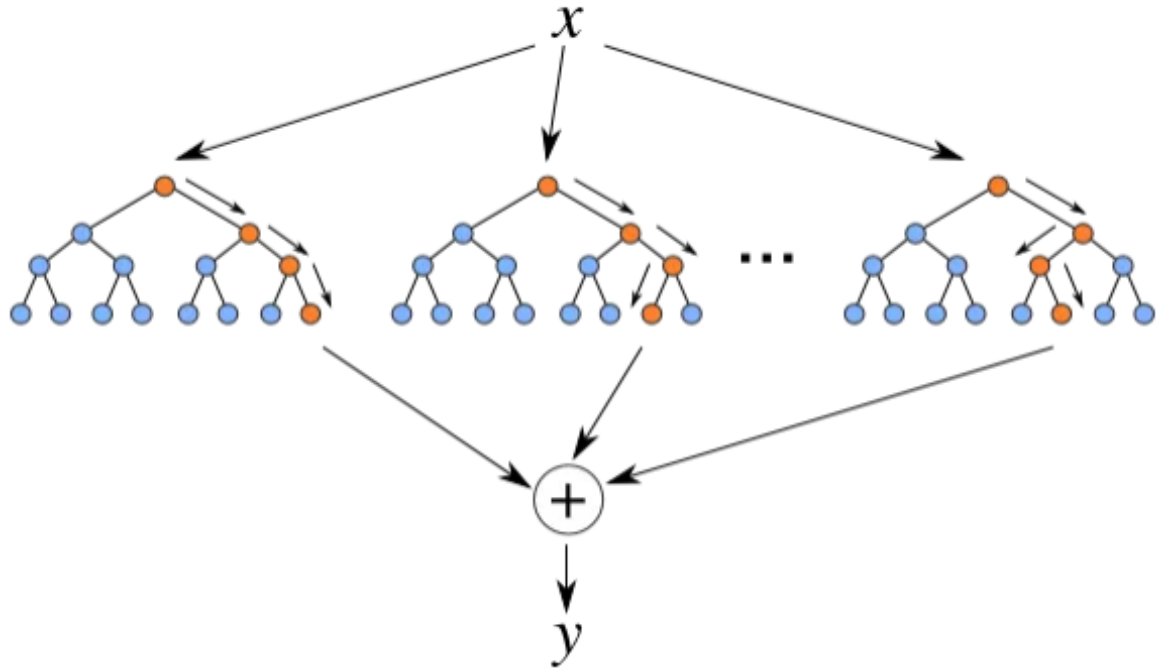


Figure 20: The structure of random forest algorithm [48]

3.3.4 K-Nearest Neighbor(KNN)

KNN is the supervised learning technique used for both regression and classification. It is a straightforward technique that stores all of the possible examples and forecasts the numerical target based on a similarity metric (such as distance functions) [50]. Working with numbers is the sole distinction. Accordingly, the `KNeighborsRegressor()` algorithm from the Sklearn library will calculate the regression for the dataset, take the `n-neighbors` parameter with the desired value, check the results of those neighbors, and average the results to give you an estimated result.

The KNN algorithm believes that related things are located nearby. In other words, related things are located close to one another. KNN uses the arithmetic to encapsulate the idea of similarity (also called distance, proximity, or closeness). Other methods of estimating distance exist, and depending on the issue at hand, one method may be better. The common and well-known option, however, is the direct distance (also known as the

Euclidean distance).

3.4 Performance Metrics

The trained model's performance is evaluated using appropriate regression evaluation metrics such as accuracy, RMSE, MAE and speed of convergence (running time). The developed Machine learning models is a crucial stage in determining the effectiveness and correctness of the model based on the metric scores.

3.4.1 Accuracy

Measures how often the model predicts accurately.

3.4.2 Root Mean Square Error(RMSE)

The square root of the average squared difference between the actual and predicted output values is known as the mean squared error, or RMSE.

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (y_i - y_p)^2} \quad (7)$$

3.4.3 Mean Absolute Error(MAE)

MAE is the absolute difference between the actual and predict values that is the absolute error but it is the mean absolute of the complete dataset.

$$MAE = \frac{1}{N} \sum_{i=1}^N |(y_i - y_p)| \quad (8)$$

3.4.4 Mean Absolute Percentage Error(MAPE)

MAPE is another measure used to assess the effectiveness of the regression model. As a ratio, it is used to express precision according to the following

equation:

$$MAPE = \frac{1}{N} \sum_{i=1}^N \frac{|(y_i - y_p)|}{y_i} \times 100\% \quad (9)$$

where N is the number of data series, y_i and y_p are the actual test sample and predicted value at time t, respectively

4 Results and Discussions

The goal of this study is to better understand how machine learning-based resource allocation algorithms perform in LTE-A networks. The findings of this study will bring insight into which algorithms perform best in terms of accuracy, RMSE, MAE and running time when it comes to allocating resources in LTE-A networks. Before training the model, the values of the hyperparameters should be initialized, and the grid search technique is used to find the model's optimal configuration.

4.1 Results

4.1.1 Long Short Term Memory

LSTM is one of the models used to perform analysis on sequential data. Grid search optimization is used to optimize the hyperparameters of the model. Hyperparameters are parameters that are set by the user rather than being learned during training. Because the grid search approach is computationally expensive even for a small number of parameters in the sample space, some major parameters, such as the kind of optimizer, loss function, and activation function, are not included in the grid search to reduce the amount of time spent searching. These hyperparameters were chosen using the following criteria.

- **Optimizer:** Due to the Adam optimizer's efficiency for stochastic gradient optimization, employs adaptive learning rates, is computationally effective, and requires less memory, which promotes rapid convergence.
- **Loss function:** MSE was utilized as the loss function in this study since it is the most used loss function for regression issues and minimizes the variance between real and predicted values.
- **Activation function:** It is selected based on regression that is Relu

Dropout

There should be a dropout layer between each LSTM layer. By omitting neurons that were chosen at random during training, such a layer reduces the sensitivity to the particular neurons' weights and helps prevent overfitting. The dropout rate should be kept low (up to 50%), but 20% is an acceptable starting point. The 20% number is generally regarded as the ideal balance between avoiding model overfitting and maintaining model correctness [28, 25].

Epochs

The number of epochs refers to how many times the neural network has processed the complete training dataset. Similar to how hidden layers and units are counted, the best number of epochs can be calculated. Once the performance on a validation set stops getting better, progressively increase the number of epochs from a low starting point [32, 28, 26]. Table 4 the number of epochs is selected based on this by trail and error.

Learning Rate

How frequently the network updates its parameters is determined by this hyperparameter. A greater learning rate speeds up the learning process, but there is a chance that the model won't converge or even diverge (a situation during training where the loss settles to within an error range around the final result) [22].

Batch size

This hyperparameter specifies how many samples must be examined before the model's internal parameters are changed. In comparison to smaller ones for the same amount of samples, larger sizes provide larger gradient steps. The generally accepted default batch size is 32 [32].

Hidden Layer

As a general rule, one hidden layer will solve the majority of easy issues, while two layers will solve most moderately complicated ones [34]. As in Table 4 this paper tries to evaluate using three values.

Table 4: Hyperparameter and its value

hyperparameter	value	optimal value
Dropout	0.0, 0.2, 0.4	0.2
epochs	10, 20, 100	20
Learning Rate	0.01, 0.001, 0.0001	0.0001
neuron	32, 64, 128	128
Batch size	16, 32, 64	32
hidden layer	1, 2, 3	2

The results of the model obtained using the optimal value are presented in Table 5.

Table 5: LSTM model performance metrics value

Metrics	value
Accuracy	98.56%
RMSE	1.0520
MAE	0.7654
Running time	278.96sec

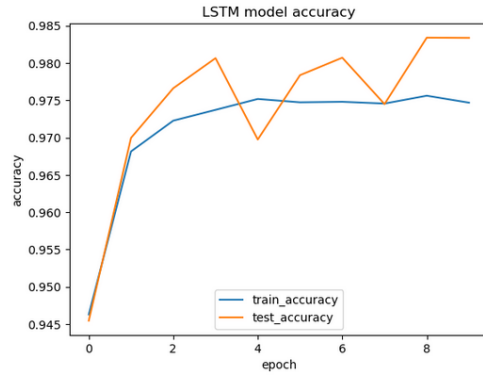
```

Epoch 1/20
7905/7905 - 21s - loss: 0.0152 - accuracy: 0.9376 - val_loss: 5.0485e-04 - val_accuracy: 0.9671 - 21s/epoch - 3ms/step
Epoch 2/20
7905/7905 - 21s - loss: 0.0023 - accuracy: 0.9600 - val_loss: 5.1129e-04 - val_accuracy: 0.9717 - 21s/epoch - 3ms/step
Epoch 3/20
7905/7905 - 22s - loss: 0.0016 - accuracy: 0.9673 - val_loss: 3.7162e-04 - val_accuracy: 0.9798 - 22s/epoch - 3ms/step
Epoch 4/20
7905/7905 - 22s - loss: 0.0012 - accuracy: 0.9698 - val_loss: 3.9429e-04 - val_accuracy: 0.9771 - 22s/epoch - 3ms/step
Epoch 5/20
7905/7905 - 22s - loss: 9.6146e-04 - accuracy: 0.9715 - val_loss: 3.8719e-04 - val_accuracy: 0.9729 - 22s/epoch - 3ms/step
Epoch 6/20
7905/7905 - 22s - loss: 8.0696e-04 - accuracy: 0.9726 - val_loss: 3.5299e-04 - val_accuracy: 0.9782 - 22s/epoch - 3ms/step
Epoch 7/20
7905/7905 - 22s - loss: 7.2177e-04 - accuracy: 0.9727 - val_loss: 3.6032e-04 - val_accuracy: 0.9790 - 22s/epoch - 3ms/step
Epoch 8/20
7905/7905 - 22s - loss: 6.6527e-04 - accuracy: 0.9730 - val_loss: 3.7629e-04 - val_accuracy: 0.9814 - 22s/epoch - 3ms/step
Epoch 9/20
7905/7905 - 22s - loss: 6.3351e-04 - accuracy: 0.9738 - val_loss: 3.9985e-04 - val_accuracy: 0.9832 - 22s/epoch - 3ms/step
Epoch 10/20
7905/7905 - 21s - loss: 6.0664e-04 - accuracy: 0.9726 - val_loss: 3.7053e-04 - val_accuracy: 0.9814 - 21s/epoch - 3ms/step
Epoch 11/20
7905/7905 - 22s - loss: 5.8578e-04 - accuracy: 0.9737 - val_loss: 4.2660e-04 - val_accuracy: 0.9832 - 22s/epoch - 3ms/step
Epoch 12/20
7905/7905 - 22s - loss: 5.6743e-04 - accuracy: 0.9735 - val_loss: 4.3800e-04 - val_accuracy: 0.9799 - 22s/epoch - 3ms/step
Epoch 13/20
7905/7905 - 22s - loss: 5.5155e-04 - accuracy: 0.9734 - val_loss: 4.1124e-04 - val_accuracy: 0.9831 - 22s/epoch - 3ms/step
Epoch 14/20
7905/7905 - 22s - loss: 5.3103e-04 - accuracy: 0.9741 - val_loss: 3.9927e-04 - val_accuracy: 0.9841 - 22s/epoch - 3ms/step
Epoch 15/20
7905/7905 - 22s - loss: 5.2064e-04 - accuracy: 0.9746 - val_loss: 3.3313e-04 - val_accuracy: 0.9847 - 22s/epoch - 3ms/step
Epoch 16/20
7905/7905 - 22s - loss: 5.1602e-04 - accuracy: 0.9739 - val_loss: 3.6130e-04 - val_accuracy: 0.9844 - 22s/epoch - 3ms/step
Epoch 17/20
7905/7905 - 22s - loss: 5.0662e-04 - accuracy: 0.9738 - val_loss: 3.4115e-04 - val_accuracy: 0.9848 - 22s/epoch - 3ms/step
Epoch 18/20
7905/7905 - 22s - loss: 4.9398e-04 - accuracy: 0.9739 - val_loss: 3.2853e-04 - val_accuracy: 0.9848 - 22s/epoch - 3ms/step
Epoch 19/20
7905/7905 - 22s - loss: 4.8628e-04 - accuracy: 0.9740 - val_loss: 3.4225e-04 - val_accuracy: 0.9815 - 22s/epoch - 3ms/step
Epoch 20/20
7905/7905 - 22s - loss: 4.7468e-04 - accuracy: 0.9744 - val_loss: 3.4015e-04 - val_accuracy: 0.9859 - 22s/epoch - 3ms/step

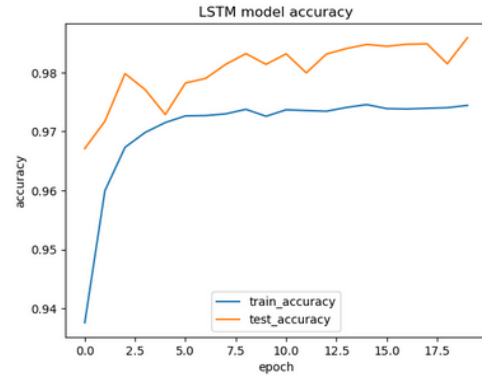
```

Figure 21: LSTM model output of each epoch train accuracy, test accuracy, train loss and test loss

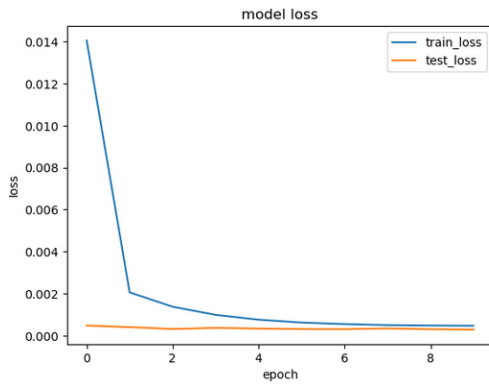
As in Figure 21 the train and validation loss decreases and the train and validation accuracy is increasing as the number of epoch increases up to some point, after that the model learns noisy data and decreases its performance.



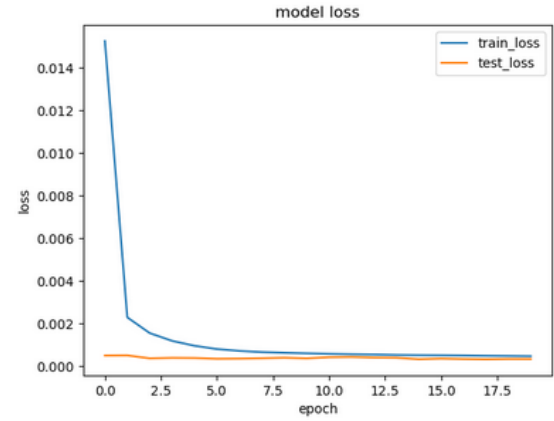
(a)



(b)



(c)



(d)

Figure 22: (a) LSTM accuracy at 10 epochs (b) LSTM accuracy at 20 epochs (c) LSTM train test loss at 10 epochs (d) LSTM train test loss at 20 epochs

As shown in Figure 22 train loss is high at the first epoch as compared to the test loss but as the number of epoch increases the train and test loss are almost similar that is below 0.002. As in Figure 22 a and b the accuracy at epoch 20 is higher than at epoch 10. This indicates that the accuracy increases upto some level of epoch, after that it decreases because the model can not learn new idea or learns noise.

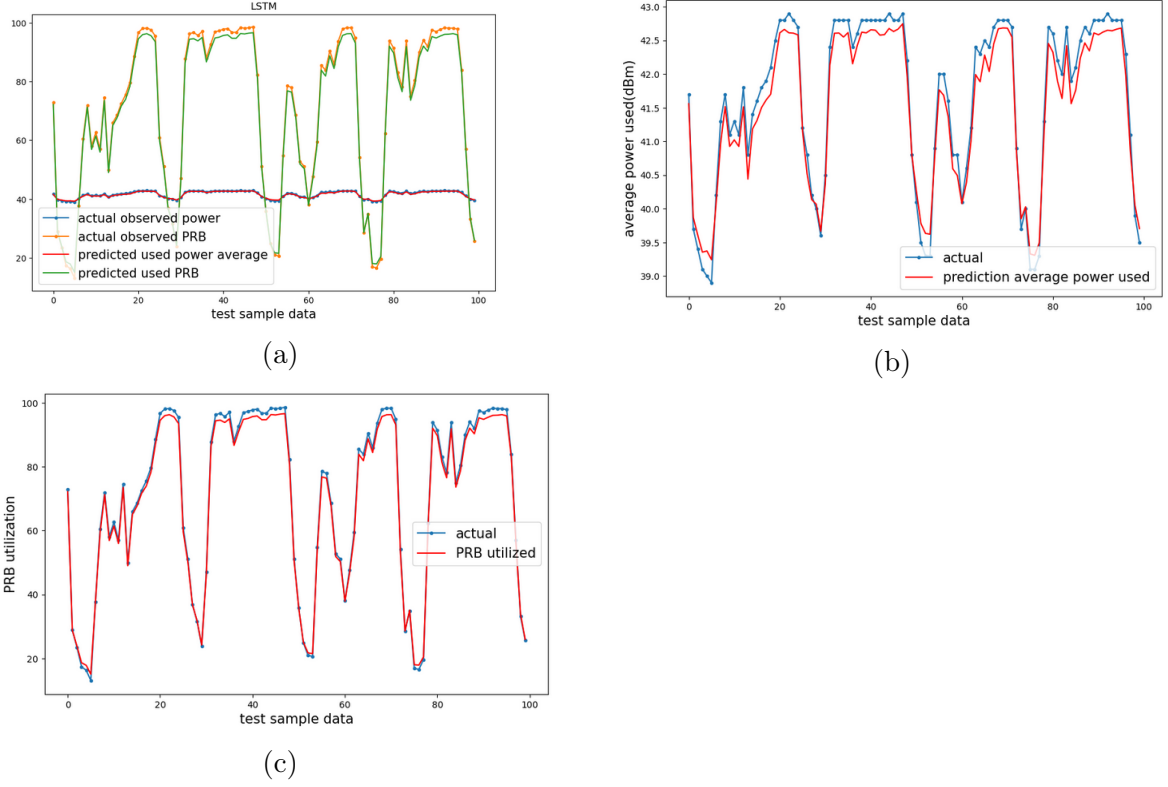


Figure 23: (a) Predicted using LSTM model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized

The Figure 23 shows the graphical comparison of the predicted and observed value of the target variable. The predicted value almost fit the actual observed value with the root Mean Square Error (RMSE) and Mean Absolute Error(MAE) of 1.0520 and 0.7654 as shown in Table 5.

4.1.2 Conventional Neural Network-Long Short Term Memory (CNN-LSTM)

It is possible to combine the CNN model's ability to automatically learn and extract features from raw sequence data with the LSTM model, which more effectively captures long-term and short-term dependency of temporal features as discussed in the previous section. The CNN model accepts input data sequences and extracts significant feature information, while the LSTM model connected in tandem determines and provides an explanation. Because the grid search approach is computationally expensive

even for a small number of parameters in the sample space, some major parameters, such as the kind of optimizer, loss function, and activation function, are not included in the grid search to reduce the amount of time spent searching. The hyperparameters were chosen using the following criteria.

- Optimizer: Adam
- Loss function: MSE
- Activation function: relu

In this section, the CNN-LSTM prediction model results are discussed, and the model performance is evaluated using the test dataset. The set of hyperparameters, possible values and optimal values are listed in Table 6. The results listed on Figure 24, 25 and Table 7 below are based on the optimal value.

Table 6: CNN-LSTM model hyperparameter possible value

Hyperparameters	hyperparameter values	Optimal values
Filters	32, 64, 128	64
Neurons	32,64,128	128
Batch size	16, 32, 64	32
Epochs	10,20,100	10

The simulation results value at the optimal values of the hyperparameters are presented in Table 7

Table 7: CNN-LSTM model performance metrics value

Metrics	value
Accuracy	98.49%
RMSE	0.95
MAE	0.66
Running time	743.07sec

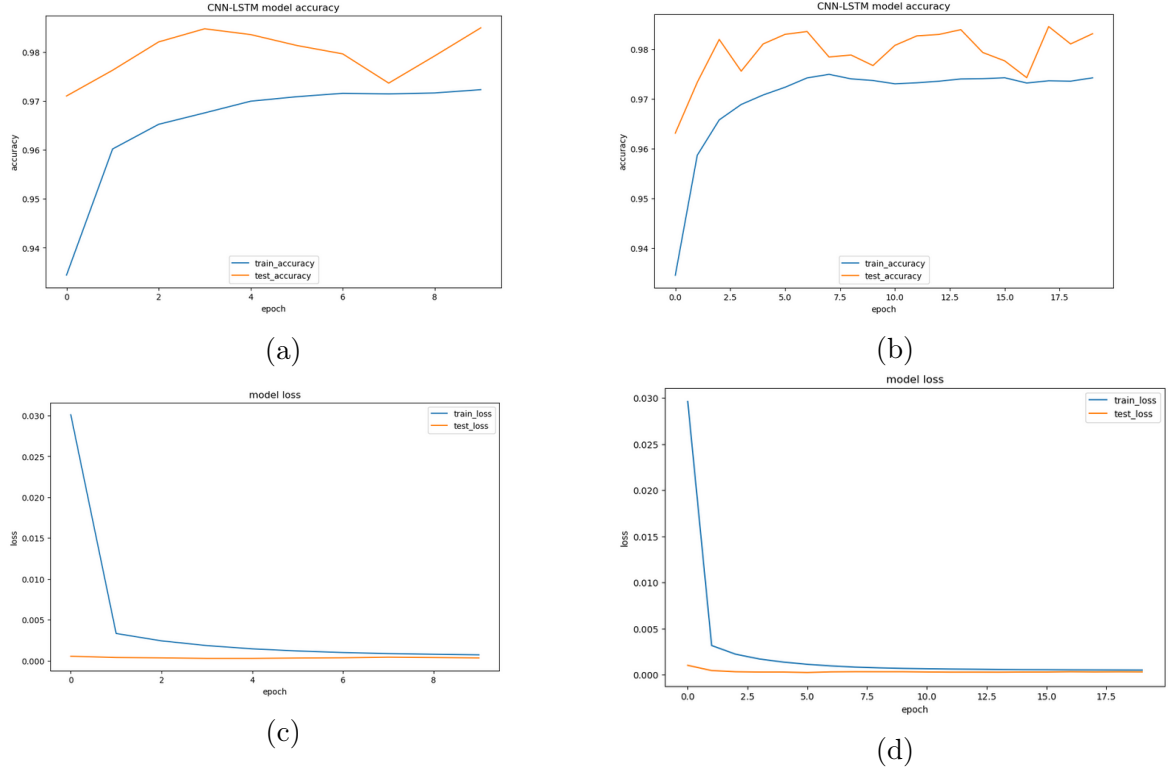


Figure 24: (a) CNN-LSTM accuracy at 10 epochs (b) CNN-LSTM accuracy at 20 epochs (c) CNN-LSTM train test loss at 10 epochs (d) CNN-LSTM train test loss at 20 epochs

Figure 24 above indicates that train and test accuracy at 10 epochs and 20 epochs, at the start train accuracy is much lower than test accuracy. After the number of epochs increase to some level the train and test accuracy increases slowly this is due to as the number of epochs increases the model does not get new idea to learn it only learns noise with the accuracy of 98.49% and 98.31%. However, adding more hidden layers over time will not dramatically increase accuracy, and in some situations, the model may begin to pick up on noisy features. In the case of loss, train loss is high at the first epoch as compared to the test loss but as the number of epoch increases the train and test loss are almost similar that is below 0.002.

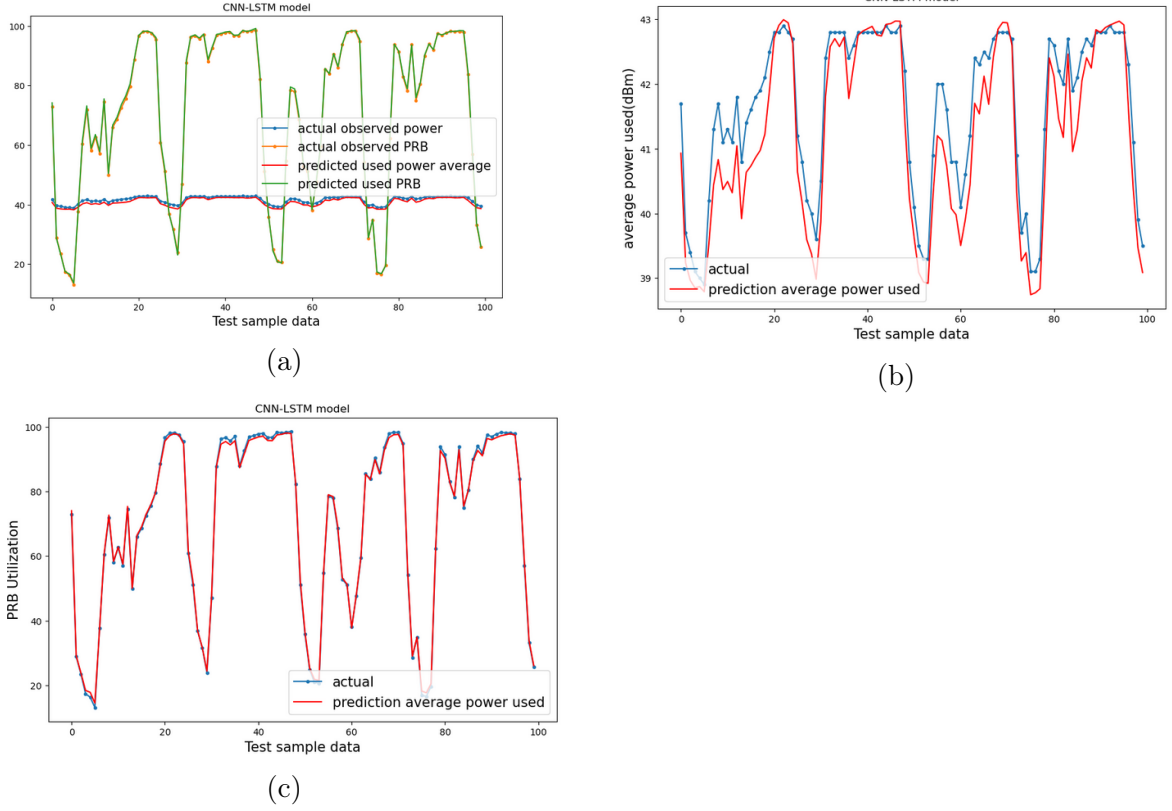


Figure 25: (a) Predicted using CNN-LSTM model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized

Figure 25 above shows as the predicted value almost fit the actual observed value with the root mean square error (RMSE) and mean absolute error (MAE) of 0.95 and 0.66 as shown in Table 7.

4.1.3 Random Forest

When the hyperparameters are adjusted, the flexible random forest algorithm among other ML algorithms produces better results. A forest is created by combining various decision trees in a sort of ensemble learning.

Table 8: Random forest model hyperparameter possible value

Hyperparameters	hyperparameter values	Optimal values
n-estimators	50,100,150,200,300	50

Table 9: RF model evaluation result of possible value

n-estimators	Accuracy	RMSE	MAE	Runing time
50	98.22%	0.296	0.09	113.05
100	98.25%	0.294	0.08	230.26
150	98.24%	0.294	0.08	361.03
200	98.22%	0.296	0.08	457.44
300	98.24%	0.294	0.08	693.82

From Table 9 the n-estimator that give better performance from the given possible combinations is at n-estimator 100. The value of each metrics is accuracy 98.25%, RMSE 0.294, MAE 0.08 and the running time 230.26sec.

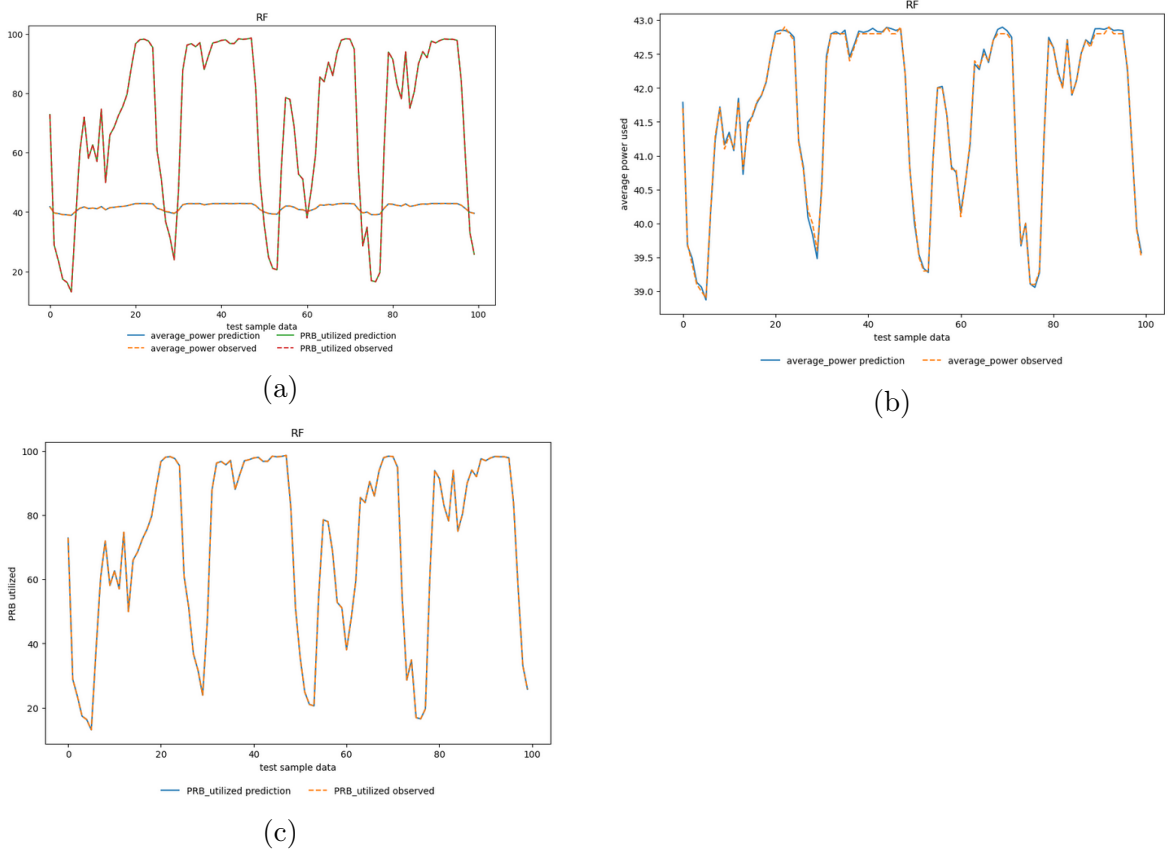


Figure 26: (a) Predicted using RF model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized

4.1.4 K-Nearest Neighbor(KNN)

Based on how similar the samples are, the kNN algorithm generates decisions. A distance function is used to measure similarity. Euclidean distance was employed as the similarity metric because we are forecasting a real-valued number and we used kNN as the regressor. Two hyperparameters such as n-neighbors from 1 to 20 and weight uniform and Euclidean distance are used and select best combination using Sklearn function GridSearchCV(). It selects best combinations n_neighbors 4 and weight distance.

Table 10: KNN model performance metrics value

Metrics	value
Accuracy	90%
RMSE	0.95
MAE	0.56
Running time	1.58sec

As shown in Figure 27b the predicted average power used is shifted from the observed average power used but in Figure 27c the predicted PRB is almost fitted and the error between predicted and actual observed is very small.

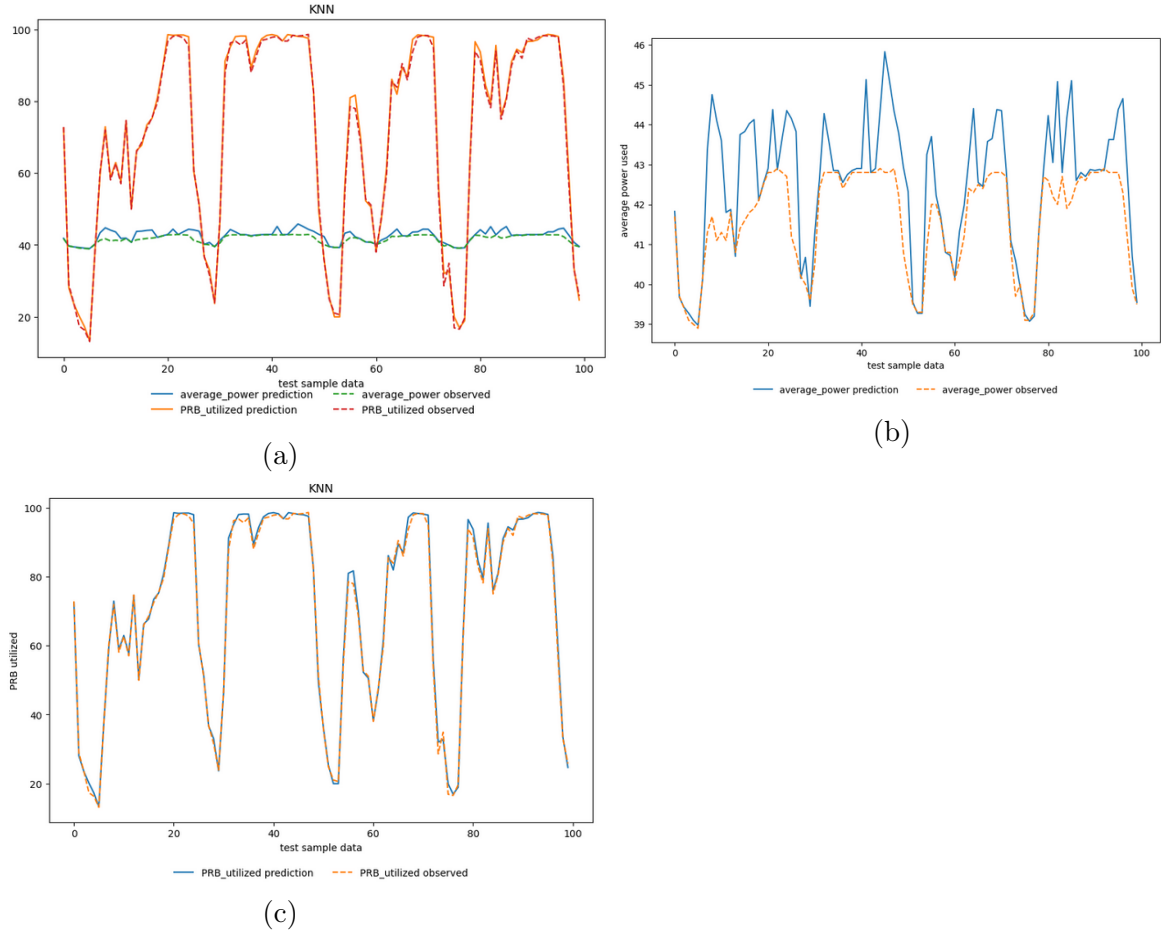


Figure 27: (a) Predicted using KNN model Vs observed value of power and PRB utilized (b) Average power used (c) PRB utilized

4.2 Performance Comparison

The performance metrics used to evaluate each model is, accuracy of resource allocation, RMSE, MAE and speed of convergence are used. Table 11 summarizes the evaluation of each model.

Table 11: Summary of performance evaluation of each model

Model	Accuracy	RMSE	MAE	Running time(sec)
LSTM	98.56	1.0520	0.08	278.96
CNN-LSTM	98.49	0.95	0.66	743.07
Random Forest	98	.294	0.08	230.22
KNN	90	14.84	8.47	1.58

From the evaluation results of each model, in Table 11 when comparing in terms of accuracy of resource allocation, RMSE, MAE and time complexity, LSTM model has higher accuracy of 98.56%, LSTM excels in handling time-series data and capturing temporal patterns in network conditions. It showed promising results with time-varying user densities and fluctuating network loads. LSTM models can automatically extract features from the sequential input where as the other models, however, could need more hand feature engineering.

In the case of RMSE and MAE, random forest has minimum error of RMSE and MAE 0.294, 0.08. This is due to several decision trees are combined to produce predictions using the ensemble learning technique of random forest. Random Forest's averaging system aids in identifying different patterns in the data and reducing overfitting. This results in increased generalization and fewer test data errors. In addition, it effectively capturing the nonlinear interactions in the data, has less hyperparameters to adjust as compared to LSTM and CNN-LSTM. These models may experience overfitting or insufficient learning if inappropriately tuned, which could lead to greater error values.

In terms of running time KNN has shortest time to fit the model that is 1.58sec this is due to its straightforward computing, lack of an explicit training phase, flexible learning methodology, and limited model complexity. Predictions are based on the similarity of new data points to previous ones, and the model stores the complete training dataset. As a result, no effort is wasted fitting or adjusting the model, which reduces computing time but increases error of prediction (Figure 28).

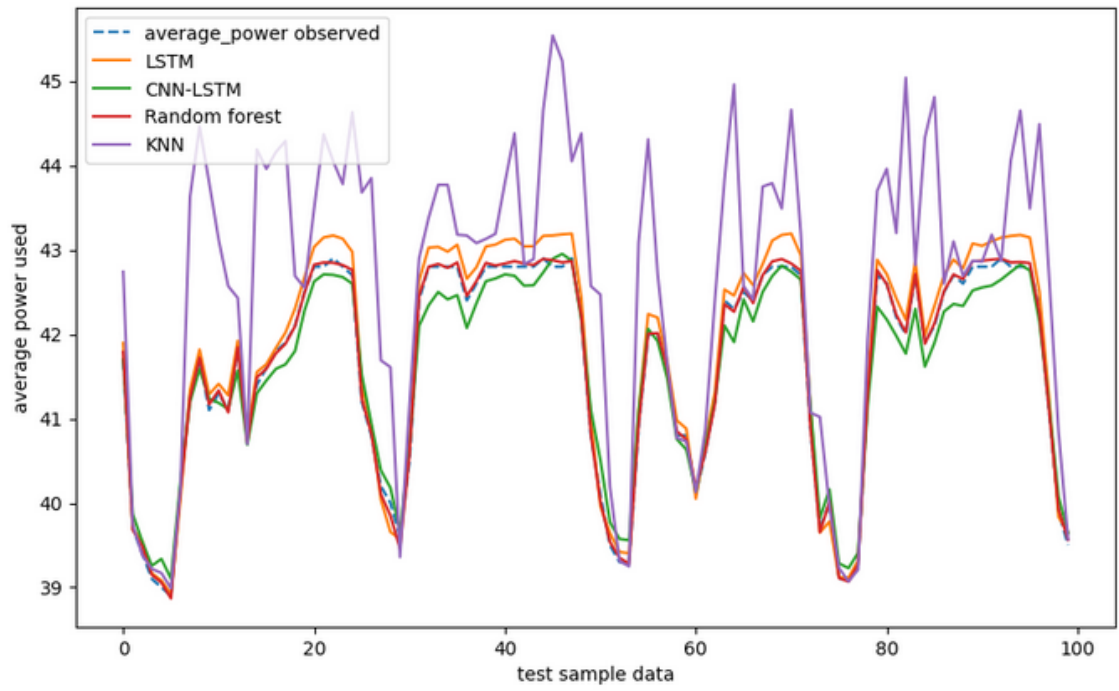


Figure 28: Predicted Vs observed value of average power used for each model

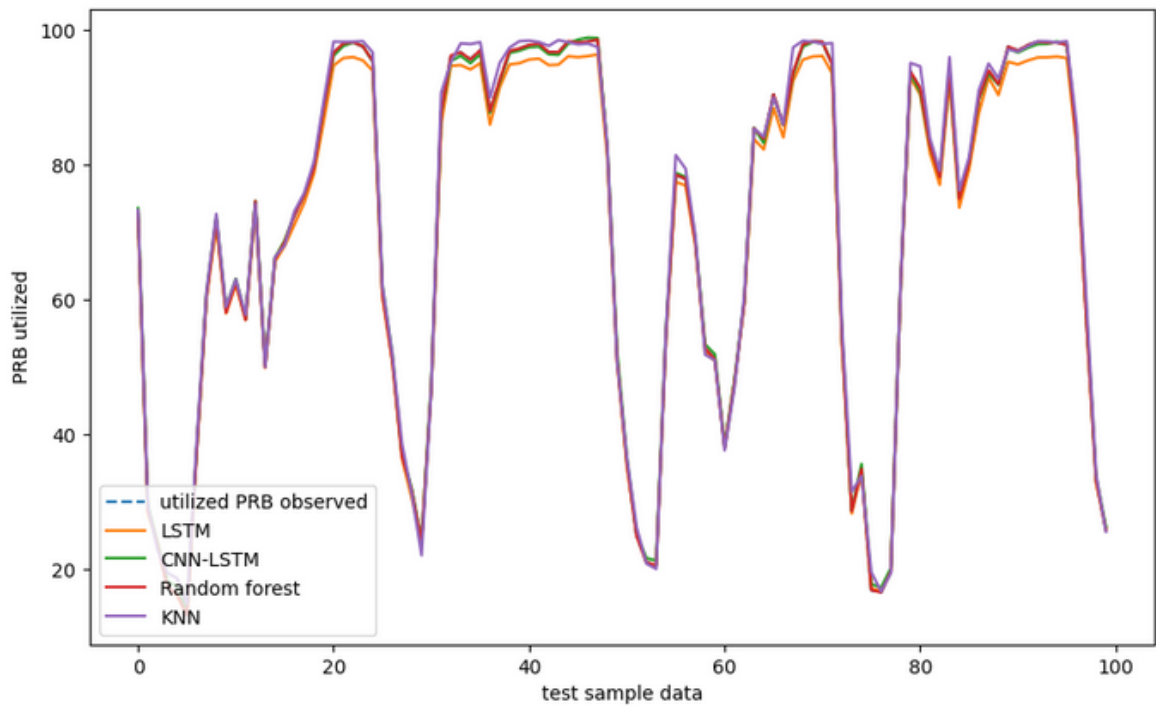


Figure 29: Predicted Vs observed value of PRB utilized for each model

5 Conclusions and Recommendations

5.1 Conclusions

This thesis demonstrates how machine learning techniques can be utilized to allocate resources for LTE-A networks in an efficient manner. Machine learning models have the ability to optimize resource consumption, minimize congestion, and improve overall network performance by dynamically adjusting resources based on conditions on the network and user needs. Furthermore, the studies' use of realistic data has made it possible for the results to be immediately applied to actual LTE-A network. The paper evaluated different machine learning and deep learning models, including long short term memory, conventional neural network-long short term memory, Random forest and k nearest neighbor to find their performance in dynamically allocating radio resource jointly, such as average power used and physical resource block with in LTE-A network. The allocation of radio resources is based on the channel quality indicator, QoS class identifier, total traffic volume of each user, average active user traffic, available physical resource block, maximum power of each cell. The main objective is to boost the total network throughput.

The proposed models takes realistic data of ethiotelecom LTE-A network to train the model offline and then after training, it allocates resources in online fashion. The performance of each model is evaluated using Accuracy score, RMSE, MAE and running time. LSTM has 98.56% accuracy score, 1.0520 RMSE, 0.76 MAE and 278.96sec running time, CNN-LSTM accuracy score 98.49%, RMSE of 0.95, MAE of 0.66 and running time 743.07sec, RF has accuracy score 98.25%, RMSE 0.294, MAE 0.08, running time 230.26sec and in the case of KNN 90% accuracy score, 0.95 RMSE, 0.56 MAE and running time 1.58 sec. From the performance evaluation of each model LSTM has highest accuracy score of 98.56%, in terms of errors RF has highest performance of lowest error that is 0.294 RMSE

and 0.08 MAE and in the case of running time KNN has shortest time of 1.58sec. Based on the results for implementation of the models and examining the realistic systems, considering the difficulties caused by latency and computational complexity.

5.2 Recommendations

Based on the research results in this thesis, a number of suggestions may be made to improve the use of machine learning in the dynamic allocation of radio resources for LTE-A networks. The thesis concentrated on individual machine learning algorithms, but it is important to look into the possible advantages of ensemble techniques. When dealing with fluctuating network conditions and user behavior, combining different models can increase performance and robustness. Conduct real-world evaluations and implementation in partnership with telecommunication providers to verify the efficiency of machine learning-based resource allocation in real LTE-A networks. This will make it easier to identify any practical difficulties and improve the models' performance in real-world settings. To ensure the viability and scalability of the suggested solutions, take into account the difficulties caused by latency and computational complexity.

References

- [1] R. Gatti *et al.*, “Improved resource allocation scheme for optimizing the performance of cell-edge users in lte-a system,” *Journal of Ambient Intelligence and Humanized Computing*, vol. 12, no. 1, pp. 811–819, 2021.
- [2] Y.-T. Mai and C.-C. Hu, “Design of dynamic resource allocation scheme for real-time traffic in the lte network,” *EURASIP Journal on Wireless Communications and Networking*, vol. 2022, no. 1, pp. 1–24, 2022.
- [3] S. Hussain, “Dynamic radio resource management in 3gpp lte,” 2009.
- [4] T. Erpek, A. Abdelhadi, and T. C. Clancy, “Application-aware resource block and power allocation for lte,” in *2016 Annual IEEE Systems Conference (SysCon)*, pp. 1–5, IEEE, 2016.
- [5] R. Khdhir, K. Mnif, A. Belghith, and L. Kamoun, “New adaptive resource allocation scheme in lte-advanced,” in *International Conference on Intelligent Systems Design and Applications*, pp. 749–759, Springer, 2016.
- [6] M. A. Abd El-Gawad, M. M. Tantawy, and M. El-Mahallawy, “Lte qos dynamic resource block allocation with power source limitation and queue stability constraints,” *International Journal of Computer Networks & Communications (IJCNC)*, vol. 5, no. 3, 2013.
- [7] H. B. Rekhissa, C. Belleudy, and P. Bessaguet, “Energy efficient resource allocation for m2m devices in lte/lte-a,” *Sensors*, vol. 19, no. 24, p. 5337, 2019.
- [8] U. Nwawelu, C. Ani, and M. Ahaneku, “Comparative analysis of the performance of resource allocation algorithms in long term evolution

- networks,” *Nigerian Journal of Technology*, vol. 36, no. 1, pp. 163–171, 2017.
- [9] F. W. Alsaade and M. Hmoud Al-Adhaileh, “Cellular traffic prediction based on an intelligent model,” *Mobile Information Systems*, vol. 2021, 2021.
- [10] A. Azari, P. Papapetrou, S. Denic, and G. Peters, “User traffic prediction for proactive resource management: Learning-powered approaches,” in *2019 IEEE Global Communications Conference (GLOBECOM)*, pp. 1–6, IEEE, 2019.
- [11] J. Lorincz, T. Garma, and G. Petrovic, “Measurements and modelling of base station power consumption under real traffic loads,” *Sensors*, vol. 12, no. 4, pp. 4281–4310, 2012.
- [12] A. Noliya and S. Kumar, “Performance analysis of resource scheduling techniques in homogeneous and heterogeneous small cell lte-a networks,” *Wireless personal communications*, vol. 112, no. 4, pp. 2393–2422, 2020.
- [13] A. Asheralieva, J. Y. Khan, and K. Mahata, “Dynamic resource allocation in a lte/wlan heterogeneous network,” in *2012 IV International Congress on Ultra Modern Telecommunications and Control Systems*, pp. 89–95, IEEE, 2012.
- [14] A. A. Neeraj Kumar, “Resource allocation in lte advanced for different services,” *Amity Journal of Computational Sciences (AJCS)*, 2015.
- [15] S. Feki, F. Zarai, and A. Belghith, “A q-learning-based scheduler technique for lte and lte-advanced network,” in *WINSYS*, pp. 27–35, 2017.
- [16] G. Piro, L. A. Grieco, G. Boggia, F. Capozzi, and P. Camarda, “Sim-

- ulating lte cellular systems: An open-source framework,” *IEEE transactions on vehicular technology*, vol. 60, no. 2, pp. 498–513, 2010.
- [17] A. Saxena and R. Sindal, “Strategy for resource allocation in lte-a,” in *2016 International Conference on Signal Processing, Communication, Power and Embedded System (SCOPEs)*, pp. 29–34, IEEE, 2016.
- [18] S. Kedir, “Comparative analysis of resource scheduling algorithms for lte,” Master’s thesis, Addis Ababa University, School of Electrical and Computer Engineering, 2021.
- [19] W. Miao, G. Min, Y. Jiang, X. Jin, and H. Wang, “Qos-aware resource allocation for lte-a systems with carrier aggregation,” in *2014 IEEE Wireless Communications and Networking Conference (WCNC)*, pp. 1403–1408, IEEE, 2014.
- [20] D. Wu, X. Zhou, Y. Xie, H. Yan, and R. Wang, “Dynamic resource allocation scheme using cooperative game for multimedia services in lte advanced system,” *EAI Endorsed Transactions on Creative Technologies*, vol. 2, no. 4, p. e2, 2015.
- [21] W. Jiang, “Cellular traffic prediction with machine learning: A survey,” *Expert Systems with Applications*, pp. 117–163, 2022.
- [22] H. D. Trinh, L. Giupponi, and P. Dini, “Mobile traffic prediction from raw data using lstm networks,” in *2018 IEEE 29th Annual International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC)*, pp. 1827–1832, IEEE, 2018.
- [23] A. Azari, F. Salehi, P. Papapetrou, and C. Cavdar, “Energy and resource efficiency by user traffic prediction and classification in cellular networks,” *IEEE Transactions on Green Communications and Networking*, vol. 6, no. 2, pp. 1082–1095, 2021.

- [24] J. Ye and Y.-J. A. Zhang, “Drag: Deep reinforcement learning based base station activation in heterogeneous networks,” *IEEE Transactions on Mobile Computing*, vol. 19, no. 9, pp. 2076–2087, 2019.
- [25] J. Wang, J. Tang, Z. Xu, Y. Wang, G. Xue, X. Zhang, and D. Yang, “Spatiotemporal modeling and prediction in cellular networks: A big data enabled deep learning approach,” in *IEEE INFOCOM 2017-IEEE conference on computer communications*, pp. 1–9, IEEE, 2017.
- [26] D. Hailemariam, *Spatiotemporal Mobile Data Traffic Prediction Using Convolutional Long Short-Term Memory: The case of Addis Ababa, Ethiopia*. PhD thesis, Addis Ababa University, 2019.
- [27] B. Tamene, “Video streaming data traffic prediction by using long short term memory (lstm) model: In the case of umts network in addis ababa,” Master’s thesis, Addis Ababa University, School of Electrical and Computer Engineering, 2020.
- [28] E. Mare, “Mobile data traffic prediction using multivariate time series data: The case of lte network in addis ababa,” Master’s thesis, Addis Ababa University, SCHOOL OF ELECTRICAL AND COMPUTER ENGINEERING, 2021.
- [29] K. I. Ahmed, H. Tabassum, and E. Hossain, “Deep learning for radio resource allocation in multi-cell networks,” *IEEE Network*, vol. 33, no. 6, pp. 188–195, 2019.
- [30] W. Lee, M. Kim, and D.-H. Cho, “Deep power control: Transmit power control scheme based on convolutional neural network,” *IEEE Communications Letters*, vol. 22, no. 6, pp. 1276–1279, 2018.
- [31] F. Meng, P. Chen, and L. Wu, “Power allocation in multi-user cellular networks with deep q learning approach,” in *ICC 2019-2019 IEEE International Conference on Communications (ICC)*, pp. 1–6, IEEE, 2019.

- [32] H. Hassan, I. Ahmed, R. Ahmad, H. Khammari, G. Bhatti, W. Ahmed, and M. M. Alam, “A machine learning approach to achieving energy efficiency in relay-assisted lte-a downlink system,” *Sensors*, vol. 19, no. 16, p. 3461, 2019.
- [33] Z. Xu, Y. Wang, J. Tang, J. Wang, and M. C. Gursoy, “A deep reinforcement learning based framework for power-efficient resource allocation in cloud rans,” in *2017 IEEE International Conference on Communications (ICC)*, pp. 1–6, IEEE, 2017.
- [34] K. R. Balmuri, S. Konda, W.-C. Lai, P. B. Divakarachari, K. M. V. Gowda, and H. Kivudujogappa Lingappa, “A long short-term memory network-based radio resource management for 5g network,” *Future Internet*, vol. 14, no. 6, p. 184, 2022.
- [35] M. C. Thayammal and M. Mary Linda, “Utility-based optimal resource allocation in lte-a networks by hybrid aco-ts with mfa scheme,” *The Computer Journal*, vol. 62, no. 6, pp. 931–942, 2019.
- [36] K. Ashfaq, G. A. Safdar, and M. Ur-Rehman, “Comparative analysis of scheduling algorithms for radio resource allocation in future communication networks,” *PeerJ Computer Science*, vol. 7, p. e546, 2021.
- [37] I.-F. Chao and C.-S. Chiou, “An enhanced proportional fair scheduling algorithm to maximize qos traffic in downlink ofdma systems,” in *2013 IEEE Wireless Communications and Networking Conference (WCNC)*, pp. 239–243, IEEE, 2013.
- [38] M. Elsayed, *Machine Learning-Enabled Radio Resource Management for Next-Generation Wireless Networks*. PhD thesis, Université d’Ottawa/University of Ottawa, 2021.
- [39] A. Zappone, M. Debbah, and Z. Altman, “Online energy-efficient power control in wireless networks by deep neural networks,” in *2018*

IEEE 19th International Workshop on Signal Processing Advances in Wireless Communications (SPAWC), pp. 1–5, IEEE, 2018.

- [40] L. Liang, H. Ye, G. Yu, and G. Y. Li, “Deep-learning-based wireless resource allocation with application to vehicular networks,” *Proceedings of the IEEE*, vol. 108, no. 2, pp. 341–356, 2019.
- [41] J. Sharma, “Introduction to supervised deep learning algorithms,” *Data Science Blogathon*, 2021.
- [42] F. Liang, W. Yu, X. Liu, D. Griffith, and N. Golmie, “Toward deep q-network-based resource allocation in industrial internet of things,” *IEEE Internet of Things Journal*, vol. 9, no. 12, pp. 9138–9150, 2021.
- [43] S. Feki and F. Zarai, “Cell performance-optimization scheduling algorithm using reinforcement learning for lte-advanced network,” in *2017 IEEE/ACS 14th International Conference on Computer Systems and Applications (AICCSA)*, pp. 1075–1081, IEEE, 2017.
- [44] B. Bojović, E. Meshkova, N. Baldo, J. Riihijärvi, and M. Petrova, “Machine learning-based dynamic frequency and bandwidth allocation in self-organized lte dense small cell deployments,” *EURASIP Journal on Wireless Communications and Networking*, vol. 2016, no. 1, pp. 1–16, 2016.
- [45] H. Na, Y. Shin, D. Lee, and J. Lee, “Lstm-based throughput prediction for lte networks,” *ICT Express*, 2021.
- [46] K. Zarzycki and M. Ławryńczuk, “Lstm and gru neural networks as models of dynamical processes used in predictive control: A comparison of models developed for two chemical reactors,” *Sensors*, vol. 21, no. 16, p. 5625, 2021.
- [47] F. Aksan, Y. Li, V. Suresh, and P. Janik, “Cnn-lstm vs. lstm-cnn to predict power flow direction: A case study of the high-voltage subnet

of northeast germany,” *Sensors*, vol. 23, no. 2, p. 901, 2023.

- [48] M. Castán-Lascorz, P. Jiménez-Herrera, A. Troncoso, and G. Asencio-Cortés, “A new hybrid method for predicting univariate and multivariate time series based on pattern forecasting,” *Information Sciences*, vol. 586, pp. 611–627, 2022.
- [49] L. Munkhdalai, T. Munkhdalai, K. H. Park, T. Amarbayasgalan, E. Batbaatar, H. W. Park, and K. H. Ryu, “An end-to-end adaptive input selection with dynamic weights for forecasting multivariate time series,” *IEEE Access*, vol. 7, pp. 99099–99114, 2019.
- [50] F. H. Al-Qahtani and S. F. Crone, “Multivariate k-nearest neighbour regression for time series data—a novel algorithm for forecasting uk electricity demand,” in *The 2013 international joint conference on neural networks (IJCNN)*, pp. 1–8, IEEE, 2013.

Performance Comparison of Machine Learning-Based Dynamic Resource Allocation Methods for LTE-A Using Realistic Data

Tiblets Kinfе Araya

Addis Ababa Institute of Technology
Addis Ababa University, Addis Ababa, Ethiopia
Email: tiblets21@gmail.com

Yihenew Wondie Marye

Addis Ababa Institute of Technology
Addis Ababa University, Addis Ababa, Ethiopia
Email: yihenew.wondie@aait.edu.et

Abstract—In the last few years, fourth generation (4G) mobile networks have become much more complex due to deployments of macro, micro, pico, and femtocells to boost network capacity and quality of Services(QoS). Long-Term Evolution Advanced (LTE-A) system is a 4G wireless communication technology that offers mobile devices high-speed data transmission, increased capacity, and improved network performance. LTE-A poses unique resource allocation problems. The following are some difficulties in allocating LTE-A resources: heterogeneous network deployment, diverse QoS requirements and dynamic traffic patterns. Traditional resource management approaches usually use static strategies or predefined rules and are insensitive to the current demand on the network. To allocate resources effectively based on the current network conditions, resource allocation must be adaptable and dynamic. Adaptive resource allocation methods that can handle the dynamic nature of LTE-A networks can be developed with the help of machine learning and artificial intelligence. This thesis evaluates the performance of machine learning-based resource allocation strategies in LTE-A networks using realistic data. Four different machine learning-based resource allocation strategies are compared: Long Short Term Memory (LSTM), Conventional Neural Network and Long Short Term Memory (CNN-LSTM), Random Forest (RF) and K Nearest Neighbor(KNN) algorithm. Performance measures include the accuracy of resource allocation, Root Mean Square Error (RMSE), Mean Absolute Error(MAE) and speed of convergence(running time). The results show that LSTM is the model with

higher accuracy score of resource allocation that is 98.56% and in terms RMSE and MAE random forest has lowest value of 0.294 and 0.08. In the case of running time KNN takes shortest time of 1.58 second.

Keywords: LTE, LTE-Advanced, Machine Learning, Radio Resource Allocation, Channel quality indicator

I. INTRODUCTION

The scale and complexity of fourth generation (4G) mobile networks have grown significantly during the past few years. By adding more cells of various types to the present deployments, which include macro-, micro-, pico-, and femtocells, operators are constantly looking to increase network capacity and Quality of Service (QoS). The Operating and Capital Expenditures (OPEX/CAPEX) of the mobile network providers significantly increase as a result of this increase in complexity. To permanently decrease these costs, operators are looking for network solutions that would automatically configure, manage, and optimize networks while reducing the need for human interaction based on the traffic demand. Radio Resource Management (RRM) making use of the available adaptation techniques to serve users depending on their QoS parameters to ensure the efficient use of available radio resources [1], [2], [3]. To enhance the capacity of the

network and absorb the mobile data traffic increment, ethiotelecom has been conducting a lot of expansion projects, base station densifications on a greenfield or rooftop, load balancing and optimizations.

In order to inform the User Equipment (UE) about its allotted time/frequency resources and transmission formats to be employed by UE, eNodeB performs scheduling on the available radio resources. UE capability, QoS, and measurement reports from the UE are used to inform Radio Resource (RR) scheduling [4]. Each eNodeB must distribute transmission powers throughout each resource block while taking into account the impact on UE throughput and interference on other UEs connected to other base stations. A resource allocation framework in the time and frequency domains is defined by the Long-Term Evolution Advanced (LTE-A) standard. The frequency dimension is broken into sub-carriers that are separated by 15 KHz. The time dimension is split into ten 1 ms sub-frames and then divided into 10 radio frames of 1ms each. It divides each sub-frame into two 0.5 ms slots. The lowest unit of resource, known as a resource element, is one sub carrier lasting one Orthogonal Frequency Division Multiplexing (OFDM) signal. The length of a resource block is one millisecond, with 12 constant sub-carriers. A centralized resource block scheduling method chooses which UE will be assigned to each resource block. Only one user may be given access to a resource block [3], [5].

A resource block (RB) is the smallest unit of resource allocation for any one user. Twelve sub-carriers and six or seven OFDM symbols are used by one resource block. Consequently, a resource block may contain a total of 72 or 84 resource elements. Radio resource elements are separated into time and frequency resource elements in LTE. As in Figure 1 resource elements are separated into a certain number of sub-carriers on the frequency axis and a certain number of OFDM symbols on the time axis. The smallest resource unit, known as a resource element (RE), uses both an OFDM

symbol and a single sub-carrier. Flexible bandwidths supported by the technology include 1.4MHz, 3MHz, 5MHz, 10MHz, 15MHz, and 20MHz. Each bandwidth can have a variable number of resource blocks, namely 6, 15, 25, 50, 75, and 100 RBs. The five physical channels Physical Downlink Control Channel (PDCCH), Physical Downlink Shared Channel (PDSCH), Physical Broadcast Channel (PBCH), Physical Control Format Indicator Channel (PCFICH), and Physical Hybrid ARQ Indicator Channel (PHICH) suggested by the LTE standard are for use in the downlink. The shared data traffic channel is called PDSCH, and the control channel is called PDCCH. In a 1ms transmit time interval (TTI) sub-frame, a number of resource blocks are allotted to these channels. 10 sub-frames, each ranging 1ms, make up a single 10ms LTE radio frame as in Figure 1. Two resource blocks (RBs) of 0.5ms each are included in each sub-frame [6], [7], [8].

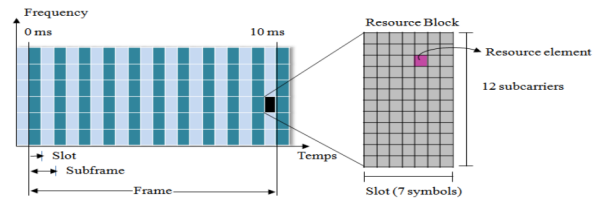


Fig. 1: Frame of long term evolution(LTE) [9]

The user measurement entity reports the Channel Quality Indicator (CQI) to the base station (BS) for every TTI to give time and frequency channel quality information for improved spectral efficiency and resource allocation. Users can alert the BS about the channel quality for downlink RBs by using the Physical Uplink Control Channel (PUCCH). By employing the PDCCH as shown in Figure 2, BS distributes downlink RB allocations and Modulation and Coding scheme(MCS) assignments to all users [10].

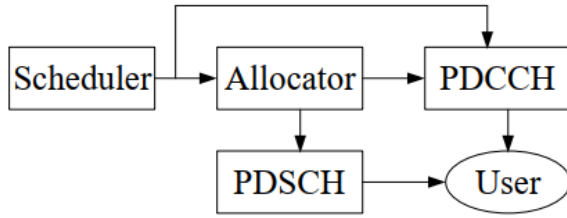


Fig. 2: Physical model of a downlink resource Allocation [11]

The International Telecommunication Union's (ITU) guideline is not met by the LTE standard. Since, it is a 3.9G rather than a full 4G technology. As a result, LTE was further developed in accordance with ITU rules to become LTE-A, which is a true 4G technology that offers up to 500 MBPS in the uplink and up to 1 GBPS in the downlink. It enables to have such increased data rates and low latency using higher order Multiple Input Multiple Output(MIMO), relay technology, Cooperative Multipoint(CoMP), carrier aggregation, and other technologies [6], [1].

To obtain high performance, it is vital to utilize limited network resources as effectively as possible in the most advantageous way possible while also maintaining the best network utilization, to model and predict exactly cellular networks will load [12]. Besides, to assign the available network resources to the devices requesting service while controlling the level of network resources, such as the quantity of radio resources, the Buffer Status Report (BSR) of users is typically used in traditional approaches to network resource management to provision and allocate resources as shown in Figure 3. However, in this case, take a proactive approach and look for chances to implement resource provisioning and allocation. Then, based on the present and past states of user traffic behavior, channel condition, required quality of service, tries to decide how to serve a group of users who coexist. It must do a through analysis of each user's traffic behavior [13]. This thesis compares the performance of different distinct machine learning based dynamic

resource allocating strategies, to jointly allocate downlink Physical Resource block(PRB) utilization and average power used by each user using realistic data collected from Addis Ababa LTE-A network.

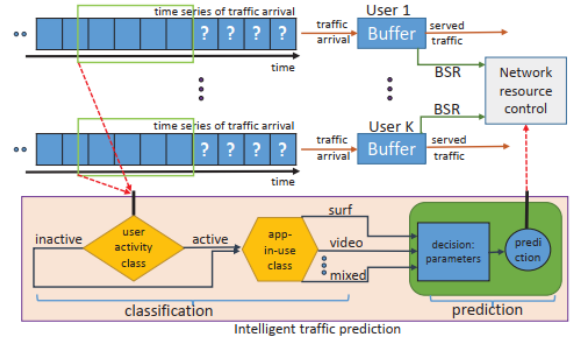


Fig. 3: Intelligent module for better decision-making in managing network resources to serve users [13]

To utilize network resources effectively, considering traffic volume, active user traffic, available resources is necessary. Accurate network traffic prediction at access points provides effective resource allocation to guarantee acceptable service quality [14].

II. RELATED WORKS

Researchers are addressing the challenge of resource allocation by utilizing innovative traffic forecasting techniques, by developing intelligent systems that considers channel and user conditions to optimize resource allocation following the traffic data flow during network operation. Some of the related literature's are presented as follows.

A. Dynamic Resource Allocation Strategies

A fair and effective resource allocation system is crucial for enhancing network performance [15]. The authors in [16] present that using RRM, the LTE Advanced system distributes radio resources for various services. Different service kinds, including real-time (RT), non-real-time (NRT), and control signaling, are dynamically active in networks, necessitating dynamic resource allocation for the transmission of pertinent

data packets. When findings from the suggested technique are compared to those from existing methods, it becomes clear that the current solution yields superior outcomes. There are descriptions of several service types and their schedule. It demonstrates service prioritizing optimization strategies in which the weights of service types are optimized. For resource allocation, it does not, however, take into account the volume of traffic, user count, channel status, etc.

Since it is responsible for ensuring user QoS in LTE-A networks, multi-user scheduling is one of the key components. Typically, scheduling algorithms strive for maximum throughput while retaining a certain level of fairness. The smallest resource unit in LTE-A is the RB, which has a frequency domain size of 180 KHz and a time domain division into two slots with a duration of 0.5 ms each. The eNodeB makes a scheduling decision every 1 ms, which is the size of the Transmission Time Interval (TTI). Physical resources are distributed to users in the frequency or time grid throughout time, and multi-user diversity is controlled in both of these time and frequency domains. Signaling limitations prevent subcarriers from being allocated separately, so they should be combined on an RB basis [17], [18].

The efficiency of Universal Terrestrial Radio Access Network (EUTRAN) when using various resource allocation algorithms, including round robin, proportional fair, best channel quality indicator (CQI), maximum transmission point (TP), and resource fair, was evaluated using MATLAB. Fair resource distribution among UE is provided via round robin, proportional fair, resource fair, and best CQI. High cell throughput with selective resource allocation is provided by MAX TP. The best CQI and resource fair strategy offers a network with degraded resource occupancy high throughput and fair resource distribution. Although the UE throughput is relatively smaller with the round robin technique, resource allocation and occupancy are quite fair. The authors recommend adapting mixed

strategies in addition to channel imperfection in which the performance of resource allocation depends [7], [19]. The paper [20] compares how well two fundamental techniques proportional fairness (PF) and round robin perform. According to this, the PF approach offers great fairness but just average throughput. However, from a scheduling standpoint, it does not distinguish between different service types; instead, a technique for raising QoS called proportional fairness is offered. Similar to this, [21] also describes adaptive QoS for PF methods. These methods, however, cannot have the capability to determine the current service type prior to resource scheduling. Dynamic resource allocation is necessary for dynamic service behavior. This research focuses on performance comparison of machine learning-based dynamic resource allocation methods for LTE-A using realistic data.

B. Machine Learning Based Dynamic Resource Allocation

In [22], the author presented supervised deep learning model for allocation of sub-band and power for multi-cell wireless network to maximize the total network throughput. The proposed model uses CQI values of each user of the multicell network as input and allocates the required radio resources. The proposed model needs only one-time offline training and after training, it can allocate radio resources in an online fashion. The optimality of this work is no guarantee because it uses Genetic algorithm to generate training and testing data. In [23], the authors propose a method for distributed power regulation based on a data-driven Convolutional Neural Network (CNN) that uses spatial data in channel gain to estimate the transmit power. Maximizing throughput through power allocation is the major goal. There is no assurance that the suggested approach would produce the centrally optimum result. The authors in [24] propose two-step training framework that is a Deep Q Network (DQN) is trained with Deep Q Learning (DQL) algorithm with the offline learning environment and DQN is fine-tuned in

on-line training with real time data. They considered Interfering Multiple-Access Channel (IMAC) cellular network with the problem of power allocation. This paper focuses on performance comparison of machine learning-based dynamic resource allocation using the realistic data of Addis Ababa LTE-A network. It considers the actual active user traffic, traffic volume, channel quality indicator and service types rather than predicting traffic first. Besides, traditional resource allocation strategies follow predefined rules so that cannot adapt the dynamic nature of the network.

III. SYSTEM MODEL AND DATA PREPARATION

A. System Model

The time and frequency domain resource allocation framework is defined by the LTE standard. The sub-carriers that make up the frequency dimension are separated by 15 kHz. The time dimension is first split into ten 1 ms subframes, followed by ten 10 ms radio frames. Two 0.5 ms slots separate each subframe. One subcarrier for one OFDM symbol makes up a resource element, which is the lowest unit of resource. For the duration of one slot, a resource block has 12 constant subcarriers [3].

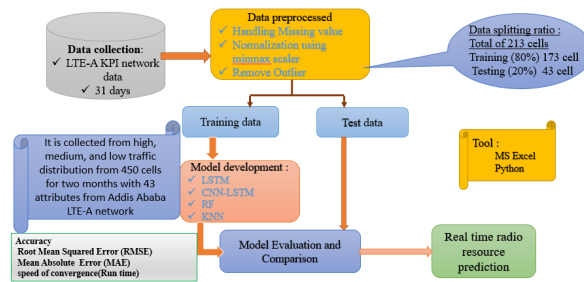


Fig. 4: Overall system model

B. Data Preparation and Feature Extraction

1) *Data Collection:* Hourly data is collected from the Key Performance Indicator(KPI) that indicates the performance of real deployment of the LTE-A network in Addis Ababa from 450 cells for two months. The collected data is taken from high, medium and

low traffic areas of Addis ababa. This data includes channel state information(CQI), quality of service class indicator(QCI), MCs, traffic volume, average active user traffic, average user traffic, available PRB and maximum power per cell with 16 features 14 input and 2 output features. Finally, by removing the sites with missing hours one month data of 213 cell is used.

C. Data Preprocessing and Feature Extraction

Data preprocessing perform steps like data cleaning, filtering, and normalizing the dataset. Data cleaning is the process of clean the dataset by removing any missing or irrelevant data points. Normalize the data to ensure uniformity and prevent biases toward specific attributes. In case of feature extraction, extract relevant features from the dataset that can serve as input to the machine learning models and affect the resource allocation decisions. This study extracts 16 features with 14 input and 2 target features from 43 total collected features. On the basis of prior data, machine learning models are employed to forecast UEs' resource needs. The LTE-A network's realistic data is used to train the machine learning algorithms.

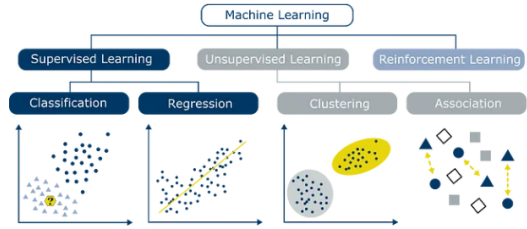


Fig. 5: Types of machine learning algorithm [25]

Regression and Classification are subcategories that simply differ in terms of output value. While classification categorizes the dataset, regression produces continuous results. This thesis uses regression supervised machine learning algorithm for jointly allocating resource. The algorithms used in this thesis are Long short term memory(LSTM), the hybrid model of Conventional Neural Network and Long Short Term

Memory(CNN-LSTM), Random Forest(RF) and K-Nearest Neighbor(KNN) algorithm.

D. Performance Metrics

The trained model's performance is evaluated using appropriate regression evaluation metrics such as accuracy, RMSE, MAE and speed of convergence (running time). The developed Machine learning models is a crucial stage in determining the effectiveness and correctness of the model based on the metric scores.

IV. RESULTS AND DISCUSSIONS

The goal of this study is to better understand how machine learning-based resource allocation algorithms perform in LTE-A networks. The findings of this study will bring insight into which algorithms perform best in terms of accuracy, RMSE, MAE and running time when it comes to allocating resources in LTE-A networks. Before training the model, the values of the hyperparameters should be initialized, and the grid search technique is used to find the model's optimal configuration.

The performance metrics used to evaluate each model is, accuracy of resource allocation, RMSE, MAE and speed of convergence are used. Table I summarizes the evaluation of each model.

TABLE I: Summary of performance evaluation of each model

Model	Accuracy	RMSE	MAE	Running time(sec)
LSTM	98.56	1.0520	0.08	278.96
CNN-LSTM	98.49	0.95	0.66	743.07
Random Forest	98	.294	0.08	230.22
KNN	90	14.84	8.47	1.58

From the evaluation results of each model, in Table I when comparing in terms of accuracy of resource allocation, RMSE, MAE and time complexity, LSTM model has higher accuracy of 98.56%, LSTM excels in handling time-series data and capturing temporal

patterns in network conditions. It showed promising results with time-varying user densities and fluctuating network loads. LSTM models can automatically extract features from the sequential input where as the other models, however, could need more hand feature engineering.

In the case of RMSE and MAE, random forest has minimum error of RMSE and MAE 0.294, 0.08. This is due to several decision trees are combined to produce predictions using the ensemble learning technique of random forest. Random Forest's averaging system aids in identifying different patterns in the data and reducing overfitting. This results in increased generalization and fewer test data errors. In addition, it effectively capturing the nonlinear interactions in the data, has less hyperparameters to adjust as compared to LSTM and CNN-LSTM. These models may experience overfitting or insufficient learning if inappropriately tuned, which could lead to greater error values.

In terms of running time KNN has shortest time to fit the model that is 1.58sec this is due to its straightforward computing, lack of an explicit training phase, flexible learning methodology, and limited model complexity. Predictions are based on the similarity of new data points to previous ones, and the model stores the complete training dataset. As a result, no effort is wasted fitting or adjusting the model, which reduces computing time but increases error of prediction (Figure 6).

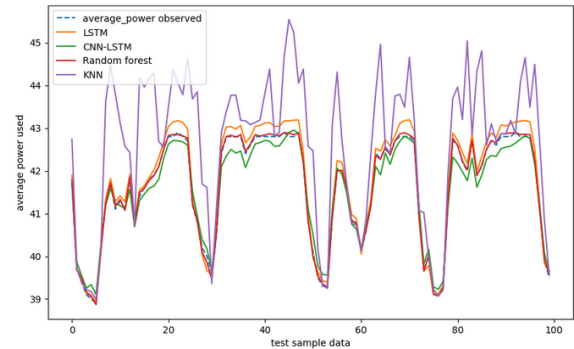


Fig. 6: Predicted Vs observed value of average power used for each model

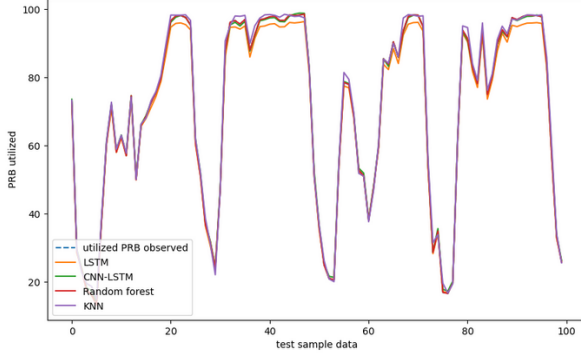


Fig. 7: Predicted Vs observed value of PRB utilized for each model

V. CONCLUSIONS

This thesis demonstrates how machine learning techniques can be utilized to allocate resources for LTE-A networks in an efficient manner. Machine learning models have the ability to optimize resource consumption, minimize congestion, and improve overall network performance by dynamically adjusting resources based on conditions on the network and user needs. Furthermore, the studies' use of realistic data has made it possible for the results to be immediately applied to actual LTE-A network. The paper evaluated different machine learning and deep learning models, including long short term memory, conventional neural network-long short term memory, Random forest and k nearest neighbor to find their performance in dynamically allocating radio resource jointly, such as average power used and physical resource block with in LTE-A network. The allocation of radio resources is based on the channel quality indicator, QoS class identifier, total traffic volume of each user, average active user traffic, available physical resource block, maximum power of each cell. The main objective is to boost the total network throughput.

The proposed models takes realistic data of ethiotelecom LTE-A network to train the model offline and then after training, it allocates resources in online fashion. The performance of each model is evaluated

using Accuracy score, RMSE, MAE and running time. LSTM has 98.56% accuracy score, 1.0520 RMSE, 0.76 MAE and 278.96sec running time, CNN-LSTM accuracy score 98.49%, RMSE of 0.95, MAE of 0.66 and running time 743.07sec, RF has accuracy score 98.25%, RMSE 0.294, MAE 0.08, running time 230.26sec and in the case of KNN 90% accuracy score, 0.95 RMSE, 0.56 MAE and running time 1.58 sec. From the performance evaluation of each model LSTM has highest accuracy score of 98.56%, in terms of errors RF has highest performance of lowest error that is 0.294 RMSE and 0.08 MAE and in the case of running time KNN has shortest time of 1.58sec. Based on the results for implementation of the models and examining the realistic systems, considering the difficulties caused by latency and computational complexity.

REFERENCES

- [1] R. Gatti *et al.*, "Improved resource allocation scheme for optimizing the performance of cell-edge users in LTE-A SYSTEM," *Journal of Ambient Intelligence and Humanized Computing* AAA, VOL. 12, NO. 1, PP. 811–819, 2021.
- [2] Y.-T. MAI AND C.-C. HU, "DESIGN OF DYNAMIC RESOURCE ALLOCATION SCHEME FOR REAL-TIME TRAFFIC IN THE LTE NETWORK," *EURASIP Journal on Wireless Communications and Networking*, VOL. 2022, NO. 1, PP. 1–24, 2022.
- [3] T. ERPEK, A. ABDELHADI, AND T. C. CLANCY, "APPLICATION-AWARE RESOURCE BLOCK AND POWER ALLOCATION FOR LTE," in *2016 Annual IEEE Systems Conference (SysCon)*, PP. 1–5, IEEE, 2016.
- [4] S. HUSSAIN, "DYNAMIC RADIO RESOURCE MANAGEMENT IN 3GPP LTE," 2009.
- [5] I.-F. CHAO AND C.-S. CHIOU, "AN ENHANCED PROPORTIONAL FAIR SCHEDULING ALGORITHM TO MAXIMIZE QOS TRAFFIC IN DOWNLINK OFDMA SYSTEMS," in *2013 IEEE Wireless Communications and Networking Conference (WCNC)*, PP. 239–243, IEEE, 2013.
- [6] R. KHDHIR, K. MNIF, A. BELGHITH, AND L. KAMOUN, "NEW ADAPTIVE RESOURCE ALLOCATION SCHEME IN LTE-ADVANCED," in *International Conference on Intelligent Systems Design and Applications*, PP. 749–759, SPRINGER, 2016.

- [7] A. SAXENA AND R. SINDAL, "STRATEGY FOR RESOURCE ALLOCATION IN LTE-A," IN *2016 International Conference on Signal Processing, Communication, Power and Embedded System (SCOPE)*, PP. 29–34, IEEE, 2016.
- [8] M. C. THAYAMMAL AND M. MARY LINDA, "UTILITY-BASED OPTIMAL RESOURCE ALLOCATION IN LTE-A NETWORKS BY HYBRID ACO-TS WITH MFA SCHEME," *The Computer Journal*, VOL. 62, NO. 6, PP. 931–942, 2019.
- [9] H. B. REKHISSE, C. BELLEUDY, AND P. BESSAGUET, "ENERGY EFFICIENT RESOURCE ALLOCATION FOR M2M DEVICES IN LTE/LTE-A," *Sensors*, VOL. 19, NO. 24, P. 5337, 2019.
- [10] M. A. ABD EL-GAWAD, M. M. TANTAWY, AND M. ELMAHALLAWY, "LTE QOS DYNAMIC RESOURCE BLOCK ALLOCATION WITH POWER SOURCE LIMITATION AND QUEUE STABILITY CONSTRAINTS," *International Journal of Computer Networks & Communications (IJNC)*, VOL. 5, NO. 3, 2013.
- [11] U. NWAWELU, C. ANI, AND M. AHANEKU, "COMPARATIVE ANALYSIS OF THE PERFORMANCE OF RESOURCE ALLOCATION ALGORITHMS IN LONG TERM EVOLUTION NETWORKS," *Nigerian Journal of Technology*, VOL. 36, NO. 1, PP. 163–171, 2017.
- [12] F. W. ALSAADE AND M. HMOUD AL-ADHAILEH, "CELLULAR TRAFFIC PREDICTION BASED ON AN INTELLIGENT MODEL," *Mobile Information Systems*, VOL. 2021, 2021.
- [13] A. AZARI, P. PAPAPETROU, S. DENIC, AND G. PETERS, "USER TRAFFIC PREDICTION FOR PROACTIVE RESOURCE MANAGEMENT: LEARNING-POWERED APPROACHES," IN *2019 IEEE Global Communications Conference (GLOBECOM)*, PP. 1–6, IEEE, 2019.
- [14] J. LORINCZ, T. GARMA, AND G. PETROVIC, "MEASUREMENTS AND MODELLING OF BASE STATION POWER CONSUMPTION UNDER REAL TRAFFIC LOADS," *Sensors*, VOL. 12, NO. 4, PP. 4281–4310, 2012.
- [15] A. NOLIYA AND S. KUMAR, "PERFORMANCE ANALYSIS OF RESOURCE SCHEDULING TECHNIQUES IN HOMOGENEOUS AND HETEROGENEOUS SMALL CELL LTE-A NETWORKS," *Wireless personal communications*, VOL. 112, NO. 4, PP. 2393–2422, 2020.
- [16] A. A. NEERAJ KUMAR, "RESOURCE ALLOCATION IN LTE ADVANCED FOR DIFFERENT SERVICES," *Amity Journal of Computational Sciences (AJCS)*, 2015.
- [17] S. FEKI, F. ZARAI, AND A. BELGHITH, "AA Q-LEARNING-BASED SCHEDULER TECHNIQUE FOR LTE AND LTE-ADVANCED NETWORKS," IN *WINSYS*, PP. 27–35, 2017.
- [18] G. PIRO, L. A. GRIECO, G. BOGGIA, F. CAPOZZI, AND P. CAMARDA, "SIMULATING LTE CELLULAR SYSTEMS: AN OPEN-SOURCE FRAMEWORK," *IEEE transactions on vehicular technology*, VOL. 60, NO. 2, PP. 498–513, 2010.
- [19] S. KEDIR, "COMPARATIVE ANALYSIS OF RESOURCE SCHEDULING ALGORITHMS FOR LTE," MASTER'S THESIS, ADDIS ABABA UNIVERSITY, SCHOOL OF ELECTRICAL AND COMPUTER ENGINEERING, 2021.
- [20] W. MIAO, G. MIN, Y. JIANG, X. JIN, AND H. WANG, "QOS-AWARE RESOURCE ALLOCATION FOR LTE-A SYSTEMS WITH CARRIER AGGREGATION," IN *2014 IEEE Wireless Communications and Networking Conference (WCNC)*, PP. 1403–1408, IEEE, 2014.
- [21] D. WU, X. ZHOU, Y. XIE, H. YAN, AND R. WANG, "DYNAMIC RESOURCE ALLOCATION SCHEME USING COOPERATIVE GAME FOR MULTIMEDIA SERVICES IN LTE ADVANCED SYSTEM," *EAI Endorsed Transactions on Creative Technologies*, VOL. 2, NO. 4, P. E2, 2015.
- [22] K. I. AHMED, H. TABASSUM, AND E. HOSSAIN, "DEEP LEARNING FOR RADIO RESOURCE ALLOCATION IN MULTI-CELL NETWORKS," *IEEE Network*, VOL. 33, NO. 6, PP. 188–195, 2019.
- [23] W. LEE, M. KIM, AND D.-H. CHO, "DEEP POWER CONTROL: TRANSMIT POWER CONTROL SCHEME BASED ON CONVOLUTIONAL NEURAL NETWORK," *IEEE Communications Letters*, VOL. 22, NO. 6, PP. 1276–1279, 2018.
- [24] F. MENG, P. CHEN, AND L. WU, "POWER ALLOCATION IN MULTI-USER CELLULAR NETWORKS WITH DEEP Q LEARNING APPROACH," IN *ICC 2019-2019 IEEE International Conference on Communications (ICC)*, PP. 1–6, IEEE, 2019.
- [25] M. ELSAYED, *Machine Learning-Enabled Radio Resource Management for Next-Generation Wireless Networks*. PHD THESIS, UNIVERSITÉ D'OTTAWA/UNIVERSITY OF OTTAWA, 2021.