

Machine-to-Machine Communications over an IEEE 802.16-based WiMAX Network in the Smart Grid



THE UNIVERSITY OF
NEWCASTLE
AUSTRALIA

Reduan H. Khan, B.Sc. (Hons)

School of Electrical Engineering and Computer Science

The University of Newcastle, Australia

A thesis submitted for the degree of

Doctor of Philosophy

July 2014

STATEMENT OF ORIGINALITY

The thesis contains no material which has been accepted for the award of any other degree or diploma in any university or other tertiary institution and, to the best of my knowledge and belief, contains no material previously published or written by another person, except where due reference has been made in the text. I give consent to the final version of my thesis being made available worldwide when deposited in the University 's Digital Repository, subject to the provisions of the Copyright Act 1968.

STATEMENT OF AUTHORSHIP

I hereby certify that the work embodied in this thesis is the result of original research and contains several published papers/scholarly work which I am a joint author. I have included as part of the thesis a written statement, endorsed by my supervisor, attesting to my contribution to the joint publication/s/scholarly work.

Reduan H. Khan

CONTRIBUTION TO JOINT PUBLICATION/S/SCHOLARLY WORK

R. Khan and J. Khan, A comprehensive review of the application characteristics and traffic requirements of a Smart Grid communications network, *Computer Networks*, 57(3):825-845, Feb 2013.

For this publication, R. Khan conducted the study and prepared the manuscript.

R. Khan, K. Mahata, and J. Khan, A novel scheduling service for distribution pilot protection over a WiMAX network, *Recent Patents on Telecommunication*, 2(1):26-40, Oct 2013.

For this publication, R. Khan designed the scheduling framework, developed the discrete event simulation model, and prepared the manuscript.

R. Khan and K. Mahata, Random access performance of a WiMAX network for M2M communications in the Smart Grid, in *Smart Grid Communications (SmartGridComm), 2014 IEEE International Conference on*, Nov 2014 (in proceedings).

For this publication, R. Khan contributed in developing the analytical model, developed the discrete event simulation model, and prepared the manuscript.

M. Hyder, **R. Khan**, and K. Mahata, An enhanced random access mechanism for Smart Grid M2M communications in WiMAX networks, in *Smart Grid Communications (SmartGridComm), 2014 IEEE International Conference on*, Nov 2014 (in proceedings).

For this publication, R. Khan contributed in developing the conceptual framework and the analytical model, ran statistical analyses, and prepared the manuscript.

R. Khan, K. Mahata, and J. Brown, An adaptive RRM scheme for Smart Grid M2M applications over a WiMAX network, in *Communication Systems, Networks, and Digital Signal Processing (CSNDSP), IEEE/IET International Symposium on*, July 2014 (in proceedings).

For this publication, R. Khan designed the radio resource management framework, developed the discrete event simulation model, and prepared the manuscript.

R. Khan and J. Khan, A heterogeneous WiMAX-WLAN network for AMI communications in the Smart Grid, in *Smart Grid Communications (SmartGridComm), 2012 IEEE Third International Conference on*, Nov 2012, pp. 710-715.

For this publication, R. Khan designed the network architecture, developed the discrete event simulation model, and prepared the manuscript.

R. Khan and J. Khan, Wide area PMU communication over a WiMAX network in the Smart Grid, in *Smart Grid Communications (SmartGridComm), 2012 IEEE Third International Conference on*, Nov 2012, pp. 187-192.

For this publication, R. Khan designed the traffic models, developed the discrete event simulation model, and prepared the manuscript.

R. Khan, J. Brown, and J. Khan, Pilot protection schemes over a multi-service WiMAX network in the Smart Grid, in *Communications Workshops (ICC), 2013 IEEE International Conference on*, June 2013, pp. 994-999.

For this publication, R. Khan designed the concept, developed the discrete event simulation model, and prepared the manuscript.

R. Khan, T. Ustun, and J. Khan, Differential protection of microgrids over a WiMAX network, in *Smart Grid Communications (SmartGridComm), 2013 IEEE International Conference on*, Oct 2013, pp. 732-737.

For this publication, R. Khan designed the concept, developed the discrete event simulation model, and prepared the manuscript.

M. Hyder, **R. Khan**, K. Mahata, and J. Khan, A predictive protection scheme based on adaptive synchrophasor communications, in *Smart Grid Communications (SmartGridComm), 2013 IEEE International Conference on*, Oct 2013, pp. 762-767.

For this publication, R. Khan designed the communications scheme, developed the discrete event simulation model, and prepared parts of the manuscript.

R. Khan, S. Stüdli, and J. Khan, A network controlled load management scheme for domestic charging of Electric Vehicles, in *Power Engineering Conference (AUPEC), 2013 Australasian Universities*, Sept 2013, pp. 1-6.

For this publication, R. Khan designed the load management framework, developed the discrete event simulation model, and prepared the manuscript.

T. Ustun and **R. Khan**, Multi-terminal hybrid protection of microgrids over wireless communications networks, *IEEE Transactions on Smart Grid*, May 2014 (submitted for publication).

For this publication, R. Khan designed the communications scheme, developed the discrete event simulation model, and prepared parts of the manuscript.

Reduan H. Khan

(Candidate)

Kaushik Mahata

(Principal Supervisor)

To my parents,
Md. Hasanuzzaman Khan and Abe Kawsar Begum

Acknowledgements

I would first like to express my deepest appreciation and heartfelt gratitude to my principal supervisor A/Prof. Kaushik Mahata. Without his invaluable guidance and dedicated supervision, this dissertation would not have been possible. I will always remain grateful for his continuous encouragement and support throughout this journey. A special gratitude and thanks to A/Prof. Jamil Khan for giving me the opportunity to work in the ‘Smart Grid Communications’ project and providing me with valuable ideas and advices which inspired me to write a number of chapters in this thesis.

I would like to thank my co-supervisor Dr. Mashud Hyder for teaching me about the intricate concepts of multicarrier wireless communications with ease and clarity. Without his help and support, the work in Chapter 5 of this dissertation would not have been possible. I am grateful to my mentor Dr. Jason Brown for his gracious help and advice throughout my doctoral studies, including his thoughtful and sincere comments while preparing this thesis.

I would like to thank my good friend Dr. Taha Selim Ustun of Carnegie Mellon University, U.S. for many enlightening discussions which inspired my works on the Smart Grid protection applications. I am especially thankful to Sonja Stüdl for her assiduous efforts in helping me to understand the interesting concepts of Electric Vehicles. Also, thanks to Prof. Victor Sreeram of the University of Western Australia, and Prof. Rick Middleton and Prof. Minyue Fu of the University of Newcastle for their advices and suggestions on various topics of Smart Grid communications.

A special thanks to my friend Dr. Ashraf Islam of Intel Corp. for introducing me to the exciting field of Smart Grid, which suited my interest and expertise. I am also thankful to my friend M. I. Sakib Chowdhury of the Pennsylvania State University, U.S. for his rigorous reviews which helped me to write this thesis with more clarity.

I wish to acknowledge the Australian Research Council (ARC) for providing me the APA Scholarships to fund this research, and thus, helping me in fulfilling my wish of pursuing a doctoral degree. Also, I acknowledge the contribution of the utility operator Ausgrid for providing me with valuable information from its WiMAX trial and supporting me with a top-up scholarship in the final year of my candidature.

I am forever indebted to my parents for their unconditional love, support and encouragements that made me who I am. I am grateful to my wife Tanzim Aditi for her relentless support during this arduous journey. Last but not least,

a very special thanks to my daughter Raifa whose smile always washes away
all my fatigues.

Abstract

A robust communication infrastructure is the touchstone of a Smart Grid that differentiates it from the conventional electrical grid by transforming it into an intelligent and adaptive energy delivery network. To cope with the rising penetration of renewable energy sources and expected widespread adoption of electric vehicles, the future Smart Grid needs to implement efficient monitoring and control technologies to improve its operational efficiency. However, legacy communication infrastructures in the existing grid are quite insufficient, if not incapable of meeting the diverse communication requirements of the Smart Grid.

Being one of the two available 4G technologies in the world, the IEEE 802.16-based WiMAX (Worldwide Interoperability for Microwave Access) networks can significantly extend the reach of a Smart Grid by allowing fast and reliable communications over a wide coverage area. However, the unique characteristics of machine-to-machine (M2M) communications in the Smart Grid pose a number of interesting challenges to the conventional telecommunications networks, including WiMAX. Hence, considerable uncertainties exist about the applicability of WiMAX in a Smart Grid environment.

The aim of this thesis is to offer an in-depth study of M2M communications over a WiMAX network in the Smart Grid. To fulfil this aim, it first conducts a detailed, technology agnostic review on the application characteristics and traffic requirements of several major Smart Grid applications and highlights their key communication challenges. Based on this review, it develops a range of traffic models for some key Smart Grid traffic sources, namely, smart meters, synchrophasors, and protective relays. Through a series of simulation studies, it then highlights a number of quality of service (QoS) and capacity issues that these applications may face within a conventional WiMAX network. A key observation from these studies is that the random access plane is the key bottleneck for supporting many of these M2M applications over a conventional WiMAX network.

To further analyse the performance of the random access plane, this thesis develops a comprehensive analytical model that incorporates all the key features of the code division multiple access (CDMA) based random access mechanism, such as multi-user multi-code transmission, parallel code detection, and back-off and retransmissions. Through this model, it formulates a number of solutions, such as an enhanced random access scheme to detect a large number of random access codes, a differentiated random access strategy to provide QoS-aware access service to various M2M devices, and an adaptive radio resource management scheme to ensure an efficient utilisation of the random access resources. Moreover, it proposes and investigates a heterogeneous network (HetNet) architecture to reap maximum benefits from a WiMAX network by improving its coverage and allowing flexible data aggregation. Finally, it presents a number of application-specific optimisations to reduce radio resource utilisation and/or improve the performance of a WiMAX-based Smart Grid communications network.

Many of the WiMAX specific analyses, results, and solutions in this thesis can be applied to other M2M applications beyond the Smart Grid. In addition, most of the traffic models developed and application-specific optimisations performed are technology agnostic; therefore, they are equally applicable to other wireless technologies, such as the 3GPP (3rd Generation Partnership Project) based LTE (Long Term Evolution).

List of Common Acronyms

3GPP	3rd Generation Partnership Project
ACSI	Abstract Communication Service Interface
AMC	Adaptive Modulation and Coding
AMI	Advanced Metering Infrastructure
AMP	Adaptive Microgrid Protection
AMR	Automatic Meter Reading
ANSI	American National Standards Institute
AP	Access Point
APDU	Application Protocol Data Unit
ARQ	Automatic Repeat Request
BE	Best Effort
BER	Bit Error Rate
BR	Bandwidth Request
BS	Base Station
CCN	Central Control Network
CDF	Cumulative Distribution Function
CDMA	Code Division Multiple Access
CFP	Contention Free Period
CFR	Channel Frequency Response
CID	Connection ID
CP	Contention Period

CPS	Cyber-Physical System
CQI	Channel Quality Indicator
CRA	CDMA-based Random Access
CRC	Cyclic Redundancy Check
CS	Convergence Sublayer
CT	Current Transformer
DCB	Directional Comparison Blocking
DCF	Distributed Coordination Function
DER	Distributed Energy Resource
DG	Distributed Generation
DIFS	Distributed Inter Frame Space
DL	Downlink
DL-MAP	Downlink - Medium Access Protocol
DMP	Differential Microgrid Protection
DNP3	Distributed Network Protocol, version 3
DR	Demand Response
DRA	Differentiated Random Access
DS	Distributed Storage
DSA	Dynamic Service Addition
DSCP	Differentiated Services Code Point
E2E	End-to-End
EV	Electric Vehicle
EVCC	EV Charging Controller
EVCS	EV Charging Station
FAN	Field Area Network
FCH	Frame Control Header
FEC	Forward Error Correction

FFT	Fast Fourier Transform
FL	Frame Latency
FLI	Frame Latency Indication
FTP	File Transfer Protocol
GOOSE	Generic Object Oriented Substation Event
GPS	Global Positioning System
GSM	Global System for Mobile Communications
GSSE	Generic Substation Status Event
HAN	Home Area Network
HARQ	Hybrid Automatic Repeat Request
HetNet	Heterogeneous Network
HMI	Human to Machine Interface
HMP	Hybrid Microgrid Protection
HTTP	Hypertext Transfer Protocol
HV	High Voltage
I/O	Input/Output
IE	Information Element
IEC	International Electrotechnical Commission
IED	Intelligent Electronic Device
IEEE	Institute of Electrical and Electronics Engineers
IP	Internet Protocol
IPv4	Internet Protocol, version 4
IPv6	Internet Protocol, version 6
IR	Initial Ranging
LAN	Local Area Network
LCDP	Line Current Differential Protection
LTE	Long Term Evolution

M2M	Machine-to-Machine
MAC	Medium Access Control
MAI	Multiple Access Interference
MCPU	Microgrid Central Protection Unit
MCS	Modulation and Coding Scheme
MDMS	Meter Data Management System
MDSR	Message Delivery Success Rate
MIMO	Multiple-Input Multiple-Output
ML	Maximum Latency
MPDU	MAC Protocol Data Unit
MSTR	Maximum Sustainable Traffic Rate
MTT	Message Transmission Time
MU	Merging Unit
NAN	Neighbourhood Area Network
NIST	National Institute of Standards and Technology
nrtPS	non real time Polling Service
OFDM	Orthogonal Frequency-Division Multiplexing
OFDMA	Orthogonal Frequency-Division Multiple Access
PCC	Point of Common Coupling
PCF	Point Coordination Function
PDC	Phasor Data Concentrator
PMU	Phasor Measurement Unit
POTT	Permissive Over-reaching Transfer Trip
PQ	Priority Queue
PT	Potential (voltage) Transformer
PUSC	Partial Usage of Subchannels
QoE	Quality of Experience

QoS	Quality of Support
QPSK	Quadrature Phase Shift Keying
RLC	Remote Load Control
ROHC	Robust Header Compression
RR	Round Robin
RRA	Radio Resource Agent
RRM	Radio Resource Management
rtPS	real time Polling Service
RTU	Remote Terminal Unit
SF	Service Flow
SFQ	Simple Fair Queuing
SG-WAN	Smart Grid Wide Area Network
SGIRM	Smart Grid Interoperability Reference Model
SIFS	Short Inter Frame Space
SMV	Sampled Measured Value
SNR	Signal-to-Noise Ratio
SoC	State of Charge
SPDC	Super Phasor Data Concerntrator
SS	Subscriber Station
TBTT	Target Beacon Transmission Time
TCP	Transmission Control Protocol
TDD	Time Division Duplex
TDM	Time Division Multiplexing
TS	Time Synchronization
TTL	Time-To-Live
TVE	Total Vector Error
UDP	User Datagram Protocol

UGS	Unsolicited Grant Service
UL	Uplink
UL-MAP	Uplink - Medium Access Protocol
UMTS	Universal Mobile Telecommunications System
V2G	Vehicle-to-Grid
VoIP	Voice over IP
VRC	Virtual Ranging Channel
WAMS	Wide Area Measurement System
WAN	Wide Area Network
WiMAX	Worldwide Interoperability for Microwave Access
WLAN	Wireless Local Area Network
WMN	Workforce Mobile Network
WWR	WiMAX-WLAN Router

Contents

1	Introduction	1
1.1	Background	3
1.2	Research Motivation	5
1.3	Aim of this Thesis	6
1.4	Contributions of this Thesis	8
1.5	Thesis Outline	9
1.6	List of Publications	12
2	M2M Communications in the Smart Grid: A Review	15
2.1	Communications Architecture of the Smart Grid	16
2.1.1	Neighbourhood Area Networks	18
2.1.2	Field Area Networks	19
2.1.3	Workforce Mobile Networks	20
2.2	Key Smart Grid Applications	21
2.2.1	Automatic Meter Reading	21
2.2.2	Demand Response	23
2.2.3	Electric Vehicles	25
2.2.4	Substation Automation	30
2.2.5	Wide Area Measurement	34
2.2.6	Microgrids	36
2.3	Performance Requirements & Key Research Challenges	40
2.4	Chapter Summary	42

3	Traffic Modelling and Performance Analysis	43
3.1	Key WiMAX Tenets	45
3.1.1	Frame Structure	45
3.1.2	QoS Management	46
3.1.3	Radio Resource Scheduling	47
3.2	AMI Communications over WiMAX	51
3.2.1	Estimated Number of Smart Meters	51
3.2.2	Simulation Study	55
3.2.3	Key Findings	58
3.3	Synchrophasor Communications over WiMAX	59
3.3.1	Traffic Modelling	59
3.3.2	Key RRM Issues	62
3.3.3	Capacity Analysis	66
3.3.4	Simulation Study	70
3.3.5	Key Findings	75
3.4	Protective Relaying over WiMAX	75
3.4.1	Traffic Model	76
3.4.2	QoS Mapping	77
3.4.3	Simulation Study	79
3.4.4	Key Findings	83
3.5	Chapter Summary	83
4	Random Access: The Key Bottleneck	85
4.1	System Description	86
4.2	Related Works	89
4.3	Analytical Modelling	91
4.3.1	Modelling Assumptions	91
4.3.2	Preliminaries and Notation	93
4.3.3	Markov Chain Analysis	95
4.3.4	Queuing Analysis	99
4.3.5	Back-off Service Time Analysis	102
4.3.6	Access Probability Analysis	108
4.4	Model Validation and Performance Analysis	109
4.5	Chapter Summary	114

5	Random Access Enhancement	115
5.1	The Proposed Random Access Scheme	117
5.1.1	Pseudo-random Code Matrix	120
5.1.2	Hadamard Code Matrix	123
5.1.3	Simulation Results	126
5.2	Differentiated Random Access	128
5.3	Adaptive Radio Resource Management	130
5.3.1	Simulation Results	134
5.3.2	Illustrative Example: Pilot Protection Scheme	136
5.4	Chapter Summary	141
6	Heterogeneous Network: The Ultimate Solution?	143
6.1	WiMAX-WLAN HetNet: An Overview	144
6.1.1	Architecture and Interworking	145
6.1.2	Modes of Operation	146
6.1.3	QoS Management	147
6.2	AMI Communications over the HetNet	150
6.2.1	Data Aggregation	151
6.2.2	Performance Analysis	152
6.3	Synchrophasor Communications over the HetNet	156
6.3.1	Data Aggregation	157
6.3.2	Delay Budget	159
6.3.3	Performance Analysis	161
6.4	Chapter Summary	164
7	Application-Specific Optimisations	165
7.1	Network Controlled EV Load Management	166
7.1.1	Basic Concept	167
7.1.2	Communications Architecture	169
7.1.3	Modelling & Assumptions	172
7.1.4	Results & Analysis	174
7.2	Predictive Synchrophasor Communications	176
7.2.1	Prediction of Voltage Instability	176

7.2.2	Adaptive PMU Reporting	179
7.2.3	Simulation Results	182
7.3	Hybrid Microgrid Protection Scheme	184
7.3.1	Adaptive Microgrid Protection Scheme	185
7.3.2	Differential Microgrid Protection Scheme	186
7.3.3	Key Communications Issues	188
7.3.4	The Hybrid Protection Scheme	190
7.3.5	Performance Analysis	192
7.4	Chapter Summary	193
8	Conclusions	194
8.1	Summary of Contributions	195
8.2	Future Research Directions	198
A	Discrete Event Simulation	200
A.1	The OPNET Modeller	200
A.2	Simulation Models	201
A.3	Statistical Validity of the Results	204
B	Multicarrier Transmission with OFDM	205
B.1	OFDM Modulation	206
B.1.1	Multiplexing with OFDM	206
B.1.2	Pulse Shaping and Transmission	208
B.2	Multipath Propagation Channel	210
B.3	OFDM Demodulation	211
B.3.1	Match Filtering and Sampling	211
B.3.2	DE-Multiplexing and Recovery	211
C	Summation Distribution of Lambda	215
C.1	Probability Density of λ	216
C.2	Mathematical Variance of λ	216
C.3	Probability Density of μ	219
C.4	Simulation	220
C.4.1	The Distribution of μ	221
C.4.2	Comparing Calculated Value and Simulated Value	222

D Voltage Instability Prediction Algorithm	225
D.1 The Prediction Algorithm	225
D.2 Comparison of Prediction Accuracy	229
D.3 Voltage Instability Predication Procedure	230
D.4 Tuning Parameter τ	230
Bibliography	231

Chapter 1

Introduction

Since its inception, the basic structure of today's electricity grid has remained unchanged. For decades, the electrical grid has been a strictly hierarchical system where electric power flows unidirectionally from generating plants towards consumer loads (see Figure 1.1). Due to the absence of appropriate communication infrastructures in the distribution domain, the existing grid effectively operates in an open-loop manner, where the control centre has no or limited real time information about the dynamic change in load and system operating conditions [1]. This poor visibility and lack of situational awareness have made the grid susceptible to frequent disturbances, which often turn into major blackouts and brownouts due to cascading failures [2].

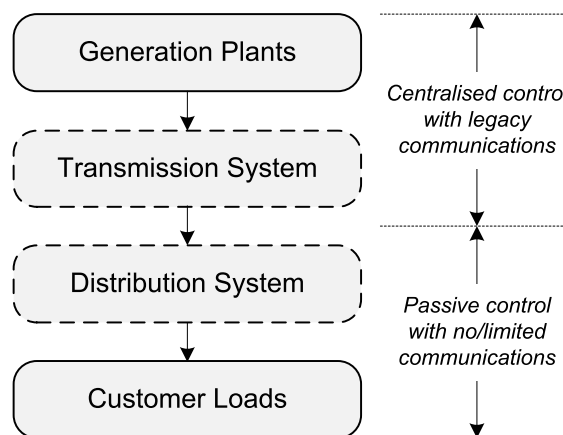


Figure 1.1: The existing electricity grid with unidirectional power flow.

Over the last decade, rising awareness about the adverse effects of climate change has prompted governments across the world to reduce their greenhouse gas emissions by promoting renewable energy sources like solar and wind power. According to a report published by the Clean Energy Council of Australia, a record 13.14 per cent of Australia's electricity generation was produced by renewable sources in 2012, which is on track to meet the national 20 per cent Renewable Energy Target by 2020 [3]. Moreover, electric vehicles (EVs) are being considered as one of the promising solutions to reduce carbon emissions and fossil fuel dependence. With the increasing adoption and use of EVs, they are expected to become a major load to the grid in the near future. For instance, according to a study by the Australian Energy Market Commission (AEMC), EV sales are projected to be around 20 per cent of total sales by 2020, and up to 45 per cent by 2030, considering a moderate uptake; this will impose an additional peak demand of 8.2 per cent (1900 MW) on the Australian national grid [4].

As most renewable energy resources are connected to the distribution domain of the grid, their wide-scale integration with the existing grid reverses the traditional direction of power flow. Moreover, EVs may also act as storage devices and feed power back to the grid, smoothing out the natural intermittency of renewable energy sources; this further strengthens the case for reverse power flows in the distribution grid. To enable bidirectional energy transfer and improve its operational efficiency, the existing grid needs advanced monitoring and control technologies to stabilise its operating parameters (e.g., voltage, current and frequency) and optimally balance its load-supply profile. Against the backdrop of these challenges, the concept of the next generation electricity grids, known as 'Smart Grid', has emerged to address the shortcomings of the existing grid. According to the U.S. Department of Energy (DOE), one of the major advocates of the concept, the Smart Grid is a distributed and automated energy delivery network that provides two-way flow of electricity and information, and enables near-instantaneous balance of supply and demand by incorporating the benefits of distributed computing and communications [2]. Thus, a robust communications infrastructure acts as a key enabler for the Smart Grid, differentiating it from the conventional grid by allowing information exchange among its different entities for data acquisition, monitoring, control and protection applications [5].

Although global Smart Grid implementation activities are still in the early stages and most commercially deployed applications are limited to automatic meter reading (AMR), utilities need to consider other prospective applications to deploy a future-proof communications network. Hence, utilities from all over the world are facing a major challenge of finding the most appropriate technology that can satisfy their current and future communication needs. Therefore, it is vital to conduct a comprehensive study to properly assess the feasibility and performance of a candidate communications technology for the Smart Grid.

1.1 Background

While the notion of a Smart Grid is relatively new, many utilities already have their own telecommunications infrastructure, mostly based on wireline networks, such as fibre-optic and cable. They are typically used for narrowband applications in substations and the control centre. However, with the ever increasing portfolio of Smart Grid applications, the utilities need a ubiquitous communications network that can provide sufficient bandwidth to support a large number of devices with a wide range of traffic requirements. A possible solution may be to extend the existing wireline network. However, extending this to hundreds or thousands of homes and businesses for utility purposes alone is time and cost prohibitive. In contrast, contemporary wireless technologies have matured enough to provide reliable coverage to a vast number of users. Moreover, they are fast, reliable, and relatively easy to install, have a low cost per bit, and do not require major construction effort or structural intervention in buildings. Further, many support mobility, and can provide mobile workforce connectivity to employees for logistics and asset management services [6].

Among the contemporary wireless technologies, the fourth generation (4G) cellular wireless networks are particularly well suited to meet the existing and future communication needs of a Smart Grid. They offer several key benefits/advantages, such as ubiquitous coverage, broadband capacity, low latency, enhanced quality of support (QoS), and built-in security capabilities. Other technologies, such as narrowband digital radio networks, sec-

ond/third generation (2G/3G) cellular networks, and wireless local area networks (WLANs) can support some applications, but not all. For example, the 2G/3G networks are well-known for their good coverage and have enough bandwidth to support metering applications, but they do not have the required capacity for high-bandwidth data applications, such as synchrophasor reporting and remote surveillance. In contrast, WLAN may support some of these applications, but typically lacks the ubiquitous coverage needed to support metering applications and the QoS mechanisms required to support mobile workforce applications.

Of the two recognised 4G wireless technologies by the international telecommunication union radiocommunication (ITU-R) sector, both WiMAX (Worldwide Interoperability for Microwave Access) and 3GPP (3rd Generation Partnership Project) based LTE (Long Term Evolution) are strong contenders to build the next generation Smart Grid communications networks. While WiMAX is based on the IEEE 802.16 family of standards, LTE evolved from the previous generations of 3GPP standards, such as GSM and UMTS. Despite fierce competition in the global 4G market, both WiMAX and LTE have many similarities in their technical features. For example, both of them use orthogonal frequency division multiplexing (OFDM) to combat multipath fading, multiple-input multiple-output (MIMO) antenna system for increased throughput, and scalable bandwidth allocation for flexible spectrum usage.

In recent years, many utilities around the world have already started (or in the process of) deploying 4G wireless based advanced metering infrastructure (AMI) networks as a first step towards the Smart Grid. For example, the world's first WiMAX-based AMI network is under deployment by SP-AusNet in Victoria, Australia; UQ communications in Japan is deploying about seven million WiMAX-enabled smart meters over the next 10 years; Center point energy in Texas, U.S. is building an AMI network of 2.3 million WiMAX-enabled smart meters from General Electric (GE). Other major WiMAX-based AMI projects include Ausgrid (Australia), Salzburg AG (Austria), PPL Electric (U.S.), and A and N Electric Cooperative (U.S.) [7]. Conversely, several other utilities and network operators in the USA are planning to deploy LTE, instead of WiMAX for their Smart Grid/AMI commu-

nications network. However, while most telecommunications operators around the world have chosen LTE over WiMAX due to its ability to collaborate and coexist with the 2G/3G networks, WiMAX is finding a niche among the utility operators as their preferred Smart Grid solution.

1.2 Research Motivation

The IEEE 802.16-based WiMAX networks can significantly extend the reach of the Smart Grid by allowing fast and reliable communications over a wide coverage area. However, it must be noted that the radio resource management (RRM) techniques of WiMAX, like other conventional telecommunication technologies, are optimised to support various multimedia applications, such as voice and video. In contrast, the communications requirements of the Smart Grid are quite different. In a Smart Grid environment, a communications network has to support information exchange among a large number of smart meters, intelligent electronic devices (IEDs), and sensors/actuators, without (or with very limited) human intervention. This form of communication is often known as machine-to-machine (M2M) communications; it is autonomous and triggered either by time or events. Each of these applications has different characteristics and traffic requirements, depending on the type of the M2M device (e.g., meter, sensor, and controller) and its modes of operation (e.g., normal, alert, and fault). For example, while the meter reading traffic from a smart meter is fairly delay tolerant, demand management traffic from the same device is much more delay sensitive; similarly, traffic priority of a demand management event and a substation event are quite different. Thus, the coexistence of protection, control, monitoring and reporting applications poses the additional challenge of strict QoS differentiation within a multi-service Smart Grid communications network. Therefore, significant uncertainty exists concerning the applicability of WiMAX in Smart Grid environments.

So far, only a modest amount of research has been conducted into the performance of a WiMAX network in a Smart Grid environment, mostly by the IEEE 802.16p task groups within the general framework of M2M communications. In 2010, the IEEE 802.16p work-

ing group released the ‘M2M communications technical report’, defining the high level communications requirements for M2M applications over a WiMAX network [8]. Later in 2012, the IEEE 802.16p amendment was ratified to optimise the IEEE 802.16 air interface to support a number of basic M2M features, such as low power transmission, random access regulation, small data bursts, and device authentication [9]. However, the amendment mainly focused on the metering applications to support ongoing deployment activities around the utility world.

It is worthwhile to note that, although AMR is a basic application of the Smart Grid, it is not the only one. A fully-functional Smart Grid is envisaged as a highly cross-functional system, which will embed computation, communications and control technologies to support a number of other advanced applications, such as wide area situational awareness and protective relaying. This will create a so-called ‘cyber-physical system’ (CPS), where the communications network will bridge the cyber and physical processes of the grid [10]. Consequently, communications requirements (e.g., throughput, latency, and packet loss) may vary with the state of these underlying physical processes. For instance, during a fault event in the grid, a protective relay may trigger trip/block signals that need to be transferred as quickly as possible to allow fast fault detection and isolation. Hence, the communications network needs to meet the newly generated QoS requirements for these packets in a multi-service Smart Grid environment, where multiple applications will contend for network resources with their diverse QoS requirements. Thus, the cyber-physical dynamics of the Smart Grid creates a new set of problems and research opportunities for each individual application. This is an area where this thesis makes a number of novel contributions.

1.3 Aim of this Thesis

The aim of this thesis is to offer an in-depth study of M2M communications over a WiMAX network in the Smart Grid¹. To fulfil this aim, it first conducts a detailed, technology agnostic review on the application characteristics and traffic requirements of several major

¹The terms IEEE 802.16 and WiMAX are used synonymously throughout this thesis to refer to the current generation of WiMAX networks based on the IEEE 802.16-2009 standard.

Smart Grid applications and highlights their key communications challenges. Based on this review, it develops a range of traffic models for some key Smart Grid traffic sources, namely, smart meters, synchrophasors, and protective relays. Through a series of simulation studies, it then highlights a number of QoS and capacity issues that these applications may face within a conventional WiMAX network. A key observation from these studies is that the random access plane is the key bottleneck for supporting many of these M2M applications over a conventional WiMAX network. To further analyse the performance of the random access plane, this thesis develops a comprehensive analytical model that incorporates all key features of the code division multiple access (CDMA) based random access mechanism, such as multi-user multi-code transmission, parallel code detection, and back-off and retransmissions. Through this model, it formulates a number of solutions, such as an enhanced random access scheme to detect a large number of random access codes, a differentiated random access (DRA) strategy to provide QoS-aware access service to various M2M devices, and an adaptive RRM framework to ensure an efficient utilisation of the random access resources. Further, it proposes and investigates a heterogeneous network (HetNet) architecture based on WiMAX-WLAN synergy to reap maximum benefits from a WiMAX network by improving coverage and allowing flexible data aggregation. Lastly, it presents a number of application-specific optimisations to reduce radio resource utilisation and/or improve performance of a WiMAX-based Smart Grid communications network.

Note that this thesis focuses only on the M2M applications of the Smart Grid. The conventional telecommunications applications, such as voice over internet protocol (VoIP), streaming multimedia and internet, which will also be present in a Smart Grid communications network, are not considered. Moreover, while security is a key issue in a large and complex cyber-physical system such as the Smart Grid, it is not within the scope of this thesis.

1.4 Contributions of this Thesis

The specific contributions of each chapter are listed in their respective chapters throughout the thesis. The key contributions of this thesis are:

- Provides an in-depth review on the characteristics and traffic requirements of several major Smart Grid applications, namely, AMR, demand response, electric vehicles, substation automation, wide area measurement system (WAMS), and microgrids, and summarises their key performance requirements in the context of a multi-service Smart Grid communications network.
- Develops traffic models for some key Smart Grid entities, such as smart meters, synchrophasors, and protective relays; conducts a number of simulation studies to examine their performance over a conventional WiMAX network; and highlights some key research challenges present in this arena.
- Develops a unified analytical model to evaluate the performance of the WiMAX random access plane, which was identified as a key bottleneck during simulation studies of the key Smart Grid traffic sources. The proposed model integrates an extended Markov chain model with a detailed queuing model to obtain closed-form expressions for the random access delay and throughput, under both saturated and unsaturated conditions. Compared to the literature, it incorporates all key features of CDMA-based random access (CRA) procedures, such as multi-user multi-code transmission, parallel code detection, and back-off parameters in the performance analysis matrix.
- Provides an enhanced random access code transmission scheme, which is demonstrated to detect a large number of codes using a simple, low-complexity detector. Based on this scheme, a DRA strategy is formulated to provide QoS-aware access service to various M2M devices in a WiMAX network. Further, an adaptive RRM framework is proposed based on this strategy that efficiently allocates the random access resources, which in turn is based on the latency requirement and priority of a Smart Grid application.

- Introduces a multi-tier HetNet architecture based on WiMAX-WLAN synergy to improve coverage and to facilitate hierarchical data aggregation; conducts performance analysis of the proposed HetNet architecture; and benchmarks the results against a standalone WiMAX network.
- Presents a number of application-specific optimisations, including a novel EV load management scheme, an adaptive synchrophasor reporting scheme, and a hybrid microgrid protection scheme to reduce the communications load of the WiMAX-based Smart Grid communications network.

Before moving on, it should be noted that many WiMAX specific analyses, results and in this thesis can be applied to other M2M applications beyond the Smart Grid. In addition, most of the traffic models developed, and the application-specific optimisations performed, are technology agnostic; therefore, they are equally applicable to other wireless technologies, such as the 3GPP-based LTE.

1.5 Thesis Outline

The remainder of this thesis is organised into one chapter of overview, five chapters of contributions, and a chapter of conclusions. The thematic connections among the chapters of contributions are illustrated in Figure 1.2.

Chapter 3 forms the basis of this study. It develops traffic models and conducts simulation studies based on the inputs from Chapter 2. Chapter 4 extensively analyses the performance of the WiMAX random access plane using a unified analytical model, which serves as a precursor for the performance enhancement activities conducted in Chapter 5. Based on the capacity and coverage issues identified in Chapter 2 and the random access bottleneck examined in Chapter 4, Chapter 6 introduces the HetNet architecture to further improve the utility of a WiMAX network for Smart Grid communications. Lastly, Chapter 7 presents a number of application-specific optimisations to reduce the radio resource usage of the WiMAX network for EV charging, synchrophasor and microgrid applications reviewed in Chapter 3.

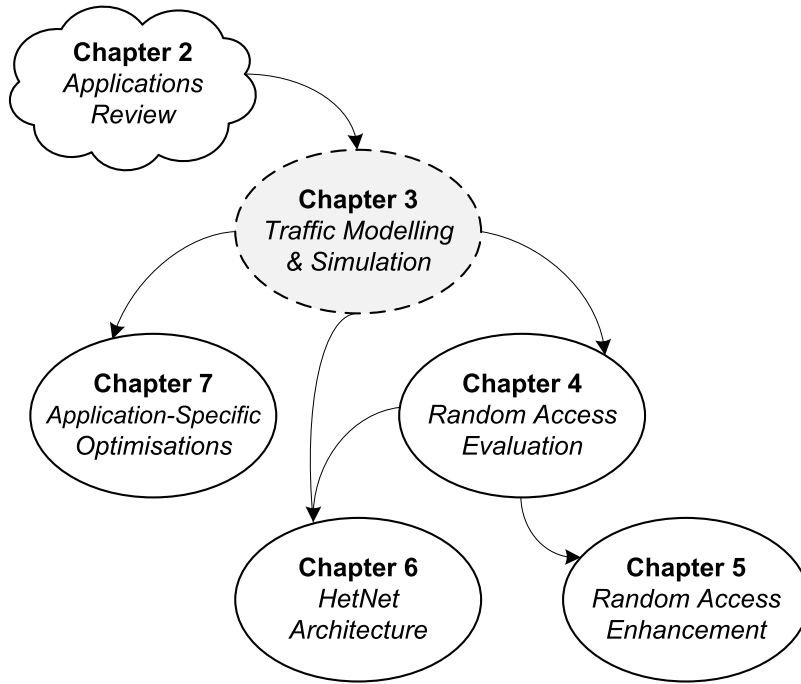


Figure 1.2: Thematic connections among the contributing chapters of this thesis.

Excerpts of each chapter are elucidated as follows:

Chapter 2 presents an overview of the communications architecture of the Smart Grid based on the IEEE 2030-2011 standard for Smart Grid interoperability. In particular, it describes the three key communications entities of a Smart Grid wide area network (SG-WAN), which is the focus of this thesis; these are: the neighbourhood area network (NAN), the field area network (FAN), and the workforce mobile network (WMN). This is followed by an in-depth review of the characteristics and traffic requirements of some major Smart Grid applications and a summary of their key performance requirements in the context of a Smart Grid communications network.

Chapter 3 develops traffic models and conducts OPNET-based simulation studies of the three key Smart Grid M2M traffic sources, namely, AMI/smart meters, synchrophasors/phasor measurement units (PMUs), and protective relays over a conventional WiMAX network. These studies provide unprecedented insight into the characteristics and performance of these applications and identify a number of key capacity and QoS issues that require addressing to further optimise and improve their performances.

Chapter 4 further examines the WiMAX random access plane, which was identified as a key bottleneck for the Smart Grid M2M applications during the simulation studies of the previous chapter. A comprehensive analytical model is developed to accurately analyse the performance of the CRA mechanism in WiMAX. Compared with the literature, the model incorporates unique features of the CRA procedure, such as multi-user multi-code transmission and parallel code detection in the performance analysis matrix. The unsaturated conditions are captured through a detailed queuing analysis and integrated with an extended Markov chain model to obtain closed-form expressions for the random access delay and throughput, under both saturation and non-saturation conditions. The accuracy of the model is validated by an extensive set of simulation results, based on OPNET.

Chapter 5 is devoted to the random access bottleneck of the WiMAX network through a number of cross-layer approaches. First, an efficient code transmission scheme is proposed, where the fixed/low-mobility M2M devices pre-equalise their random access codes using the estimated frequency response of the slowly-varying wireless channel. Consequently, the base station (BS) can detect a large number of codes using a simple, low-complexity detector, as their mutual orthogonality remains preserved. Next, a DRA strategy is proposed to provide QoS-aware access service to various M2M devices. Finally, an adaptive RRM framework is proposed based on this strategy that efficiently allocates the random access resources based on the latency requirement and priority of a Smart Grid application.

Chapter 6 argues the case for a multi-tier HetNet to better serve the Smart Grid applications using a WiMAX network. WLAN is used as the partner technology for the HetNet, alongside WiMAX, due to their technical similarities, as well as the commercial availability of the dual radio WiMAX-WLAN routers. The behaviour and performance of the HetNet are examined through a number of simulation studies based on OPNET. The results are benchmarked against a standalone WiMAX network.

Chapter 7 presents a number of application-specific optimisations, such as a novel EV load management scheme, an adaptive synchrophasor reporting scheme, and a hybrid microgrid protection scheme, to reduce the communications load of the WiMAX-based Smart Grid communications network. The performance of these applications is demonstrated using a

number of simulation studies based on OPNET.

The last chapter summarises the main contributions and results presented throughout this thesis, followed by some thoughts on future research directions.

1.6 List of Publications

During the start of this research in early 2011, Smart Grid was still a fledgling concept. Therefore, the author had to publish a number of papers to establish possible Smart Grid application scenarios and analyse their effects on a broadband wireless network (i.e. WiMAX). The majority of the papers were published in peer-reviewed conferences to keep pace with ongoing research activities from the Smart Grid community. The complete list of publications related to this research is given below.

Refereed Journal Articles

- **R. Khan** and J. Khan, A comprehensive review of the application characteristics and traffic requirements of a Smart Grid communications network, *Computer Networks*, 57(3):825-845, Feb 2013.
- **R. Khan**, K. Mahata, and J. Khan, A novel scheduling service for distribution pilot protection over a WiMAX network, *Recent Patents on Telecommunication*, 2(1):26-40, Oct 2013.
- T. Ustun and **R. Khan**, Multi-terminal hybrid protection of microgrids over wireless communications networks, *IEEE Transactions on Smart Grid*, May 2014 (submitted for publication).

International Refereed Conference Proceedings

- **R. Khan** and J. Khan, A heterogeneous WiMAX-WLAN network for AMI communications in the Smart Grid, in *Smart Grid Communications (SmartGridComm)*, 2012 IEEE Third International Conference on, Nov 2012, pp. 710-715.

- **R. Khan** and J. Khan, Wide area PMU communication over a WiMAX network in the Smart Grid, in *Smart Grid Communications (SmartGridComm), 2012 IEEE Third International Conference on*, Nov 2012, pp. 187-192.
- **R. Khan**, J. Brown, and J. Khan, Pilot protection schemes over a multi-service WiMAX network in the Smart Grid, in *Communications Workshops (ICC), 2013 IEEE International Conference on*, June 2013, pp. 994-999.
- **R. Khan**, T. Ustun, and J. Khan, Differential protection of microgrids over a WiMAX network, in *Smart Grid Communications (SmartGridComm), 2013 IEEE International Conference on*, Oct 2013, pp. 732-737.
- M. Hyder, **R. Khan**, K. Mahata, and J. Khan, A predictive protection scheme based on adaptive synchrophasor communications, in *Smart Grid Communications (SmartGridComm), 2013 IEEE International Conference on*, Oct 2013, pp. 762-767.
- **R. Khan**, S. Stüdli, and J. Khan, A network controlled load management scheme for domestic charging of electric vehicles, in *Power Engineering Conference (AUPEC), 2013 Australasian Universities*, Sept 2013, pp. 1-6.
- **R. Khan**, K. Mahata, and J. Brown, An adaptive RRM scheme for Smart Grid M2M applications over a WiMAX network, in *Communication Systems, Networks, and Digital Signal Processing (CSNDSP), IEEE/IET International Symposium on*, July 2014 (in proceedings).
- **R. Khan**, M. Hyder, and K. Mahata, Random access performance of a WiMAX network for M2M communications in the Smart Grid, in *Smart Grid Communications (SmartGridComm), 2014 IEEE International Conference on*, Nov 2014 (in proceedings).
- M. Hyder, **R. Khan**, and K. Mahata, An enhanced random access mechanism for Smart Grid M2M communications in WiMAX networks, in *Smart Grid Communications (SmartGridComm), 2014 IEEE International Conference on*, Nov 2014 (in proceedings).

Others (related to the research, but not included in this thesis)

- S. Stüdli, **R. Khan**, R. Middleton, and J. Khan, Performance analysis of an AIMD based EV charging algorithm over a wireless network, in *Power Engineering Conference (AUPEC), 2013 Australasian Universities*, Sept 2013, pp. 1-6.
- T. Ustun, **R. Khan**, A. Hadbah, and A. Kalam, An adaptive microgrid protection scheme based on a wide-area Smart Grid communications network, in *Communications (LATINCOM), 2013 IEEE Latin-America Conference on*, Nov 2013, pp. 1-5.
- M. Hyder, K. Mahata, and **R. Khan**, Online monitoring of steady-state voltage stability limit using wide area PMU measurements, *IEEE Transactions on Smart Grid*, Mar 2014 (submitted for Publication).

Chapter 2

M2M Communications in the Smart Grid: A Review

The next generation Smart Grid is envisaged to use advanced monitoring, control and protection technologies to enable active customer participation, integrate distributed energy resources and implement self-healing functionalities. A ubiquitous communications network is the key enabler of such a complex, multidimensional energy delivery system, allowing information exchange among a large number of distributed devices over a vast geographic area. Much information exchange will take place between: two or more end devices; or an end device and a remote server located either in a local/regional substation or the central control centre, with no or minimal human interaction. This form of communication is often characterised as M2M and is considered a building block of the so-called ‘Internet of Things’ (IoT).

Although at present, most of the commercially available Smart Grid applications are limited to smart metering and non-real time demand side management, utilities need to consider other prospective applications. They need to develop a future-proof network that can satisfy their current and future communications needs. Therefore, to properly study the use of a 4G broadband wireless network such as WiMAX for Smart Grid communications, it is vital to acquire detailed knowledge about its architecture and applications.

The aim of this chapter is to provide a detailed overview of the characteristics and traffic

requirements of some of the key Smart Grid M2M applications, namely: AMR, demand response, EVs, substation automation, microgrids and wide area measurement, and highlight some of their key performance requirements in the context of a Smart Grid communications network.

The rest of the chapter is organised as follows. Section 2.1 reviews the IEEE 2030-2011 standard for Smart Grid interoperability to develop an understanding of the Smart Grid communications architecture and identify its key components. Section 2.2 describes the application characteristics and traffic requirements of the aforementioned Smart Grid applications, and highlights their key communications challenges. Section 2.3 summarises the performance requirements of these applications. Finally, Section 2.4 concludes the chapter ¹.

2.1 Communications Architecture of the Smart Grid

To study the applications of a Smart Grid, it is important to develop an understanding of its architecture first. Although it is very difficult to reach a consensus about the architecture and scope of a highly cross-functional infrastructure as the Smart Grid, the IEEE 2030–2011 standards are widely accepted as the industry’s first guideline regarding its architecture and interoperability [12]. It provides the Smart Grid interoperability reference model (SGIRM) that uses a systems-level approach to provide guidance on interoperability among various components of communications, power systems and information technology platforms in the Smart Grid. The communications technology perspective of SGIRM provides a broad set of communications technologies to interconnect Smart Grid generation, transmission, distribution and customer domains. Together they constitute an end-to-end (E2E) Smart Grid network. A representative model of the E2E Smart Grid network architecture, based on the IEEE 2030 standard is provided in Figure 2.1.

The power generation and transmission domains in Figure 2.1 are generally comprised of large substations that already have legacy communication infrastructures in place. A

¹The contents of this chapter has been published as a review paper in vol. 57, no. 3 of the Elsevier journal *Computer Networks* in Feb 2012 [11].

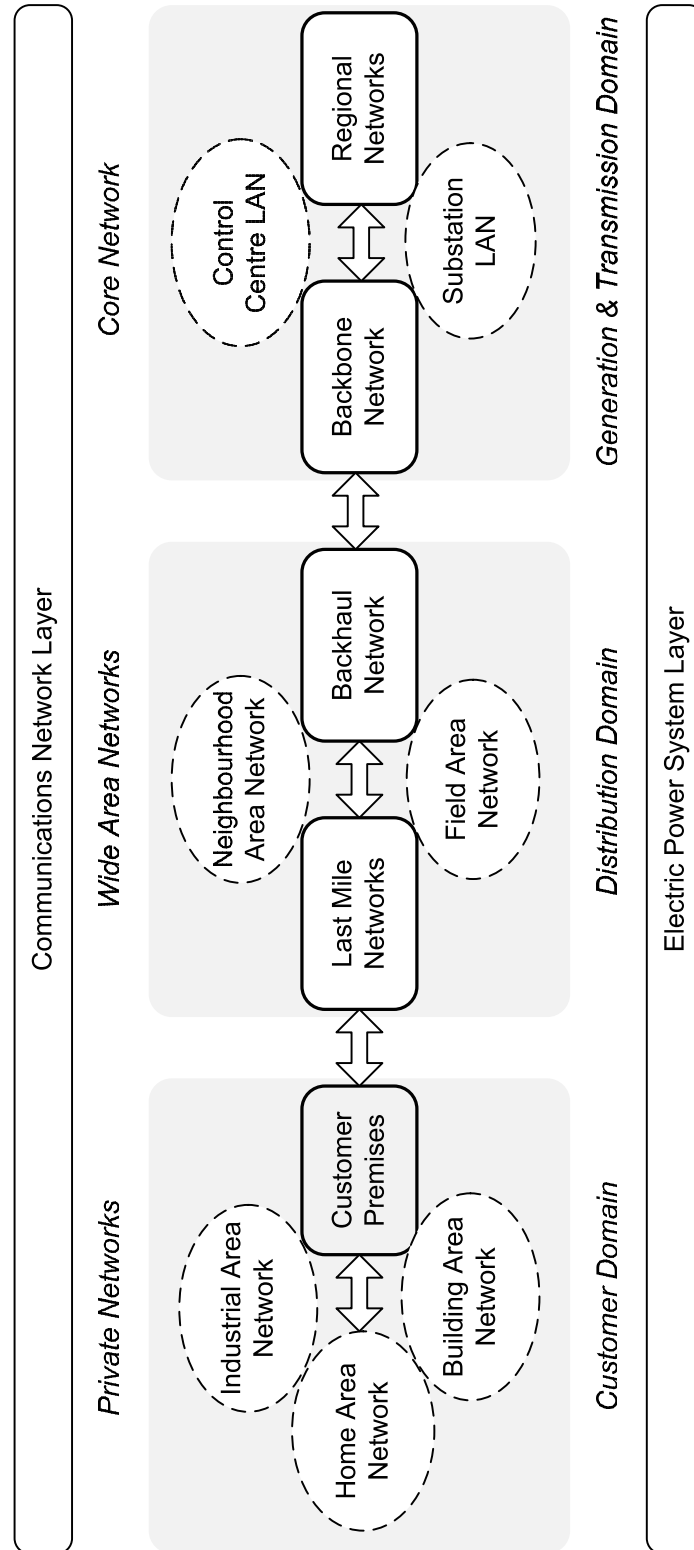


Figure 2.1: E2E Smart Grid network architecture based on the IEEE 2030 standard [12].

backbone network connects these substations with the utility control center and third party networks using mostly wireline infrastructures, such as digital subscriber line (DSL), fibre and cables. The power distribution domain contains substation and feeder equipments, as well as end-user loads over a vast geographic area. A wide area communications network, such as WiMAX, connects these infrastructures with the utility control centre. Additionally, it provides ‘last mile’ connectivity to the customer premises to support various end-user applications within a home area network (HAN) or a building area network (BAN) or an industrial area network (IAN). A backhaul network connects this wide area network (WAN) with the utility’s backbone network.

Figure 2.1 clearly shows that the WAN acts as a communications hub for the E2E Smart Grid network, connecting the power generation, transmission and distribution domains to enable grid-wise monitoring and control applications. Since the aim of this thesis is to study the application of a wide area wireless network (i.e., WiMAX) in the Smart Grid, the terms ‘Smart Grid WAN’ and ‘Smart Grid’ communications network are used synonymously throughout the thesis.

The IEEE 2030 standard further specifies a logical architecture for the Smart Grid WAN, comprised of the following three subnetworks: (i) NAN for customer premises; (ii) FAN for the utility infrastructures; and (iii) WMN for the utility workforce, as shown in Figure 2.2. A brief description of these logical subnetworks is given in the following subsections.

2.1.1 Neighbourhood Area Networks

A NAN is the logical representation of an AMI system that connects customer premises with the utility control centre. The basic constituent of a NAN/AMI network is a smart meter that performs a variety of intelligent metering tasks, such as consumption metering, power quality monitoring, and can also optionally perform load control activities. Additionally, a smart meter may act as an energy services interface (ESI) that allows a private network (e.g., a HAN) to exchange information with the AMI system. It supports a number of advanced applications, such as remote load control, monitoring and control of distributed energy resources and EVs, in-home display of customer usage, and reading of non-energy

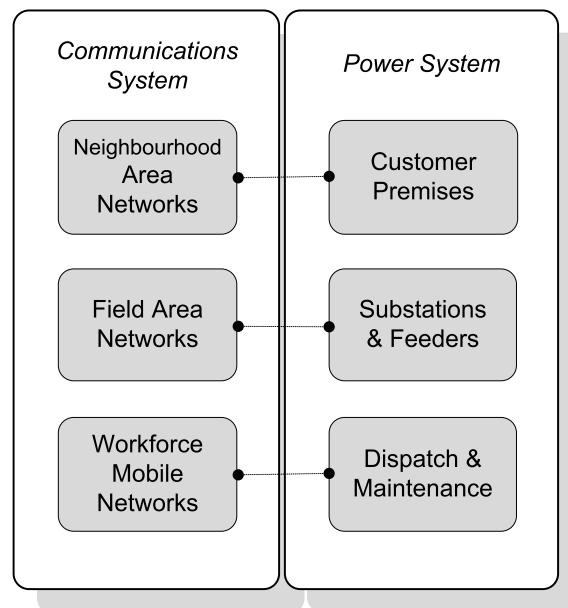


Figure 2.2: Logical architecture of a Smart Grid wide area communications network.

meters (e.g., water and gas). Figure 2.3 illustrates the types of connectivity within a Smart Grid NAN. The end-points of a NAN can either be a standalone smart meter or data aggregation point (DAP) that collects data to and from a group of smart meters and sends the aggregated information to the meter data management system (MDMS) via a backbone network.

2.1.2 Field Area Networks

The FAN is responsible for facilitating information exchange between the utility control centre and the distribution substation and feeder equipments for various monitoring, control and protection applications. The distribution substations convert high voltage (HV) electricity into low voltage, as required by homes and businesses. In addition to voltage transformation, they also isolate faults and are used as a point of voltage regulation. In a Smart Grid environment, they will be equipped with advanced monitoring and control equipments, such as remote terminal units (RTUs), phasor measurement units (PMUs) and IEDs, to perform various substation automation functions. The distribution feeders include

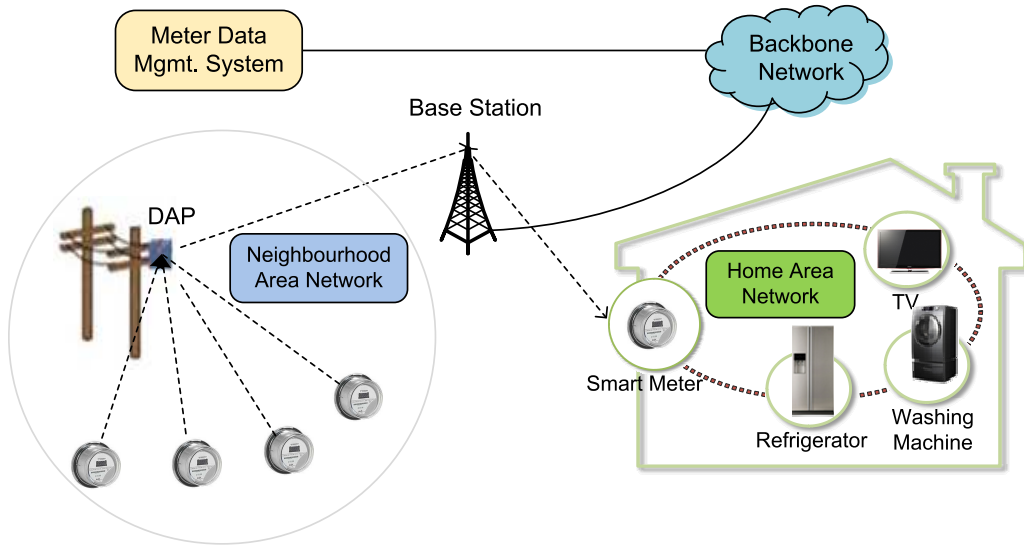


Figure 2.3: Types of connectivity within a Smart Grid NAN.

transmission lines, cable poles and towers to provide electricity connection to customer premises. They also act as a point of common coupling (PCC) for the distributed energy resources (DERs) and microgrids, which are connected with the distribution side of the grid. Moreover, there might be a number of sensor and actuator networks overlaid along the distribution feeders for distribution supervision and monitoring applications.

2.1.3 Workforce Mobile Networks

The WMN is used by the utility's workforce to provide dispatch, maintenance and normal day-to-day operations. Typical requirements of a WMN include broadband connectivity to employees, including virtual private network (VPN), VoIP and geographic information system (GIS) based applications for logistics or asset management. Moreover, in-vehicle applications and fleet telematics, such as location-based services (LBS) with global positioning system (GPS) based tracking and navigation and automatic vehicle locating (AVL), are also expected to be integrated with the WMN system [6]. The WMN should be able to access both NAN and FAN via the WAN to collect information from the smart meters and field devices.

The typical communication requirements of a WMN are similar to that of standard non-

M2M telecommunication services, such as voice, video and internet applications. In addition, mobility support is typically required for many of these applications. However, since the focus of this thesis is mainly the M2M communications in the Smart Grid, applications related to the WMN are out of its scope.

2.2 Key Smart Grid Applications

Based on the two logical subnetworks of a Smart Grid WAN—for instance, NAN and FAN—a plethora of applications with different requirements and features are expected to emerge. However, to limit the scope of this thesis, only a selected set of applications are considered, which have drawn significant attention from the industry and research communities. In the rest of this section, we discuss the characteristics and traffic requirements of these applications and highlight some key challenges in the context of a Smart Grid communications network.

2.2.1 Automatic Meter Reading

AMR uses communication systems to collect meter readings, as well as events and alarm data from meters. It is the most basic and simplest Smart Grid application. Several published standards are available for AMR applications. Of these, American National Standards Institute (ANSI) C12.19-2008, IEEE 1377-2012, and International Electrotechnical Commission (IEC) 61968-9 are the prominent ones [13, 14, 15]. Both the ANSI C12.19 and the IEEE 1377-2012 standards provide specifications for the communication syntax for data exchange between the end device and the utility server, using binary codes and extensible markup language (XML) contents. However, the scope of the IEC 61968-9 standard is more general and covers various aspects of AMR communications, including meter reading, meter connect/disconnect, meter data management, outage detection, and prepaid metering. Based on the IEC 61968-9 standard, some major AMR communications scenarios are listed below:

- **Meter Reading.** Smart meters may send meter readings according to a schedule defined by the MDMS. In addition, both the customer and the MDMS may request meter readings on demand for billing inquiries, outage extent verification, and verification of the restoration;
- **Meter Events and Alarms.** Smart meters may report various meter health events and alarms to notify of hardware, configuration or connection issues. Examples of such communication include: diagnostic alarms, tamper alarms or other unusual conditions. In addition, the meters can participate in periodic software and firmware upgrade activities;
- **Grid Events and Alarms.** Smart meters may collect information related to grid events such as momentary outage, sustained outage, low or high voltage, and high distortion. They can then send the report to the MDMS as an event or alarm. This information could be used for outage analysis, maintenance scheduling or capacity planning;
- **Others.** Smart meters may support other advanced applications such as: receiving pricing information, supporting prepaid services, and supporting customer switch between energy suppliers.

Typically, an AMR network comprises a large number of devices, as it has to collect measurements from each residential and commercial meter within its coverage. For example, according to the *Smart Grid Priority Action Plan 2* (PAP2) report released by the U.S. National Institute of Standards and Technology (NIST), meter density is 100, 800 and 2000 per Km^2 for rural, suburban and urban areas respectively [16].

Note that in the event of a widespread power outage, affected meters may be required to send a ‘last gasp’ alarm to the control centre. As meters are expected to operate without battery backup and rely on the charge stored in a capacitor to send this alarm, it needs to be sent out within a few hundred milliseconds [8]. Hence, providing network access to a large number of devices within a short interval is an important requirement for the AMR applications.

The AMR applications also have a very small data rate. For example, a typical meter read-

ing report is 100–200 bytes [16]. This requires a very high signalling overhead per packet due to the signalling messages associated with establishing network connectivity. Moreover, the protocol overhead associated with each packet is also very high. For example, a 100 byte meter reading payload associated with a 40 byte IPv6 (IP version 6) header and a 20 byte transmission control protocol (TCP) header, may yield a protocol overhead of 60 per cent. Hence, increasing data transmission efficiency by minimising signalling and protocol overheads is another key requirement for AMR applications [8].

2.2.2 Demand Response

Demand response (DR) enables a utility operator to optimally balance between power generation and consumption, either by offering dynamic pricing or by implementing various load control programs.

The dynamic pricing programs are intended to reduce energy consumption during peak hours by encouraging customers (through various incentives) to limit energy use or shift use to other periods. In such supplications, the primary role of a communications network is to convey pricing information to the smart meter/ESI of the customer premises in the form of a price signal. The customer, upon receiving the price signals, takes necessary steps to regulate his or her energy consumption. Currently, several dynamic pricing programs exist in the market [17]:

- **Time-of-use (TOU)** - The day is divided into contiguous blocks of hours with varying prices, with the highest price for the on-peak block,
- **Real time pricing (RTP)** - The price may vary hourly and is tied to the real market cost of delivering electricity,
- **Critical peak pricing (CPP)** - This is the same as TOU, except it is only applied on a relatively small number of ‘event’ days, and
- **Peak time rebates (PTR)** - Customers receive electricity bill rebates for not using power (during peak periods).

A more advanced DR operation executes remote load control (RLC) programs, where the response of household appliances to price signals is automated with the help of M2M communications and intelligent appliance control techniques. For example, a thermostat can be programmed by a remote server to increase set-point temperatures in response to a critical peak pricing event signal.

For RLC programs, the loads can be classified into the following three categories based on their ability to be regulated:

- **Interruptible Loads.** These loads can be interrupted during the peak period and shifted to another time. Examples of such loads include: water pumps, dryers and dishwashers. However, these loads will again come to the grid when the waiting period is over. For instance, the load from the water pumping system can easily be shed for as long as the water is available in tanks. However, after load shedding, the tanks have to be filled again, causing a rebound of the load to the grid [18]. These loads require a simple load control signal to interrupt and re-schedule the process;
- **Reducible Loads.** These loads can be reduced to a lower level for a certain amount of time. For example, during the peak hour, refrigerator or air-conditioner thermostats can be set at a higher temperature, which will reduce the overall load. These loads need periodic interaction with the remote DR server during the load regulation time;
- **Partially Interruptible Loads.** These loads can be partially interrupted over the peak period by limiting the run-time cycle. For example, a 50 per cent cycling would result in 30 minutes of run-time per hour. Examples of such loads include: air-conditioners, electric heaters, washing machines and electric cookers. These loads require two load control signals to start and end the limited cycle mode.

Among DR applications, communication loads varies according to the connected loads and the DR method used. For price-based programs, the communications load is very low, as an event notification followed by an acknowledgment message is sufficient for each session. A more intense information exchange is required for RLC programs. Among these, the

communications load for the interruptible loads are smaller, as a few control signals are sufficient to interrupt and resume the load for a certain time. Conversely, for reducible and partially interruptible loads, the communications load is higher, as they need frequent control information exchange during the load regulation period.

The overall communications requirements for the DR applications depend on a range of factors, such as the operator's policy, system loading, wholesale energy price and weather conditions. Moreover, many of these factors are inter-related. For example, on a hot summer day, electricity demand can increase significantly, along with the wholesale energy price. As a result, the DR controller may have to execute a number of RLC sessions simultaneously. In turn, this will increase the traffic load in the underlying Smart Grid communications network. Hence, the communications network should be robust enough to cope with such busy traffic. In contrast, most price-based DR programs are based on publisher/subscriber mechanisms, where a remote server sends pricing signals to the subscribed nodes using multicast techniques. Although the transmission of pricing information is generally delay tolerant, they require reliable communication and are therefore very sensitive to packet loss.

2.2.3 Electric Vehicles

EVs are the motor vehicles with rechargeable battery packs that can be charged from the distribution feeder. While EVs may be charged at both public charging stations and home, domestic charging is expected to be the most popular choice, considering its flexibility and low overhead cost [19]. Charging EVs at home can operate either in a controlled or uncontrolled manner. According to the analysis presented in [20], most EV users will return home and plug the EV into the charging system in the evening. With uncontrolled charging, this may coincide with the normal evening peak of the grid, which may significantly affect the stability and reliability of the overall power system [21].

To mitigate the above problem, the concept of 'Smart Charging' has emerged; this enables coordinated charging of EVs, using the two-way communication capabilities of the Smart Grid. Since EVs remain parked for a longer period, and the trip schedule is often known

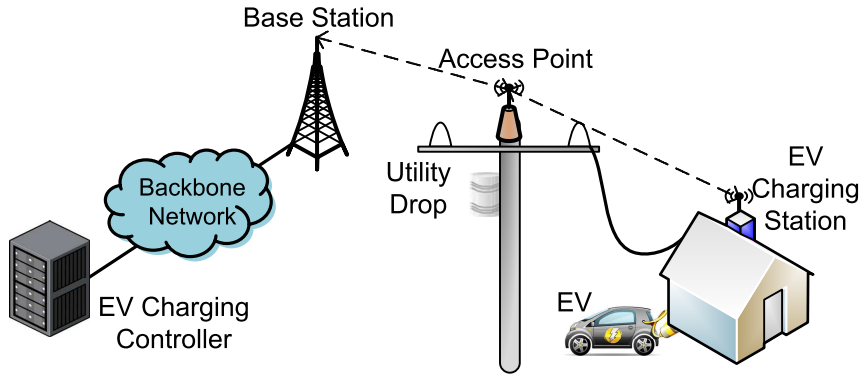


Figure 2.4: Illustration of smart EV charging using the Smart Grid communications infrastructure.

ahead, they provide a higher degree of load shifting ability. Moreover, EVs can be used as distributed storage devices, absorbing excess energy from renewable energy sources, and feeding power back to the grid under the vehicle-to-grid (V2G) operation [22]. Hence, an obvious extension of the smart charging concept is to incorporate real time DR programs known as ‘Demand Dispatch’ [23]. Demand dispatch aims to aggregate a large number of controllable loads, like EVs, to improve grid energy efficiency by optimally balancing its load-supply profile. One promising way to achieve this goal is to use EV fleets for night time valley-filling of the daily load curve [24]. A simple illustration of the smart EV charging concept is provided in Figure 2.4.

Smart Charging Architecture

The key instrument behind the ‘Smart Charging’ concepts is a centralised EV charging controller (EVCC) located at the utility’s control centre, which is responsible for coordinating each energy transfer session in real time to accommodate the time-varying nature of the total available power and the number of EVs being charged. To accomplish this, the EVCC sends control signals to and receives the state of charge (SoC) of the battery from the EV charging station (EVCS) using the Smart Grid communications network. The message exchanges during an EV charging session are depicted in Figure 2.5 (based on [25]).

As shown in Figure 2.5, once the vehicle has been plugged in, the EVCS sends an initial charging request to the EVCC, specifying the battery parameters (e.g., battery size, rated

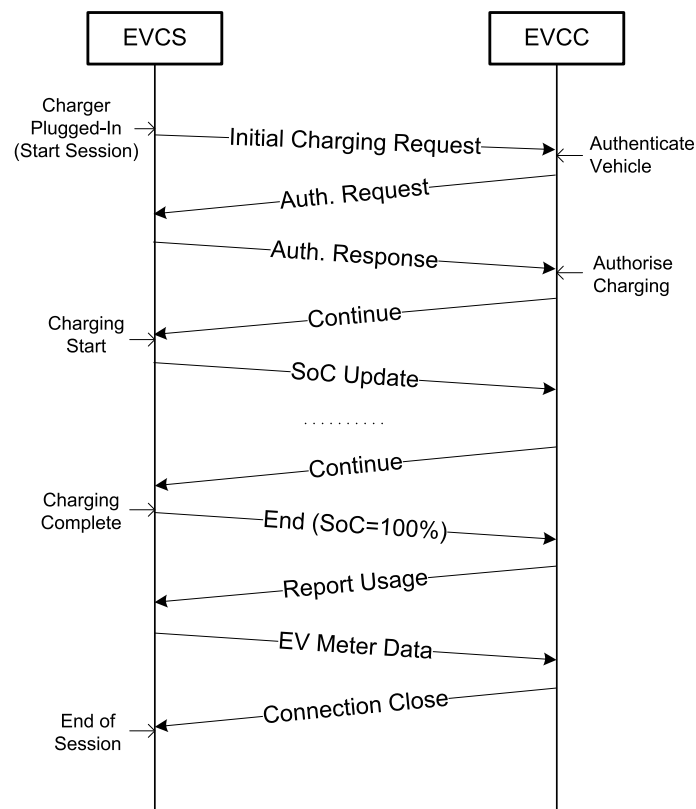


Figure 2.5: Application message exchanges during a smart EV charging session.

capacity, SoC) and the charger properties (e.g., voltage, current, maximum charging rate, charger efficiency). The EVCC then authenticates the vehicle, authorises the charging request, and starts the charging control process through a series of ‘continue’ messages. Note that the instructions passed in the ‘continue’ message depend on the algorithm used for the EV load management purpose. For example, the algorithm might control the charging process either by controlling charging start times or by regulating charging rates, or both. Against each ‘continue’ message, the EVCS sends a ‘SoC update’ message to inform the EVCC about the current SoC battery level. The process is repeated until an ‘end’ message is transmitted, either by the EVCC (battery charged fully) or by the EV (plugs-out from the system). The ‘end’ message is followed by an energy consumption report and a ‘connection close’ message.

Communications Load

According to the application model in Figure 2.5, each EV charging session comprises a number of fixed messages for initialization, vehicle authentication, energy transfer authorisation, and metering purposes. In addition, the EV uses the SoC update message to periodically update its charging status. Thus, the communications load (in bytes) for the i^{th} EV charging session is given by

$$l_i = l_{\text{Fixed}} + l_{\text{Variable}} = l_{\text{Fixed}} + n_i * l_{\text{SoC}}, \quad (2.1)$$

where l_{Fixed} denotes the number of bytes for the fixed messages, l_{Variable} denotes the number of bytes for the variable messages, n denotes the number of SoC update messages, and l_{SoC} denotes the length of each SoC message in bytes. Note that the above load calculation considers the uplink (UL) communications only. The downlink (DL) communications load would be almost the same, as each message transaction is paired, that is, comprised of a UL and a DL component.

To calculate the total communications load (in bytes) of an EV charging application, let us assume that the charging requests follow a Poisson arrival process, with arrival rate λ , and that the charging duration (i.e., service time) is exponentially distributed with mean μ [26]. Thus, the traffic intensity ρ of the application is given by

$$\rho = \frac{\lambda}{\mu} \beta, \quad (2.2)$$

where β is the percentage of the battery to be charged, which depends on the SoC level [27].

Similar to [27], we can calculate the blocking probability of the system for M number of contending EVs using the Erlang B formula as

$$B = \frac{(\rho^M / M!)}{\sum_{i=0}^M (\rho^i / i!)} \quad (2.3)$$

Using (2.2) and (2.3), we can calculate the total energy requirement E of the system as

$$E = \rho(1 - B).$$

From the DR's perspective, EVs are assumed as a partially interruptible load, whose charging cycle can be reduced during peak periods. Instead of strict admission control (i.e., either 'reject' or 'accept'), the EVCC can allow the contending vehicles to be partially charged, based on their priority and fairness. For example, the EVCC may allow emergency vehicles (e.g., ambulance, police cars) to be charged fully, as well as provide prioritised charging to vehicles whose battery charge has dropped below a certain threshold.

To incorporate the DR effect in the charging process, let us consider a load regulation factor α , where $0 \leq \alpha \leq 1$, such that it denotes the degree of partial charging among the contending vehicles. The concept is similar to the 'activity factor' introduced in [27]. With this assumption, the energy budget of the system $\hat{E}$ is given by

$$\hat{E} = \alpha\rho(1 - B). \quad (2.4)$$

From (2.4), we see that the energy budget of the system can be tuned by the parameter α , given that the maximum number of allowed EVs remains fixed. The number of SoC messages for a particular vehicle i is given by

$$n_i = \frac{\alpha\beta_i T_i}{t_{SoC}}, \quad (2.5)$$

where T is the time required to reach full battery charge and t_{SoC} is the SoC reporting interval. Note that the SoC reporting interval depends on the load management algorithm and may vary between five to 30 minutes. Thus, the overall communications load of an EV charging system with M EVs can be calculated by using (2.1) and (2.5) as

$$L_{Total} = \sum_{i=0}^M l_i = M l_{Fixed} + \alpha \sum_{i=0}^M \frac{\beta T_i}{t_{SoC}} l_{SoC}. \quad (2.6)$$

From (2.6), we see that the total communications load of the EV charging system is tightly

coupled with the available energy of the system, which can be scaled by the load regulation factor α .

Key Issues

Note that the SoC update messages are very critical for the EV charging applications, as the charging controller relies on them to adjust the charging rates. Moreover, they are delay sensitive, since the charger may remain idle and energy may not be transferred (hence, wasted) until it receives charging instructions against a SoC update message (see Figure 2.5). Hence, a fast and reliable message transfer is a key requirement for EV charging applications. Therefore, the SoC update messages from different vehicles need to be scheduled in a way that distributes them over the entire scheduling interval of an EV charging system. Otherwise, if all vehicles send SoC update messages simultaneously, the underlying communications network may experience congestion.

2.2.4 Substation Automation

Substation automation refers to the monitoring, protection and control functionalities performed on transmission and distribution substations and feeder equipments. In the substation automation domain, the IEC 61850 and DNP3/IEEE 1815 are the most widely adopted standards [28, 29]. While the DNP3 (Distributed Network Protocol, version 3) standard only provides communication specifications for low-bandwidth monitoring and control operations, the IEC 61850 standard covers almost all aspects of substation automation, including real time high-bandwidth protection and control applications. Therefore, IEC 61850 is gradually becoming the dominant protocol in this field.

Traffic Types

The IEC 61850 standard is based on interoperable IEDs that interact with each other, either within a substation (e.g., protection signals to circuit breakers) or on feeders (e.g., automated reclosers and switches along a feeder responding to isolate a fault). The IEC 61850

protocol is designed to run over standard communication networks based on the Ethernet and IP standards. To differentiate among various applications, and to prioritise their traffic flows, the standard defines five types of communication services:

1. Abstract Communication Service Interface (ACSI)
2. Generic Object Oriented Substation Event (GOOSE)
3. Generic Substation Status Event (GSSE)
4. Sampled Measured Value (SMV)
5. Time Synchronization (TS)

The ACSI services include querying device status, setting parameters and reporting and logging. The GOOSE and GSSE services, together often called generic substation events (GSE), exchange event and status information (e.g., a binary change of state or an analogue value crossing the reporting threshold) in real time. The SMV services transfer sampled analogue signals and status information. The TS service broadcasts the system clock information to the IEDs, ensuring measurement accuracy.

Most communications traffic from the IEC 61850-based applications is strictly delay sensitive, acting as a lifeline for underlying protection and control systems. For example, each GOOSE message contains a time-to-live (TTL) field, which implies that the message will lose relevance if not delivered within the specified time. To ensure appropriate QoS, the IEC 61850 standard specifies seven message types, based on their transfer time requirements. They are listed in Table 2.1.

Intra-substation Communications

The IEC 61850 protocol was originally designed to support intra-substation communications only. Hence, the communications architecture comprises three hierarchical levels: station, bay and process, as shown in Figure 2.6. The process level includes: various switch-yard equipments of the substation, such as CT/PT, I/O devices, sensors and actuators as well as breaker, switch, transformer, and merging unit (MU) IEDs. The MU IEDs

Table 2.1: IEC 61850 Message Types and Transfer Time Requirements [28]

Msg. Type	Application	Services	Transfer Time (ms)
1A	Fast Message (Trip)	GOOSE, GSSE	3-100
1B	Fast Message (Other)		20-100
2	Medium Speed	ACSI	100
3	Low Speed		500
4	Raw Data	SMV	3-10
5	File Transfer	ACSI	>1000
6	Time Synchronization	TS	(Accuracy)

collect the analogue voltage and current signals from field CT and PT, converts them into digital format and then transmits to the P&C (Protection and Control) IEDs in the bay level. The station level contains the human to machine interface (HMI) devices and station controller computers. The standard also defines two separate Ethernet subnetworks (called ‘buses’) to facilitate QoS implementations. While the process bus handles the delay sensitive communication between P&C IEDs and switch-yard devices, including breaker and switch IEDs, the station bus handles communication among different bays and with the station controller. The station bus terminates in an Ethernet-gateway that connects the substation local area network (LAN) with the external networks, including other substations and the utility control centre.

Inter-substation Communications

Over the last decade, the use of IEC 61850-based GOOSE messaging has significantly improved cost and reliability factors of intra-substation protection schemes, by replacing the legacy point-to-point wiring with an Ethernet-based LAN(similar to the one in Figure 2.6) [30]. In recent years, the widespread availability and use of microprocessor-based protective relaying techniques has renewed the potential for a wide area monitoring, protection and control (WAMPC) system [31], especially in the context of a ubiquitous Smart Grid communications network [32, 33, 34, 35]. More recently, IP routing functionality has been incorporated in the GOOSE profile based on the IEC 61850-90-5 standard [36]. This allows the Ethernet-based GOOSE messages to be exchanged among different substations under

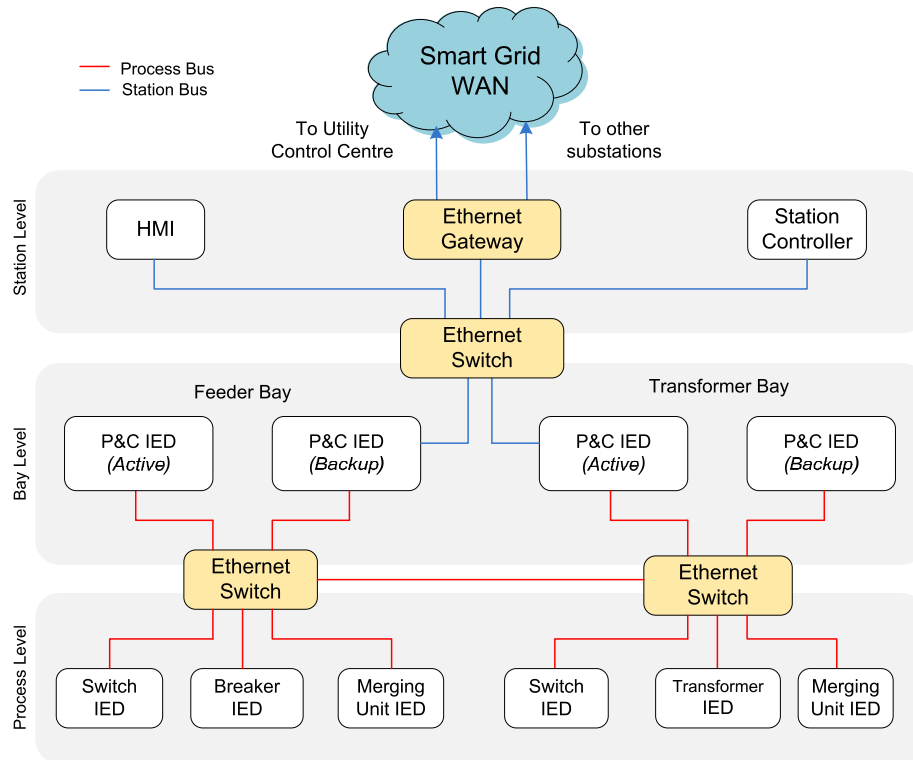


Figure 2.6: Architecture of an IEC 61850-based substation automation system.

different IP subnetworks using both unicast and multicast techniques. This important advancement in inter-substation communications paves the way to use an IP-based integrated Smart Grid communications network for advanced protection and control schemes in the power transmission and distribution grid.

Examples of these include distribution pilot protection schemes, such as directional comparison blocking (DCB) and permissive over-reaching transfer trip (POTT). They are widely considered as very effective methods of protecting networked distribution lines in urban and suburban areas [35]. They use high-speed, peer-to-peer communications among the relays to facilitate fast fault detection, isolation and restoration. Unlike other protection schemes, such as line current differential, the pilot relays need to communicate only when there is a fault, and a simple trip/block signal is sufficient to perform the relaying operation. This is quite advantageous in terms of communications load.

A ubiquitous Smart Grid communications network can significantly extend the reach of such protection schemes in a large number of substations over a vast geographic area. Sev-

eral works have proposed and investigated the use of digital communications in such applications [32, 33, 34, 35]. The use of wireless communication media such as microwave, narrow band radio, and spread spectrum radio for pilot protection schemes has been discussed [34]. In 2005, the IEEE Power System Relaying Committee (PSRC) released a detailed report outlining the use of spread spectrum radio for protective relaying applications. More recently in 2012, PSRC released another report on protective relaying applications using the forthcoming Smart Grid communications infrastructure [37]. The use of IEC 61850-based GOOSE messaging for pilot protection application has also been investigated [38, 39]. More recently, a distribution protection scheme based on GOOSE messaging and high-speed wireless communication systems such as WiMAX has been proposed [40].

The key communications requirement for such applications is to transfer the protection signals (e.g., trip, block) within a specified time. Another important requirement is to send these signals as quickly as possible. This is because operation of the associated switch/circuit breaker is often delayed by the data communication time, plus some margin [39]. Hence, the lower the communication delay, the faster the protection scheme will operate. Depending on the type of protection scheme, the total operating time may vary from two to six cycles for a distribution line, ignoring the breaker operating time [35]. Considering the operating time of high-speed digital relays (i.e., less than 5 ms), such a protection signal has a delay budget of around two cycles for the communications network.

2.2.5 Wide Area Measurement

WAMS refers to an advanced sensing and measurement system that continuously monitors the health of the power grid. In a WAMS, the system state and power quality information are obtained from the state measurement modules embedded in the PMUs. Unlike conventional measurement systems, the PMUs provide accurate system state measurements in real time by using GPS to provide a time stamp for each measurement [41]. Due to the precise synchronisation of the measurements, the utility control centre can obtain high resolution phase information, which enables them to initiate a proper response within seconds, to protect a whole WAN from black-out events [42]. The PMUs were traditionally installed

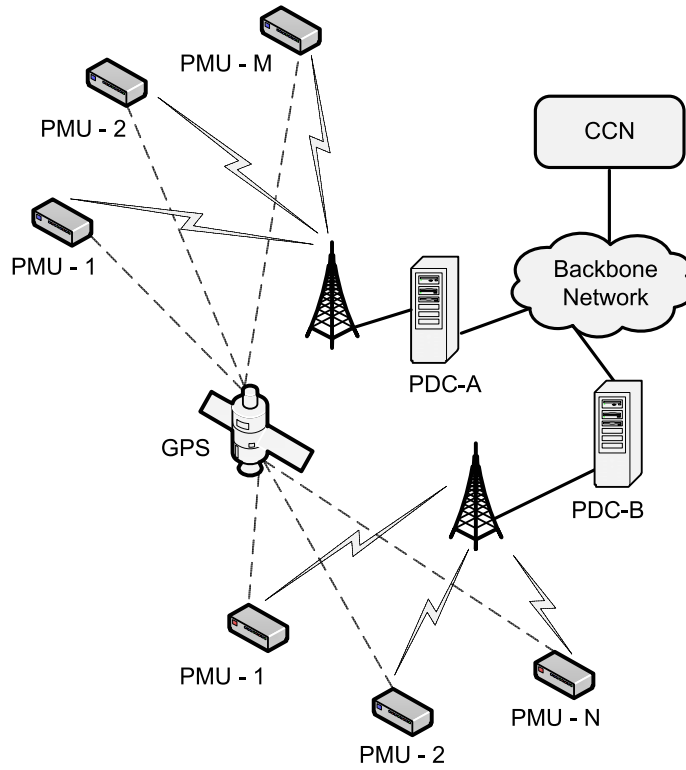


Figure 2.7: Communications architecture of a wide area measurement system in the Smart Grid.

on generation and transmission domains (HV level) because of the conventional top-down direction of the power flow. However, in a Smart Grid, the PMUs are also expected to be deployed at the distribution domain (medium and low voltage levels) to enable real time monitoring of the overall power system [43].

A complete WAMS system comprises hundreds of PMUs deployed at various locations in the national or regional electrical grid. The PMUs are locally grouped, and the measurements from them are first collected by a phasor data concentrator (PDC) via the local communication network, as shown in Figure 2.7. The data is then routed to a central control network (CCN) located at the utility's core network, via a backhaul network [44]. Here, the key challenge for a Smart Grid communications network is to provide reliable communication paths between the PMUs and the PDC within a strict delay bound.

The IEEE C37.118.2–2011 standard provides data communication specifications for the PMUs [45]. According to the standard, each PMU data packet contains a fixed 16 byte

Table 2.2: Payload Structure of a Sample Phasor Data Frame

No.	Field	Size (bytes)	Comments
1	SYNC	2	Synchronisation byte
2	FRAME SIZE	2	Number of bytes in frame
3	ID CODE	2	PMU ID
4	SoC	4	Second of Century time stamp
5	FRASEC	4	Time fraction and quality flag
6	STAT	2	Bitmapped Flag
7	PHASORS	$4 \times 4^*$	No. of Phasors
8	FREQ	2^*	Frequency
9	DFREQ	2^*	Rate of change of frequency
10	ANALOG	$2 \times 4^{**}$	2 Analogue Data
11	DIGITAL	1×2	1 Digital data (16 bit field)
12	CHK	2	Cyclic Redundancy Checks

* 16-bit signed integer, ** 32-bit floating-point.

header followed by a variable length message body depending on the number of phasor readings. Table 2.2 shows a typical payload structure of a sample synchrophasor data frame.

The overall communications load of a PMU depends on its reporting frequency f_s . The IEEE C37.118.2-2011 standard specifies reporting frequency of 10, 25 Hz and 10, 12, 15, 20, 30 Hz for a 50 Hz and 60 Hz based power systems respectively. The maximum communication delay $\bar{d}_{PMU}$ between the PMU and the local PDC should be such that a measurement is received before the next one is available. Thus, it can be expressed as

$$\bar{d}_{PMU} = \frac{1}{f_s} \quad (2.7)$$

2.2.6 Microgrids

Microgrid is a small local electric power system having one or more DER units and loads. The DERs are small sources of power generation and/or storage connected to the distribution grid. A DER can be from both renewable and non-renewable sources, and can either

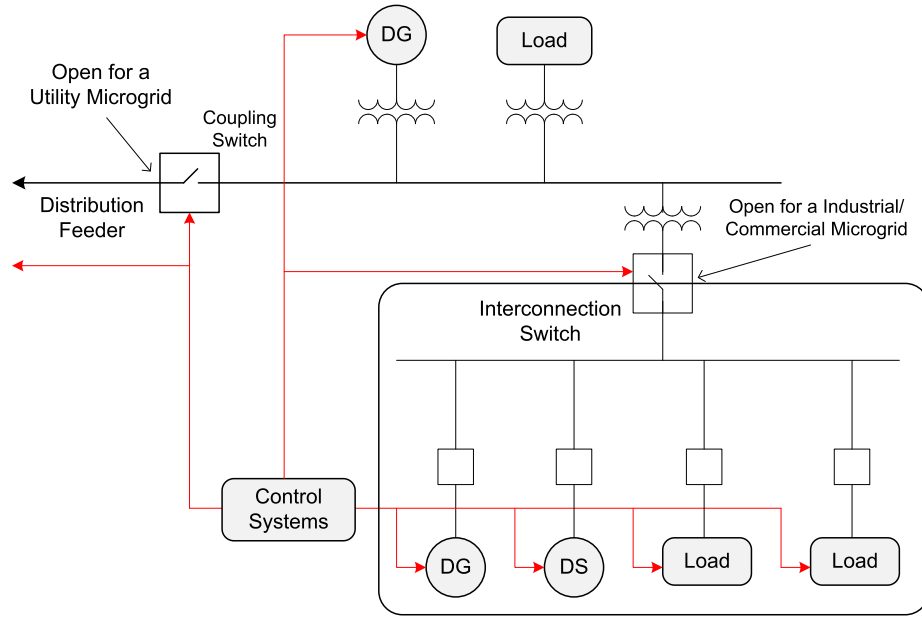


Figure 2.8: Diagram of a networked microgrid [46].

be a distributed generation (DG) source or distributed storage (DS) source, or both. Examples of DER include solar panels, wind turbines, combustion turbines, fuel cells and battery storage systems.

During normal operation, a microgrid is connected to the grid and operates in a synchronised mode. However, in case of any fault or maintenance event, it can operate autonomously in an island mode and is capable of supporting its own load [46]. Microgrids can be two types, based on their purpose: utility microgrids that serve parts of the utility; and industrial/commercial microgrids that only serve customer facilities, as shown in Figure 2.8. In a microgrid, loads and energy sources can be disconnected from and reconnected to the main grid with minimal disruption to the local loads.

Communications Requirements

The IEC 61850-7-420 extension incorporates different parts of DER/microgrid system in the IEC 61850 standard. For seamless operation, the DER/microgrid controller needs to exchange various operating parameters frequently with the utility control centre in real time. The communication requirements of a DER/microgrid controller include the follow-

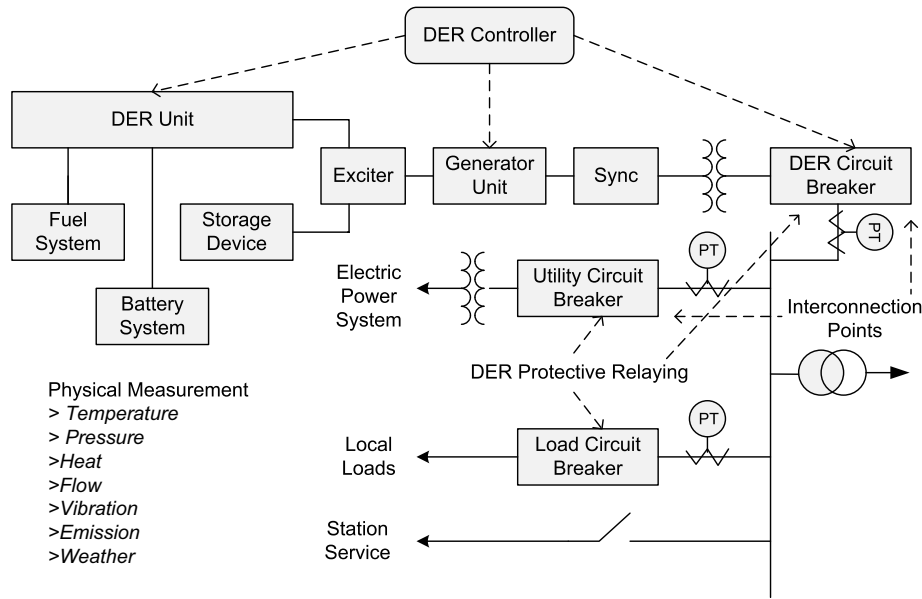


Figure 2.9: A DER/microgrid monitoring, protection and control system based on the IEC 61850-7-420 standard [47].

ing [47]:

- Management of the interconnection points between the DER units and the power systems they connect to, including local power systems, switches and circuit breakers, and protection;
- Monitoring and controlling the DER units and the associated the energy conversion systems, such as reciprocating engines (e.g., diesel engines), fuel cells, photovoltaic systems, and combined heat and power systems;
- Monitoring and controlling the individual generators, excitation systems, inverters/converters and auxiliary systems, such as interval meters, fuel systems and batteries; and
- Monitoring the physical characteristics of equipment, such as temperature, pressure, heat, vibration, flow, emissions and meteorological information.

A representative model of a DER/microgrid monitoring, protection, and control system based on the IEC 61850-7-420 standard is illustrated in Figure 2.9.

Protection Schemes

Of the above functionalities, the most challenging are the microgrid protection schemes. This is because the microgrids have very dynamic behaviour. A DER can act both as a source and a load (e.g., solar-based DERs). Consequently, the associated fault currents may no longer flow in one direction [48]. In such a system, differential protection schemes are more suitable as they are able to operate without prior knowledge of the direction and level of the fault current directions [49, 48].

The differential protection schemes, such as the line current differential protection (LCDP) scheme, are well-regarded for their high selectivity and sensitivity with a very low configuration complexity[50]. An LCDP relay works by continuously measuring its local line-current and comparing it to that of a remote terminal to detect a vector difference in current. If the difference exceeds a certain threshold, a fault is detected, and the relay operates. Communications networks play a vital role in such a scheme, as the local and remote line terminals must exchange their current elements in real time to perform the differential calculation.

In most modern LCDP schemes, the current elements are exchanged either as digitised samples of the analogue current or as current phasors with magnitude and phase angle [49]. If current samples are used, synchronisation is typically provided by measuring the round-trip delay of the communications channel, and then shifting the phase angle of the measured current accordingly [51]. In contrast, the phasor measurements are time-stamped using a common timing reference, such as GPS, that eliminates the requirement for a channel-based synchronisation technique [41]. Moreover, phasor data communication is regulated by international standards such as IEEE C37.118.2-2011 and IEC 61850-90-5; these allow the use of commercially available IP and Ethernet-based networks. Hence, phasor-based LCDP schemes are expected to be the most prevalent in the next generation Smart Grid.

From a communications network's perspective, the key requirement is to transfer the real time line current measurements (i.e., current phasors) reliably within a specific delay bound. The overall communications load from a differential relay depends on two factors: the total size of the measurement packet, and the rate of measurement. Although a higher measure-

ment rate may allow a faster response, it comes at the cost of a higher communications load. In contrast, while a HV transmission feeder requires a high measurement rate (e.g., 2–4 measurements per cycle) the requirements can be slightly relaxed for a distribution feeder, considering the relative effects of an outage.

2.3 Performance Requirements & Key Research Challenges

In the previous section, several major Smart Grid applications and their basic networking architectures and traffic characteristics were discussed. It is clearly evident from those discussions that a potential Smart Grid communications network needs to support a wide range of traffic sources with significantly varying QoS requirements. In this section, we summarise the key performance requirements of these applications, and highlight some key research challenges in relation to them.

In the case of multimedia applications, while it is important to ensure proper QoS, an even more vital issue is how the end-user perceives and experiences the service, known as quality of experience (QoE). Although it is a subjective measure, in the end it determines how satisfied the user is [52]. For example, a user can tolerate a certain amount of performance degradation during a voice call, while still having a fair degree of QoE at the end of that session. Hence, radio resource scheduling techniques for these applications try to optimise the throughput and/or fairness of the overall session, instead of a single data packet.

In contrast, session duration for most M2M applications in the Smart Grid is typically brief and involves only a handful of message exchanges. However, many require highly reliable message delivery within a strict delay bound, acting as the triggering points for the underlying protection and control systems. Consequently, the performance of these applications has to be determined in objective terms (e.g., message delivery success rate) against a set of pre-defined QoS attributes, such as delay and packet loss for each individual packet. For such applications, the key challenge for a Smart Grid communications network is to efficiently allocate radio resources among different classes of traffic so that their end-to-end QoS requirements are met.

Based on these discussions in the previous section, we can classify the key Smart Grid M2M applications/traffic sources into the following six generic categories:

- Protection traffic that can operate both at the substation level and in distribution systems (e.g., microgrids),
- Control traffic offering services to substations, high and low voltage distribution networks,
- Monitoring applications over a large geographic area covering the electricity distribution network,
- Metering and billing applications, mostly servicing residential and industrial premises through smart meters, and
- Demand management traffic serving different applications such as EV charging, energy storage systems, scheduling renewable energy sources.

Data traffic for the above applications can be further classified as periodic/deterministic, semi-periodic, random and event-based. Table 2.3 lists the traffic profiles for these classes for a typical Smart Grid communications network. It also indicates the range of delay variability of the Smart Grid applications that needs to be covered by a multi-service communications network.

Table 2.3: List of Traffic Profiles in a Typical Smart Grid Communications Network

Profile	Characteristics	Example Applications	Priority	Latency
1	Periodic, Delay Sensitive	Synchrophasor reporting	High	40-100 ms
2	Periodic, Delay Tolerant	Meter/ Sensor report	Low	1-60 sec
3	Semi-Periodic, Delay Sensitive	Distribution Automation	Medium	<1 sec
4	Semi-Periodic, Delay Tolerant	Demand Management	Medium	1-5 sec
5	Event-based, Delay Sensitive	Trip/Block Signal	High	<50 ms
6	Event-based, Delay Tolerant	Event/Alarm Reporting	Medium	< 1 sec

Note that in conventional communication networks, well defined resource allocation and QoS models exist to serve traffic, such as voice, data, constant bit rate (CBR) and variable bit rate (VBR) coded video. However, besides periodic and semi-periodic traffic, event-based traffic sources are hardly seen in these applications, other than some specialised

industrial sensor networks. Delay requirements in such events could be very short, as low as tens of milliseconds. For example, delay requirement for a pilot protection event is only two cycles, such as 40 ms for a 50 Hz power system. Other event-based traffic, such as the ‘last gasp’ alarms, may also have a stringent delay requirement: the alarm needs to be detected and transmitted with the shortest possible delay due to possible energy constraints after a power outage.

2.4 Chapter Summary

This chapter began by studying the communications architecture of the Smart Grid, followed by a comprehensive review of several key smart grid applications and their communication requirements. Section 2.3 highlighted some key research challenges pertinent to M2M applications in the Smart Grid, and thus, outlined the research focus of this thesis. The discussions undertaken in this chapter serve as a precursor to the subsequent analyses and solutions in the remainder of this thesis.

Chapter 3

Traffic Modelling and Performance Analysis

It is widely acknowledged that a 4G broadband wireless technology, in particular WiMAX, is well suited to the requirements of a Smart Grid communications network, thanks to its extended coverage, low latency, high throughput and advanced QoS functionalities. However, it is important to remember that network access and RRM techniques in a conventional WiMAX network were developed to support various multimedia applications, such as voice, video and web browsing. In contrast, the performance requirements of the Smart Grid applications are quite different (as discussed in Section 2.3). Hence, it is important to undertake a detailed performance analysis to investigate whether a conventional WiMAX network can meet these requirements.

In the previous chapter, we reviewed several major Smart Grid M2M applications and summarised their key traffic and performance requirements. However, to limit the scope of this work, we will select only three key Smart Grid traffic sources for this study: AMI/smart meters, synchrophasors/PMUs and protective relays. The reason behind this selection is that each of these traffic sources represents a distinct traffic class within a Smart Grid communications network, as described in Section 2.3. For example, AMI communications represents semi-periodic traffic from the Smart Grid NAN and FAN applications, such as AMR, DR, EV and microgrid management; synchrophasor communication represents

periodic traffic from the WAMS applications and microgrid protection schemes, such as LCDP; protective relaying applications represent event-based traffic from the substation protection and control applications, such as distribution pilot protection.

For AMI communications, we first estimate the number of smart meters within a typical suburban WiMAX cell, and then conduct a simulation study to investigate its performance under regular and fault scenarios. For synchrophasor communications, we first develop the basic traffic models and map them with existing radio resource scheduling services of WiMAX; we then conduct a capacity analysis followed by a simulation study to examine its communications load and delay performance. For protective relaying, we first develop a traffic model considering the distribution pilot protection application described in Section 2.2, map it with the existing scheduling services of WiMAX, and then perform a simulation study to examine its performance in terms of message transmission time and message delivery success rate.

The simulation models are developed using the well-known discrete event simulator OPNET. OPNET provides a specialised WiMAX model, which is widely used by vendors and network operators for planning and designing WiMAX networks and devices [53]. For this study, the OPNET's built-in WiMAX model is re-coded to incorporate the cyber-physical dynamics of the selected Smart Grid entities. For example, the built-in WiMAX subscriber station (SS) node model was embedded with smart meter, PMUs, EVs and protective relay functionalities by augmenting it with relevant customised application modules. More details regarding the simulation models are included in Appendix A.

The rest of the chapter is outlined as follows. Section 3.1 gives a brief overview of the key WiMAX tenets relevant to the modelling and analysis performed in this chapter. Section 3.2, 3.3 and 3.4 develop traffic models, conduct simulation studies, and describe their key outcomes. Finally, Section 3.5 concludes this chapter ¹.

¹The contents of this chapter has been published in four conference papers: three of these in proceedings of the *IEEE International Conference on Smart Grid Communications (SmartGridComm)* in Nov 2012/2013 [54, 55, 56], and another in proceedings of the *IEEE International Conference on Communications (ICC) workshop* in June 2013 [57].

3.1 Key WiMAX Tenets

In this section, we familiarise ourselves with some key WiMAX tenets essential for traffic modelling, QoS mapping, and performance analysis activities conducted in the remainder of this chapter.

3.1.1 Frame Structure

In an orthogonal frequency division multiple access (OFDMA) based system such as WiMAX, radio resources are allocated both in time and in frequency domains. The minimum possible data allocation unit is called an OFDM symbol, which is comprised of a data subcarrier and an OFDM symbol time. Typically, data is allocated using a group of contiguous OFDM symbols called the subchannels. Thus, each allocation can be visualised as rectangle with the number of OFDM symbols in the vertical axis and the number of subchannels in the horizontal axis. Figure 3.1 shows the frame structure of a time division duplex (TDD) based WiMAX system. Each frame is divided into two subframes: the DL subframe and the UL subframe, separated by a small guard-time that allows the BS to switch from transmit to receive mode.

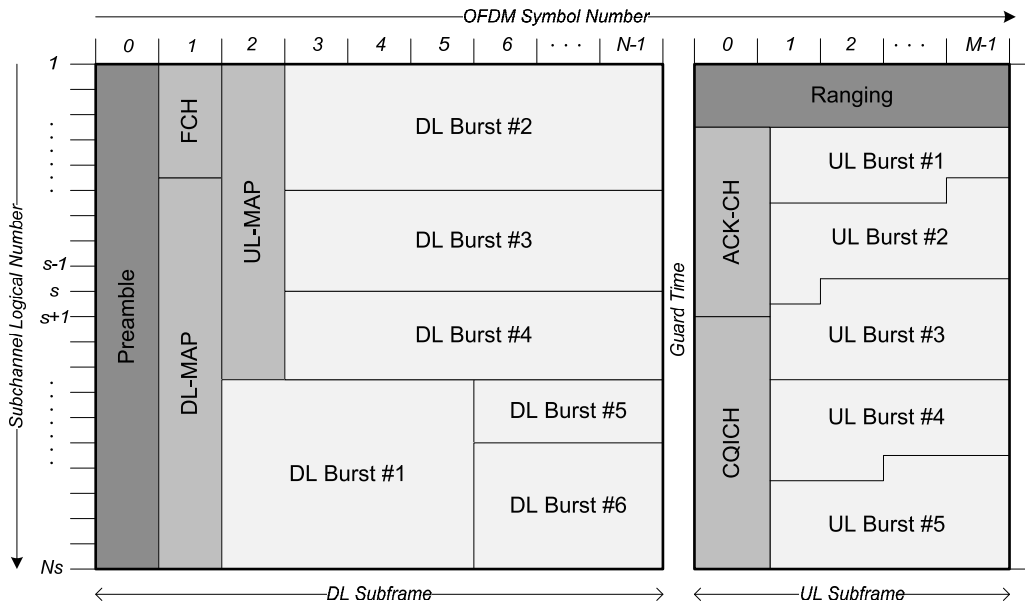


Figure 3.1: The frame structure of a OFDMA/TDD WiMAX System [58].

As seen from Figure 3.1, the DL subframe is comprised of a preamble followed by a frame control header (FCH), downlink medium access protocol (DL-MAP) messages, uplink medium access protocol (UL-MAP) messages, and DL data bursts. On the other hand, the UL subframe contains the initial ranging (IR) and bandwidth request (BR) channels, channel quality indicator (CQI) and fast feedback channels, and UL data bursts. The preamble is used for initial synchronisation, cell/sector identification, and channel estimation purposes. The FCH contains the information regarding the size and position of the DL-MAP and UL-MAP messages. The DL and UL-MAP messages are used to specify OFDM resource allocations over the DL and UL subframes. The information elements (IEs) of these messages indicate the start time and the OFDMA channel details of an UL/DL data burst.

3.1.2 QoS Management

Resource allocation and QoS control in WiMAX is provided via service flows (SFs). An SF is a logical unidirectional flow of packets between the BS and the SS, based on a particular set of QoS attributes (e.g., throughput, latency and jitter), identified by a connection ID (CID). The QoS parameters associated with the service flow define the transmission ordering and scheduling on the air interface. Thus, the connection-oriented QoS can provide accurate control over the air interface.

According to the IEEE 802.16-2009 standard, any packet traversing through the medium access control (MAC) interface is associated with a SF according to a set of classification rules. Traffic classification and mapping from application packets onto SFs in WiMAX is done at the convergence sublayer (CS) (located between the MAC and the IP layer) based on protocol-specific packet matching criteria, such as source and destination IP addresses, source and destination port address, protocol, and differentiated services code point (DSCP). The mapping is done at the BS for DL and at the SS for UL directions, respectively. WiMAX supports both network-initiated and terminal-initiated service flow. While the network-initiated SF creation is mandatory, terminal-initiated SF creation is an optional capability [58]. SFs may be created, changed or deleted through a series of MAC management messages, referred to as DSX: DSA (dynamic service addition), DSC (dynamic

service change) and DSD (dynamic service deletion) messages.

For QoS management among SFs, the IEEE 802.16 standard has defined five scheduling services/ QoS classes, as listed in Table 3.1.

Table 3.1: Five QoS Scheduling Services of WiMAX

QoS Class	Traffic Characteristics	Example
Unsolicited Grant Service (UGS)	Periodic, fixed-size data packets	TDM voice
Real-time Polling Service (rtPS)	Real-time, periodic, variable-size data packets	Streaming audio/video
Extended Real-time Polling Service (ertPS)	Same as above with ON/OFF intervals	Voice over IP
Non Real-time Polling Service (nrtPS)	Delay tolerant applications	File transfer
Best Effort (BE)	Regular data packets	Web browsing

3.1.3 Radio Resource Scheduling

In WiMAX, the radio resource scheduler is located at the BS, enabling rapid response to traffic requirements and channel conditions. It determines how radio resources will be allocated among multiple SFs based on their QoS attributes. Details of the radio resource scheduling operation are not specified by the standard; instead, they are vendor-specific.

The radio resource scheduler at the BS provides scheduling services for both DL and UL traffic. In the DL, the BS scheduler knows the bandwidth requirements of each connection to make an efficient resource allocation. However, in the UL bandwidth requirement information is distributed among the SSs. Therefore, a BR mechanism is required by the SSs to inform the BS about their bandwidth needs. For UL bandwidth allocation, the IEEE 802.16 standard uses three main techniques: contention-based random access; contention free polling; and a reservation-based unsolicited grant service. They are described below.

1. Contention

Under the contention mechanism, a SS uses the CDMA based random access procedure to send a BR. Under this procedure, whenever a SS intends to transmit a packet, it waits for a random number of frames. This number is chosen uniformly at random from the set $\{0, 1, 2, \dots, W_0 - 1\}$, where W_0 denotes the initial back-off window. Once the back-off counter reaches zero, it selects a single BR random access code with equal probability from a bank of K codes, and transmits it onto the ranging channel. If the code is received successfully, the BS provides the SS with a small bandwidth grant using the CDMA allocation IE to send the BR information. Once the BR information has been received by the BS, it allocates the actual bandwidth grant to the SS, which is then used to send the data packet. A sequence diagram of the random access based BR procedure is depicted in Figure 3.2.

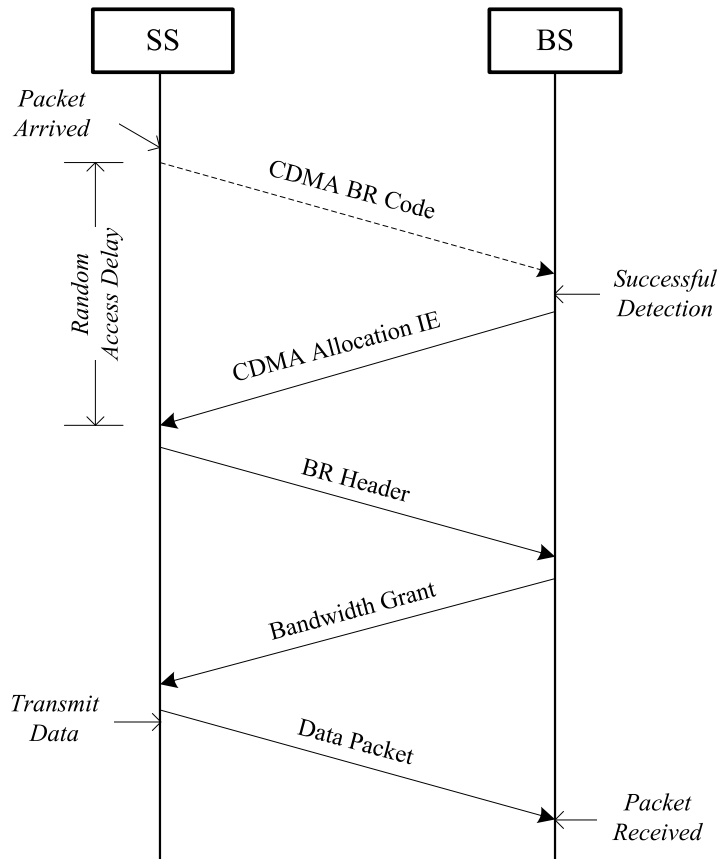


Figure 3.2: The WiMAX CRA procedure for uplink BR.

Since the ranging channel has no carrier sensing mechanism, the SS can not immediately

determine whether its code transmission was successful or not. Therefore, it waits for a number of frames specified by the T_{16} timer. The SS considers its CDMA code lost (either due to collision or failed detection) if no CDMA allocation has been received within this period. To resolve contention, WiMAX uses the truncated binary exponential back-off (TBEB) algorithm. Under the algorithm, the SS increases its current back-off window by a factor of two and repeats the code transmission process. Thus, after the i^{th} attempt the back-off window becomes $W_i = 2^i W_0$. The doubling of the back-off window continues until it reaches a maximum value denoted as W_f such that $W_f = 2^f W_0$. The retransmission continues with the back-off window size W_f until the maximum retry limit m has been reached, where $m \geq f$.

Thus, the overall delay under the contention mechanism is comprised of three components: i) time incurred in the CDMA contention process; ii) time required to request bandwidth and receive grant; and iii) time required to send the actual data packet. It can be expressed as

$$d_{\text{Contention}} = d_{\text{RandomAccess}} + d_{\text{Grant}} + d_{Tx}, \quad (3.1)$$

where $d_{\text{RandomAccess}}$ is the delay incurred in the random access based BR procedure, d_{Grant} is the time required to obtain the actual data grant, and d_{Tx} is the time required to send the data packet. Note that the latter two delay components remain fairly constant, and require roughly two WiMAX frames for a moderately loaded WiMAX network. Hence, the overall delay is mainly determined by $d_{\text{RandomAccess}}$, which further depends on the random access based BR load of the system and the contention resolution parameters of the TBEB algorithm.

2. Polling

When polling is used, the BS maintains a list of registered SSs and polls them according to a pre-defined schedule. A SS is only allowed to transmit the BR message after it has been polled. Thus, the overall delay under the polling mechanism is same as (3.1), except the

contention delay is replaced by the polling delay. It is given by

$$d_{polling} = d_{poll} + d_{Grant} + d_{Tx}, \quad (3.2)$$

where d_{poll} is the delay between subsequent polls, which depends on the arrival time of the first poll. This is because the subsequent poll arrival times can be calculated by adding the polling interval Δt with the previous poll arrival time. Typically, the time of first poll is a random value uniformly chosen by the BS from the polling interval Δt . Hence, the average poll delay can be given by

$$\bar{d}_{poll} = \mathbf{Uniform}(0, \Delta t) = \frac{1}{2}\Delta t. \quad (3.3)$$

The other delay component, d_{Tx} in (3.2), depends on the number of frames across which the bandwidth grant from the BS is distributed. The larger the packet, the more MAC frames are required by the BS to spread the bandwidth grant; therefore, the more delay is incurred during the packet transmission process.

3. Unsolicited Grant

Under the unsolicited grant mechanism, the BS allocates fixed-size data grants on a periodic basis without any explicit BR procedure. The bandwidth allocation process under the UGS scheduling is regulated by the two fixed parameters: maximum sustainable traffic rate (MSTR) and maximum latency (ML). While MSTR defines the peak data rate required by the application, maximum latency determines the unsolicited grant interval T_{ugs} , which is the time interval between two successive UGS grants. The MSTR is calculated by the following formula:

$$MSTR \text{ (in bps)} = \frac{B}{T_{ugs}}, \quad (3.4)$$

where B is the fixed packet size in bits.

The delay components under the UGS scheduling is comprised of just two components:

i) time required by the BS to send bandwidth grants; and ii) time required by the SS to

transmit data to the BS. To meet the delay requirement, T_{ugs} should be less than the ML of the application.

3.2 AMI Communications over WiMAX

AMI applications, such as AMR, DR and EV, use real time two-way communications between the utility and the customers to ensure optimal energy usage, enabling customer participation to enhance the overall utilisation and efficiency of the grid. They are considered the most basic applications of the Smart Grid. To enable AMI communications, a WiMAX network has to support a very large number of M2M devices per cell, as it has to collect measurements from each residential and commercial meter within its coverage area. While the AMI applications typically require a very low data rate, they have to perform frequent network entry (and re-entry) to transmit small data bursts, yielding a high random access load in the WiMAX BR ranging channels [8]. Moreover, during a grid event such as a power outage, several M2M devices may try to access the network simultaneously to send an alarm/event report. This may result in a sudden overload in the ranging channels, which may adversely affect the performance of the overall WiMAX network.

In this section we analyse the performance of AMI communications over a WiMAX network using a generic application model. In particular, we concentrate on the random access performance of the WiMAX network under regular and fault scenarios.

3.2.1 Estimated Number of Smart Meters

To conduct the performance analysis, we first need to estimate the possible number of smart meters in a typical WiMAX cell. We perform this estimation in the following two steps.

a) Coverage Analysis: The aim of the coverage analysis is to predict the maximum cell radius of a WiMAX BS under a given set of operating parameters. As most M2M devices in an AMI network will remain stationary, the coverage area is projected according to the

Erceg path-loss model, which is recommended by the IEEE 802.16 group for fixed broadband applications in suburban and rural area deployments [59]. The model is applicable from 1800 to 2700 MHz and defines three terrain types: A, B and C [60]. The median path-loss under the Erceg model is given by

$$PL_{dB} = 20 \log\left(\frac{4\pi d_0}{\lambda}\right) + 10 \left(a - b h_b + \frac{c}{h_b} \right) \log\left(\frac{d}{d_0}\right) + 6 \log 10 \left(\frac{f}{2000} \right) - X \log\left(\frac{h_s}{2}\right), \quad (3.5)$$

where $d_0 = 100m$, d is the distance between the BS and the SS antennae, such that $d > d_0$, λ is the wavelength in metres, f is the carrier frequency in MHz, h_b is the BS antenna height in metres, h_s is the SS antenna height in metres; the constant values a , b , c , and X can be found in [60].

For the coverage analysis, two types of WiMAX SSs are considered: i) fixed indoor SS, representing WiMAX-enabled smart meters and field devices; and ii) fixed outdoor SS, representing WiMAX-enabled DAPs, described in Section 2.1 [16]. The corresponding WiMAX BS and SS parameters are listed in Table 3.2. The parameters are chosen based on the deployment guidelines specified by the NIST PAP2 report [16].

Table 3.2: WiMAX BS and SS Parameters for AMI Coverage Analysis

Parameters	BS	SS (Fixed Indoor)	SS (Fixed Outdoor)
Antenna Height	30 m	2.5 m	10 m
Antenna Gain	17 dBi	6 dBi	14 dBi
Transmit Power	10 W	200 mW	300 mW
Noise Figure	7 dB	4 dB	4 dB

Note that in Table 3.2, the fixed outdoor SSs have a relatively high transmit power, as well as greater antenna height and gain; this is because they are expected to be mounted on high utility structures, such as transmission poles and substation structures. Moreover, they are free from building penetration losses as they are located outside the building.

The link-budget equation in terms of the maximum allowable path-loss is given by

$$PL_{Max} = P_{EIRP} + G_{Rx} - R_{SS} - \text{PenLoss}, \quad (3.6)$$

where P_{EIRP} is the effective isotropically radiated power (EIRP) level in dBm, which is given by the sum of transmit power and transmit antenna gain, minus the coupling losses; G_{Rx} is the receiver gain in dBi; PenLoss is the building penetration loss in dB; and R_{ss} is the receiver sensitivity level in dB, which is derived according to the equation specified in the IEEE 802.16-2009 standard:

$$R_{SS} = -114 + SNR_{Rx} + 10\log R + 10\log \left(\frac{F_S \times N_{Used} \times 10^{-6}}{N_{FFT}} \right) + \text{ImpLoss} + NF, \quad (3.7)$$

where,

SNR_{Rx} is the receiver SNR requirement,

R is the repetition factor, typically 1 for data traffic,

F_S is the OFDM sampling frequency in Hz,

ImpLoss is the implementation loss, which includes non-ideal receiver effects such as channel estimation errors, tracking errors, quantisation errors, and phase noise [58],

NF is the receiver noise figure, referenced to the antenna port,

N_{FFT} is the total number of OFDM subcarriers or the fast Fourier transform (FFT) size, and

N_{used} is the number of used subcarriers, i.e. data and pilot subcarriers.

Note that in the UL, a SS can only transmit over a limited number of subchannels. Thus, the transmit power is spread over a smaller subset of subcarriers instead of the entire UL subframe. This yields a sub-channelisation Gain, which is given by

$$G_{sch} = 10\log \left(\frac{N_{sch}^{UL}[\text{used}]}{N_{sch}^{UL}} \right), \quad (3.8)$$

where N_{sch}^{UL} is the number of total subchannels in the UL and $N_{sch}^{UL}[\text{used}]$ is the number of used subchannels for UL data transmission.

The OFDMA parameters used in the analysis are listed in Table 3.3.

The coverage analysis is summarised in Table 3.4 for both the indoor and the outdoor SSs. Here, the maximum cell radius is projected based on a SNR requirement of 5 dB, which

Table 3.3: OFDMA Parameters for the Coverage Analysis

Parameters	Notation	Values
Base Frequency	f_{sys}	2.3 GHz
Channel Bandwidth	BW	5 MHz
FFT Size	N_{FFT}	512
Sampling Factor	n	28/25
Sampling Frequency	$F_s = n * BW$	5.6 MHz
Subcarrier Spacing	$\Delta f = F_s / N_{FFT}$	10.94 KHz
Useful symbol time	$T_b = 1 / \Delta f$	91.4 μs
Cyclic Prefix Time	T_g	$T_b / 8$
OFDM Symbol Time	$T_s = T_b + T_g$	100.8 μs
Frame Duration	T_f	5 ms
No. of used Subcarriers	N_{used}	421
No. of DL Data Subcarrier	$N_{sc}^{DL}(\text{PUSC})$	360
No. of UL Data Subcarrier	$N_{sc}^{UL}(\text{PUSC})$	272
No. of UL Subchannels	N_{sch}^{UL}	15
No. of DL Subchannels	N_{sch}^{DL}	15
UL:DL Ratio	r_{fame}	1:1

is the minimum SNR requirement for the data traffic as specified in the IEEE 802.16-2009 standard, i.e. quadrature phase shift keying (QPSK) with half-rate forward error correction (FEC) coding [58].

From the results, we see that the cell radius of a WiMAX BS can be quadrupled if only the outdoor SSs are deployed (as DAPs). However, since the focus of this chapter is to analyze the performance of a standalone WiMAX network, we consider the case of indoor SSs only.

b) Estimation: Based on the projected cell radius R , the estimated number of smart meters under a given BS can be calculated as

$$N = \frac{\rho \pi R^2}{s}, \quad (3.9)$$

where ρ is the meter density per km^2 and s is the number of sectors under the BS.

Thus, considering a single sector BS with a meter density of 800 meters/ km^2 [16], the number of smart meters for a cell radius of 2.17 km is around 12,000.

Table 3.4: Summary of Coverage Analysis

Parameters	Fixed Indoor SS		Fixed Outdoor SS	
	DL	UL	DL	UL
Tx Power (dBm)	40	23	40	25
No. of Tx Antennae	2	1	2	1
Tx Antenna Gain (dBi)	17	6	17	14
Combination Gain (dB)	3	0	3	0
Coupling Loss (dB)	3	0	3	0
EIRP (dB)	57	29	57	39
Imp. Loss (dB)	2	2	2	2
Noise Figure (dB)	7	4	7	4
Required SNR (dB)	5	5	5	5
No. of Subchannels	15	17	15	17
Subch. Gain (dB)	0	-12.3	0	-12.3
Rx Sensitivity (dBm)	-93.37	-108.67	-93.37	-108.67
Rx Antenna Gain (dBi)	6	17	14	17
Rx Diversity Gain (dB)	3	3	3	3
System Gain (dB)	159.37	157.67	167.37	167.67
Interference Margin (dB)	2	2	2	2
Shadow Fade Margin (dB)	8.2	8.2	8.2	8.2
Fast Fade Margin (dB)	2	2	2	2
Building Pen. Loss (dB)	10	10	0	0
Link Margin (dB)	139.17	137.47	157.17	157.47
Cell Radius (km)	2.17		8.62	

3.2.2 Simulation Study

For the performance analysis, a simulation model is developed using the OPNET Modeller 16.0. The model is based on a single sector WiMAX BS serving 12,000 smart meters, as estimated in the previous subsection. The OFDMA parameters of the WiMAX network are set according to Table 3.3. Moreover, the random access parameters are listed in Table 3.5, which are set according to the recommendations made by the IEEE 802.16p working group [61].

The following two simulation trials are conducted based on the application scenarios described in [8].

1. A generic AMI application performing random access to transmit short data bursts

Table 3.5: WiMAX Simulation Parameters

Parameters	Values
No. of BR Ranging Channels	1 per frame
Initial Back-off Window (W_0)	2^2
Final Back-off Window (W_f)	2^{15}
Max. Retries (m)	16
No. of Detectable Codes /Channel	1
Probability of Misdetetection	0%
Contention Time-out (T_{16})	8 frames

(100 bytes) according to an exponentially distributed reporting cycle. The reporting interval is further varied from one to five minutes to examine its effect on the performance of the WiMAX network. While such a reporting interval is quite aggressive considering AMR applications only, we assume that the reporting interval also accounts for other delay-sensitive AMI applications, such as DR and EV (see Table 2.3) of the Smart Grid.

2. A group of smart meters are accessing the network asynchronously within a small interval (10 sec) to send the ‘last gasp’ alarm after a power outage [62]. The number of smart meters is varied between 200 to 2,000, with a step of 200, to examine its effect on the performance of the WiMAX network.

For the first simulation trial, the corresponding random access delay and access success rate are plotted in Figure 3.3. From the results, we see that when reporting intervals are smaller, i.e. less than three minutes, the access delay starts to increase exponentially and the corresponding access success rate drops. This is because at the smaller reporting intervals more access attempts are made, which increases the number of collisions in the ranging channel. Consequently, the SSs have to try multiple times to successfully transmit their random access codes to the BS. Figure 3.4 provides more insight into this scenario by plotting the cumulative distribution function (CDF) of the contention retries at different reporting intervals. Figure 3.4 shows that at the highest arrival rate, for instance, at a reporting interval of one minute, some devices reach their maximum retry limit (i.e. $m = 16$), and are discarded. As a result, the access success rate decreases, as shown in Figure

3.3. A more detailed analysis on the performance of the random access plane is provided in Chapter 4.

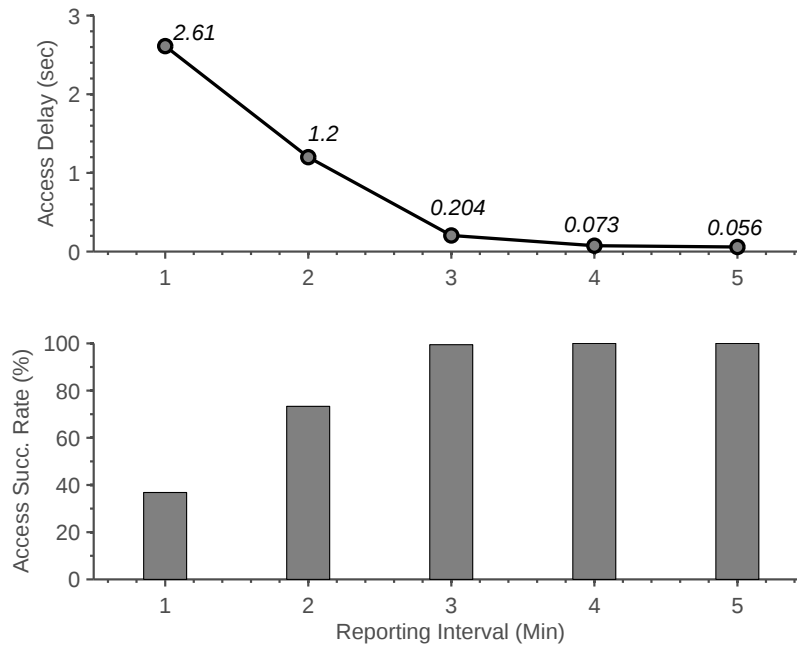


Figure 3.3: Mean access success rate and access delay of a WiMAX ranging channel under the generic AMI traffic model.

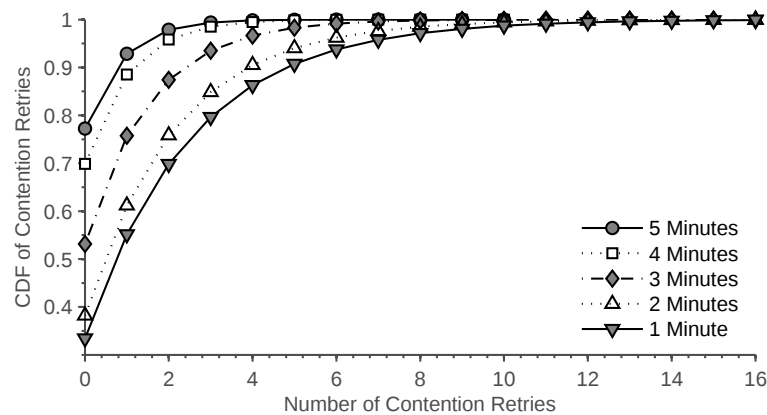


Figure 3.4: CDF of contention retries in the WiMAX ranging channel under the generic AMI traffic model.

For the second simulation trial, the corresponding alarm dispatch time and alarm success rate are shown in Figure 3.5. Here, we make a simplified assumption that a successful

random access attempt is sufficient to convey the power outage information. The results are similar to those obtained for the first simulation trial; the more the number of meters trying to access the network, the more the random access load, the less the number of successful alarms.

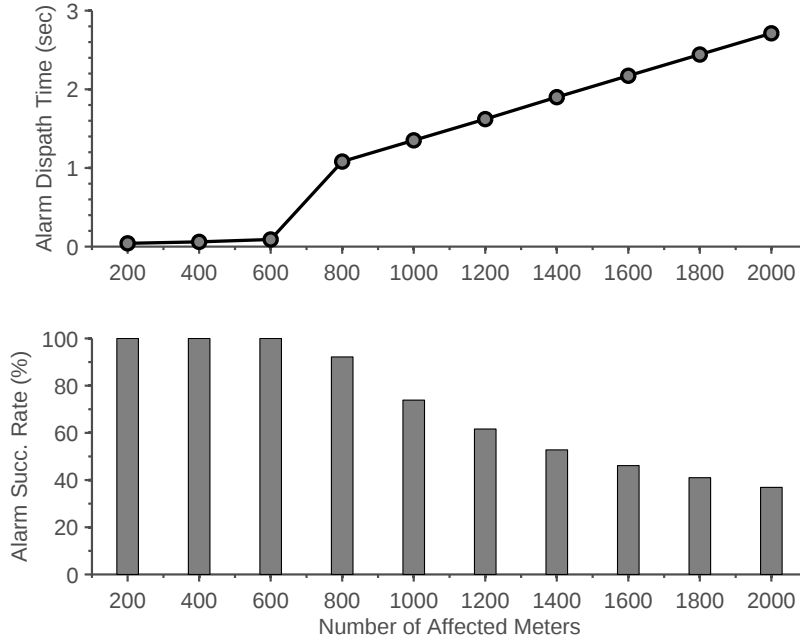


Figure 3.5: Alarm dispatch time and alarm success rate for the outage scenario.

Note that since the smart meters are expected to operate without battery backup and rely on charge stored in a capacitor to send this alarm, the delay budget is very stringent here, for instance, 100 to 250 ms. Thus, when a large number of meters are trying to access the network simultaneously, there is a possibility that none will be able to send an alarm.

3.2.3 Key Findings

From the above analyses it is clear that (apart from the high perennial random access load from the bursty AMI applications) when a large number of M2M devices try to access the network simultaneously after a grid event, the number of collisions can be very high in the ranging channel. This results in devices re-trying to access the network. In turn, this creates even more collisions and quickly reduces the random access success rate. Without

necessary mechanisms to handle such traffic, the entire ranging channel could become congested. In addition to prolonged access delay, a congested ranging channel can significantly increase the power consumption of the contending devices, due to multiple retries. Hence, it is quite imperative to enhance the performance of the WiMAX random access plane to efficiently support M2M devices in an AMI environment.

3.3 Synchrophasor Communications over WiMAX

In a typical WAMS, all PMUs transfer data directly to the nearest PDC (see Figure 2.7). On the other hand, in many wide-area protection applications, such as the phasor-based differential protection scheme described in Subsection 2.2.6, a PMU-enabled relay needs to exchange its phasor information with one or more neighbouring PMUs (or PDCs). Thus, in a Smart Grid communications network, PMU-based synchrophasor communication schemes can be classified into two major types: i) PMU to PDC; and ii) PMU to PMU. In the remainder of this section, we analyse the capacity and performance of these schemes over a conventional WiMAX network.

3.3.1 Traffic Modelling

To analyse the performance of the synchrophasor communication schemes, we first need to develop the necessary traffic models. The following two traffic models are considered:

i) PMU to PDC: This is the most common synchrophasor communication scheme, where the PMUs send periodic measurement frames to the PDC according to a pre-defined measurement cycle [45]. However, if all PMUs transmit data simultaneously, there will be a sudden burst of communications load, which may lead to congestion in the network. Hence, the PMU data transfer should be intelligently scheduled so that the communications load and the communications delay remain within an acceptable level. For this study, we define two basic application models, namely PULL and PUSH, as depicted in Figure 3.6.

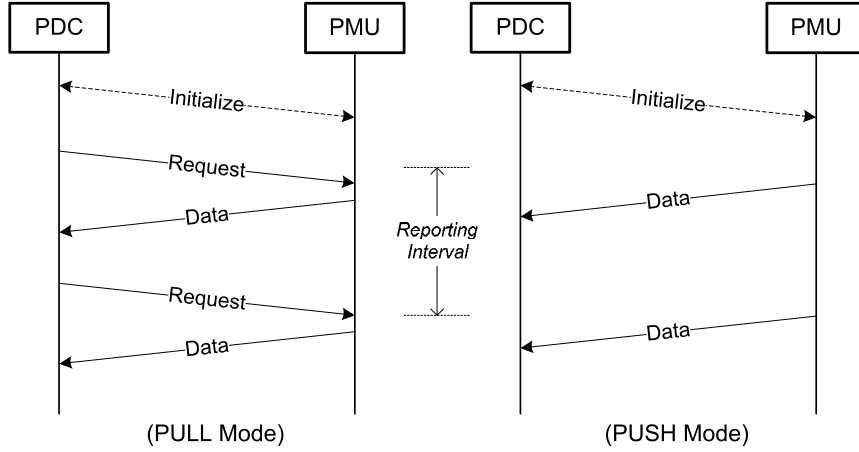


Figure 3.6: Application models for PMU to PDC communication.

Under the PULL mode, the PDC sends periodic requests to each of its member PMUs to fetch their individual measurements. Conversely, under the PUSH mode, the PMUs periodically send their measurements to the PDC without any explicit request. Thus, while the PULL mode offers a PDC to manage the data acquisition process (in coordination with the WiMAX BS) on a frame-by-frame basis, the PUSH mode is simple and requires less message exchanges.

For this study, we use the same PMU payload structure as shown in Table 2.2. Thus, considering the lower layer protocol overheads, such as user datagram protocol (UDP) (8 byte) and IP (20 byte) headers, WiMAX MAC header (6 byte), and cyclic redundancy check (CRC) bits (4 byte), the total MAC protocol data unit (MPDU) size becomes 86 bytes.

ii) PMU to PMU: These communication schemes are typically application-specific. For this study, we consider the phasor based LCDP scheme described in Subsection 2.2.6. Figure 3.7 shows a generic LCDP scheme with two line terminals. Both of the PMU-enabled relays at the end of the protected line continuously exchange their current phasors I_1 and I_2 via the communications network. Under normal circumstances, the sum of these current phasors at each relay should be zero. Should there be a fault, the relays will detect a differential current component $I_D = |I_1 + I_2|$; therefore, they will trip the associated circuit breaker. Note that in practice, the vector sum of the current phasors may not always be

zero due to measurement errors. Therefore, a restrained current I_R is used to compensate for such an error, which is typically a function of the individual current magnitudes of the relays, i.e. $I_R = f(|I_1|, |I_2|)$. In case $I_D > I_R$, the relay operates.

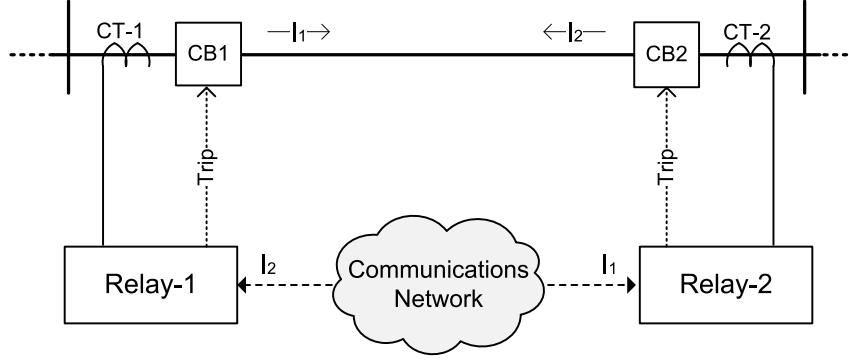


Figure 3.7: A generic LCDP scheme for a two-terminal feeder.

Figure 3.8 shows the sample MPDU structure of a PMU data packet. The measurement payload comprises 4 fields. The first two fields are used to identify the individual relay station and its operating zone, followed by a time-tag field and the current phasor readings (three phasors for a three phase line). The field sizes are set according to the IEEE C37.118.2 standard [45].

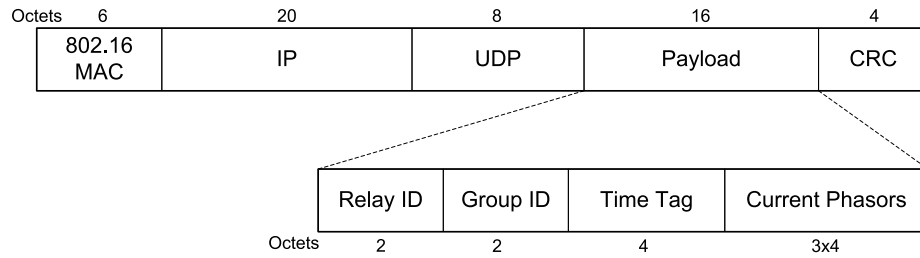


Figure 3.8: Sample MPDU structure of a PMU-enabled Differential Relay.

Although the LCDP scheme require peer-to-peer communication among the relays, in a cellular wireless environment like WiMAX, all the data packets must have to traverse via the BS. This eliminates the requirement of mesh connections (i.e., $n - 1$ links for each of the n PMUs) by providing a central point of communication. Here, the BS will forward the incoming data packets to the destination relays based on their destination IP addresses. Thus, each PMU data packet is associated with a UL component, as from the source relay to the BS; and a DL component, as from the BS to the destination relay.

3.3.2 Key RRM Issues

In this subsection, we describe the key RRM issues related to the synchrophasor communication schemes over a WiMAX network, and discuss the additional WiMAX features that can be used to support these schemes more efficiently.

a) Synchronisation: The PMUs have a precise synchronisation requirement to maintain the accuracy of their measurements. This accuracy is often measured in terms of total vector error (TVE), which estimates the difference between the theoretical actual phasor and the estimated phasor [63]. The IEEE C37.188.2 standard for synchrophasor measurement allows a TVE of one per cent that corresponds to a time error of $\pm 31 \mu\text{s}$ for a 50 Hz power system.

A typical TDD based WiMAX network requires precise synchronisation and timing to assure that the SSs can access their UL and DL time-slots without interfering with each other. Typically, a WiMAX BS receives synchronisation information either directly from a GPS receiver or from a master clock located in the IP backbone network using the IEEE1588 based precision time protocol (PTP). During the network entry procedure, a SS first synchronises itself with the DL preamble (see Figure 3.1) transmitted at the beginning of each WiMAX frame. However, synchronising with the DL preamble does not guarantee precise time synchronisation with the BS. This is because the SSs are placed at random locations within the BS's coverage area and their signal arrival times depend on their relative distance from the BS. Therefore, the next level of synchronisation is obtained by the ranging procedure, which adjusts each SS's timing-offset such that it appears co-located with the BS. The BS calculates the amount of timing offsets based on the round-trip delay between itself and the SS. The reference time tolerance specified in the IEEE 802.16-2009 standard is $\pm(T_b/32)/4$, where T_b is the OFDM symbol time ($\sim 91.4 \mu\text{s}$) [58]. This yields a timing error tolerance of $\pm 0.714 \mu\text{s}$, which is far below the PMU timing error margin of $\pm 31 \mu\text{s}$.

Thus, we can safely conclude that a WiMAX network can provide adequate synchronisation to the PMUs. This removes the need to install a separate GPS receiver for the relays, which may significantly reduce the deployment cost of a WAMS.

b) QoS Mapping: To support PMU communications over a WiMAX network, the corresponding traffic flows need to be mapped with the appropriate scheduling services of WiMAX. To facilitate radio resource sharing among different users, the IEEE 802.16 standard provides three key mechanisms: contention-based BE service, contention-free polling service (PS), and reservation-based UGS, as described in Subsection 3.1.3. Under the UGS, the BS periodically assigns fixed-size bandwidth grants to the SS; under the PS such as rtPS, the BS polls the SSs at fixed intervals to identify their current bandwidth requirements and then allocate grants accordingly.

For PMU to PDC communication, since the PULL mode requires the PDC to periodically request data from each PMU, we can map this application over both the UGS and rtPS classes. In that case, while the rtPS explicitly polls the PMUs to send their measurements, the UGS sends implicit requests via the periodic bandwidth grants. Conversely, in PUSH mode the PMUs try to send their measurements as soon as they are available. Hence, the BE scheduling mode is the most appropriate for this case. The delay components of the UGS and the rtPS services under the PMU communications are illustrated in Figure 3.9 and Figure 3.10, respectively. The delay components under the BE scheduling are the same as the rtPS, except the poll injection delay is replaced by the BR delay.

For PMU to PMU communication, we map the corresponding traffic flows with the UGS scheduling service only. This is because such information exchange emulates a time division multiplexing (TDM) connection between the peer PMUs/relays for which the UGS is the most appropriate scheduling service.

Note that the IEEE 802.16 standard provides mechanism to synchronise the UGS data grant and packet transmission time using the frame latency (FL) and the frame latency indication (FLI) fields embedded in a grant management subheader. Using these fields, the data transmission from the relays can be synchronised such that the relays from the i^{th} protection zone generates data at the x^{th} frame and the relays from the $(i + 1)^{th}$ zone generates data at the $(x + 1)^{th}$ frame. This ensures minimum UGS delay by minimising the time gap between the packet generation and the packet transmission time.

Moreover, WiMAX allows persistent scheduling technique that significantly reduces MAP

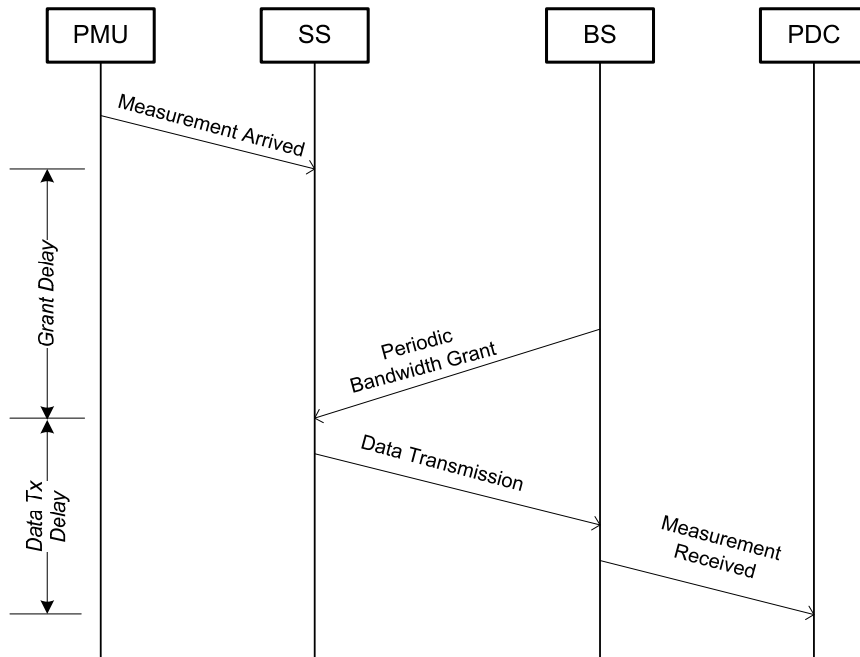


Figure 3.9: Delay components of PMU communications under UGS scheduling.

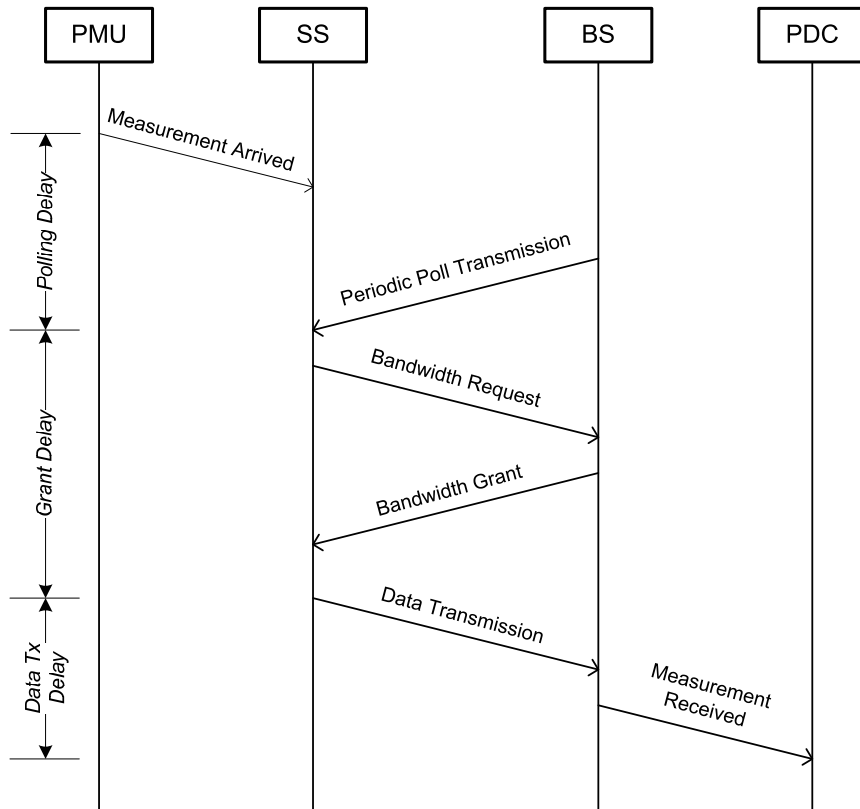


Figure 3.10: Delay components of PMU communications under rtPS scheduling.

signalling overhead for the UGS connections. Under persistent scheduling, the UL and DL burst information is sent once in a persistent MAP element and is not repeated unless any parameter associated with the connection changes [58].

c) Packet Loss & Retransmission: Another key issue for synchrophasor communications is the random packet-losses that occur in the multipath wireless channel due to random noise and fading. This is even more crucial for a PMU-based protection scheme, as the loss of a measurement packet may significantly affect the accuracy and speed of the relaying operation. To recover the lost packets, WiMAX allows both automatic repeat request (ARQ) and hybrid automatic repeat request (HARQ) retransmission schemes. However, since each PMU measurement has a small TTL period, it is difficult to allow more than one retransmission.

Between the ARQ and the HARQ schemes, the HARQ is particularly suitable for a differential protection scheme. This is because the ARQ relies on a feedback mechanism to detect a packet error and wait for a certain time-out parameter for the next retransmission opportunity. In contrast, the HARQ sends each data packet with FEC coding, and the receiver uses both retransmitted packet and packet received with errors to reconstruct the original packet which reduces the number of retransmissions. Moreover, the HARQ scheme with chase combining (CC) provides fast retransmission opportunities through one of the dedicated acknowledgment channels (ACK-CH) in the UL subframe (see Figure 3.1).

Note that the overall packet delay of a PMU measurement including the retransmission attempts should remain below the current measurement interval, as a measurement becomes obsolete when a new one is available. Hence, the maximum number of HARQ retransmissions should be limited to

$$R_{max} = \left\lfloor \frac{T - d_{max}}{\Delta t} \right\rfloor, \quad (3.10)$$

where T is the current measurement interval, d_{max} is the maximum network delay, and Δt is the minimum duration for a HARQ retransmission (i.e. one WiMAX frame for the HARQ with CC).

3.3.3 Capacity Analysis

Before moving on to the performance analysis, we conduct a basic capacity analysis of the WiMAX network in terms of the number of PMUs it can support. We are particularly interested to see the effect of persistent scheduling and robust header compression (ROHC) on the overall capacity and utilisation of the network. The analysis is conducted for a generic WiMAX network based on the OFDMA parameters specified in Table 3.3. Only PMU to PMU communication is considered, as it comprises both UL and DL data components.

Let us consider the OFDMA TDD frame structure in Figure 3.1. The total number of OFDM symbol-times available in the frame is given by

$$N_s = \left\lfloor \frac{T_f - T_{TTG} - T_{RTG}}{T_s} \right\rfloor. \quad (3.11)$$

This yields a total 47 symbol-times for the assumed configuration in Table 3.3. For a UL/DL subframe ratio of 1:1, the first DL symbol time is used as preamble, the next 22 symbols are allocated to the DL subframe, and the remaining 24 symbols are allocated to the UL subframe. Thus, the number of OFDM symbol-times available in the UL is 6528 ($= N_{sc}^{UL} \times 24$) and in the DL is 7920 ($= N_{sc}^{UL} \times 22$).

As seen from Figure 3.1, the DL subframe hosts the DL-MAP and the UL-MAP signalling messages. Each of these messages comprises a fixed header followed by a number of IEs. According to the IEEE 802.16-2009 standard, a typical DL-MAP header size is 88 bits and an IE size is 32 bits. In contrast, a typical UL-MAP header size is 48 bits and IE size is 48 bits. Each UL-MAP message contains three fixed IEs for ranging and CQICH allocation areas (total 144 bits). Both DL and UL MAP messages are preceded by a WiMAX MAC header (48 bits). Thus, for an unloaded network (i.e., number of data IE=0), the DL-MAP message size is $L_{DLMAP} = 136$ bits and the UL-MAP message size is $L_{ULMAP} = 240$ bits (in total 376 bits). Moreover, the MAP messages are often sent with two or four times of repetition coding to ensure robust transmission over the air-interface. Assuming a repetition coding rate of 4, the total MAP message size is $(376 \times 4) = 1504$ symbols per frame. Thus, the number of available data symbols in the DL subframe becomes $N_{DL}^{Data} =$

$$(7920 - 1504) = 6416.$$

Conversely, the UL subframe contains the IR channel (6 subchannels x 1 symbol-times), the BR channel (6 subchannels x 2 symbol-times), and the CQICH (1 subchannel x 6 symbol-times) channel. Considering partial usage of subchannels (PUSC) in the UL, the number of data subcarriers in each subchannel is 16. Thus, the number of available data symbols in the UL subframe becomes $N_{UL}^{Data} = [6528 - (6 \times 1 \times 12 + 6 \times 2 \times 16 + 1 \times 6 \times 16)] = 6144$.

Now, consider a synchrophasor application where each measurement packet has a DL and a UL component and requires both DL-MAP and UL-MAP signalling elements to be transported through the WiMAX BS (similar to the LCDP application described earlier). From Figure 3.8, we see that the MPDU size (L_{MPDU}) for this application is 58 bytes with 20 bytes of payload and 38 bytes of protocol overheads. The MPDU size can be further reduced by using an IP header compression techniques, such as ROHC. WiMAX supports ROHC over both UL and DL data connections. For this study, let us assume that the use of ROHC reduces the size of UDP/IP overhead to six bytes [64]. This yields a reduced MPDU size of 36 bytes.

Thus, the overall radio resource utilisation of a single MPDU within one WiMAX frame is given by

$$U = U_{DL} + U_{UL} = \frac{(L_{MPDU} + L_{DLMAP})}{N_{DL}^{Data}} + \frac{(L_{MPDU} + L_{ULMAP})}{N_{UL}^{Data}}. \quad (3.12)$$

Note that the above equation assumes QPSK with 1/2 rate FEC as the default modulation and coding scheme (MCS) so that the data usage can be evenly compared. However, the actual resource utilisation might be slightly higher than the one obtained by (3.12) due to the wastage of symbols during rectangular resource allocation in the OFDMA subchannels [58]. Next, a number of numerical results are provided to illustrate the outcome of the above capacity analysis.

First, we compare the effect of ROHC and persistent scheduling techniques with the baseline UGS configuration. Figure 3.11 shows the overall utilisation of a WiMAX frame for a single PMU under these configurations.

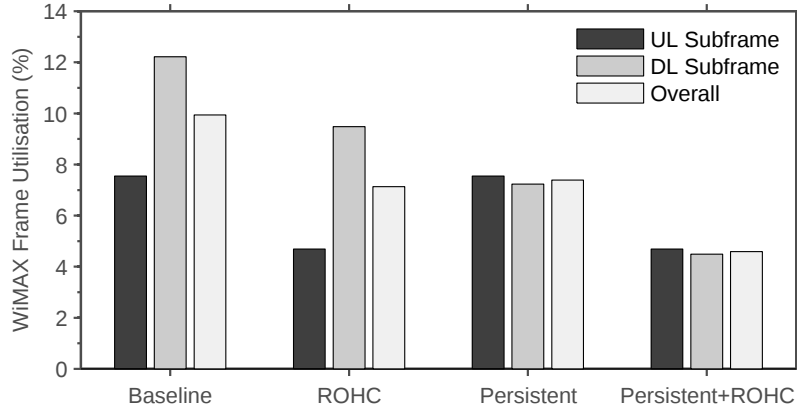


Figure 3.11: WiMAX frame utilisation for a single PMU under different UGS configurations

From the results, we see that the baseline configuration requires the highest amount of radio resources among all configurations. Also, the DL subframe utilisation is higher than that of the UL. This is because the DL subframe contains both DL and UL signalling components, which substantially uses its available resources. This is further evident in Fig 3.12 for the baseline scenario, that is, as the number of PMUs increases, the signalling overheads increase at a higher rate than the data burst usage, which in turn reduces the overall capacity of the frame. Note that under the baseline configuration, the maximum number of PMUs that can be accommodated in a WiMAX frame is eight (based on the combined MAP and data usage in the DL subframe and without considering fragmentation and packing).

Although the use of ROHC improves the UL frame utilisation (see Figure 3.11), the problem of higher DL subframe utilisation still remains. This is solved when persistent scheduling is used, which removes the signalling overheads associated with each packet. However, the best utilisation is achieved when ROHC and persistent scheduling are combined. To illustrate this, let us calculate the maximum number of PMUs that can be supported by a WiMAX network. It is given by

$$N_{PMU} = n_{PMU} \frac{T}{T_f}, \quad (3.13)$$

where n_{PMU} is the estimated number of PMUs per frame (see Figure 3.12), T_f is the

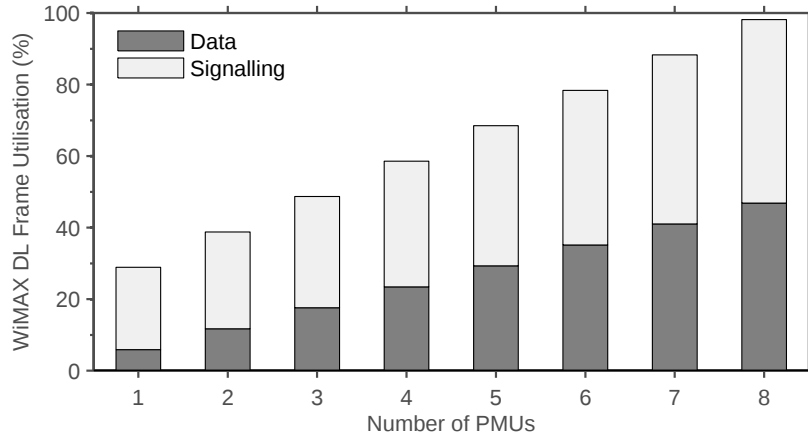


Figure 3.12: WiMAX DL subframe utilisation at different number of PMUs under the baseline UGS configuration

WiMAX frame duration, and T is the measurement interval of the PMUs.

For this calculation, let us consider both PMU to PDC and PMU to PMU communication schemes. The measurement cycles are assumed to be 40 ms and 20 ms for these two communication schemes respectively, and their payload size is set based on the discussions in Subsection 3.3.1. The corresponding results are listed in Figure 3.13.

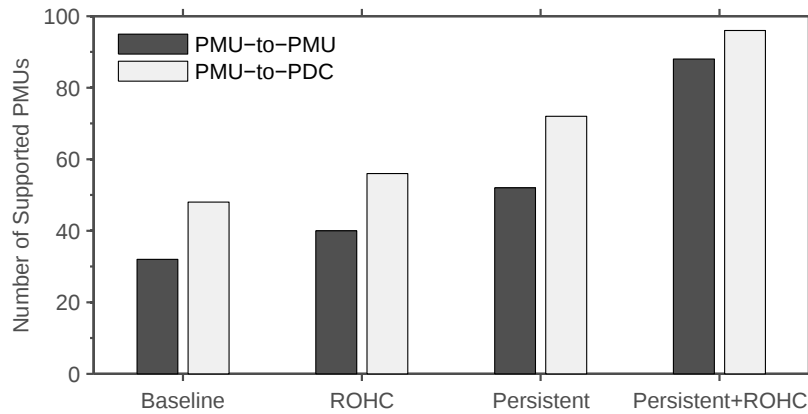


Figure 3.13: Number of supported PMUs under different UGS configurations.

From Figure 3.13, we see that in both these cases, the combined use of ROHC and persistent scheduling significantly reduces the PMU communication traffic, compared to the baseline configuration.

3.3.4 Simulation Study

To evaluate the performance of the synchrophasor communication schemes described in Subsection 3.3.1, a single cell WiMAX model is developed in OPNET. The following two simulation scenarios are considered:

1. PMU to PDC data transfer with a measurement cycle of 40 ms (i.e., a reporting rate of 25 Hz); and
2. PMU to PMU data transfer under the LCDP scheme with a current measurement cycle of 20 ms (i.e. a reporting rate of 50 Hz).

For both of the simulation scenarios, the propagation channel is modelled using the Erceg (Type-B) path-loss model with an additional shadow fading margin of 10 dB. We further assumed that the PMUs (i.e., WiMAX SSs) would be mounted at the existing utility poles and the substation structures to take advantage of high antenna heights [16]. The OFDMA parameters are same as in Table 3.3. The QPSK with $\frac{1}{2}$ rate FEC used as the default MCS level.

For the first simulation scenario, the contention parameters for the BE scheduling is set as $T_{16} = 8$ frames, $W_0 = 2^2$, $W_f = 2^4$ and $m = 16$. In addition, both the unsolicited grant interval of the UGS scheduling and the polling interval of the rtPS scheduling is set to 40 ms to meet the maximum delay requirement of the PMU application. To examine the effect of PMU traffic load on the application performance, the number of PMUs was increased from 10 to 50, based on the capacity analysis in the previous subsection for the baseline UGS configuration. The corresponding delay performance of the PMU traffic under the three scheduling services of WiMAX is summarised in Table 3.6.

From the results in Table 3.6, we see that both the UGS and the rtPS scheduling services can meet the maximum delay requirement of 40 ms. However, in both cases, delay starts to increase slightly with the number of PMUs. This is because as the number of PMU increases, the UL-MAP size increases to accommodate more number of IEs. As a result, the PMUs need to wait longer to extract the UL grant information, which in turn increases

Table 3.6: Summary of Delay Performance for PMU Communications

Service	No.of PMUs	Avg. (ms)	Max. (ms)	Std. Dev.
UGS	10	16.92	34.92	8.43
	20	18.92	39.92	10.68
	30	21.92	39.92	9.97
	40	23.42	39.92	11.52
	50	22.40	39.90	11.24
rtPS	10	22.06	39.92	10.25
	20	26.90	39.92	9.48
	30	24.00	34.92	7.84
	40	28.42	39.92	9.05
	50	31.61	39.92	7.97
BE	10	98.33	102.78	1.60
	20	167.60	173.92	2.92
	30	288.83	303.04	5.55
	40	143745.46	167104.46	46039.83
	50	253798.37	302609.60	88883.95

the delay. Note that the UGS provides better delay performance than the rtPS scheduling. This is because the polling process in the rtPS scheduling adds an additional delay of one WiMAX frame to receive poll and send the BR.

Under the BE scheduling service, the delay increases exponentially with the increase in the number of PMUs. This is because all the PMUs produce synchronised measurements that increase the instantaneous contention level of the system during a reporting instant. Higher the PMU traffic results in higher the contention levels and longer delays. This is further depicted in Figure 3.14, which shows the CDF of contention retries at different PMU traffic loads. Note that at low PMU loads, the BE scheduling also experiences large delays due to higher instantaneous contention level in the system. Hence, it is quite clear that the conventional BE scheduling service cannot meet the QoS requirements of PMU communication.

Next, we compare the performance of the UGS scheduling at two delay bounds with a moderate number of PMUs (20 in this case). The results are summarised in Table 3.7. The results indicate that although the UGS scheduling can meet the tighter delay bounds, it consumes an additional 12.5 per cent signalling overhead, and takes up 27.7 per cent more

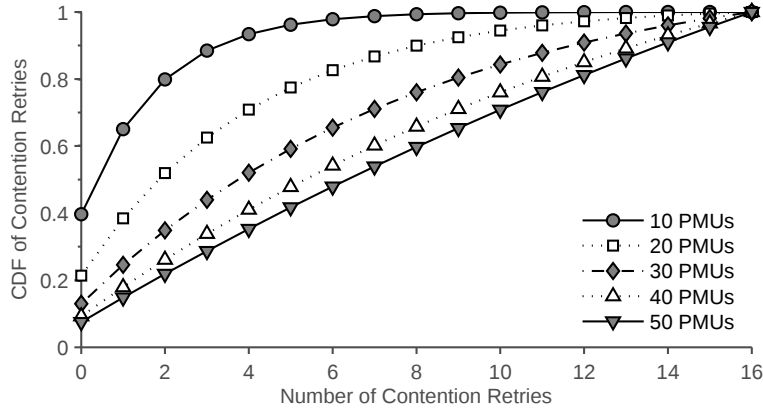


Figure 3.14: CDF of contention retries under the BE scheduling for PMU communications.

uplink data resources. The reason is that when the delay bound is less than the duration between two consecutive measurements, one in every two grants is wasted, as the PMU has no data to send.

Table 3.7: UGS Scheduling Performance with Different Delay Bounds

Parameters	Delay Bound	
	40 ms	20 ms
Min. Delay	24.19 ms	15.69 ms
Max. Delay	39.92 ms	19.92 ms
% MAP Usage	37.6 %	50.14 %
% UL Data Usage	23.74 %	51.47 %

In the second simulation scenario, we look into the delay performance of a PMU-based differential protection application under the baseline UGS service, with regular and synchronised grant allocation. We also vary the number of PMUs from 8 to 32, to examine the effect on the overall delay performance. The corresponding delay statistics are listed in Table 3.8.

From the results, we see that under both of the UGS allocation modes, the WiMAX network transfers the current measurements within the stipulated 20 ms delay bound. However, grant synchronisation significantly improves the UL delay, while the DL delay remains the same. This occurs as the UGS grants are synchronised with the packet generation times,

Table 3.8: WiMAX UGS Delay Statistics

(a) Regular UGS Grant				
No. of PMUs	Uplink		Downlink	
	Mean	S.Dev.	Mean	S.Dev.
8	9.15	1.81	5.98	0.11
16	9.87	1.58	6.15	0.21
24	11.55	1.06	6.33	0.33
32	13.02	0.79	6.51	0.44

(b) Synchronous UGS Grant				
No. of PMUs	Uplink		Downlink	
	Mean	S.Dev.	Mean	S.Dev.
8	5.01	1.48	5.98	0.11
16	5.04	1.48	6.15	0.21
24	5.06	1.48	6.33	0.33
32	5.08	1.48	6.51	0.44

the PMUs are able to send their measurements immediately without any additional waiting period. This extra delay margin can be used to allow retransmissions in the network. Note that as the number of PMUs increases, the amount of delay increases for all cases. This is because the BS has to allocate data bursts in the UL/DL subframe over more OFDM symbol-times, and a SS has to wait longer to find its data grant.

Next, we examine the effect of packet loss and retransmission on the network performance. The simulation scenario comprises 16 PMU-enabled relays, each sending measurements using the baseline UGS service with synchronised grant. A fast fading margin of 5 dB was introduced in the network to simulate the effect of random packet loss. From the previous results in Table 3.8, we see that the end-to-end packet delay for such a configuration is around 11 ms, considering both UL and DL packet delay. Hence, to meet the delay bound of 20 ms, only one HARQ retransmission can be allowed as per (3.10). The corresponding throughput and delay performance of the WiMAX network is plotted in Figures 3.15 and 3.16, respectively.

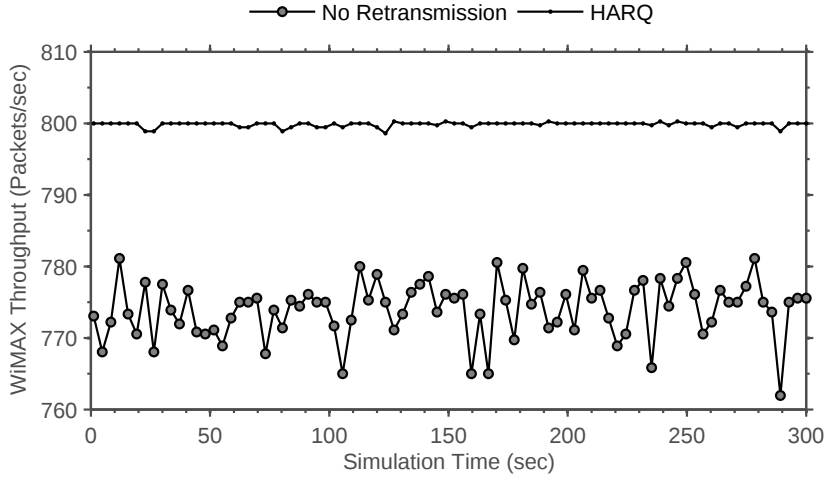


Figure 3.15: WiMAX network throughput for the LCDP scheme under no retransmission and HARQ retransmission.

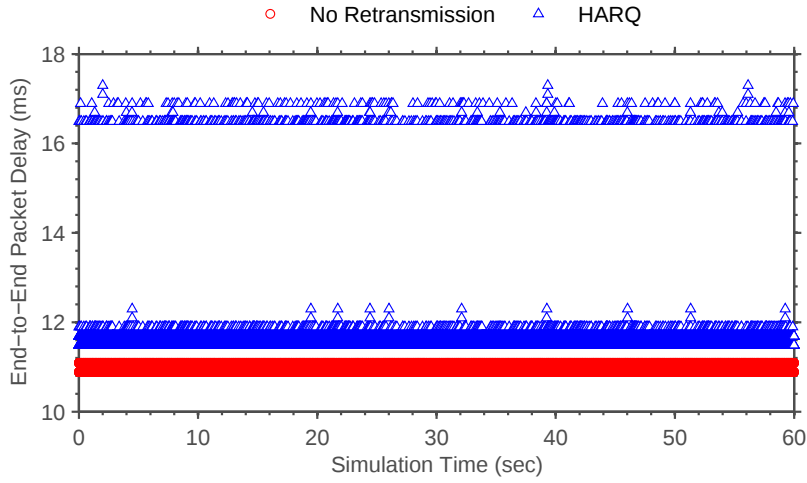


Figure 3.16: Distribution of packet delays for the LCDP scheme under no retransmission and HARQ retransmission (1 out of 5 minutes of simulation run-time).

From the results, we see that using HARQ allows the WiMAX network to recover most of the lost packets. However, it increases the end-to-end packet delay due to an extra retransmission delay component, which is one WiMAX frame, as per (3.10). Nonetheless, such a delay is acceptable as the end-to-end delay still remains below the required 20 ms bound.

3.3.5 Key Findings

The simulation studies clearly indicate that conventional BE scheduling fails to meet the latency requirement of PMU traffic, as it incurs large delays due to contention in the random access channel. In contrast, both the UGS and the rtPS scheduling services meet the latency requirement of PMU traffic. Of these two scheduling services, the UGS provides better delay performance and consumes relatively less radio resources.

Under the UGS scheduling mode, more PMUs can be accommodated by the combined use of persistent scheduling and the ROHC technique. In addition, by synchronising the PMU measurement time with the UGS allocation time, the WiMAX BS can significantly improve the delay performance of the PMU applications. This additional delay margin can allow a fast retransmission opportunity, based on the HARQ-CC technique, which may further improve the reliability of the PMU communication.

In summary, the conventional UGS scheduling service, along with some additional features such as persistent scheduling and ROHC, can provide a optimum platform for PMU communications over a WiMAX network. Nevertheless, although the individual size of a PMU measurement packet is small, their continuous high frequency data transmission takes up a significant amount of radio resource from the WiMAX network. This may adversely affect the performance of other applications in a multi-service Smart Grid communications network. Hence, application-specific performance optimisation is required to further reduce the synchrophasor communications load in such a network.

3.4 Protective Relaying over WiMAX

Protective relaying is considered as one of the most challenging applications in a Smart Grid communication network due to their stringent delay and reliability requirements. In this section, we examine the use of a conventional WiMAX network to support such an application. In particular, we consider the case of distribution pilot protection scheme described in Subsection 2.2.4. Such a scheme requires high-speed, peer-to-peer communications among the relays to perform the tripping operation. Here, the key challenge for

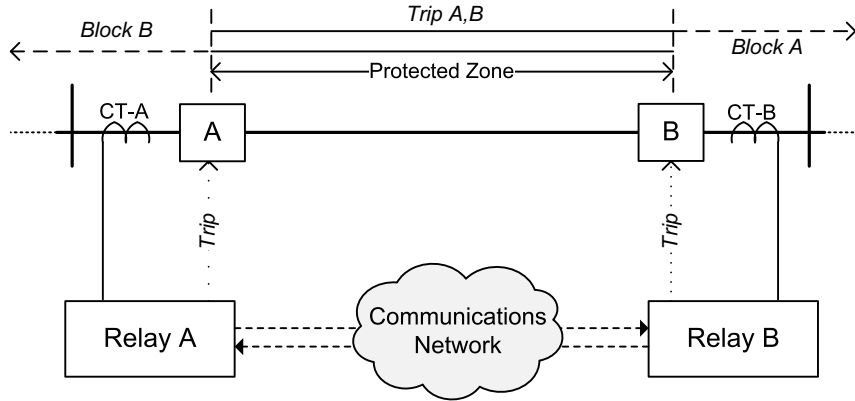


Figure 3.17: A generic pilot protection scheme in a two-terminal feeder.

a WiMAX network is to transfer the event-driven, time-critical pilot signals within the required delay bound in a multi-service Smart Grid environment.

3.4.1 Traffic Model

Under a pilot protection scheme, a group of two or more pilot relays form a protection zone, where each relay measures the voltage and current of its own terminal to calculate the impedance to a forward or reverse fault. This information is then exchanged between the peer relays in the form of a blocking signal if the fault is behind the relay (for the DCB scheme), or a permissive trip signal if the fault is in front of the relay (for the POTT scheme). A pilot trip occurs only if a relay detects a fault and a permissive trip signal is received (or a block signal has not been received) from the remote end.

Figure 3.17 shows the operating principles of a generic pilot protection scheme in a two-terminal feeder. When a fault occurs, each relay sends a pilot (trip or block) signal to its remote counterpart based on type of the protection scheme used. For example, if the POTT scheme is used and the fault is located between terminal A and B, a trip signal is exchanged between terminal A and B that allows permissive tripping in both relays. Conversely, if the DCB scheme is used and the fault is outside the protection zone, for instance to the left of terminal A, A sends a block signal that prevents tripping in B. For more details about the pilot protection schemes, please refer to [65].

Note that a pilot protection scheme can also be deployed in a multi-terminal protection

zone. Although the principle of operation is similar to a two-terminal zone, it involves exchange of pilot signals among multiple relays. Alternatively, a protected transmission line may have more than one protection zones. Within each zone, the member relays pick up a fault based on their *zone characteristics functions*. Often the zones over-reach with one-another to constitute a hierarchical protection system. However, in such cases the relays are time-graded using a *zone coordination scheme* so the successive zones operate sequentially [66].

For this study, we assume that the pilot signals are encapsulated in an IEC61850-90-5 based routed GOOSE (R-GOOSE) protocol and sent over the UDP/IP transport to one or more peer relays located in different IP subnetworks [36]. In addition, a delay margin of 40 ms (i.e., 2 cycles) is assumed for the pilot signals, considering a 50 Hz power system.

3.4.2 QoS Mapping

Figure 3.18 illustrates the protocol stacks for pilot protection communications over a WiMAX network. The IP specific convergence sublayer (IPCS) of the WiMAX MAC layer receives the incoming UDP/IP packets and maps them to an IEEE 802.16 MPDU to be sent over the IEEE 802.16 air-interface. The BS forwards the incoming pilot signal data packets to the destination relays based on their destination IP addresses. Thus, each pilot signal data packet has a UL and a DL component.

Note that for a WiMAX network, the multi-terminal pilot protection scheme can be considered as special case of a two-terminal protection scheme, where the BS multicasts the DL component of the pilot signal data packet to all other member relays within the same protection zone.

As discussed in Subsection 3.1.3, the WiMAX radio resource scheduler is located at the BS that provides bandwidth allocations to both UL and DL connections based on their QoS attributes. In the UL, a BR mechanism allows the SSs to inform the BS about their bandwidth needs. In case of pilot protection, the trip/block signals are generated only after a fault event which is a highly random phenomenon caused by various external factors such

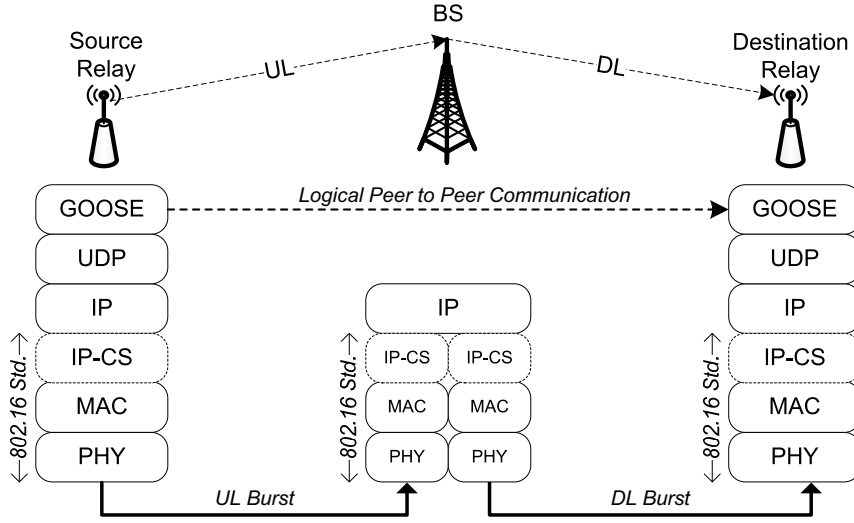


Figure 3.18: Protocol stacks for pilot protection communications over a WiMAX network.

as storms, bushfire, lightning, trees, and animals. A protected line can be without faults for days and even weeks. Hence, the BR mechanisms, such as polling and unsolicited grants, are not suitable since the allocations will be wasted during the normal operation of the line. This leaves us with only one scheduling choice - the random access based BE service that supports delay-tolerant bursty data traffic. Figure 3.19 depicts the various delay components of the pilot signal transmission under the conventional BE service.

Transmission time of a pilot signal message depends heavily on the random access delay due to stochastic variability in the random access channel, as described in Subsection 3.1.3. The other delay components such as the UL and the DL delays (see Figure 3.19) depend on the grant scheduling algorithm used by the BS. As the BE service is typically used to serve delay tolerant applications, it has the least priority in the scheduler and uses a simple round-robin (RR) algorithm for grant scheduling.

The performance of the RR scheduler can be easily analyzed using an M/D/1 queuing model with an exponentially distributed arrival rate λ , a deterministic service rate μ (of one WiMAX frame), and a single server queue. According to the queuing theory, the average time spent by a packet in an M/D/1 queue is given by

$$E(W) = \frac{2 - \rho}{2\mu(1 - \rho)}, \quad (3.14)$$

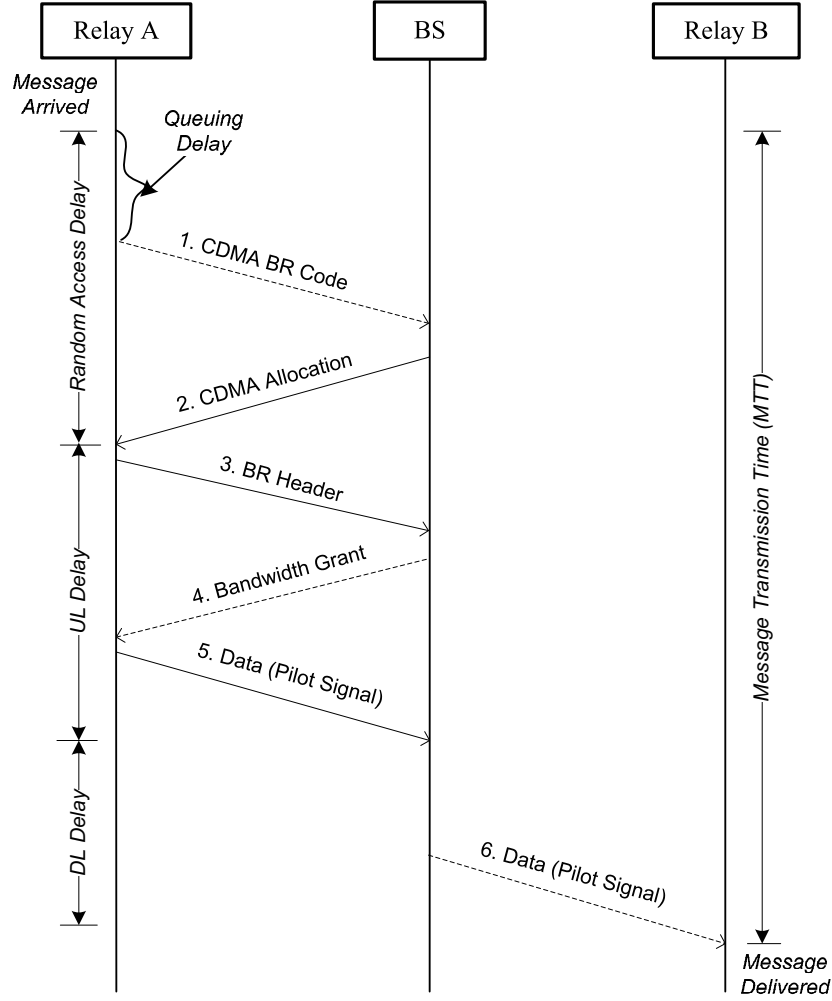


Figure 3.19: Pilot signal message transmission using the conventional BE scheduling service.

where ρ is the traffic load of the system ($= \lambda/\mu$). From (3.14) we see that the grant delay increases with the increase in the traffic load, as a BR needs to wait longer for its turn. Note that an increased grant delay may also increase the random access delay. This is because if a device does not receive a grant within T_{16} period, it re-starts the back-off process.

3.4.3 Simulation Study

To evaluate the communications performance of a pilot protection scheme over a WiMAX network, we have developed a single cell simulation model with a cell radius of 5 km. The OFDMA parameters are same as specified in Table 3.3. Under the simulation scenario, we

assume there are four pilot relays per km^2 , a total of 320 relays within the coverage area of 78.5 km^2 (approx.). Each relay is assumed to communicate with its peer relay only, i.e. two relays per protection zone. Each protection zone is assumed to experience a single fault during the simulation run-time of one hour, and each fault is associated with two pilot signals (see Figure 3.17). Thus, there are 640 pilot signals during the simulation run-time. All the relays are assumed to be in radio resource connected mode during the course of the simulation. This can be achieved by regularly exchanging either periodic ranging messages or explicit status update messages at a pre-defined interval.

To ensure the least possible delay for the pilot protection application under the conventional BE scheduling service, an aggressive contention parameters was assumed, as listed in Table 3.9.

Table 3.9: WiMAX Simulation Parameters

Parameters	Values
No. of BR Opportunities	4 per frame
Initial Back-off Window (W_0)	2^2
Final Back-off Window (W_f)	2^8
No. of Max. Retries (m)	8
No. of Detectable Codes/channel	2
Probability of Misdetection	0%
Contention Time-out (T_{16})	4 frames

According to the IEC 61850 standard, each GOOSE message is comprised of a 26 byte Ethernet-based MAC header followed by an application-specific payload, called the application protocol data unit (APDU) [28]. However, since we are assuming that the relays are WiMAX enabled, the Ethernet-based MAC header is not required for the GOOSE message routing. Instead, each GOOSE APDU is tagged with a 6 byte WiMAX generic MAC header and a 4 byte CRC. Besides, each pilot signal MPDU contains an 8 byte UDP header and a 20 byte IPv4 (IP version 4) header. Figure (3.20) shows the MPDU structure of the pilot signal data packet used for this study.

For application level performance analysis, we use the following two performance indicators: message transmission time (MTT) and message delivery success rate (MDSR). The

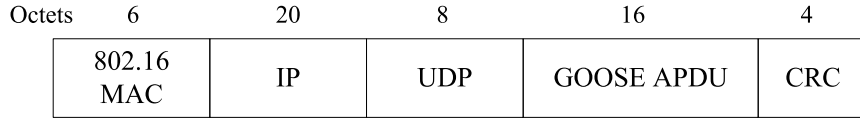


Figure 3.20: MPDU structure of the pilot signal data packet

MTT represents the end-to-end delay for the pilot signal data packets (see Figure 3.19). The MDSR represents the percentage of total pilot signal messages that were received before their TTL period expired. As mentioned earlier, a message expiry time of 40 ms is assumed for the pilot signals. Moreover, to imitate a multi-service environment, we analyse the performance of the pilot protection application with variable background BE traffic between 100 to 400 packets/sec. A small packet size of 100 bytes with a Poisson arrival process is considered for the background traffic that represents various bursty AMI applications, such as AMR, DR and EV. The corresponding delay components of the pilot signal data packets are listed in Table 3.10.

From the results, we see that as the amount of background BE traffic load increases, the random access delay increases, with a large standard deviation from the mean. This is because as more devices transmit, the probability of collision increases and the devices need to undergo more back-off stages for a successful BR. On the other hand, both the UL and the DL delay components increase with an increase in the amount of background BE traffic loads. This is because as the pilot signals data packets have the same priority as other BE packets, they experience the same mean waiting time in the RR scheduling queue, which also increases with load.

Figure 3.21 shows CDF of the MTTs under various background BE traffic loads. The

Table 3.10: Pilot Signal Delay Components (in ms) under the BE Service

Load (Packets/sec)	Random Access		Uplink		Downlink	
	Mean	S.Dev.	Mean	S.Dev.	Mean	S.Dev.
100	39.95	23.63	5.02	0.28	6.18	0.18
200	45.95	35.12	5.36	3.38	6.43	0.19
300	53.86	84.22	5.51	3.54	6.51	0.37
400	68.81	92.65	6.15	4.40	6.60	0.49

results show the aggregate effect of increase in various delay components, as listed in Table 3.10. Due to the increased MTT, a large number of pilot signal packets fail to reach the destination within the required 40ms delay bound². In turn, this decreases the MDSR of the application as shown in Figure 3.22. Note that, even at minimum background BE traffic with aggressive contention parameters, the MDSR trails below 50 per cent.

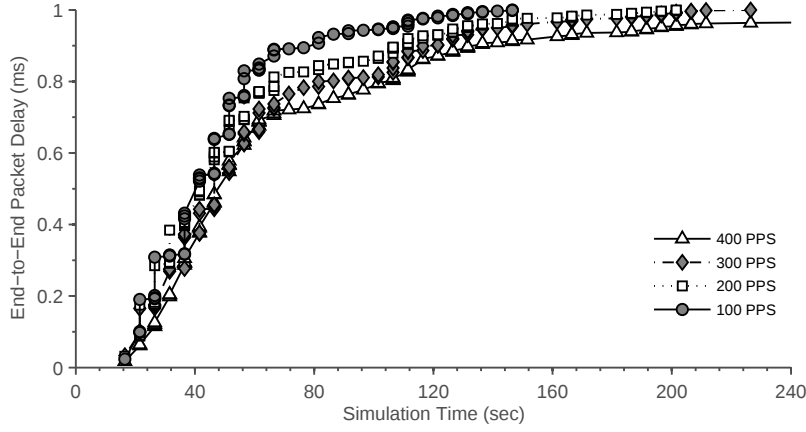


Figure 3.21: CDF of MTT for the pilot signals under the conventional BE service at different background traffic loads.

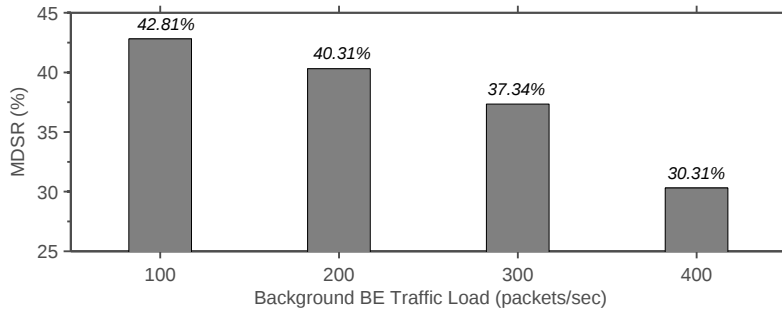


Figure 3.22: MDSR for the pilot signals under the conventional BE service at different background traffic loads.

²MDSR represents the number of packets delivered within the expiry time. It does not indicate any packet loss in the system.

3.4.4 Key Findings

From the above analysis, it is quite clear that the conventional BE scheduling service of WiMAX fails meet the QoS requirement of the pilot protection traffic. This is mainly due to the fact that the random access based BR mechanism in WiMAX uses a shared contention media without any traffic differentiation. Moreover, all the packets share the same scheduling queue for a BR grant after a successful access. As a result, the high priority pilot signal data packets experience the same average delay as the other low priority packets. Hence, random access differentiation and prioritised grant scheduling are required for such a mission-critical application to be supported by a WiMAX network.

3.5 Chapter Summary

In this chapter, we conducted a detailed study on AMI, synchrophasor and protective relaying communications over a conventional WiMAX network. While a generic traffic model was used for the performance analysis of the AMI traffic, an application-specific traffic model was used for the protective relaying traffic. In contrast, both a generic and an application-specific model were used for the synchrophasor traffic.

The coverage analysis in Section 3.2 shows that a typical WiMAX cell has a moderate range within a suburban AMI/Smart Grid network with predominantly indoor SSs. However, the coverage can be greatly improved if outdoor SSs are used due to improved antenna height/gain, as well as immunity from the building penetration losses. In Chapter 6, we discuss this in details and propose a hybrid network architecture for improved network coverage and utilisation.

The simulation studies in Sections 3.2 and 3.4 reveal that the existing random access mechanism of WiMAX can significantly degrade the performance of bursty AMI applications, especially when the access load is high. The performance is even worse for event-based traffic, such as outage alarms from the smart meters and the protection signals from the relays. Hence, further investigation of the random access plane is required, followed by the

development of appropriate load control and traffic differentiation techniques. These are carried out in Chapter 4 and 5, respectively.

The synchrophasor traffic can be well-supported by the conventional UGS scheduling service of WiMAX. Moreover, the use of additional WiMAX features such as persistent scheduling and ROHC can significantly reduce the radio resource usage of PMU traffic, as demonstrated in the capacity analysis of Subsection 3.3.3. Although no additional enhancement is required, application-specific optimisation may further reduce traffic load for such applications. This may free-up valuable radio resources from the WiMAX network. To illustrate this, a number of application-specific optimisations are conducted in Chapter 7 related to synchrophasor communications.

Chapter 4

Random Access: The Key Bottleneck

In the previous chapter, a number of simulation studies revealed that the random access plane of a conventional WiMAX network is the key bottleneck for supporting bursty M2M applications in the Smart Grid. Such a notion has also been acknowledged by the IEEE 802.16p working group on M2M communications [8]. Therefore, it is imperative that we develop detailed understanding of the WiMAX random access procedure so it can be enhanced and optimised to efficiently support M2M applications in the Smart Grid.

The mandatory random access mechanism for the current generation WiMAX networks (IEEE 802.16-2009 and beyond) is based on a multicarrier CDMA technique, where multiple users are allowed to collide with multiple codes [58]. However, the BS is able to detect only a limited number of codes in the presence of multiple access interference (MAI) from different codes [67]. Besides this, the code detection performance is affected by random noise and frequency-selective fading in the multipath wireless channel. To retransmit a collided or a misdetected code, the SSs use a distributed contention resolution protocol based on the TBEB algorithm. Thus, the overall random access performance depends on not only the back-off parameters of the system, but also the physical parameters of the channel, such as the number of available codes to the SSs and the number of detectable codes by the BS.

Unfortunately, while a number of analytical models are available for the legacy message-based contention procedure, only a few have considered the dynamics of the CDMA-based contention procedure. Therefore, to allocate optimum random access resources to both

M2M and conventional traffic in a Smart Grid communications network, it is quite important to develop a tractable model that can accurately evaluate the key performance metrics of the CRA procedure of WiMAX.

To this aim, in this chapter we present a comprehensive analytical model to evaluate the performance of CRA procedure for the contemporary WiMAX networks. Compared to the literature, our model incorporates all key features of the CRA procedure, such as multi-user multi-code transmission, parallel code detection, and back-off parameters in the performance analysis matrix. The unsaturated conditions are captured through a detailed queuing analysis and integrated with an extended Markov chain model to obtain closed-form expressions for the random access delay and throughput under both saturation and non-saturation conditions. The accuracy of the model was validated by an extensive set of simulation results using the well-known OPNET simulator.

The rest of the chapter is organised as follows. Section 4.1 describes the system model including the code detection process. Section 4.2 describes the related works and highlights our key contributions. Section 4.3 presents the proposed analytical model. Section 4.4 validates the proposed model by comparing simulation and theoretical results and provides some interesting insights on the performance of the CRA procedure. Finally, Section 4.5 concludes this chapter ¹.

4.1 System Description

Under the CRA mechanism, the BS allocates a group of random access channels and codes (also known as ranging channels and ranging codes) via the UL-MAP broadcast message. A random access channel is comprised of six adjacent OFDMA subchannels over $L = 144$ randomly chosen subcarriers. Separate channels are used for network entry (i.e., ranging) and the BR procedure. A random access code is an L -bit pseudo-random binary sequence (PRBS), chosen with equal probability from a bank of K codes, where $K \leq 256$ [58]. The available random access codes are further divided into four subgroups for ranging (initial,

¹A part of this chapter has been accepted for publication in proceedings of the *IEEE International Conference on Smart Grid Communications (SmartGridComm)* in Nov 2014 [68].

periodic, and handover) and BR purposes. Within each random access channel, multiple users are allowed to collide using multiple codes. The BS can detect whether a code is present or not in a random access channel by correlating the received signal with all the available random access codes.

Let us denote the BR code-matrix as $\mathbf{C} \in \mathbb{R}^{L \times K}$. A component of $\mathbf{C}$ is indicated by $\mathbf{C}_{l,k}$, where $l = 1, 2, \dots, L$ and $k = 1, 2, \dots, K$; and the k^{th} column of $\mathbf{C}$ represents an independent code denoted by $\mathbf{C}_k$. As described in Subsection 3.1.3, for a random access, a SS picks a random column from the code matrix $\mathbf{C} \in \mathbb{R}^{L \times K}$ and transmits it onto the BR channel by BPSK modulating each of its L subcarriers, i.e. $\mathbf{C}_{l,k} \in \{-1, +1\}$.

Now, consider a time instant where M number of SSs are simultaneously contending on the same BR channel. Let, the m^{th} SS selects the k_m^{th} column of the code matrix $\mathbf{C} \in \mathbb{R}^{L \times K}$, where $m = 1, 2, \dots, M$. The transmitted code over the l^{th} subcarrier from the m^{th} SS be

$$x_{l,m} = \mathbf{C}_{l,k_m} \text{ for } \forall l \in \{1, 2, \dots, L\}. \quad (4.1)$$

More details on OFDM modulation and demodulation process is provided in Appendix B.

In the BS, after down-conversion to baseband and OFDM demodulation, the received signal from the m^{th} SS over the l^{th} subcarrier be

$$y_{l,m} = x_{l,m}h_{l,m} + e_{l,m} = \mathbf{C}_{l,k_m}h_{l,m} + \eta_{l,m} \text{ for } \forall l \in \{1, 2, \dots, L\}, \quad (4.2)$$

where $h_{l,m}$ is the frequency response of the l^{th} subcarrier and $\eta_{l,m} \sim \mathcal{CN}(0, \sigma_{\eta,m}^2)$ is a zero-mean complex Gaussian noise.

Considering the above phenomena, the combined received signal at BS from all M stations over the l^{th} subcarrier be

$$y_l = \sum_{m=1}^M \mathbf{C}_{l,k_m}h_{l,m} + \vartheta_l, \quad (4.3)$$

where $\vartheta_l = \sum_{m=1}^M \eta_{l,m}$, which can be modelled as a zero mean complex Gaussian noise, i.e. $\vartheta_l \sim \mathcal{CN}(0, \sigma_{\vartheta}^2)$.

Let us construct:

$$\mathbf{y} = [y_1, y_2, \dots, y_L]^T, \quad (4.4)$$

where $(\cdot)^T$ denotes transpose of a vector.

To detect the presence of a BR code, the code detector in the BS correlates the received signal with the original code matrix $\mathbf{C} \in \mathbb{R}^{L \times K}$. The detector takes every column of $\mathbf{C}$ and computes its cross-product with $\mathbf{y}$. Let the detector picks a random column p denoted as $\mathbf{C}_p$, where $p = 1, 2, \dots, K$. The corresponding result of cross-product would be

$$\mathbf{C}_p^T \mathbf{y} = \sum_{l=1}^L \left\{ \sum_{m=1}^M \mathbf{C}_{l,k_m} \mathbf{C}_{l,p} h_{l,m} + \vartheta_l \mathbf{C}_{l,p} \right\} \quad (4.5)$$

If the code $\mathbf{C}_p$ is present in $\mathbf{y}$, i.e. $\mathbf{C}_{k_m} = \mathbf{C}_p$ for some m so that $\mathbf{C}_{k_m} \cdot \mathbf{C}_p = 1$; for all $\forall l$, we get

$$\mathbf{C}_p^T \mathbf{y} = \sum_{l=1}^L h_{l,m} + \sum_{l=1}^L \sum_{m=1, k_m \neq p}^M \mathbf{C}_{l,k_m} \mathbf{C}_{l,p} h_{l,m} + \sum_{l=1}^L \vartheta_l \mathbf{C}_{l,p}. \quad (4.6)$$

The first term of (4.6) represents the auto-correlation peak for the code $\mathbf{C}_p$, which is the sum of complex path gains for the station m . The second term represents the MAI from other users, which can be modelled as random variable with complex Gaussian distribution with zero mean and a variance σ_I^2 . The last term represents the noise floor of the correlator with zero mean and a variance σ_ϑ^2 . If $\mathbf{C}_p^T \mathbf{y}$ is greater than a pre-defined threshold ξ , the code $\mathbf{C}_p$ is assumed to be present in $\mathbf{y}$.

In the receiver, the mutual orthogonality of the codes are partially lost as the cross-correlation is performed between $h_{l,m} \mathbf{C}_{l,k_m}$ and $\mathbf{C}_{l,p}$, instead of $\mathbf{C}_{l,k_m}$ and $\mathbf{C}_{l,p}$. Therefore, as the number of contending codes increases, the MAI level in (4.6) start to increase, due to the residual cross-products. In turn, this reduces the probability of a successful detection. Thus, depending on the detection algorithm used, the BS can detect only a limited number of codes D , such that $D \leq K$.

Contention occurs when two or more SSs select the same random access code, as the BS

can not provide a unique bandwidth grant: it has no information about their connection IDs. The SS considers its code transmission lost (either due to contention or failed detection), if no grant has been received within the n waiting frames of the T_{16} time-out period. To resolve contention, WiMAX uses the TBEB algorithm with back-off parameters W_0 , W_f and m , as described in Subsection 3.1.3.

4.2 Related Works

Performance analysis and optimisation of the random access plane has been an intensely researched topic in WiMAX. A vast number of analytical models are available, mostly for the legacy message-based contention procedure. Under the message-based contention procedure, the BS allocates a number of fixed-size contention-slots and allows the SSs to send their BR headers directly. Thus, the message-based contention mechanism is a three step procedure, in contrast to the five step procedure of the CRA mechanism (see Figure 3.2). While there are fundamental differences between these two mechanisms in terms of channel configuration, both use the same TBEB contention resolution protocol.

One of the early works in this arena was reported by Vinel *et al.* (2005) in which they proposed an analytical model to calculate random access delay under saturated condition by adapting the Markov chain model for IEEE 802.11 proposed by Bianchi (2001) [69, 70]. An alternative model was proposed by He *et al.* (2007) based on fixed-point analysis to calculate both saturation throughput and delay [71]. They assumed that only a fixed number of grants can be served per frame and the successful BRs are either served in the next frame or dropped. Perera *et al.* (2007) proposed a Markov chain based model considering the effect of T_{16} period at the end of each back-off stage [72]. They assumed that the successful BRs are allocated in the next frame, while the unsuccessful BRs have to wait until the T_{16} timer expires.

Another notable work was published by Fallah *et al.* (2008) considering the effect of waiting time and timeouts in the T_{16} period. They assumed that the SSs have a constant probability of receiving a grant from the BS at each waiting frame depending on the admission

control strategy [73]. However, in practice the process has a conditional probability, that is, the probability that a SS will receive a grant in the current waiting frame decreases as it approaches the end of the waiting period.

Note that all the above models have considered saturated conditions for the performance analysis. However, at saturation the random access delay increases abruptly due to long waiting times in the SS queue and the throughput starts to drop, as some packets reach their retry limits. Therefore, for most of the time a typical WiMAX network operates under unsaturated conditions, where the overall random access load is low and the SSs generate requests intermittently.

To incorporate the unsaturated conditions, Ni *et al.* (2010) first proposed a Markov chain based model by considering the effect of the SS queues [74]. They assumed that a BR may fail either due to collision or channel impairment and evaluated the throughput and delay under various load and channel conditions. However, their assumption about the SS queuing behavior is not applicable for a request-based contention process, such as the CRA procedure (more details in the next section). Moreover, the effect of T_{16} period was not considered.

Recently, another Markov chain based model was proposed by Giovanni and Hadzic-Puzovic (2012) adopting a similar approach but with a more accurate analysis of the SS queues [75]. Moreover, they obtained the probability distribution for the grant allocation process by modelling the BS as a RR server that serves the BE traffic only. They further assumed a fixed packet size for all the requests (equal to the capacity of one UL subframe). However, in practice the BS serves multiple classes of traffic with variable packet sizes using vendor-dependent scheduling algorithms [76]. Further, multiple grants may be allocated in the same frame or a single grant may spread over multiple frames depending on the packet-size and QoS profile of the connection. With the CRA procedure, the BS only allocates a nominal bandwidth grant (the CDMA allocation IE), which decouples the grant scheduling process from the random access procedure.

Of the notable works on the CRA procedure, Seo *et al.* (2006) investigated its performance considering fixed delays in the T_{16} period [77]. Later, Seo and Leung (2011) further en-

hanced the work and proposed an access prioritisation scheme [78]. However, the request arrival process was not modelled and the effect of T_{16} period was not considered. A key assumption in their analysis was that the transmitted CDMA codes maintain perfect orthogonality, that is, the BS is able to detect all the codes that it may receive. However, in practice, the BS has a limited code-detection capability depending on the performance of the code detection algorithm under random noise and multipath fading.

Compared to the aforementioned works, the main contributions of the proposed model are as follows:

- a detailed queuing analysis to accurately evaluate the mean and probability distribution of the random access delays for both successful and unsuccessful requests;
- an extended Markov chain analysis to evaluate the access throughput and access success probability considering the stochastic variability during the T_{16} time-out period;
- analysis of the effect of multi-user code transmission and parallel code detection on the performance of the random access channel; and
- accurate prediction of the transition from non-saturation to saturation mode under a given set of back-off parameters.

4.3 Analytical Modelling

In this section, we present a comprehensive analytical model for accurate performance analysis of the WiMAX CRA procedure under both saturated and unsaturated conditions. While the model is developed by considering the BR aspect of the CRA procedure, it is equally applicable to the ranging process.

4.3.1 Modelling Assumptions

Let us consider a single cell WiMAX network with N number of SSs. The following assumptions are made for the proposed analytical model.

1. The request arrival process follows a Poisson distribution (exponentially distributed interarrival time) with an arrival rate λ . Such an assumption is consistent with the ranging evaluation guidelines specified by the IEEE 802.16p working group [61].
2. The smallest back-off period in the SS is one WiMAX frame T_f . We choose T_f as the unit of time. All other relevant quantities are expressed in this unit. For instance λ is the average number of request arrivals in one frame.
3. For each higher layer data packet, the SS MAC registers a request-pointer in the BR queue to execute the random access procedure. A SS does not initiate a new random access procedure until the last one has been completed.
4. Since the BR queue only stores request-pointers and removes a request-pointer at the end of the random access procedure regardless of its outcome, an infinite queue length is assumed.
5. As described in the previous section, the BS can simultaneously detect only a limited number of codes (denoted by D) in the presence of random noise, multipath fading, and MAI from different codes. If more than D codes collide on the same ranging channel, the BS will fail to detect any of them [61]. Also, the probability that the BS fails to detect an un-collided code remains constant independent of the back-off process.
6. Collision occurs if more than one SSs transmit the same code on the same ranging channel. In this case, all the code transmissions are assumed to be failed. This is because the BS can detect only one code (due to the capture effect) and schedules a CDMA allocation to that code-channel pair. However, multiple devices will recognise the same allocation; hence, their BR headers will collide. Moreover, the probability that a code transmission from a given SS will suffer a collision remains constant and does not depend on the previous retransmission history [70].
7. Upon successful reception of a code, the BS allocates a CDMA grant within the timeout period based on an implementation-specific algorithm, which is independent of

the contention process [73, 79]. However, to develop a generic model incorporating its effect, we assume that the BS allocates a bandwidth grant according to the probability mass function $f(x)$ within the n waiting frames of the T_{16} time-out period, where $x = 1, 2, \dots, n$. Also, we denote the mathematical expectation of $f(x)$ as μ such that $\mu = \mathbb{E}(x) = \sum_{x=1}^n x f(x)$.

4.3.2 Preliminaries and Notation

In the sequel we encounter random variables which take non-negative integer values. Here, we present a summary of the notation, conventions and basic results used later.

A random variable is written in boldface font. The probability distribution function of a integer valued random variable $\mathbf{v}$ is denoted as $P_{\mathbf{v}}(\cdot)$. Thus, the probability that $\mathbf{v} = \ell$ is $P_{\mathbf{v}}(\ell)$, where $0 \leq P_{\mathbf{v}}(\ell) \leq 1$, $\forall \ell$. Since all the probabilities must sum up to one, we can write

$$\sum_{\ell=0}^{\infty} P_{\mathbf{v}}(\ell) = 1. \quad (4.7)$$

The mean or mathematical expectation of $\mathbf{v}$ is defined as

$$E(\mathbf{v}) = \sum_{\ell=0}^{\infty} \ell P_{\mathbf{v}}(\ell). \quad (4.8)$$

In the sequel we make extensive use of characteristic functions. The characteristic function of $\mathbf{v}$ is denoted by $C_{\mathbf{v}}$ and defined as

$$C_{\mathbf{v}}(z) := E(z^{\mathbf{v}}) = \sum_{\ell=0}^{\infty} P_{\mathbf{v}}(\ell) z^{\ell}. \quad (4.9)$$

where z is a complex number in general. Note that since $\mathbf{v} \geq 0$, $C_{\mathbf{v}}(z)$ is analytic within the closed unit disc in the standard complex plane, and we can evaluate $C_{\mathbf{v}}(z)$ for any z satisfying $|z| = 1$. In particular, by setting $z = 1$ in (4.40) and using (1) note that

$$C_{\mathbf{v}}(1) = 1. \quad (4.10)$$

Consequently, we can compute the entire probability distribution $P_v(\ell)$ via inverse discrete-time Fourier transform:

$$P_v(\ell) = \frac{1}{2\pi} \int_{-\pi}^{\pi} C_v(e^{-i\omega}) e^{i\ell\omega} d\omega, \quad \text{where } z = e^{-i\omega}.$$

Now consider another random variable $\mathbf{u}$, which is potentially correlated with v . Let $P_{v|\mathbf{u}}(\ell|j)$ be the conditional probability that $v = \ell$, given that we know $\mathbf{u} = j$. If we know the distribution of $\mathbf{u}$ and the conditional distribution of v given $\mathbf{u}$, then we can calculate the distribution of v using the total probability theorem:

$$P_v(\ell) = \sum_{j=0}^{\infty} P_{\mathbf{u}}(j) P_{v|\mathbf{u}}(\ell|j). \quad (4.11)$$

The conditional characteristic function $C_{v|\mathbf{u}}(z|j)$ of v given that $\mathbf{u} = j$ is defined as the z-transform of $P_{v|\mathbf{u}}(\ell|j)$:

$$C_{v|\mathbf{u}}(z|j) := \sum_{\ell=0}^{\infty} P_{v|\mathbf{u}}(\ell|j) z^{\ell}. \quad (4.12)$$

As before, for all j the conditional characteristic function $C_{v|\mathbf{u}}(z|j)$ is analytic when $|z| \leq 1$, and we can recover $P_{v|\mathbf{u}}$ from the $C_{v|\mathbf{u}}$ via inverse discrete time Fourier transform of $C_{v|\mathbf{u}}$. It is also readily verified that the following identities hold for all j :

$$C_{v|\mathbf{u}}(1|j) = 1, \quad C(0|j) = P_{v|\mathbf{u}}(0, j), \quad \dot{C}(1|j) = E_{v|\mathbf{u}}(j).$$

Moreover, using (4.9), (4.11) and (4.12), we deduce that

$$C_v(z) = \sum_{j=0}^{\infty} P_{\mathbf{u}}(j) C_{v|\mathbf{u}}(1|j). \quad (4.13)$$

We note in passing that if v and $\mathbf{u}$ are independent and $w = \mathbf{u} + v$, then

$$C_{\mathbf{w}}(z) = E(z^{\mathbf{w}}) = E(z^{\mathbf{u}} z^v) = C_{\mathbf{u}}(z) C_v(z). \quad (4.14)$$



Let, the probability that code transmission fails be denoted by p . A discrete time 2-D Markov chain shown in Figure 4.1 is used to study the back-off process of a SS, where the vertical axis represents the stochastic processes at different back-off stages, i.e. $i \in \{0, 1, 2, \dots, m\}$, and the horizontal axis represents the stochastic processes within each back-off stage, i.e. $k \in \{-n, \dots, 0, 1, \dots, W_i - 1\}$. Within each back-off stage, the random states are divided into two phases: i) the back-off waiting phase before a code transmission, where $k \in \{0, 1, \dots, W_i - 1\}$, and ii) the grant waiting phase after a code transmission, where $k \in \{-n, \dots, -2, -1\}$. In addition, the two rhomboidal states at the beginning of the chain are used to represent *Idle* and *Active* states of the SS queue, respectively. Let q be the probability that a station is in the *Active* state.

Let $r_j : j \in \{1, 2, \dots, n\}$ be the steady-state transition probabilities of the successive grant

waiting states. In other words, the probability that the system jumps from state $(i, -j)$ to $(i, -j-1)$ is r_j . As can be seen from Figure 4.1, if a CDMA allocation has been received while the system is in state $(i, -j)$, the system goes back to the *Idle* state with probability $1-r_j$. On the other hand, if the system is in state $(i, -n)$, and it does not receive a grant while it is there, then the back-off index will increase and a new back-off counter will be uniformly chosen from the range $\{0, W_i - 1\}$, where $i < m$.

We apply the laws of conditional probability to find r_j . Firstly, the probability that the system is in the state $(i, -j)$ is

$$b_{i,-j} = b_{i,0} \prod_{k=1}^{j-1} r_k.$$

On the other hand, the system moves on to the state $(i, -j)$ when a transmission attempt is made from the state $(i, 0)$, but it does not receive a grant from the BS until the state $(i, -j-1)$. This happens if,

- either the code transmission attempt fails (probability of this is p), or
- the transmission attempt succeeds, but the grant is not received until time j . The probability that this happens is $(1-p) \sum_{x=j+1}^n f(x)$, where $x = 1, 2, \dots, n$.

Hence, the probability that the system is in the state $(i, -j-1)$ is

$$b_{i,-j-1} = b_{i,0} \prod_{k=1}^j r_k = b_{i,0} \left\{ (1-p) \sum_{x=j+1}^n f(x) + p \right\}.$$

After a few steps of algebra, the above equation yields

$$\begin{aligned} r_j &= \frac{\sum_{x=j+1}^n f(x) + p \sum_{x=1}^j f(x)}{\prod_{k=1}^{j-1} r_k} \\ &= \frac{\sum_{x=j+1}^n f(x) + p \left(1 - \sum_{x=j+1}^n f(x) \right)}{\prod_{k=1}^{j-1} r_k} \\ &= \frac{p + (1-p) \sum_{x=j+1}^n f(x)}{\prod_{k=1}^{j-1} r_k} \end{aligned} \tag{4.15}$$

In addition, we note that $r_1 r_2 \dots r_j = p$.

At the end of the waiting phase of the final back-off stage, the request will either be transmitted successfully or discarded by the system. That is, if the system is in state $(m, -n)$ and it does not receive a grant while it is there, then the packet transmission request is denied altogether, and the system goes back to the *Idle* state. After entering the *Idle* state, the SS starts a new back-off process if there is a packet(s) waiting in the queue. Otherwise, it continuously checks for new packet arrivals and goes to the *Active* state once a new packet arrives.

At steady-state, we can write

$$b_{i,k} = \begin{cases} r_n b_{i-1,-n}/W_i, & k = W_i - 1 \\ b_{i,k-1} + r_n b_{i-1,-n}/W_i, & k = 0, 1, \dots, W_i - 2 \end{cases} \quad (4.16)$$

However, $r_n b_{i-1,-n} = b_{i,0} \cdot (r_1 \dots r_n) = p b_{i,0}$. Hence,

$$b_{i,k} = \frac{p(W_i - k)}{W_i} b_{i-1,0} \quad (4.17)$$

This implies

$$b_{i,0} = p b_{i-1,0} = p^i b_{0,0}. \quad (4.18)$$

Combining (4.17) and (4.18), we get

$$b_{i,k} = \frac{p^i(W_i - k)}{W_i} b_{0,0}. \quad (4.19)$$

The normalisation condition of the Markov chain requires

$$\sum_{i=0}^m \left(\sum_{k=0}^{W_i-1} b_{i,k} + \sum_{k=-1}^{-n} b_{i,k} \right) = 1. \quad (4.20)$$

Consider the first term of (4.20):

$$\sum_{i=0}^m \sum_{k=0}^{W_i-1} b_{i,k} = \sum_{i=0}^m b_{0,0} \sum_{k=0}^{W_i-1} (W_i - k) \frac{p^i}{W_i}$$

Recalling $W_i = 2^i W_0$, and assuming $\alpha = (1 - p^{m+1})/(1 - p)$ and $\beta = \{1 - (2p)^{m+1}\}/(1 - 2p)$ in above, we get

$$\sum_{i=0}^m \sum_{k=0}^{W_i-1} b_{i,k} = \frac{b_{0,0}}{2} (W_0 \beta + \alpha). \quad (4.21)$$

Consider the second term of (4.20):

$$\begin{aligned} \sum_{i=0}^m \sum_{k=-1}^{-n} b_{i,k} &= \sum_{i=0}^m b_{i,0} (1 + r_1 + r_1 r_2 + \dots + r_1 r_2 \dots r_{j-1}) \\ &= b_{0,0} \sum_{i=0}^m p^i \sum_{j=1}^n \left\{ p + (1-p) \sum_{k=j}^n f(k) \right\} \\ &= b_{0,0} \alpha \left\{ np + (1-p) \sum_{j=1}^n j f(j) \right\} \\ &= b_{0,0} \alpha \{ np + (1-p) \mu \} \end{aligned} \quad (4.22)$$

Combining (4.21) and (4.22), we get

$$b_{0,0} \{ 0.5 W_0 \beta + \alpha [0.5 + np + (1-p) \mu] \} = 1. \quad (4.23)$$

While in *Active* state, the SS transmits only when its back-off counter reaches zero, i.e. in the $(i, 0)$ states. Hence, the probability τ that a station transmits in a ranging channel is given by

$$\tau = q \sum_{i=0}^m b_{i,0}. \quad (4.24)$$

Using (4.18) and (4.23) in above, we get

$$\tau = \frac{q \alpha}{0.5 W_0 \beta + \alpha [0.5 + np + (1-p) \mu]}. \quad (4.25)$$

Since τ is a function of p , we need to determine τ to find the value of p . However, (4.24) implies that in order to determine τ , we first need to find the value of q .

4.3.4 Queuing Analysis

The aim of this subsection is to determine the probability q that the Markov Chain in Figure 4.1 remains in *Active* state. Here, we consider an infinite SS queue length as per *Assumption 4* in Subsection .

Let, a SS spends $\mathbf{t}$ amount of time to complete the random access procedure. While the SS is busy serving a request, additional requests may arrive at the SS queue. Let the random variable $\mathbf{m}$ denotes the number of requests arrived within a time interval of $\mathbf{t}$ and the random variable $\mathbf{r}$ denotes the number of requests waiting to be processed at a given time instant.

We further denote the offered load as

$$\rho := \lambda E(\mathbf{t}), \quad (4.26)$$

where $E(\mathbf{t})$ is the mean back-off service time. In other words, it is the mean random access delay of the system.

Now, based on the preliminaries in Subsection 4.3.2, we make the following propositions:

Proposition 1. The characteristic function of $\mathbf{m}$ is given by

$$C_{\mathbf{m}}(z) = C_{\mathbf{t}}[\exp\{\lambda(z-1)\}]. \quad (4.27)$$

In addition,

$$\dot{C}_{\mathbf{m}}(1) = \rho.$$

Proof. Let $P_{\mathbf{m}|\mathbf{t}}(\ell|t)$ denotes the conditional probability that $\mathbf{m} = \ell$ given $\mathbf{t} = t$. Using the total probability theorem, we have

$$P_{\mathbf{m}}(\ell) = \sum_{t=0}^{\infty} P_{\mathbf{t}}(t) P_{\mathbf{m}|\mathbf{t}}(\ell|t). \quad (4.28)$$

However, $P_{\mathbf{m}|\mathbf{t}}(\ell|t)$ is same as the probability that ℓ requests arrive within an interval t .

Hence, by Markovian assumption on the arrival process we have

$$P_{\mathbf{m}|\mathbf{t}}(\ell|t) = \frac{(\lambda t)^\ell}{\ell!} \exp(-\lambda t). \quad (4.29)$$

Combining (4.28) and (4.29) according to (4.9), we get

$$\begin{aligned} C_{\mathbf{m}}(z) &= \sum_{\ell=0}^{\infty} P_{\mathbf{m}}(\ell) z^\ell = \sum_{t=0}^{\infty} P_{\mathbf{t}}(t) \sum_{\ell=0}^{\infty} \frac{(\lambda t z)^\ell}{\ell!} \exp(-\lambda t) \\ &= \sum_{t=0}^{\infty} P_{\mathbf{t}}(t) \exp\{\lambda t(z-1)\} = \sum_{t=0}^{\infty} P_{\mathbf{t}}(t) [\exp\{\lambda(z-1)\}]^t \\ &= C_{\mathbf{t}}[\exp\{\lambda(z-1)\}]. \end{aligned}$$

Now differentiate both sides of the above to get

$$\begin{aligned} \dot{C}_{\mathbf{m}}(z) &= \lambda \exp\{\lambda(z-1)\} \dot{C}_{\mathbf{t}}[\exp\{\lambda(z-1)\}], \\ \Rightarrow \dot{C}_{\mathbf{m}}(1) &= \lambda \dot{C}_{\mathbf{t}}(1) \end{aligned}$$

Using preliminaries (4.8) and (4.9), it can be readily verified that $\dot{C}_{\mathbf{t}}(1) = E(\mathbf{t})$. Using this and (4.26) in above we can write

$$\dot{C}_{\mathbf{m}}(1) = \lambda E(\mathbf{t}) = \rho.$$

Proposition 2. The characteristic function of $\mathbf{r}$ is given by

$$C_{\mathbf{r}}(z) = \frac{(1-\rho)(1-z)}{1-z/C_{\mathbf{m}}(z)}. \quad (4.30)$$

In addition,

$$P_{\mathbf{r}}(0) = C_{\mathbf{r}}(0) = 1 - \rho.$$

Proof. Let, $\mathbf{r}_i$ denotes the number of pending requests immediately after the i^{th} request is served. By stationary $P_{\mathbf{r}_i}(\ell) = P_{\mathbf{r}}(\ell)$ for all i sufficiently large. If it turns out that $j > 0$ transmission requests are pending immediately after the $(i-1)^{th}$ request is served, i.e. $\mathbf{r}_{i-1} = j > 0$, then the SS queue will have no idle time. It will start serving the i^{th} request

immediately. If ℓ requests arrive while the SS serves the i^{th} request, then $\mathbf{r}_i$ will be $j - 1 + \ell$. The probability of this event is $P_{\mathbf{m}}(\ell)$. Therefore, $\forall j > 0$,

$$P_{\mathbf{r}_i|\mathbf{r}_{i-1}}(j-1+\ell|j) = \begin{cases} P_{\mathbf{m}}(\ell), & \ell \geq 0, \\ 0 & \ell < 0. \end{cases}$$

Hence $\forall j > 0$, using the preliminary (4.12), we can write

$$C_{\mathbf{r}_i|\mathbf{r}_{i-1}}(z|j) = z^{j-1} C_{\mathbf{m}}(z). \quad (4.31)$$

On the other hand, if it turns out that there are no requests pending immediately after the $(i-1)^{th}$ request is served, then the SS will sit idle until the next request arrives. In this case, $\mathbf{r}_i = \ell$ only if ℓ requests arrive while the i^{th} request is being served and the probability of that is $P_{\mathbf{m}}(\ell)$. Therefore,

$$P_{\mathbf{r}_i|\mathbf{r}_{i-1}}(\ell|0) = P_{\mathbf{m}}(\ell), \Rightarrow C_{\mathbf{r}_i|\mathbf{r}_{i-1}}(z|0) = C_{\mathbf{m}}(z). \quad (4.32)$$

If a steady state is reached, then $C_{\mathbf{r}}(z) = C_{\mathbf{r}_i}(z)$. Hence, by total probability theorem we have

$$C_{\mathbf{r}}(z) = C_{\mathbf{r}_i}(z) = \sum_{j=0}^{\infty} P_{\mathbf{r}_{k-1}}(j) C_{\mathbf{r}_i|\mathbf{r}_{i-1}}(z|j).$$

Using (4.31) and (4.32) in above, we can write

$$C_{\mathbf{r}}(z) = P_{\mathbf{r}}(0) C_{\mathbf{m}}(z) + \sum_{j=1}^{\infty} P_{\mathbf{r}}(j) z^{j-1} C_{\mathbf{m}}(z).$$

After re-arranging, we get

$$\begin{aligned} \frac{C_{\mathbf{r}}(z)}{P_{\mathbf{r}}(0)} &= C_{\mathbf{m}}(z) \left\{ 1 + \frac{1}{z} \sum_{j=1}^{\infty} \frac{P_{\mathbf{r}}(j)}{P_{\mathbf{r}}(0)} z^j \right\} \\ &= C_{\mathbf{m}}(z) \left\{ 1 + \frac{1}{z} \left(\frac{C_{\mathbf{r}}(z)}{P_{\mathbf{r}}(0)} - 1 \right) \right\}. \end{aligned}$$

After a few steps of algebra the above yields

$$\frac{C_{\mathbf{r}}(z)}{P_{\mathbf{r}}(0)} = \frac{z-1}{z/\{C_{\mathbf{m}}(z)-1\}}. \quad (4.33)$$

Next we find $P_{\mathbf{r}}(0)$. By differentiating the above, we get

$$\begin{aligned} \frac{\dot{C}_{\mathbf{r}}(z)}{P_{\mathbf{r}}(0)} \{z - C_{\mathbf{m}}(z)\} + \frac{2C_{\mathbf{r}}(z)}{P_{\mathbf{r}}(0)} \{1 - \dot{C}_{\mathbf{m}}(z)\} \\ = \dot{C}_{\mathbf{m}}(z) \{z-1\} + 2\dot{C}_{\mathbf{m}}(z). \end{aligned}$$

By setting $z = 1$ in above and considering the fact that $C_{\mathbf{m}}(1) = 1$ as per (4.10), we get

$$\frac{1}{P_{\mathbf{r}}(0)} \{1 - \dot{C}_{\mathbf{m}}(1)\} = 1.$$

But *Proposition 1* yields $\dot{C}_{\mathbf{m}}(1) = \rho$. Hence, we have

$$P_{\mathbf{r}}(0) = 1 - \rho. \quad (4.34)$$

By combining (4.33) and (4.34), we get

$$C_{\mathbf{r}}(z) = \frac{(1-\rho)(1-z)}{1-z/C_{\mathbf{m}}(z)}.$$

Note that $P_{\mathbf{r}}(0)$ is the probability that the queue is empty and it is the probability that the system is idle. Thus, using (4.26) we can write

$$q = 1 - P_{\mathbf{r}}(0) = \lambda E(\mathbf{t}). \quad (4.35)$$

From (4.35) we see that in order to determine q , we first need to find the value of $E(\mathbf{t})$.

4.3.5 Back-off Service Time Analysis

In this subsection, we derive the mean back-off service time $E(\mathbf{t})$ of the system. The analysis is carried out in the following two steps:

i) Delay in the i^{th} back-off Stage: Let the back-off stage i be associated with a binary random variable s_i such that s_i takes a value 1 when a request is transmitted successfully; and otherwise $s_i = 0$. It is readily verified that

$$P_{s_i}(0) = p, \quad P_{s_i}(1) = 1 - p. \quad (4.36)$$

Let $\mathbf{d}_i$ be the time that the process needs to spend in the back-off stage i . This delay consists of two independent components, i.e.

$$\mathbf{d}_i = \mathbf{b}_i + \delta. \quad (4.37)$$

The first component $\mathbf{b}_i$ denotes the time spent in the back-off waiting phase, and the second component δ denotes the time spent in the grant waiting phase in Figure 4.1. Note that,

$$P_{\mathbf{b}_i}(k) = \begin{cases} 1/W_i, & k \in \{0, 1, \dots, W_i - 1\} \\ 0, & \text{Otherwise} \end{cases} \quad (4.38)$$

In terms of characteristics function

$$C_{\mathbf{b}_i}(z) = \frac{1}{W_i} \sum_{k=0}^{W_i-1} z^k. \quad (4.39)$$

In addition,

$$E(\mathbf{b}_i) = \frac{1}{W_i} \{0 + 1 + \dots + W_i - 1\} = \frac{W_i - 1}{2},$$

$$E(\mathbf{b}_i^2) = \frac{1}{W_i} \{0^2 + 1^2 + \dots + (W_i - 1)^2\} = \frac{(W_i - 1)(2W_i - 1)}{6}.$$

If the transmission attempt is unsuccessful, i.e. $s_i = 0$, then $\delta = n$, where n is the T_{16} time-out interval. This yields

$$E(\delta | s_i = 0) = n, \quad E(\delta^2 | s_i = 0) = n^2.$$

Also, using (4.12) we have

$$C_{\delta | s_i}(z | 0) = z^n.$$

The conditional distribution $P_{\delta|s_i}$ of a given successful transmission, i.e. $s_i = 1$ depends on the probability mass function $f(\mathbf{x})$, as defined in *Assumption 7*. Thus, we can write

$$E(\delta|s_i = 1) = \mu, \quad E(\delta^2|s_i = 1) = \sigma^2,$$

where $\mu = E(\mathbf{x}) = \sum_{\mathbf{x}=1}^n \mathbf{x}f(\mathbf{x})$ and $\sigma^2 = \text{Var}(\mathbf{x})$.

We further assume:

$$C_{\delta|s_i}(z|1) = D(z).$$

Now, based on (4.37) and using the preliminary (4.14), we find the conditional distribution and mean of $\mathbf{d}_i$ for a successful code transmission attempt from stage i as

$$C_{\mathbf{d}_i|s_i}(z|1) = C_{\mathbf{b}_z}(z)C_{\delta|s_i}(z|1) = C_{\mathbf{b}_z}(z)D(z). \quad (4.40)$$

Also,

$$\begin{aligned} E\{\mathbf{d}_i|s_i = 1\} &= \frac{W_i - 1}{2} + \mu, \\ E\{\mathbf{d}_i^2|s_i = 1\} &= \frac{(W_i - 1)(2W_i - 1)}{6} + \sigma^2. \end{aligned} \quad (4.41)$$

Similarly, the conditional distribution and mean of $\mathbf{d}_i$ for a failed code transmission attempt from stage i are given by

$$C_{\mathbf{d}_i|s_i}(z|0) = C_{\mathbf{b}_z}(z)C_{\delta|s_i}(z|0) = C_{\mathbf{b}_z}(z)z^n. \quad (4.42)$$

Also,

$$\begin{aligned} E\{\mathbf{d}_i|s_i = 0\} &= \frac{W_i - 1}{2} + n, \\ E\{\mathbf{d}_i^2|s_i = 0\} &= \frac{(W_i - 1)(2W_i - 1)}{6} + n^2. \end{aligned} \quad (4.43)$$

ii) Service Time of the Back-off Process: Define a binary random variable $\mathbf{s}$ such that $\mathbf{s}$ takes a value 1 when a code is transmitted successfully; and otherwise $\mathbf{s} = 0$. Note that the number of back-off stages is $m + 1$ and they are indexed as $0, 1, 2, \dots, m$. Now, let us define another random variable $\mathbf{n}$ such that

- $\mathbf{n} = i \leq m$ if a code is transmitted successfully from back-off stage $i < m$.
- $\mathbf{n} = m + 1$ if the all the attempts for transmitting the code fail.

First, we wish to find the conditional distribution of the total delay $\mathbf{t}$ given that a request is transmitted successfully from back-off stage i . The total delay is

$$\mathbf{t} = \mathbf{d}_0 + \mathbf{d}_1 + \dots + \mathbf{d}_i.$$

Since the delay incurred in a back-off stage is independent of the delay incurred in the other stages, we have

$$C_{\mathbf{t}|\mathbf{n}}(z|k) = C_{\mathbf{d}_i|\mathbf{s}_i}(z|1) \prod_{j=0}^{i-1} C_{\mathbf{d}_j|\mathbf{s}_j}(z|0).$$

Using (4.40) and (4.42), and then (4.39) in above, we get

$$C_{\mathbf{t}|\mathbf{n}}(z|k) = z^{ni} D(z) \prod_{j=0}^i \left\{ \sum_{\ell=0}^{W_j-1} \frac{z^\ell}{W_j} \right\}. \quad (4.44)$$

In addition,

$$E(\mathbf{t}|\mathbf{n} = i) = ni + \mu + \sum_{j=0}^i \frac{W_j - 1}{2}, \quad (4.45)$$

and

$$E(\mathbf{t}^2|\mathbf{n} = i) = ni + \mu + \sum_{j=0}^i \frac{(W_j - 1)(2W_j - 1)}{6}. \quad (4.46)$$

Next we find the conditional distribution of $\mathbf{n}$ given that $\mathbf{s} = 1$, that is, the code is transmitted successfully. Recalling (4.36) for $i \leq m$, we have

$$P_{\mathbf{n}}(k) = P_{\mathbf{s}_i}(1) \prod_{j=0}^{i-1} P_{\mathbf{s}_j}(0) = p^i (1 - p).$$

Now $\mathbf{s} = 1$ if and only if $\mathbf{n} \leq m$. Hence,

$$P_{\mathbf{s}}(1) = \sum_{i=0}^m P_{\mathbf{n}}(k) = (1-p) \sum_{i=0}^m p^i = 1 - p^{m+1}.$$

Clearly, the probability that all m attempts to transmit a code fails is $P_{\mathbf{s}}(0) = 1 - P_{\mathbf{s}}(1) = p^{m+1}$. Hence, we have

$$P_{\mathbf{n}|\mathbf{s}}(i|1) = \frac{p^i(1-p)}{1-p^{m+1}} = p^i \left(\sum_{i=0}^m p^i \right)^{-1}.$$

Thus, the conditional characteristic function of the delay $\mathbf{t}$ incurred by a request given that it has been transmitted successfully is

$$C_{\mathbf{t}|\mathbf{s}}(z|1) = \left(\sum_{i=0}^m p^i \right)^{-1} \sum_{i=0}^m p^i C_{\mathbf{t}|\mathbf{n}}(z|i). \quad (4.47)$$

Using (4.44) in above, we get

$$C_{\mathbf{t}|\mathbf{s}}(z|1) = \left(\sum_{i=0}^m p^i \right)^{-1} D(z) \sum_{i=0}^m p^i z^{ni} \prod_{j=0}^i \left\{ \sum_{\ell=0}^{W_j-1} \frac{z^\ell}{W_j} \right\}. \quad (4.48)$$

In addition, the mean delay experienced by a successfully transmitted code is given by

$$E(\mathbf{t}|\mathbf{s} = 1) = \left(\sum_{i=0}^m p^k \right)^{-1} \sum_{i=0}^m p^i E(\mathbf{t}|\mathbf{n} = 1).$$

Using (4.45) in above, we get

$$E(\mathbf{t}|\mathbf{s} = 1) = \mu + \left(\sum_{i=0}^m p^i \right)^{-1} \sum_{i=0}^m p^i \left\{ ni + \sum_{j=0}^{W_j-1} \frac{W_j - 1}{2} \right\}. \quad (4.49)$$

Also,

$$E(\mathbf{t}^2|\mathbf{s} = 1) = \left(\sum_{i=0}^m p^i \right)^{-1} \sum_{i=0}^m p^i E(\mathbf{t}^2|\mathbf{n} = 1).$$

Using (4.46) in above, we get

$$E(\mathbf{t}^2|\mathbf{s} = 1) = \sigma^2 + \left(\sum_{i=0}^m p^i \right)^{-1} \cdot \sum_{i=0}^m p^i \left\{ n^2 i + \sum_{j=0}^{W_j-1} \frac{(W_j - 1)(2W_j - 1)}{6} \right\}. \quad (4.50)$$

Now, similar to (4.47), the conditional characteristic function of the delay $\mathbf{t}$ incurred by a failed request is

$$C_{\mathbf{t}|\mathbf{s}}(z|0) = \prod_{i=0}^m C_{\mathbf{d}_i|\mathbf{s}_i}(z|0) = z^{nm} \prod_{i=0}^m \left\{ \sum_{\ell=0}^{W_j-1} \frac{z^\ell}{W_j} \right\}. \quad (4.51)$$

In addition, the mean delay incurred by a failed is given by

$$E(\mathbf{t}|\mathbf{s} = 0) = mn + \sum_{i=0}^m \frac{W_i - 1}{2}. \quad (4.52)$$

Also,

$$E(\mathbf{t}^2|\mathbf{s} = 0) = mn^2 + \sum_{i=0}^m \frac{(W_i - 1)(2W_i - 1)}{6}. \quad (4.53)$$

Combining above formulae, we have

$$C_{\mathbf{t}}(z) = (1 - p^m)C_{\mathbf{t}|\mathbf{s}}(z|1) + p^m C_{\mathbf{t}|\mathbf{s}}(z|0).$$

Using (4.48) and (4.51) in above, we get

$$C_{\mathbf{t}}(z) = (1 - p)D(z) \sum_{i=0}^m p^i z^{ni} \prod_{j=0}^i \left\{ \sum_{\ell=0}^{W_j-1} \frac{z^\ell}{W_j} \right\} C_{\mathbf{t}}(z) + p^i z^{nm} \prod_{i=0}^m \left\{ \sum_{\ell=0}^{W_j-1} \frac{z^\ell}{W_j} \right\}, \quad (4.54)$$

while the mean back-off service time is

$$E(\mathbf{t}) = P_s(0)E(\mathbf{t}|\mathbf{s} = 0) + P_s(1)E(\mathbf{t}|\mathbf{s} = 1).$$

Using (4.49) and (4.52) in above and after some simplifications, we get

$$E(\mathbf{t}) = \sum_{i=0}^m p^i \left\{ \frac{W_i - 1}{2} + n \cdot i + (1 - p)\mu \right\}. \quad (4.55)$$

Similarly,

$$E(\mathbf{t}^2) = P_s(1)E(\mathbf{t}^2|\mathbf{s} = 0) + P_s(1)E(\mathbf{t}^2|\mathbf{s} = 1).$$

Using (4.50) and (4.53) in above and after some simplifications, we get

$$\begin{aligned} E(\mathbf{t}^2) &= (1 - p^{m+1})\sigma^2 + \frac{pn^2(1 - p^{m+1})}{1 - p} \\ &\quad + \sum_{i=0}^m p^i \frac{(W_i - 1)(2W_i - 1)}{6}. \end{aligned} \quad (4.56)$$

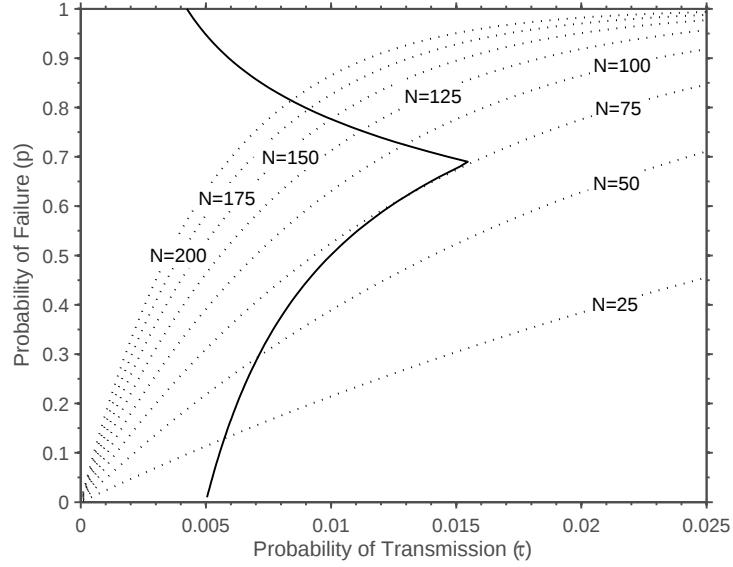
4.3.6 Access Probability Analysis

In this section, we calculate the access probability of the system in terms of the conditional probability p that a code transmission fails. To do so, we exploit the fact that for a successful code transmission, maximum D stations can transmit with different codes out of K number of available codes. Thus, we get

$$p = 1 - \sum_{i=0}^{D-1} \binom{N-i}{i} \tau^i (1 - \tau)^{N-1-i} \left(1 - \frac{1}{K}\right)^i. \quad (4.57)$$

Equation (4.25) and (4.57) present a non-linear system that can be solved for p in the range of $0 \leq p \leq 1$ using numerical techniques. To verify whether the system has a unique solution, we plot p versus τ for different values of N in Figure 4.2. The plot shows that both p and τ are continuous functions having only one intersection for a given value of N , which represents the unique solution for the system.

Before moving on to the next section, we introduce two performance indicators for a random access channel. First, the *normalised access load* is defined as the average number of access requests over a given ranging channel. Since we assumed a mean request arrival rate of λ per frame per SS, the normalised load is $N\lambda$ per random access channel. Second, the *normalised access throughput* is defined as the average number of successfully transmitted

Figure 4.2: τ vs. p for different values of N .

codes over a given ranging channel. It is obtained from the conditional probability that a given SS transmits a code and the code is successfully received by the BS, which is given by

$$\theta = N\tau(1 - p). \quad (4.58)$$

4.4 Model Validation and Performance Analysis

First, we evaluate the code detection performance of the WiMAX random access plane. To do this, we develop a Monte-Carlo simulation model using MATLAB. The system is assumed to be operating at 2.3 GHz with a bandwidth of 5 MHz and a sampling frequency of 5.6 MHz. The FFT size is assumed to be 512 with a cyclic-prefix size of 128 samples. To simulate multipath channels, the Stanford University Interim (SUI) channel model 3 and 4 are used, as recommended by the IEEE 802.16p evolution methodology document (EMD) [80]. Each of these models has three taps as listed in Table 4.1 (considering an omnidirectional BS). The first tap is modelled as a Rician-fading channel and the remaining two are modelled as Rayleigh-fading channels, with a mean value described in Table 4.1. Additionally, the tap delays are varied based on a Chi-square distribution to ensure that

Table 4.1: Multipath Channel Parameters

Parameter	SUI-3 Model			SUI-4 Model		
	<i>Tap-1</i>	<i>Tap-2</i>	<i>Tap-3</i>	<i>Tap-1</i>	<i>Tap-2</i>	<i>Tap-3</i>
Delay (μ s)	0	0.4	0.9	0	1.5	4
Power (dB)	0	-5	-10	0	-4	-8
Doppler	0.4	0.3	0.5	0.2	0.15	0.25
K-factor	1	0	0	0	0	0

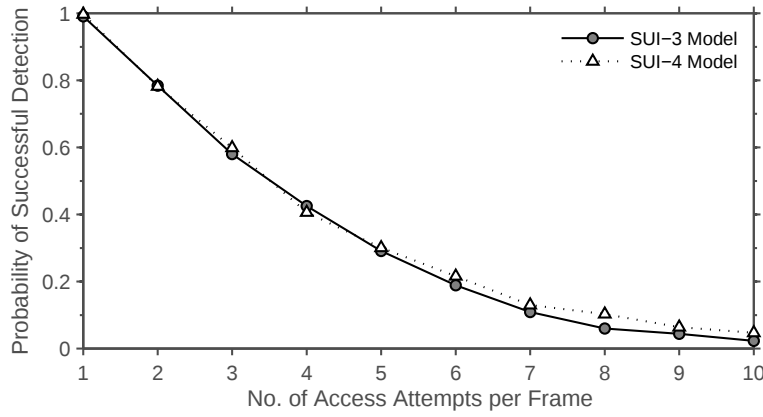


Figure 4.3: Probability of successful code detection at varying access load under SUI-3 and SUI-4 multipath channel models.

each user experiences a different channel in every WiMAX frame. A raised-cosine filter is used for OFDM pulse-shaping with a roll-off factor of 0.22. Moreover, the detection performance is measured at a SNR level of 10 dB.

The corresponding results are shown in Figure 4.3. From the results, we see that the code detector in the BS can only detect a few codes under the existing WiMAX random access scheme. For example, under the SUI-3 and 4 channel models, the code detector can detect only a single code with more than 90 per cent detection probability. Note that the detection performance can vary depending on the detection threshold used, which is outside the scope of this work.

Next, we validate our proposed analytical model by comparing numerical and simulation results. The simulation model is developed using the OPNET Modeller 16.0 based on the same network parameters specified previously. Figure 4.4-4.7 plots the simulation results

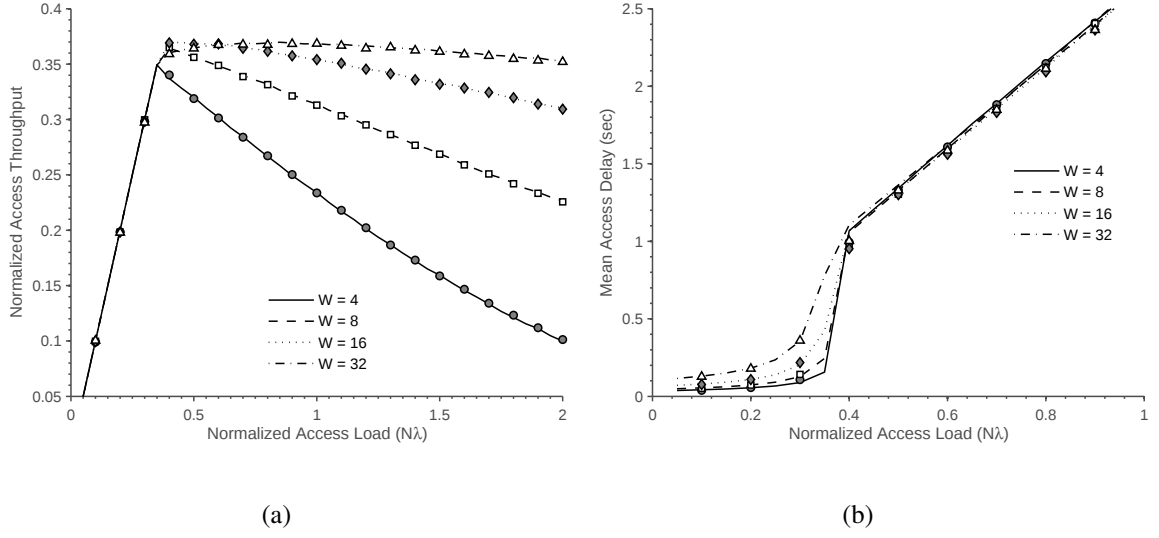


Figure 4.4: WiMAX random access performance under different initial back-off windows.

(dots) versus the numerical results (lines) in terms of mean random access throughput and mean random access delays at varying loads under different back-off and channel parameters. The figures clearly show a good agreement between the numerical and the simulation results which confirms the accuracy of the analytical model.

Figure 4.4 shows the effect of initial back-off window sizes on the mean throughput and delay of the random access channel under the fixed parameters $m = 8$, $n = 8$, $D = 1$. From Figure 4.4(a), we see that for all back-off window sizes, the throughput follows the load until the saturation point after which it degrades with further increase in the load. From Figure 4.4(b), we see that after the saturation point, the random access delay starts to increase abruptly as the SS queues become unstable: they are never empty (i.e. $q \rightarrow 1$). As a result, the requests experience higher delays and the throughput starts to drop, since some requests reach their retry limits. Note that the larger the initial back-off window, the longer the throughput remains stable after the saturation point, as the code transmission attempts are randomised over a larger interval. However, this achievement comes at a cost of higher delays, especially during the unsaturated condition.

Figure 4.5 shows the effect of different retry limits under the fixed parameters $W = 2^3$, $n = 8$, $D = 1$. From the results, we see that increasing the retry limit improves the throughput performance only after the saturation point, while the delay performance remains almost

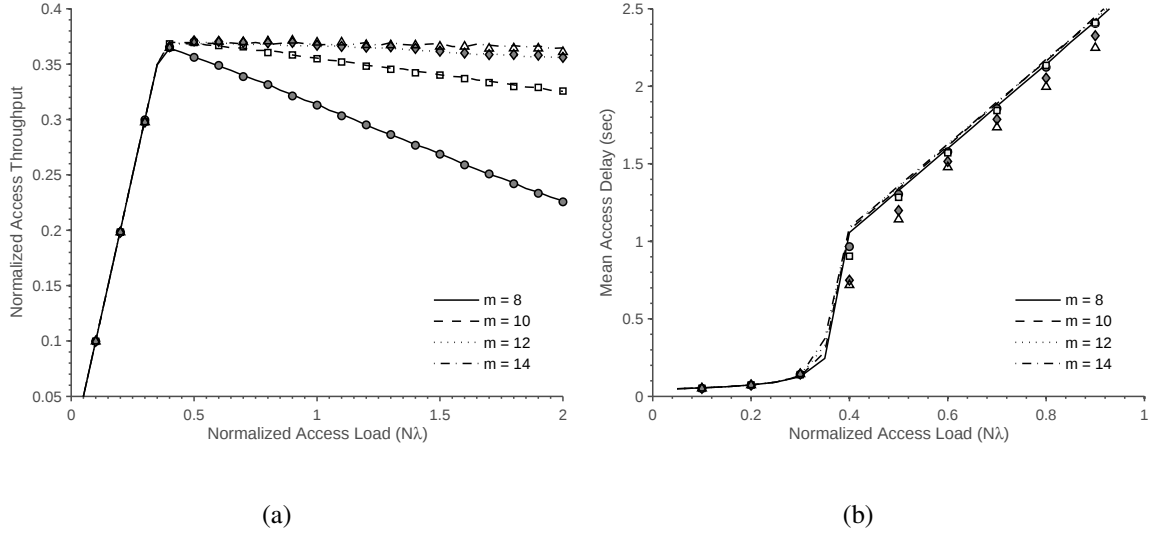


Figure 4.5: WiMAX random access performance under different back-off retry limits.

unchanged.

Next, we increase the number of detectable codes by the BS to two (i.e. $D = 2$) and vary the number of available codes K to the SSs for fixed parameters $W = 2^3$, $m = 8$, $n = 8$. The corresponding results are plotted in Figure 4.6. Comparing Figures 4.4 and 4.6, we see that when the number of available codes is high, doubling the detectable codes almost doubles the throughput performance. However, the delay remains the same. Hence, improved code detectability can significantly improve the random access performance of the network.

Lastly, we examine the relation between the available codes K and the detectable codes D . Note that while D remains fixed under a given channel condition, K can be adjusted by the system operator. However, the extent to which K can be varied is restricted, as only a limited number of codes are available for a particular random access channel.

Figure 4.7 plots the random access throughput and delay for three detectable codes (i.e. $D = 3$) with a varying number of available codes. The corresponding results are consistent with those obtained from Figure 4.6 with two detectable codes. To interpret the results, we define the parameter sustainable access throughput θ_s as the maximum throughput of the channel before entering saturation. From both Figures 4.6 and 4.7, we see that the optimum value of θ_s is obtained when $K = 2D$.

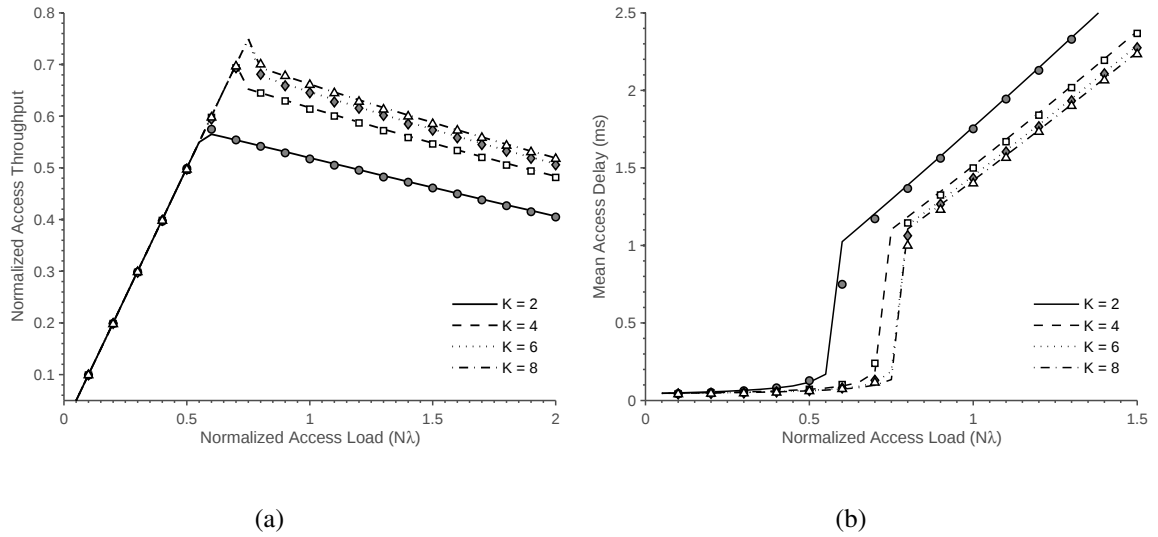


Figure 4.6: WiMAX random access performance under two detectable codes.

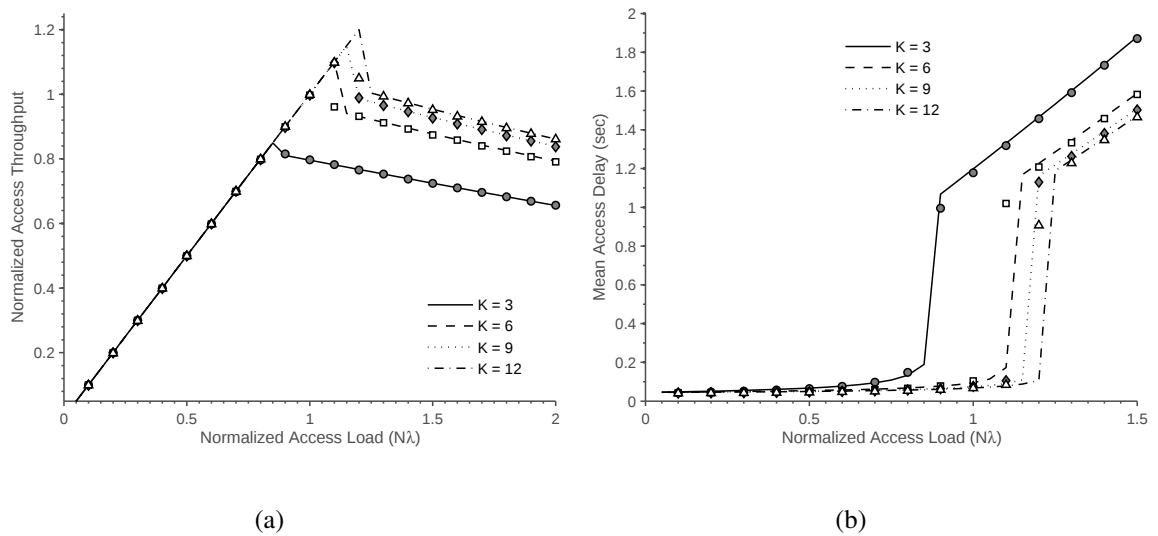


Figure 4.7: WiMAX random access performance under three detectable codes.

4.5 Chapter Summary

In this chapter, we presented a comprehensive analytical model to evaluate the performance of the CDMA-based random access procedure in WiMAX. The results obtained from the analytical model and matched by the simulation results verify the accuracy of the proposed model. Moreover, it shows that the model is able to accurately predict the transition of the random access plane from non-saturation to saturation mode under a range of contention and channel parameters. Since the proposed model is physical layer agnostic, it can be used to predict the performance of a random access channel for different code detection algorithms under varying channel conditions.

Based on this model, in the next chapter we develop a number of techniques to enhance the performance of the WiMAX random access plane for M2M communications in the Smart Grid.

Chapter 5

Random Access Enhancement

From simulation studies in Chapter 3 and the results from the analytical model in Chapter 4, we observe that the unique features of M2M applications pose the following two key challenges to the existing WiMAX random access mechanism.

First, the random access channels perform optimally when the number of devices that are contending simultaneously does not exceed a particular threshold. Otherwise, the channel becomes unstable in terms of access success rate and access delay. Moreover, apart from the high perennial random access load, the number of simultaneously contending devices may increase rapidly due to a certain fault/outage event [81]. In turn, this may congest the whole random access plane, resulting in heavy packet loss and prolonged delay.

Second, in the existing IEEE 802.16 standards, only a single QoS scheduling class, i.e. the BE service is associated with random access mechanism. Therefore, it is not possible to provide differentiated access service to various M2M devices/applications using the conventional BE service.

A simple solution to the first problem is to increase the number of BR channels to accommodate more devices and use separate channels to enforce traffic differentiation. However, from Figure 4.3 in the previous chapter, we see that under the existing schemes only a few codes can be detected reliably per channel per frame. Therefore, many BR channels are required, which would substantially reduce the payload capacity of the overall system. Moreover, allocating separate channels for different classes of traffic is not an efficient

solution as it fails to take the advantage of statistical multiplexing of BRs from different devices.

To address this issue, the recent IEEE 802.16p amendment for M2M communications has proposed several solutions, mostly based on access control over the MAC layer [9]. Moreover, several works have already appeared in the literature concerning this problem. Of the notable works, an adaptive slotted-ALOHA based access protocol is proposed in [82] to support event-driven M2M application. A CSMA/CA based MAC protocol is proposed in [83] that employs a contention period and a transmission period to support a large number of M2M devices. However, these schemes require significant modification to the existing standards and are not suitable for a network supporting both M2M and non-M2M traffic.

Since in a Smart Grid environment, most M2M devices are either fixed or have very low-mobility, their wireless channels are expected to experience only a small variation in time. The Doppler spectrum for such a channel has a rounded shape with zero-mean which yields a large coherence time [84]. Based on this unique feature of M2M traffic, in this chapter we propose an enhanced random access scheme, where the fixed/low-mobility M2M devices pre-equalise their BR codes using the estimated frequency response of their slowly-varying channels. Consequently, the BS is able to detect a large number of codes as their mutual orthogonality remains preserved. Mathematical analysis is conducted to determine the theoretical performance limit of the code detector. The analysis reveals that the default pseudo-random code matrix specified in the IEEE 802.16 standard is not quite effective for detecting a large number of codes under the proposed scheme. As a remedy, we argue that a Hadamard code matrix can significantly increase its code detection performance. Moreover, a DRA strategy is proposed to provide QoS-aware access service to various M2M devices. The theoretical performance of the proposed scheme is validated by simulation results under both of the two code matrices. Compared to the literature, our proposal is fully compliant with the existing WiMAX specifications, except that it requires a dedicated BR channel when both M2M and conventional applications need to be supported. Such a requirement is reasonable considering the volume of the M2M devices per BS and has already been provisioned in the IEEE 802.16p amendment.

Note that to provide guaranteed QoS to the bursty M2M traffic, the BR and grant mechanism must work cohesively. Hence, we further propose an adaptive RRM framework based on the above DRA strategy, which adaptively allocates the ranging resources based on load and priority of the application. The performance of the proposed scheme is demonstrated using both generic traffic profiles and the pilot protection application described in Chapter 3.

The rest of the chapter is organised as follows. Section 5.1 describes the system model and the key tenets of the proposed random access mechanism. Section 5.2 formulates the concept of DRA strategy. Section 5.3 presents the proposed radio resource scheduling framework and demonstrates its performance ¹.

5.1 The Proposed Random Access Scheme

Consider a single-cell WiMAX network with TDD-OFDMA physical layer [58]. Without loss of generality, we assume that there is only one BR random access channel in the uplink subframe. We consider the same system model described in the previous chapter, where the BR code-matrix is denoted as $\mathbf{C} \in \mathbb{R}^{L \times K}$

The first OFDM symbol of each WiMAX frame is a preamble transmitted by the BS, where the subcarriers are BPSK modulated with a boosted pilot sequence [58]. Typically, the SSs use this information to estimate the channel frequency response (CFR) for the OFDM demodulation process. A SS can pre-equalise its BR code using this estimated CFR exploiting the channel reciprocity of the TDD system [88]. A number of pre-equalisation techniques are available. For more details, please refer to [89]. A simplified block diagram of the BR code transmission scheme with pre-equalisation is illustrated in Figure 5.1. As shown in the diagram, when the pre-equalised OFDM signal passes through the multipath wireless channel, the channel equalises the pre-distorted symbols; consequently, the output of the

¹The contents of this chapter has been published in two conference papers and a journal paper: one in proceedings of the *IEEE/IET International Symposium on Communication Systems, Networks, and Digital Signal Processing (CSNDSP)* in July 2014 [85]; one to be published in proceedings of the *IEEE International Conference on Smart Grid Communications (SmartGridComm)* in Nov 2014 [86]; and another in the vol. 2, no. 1 of the journal *Recent Patents on Telecommunication* in Oct 2013 [87].

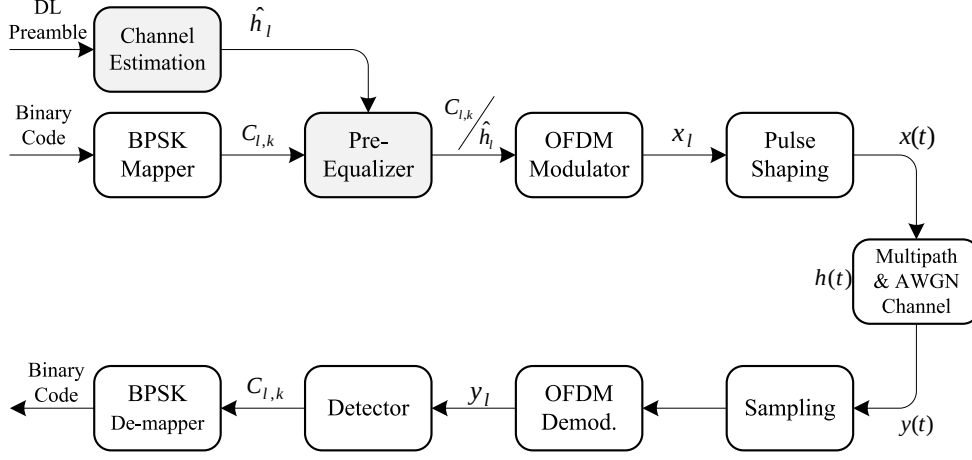


Figure 5.1: Simplified block diagram of BR code transmission with pre-equalisation.

OFDM receiver will simply be the original BPSK modulated BR code.

Now, consider a time instant where M number of SSs are simultaneously contending on the same BR channel. Let the m^{th} SS selects the k_m^{th} column of the code matrix $\mathbf{C} \in \mathbb{R}^{L \times K}$, where $m = 1, 2, \dots, M$. Considering zero-forcing (ZF) pre-equalisation [89], the transmitted code over the l^{th} subcarrier from the m^{th} SS be

$$x_{l,m} = \frac{\mathbf{C}_{l,k_m}}{\hat{h}_{l,m}} \text{ for } \forall l \in \{1, 2, \dots, L\}, \quad (5.1)$$

where $\hat{h}_{l,m}$ is the pilot-aided CFR estimated by the m^{th} SS. To be more precise, $\hat{h}_{l,m}$ can be expressed as

$$\hat{h}_{l,m} = h_{l,m} + e_{l,m}, \quad (5.2)$$

where $h_{l,m}$ is the actual frequency response of the l^{th} subcarrier and $e_{l,m} \sim \mathcal{CN}(0, \sigma_{e,m}^2)$ is a zero-mean complex Gaussian noise for the estimated channel response $\hat{h}_{l,m}$.

In the BS, after down-conversion to baseband and OFDM demodulation, the received signal from the m^{th} SS over the l^{th} subcarrier is

$$y_{l,m} = x_{l,m} \bar{h}_{l,m} = \mathbf{C}_{l,k_m} \frac{\bar{h}_{l,m}}{\hat{h}_{l,m}} \text{ for } \forall l \in \{1, 2, \dots, L\}, \quad (5.3)$$

where $\bar{h}_{l,m}$ is the effective channel experienced by the SS.

Note that although the positions of the BS and the SSs remain fixed for the stationary M2M devices, the channel is continually affected by the movement of external scatterers in the surrounding environment [84]. Consequently, $\bar{h}_{l,m}$ will differ from $h_{l,m}$. To generalise, assume the effective channel can be modelled as [90]:

$$\bar{h}_{l,m} = \alpha h_{l,m} + \eta_{l,m}, \quad (5.4)$$

where $\eta_{l,m} \sim \mathcal{CN}(0, \sigma_{\eta,m}^2)$ is a zero-mean complex Gaussian noise for the effective channel response $\bar{h}_{l,m}$ and α is some deterministic complex valued constant.

Considering the above phenomena, the combined received signal at BS from all M stations over the l^{th} subcarrier be

$$y_l = \sum_{m=1}^M \left\{ \mathbf{C}_{l,k_m} \frac{\alpha h_{l,m} + \eta_{l,m}}{h_{l,m} + e_{l,m}} \right\} + \vartheta_l, \quad (5.5)$$

where $\vartheta_l \sim \mathcal{CN}(0, \sigma_{\vartheta}^2)$ is a zero-mean complex Gaussian noise for the combined received signal y_l .

Equation (5.5) can be represented as

$$y_l = \sum_{m=1}^M \{ \mathbf{C}_{l,k_m} - \lambda_{l,m} \mathbf{C}_{l,k_m} \} + \vartheta_l, \quad (5.6)$$

where $\lambda_{l,m}$ is a ratio of complex variables, i.e.

$$\lambda_{l,m} = \frac{(1 - \alpha)h_{l,m} + e_{l,m} + \eta_{l,m}}{h_{l,m} + e_{l,m}} \quad \forall l. \quad (5.7)$$

Let us construct:

$$\mathbf{y} = [y(1), y(2), \dots, y(L)]^T, \quad (5.8)$$

where $(\cdot)^T$ denotes transpose of a vector.

To detect the presence of a BR code, the code detector in the BS correlates the received signal with a matching matrix $\mathbf{D} \in \mathbb{R}^{L \times K}$. The detector takes every column of $\mathbf{D}$ and computes its cross-product with $\mathbf{y}$. The design principle of $\mathbf{D}$ should be such that: i) the correlation

between p^{th} column of $\mathbf{D}$, i.e. $\mathbf{D}_p$ and $\mathbf{C}_p$ is maximised, and ii) the correlation between $\mathbf{D}_p$ and $\mathbf{C}_{k_m}$ with indices $k_m \neq p$ is minimised. Let the detector picks a random column p denoted as $\mathbf{D}_p$, where $p = 1, 2, \dots, K$. The corresponding result of cross-product would be

$$\mathbf{D}_p^T \mathbf{y} = \sum_{l=1}^L \left\{ \sum_{m=1}^M \{ \mathbf{C}_{l,k_m} \mathbf{D}_{l,p} - \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{D}_{l,p} \} + \vartheta_l \mathbf{D}_{l,p} \right\}. \quad (5.9)$$

For code detection, the detector checks the probability density function (PDF) of $\mathbf{D}_p^T \mathbf{y}$ and decides in favour of one of the following two hypotheses: i) $\mathcal{H}_0$: the code $\mathbf{C}_p$ is not present in $\mathbf{y}$, and ii) $\mathcal{H}_1$: the code $\mathbf{C}_p$ is present in $\mathbf{y}$. The code detector applies this strategy individually for every $\mathbf{D}_p : p \in \{1, 2, \dots, K\}$ in (5.9).

In this work, we consider two different matching matrices. The first one is the default pseudo-random code matrix defined in the IEEE 802.16 standard, which provides K number of nearly orthogonal codes. The other one is generated from the partial Hadamard code matrix. The detection performance of these two matrices are analysed in the following subsections.

5.1.1 Pseudo-random Code Matrix

We consider $\mathbf{D} = \mathbf{C}$. Under this case, if $\mathcal{H}_0$ is true, then (5.9) follows

$$\mathbf{C}_p^T \mathbf{y} = \sum_{l=1}^L \sum_{m=1}^M \mathbf{C}_{l,k_m} \mathbf{C}_{l,p} - \sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{C}_{l,p} + \sum_{l=1}^L \vartheta_l \mathbf{C}_{l,p}. \quad (5.10)$$

Otherwise, under $\mathcal{H}_1$

$$\mathbf{C}_p^T \mathbf{y} = L - \sum_{l=1}^L \lambda_{l,m} + \sum_{l=1}^L \sum_{m=1, k_m \neq p}^M \mathbf{C}_{l,k_m} \mathbf{C}_{l,p} + \sum_{l=1}^L \sum_{m=1, k_m \neq p}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{C}_{l,p} + \sum_{l=1}^L \vartheta_l \mathbf{C}_{l,p}. \quad (5.11)$$

In order to conduct the hypothesis test, we need to obtain the probability density function of $\mathbf{C}_p^T \mathbf{y}$ under the hypothesis $\mathcal{H}_\ell$ for $\ell = 0, 1$. At first, we consider $\mathcal{H}_0$. Recall (5.6). If there was no channel estimation error, i.e. $\lambda = 0$, the whole energy of the users will be

only in the real part of $\mathbf{y}$, and the imaginary part will contain only noise. Moreover, in the R.H.S. of (5.10), the first term $\sum_{l=1}^L \sum_{m=1}^M \mathbf{C}_{l,k_m} \mathbf{C}_{l,p}$ is real-valued and its variance is much larger than the other two terms. Hence, we can consider the real part of $\mathbf{C}_p^T \mathbf{y}$ only.

By using the central limit theorem, the distribution of $\sum_{l=1}^L \sum_{m=1}^M \mathbf{C}_{l,k_m} \mathbf{C}_{l,p}$ can be approximated as a normally distributed random variable with zero-mean and variance ML . Similarly, as ϑ_l is modelled as random variable with complex Gaussian distribution, we can approximate distribution of $\Re [\sum_{l=1}^L \vartheta_l \mathbf{C}_{l,p}]$ using a normally distributed random variable with zero-mean and variance $L\sigma_{\vartheta}^2/2$. According to (5.7), $\lambda_{l,m}$ is a ratio of two complex quantities. Similar to [90], the quantity $\lambda_{l,m}$ can be modelled using a random variable $\lambda = \lambda_r + i\lambda_i$ with PDF:

$$f(\lambda_r, \lambda_i) = \frac{(1 - |\rho|^2) \sigma_u^2 \sigma_v^2}{\pi} (\sigma_v^2 (\lambda_r^2 + \lambda_i^2) + \sigma_u^2 - 2\rho_r \lambda_r \sigma_u \sigma_v + 2\rho_i \lambda_i \sigma_u \sigma_v)^{-2}, \quad (5.12)$$

where $\sigma_u^2 = |1 - \alpha|^2 \sigma_{h,m}^2 + \sigma_{e,m}^2 + \sigma_{\eta,m}^2$, $\sigma_v^2 = \sigma_{h,m}^2 + \sigma_{e,m}^2$, and $\rho = (1 - \alpha) \sigma_{h,m}^2 + \sigma_{e,m}^2$. In Appendix C, we derive the mathematical expectation and variance of λ_r , i.e. μ_r and σ_r^2 respectively.

In practice, the variance of the CFR, i.e. $\sigma_{h,m}^2$ may be different for $m = 1, 2, \dots, M$. However, the BS always equalises the channel power of the active users through the ranging procedure, where the channel power can be expressed by $\frac{\|\mathbf{h}_m\|^2}{L}$ with $\mathbf{h}_m = [h_{1,m}, h_{2,m}, \dots, h_{L,m}]^T$. Again, if the CFR is modelled as a complex Gaussian random variable, the channel power can be approximated by its variance. Hence, the variance must remain in a known interval due to the ranging procedure. Consequently, we can use the value of average channel power as an estimate of $\sigma_{h,m}^2$ for all m . Accordingly, the BS can estimate $\sigma_{e,m}^2$ and $\sigma_{\eta,m}^2$.

Note that high order moment of $f(\lambda_r, \lambda_i)$ do not exist in general. However, for a bounded variance of λ , in Appendix C we show that the central limit theorem is still applicable for λ . Thus, by using the central limit theorem, the distribution of $\Re [\sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{C}_{l,p}]$ can be modelled as a normally distributed random variable with mean zero and variance $LM\sigma_r^2$. Thus, under the hypothesis $\mathcal{H}_0$, the distribution of $\Re (\mathbf{C}_p^T \mathbf{y})$ can be approximated

as a random variable with zero-mean Gaussian distribution and variance σ_0^2 , where

$$\sigma_0 = \sqrt{L\sigma_\theta^2/2 + ML + ML\sigma_r^2}.$$

Similarly, under the hypothesis $\mathcal{H}_1$, the distribution of $\Re(\mathbf{C}_p^T \mathbf{y})$ can be approximated by a normally distributed random variable with mean L and variance σ_1^2 , where

$$\sigma_1 = \sqrt{L\sigma_\theta^2/2 + (M-1)L + ML\sigma_r^2}$$

Next, we set a threshold ξ such that the decision statistics for p^{th} correlation yields $|\Re(\mathbf{C}_p^T \mathbf{y})| \leq_{\mathcal{H}_0}^{\mathcal{H}_1} \xi$. Thus, the probability of false alarm be

$$P_f^{(p)} = P(|\Re(\mathbf{C}_p^T \mathbf{y})| > \xi | \mathcal{H}_0). \quad (5.13)$$

Note that the random variable $|\Re(\mathbf{C}_p^T \mathbf{y})|$ has a folded normal distribution. Hence, its CDF can be defined as

$$F(z) = P(|\Re(\mathbf{C}_p^T \mathbf{y})| \leq z) = \text{erf}\left(z/\sqrt{2\sigma_0^2}\right),$$

where $\text{erf}(x)$ is the standard error function. Then (5.13) becomes

$$P_f^{(p)} = P(|\Re(\mathbf{C}_p^T \mathbf{y})| > \xi | \mathcal{H}_0) = 1 - \text{erf}\left(\xi/\sqrt{2\sigma_0^2}\right) \quad (5.14)$$

Let us select a threshold ξ_ϕ such that the desired false alarm probability $P_f^{(p)} = \phi$. Then, we have the relation:

$$1 - \text{erf}\left(\xi_\phi/\sqrt{2\sigma_0^2}\right) = \phi. \quad (5.15)$$

The value of ϕ in (5.15) indicates the probability that $|\Re(\mathbf{C}_p^T \mathbf{y})|$ goes above ξ_ϕ under the hypothesis $\mathcal{H}_0$. If any $\{|\Re(\mathbf{C}_p^T \mathbf{y})|\}_{p=1}^L$ that is under $\mathcal{H}_0$, goes above ξ_ϕ , we get a false alarm. Thus, the overall false alarm probability can be defined as

$$P_f = 1 - (1 - \phi)^{L-M}. \quad (5.16)$$

Similarly, let the probability of detection for the p^{th} correlation under $\mathcal{H}_1$ be

$$\begin{aligned} P_d^{(p)} &= P(|\Re(\mathbf{C}_p^T \mathbf{y})| > \xi | \mathcal{H}_1) \\ &= 1 - \frac{1}{2} \left[\text{erf} \left(\frac{\xi + L}{\sqrt{2\sigma_1^2}} \right) + \text{erf} \left(\frac{\xi - L}{\sqrt{2\sigma_1^2}} \right) \right]. \end{aligned} \quad (5.17)$$

Since a detection will only be successful when all M terms of $\{|\Re(\mathbf{C}_p^T \mathbf{y})|\}_{p=1}^L$ that are in $\mathcal{H}_1$ goes above ξ , we can write

$$P_d = \left(P_d^{(p)} \right)^M. \quad (5.18)$$

5.1.2 Hadamard Code Matrix

In general, the code detector cannot detect a large number of active users by using the pseudo-random code matrix. On the detection process under $\mathcal{H}_1$, for a particular active code $\mathbf{C}_p$, the algorithm treats the contribution of other active codes as additional noise without attempting any mitigation of MAI. Moreover, when we correlate $\mathbf{y}$ with $\mathbf{C}_p$, the variance of other noise terms $(\sum_{l=1}^L \vartheta_l \mathbf{C}_{l,p} - \sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{C}_{l,p})$ also increases. The process results in performance degradation as the number of subscribers increase. To overcome his limitation, we apply partial Hadamard matrix as the code matrix $\mathbf{C}$. We first construct a Hadamard matrix of order 256, then retain its upper left block of size $L \times L$ as $\mathbf{C}$. The matching matrix is constructed as $\mathbf{D} = (\mathbf{C})^{-1}$.

Under this case, if $\mathcal{H}_0$ is true, then (5.9) follows

$$\mathbf{D}_p^T \mathbf{y} = - \sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{D}_{l,p} + \sum_{l=1}^L \vartheta_l \mathbf{D}_{l,p}. \quad (5.19)$$

The last part in (5.19) can be approximated as normally distributed complex random variable with zero-mean and variance $\|\mathbf{D}_p\|_2^2 \sigma_v^2$. Since $\mathbf{C}_{l,k_m} \in [+1, -1]$, we can write

$$\mathbb{E} \left(\sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{D}_{l,p} \right) = 0.$$

Moreover, using the derivation in Appendix C, we can assume $\sigma_r \approx \sigma_i$. By using the central limit theorem, the distribution of real and imaginary parts of $\sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{D}_{l,p}$ can be approximated by two zero-mean uncorrelated random variables having normal distributions with the same variance $M \|\mathbf{D}_p\|_2^2 \sigma_r^2$. Hence, the quantity $\mathbf{D}_p^T \mathbf{y}$ can be characterised as a complex random variable with zero-mean and variance $2M \|\mathbf{D}_p\|_2^2 \sigma_r^2 + \|\mathbf{D}_p\|_2^2 \sigma_v^2$. Thus, under the hypothesis $\mathcal{H}_0$, $\mathbf{D}_p^T \mathbf{y}$ can be modelled as a Rayleigh distribution with mode

$$\sigma_{0,p} = \sqrt{M \|\mathbf{D}_p\|_2^2 \sigma_r^2 + \|\mathbf{D}_p\|_2^2 \sigma_v^2 / 2}.$$

Now, for a particular value of $\xi_p : p \in \{1, 2, \dots, M\}$, the probability of false alarm for the p^{th} correlation be

$$P_f^{(p)} = P(|\mathbf{D}_p^T \mathbf{y}| > \xi_p | \mathcal{H}_0). \quad (5.20)$$

Since the random variable $\mathbf{D}_p^T \mathbf{y}$ has a Rayleigh distribution, its CDF can be defined as

$$F(z) = P(|\mathbf{D}_p^T \mathbf{y}| \leq z) = 1 - e^{-z^2 / 2\sigma_0^2}. \quad (5.21)$$

Then (5.20) becomes

$$P_f^{(p)} = P(|\mathbf{D}_p^T \mathbf{y}| > \xi_p | \mathcal{H}_0) = e^{-\xi_p^2 / 2\sigma_{0,p}^2}. \quad (5.22)$$

Thus, for a given probability of false alarm ϕ , we have

$$\xi_p = \sigma_{0,p} \sqrt{-2 \log(\phi)}. \quad (5.23)$$

Conversely, if $\mathcal{H}_1$ is true, then (5.9) follows

$$\mathbf{D}_p^T \mathbf{y} = 1 - \sum_{l=1}^L \sum_{m=1}^M \lambda_{l,m} \mathbf{C}_{l,k_m} \mathbf{D}_{l,p} + \sum_{l=1}^L \vartheta_l \mathbf{D}_{l,p}. \quad (5.24)$$

The quantity $|\mathbf{D}_p^T \mathbf{y}|$ can be approximated as a random variable having Rician distribution with parameters $\mu = 1$ and $\sigma = \sigma_{0,p}$. Thus, for a particular threshold ξ_p , the detection

probability for the p^{th} correlation can be calculated as

$$\begin{aligned} P_d^{(p)} &= P(|\mathbf{D}_p^T \mathbf{y}| > \xi_p | \mathcal{H}_1) \\ &= Q_1\left(\frac{1}{\sigma_{0,p}}, \frac{\xi_p}{\sigma_{0,p}}\right). \end{aligned} \quad (5.25)$$

where $Q_1(\cdot, \cdot)$ is the Marcum-Q function.

The value of $\sigma_{0,p}$ can control the false alarm rate, as well as the probability of detection. A large value of $\sigma_{0,p}$ increases the noise contribution in the signal, and hence, the false alarm rate (see (5.22)). Further, the value of $\sigma_{0,p}$ depends on the value of $\|\mathbf{D}_p\|_2^2$. Hence, in a particular environment, the code detection probability of a user increases if it can select a code with lower $\|\mathbf{D}_p\|_2^2$. Another interesting observation is that if we set a fixed false alarm rate ϕ in (5.23) for every $p = 1, 2, \dots, K$, we may get K different values of threshold $\{\xi_p\}_{p=1}^K$, each for a specific column of $\mathbf{D}_p$; hence, the detection rate computed in (5.25) may have K different values. For illustration purpose, we set $\bar{P}_d^{(p)} = \frac{1}{K} \sum_{p=1}^K \{P_d^{(p)}\}$ in (5.18) to compute the overall detection probability.

Note that (5.16) needs to have the total number of active devices M , which is unknown in practice. The energy received at the BS from the real parts of $\mathbf{y}$ can be expressed as (using (5.6)):

$$\mathcal{G} = \Re(\mathbf{y})^T \Re(\mathbf{y}) = \sum_{\ell=1}^L \left[\sum_{m=1}^M \left\{ \mathbf{C}_{\ell,k_m} - \lambda_{\ell,m}^{(\Re)} \mathbf{C}_{\ell,k_m} \right\} + v_{\ell}^{(\Re)} \right]^2, \quad (5.26)$$

where $v_{\ell}^{(\Re)}$ is the real part of v_{ℓ} , such that $v_{\ell}^{(\Re)} \sim \mathcal{N}(0, \sigma_v^2/2)$. Then $\sum_{k=1}^L \left(v_k^{(\Re)}\right)^2$ can be described as a random variable with Chi-square distribution having mean $L\sigma_v^2/2$ and variance $L\sigma_v^4/4$. Also, for a large L , using central limit theorem it can be shown that the quantity can be approximated as a normally distributed variable with mean $L\sigma_v^2/2$ and variance $L\sigma_v^4/4$. After a few computations, it can be verified that

$$\mathbb{E}(\mathcal{G}) = ML - 2ML\mathbb{E}(\lambda^{(\Re)}) + ML\mathbb{E}\left(\left(\lambda^{(\Re)}\right)^2\right) + L\sigma_v^2/2, \quad (5.27)$$

where $\mathbb{E}(x)$ is mathematical expectation of x . Hence, one can estimate the value of M from

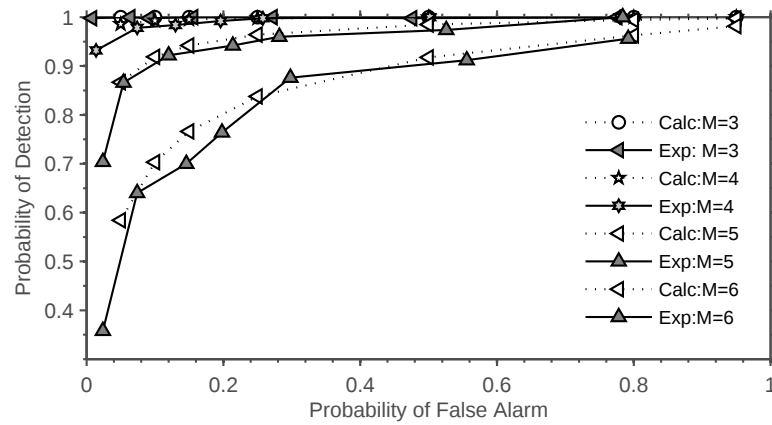
$$\Re(\mathbf{y})^T \Re(\mathbf{y}).$$

5.1.3 Simulation Results

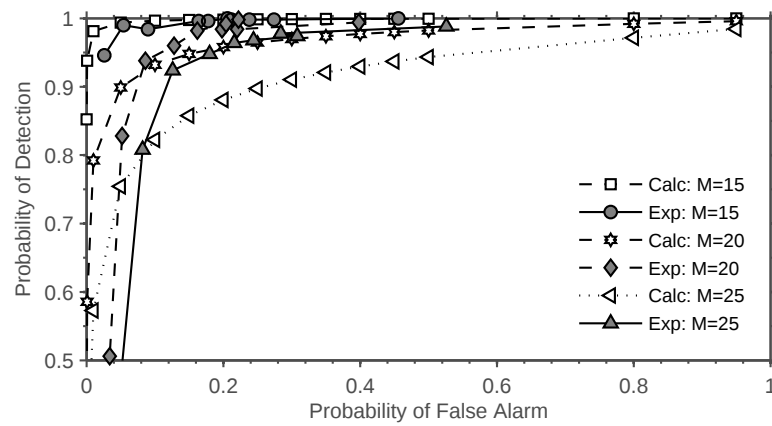
To evaluate the performance of the proposed random access scheme, we develop a Monte-Carlo simulation model using MATLAB. The simulation parameters are chosen based on the IEEE 802.16-2009 standard [58]. The system is assumed to be operating at 2.3 GHz with a bandwidth of 5 MHz and a sampling frequency of 5.6 MHz. The FFT size is assumed to be 512 with a cyclic-prefix size of 128 samples. The multipath channel parameters are same as listed in Table 4.1.

We evaluate the code detection performance under each of the two different code matrices. Only SUI-3 channel model is used for this set of simulations. We follow the direction of [90] and set $\alpha = e^{i\theta}$, where $\theta = 5^\circ$. The ROC curve for different users by applying pseudo-random matrix is shown in Figure 5.2(a). For a fixed number of users M , we produce the curves in the following way: i) set an elementary false alarm ϕ and by using (5.15) compute the corresponding threshold ξ_ϕ ; ii) use the ξ_ϕ to produce theoretical overall probability of false alarm (using (5.16)) and theoretical probability of detection (using (5.18)); iii) generate 500 different received signals $\mathbf{y}$ under the simulation environment described earlier; iv) for every $\mathbf{y}$ perform the correlation as in (5.9); v) the correlation output is checked with the threshold ξ_ϕ to compute empirical probability of false alarm and detection rate; vi) repeat procedures i)-v) for different values of ϕ . A similar strategy is applied in obtaining Figure 5.2(b), which exhibits the ROC curve for the partial Hadamard code matrix.

From the results in Figure 5.2, we see that under both code matrices, the probability of detection increases when we increase the false alarm rate. Moreover, although the default pseudo-random matrix has more available codes (i.e. $K = 256$) than the Hadamard code matrix (i.e. $L=144$), the Hadamard code matrix significantly outperforms the pseudo-random code matrix at a given probability of false alarm. This is because the Hadamard matrix provides perfectly orthogonal codes, and the effect of MAI is mitigated during the cross-correlation of $\mathbf{y}$ and $\mathbf{C}_p : p \in \{1, 2, \dots, K\}$. It also allows segmentation of the code matrix, which will be discussed in the next section.



(a)



(b)

Figure 5.2: Theoretical (Calc) and empirical (Exp) ROC curve for different active users (M) for (a) pseudo-random and (b) Hadamard code matrices under the SUI-3 channel model.

Next, we evaluate the effect of channel noise on the code detection performance by two coding matrices in Figure 5.3. The simulations are conducted using both SUI-3 and SUI-4 channel models. Let the average channel power be P . Then channel SNR is defined as $\text{SNR} = 10 \log_{10} (P/\sigma_b^2)$. We fixed $\phi = 0.05$, and calculate $\xi^{(p)}$ by using (5.15) and (5.23). The values of $\xi^{(p)}$ are used as threshold for hypothetical test. From the results, we see that under both of the matrices, as the SNR value increases, the probability of detection increases. However, the Hadamard code matrix clearly outperforms the pseudo-random matrix irrespective of the SNR values. The results are consistent for both channel models.

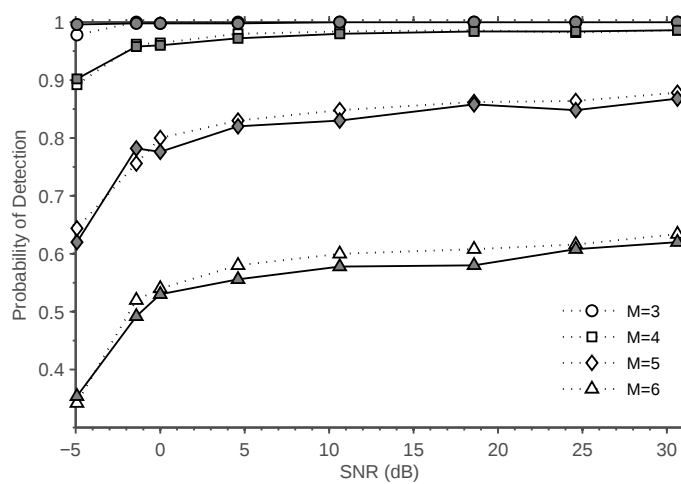
5.2 Differentiated Random Access

In a Smart Grid communications environment, it is expected that the number of low priority M2M devices (e.g., meters and sensors) would be much higher than that of the high priority devices (e.g., controllers and relays). As a result, an increase in the low priority random access requests may increase the access delay of the high priority requests resulting into QoS-degradations for the time-critical M2M applications. The aim of the differentiated random access procedure is to provide a mechanism so that the high and the low priority requests can be served separately over the same BR channel.

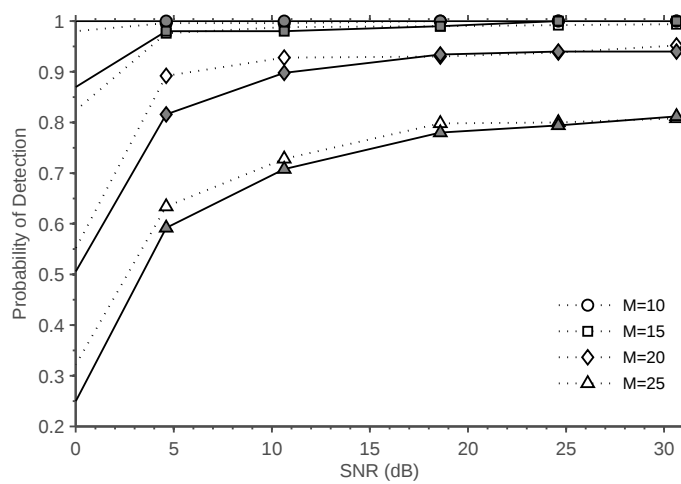
First, consider the simple case where two classes of M2M traffic exists: high and low. Let us construct the Hadamard code matrix described in the previous section as

$$\mathbf{C} = \begin{bmatrix} \mathbf{C}_G & \mathbf{0}_{G \times W} \\ \mathbf{0}_{W \times G} & \mathbf{C}_W \end{bmatrix}, \quad (5.28)$$

where $\mathbf{0}_{W \times G}$ is a null matrix of size $W \times G$, $\mathbf{C}_G \in \mathbb{R}^{G \times G}$ and $\mathbf{C}_W \in \mathbb{R}^{W \times W}$ are the two partial Hadamard code-matrices for high and low priority classes, respectively. Thus, the original $K = L$ codes are divided into two groups such that the first G codes are assigned to the N_G high-priority stations where $N_G > G$, and the last W codes are assigned to the N_W low-priority stations, where $K = G + W$. When multiple devices access the BR channel, class-wise detection can be performed by considering $\mathbf{D}_G = (\mathbf{C}_G)^{-1}$ and $\mathbf{D}_W = (\mathbf{C}_W)^{-1}$.



(a)



(b)

Figure 5.3: Detection performance as a function of channel SNR for (a) pseudo-random and (b) Hadamard code matrices under the SUI-3 (dotted line) and SUI-4 (solid line) channel models. Distribution parameters are $\sigma_e = 0.05$, $\sigma_\eta = 0.05$.

Now, if the detector is able to detect D_G high-priority codes and D_W low-priority codes from the two classes, the probability of successful access for the system is given by (see (4.57)):

$$p = \sum_{i=0}^{D_G-1} \binom{N_G-i}{i} \tau_G^i (1 - \tau_G)^{N_G-1-i} \left(1 - \frac{1}{G}\right)^i + \sum_{i=0}^{D_W} \binom{N_W-i}{i} \tau_W^i (1 - \tau_W)^{N_W-1-i} \left(1 - \frac{1}{W}\right)^i, \\ \text{or, } p = p_G + p_W. \quad (5.29)$$

From (5.29), we see that the use of separate code-matrices divides the contention domain of the physical BR channel into two logical domains (i.e. p_G and p_W) for each classes of traffic. Thus, system is able to control the access delay by changing the size of the code-matrices as well as using different sets of contention parameters for each of the logical channels.

5.3 Adaptive Radio Resource Management

As described in Section 3.1.3, in a WiMAX network, the radio resources are allocated using a request/grant mechanism. The request phase is required to reserve a sufficient amount of bandwidth from the BS before any data transmission. On the other hand, the grant phase involves allocating grant against each BS based on its QoS class. There are mainly two types of BR mechanisms specified in the IEEE 802.16 standard - contention based random access and contention free polling [5]. Of these, the random access based BR is preferable for the bursty Smart Grid M2M applications that exploits the benefits of statistical multiplexing to support a large number of devices with a fixed overhead.

As depicted in Figure 3.2, after a successful random access, the SS is provided with a bandwidth grant to send the data packet. However, a high priority packet need not only a quick BR procedure, but also a small grant scheduling delay to meet its end-to-end latency requirement. Moreover, in the existing IEEE 802.16 standard, only a single QoS scheduling

class, the BE service is associated with the random access mechanism. Typically, the BE class is used to serve delay tolerant applications, such as file transfer and web browsing; therefore, it is given the least priority in the BS scheduler and typically uses the RR algorithm for grant scheduling. Consequently, the grant delay increases with the increase in the traffic load, as a BR needs to wait longer for its turn. Moreover, an increased grant delay may also increase the random access delay. This is because if a device does not receive a grant within T_{16} period, it resumes the back-off process.

Therefore, to provide guaranteed QoS to the bursty M2M traffic in the Smart Grid, the BR and grant mechanism must work cohesively. To achieve this, we propose a novel scheduling framework that uses three separate scheduling classes to support the bursty M2M traffic in the Smart Grid, based on their priority and latency requirements. The names of these three services along with their illustrative latency margins are provided in Table 5.1. Note that these names are borrowed from the DSCP concept to indicate their relative QoS precedence.

Table 5.1: Proposed Scheduling Classes for Bursty M2M Traffic in the Smart Grid

Scheduling Class	Traffic Priority	Latency Margin
Expedited Effort (EE)	Very High	<100 ms
Assured Effort (AE)	High/Medium	100-500 ms
Best Effort (BE)	Low	> 500 ms

For the proposed RRM framework, we utilise the DRA strategy described in Section 5.2. In this concept, the physical BR channels are divided into a set of logical channels called the virtual ranging channels (VRCs). Each VRC comprises a dedicated set of ranging codes that are adaptively allocated by a radio resource agent (RRA) located in the BS. Depending on the code detection capability of the BS and the random access load in the system, there might be multiple BR channels in the system. Hence, the ranging resources are assigned with a code, channel pair information for a particular VRC (or traffic class). Moreover, the RRA may periodically update the allocation parameters after a pre-defined interval. After a successful BR, each packet is served by a dedicated data queue to meet its application-specific QoS requirements. For grant scheduling, we use a non-preemptive

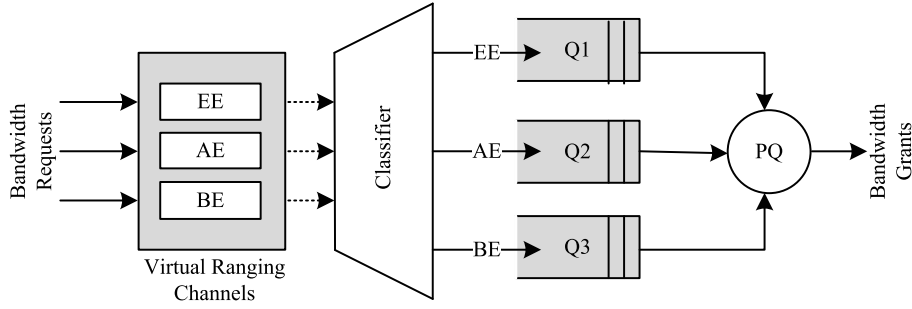


Figure 5.4: The proposed adaptive RRM framework.

priority queue (PQ) that allocates bandwidth based on the priority of each packet. The use of other scheduling algorithms is also possible. However, this is outside the scope of this work. Figure 5.4 illustrates the proposed RRM framework.

As seen from the figure, the proposed RRM framework is comprised of three planes: the random access plane, the request classification plane, and the grant scheduling plane. After a successful BR, each packet is served by a dedicated data queue to meet the application specific QoS target. Note that to meet a certain delay bound, the amount of random access load in a given VRC needs to be maintained below a certain threshold. This can be done by adaptively adjusting the number of available ranging codes and/or adjusting the back-off window. However, the exact configuration depends on the priority and latency requirement of the traffic class associated with that VRC.

The random access plane in the proposed RRM scheme works in two phases: load monitoring and parameter adjustment. They are briefly described below.

i) Load Monitoring: The RRA continuously monitors the contention load for each VRC. The monitoring process is time-sequenced, i.e. separated into fixed intervals (Δt). The fixed interval also denotes the time horizon where the ranging resource adjustment will be performed. At each Δt interval, the RRA measures the normalised contention load $l_i(t)$ for the i^{th} virtual channel at time t . The value of i also represents the priority of the channel in descending order, i.e. $i = 0$ for EE, $i = 1$ for AE, and $i = 2$ for BE.

i) Parameter Adjustment (Heuristic Method): The contention parameters for the VRCs can be heuristically derived from the random access performance analysis using the analyt-

Table 5.2: Illustrative Contention Parameters for the Adaptive RRM Scheme

Traffic Class	Contention Load ($\hat{l}$)	Contention Parameters		
		W_0, W_f	T_{16} (ms)	m
EE	15%	2,4	20	16
AE	25%	4,8	40	16
BE	35%	4,8	80	16

ical model presented in Section 4.3. For instance, consider the contention parameters listed in Table 5.2 for a single VRC with only one detectable code. Here, the time-out parameter T_{16} is adjusted depending on the delay margin. A small T_{16} ensures that the next contention retry will be faster. This is particularly useful for the delay sensitive traffic to receive a quick bandwidth grant.

The algorithm for the adaptive ranging resource allocation is described below:

1. Initialise the number of ranging codes (denoted as k) to one for the i^{th} VRC such that $k[i] = 1$.
2. Calculate the normalised access load for each VRC using the formula:

$$l_i(t) = \frac{n_i(t) - n_i(t - \Delta t)}{k[i]} \cdot \frac{T_f}{\Delta t},$$

where n denotes the number of successful requests and T_f is the WiMAX frame duration.

3. Calculate the required number of ranging codes for each VRC using the following procedure. For example, for the EE traffic class ($i=0$), **Do**: Calculate the normalised contention load

$$l_0(t) = \frac{n_0(t) - n_0(t - \Delta t)}{k[0]} \cdot \frac{T_f}{\Delta t},$$

and increase the number of ranging codes, $k[0]++$; **While**: $l_0(t) > \hat{l}_0$, where $\hat{l}_0$ is the maximum contention load for the EE traffic class.

4. Repeat Step-3 for $l_1(t), \hat{l}_1$ and $l_2(t), \hat{l}_2$ and obtain $k[1]$ and $k[2]$.

Table 5.3: Traffic Parameters for the Adaptive RRM Simulation

Traffic Priority	Base Load (Packets/sec)	Peak Load (Packets/sec)	Delay Margin
High	100	200	<100 ms
Medium	100	200	<500 ms
Low	200	200	1 sec

5. Calculate total number of ranging channels required for each VRC and allocate code, channel pair via the next UL-MAP message.

5.3.1 Simulation Results

To demonstrate the performance of the proposed scheme, we develop a simulation model using OPNET Modeller 16.0. The baseline scenario is made of three classes of exponentially distributed M2M traffic, with the parameters specified in Table 5.3. Since we are mainly interested in the RRM aspects of the WiMAX network, the free-space path-loss model is used for the study. The simulation run-time is 15 minutes. We assume that during the simulation time, a grid event drives both medium and high priority traffic to their peak levels for 5 minutes and return to base load again.

The WiMAX simulation parameters are same as Table 3.5, except here we use an enhanced ranging channel (as per Section 5.1) with 8 detectable codes, yielding a normalised contention load of around 35 per cent in terms of the aggregated peak load. For the adaptive RRM scheme, we map the high, medium and low priority traffic with EE, AE, and BE classes, respectively. The contention parameters are set according to Table 5.2. For performance evaluation, we use the KPI MDSR, which represents the percentage number of packets that were received within the required delay bound.

Figure 5.5 shows the instantaneous load of the network, including the individual loads from each class of traffic. As seen from the figure, a base load of 400 packets/sec shoots to 600 packets/sec from the 5th to the 9th minute, as the traffic from medium and high priority class increase at that duration.

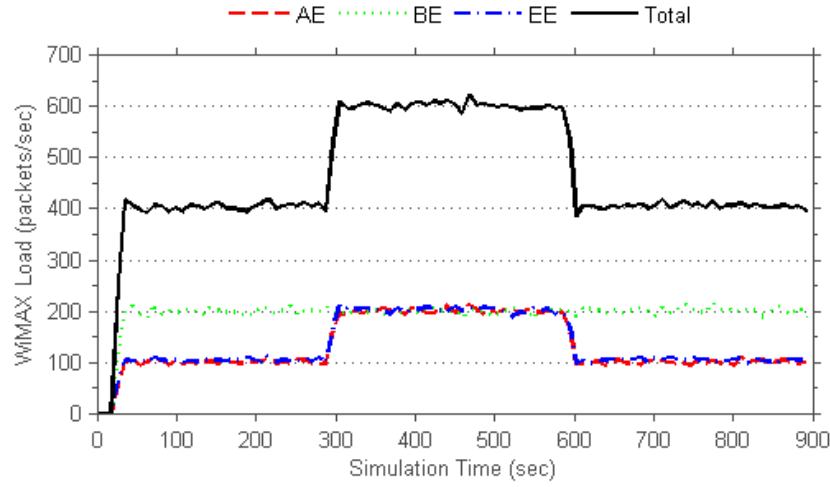


Figure 5.5: WiMAX network load for different classes of traffic under the simulation scenario.

Table 5.4: Comparative Delay Performance of the Adaptive RRM Scheme

Scheme	Traffic Class	Mean Access Delay (ms)	Max. Grant Delay (ms)
Conventional	EE	60.62	85
	AE	60.31	85
	BE	55.13	90
Proposed	EE	17.78	30
	AE	71.84	45
	BE	148.57	90

Next, we look into the comparative delay performance of the conventional and proposed schemes in terms of mean random access delay and the maximum grant delay in Table 5.4. From the results, we see that the proposed scheme supported by segregated ranging domains and adaptive ranging resources reduces the mean random access delay significantly, especially for the high and medium priority traffic. Moreover, the strict priority queuing significantly reduces the maximum grant scheduling delay for the high and medium priority traffic as compared to the simple RR queuing under the conventional scheme. However, the delay for the BE traffic class is increased under the proposed scheme as it tries to optimise its throughput using minimum number of ranging codes.

The accumulated effect of DRA and prioritised bandwidth grant substantially reduces the

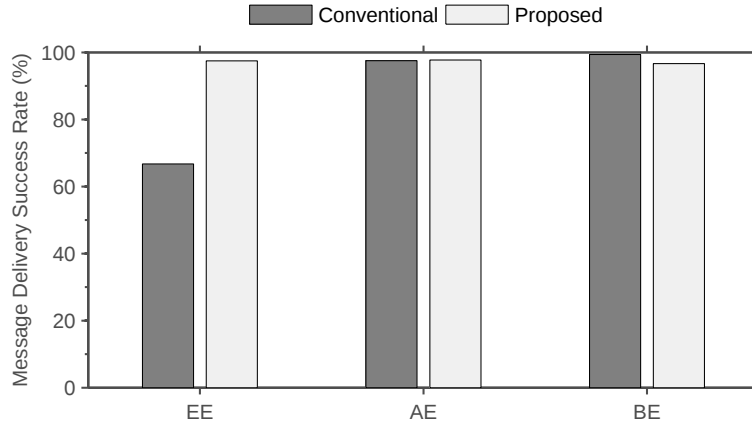


Figure 5.6: Mean MDSR for different classes of traffic under the conventional and the proposed RRM scheme.

overall uplink packet delay of the EE and AE traffic class. This is evident in Figure 5.6, which compares the MDSR for different classes of traffic under the conventional and the proposed scheme. The results show that the proposed scheme drastically improves the MDSR of the EE traffic while the MDSR for the AE traffic class remains almost the same. Note that the MDSR for the BE class is slightly reduced under the proposed scheme. However, this is acceptable since the BE service is typically used to serve delay-tolerant applications only.

Figure 5.7 plots the number of ranging codes for different classes of traffic under the proposed scheme. The results show how the allocation of ranging codes changes with the change of network load as shown in Figure 5.5.

Lastly, Figure 5.8 shows how the overall contention load remains below the threshold specified in Table 5.2 for different classes of traffic.

5.3.2 Illustrative Example: Pilot Protection Scheme

To further illustrate the performance of the proposed scheduling framework, we consider the pilot protection scheme described in Section 3.4. To support this application, the WiMAX network has to transfer the delay-sensitive pilot data packets with the required delay bound of 40 ms. Hence, we assume that the pilot protection traffic is mapped to the

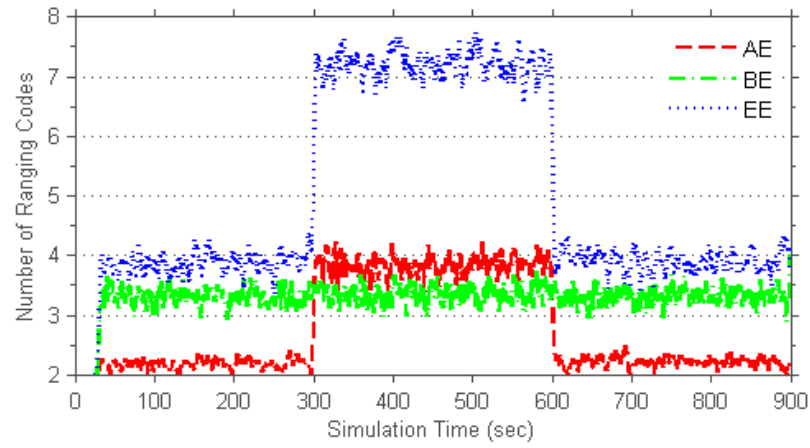


Figure 5.7: The number of allocated VRCs for different classes of traffic under the adaptive RRM scheme.

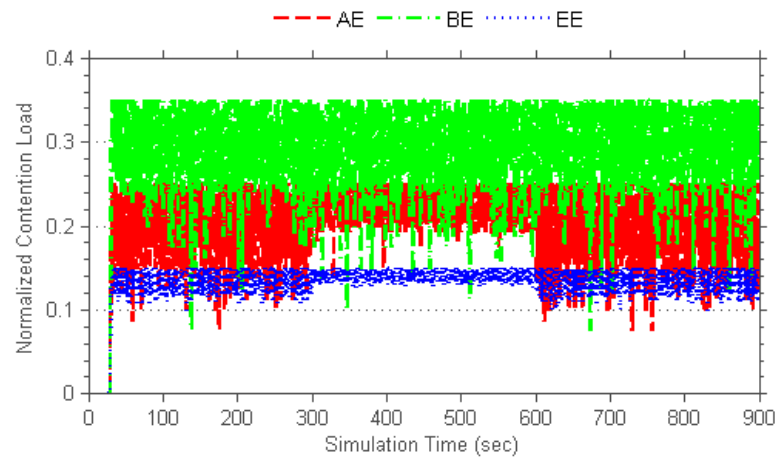


Figure 5.8: The normalised contention loads for different classes of traffic under the adaptive RRM scheme.

EE scheduling service listed in Table 5.1. As the random access attempts from such an application are very low, we assume that the corresponding VRC is allocated with a single dedicated ranging code.

It is plausible that the packet size of a pilot protection message would remain fixed as it carries a fixed set of information. Hence, the BS does not need to receive an explicit BR header to allocate the UL bandwidth grant, as it already knows the bandwidth requirement by recognising the dedicated ranging code. This eliminates step two and three under the conventional BE service as shown in Figure 3.2, and thereby, reduces the end-to-end delay by at least one WiMAX frame.

Next, we need to set appropriate contention parameters to the proposed EE service to meet the delay bounds of the pilot protection messages. Based on the application characteristics, we can make the following assumptions for a pilot protection scheme:

1. As the time of fault occurrence in a protected line is highly random and the multiple protection zones (if any) over the same protected line are time-graded by a zone coordination scheme [66], the probability of fault occurring simultaneously in more than one protection zone is very low.
2. The probability of detecting a fault simultaneously by more than one relays in the same protection zone is very low as the upstream pilot relay (i.e. the relay closer to the fault) should detect a fault earlier than the downstream relay [65].

Before moving on, let us define the probability that a code suffers a collision as p_c and the probability that a code is misdetected by the BS as p_m . Thus, the overall failure probability is given by

$$p = 1 - (1 - p_c)(1 - p_m) \quad (5.30)$$

Based on these assumptions, we can safely conclude that for most of the times, the relays will be able to request bandwidth without any collision, i.e. $p_c \rightarrow 0$. Thus, the primary factor impacting the random access delay would be the probability of misdetection p_m independent of p_c . Hence, there is no need for a random back-off during the first attempt,

Table 5.5: Contention Parameters per Traffic Class

Parameters	EE	BE
No. of BR Ranging Channels	1	
No. of Available Codes (K)	1	8
Initial Back-off Window (W_0)	0	2
Final Back-off Window (W_f)	1	15
Number of Retries (m)	8	15
Contention Time-out (n)	1	8

i.e. $W_0 = 0$. Additionally, it is reasonable to set $W_f = 1$ to keep the random access delay minimum for the subsequent retransmissions. Note that the typical size of a pilot signal packet (around 100 bytes) is much less than the UL subframe capacity of a WiMAX network. Hence, we can also assume that the successful codes are allocated CDMA grant within the next WiMAX frame (i.e. $T_{16} = 1$ frame), considering that the protection traffic will be given the highest priority under the proposed EE service.

To evaluate the communications performance of the above pilot protection scheme, we use the same simulation model as described earlier. Moreover, to account for the misdetected codes due to channel noise and fading, random failures were generated using a uniform probability distribution function. For the sake of simplicity, we assume that the network only supports two classes of traffic: the conventional BE traffic, and the pilot protection traffic under the EE service. The BE traffic is further varied between 100 to 400 packets/sec. The simulation scenarios and parameters are same as used in Section 3.4. However, we use a separate set of contention parameters as listed in Table 5.5.

Next, we look into the performance of the pilot protection application under the proposed EE service. As the probability of collision p_c is very low for such an application over a dedicated VRC, the performance is mainly affected by the stochastic variability in the PHY layer in the form of probability of misdetection p_m . The corresponding delay components of the pilot signal data packets are listed in Table 5.6, with different probabilities of misdetection.

From the results, we see that under the proposed EE service only the random access delay increases while the other delay components remains the same. This is because since the

Table 5.6: Pilot Signal Delay Components (in ms) under the EE Service

Prob. of Misdetect.	Random Access		Uplink		Downlink	
	Mean	S.Dev.	Mean	S.Dev.	Mean	S.Dev.
0%	7.59	1.67	5.61	0.00	5.97	0.00
5%	8.45	4.33	5.61	0.00	5.97	0.00
10%	9.83	6.52	5.61	0.00	5.97	0.00
15%	10.79	7.96	5.61	0.00	5.97	0.00
20%	12.49	9.89	5.61	0.00	5.97	0.00

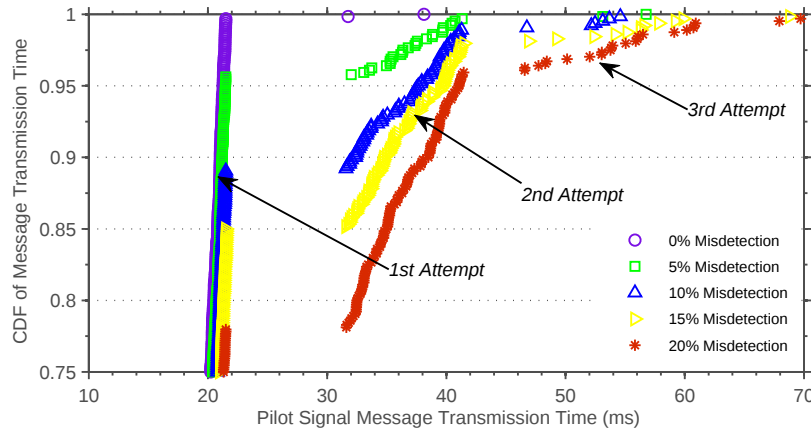


Figure 5.9: CDF of MTTs for the pilot signals under the proposed EE service at various channel error probabilities.

EE service prioritises the pilot signal data packets in the BS scheduler, they receive an immediate bandwidth grant without any further delays. The mean random access delay has a very small standard deviation under the ideal channel condition as only a few collisions were observed due to very low traffic load. However, the random access delay increases as the channel condition degrades; this is mainly due to the retransmission attempts made during code transmission failures.

This is further evident in Figure 5.9 that plots the CDF of MTTs for the pilot signals under various channel error probabilities. Also, the corresponding MDSRs are listed in Figure 5.10.

From Figure 5.9, we can clearly see the different bands of delay distribution related to the number of retransmission attempts from the relays. The results imply that a pilot signal data packet is able to meet the 40 ms delay bound even after suffering one retransmission. This

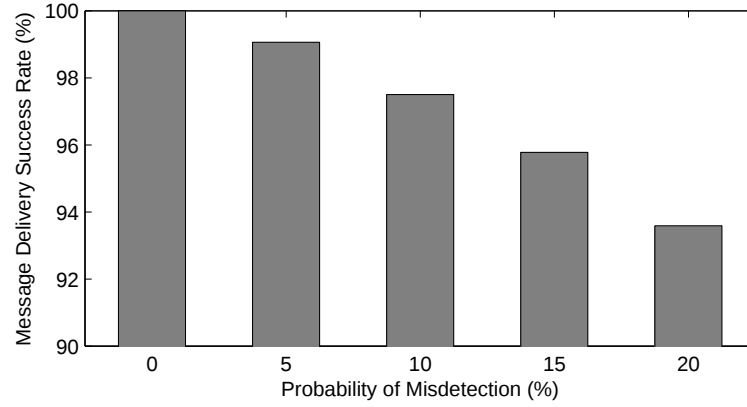


Figure 5.10: Mean MDSR for the pilot signals under the proposed EE service at various channel error probabilities.

is because a small back-off window (i.e., $W_m = 2^1$) followed by an immediate bandwidth grant (i.e., $T_{16} = 1$ frame) from the PQ scheduler ensures that the retransmission delay remains minimum. Note that the mean MTTs for all channel conditions remain at around 50 per cent of the TTL value (i.e., 20 ms). This additional delay margin may also allow MAC layer retransmissions in case of a corrupted packet transmission using ARQ or HARQ techniques. Lastly, from Figure 5.10, we see that the proposed EE service is able to provide fairly high MDSR, even in difficult channel conditions.

5.4 Chapter Summary

In this chapter, we addressed the key bottleneck for supporting M2M communications over a WiMAX network, by proposing an enhanced random access scheme based on frequency domain pre-equalisation in the physical layer. Based on this scheme, we also introduced the concept of DRA for the M2M traffic and proposed an adaptive RRM scheme to efficiently allocate the random access resources. The performance of the proposed scheme was demonstrated through a comprehensive set of simulations under different access loads, multipath channel models, and application scenarios. We believe that the proposed random access scheme, along with the DRA strategy and adaptive RRM framework, would be able to significantly improve the performance and use of a WiMAX network under an M2M

communications environment.

Chapter 6

Heterogeneous Network: The Ultimate Solution?

From the analyses and discussions in the preceding chapters, we can conclude that coverage improvement, random access regulation, and efficient transmission of small data bursts are the three key challenges that a typical WiMAX network faces in an M2M communications environment in the Smart Grid. In this chapter, we argue that a heterogeneous deployment of a WiMAX network with one or more short/medium range wireless technologies, such as IEEE 802.11 based WLAN or IEEE 802.15.4 based ZigBee, can efficiently address these challenges. Thus, it has the potential to offer the ultimate wireless solution for a Smart Grid communications network. Such a multi-tier network is fully compliant with the communications architecture of the Smart Grid described in Section 2.1. Moreover, it provides physical separation between the NAN/FAN and the WAN of the E2E Smart Grid network, which improves security and encourages the development of new distributed control and energy management applications in the NAN/FAN domain.

To validate the proof of concept, we consider a WiMAX-WLAN HetNet architecture, where the dual radio WLAN access points (APs) are embedded within the coverage umbrella of a WiMAX network. Being the two promising technologies of the next generation wireless networks, the integration of WiMAX and WLAN can significantly improve the cost and efficiency of the Smart Grid communications network by extending transmission

range, improving link quality, and allowing flexible data aggregation to reduce signalling and protocol overheads.

The remainder of this chapter is outlined as follows. Section 6.1 introduces the proposed HetNet architecture and describes its key tenets, such as network architecture, interworking, and QoS management techniques. Sections 6.2 and 6.3 examine the behaviour and performance of the HetNet architecture over the two key Smart Grid M2M entities: AMI/smart meters and synchrophasors, and compare its performance with a standalone WiMAX network. Finally, Section 6.4 concludes this chapter ¹.

6.1 WiMAX-WLAN HetNet: An Overview

In Chapter 3, we performed a basic coverage analysis of a typical suburban WiMAX network under the Smart Grid/AMI environment. The results show the coverage area is significantly reduced by a large building penetration loss, due to obstructions in the harsh utility environment. Consequently, some M2M devices may always remain under poor coverage area as they operate with the fixed wireless links. A possible solution to this problem could be a denser deployment of the BSs. However, it is cost prohibitive and may further deteriorate the propagation environment due to increased interference. Another solution could be the deployment of low cost relays in the network, such as the IEEE 802.16j based WiMAX multi-hop relays [92]. However, the relays also increase interference, transmission latency, and signalling overhead, as well complicates the network infrastructure.

In contrast, heterogeneous deployment of WiMAX and WLAN networks can significantly improve the coverage, capacity and utilisation of a Smart Grid communications network by embedding WLAN APs within the coverage umbrella of the WiMAX network. The WLAN APs can be strategically mounted on the existing utility poles and substation structures to take advantage of reasonably high antenna heights and increased power supply, which can significantly extend the range, improve link quality and eliminate dead spots in the combined wireless network.

¹A part of this chapter has been published in proceedings of the *IEEE International Conference on Smart Grid Communications (SmartGridComm)* in Nov 2012 [91].

Note that unlike the IEEE 802.16j-based multi-hop relay networks, the advantage of a WiMAX-WLAN HetNet is that the hybrid network can use spectra from both licensed and licensed-exempt bands. This could increase the capacity of the overall system without raising the interference margin. Moreover, the improved link margin will allow the WiMAX network to operate on a higher MCS level, which can significantly increase its overall data rate. Further, the WiMAX-enabled APs can aggregate data from multiple M2M devices/services into a single burst and then send it to the WiMAX BS. This could significantly improve the data transmission efficiency of the WiMAX network by reducing signalling and protocol overheads.

An additional benefit of the above solution is that the end devices may considerably save their transmit power due to improved SNR. Also, they can use relatively less expensive WLAN chips (compared with WiMAX), which may reduce the overall network deployment cost. Moreover, the WLAN APs can be deployed in mesh mode to enable redundancy and self-healing in the NAN/FAN domain. However, this is outside the scope of this work.

6.1.1 Architecture and Interworking

Since the WiMAX and the WLAN networks have different protocol architectures and QoS support mechanisms, protocol adaptation is required for their interworking [93]. The integration between WiMAX and WLAN network can occur in different layers of the communication stack. For example, with the layer 2 approach, adaptations would be required in the MAC layers of the WiMAX BS and the WLAN stations, where the integrated network will operate in the same frequency band and employ a hybrid MAC protocol [94, 95]. In contrast, with the layer 3 approach, adaptation would be performed at the IP layer, and a WLAN station would interact with the corresponding WLAN AP/router only, as shown in Figure 6.1 [96]. For the proposed HetNet architecture, we use the layer 3 approach because it allows physical separation between the two networks, and provides the full benefits of heterogeneity as described earlier.

The key enabler for the proposed network architecture is a dual radio WiMAX-WLAN router (WWR) that performs protocol adaptation in the IP layer as shown in Figure 6.1.

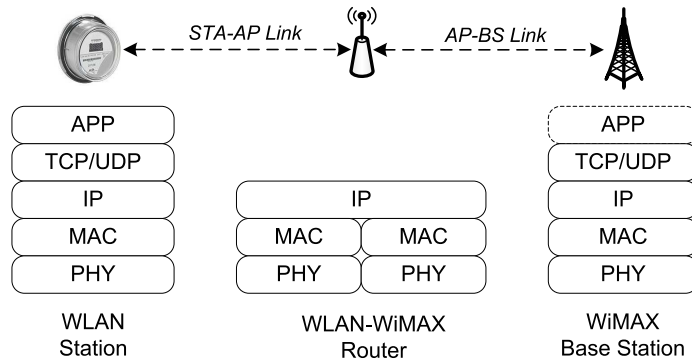


Figure 6.1: Protocol stacks within the WiMAX-WLAN HetNet architecture.

The advantage of such an interworking is that the networks are physically separated, and the end devices need to interact with the corresponding WWR only. Hence, modifications of WLAN stations and the WiMAX BS are not required [93]. It should be noted that dual-mode WiMAX-WLAN routers (WWRs) are already commercially available. For example, the Intel WiMAX/WLAN Link 5150 and 5350 that operates in the 2.5GHz spectrum for WiMAX and 2.4 GHz and 5 GHz spectra for WLAN [97].

6.1.2 Modes of Operation

Based on the working principle of the WWR, the proposed network architecture can operate in the following two different modes:

1. **Transparent Mode:** here the WWR acts as a relay that receives data from the end devices via the WLAN link and transmits it to the WiMAX BS via the WiMAX link, as shown in Figure 6.2(a). The WiMAX network is responsible for maintaining the end-to-end QoS for the applications, which requires explicit coordination between the WLAN and WiMAX network, for example, FTP and HTTP sessions for the DR applications; and
2. **Aggregation Mode:** here the WWR acts as a DAP and aggregates traffic from multiple M2M devices into a single burst, and then transmit it to the WiMAX-BS as shown in Figure 6.2(b). Note that the data aggregation process can significantly improve the data transmission efficiency of the M2M devices, as a typical M2M data

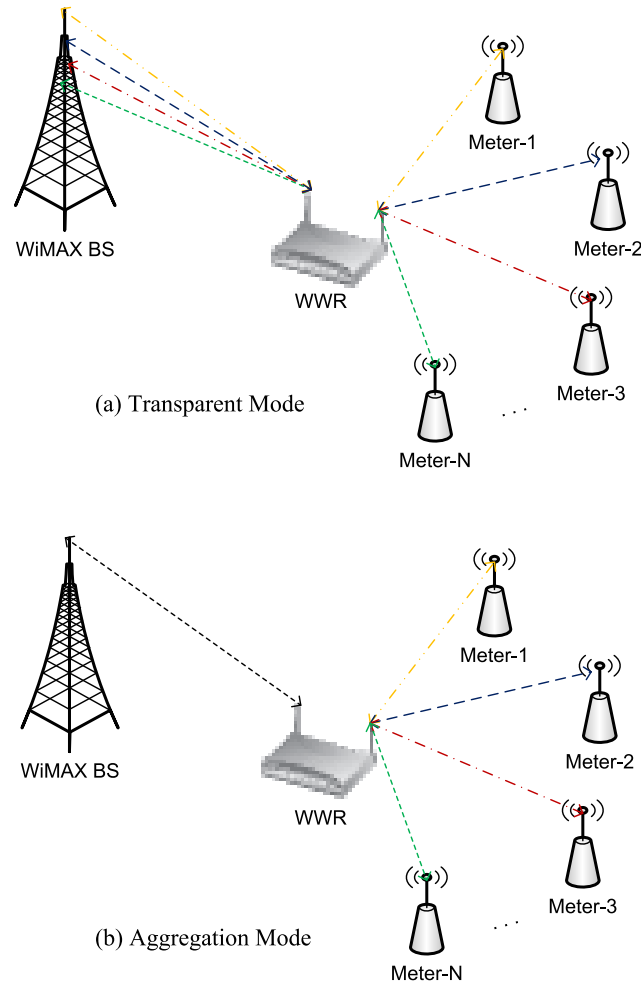


Figure 6.2: Modes of operation of the proposed WiMAX-WLAN HetNet architecture.

packet is associated with a small payload and a relatively large protocol overhead. Moreover, data aggregation can significantly reduce the MAP signalling overhead in the WiMAX network, as only a single MAP IE is required for transferring the aggregated data burst, and no QoS handshake is required between the WiMAX and the WLAN network.

6.1.3 QoS Management

Perhaps the most important challenge for the HetNet architecture is to maintain end-to-end QoS throughout the WiMAX and WLAN links, especially for the delay sensitive M2M traffic. Within the HetNet, the delay in WLAN networks mainly depends on the MAC

technique used. The WLAN standards (i.e., IEEE 802.11 a/b/g/n) supports two basic MAC mechanisms: the contention based distributed coordination function (DCF) for the bursty data traffic; and the polling-based point coordination function (PCF) for the periodic data traffic. In Chapter 3, we explained the QoS framework of WiMAX in detail. In the below paragraphs, we briefly describe the WLAN DCF and PCF mechanisms.

a) DCF Mechanism: Under the DCF mechanism, whenever a WLAN station intends to transmit a packet, it first senses the channel. If the channel remains free for a distributed inter-frame space (DIFS) time ($34\ \mu\text{s}$ in 802.11a) plus an additional random time called the back-off time, it sends the packet. The duration of this random time is determined by each station individually as a multiple of a slot time ($9\ \mu\text{s}$ in 802.11a). If the channel is sensed as busy, the station waits until the channel becomes idle for a DIFS time, after which it starts a back-off process using the TBEB algorithm discussed in Subsection 3.1.3. If the packet is received correctly, the receiving station sends an ACK frame after a short inter-frame space (SIFS) time ($9\ \mu\text{s}$ in 802.11a). If the ACK frame is not received within a time-out parameter, a collision is assumed to have occurred, and the packet transmission is rescheduled according to the TBEB algorithm. Figure 6.3 illustrates the DCF mechanism in a WLAN network.

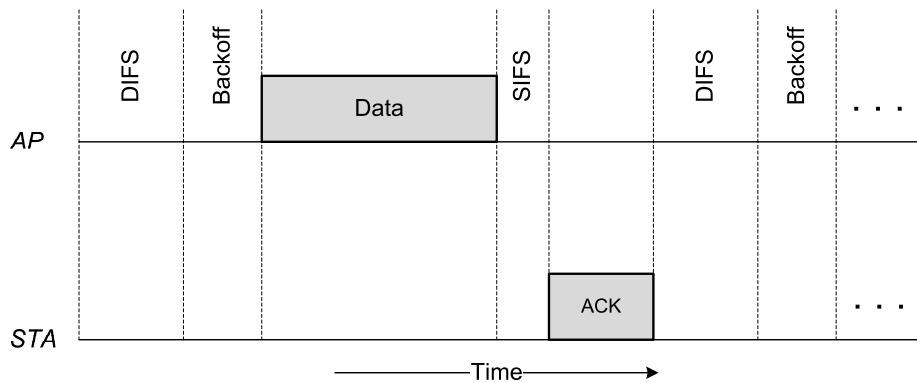


Figure 6.3: DCF mechanism in the IEEE 802.11-based WLAN networks.

Note that the current generation WLAN networks provide limited QoS support based on the hybrid coordination function (HCF) specified in the IEEE 802.11e standard [98]. However, unlike WiMAX networks, they only supports two major traffic types: the low priority BE

and background traffic; and the high priority voice and video traffic.

b) PCF Mechanism: A PCF enabled MAC has a contention free period (CFP) and a contention period (CP). Together they form a superframe, where the PCF is used for accessing the medium during the CFP and the DCF is used during the CP. A superframe starts with a periodic beacon management frame transmitted by the AP. The interval between two subsequent beacon frames is called the target beacon transmission time (TBTT). Figure 6.4 illustrates the PCF mechanism in a WLAN network.

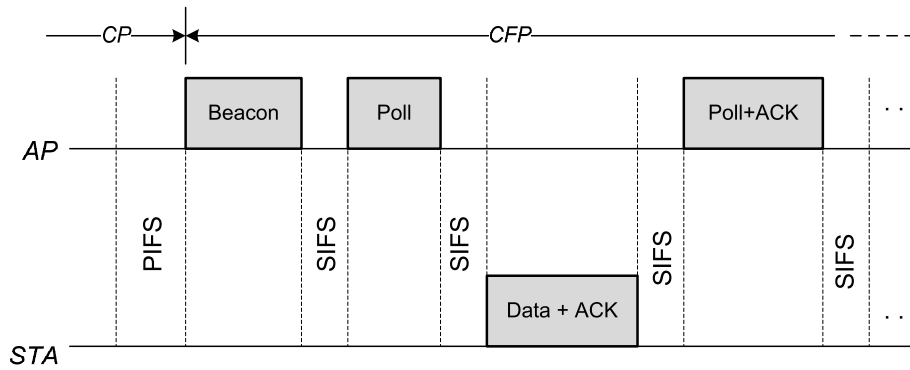


Figure 6.4: PCF mechanism in the IEEE 802.11-based WLAN networks.

During the CFP, the WLAN AP polls each station. If the AP does not receive a response from a polled station after waiting for a PCF inter-frame space (PIFS) time (25 μ s in 802.11a), it polls the next station. Upon receiving the poll, the polled station waits for a SIFS time, acknowledges the successful poll reception, and piggybacks the pending data with the ACK message.

The QoS-enabled version of PCF under the IEEE 802.11e standard is called the HCF controlled channel access (HCCA). Unlike PCF, here the AP can poll the stations during the CP intervals and takes into account the station-specific SF requirements in the grant scheduling process [98].

6.2 AMI Communications over the HetNet

To enable AMI communications over a WiMAX-WLAN HetNet, the WWR acts as a WiMAX-enabled DAP with a fixed outdoor antenna strategically mounted on a high structure. Figure 6.5 compares the projected cell range of a WiMAX network in the standalone mode (indoor SSs) and the HetNet mode (outdoor SSs) at different frequency bands, based on the coverage analysis conducted in Section 3.2. The results clearly show that the HetNet can provide significantly extended coverage compared to a standalone WiMAX network by bringing connectivity closer to the end devices. In terms of network performance, the major benefits come from the fact that the WWR can regulate random access load and perform flexible data aggregation to improve the signalling and data transmission efficiency of the WiMAX network.

In this section, we further compare the behaviour and performance of a WiMAX-WLAN HetNet with a standalone WiMAX network under the AMI applications. We particularly concentrate on the data aggregation process in the WiMAX-WLAN HetNet.

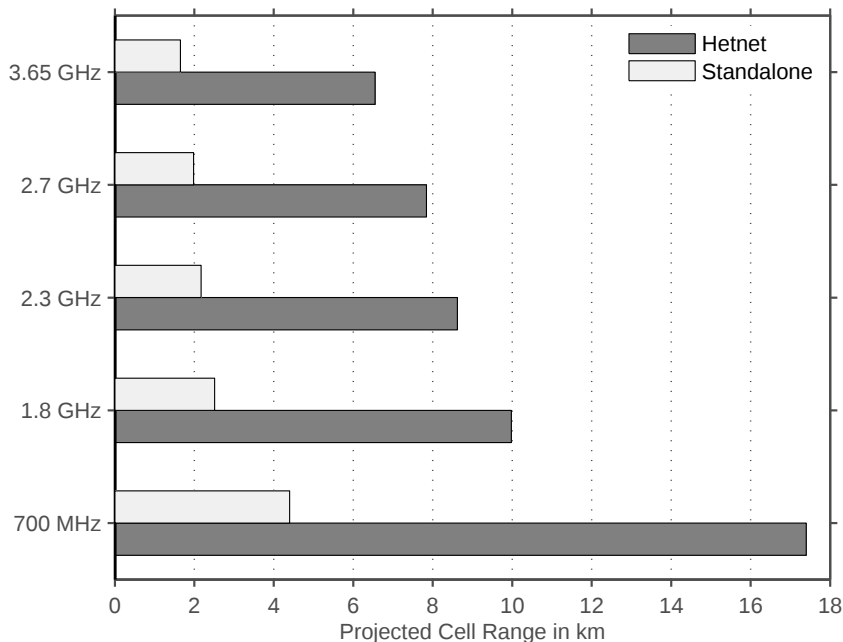


Figure 6.5: Projected cell range of a standalone WiMAX network and a WiMAX-WLAN HetNet.

6.2.1 Data Aggregation

The overall delay in the HetNet is comprised of delays in the WLAN network, in the WWR, and in the WiMAX network. Analytically,

$$d_{Hetnet} = d_{WLAN} + d_{WWR} + d_{WiMAX}. \quad (6.1)$$

The component d_{WLAN} is comprised of the MAC delay and the packet transmission delay in the WLAN network. The average MAC delay assuming DCF mode includes the time spent in back-off, DIFS/SIFS and ACK time-out intervals. The data transmission delay is determined by the ratio of the packet size (including the ACK packet) and the capacity (bit rate) of the channel. More details about the WLAN delay components can be found in [99].

The delay within WWR d_{WWR} is incurred when the MPDUs are exchanged between the WLAN and the WiMAX MAC layers (see Figure 6.1). It involves the scheduling delay in the WiMAX MAC layer, which requires at least one WiMAX frame to complete. From the WiMAX network's perspective, a WWR is a data buffer that stores multiple AMI data packets from the end devices. The BS can collect them by either using random access or polling mechanisms. Consequently, the component d_{WiMAX} depends on the scheduling service used by the WiMAX network.

First, let us consider a case where the polling mechanism is used. According to (3.2) and (3.3), the BS can easily meet the delay bound of t_{max} second for the WWRs by polling the end devices in every t_{max} second, which would generate a mean delay of $\frac{1}{2}t_{max}$. Unlike a standalone WiMAX network, the polling overhead is much less here, as the BS needs to poll only the DAPs, instead of all the end devices. Moreover, the polling mechanism allows the BS to perform flexible MPDU aggregation by adjusting the poll rate in accordance with the latency requirement. Figure 6.6 illustrates such an MPDU aggregation process in the HetNet. Here, the BS sends a regular poll to receive the aggregated MPDUs from the WWRs that were collected from the WLAN-enabled smart meters during the previous polling interval.

Next, we consider the case where the WWRs send MPDUs via the random access mech-

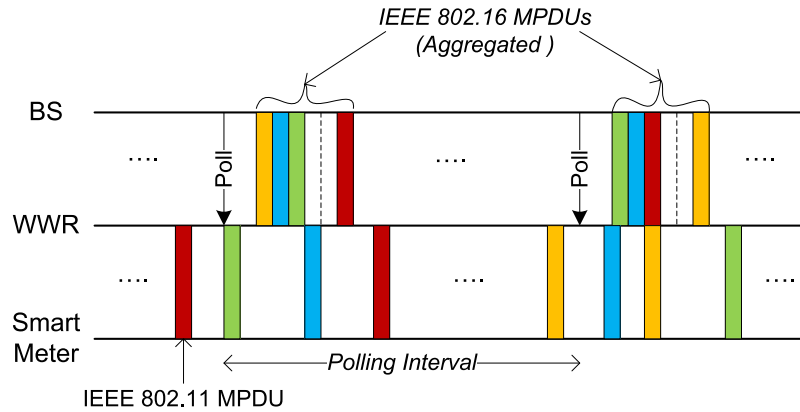


Figure 6.6: MPDU Aggregation by adjusting poll rate in the WiMAX-WLAN HetNet.

anism. This can be done either per flow or on aggregate basis [93]. Under the per flow approach, the WWR simply acts as an IP forwarding agent and initiates the BR procedure as soon as it receives an MPDU. In this case, the contention level in the WiMAX network remains the same as a standalone network, and therefore does not improve the random access delay. Under the aggregation approach, the WWR concentrates multiple MPDUs and sends a single BR for the aggregated MPDU. This reduces the contention level in the WiMAX network, and improves the random access delay. However, the concentrated MPDUs in the WWRs incur an additional aggregation delay, that is, the time elapsed between arrival of an MPDU and transmission of a BR, which may exceed the improved random access delay margin.

6.2.2 Performance Analysis

For comparative performance analysis, we use the same WiMAX network configuration and the same generic AMI application described in section 3.2. We further assume that each WLAN APs in the HetNet serves 60 smart meters as per the coverage analysis performed in [16]. Thus, there are 200 WLAN APs (i.e., WWRs) uniformly distributed throughout the coverage area of the WiMAX cell serving 12K smart meters. A simulation model is developed using the OPNET Modeller 16.0. The physical layer and the contention parameters of the WiMAX network under the HetNet mode are the same as Table 3.3 and Table 3.5, respectively. For WLAN deployment, we choose the 5 GHz unlicensed band with a 20

MHz system bandwidth, based on the IEEE 802.11a standard. The main reason for choosing this band is to reduce interference in WLAN network from the neighbouring WiMAX and home WLAN networks.

Now, we look into the performance of the AMI application with HetNet architecture. First we consider the case where the WWR uses the polling method to send the AMI data packets to the WiMAX BS. Since we assume that all smart meters are uniformly distributed among the APs, we concentrate our analysis to a single WLAN network connected with the WiMAX BS via the WWR. We further assume a poll rate of one second for the WiMAX network, considering a maximum delay bound t_{max} of one second. We set $t_{SIFS} = 16 \mu s$ and $t_{DIFS} = t_{SIFS} + 2 \times t_{slot} = 34 \mu s$ for the WLAN network according to [100].

Table 6.1: Delay Components (in ms) of the HetNet under the Polling Method

Interval (min)	d_{WLAN}	d_{WWR}	d_{WiMAX}
1	0.68	4.83	531.16
2	0.68	4.84	514.27
3	0.64	4.82	510.13
4	0.64	4.83	510.21
5	0.65	4.83	510.27

The various delay components of the HetNet at different reporting intervals are listed in Table 6.1. From the results, we see that the overall delay under the HetNet mode remains almost unchanged irrespective of the reporting intervals. Note that the delay in the WLAN network is much less than that of the WiMAX network. This is because the WiMAX delay depends on the polling rate of the system, which is one poll/sec for this case, with an average poll delay of 500 ms (see 3.3). In contrast, the slot duration in a WLAN network is just $9 \mu s$ (for IEEE 802.11a standard) and the use of carrier sensing mechanism along with the TBEB algorithm significantly reduces the collision level in a DCF-based WLAN network. Additionally, the data transmission time is very small as the channel bit rate is very high (due to relatively higher bandwidth in the unlicensed band) and the stations do not need to wait for a bandwidth grant to send the packet. While the delay in the WWR remains fixed, it increases slightly in the WiMAX network with smaller reporting intervals. This is because at small reporting intervals, more MPDUs are concentrated by the AP; therefore,

the BS needs more frames to allocate the bandwidth grants. In turn, this increases the data transmission delay, according to (3.2). However, it should be noted that both the polling delay and the polling overheads can be further optimised by using a multicast/broadcast polling technique.

To further depict the MPDU aggregation process, we conduct 5 simulation trials at different aggregation levels by varying the polling interval from 1 to 5 seconds at a fixed reporting interval of 1 minute. The average uplink MAC protocol overhead and percentage MAP signalling usage (percentage of the downlink subframe that is used by the MAP message) of the WiMAX network are plotted in Figure 6.7. Here we define the WiMAX UL MAC protocol overheads as the sum of MAC header, fragmentation and packing headers, and grant management sub-headers [58]. The corresponding mean queue size and queuing delay of a random AP is plotted in Figure 6.8.

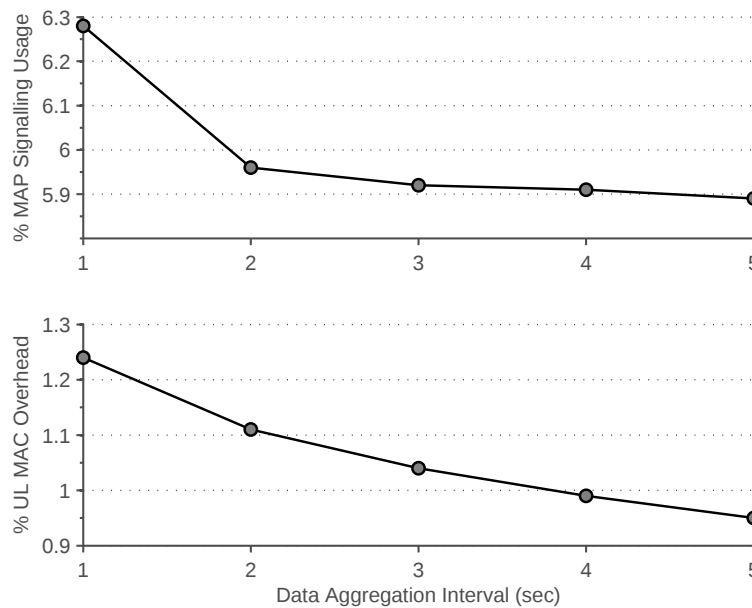


Figure 6.7: Percentage MAP signalling usage and UL MAC overhead of a WiMAX frame at various data aggregation intervals under the polling method.

From the results, we see that the MPDU aggregation by increasing the poll interval can reduce the MAC overheads and percentage MAP signalling usage at the cost of increased delay. The additional delay is due to the higher the aggregation level; more MPDUs are enqueued in the WWR, with a higher queuing delay, which is depicted in Figure 6.8. Note

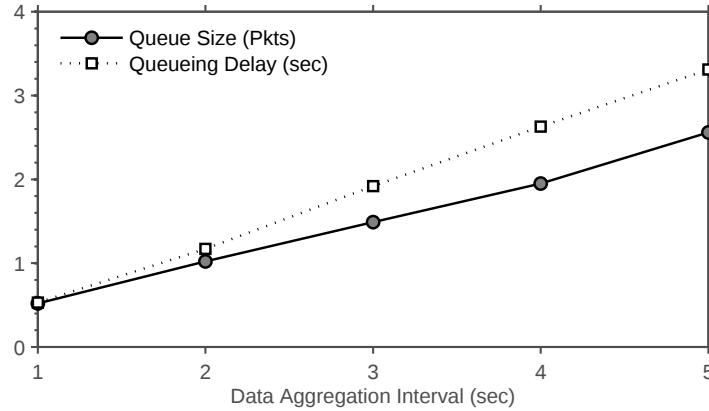


Figure 6.8: Mean queue size and queuing delay of a random WLAN AP at various data aggregation intervals [1, 2, 3, 4 and 5 seconds] under the Polling method.

that the improvement presented here is small since we have run our simulations with only one application: this could be more evident under multi-service scenarios. Moreover, here we have considered MAC layer aggregation only. Aggregation can also be done in the application layer, which can reduce both the IP and TCP/UDP overheads, and thus significantly improve the overall aggregation gain.

Now, let us consider the case where the APs send the packets using the random access mechanism. As discussed earlier, the AP can send the MPDUs in two modes, namely: transparent and aggregate. For the simulation trials, only the aggregate mode is considered as the transparent mode is same as the transmission under a standalone WiMAX network. Figure 6.9 plots the corresponding mean random access delay and mean data aggregation delay in the WiMAX network under different numbers of aggregated MPDUs for a reporting cycle of one minute.

The results show that MPDU aggregation drastically reduces the network delay as the number of contention retries falls rapidly; the higher the number of aggregated MPDUs, the lower the random access delay. When more MPDUs are aggregated, the number of BRs is reduced, which improves the random access delay in the WiMAX network. However, the higher the number of aggregated MPDUs, the higher the aggregation delays, since the MPDUs have to wait longer for the next BR. Consequently, the overall UL packet delay is high at the lower aggregation levels, due to a higher random access delay; it decreases

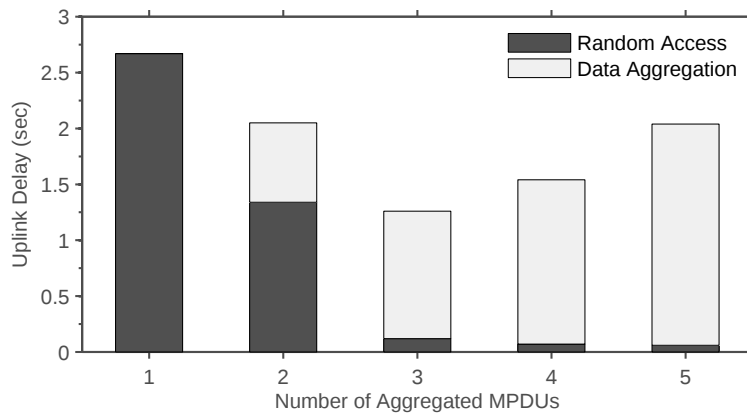


Figure 6.9: Mean uplink delay components at different MPDU aggregation levels under the random access method (a reporting cycle of one minute is considered).

in the medium aggregation levels due to a reduction in the random access delay; and then increases again at the higher aggregation levels due to an increase in aggregation delay. Note that irrespective of the aggregation levels, the polling mechanism (see Figure 6.7) outperforms random access in terms of uplink packet delay.

6.3 Synchrophasor Communications over the HetNet

As discussed in Chapter 2, a typical synchrophasor data collection network is comprised of several PMUs connected to a hierarchy of PDCs. The first level data transfer takes place between a group of PMUs and a local PDC. A local PDC is often located physically close to the PMUs where it aggregates and time-aligns the measurements from multiple PMUs and then sends it to higher level PDCs as a unified data stream [42]. Higher level PDCs (super PDCs) both aggregate and archive synchrophasor data. The hierarchical and distributed architecture of the synchrophasor network allows it to serve a hierarchy of systems in substation, utility, and interconnection levels.

A WiMAX-WLAN HetNet architecture is particularly well suited to facilitate such hierarchical data concentration and message transfer for synchrophasor communications. Similar to the AMI communications, the WLAN AP can aggregate traffic from multiple PMUs (residing within a substation area) into a single data burst, reducing the amount of ran-

dom access load and increasing the data and signalling efficiency of the WiMAX network. Moreover, the neighbouring PMUs within the substation WLAN can share synchrophasor measurements using multicast/broadcast techniques without involving the WiMAX network, saving its scarce radio resources.

6.3.1 Data Aggregation

To enable synchrophasor data communication in the HetNet, the local PDC functionality has to be incorporated in the WWR where: it receives PMU measurements via its WLAN interface; aggregates them in a single MPDU; and sends the aggregated data to the WiMAX BS via its WiMAX interface, as shown in Figure 6.2(b). Upon receiving the data, the BS transparently sends it to the next level PDC (super PDC for the two tier network here) via the IP backbone.

The state transition diagram for the data aggregation process in the WWR is illustrated in Figure 6.10. At each measurement cycle, the WWR stores the incoming PMU frames in the queue. Once measurements from all member PMUs have been received (or the measurement cycle has expired), the WWR constructs the PDC data frame by collating the PMU frames and sending it to its WiMAX MAC layer. Upon receiving an uplink bandwidth grant from the BS, the WWR then transmits the PDC data packet to a SPDC.

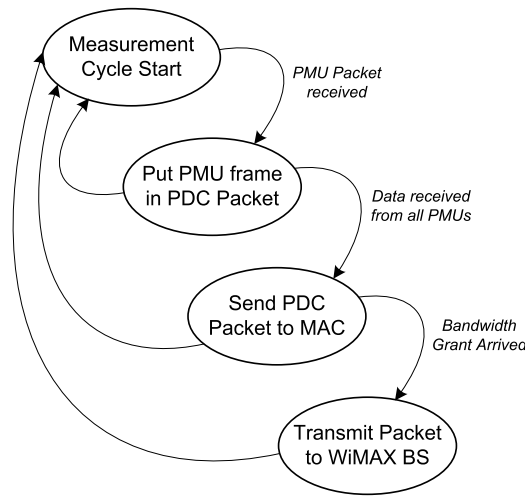


Figure 6.10: State transition diagram of the data aggregation process in the WWR.

The data aggregation gain due to protocol overhead reduction in the PDC for n member of PMUs can be expressed by the formula:

$$G_{Aggn} = 1 - \frac{L_{PDC}}{nL_{PMU}} = 1 - \frac{L_{Header}^{PDC} + n \times L_{Payload}^{PMU} + L_{Overheads}}{n \times (L_{Payload}^{PMU} + L_{Overheads})}, \quad (6.2)$$

where L_{Header}^{PDC} is the PDC header size, $L_{Payload}^{PMU}$ is the PMU payload size, and $L_{Overheads}$ is the sum of all lower-layer protocol overheads, i.e. UDP, IP, MAC headers and CRC.

Figure 6.11 plots data aggregation gain for different number of aggregated PMU data frames. Here, we assume $L_{Header}^{PDC} = 16$ bytes, $L_{Payload}^{PMU} = 48$ bytes, and $L_{Overheads} = 38$ bytes. From the plot, we see that the aggregation gain increases nonlinearly with the number of PMU data frames. This is because as the number of aggregated data frames increases, the corresponding PDC packet size increases; thus, the percentage of reduction in the protocol overhead decreases. Nonetheless, the overall aggregation gain is fairly high irrespective of the number of aggregated PMU data frames that can save a considerable amount of radio resources in the WiMAX network.

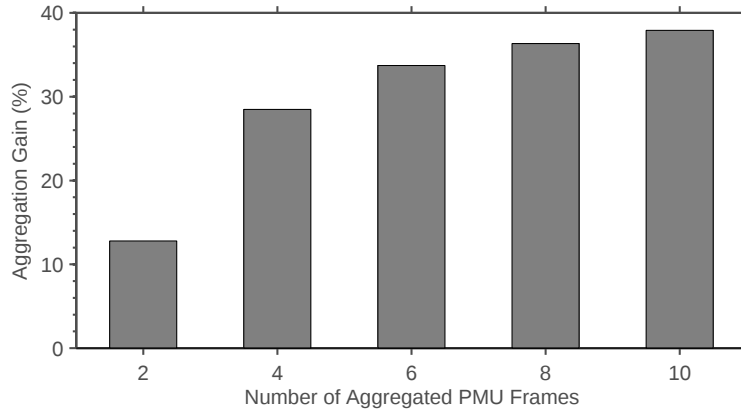


Figure 6.11: Data aggregation gain in the PDC for different number of aggregated PMU data frames.

6.3.2 Delay Budget

Since PMU communications involves fixed-size data packet transfer at a regular interval, it is expected that the PCF mechanism is the most appropriate for the first level data transfer between the PMUs and the local PDC (i.e., WWR) in the WLAN network. For the same reason, in a WiMAX network, the UGS scheduling is the obvious choice for communication between the WWR and the WiMAX BS. As the WiMAX-WLAN HetNet covers only the first two levels of data collection, the delay budget here is mainly comprised of two components: the delay between the PMU and the PDC, including the queuing delay in the PDC; and the delay between the PDC and the SPDC, as shown in Figure 6.12.

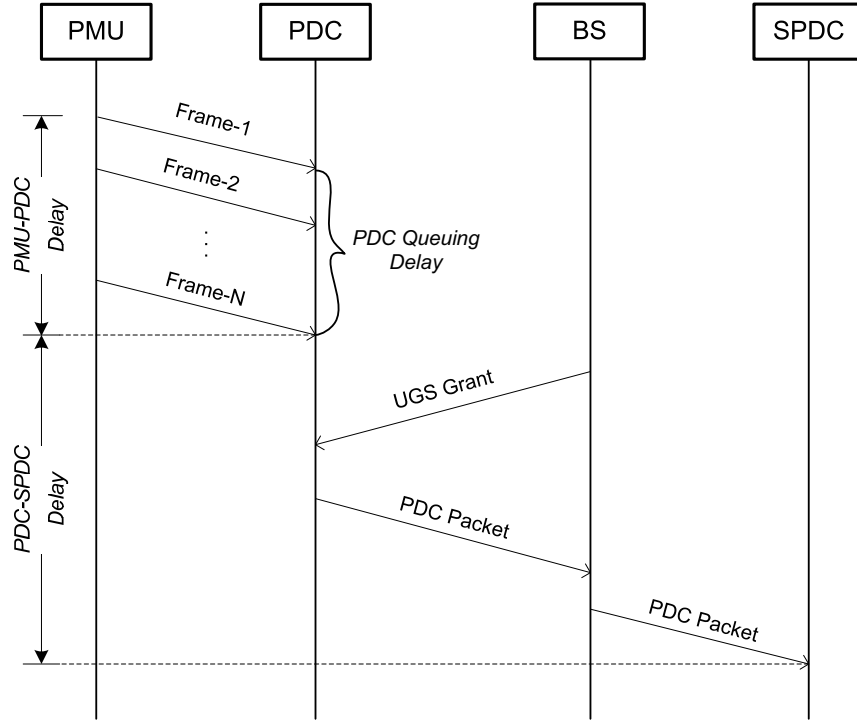


Figure 6.12: Synchrophasor data transfer within a WiMAX-WLAN HetNet.

The PMU-PDC delay $d_{PMU-PDC}$ mainly depends on the packet transmission delay in the WLAN network and the PDC queuing delay in the WiMAX network, which is simply the time difference between arrival of the first and the last PMU data frame. To ensure that the PMU measurement will be included in the current PDC data frame, the maximum PMU-PDC delay should be such that it is less than, or at least equal to the PMU reporting interval T_{PMU} . On the other hand, the delay between the PDC and the SPDC, i.e. $d_{PDC-SPDC}$ is

the sum of the times spent for packet transmission in the WiMAX network d_{WiMAX} and the delay in the IP backbone network to transfer the packets from the BS to the SPDC. As the IP backbone delay is negligible in comparison with the WiMAX delay, the end-to-end communications delay depends mainly on the delay performance of the WiMAX network. Thus, the end-to-end delay of synchrophasor communications can be expressed as

$$d_{e2e} = d_{PMU-PDC} + d_{PDC-SPDC} \leq T_{PMU} + d_{WiMAX} . \quad (6.3)$$

As such, the delay budget for the WiMAX network can be obtained by calculating the difference between the end-to-end delay bound and the PMU reporting interval as per (6.3). For example, if the PMU reporting interval is 40 ms and an end-to-end delay of 100 ms is required, the delay budget for the WiMAX network is 60 ms.

Note that a PMU may often act as an IED (e.g., a protective relay) within the substation domain (see Figure 2.6). In that case, it might need to broadcast its output streams to multiple neighbouring PMUs within the substation WLAN. In such a case, if a standalone WiMAX network is used, the corresponding PMU data packet will need to traverse through the WiMAX BS to the destination PMUs/relays. However, if a HetNet is used, the WWR can do domain specific multicasting within the substation WLAN, without involving the WiMAX network (see Figure 6.13). This will improve the end-to-end delay of PMU communications, and save significant amount of radio resources in the WiMAX network.

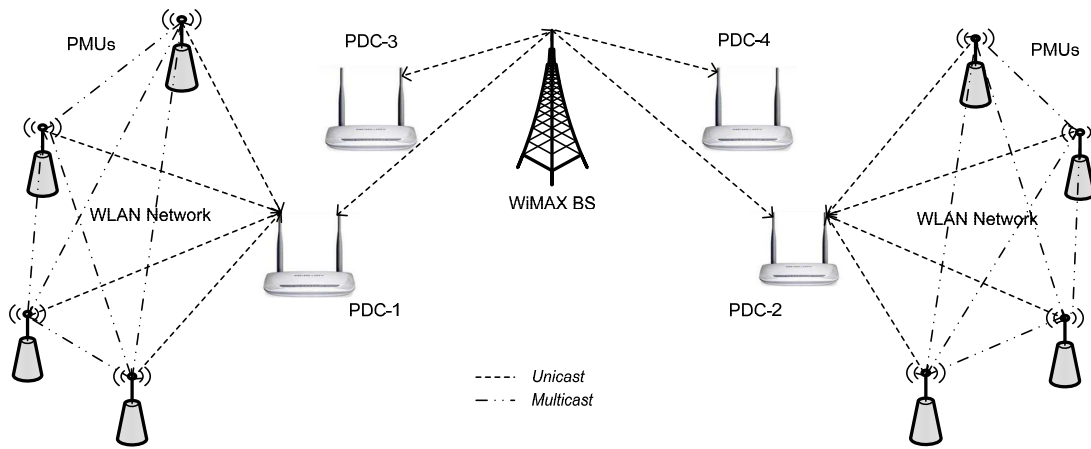


Figure 6.13: Synchrophasor multicast domains in the HetNet architecture.

6.3.3 Performance Analysis

To evaluate the behaviour and performance of synchrophasor communications over the HetNet architecture, a discrete event simulation model is developed using OPNET Modeller 16.0. The WiMAX simulation parameters are similar to those described in Subsection 3.3.4. The WLAN network is based on the IEEE 802.11a standard with 20 MHz bandwidth in the 5 GHz band.

The baseline scenario contains a local PDC (i.e. a WWR) with 10 member PMUs under the same substation WLAN network. We then increase the number of WWRs/PDCs to examine the delay performance of the application at different loads. We assume that the PDCs are randomly distributed over the entire WiMAX cell. The PMU payload size is assumed to be 48 bytes as listed in Table 2.2. The UDP/IP/MAC headers and CRC add 38 bytes to the MPDU as protocol overheads. On the other hand, each PDC data frame is comprised of a 16 byte header followed by 10 PMU data frames. This yields an MPDU size of 534 bytes (i.e., 10 PMUs x 48 byte PMU payload + 38 byte UDP/IP/MAC header + 16 byte PDC header) and a MAC data rate of 106,800 bps for each PDC connection.

We assume that every PMU under the local PDC is sending measurements at a rate of 25 frames/sec, which translates into a maximum delay bound of 40 ms for the WLAN network. We further assume an end-to-end delay requirement (i.e. from PMUs to the SPDC) of 100 ms for the synchrophasor application over the HetNet. To meet the 40 ms delay bound in the WLAN network, we set the TBTT interval to 20 ms and the CFP interval to 10 ms for the PCF mechanism in the WLAN network. In addition, we set the maximum latency for UGS scheduling in WiMAX network to 40 ms, so that the PDC can transfer the PMU data within each measurement interval. This translates into an MSTR of 104,800 bps for the WiMAX network according to (3.4).

We first look into the maximum PMU communications delay under the DCF and the PCF mechanism in WLAN network at different PMU group sizes. The results are plotted in Figure 6.14. From the results, we see that under both PCF and DCF mechanism, the delay increases linearly. However, under the DCF, the delay profile has a steeper slope, i.e. the delay increases sharply with the number of PMUs. This is because under the DCF, all the

PMUs try to access the medium simultaneously, which increases the network delay due to increased back-off and retransmissions. The problem is aggravated when more PMUs join the network further raising the instantaneous contention level. Conversely, under the PCF, the WLAN AP coordinates the sharing of the transmission media among the PMUs. This allows the PMUs to send their measurements without any contention. However, as the number of PMUs increases in the WLAN, the PMUs need to wait longer to get their chance, increasing the overall delay.

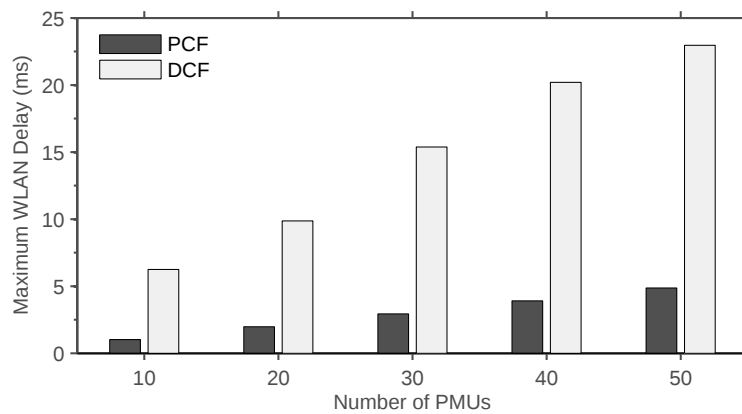


Figure 6.14: Maximum PMU delay under the DCF and the PCF mechanisms in WLAN.

Next, we look into the delay between the PDCs and the SPDC at different loads. In this part of the HetNet, the delay is mainly determined by the performance of the UGS scheduling service of the WiMAX network. The corresponding results are summarised in Table 6.2.

Table 6.2: PDC to SPDC Delay Statistics

No. of PDCs	Mean (ms)	Min. (ms)	Max. (ms)	S. Dev. (ms)
2	24.02	9.02	39.03	15.00
4	26.52	14.02	39.02	10.31
6	21.52	9.02	34.02	8.54
8	26.52	9.02	44.02	11.46
10	28.54	14.02	44.07	10.83

From the results, we see that the delay increases with the increase in the number of PDCs. In addition, for high PDC loads, although the maximum latency of the UGS scheduling is set to 40 ms, the maximum delay exceeds this margin roughly by one WiMAX frame (5

ms). This is because at high loads, the BS tries to fit as many grants as possible in one UL subframe. However, for this study the capacity of a UL subframe is 6144 symbols (as per Subsection 3.3.3), and the size of a PDC data packet is 4272 symbols (534 bytes). Hence, only one grant can be fully fitted into one UL subframe. The rest of the capacity is used to partially fit another grant that needs another UL subframe. Consequently, the UGS delay exceeds the maximum latency requirement roughly by one WiMAX frame.

Next, we look into the end-to-end delay between the PMUs and the SPDC. The results are summarised in Table 6.3.

Table 6.3: End-to-End Delay Statistics

No. of PDCs	Mean (ms)	Min. (ms)	Max. (ms)	S.Dev. (ms)
2	25.04	10.04	40.04	15.00
4	27.54	15.04	40.04	10.31
6	22.54	10.04	35.04	8.54
8	27.54	10.04	45.04	11.46
10	29.55	15.04	45.08	10.83

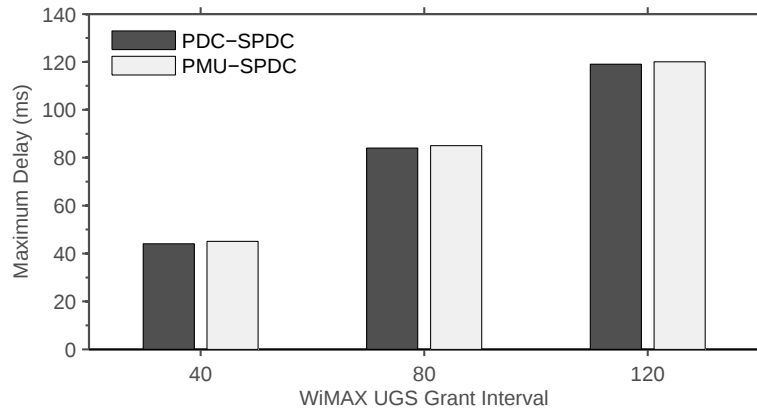


Figure 6.15: Synchrophasor network delay profiles at different WiMAX UGS grant intervals.

From the results, we see that the end-to-end delay closely follows the delay in WiMAX network. This is because, since we assumed the same number of PMUs in each PDC, the delay in the WLAN network remains fixed throughout all the scenarios. Additionally, the PDC queuing delay remains fixed at around 0.87 ms, which is the difference between

the maximum and minimum delay of the WLAN network. Nonetheless, the overall delay remains well within the maximum delay of 100 ms for the overall synchrophasor network.

It should be noted that the synchrophasor communications delay can be controlled by tuning the UGS grant interval of the WiMAX network, as shown in Figure 6.15; it shows that the end-to-end delay (i.e., PMU to SPDC) closely follows the UGS grant interval set in the WiMAX network.

6.4 Chapter Summary

In this chapter, we proposed a WiMAX-WLAN HetNet architecture to leverage the benefits of a WiMAX network in a Smart Grid M2M communications environment. The key tenets of the architecture: interworking and protocol adaptation, modes of operation, and QoS management, were also discussed. Moreover, the behaviour and performance of the HetNet was examined and compared with a standalone network for both AMI and synchrophasor communications. The results strongly advocate the use of the HetNet architecture in the Smart Grid, which may offer a range of additional benefits and improvements over the standalone WiMAX network.

To keep the scope of this work limited, we have used only a single application (i.e., AMI or synchrophasor communications) at a time for the performance analyses. However in practice, the proposed HetNet architecture will also need to support a wide range of QoS requirements in a multi-service Smart Grid communications environment. While WiMAX has an end-to-end QoS framework, the WLAN provides limited QoS support at the MAC layer based on the IEEE 802.11e standard. Hence, one of the most important challenges for the proposed HetNet architecture is to maintain end-to-end QoS for each traffic class, which is the product of QoS values within each subnetwork. We have considered this as a key future work of this thesis.

Chapter 7

Application-Specific Optimisations

From the reviews in Chapter 2, we see that the traffic patterns of the M2M devices that access the Smart Grid communications network and use its services at any given time heavily depend on their application characteristics. For instance, the AMR application can be programmed in a way that the smart meters will report their hourly meter readings at the same time. As a result, too many devices will try to access the network in a short time, and hence incur prolonged delay due to congestion in the random access plane. However, if the reporting intervals are randomised, such a situation is unlikely to occur. In short, it is possible to optimise some of the Smart Grid M2M applications in a way that their traffic profiles are adapted to the requirements of the underlying communications network. Although such an approach is relatively easy for simple applications like AMR; it is quite difficult for the other applications, such as EV charging, synchrophasors communications, and protective relaying. This is because the application programmer need to consider not only the communications impact, but also the cyber-physical characteristics of the application.

In this chapter, we present a number of application-specific optimisations to support the EV load management and synchrophasor-based protection and control schemes more efficiently over a WiMAX based Smart Grid communications network. The first application is an EV charging scheme that uses a random access aware load management technique to optimise its traffic requirements over a WiMAX network. The second application is a predictive synchrophasor communications scheme that dynamically switches between a

normal and an expedited reporting rate to maintain a minimum communications load in the WiMAX network. The last application is a hybrid microgrid protection scheme that uses a combination of differentiation protection technique and an adaptive protection technique to reduce the overall communications load of the WiMAX network.

The rest of the chapter is outlined as follows. Sections 7.1, 7.2 and 7.3 present the proposed EV load management scheme, adaptive synchrophasor communications scheme and hybrid protection scheme respectively, discuss their architectures and working principles, and conduct simulation studies to demonstrate their performances over a WiMAX network. Section 7.4 concludes the chapter ¹

7.1 Network Controlled EV Load Management

As discussed in Subsection 2.2.3, the communications requirement of a smart EV charging scheme depends directly on the load management technique used. However, although a vast number of EV charging schemes are available, their communications strategies are not sufficiently well-defined. Most of these schemes employ optimisation-based techniques, where a charging controller varies the charging rates of the connected EVs based on a set of objective functions, such as minimising power losses, voltage deviations, and charging cost. However, they assume perfect knowledge about the behaviour of the EVs to solve the optimisation objectives [102]. A key requirement here is that the energy state (e.g., charging rate, battery status) of each EV is known ahead of the optimisation period [103]. In such cases, if a communications network is used to obtain energy state reports from all the EVs synchronously right before the next optimisation period, this may overwhelm the entire network with a burst of network access attempts. This may lead to congestion in the random access plane, resulting in an increased delay. In turn, this may result in a poor optimisation outcome.

¹The contents of this chapter has been published in two conference papers, one in proceedings of the *Australasian Universities Power Engineering Conference (AUPEC)* in Sept 2013 [25], and another in proceedings of the *IEEE International Conference on Smart Grid Communications (SmartGridComm)* in Nov 2013 [101], as well as another submitted paper in the *IEEE Transaction on Smart Grid* in May 2014.

To accurately analyze the communications requirements of an EV charging application and optimise it over a WiMAX-based Smart Grid communications network, in this section we present a network controlled load management scheme for domestic charging of EVs. Motivated by the work of Brooks *et al.* [23], the proposed scheme allows only a limited number of vehicles to be charged at a maximum rate using a stream of ‘energy bursts’; thus, it uses the benefits of statistical multiplexing to accommodate a higher number of vehicles and to provide differentiated energy supply against a time-varying available power. The controller needs to communicate with only a small number of vehicles at a time. While this reduces the overall communications load of the system, it also provides an opportunity to evenly distribute the communications load by randomising the uplink packet transmission times of the EVs. Thus, the scheme is highly efficient and scalable in terms of communications overheads; therefore, it can be served by a multi-service smart grid communications network along with other applications.

7.1.1 Basic Concept

The proposed EV charging scheme adopts a request/grant mechanism to allocate energy among the contending vehicles. Energy is allocated using small bursts dimensioned in time based on a time quantum Δt and in power based on the maximum charging rate $\bar{p}$ of a given EV as shown in Figure 7.1. Thus, the available energy in a single burst is given by

$$\Delta E = \bar{p} \Delta t. \quad (7.1)$$

Once an energy burst has been consumed, the EV sends a request message and waits for the next round of allocation. Note that $\bar{p}$ is bounded by the capacity of the power outlet which typically remains fixed for a given distribution grid. Hence, the amount of energy within a single burst also remains fixed for a given time quantum. Thus, the total energy received by the i^{th} EV over the charging interval T can be obtained by

$$E_i(T) = \int_0^T p_i(t) dt = \sum_n \Delta E, \quad (7.2)$$

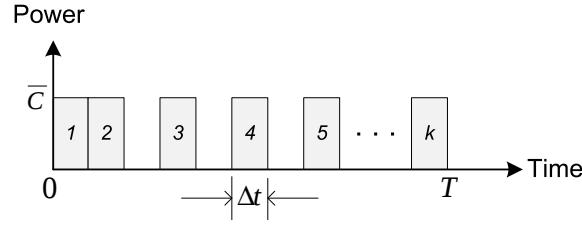


Figure 7.1: Centralised EV charging using small energy bursts.

where $p_i(t)$ is the instantaneous charging rate and n is the total number of energy bursts received by an EV.

The main objective of the proposed load management scheme is to allocate the available energy bursts in a manner that the total power consumption from the aggregated EV load remains below the time-varying available power $\bar{P}$. Assuming the available power remains constant over the interval Δt , i.e. $\bar{P}(t + \Delta t) \rightarrow \bar{P}(t)$, the number of energy bursts k at a given time t is therefore:

$$k \leq \frac{\bar{P}\Delta t}{\Delta E} = \frac{\bar{P}}{\bar{p}} \quad \text{for } \forall t. \quad (7.3)$$

Consider an energy constrained system where the number of EVs N is greater than the number of available energy bursts k . Since the number of request messages is directly related to the number of allocated energy bursts, the communications load l of the system towards a given direction in terms of packet per second is given by

$$l = \frac{k}{\Delta t} \leq \frac{\bar{P}}{\bar{p}\Delta t} \quad \text{for } \forall t. \quad (7.4)$$

From (7.4), we see that the communications load of the system does not depend on the number of vehicles, rather it depends on the available power in the system.

Now, let the system have a mean service rate μ such that $\mu = k/\Delta t$ bursts per second. Since an EV will get the next round of allocation only after all the remaining EVs have been allocated, the mean waiting time of the system is given by

$$W = \frac{N}{\mu} = \frac{N\Delta t}{k} \quad \text{for } \forall t, \quad (7.5)$$

where N is the total number of EVs in the system. From (7.5) we see that for a given

available power, as the number of vehicles increases, the waiting time for each round of allocation increases, such that the cars need more time to reach a certain SoC level.

7.1.2 Communications Architecture

The application message exchanges for the proposed EV load management scheme is similar to the one depicted in Figure 2.5. However, here the ‘SoC update’ message is used by the EVCS to request energy and the ‘continue’ message is used by the EVCC to allocate energy. The communications architecture of the scheme is illustrated in Figure 7.2. The EVCC, which is in charge of the energy management process, comprises the following four subsystems:

- **Communication Interface:** facilitates real time information exchange with the EVCSs through the Smart Grid communications network (i.e., WiMAX),
- **Request Classifier:** extracts the energy request information from the initial charging request and SOC update messages (see Figure 2.5) and puts them on one or more service queues based on the energy scheduling criteria,
- **Service Queues:** stores the energy request information of the contending vehicles, and
- **Energy Scheduler:** allocates energy bursts among the contending vehicles based on their energy requirements and the energy scheduling algorithm.

For energy scheduling, we use the simple fair queuing (SFQ) algorithm, depicted in Figure 7.3. As seen from the figure, a request classifier files the incoming energy requests from both the newly arrived (initial charging request) vehicles and the already connected vehicles (SoC update message) into a service queue in the FIFO (first-in-first-out) manner. The energy scheduling process is triggered upon arrival of an energy request in the service queue. At the beginning of the process, the scheduler first checks whether sufficient power is available for sending at least one energy burst; should the condition be met, the scheduler then looks into the service queue and serves the head of the queue. At the end of the charging cycle, the request is removed from the queue, and an end message is transmitted.

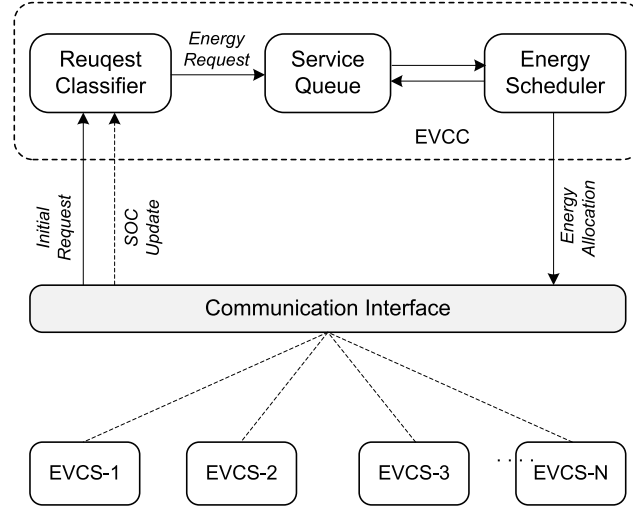


Figure 7.2: Communication architecture of the proposed EV load management scheme.

The process is repeated as long as power is available in the system. Thus, under the SFQ scheduling, the vehicle that arrives earlier receives the higher number of energy bursts.

To prevent excess energy delivery after an EV battery has reached its full capacity $\bar{E}$, the scheduler computes the expected SoC before each round of allocation. If the expected SoC exceeds 100%, the scheduler adjusts the duration of the last energy burst using the below formula:

$$\Delta\hat{t} = \frac{1}{\bar{p}}[100 - \bar{E} \cdot SoC(t)], \quad (7.6)$$

where $\Delta\hat{t}$ is the truncated time quantum for the last energy burst. In the case where multiple EVs send energy requests simultaneously or the available power increases sharply, the controller may have to allocate a number of energy bursts at the same time which may increase the instantaneous communications load. To avoid such a scenario, a small randomisation interval $\delta t (= 1 \text{ sec})$ is used, such that if the time between two successive energy burst allocations is less than δt , the controller will delay the next allocation by δt . By using the randomisation interval, the controller is able to regulate the maximum communications load of the system, which is given by $1/\delta t$ packets/sec.

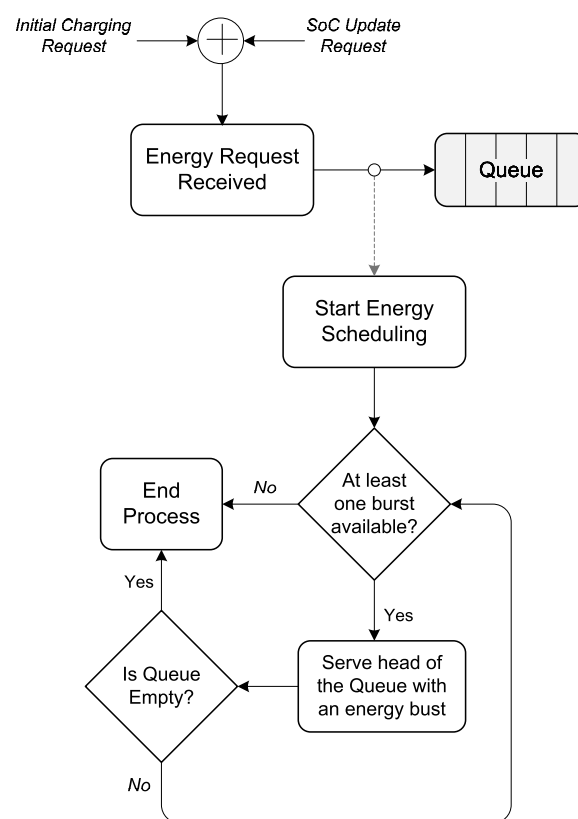


Figure 7.3: Flow chart of the SFQ scheduling algorithm.

7.1.3 Modelling & Assumptions

To examine the behaviour and performance of the proposed EV load management scheme, a discrete event simulation is developed using OPNET Modeller 16.0. In particular, we consider the case of using the EV fleets for night-time valley-filling of the daily load curve, where users are subscribed to a flat-rate contract and the charging process is independently controlled. The simulation scenarios are based on a small distribution substation serving 2,000 cars. The baseline scenario assumes an EV penetration level of 25 per cent that yields an EV fleet size of 500. The number of EVs is then increased to 750 (37.5 % penetration) and 1,000 (50 % penetration) to examine the effect of higher adoption rate. The following assumptions are made:

a) Charging Window & Vehicle Arrival Patten: A charging window of 8 hours is considered that starts at 10 pm at night and ends at 6 am in the morning. Since by 10 pm most EVs will be plugged into the charging system, we use a heavy tailed (shape, $\alpha = 5$) Pareto distribution function for modelling the vehicle arrival pattern. According to the distribution, 90 per cent of the cars will be connected to the system within the first hour of the charging window. The remaining cars will arrive late with the last one coming as late as 4 am in the morning.

b) Charging Requirements: For the sake of simplicity, we assume that all vehicles have a similar battery size of 16 kW, with a capacity rating of 80 per cent and the initial SoC of the batteries are uniformly distributed between zero to 50 per cent. We further assume that the cars will be charged via a regular 240V/15A power outlet (considering the case of Australia) and the charger has an energy transfer efficiency of 90 per cent to compensate for losses during the charging session.

c) Energy Budget: We assume that in the Smart Grid, the utility control centre will provide each distribution substation with a certain energy budget to satisfy the EV charging loads at a given period. Since the loads remain fairly constant during the off-peak hours

at night, the grid is mainly supplied by the base load power plants. Hence, we model the available energy budget as a three-step function, such that if the maximum available power is P , then the energy budget at time t is given by

$$E_{Budget} = \begin{cases} \frac{1}{2}P, & \text{from 10 p.m. to 12 a.m.} \\ P, & \text{from 12 a.m. to 4 a.m.} \\ \frac{1}{2}P, & \text{from 4 a.m. to 6 a.m.} \end{cases} \quad (7.7)$$

Without any loss of generality, we assume that the system has enough power to serve the baseline scenario of 500 EVs within the stipulated charging window. Now, the total energy requirement for charging a given EV fleet can be roughly calculated by

$$E_{Total} = \sum_{i=1}^N (1 - \alpha_i) E_i \frac{\beta_i}{\eta_i}, \quad (7.8)$$

where N is the number of EVs in the fleet, α is the initial SoC, E is the battery size, β is the battery rating, and η is the efficiency of the charger. Assuming a flat charging profile with the parameters specified and using (7.8), the total energy requirement for 500 EVs becomes 5334 kWh (approximately). Thus, over the charging window of eight hours, the available power is distributed into three segments: 450 kW between 10 pm to 12 am; 900 kW between 12 am – 4 am; and 450 kW between 4 am – 6 am. In addition, a time quantum Δt of 5 minutes is assumed for this study.

d) Communications Network: We design the communications part of the model as a single cell suburban WiMAX network based on the IEEE 802.16-2009 standard. The EVCC is modelled as a remote server located in the core network. The EVs are equipped with a WiMAX SS to perform IP based communications with the EVCC via the BS. For UL data transmission, we use the BE scheduling service flow considering that the bursty SoC update messages from the EVs. To imitate a multi-service smart grid network, we simulate an exponentially distributed background BE traffic of 100 pps (packets/sec). The key WiMAX parameters are same as in Table 3.3. In addition, we examine the use of both

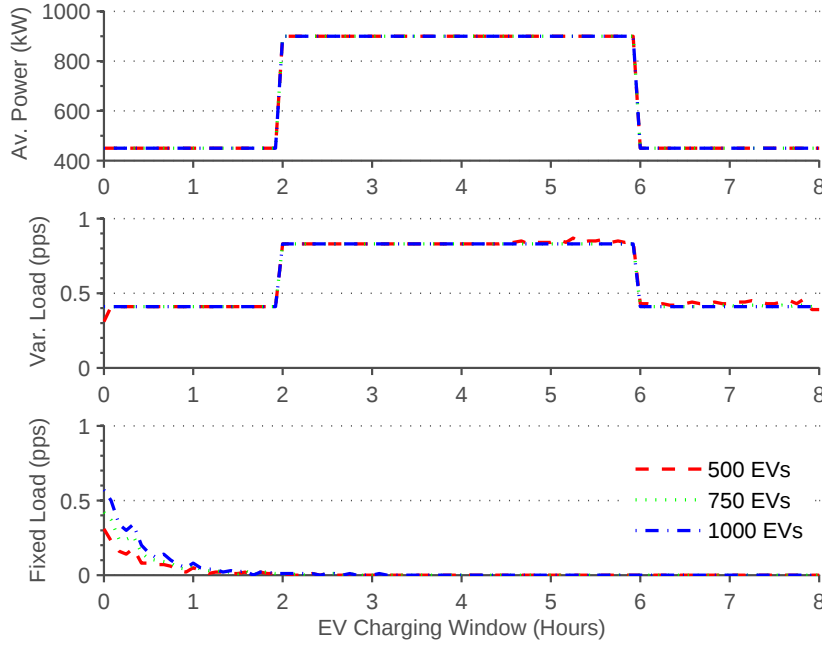


Figure 7.4: Initial charging request and SoC update messaging load for various EV fleet sizes.

ARQ and HARQ retransmission schemes to recover the lost packets. The ARQ retransmission interval is set to 0.1s; 8 HARQ channels are assumed with maximum 4 retransmissions per packet. The contention parameters for the BE scheduling service is set as $T_{16} = 4$ frames, $W_0 = 2^2$, $W_f = 2^8$ and $m = 16$.

7.1.4 Results & Analysis

As discussed in Chapter 2, the overall communications load of the EV charging application comprises a fixed and variable messaging components. Figure 7.4 plots the WiMAX network load for these two messages at varying EV charging loads. Note that as both of these messages are associated with a DL reply message (see Figure 2.5), the actual load is twice that of the one shown in the figure. From the results, we see that the initial charging requests from the EVs follows a Pareto distribution as assumed in Subsection 7.1.3.

Moreover, the initial request load increases as the number of vehicles increases. In contrast, the SoC update load remains the same independent of the number of vehicles in the system

Table 7.1: SoC Update Message latency (For 500 EVs)

Retrans. Scheme	Mean Latency (ms)	Max. Latency (ms)
ARQ	34.23	315.0
HARQ	29.95	330.02
Error Free	26.37	300.0

and closely follows the available power as indicated by (7.4). In addition, the load is very small (less than one packet/sec) and is evenly distributed throughout the charging window, as the controller scatters the allocation of the energy bursts through the randomisation interval δt . This corresponds to very low radio resource utilisation, which is only around 0.04 per cent for the WiMAX network.

Table 7.1 lists the latency distribution of the SoC update messages over an error-free channel and an error-prone channel with ARQ and HARQ retransmission schemes. The results show that the overall message latency increases due to retransmissions of the corrupted packet. Between ARQ and HARQ, the latter provides a better delay performance. This is because the ARQ relies on a feedback mechanism to detect a packet error and initiate retransmission. Conversely, the HARQ sends each data packet with FEC and the receiver uses both the retransmitted and the erred packets to reconstruct the original packet which significantly reduces the number of retransmissions.

Since at steady-state, the controller allocates another energy burst only after one has been consumed, the latency of the SoC update message $\bar{d}_{SoC}$ causes a slight loss of energy at each round of allocation which is given by

$$E_{Lost} = \bar{d}_{SoC} \cdot \Delta E \cdot l(t).$$

However, since $\bar{d}_{SoC} \ll \Delta t$, the amount of this loss is negligible. For example, the loss is 0.044 kWh, which is less than 0.01% for an available power of 450 kWh, an SoC update load of 0.41 packets/sec (as in Figure 7.4), and a mean latency of 30 ms (as in Table 7.1).

7.2 Predictive Synchrophasor Communications

One of the key objectives of a WAMS is to support predictive protection and control schemes in the Smart Grid to prevent cascading failures due to voltage instability events, such as power swing and voltage oscillation [42]. Such schemes utilise synchrophasor measurements from the PMUs to determine the characteristics of a voltage instability event and predict an out-of-step condition using a time-series expansion algorithm. When the voltage of a load-bus reaches close to its instability point, an aggressive synchrophasor reporting rate may enhance the performance of the time-series algorithm to better predict the transient response of the system, and thereby allow the control system to execute appropriate actions within the incubation period [104]. However, an aggressive reporting rate comes at a cost of higher bandwidth requirement. This is particularly challenging for a multi-service Smart Grid communications network, where the locked radio resources may cause bandwidth-starvation to the other applications. This problem can be overcome by using an adaptive reporting scheme where the controller dynamically switches between a normal and an expedited reporting rate based on the output of the prediction algorithm. Such a scheme is reasonable since a voltage instability event in a power system is a highly random phenomenon. For most of the time, the system will operate normally and the synchrophasor reporting rates can be maintained at an optimum level.

In this section, we propose an adaptive synchrophasor reporting scheme that dynamically switches between a normal and an expedited reporting rate to keep its communications load minimum. The switching of the reporting intervals is performed based on a prediction framework that periodically estimates the maximum loadability and voltage instability point of the system by using the local synchrophasor measurements.

7.2.1 Prediction of Voltage Instability

Let us consider a power system network where a local bus consuming power $P_t + jQ_t$ at time t . The local voltage stability monitoring techniques [105] generally convert the whole power system network into a time dependent two-bus Thevenin equivalent system

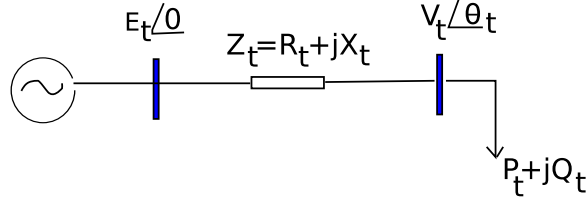


Figure 7.5: Two-bus Thevenin equivalent system. Load bus and the rest of the system is represented by a voltage source (E) and transmission line.

(see Figure 7.5) where a generator $E_t \angle 0$ supplies power $P_t + jQ_t$ to the local bus through a transmission line $Z_t = R_t + jX_t$, where R_t is the resistance and X_t is the reactance of the bus. The voltage phasor of local bus is $V_t \angle \theta_t$. The real and reactive power transfer to the local bus can be expressed as [106]:

$$P_t = \frac{V_t}{|Z_t|^2} [(E_t \cos \theta_t - V_t)R_t - E_t \sin \theta_t X_t] \quad (7.9)$$

$$Q_t = \frac{V_t}{|Z_t|^2} [(E_t \cos \theta_t - V_t)X_t + E_t \sin \theta_t R_t] \quad (7.10)$$

In the above two equations, both P_t and Q_t are functions of five variables; they are E_t, V_t, θ_t, X_t and R_t . However, it is very difficult to predict the future behaviour of those variables, and hence the profile of P_t, Q_t . Different assumptions have been used to estimate the future profiles of P_t and Q_t [105, 106]. A compact expression of maximum active and reactive power transfer limit to a local bus has been developed in [106] by assuming R_t/X_t is small and E_t, X_t and Q_t or P_t are constant:

$$P_{\max} = \sqrt{\frac{E_t^4}{4X_t^2} - Q_t \frac{E_t^2}{X_t}}; \quad \text{and} \quad Q_{\max} = \frac{E_t^2}{4X_t} - P_t^2 \frac{X_t}{E_t^2}.$$

Different assumptions have been used to estimate the future profiles of P_t and Q_t [106, 105]. However, the method cannot perform well when the load in a local bus is voltage-sensitive,

i.e. ZIP load [105]. The ZIP load can be modelled as [107, 105]:

$$P_t = k_t(P_c + a_1 V_t + a_2 V_t^2), \text{ and} \quad (7.11)$$

$$Q_t = k_t(Q_c + b_1 V_t + b_2 V_t^2), \quad (7.12)$$

where $P_c, Q_c, a_1, a_2, b_1, b_2$ are constants and the value of k_t is an independent demand variable called ‘loading factor’ [105]. The loading factor k_t can be used as an indication of the total load demand changing with time. By combining (7.11) and (7.12) a relationship is developed in [105]:

$$V_t^4 + V_t^2[2(P_t R_t + Q_t X_t) - E_t^2] + Z_t^2(P_t^2 + Q_t^2) = 0. \quad (7.13)$$

Let us draw the P-V characteristics curve from (7.13) by assuming Q_t, E_t, X_t are constant (see Figure 7.6). We can also draw a set of ZIP load characteristics curves (P_t - V_t), by using (7.11) for different values of k_t . It has been shown in [105] that the system remains voltage stable for a loading factor k_t if the ZIP load characteristic curve (7.11) for k_t intersects the P-V curve at least two points. The system will be voltage unstable, when for some critical loading factor k_* , the ZIP curve becomes tangent to the system P-V curve.

In Figure 7.6 we see that the voltage of the system will be unstable when the load demand of the local bus reaches the point when the voltage and active power supply become V_C and P_C p.u., respectively. Let the power transfer to the local bus at time t is denoted by P_t . When the load demand increases, the power transfer to the local bus increases and the voltage decreases continuously. At some point of load demand, the system transfer maximum power $P_{\max}$ to the local bus. After that point, both power transfer to the local bus and bus voltage decreases simultaneously with increasing load demand. Finally, the voltage collapse occurs at V_C and P_C . The interesting observation is that the voltage instability occurs beyond the maximum power transfer point $P_{\max}$.

Let us assume the voltage and current phasors of a load bus are sampled using PMU at the discrete sampling instants $\{t_r\}_{r=1,2,\dots}$. Also, the load demand of the bus is increasing with time. Consequently, the system voltage will be collapsed at time C with power and

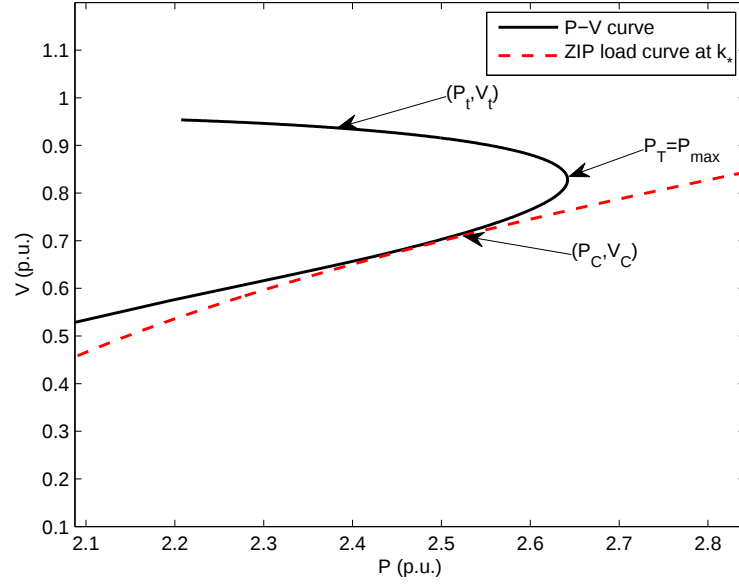


Figure 7.6: PV curve of a typical load bus and ZIP load curve at critical k_* .

voltage being (P_C, V_C) . Assume that the current time instant be t_ℓ , and a local bus is likely to consume maximum power $P_{\max}$ at time $T > t_\ell$ and will be collapsed at time $C > T, t_\ell$. Our target is to estimate the future incidents of $P_{\max}, T$ and C from current time t_ℓ . Suppose the prediction of $P_{\max}, V_{pm}, T$ and C at time t_ℓ are $P_{\max}^\ell, V_{pm}^\ell, t_{pm}^\ell$ and t_*^ℓ , respectively. We want to develop an efficient prediction algorithm such that

- $\|P_{\max} - P_{\max}^\ell\|, \|V_{pm} - V_{pm}^\ell\|, \|T - t_{pm}^\ell\|$ and $\|C - t_*^\ell\|$ are small, and
- $(T - t_\ell)$ and $(C - t_\ell)$ are large.

The details of the voltage instability prediction algorithm is provided in Appendix D.

7.2.2 Adaptive PMU Reporting

Due to the complex dynamic behavior of the power system, the change of voltage of a bus ΔV does not remain same over time which makes it quite difficult to predict. Moreover, the value of ΔV changes rapidly when the voltage of a load bus approaches the voltage instability. Hence, an aggressive PMU reporting rate is helpful to precisely predict a voltage

instability condition. Note that such an adaptive PMU reporting scheme is consistent with the IEEE C37.118.2 standard for synchrophasor data transfer [45]. The standard allows the PDC to send commands to the PMUs using a command (CMD) frame. The frame has a 2 byte field that can be used to change the synchrophasor reporting states.

From our previous study in Section 3.3, it is clear that the UGS scheduling service of WiMAX is the best candidate to transfer the PMU data packets. However, to adaptively change the PMU reporting intervals, first the concerned PMU(s) needs to be notified through a CMD message to change its data output. After that, the service flow parameters of the corresponding UGS data connection need to be dynamically modified so that the newly generated bandwidth requirements are met. Such a dynamic negotiation is allowed in the existing WiMAX QoS framework by using the DSC procedure described in Subsection 3.1.2. The DSC procedure supports both the SS initiated DSC and the BS initiated DSC.

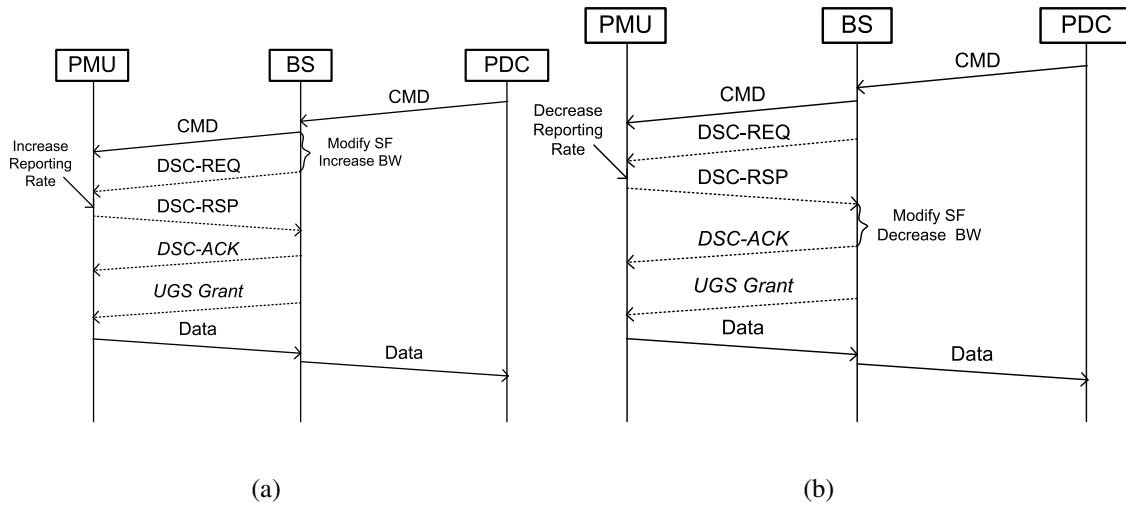


Figure 7.7: DSC negotiation for: (a) increasing PMU reporting rates, (b) decreasing PMU reporting rates within a WiMAX UGS connection.

For adaptive PMU reporting, as the upstream PDC is changing the connection properties, the BS needs to initiate the DSC procedure. An important consideration is that the required bandwidth change needs to be sequenced between the PMU and the BS to prevent any packet-loss. According to the IEEE 802.16 standard, if a UL service flow's bandwidth is being increased, the BS increases the bandwidth scheduled for the service flow first and then the SS increases its payload bandwidth. In contrast, if a UL service flow's bandwidth

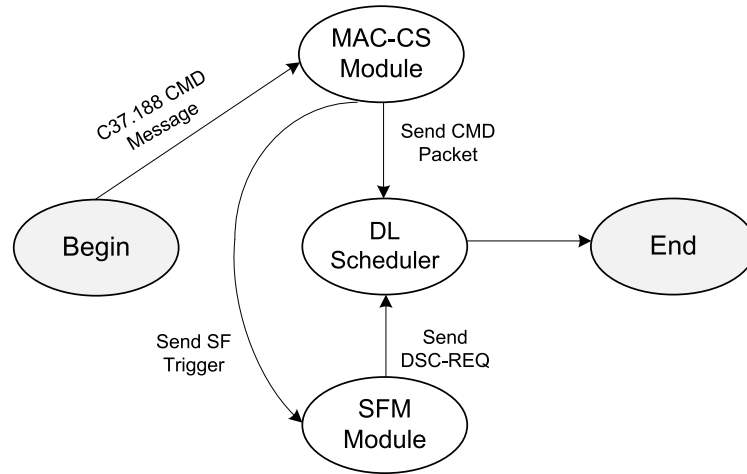


Figure 7.8: State transition diagram of SF negotiation in the WiMAX BS for the adaptive PMU reporting scheme.

is being reduced, the SS reduces its payload bandwidth first and then the BS reduces the bandwidth scheduled for the service flow. The process is depicted for the adaptive PMU reporting scheme in Figure 7.7.

As seen from Figure 7.7, the PDC initiates the change procedure by issuing a CMD message to the PMU via the WiMAX BS. When the BS receives the packet, its MAC CS distinguishes the CMD message using a protocol matching criteria such as IP address, port, or type-of-service. Upon recognising the CMD message, the CS sends a SF trigger to the service flow management (SFM) entity in the BS according to the procedure specified in [108]. The CMD packet is processed normally and transmitted to the desired PMU. On the other hand, the SFM extracts the CID of the particular PMU and the new reporting interval from the CMD message and determines the required SF parameters (data rate and latency). The BS then sends a DSC request (DSC-REQ) message to the SS of the corresponding PMU. The state transition diagram of the SF negotiation in the WiMAX BS is depicted in Figure 7.8.

Upon receiving the DSC-REQ message, the SS modifies the SF parameters and notifies the change via a DSC response (DSC-RSP) message. Finally, the BS confirms the change via a DSC acknowledgment (DSC-ACK) message. The BS increases the bandwidth of the UGS connection before sending the DSC-REQ message and decreases it after receiving the DSC-RSP message. This is required to ensure seamless changeover between two reporting

rates without any loss of packets.

7.2.3 Simulation Results

To examine the impact of the proposed adaptive reporting scheme on WiMAX network, a single cell simulation model is developed using the OPNET Modeller 16.0. The simulation scenario is based on a New England 39-bus test system covered by the WiMAX cell. Within the 39-bus system, only 8 PMUs are required to make the system observable [109]. The simulation parameters are the same as described in Subsection 3.3.4. We further assume that the PMU data transmission and UGS grants are synchronised using the FL/FLI technique, described in Section 3.3.

While, at a given instance, only a particular bus or a set of buses may be at risk of voltage instability. Considering the worst case scenario, we assume that the PDC increases the reporting rates for all 8 PMUs simultaneously from 10 Hz to 50 Hz during the 3rd minute of the 5 minute simulation period. The expedited reporting rate exists for 30 seconds after which the system becomes stable again. Figure 7.9 shows the bandwidth utilisation of the WiMAX uplink subframe under the adaptive and the conventional reporting schemes. The results show that the proposed scheme consumes bandwidth more efficiently than the conventional non-adaptive scheme.

Next, we look at the delay performance of the WiMAX UGS connection under the adaptive reporting scheme as shown in Figure 7.10. From the results, we see that no additional delay is incurred during the changeover time, as the BS allocates resources synchronously to the PMUs ensuring a seamless migration from low to high reporting rates. Nevertheless, when the reporting rate changes from high to low, there is a sharp spike in the delay profile. This is due to the UGS grant dispersion technique that adjusts the PMU data transmission and UGS grant times to ensure minimum UGS delay. Note that during the spike, the delay increases to 40 ms which yields that the data transfer from the eight PMUs are spread over 8 WiMAX frames using the FL/FLI method.

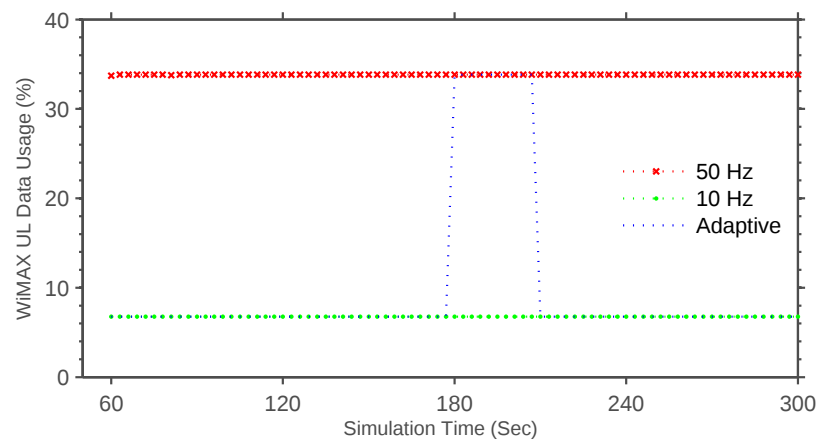


Figure 7.9: WiMAX UL bandwidth usage under the adaptive and the non-adaptive synchrophasor reporting schemes.

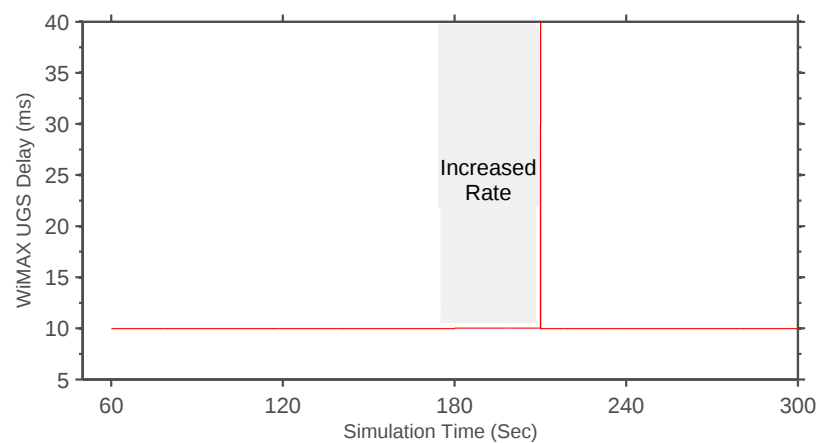


Figure 7.10: Delay profile of the WiMAX UGS service for the adaptive synchrophasor reporting scheme.

7.3 Hybrid Microgrid Protection Scheme

As discussed in Subsection 2.2.6, communications networks play a key role in a microgrid protection scheme as the protective relays need to communicate with each other to dynamically track their fault-current levels due to the time-varying nature of the loads and the DGs. A number of protection schemes have so far been proposed for the microgrids. Of them, some are highly dependent on real time communications; for example, the differential protection scheme described in Subsection 2.2.6. Others are based on fault current estimations and require communications only when there is an event in the system, for example, a system fault or change of topology [110]. Typically, communications-intensive schemes like differential protection are more sensitive as they work with the actual fault-current values. However, they tend to be less reliable, especially under a wireless media where the channels are often affected by variable propagation conditions, such as random noise and multipath-fading.

In this section, we present the concept of a hybrid microgrid protection (HMP) scheme that implements the traditional differential microgrid protection (DMP) scheme, along with an adaptive microgrid protection (AMP) scheme based on a central controller, denoted as the microgrid central protection unit (MCPU). While the differential scheme is more sensitive, it requires extensive communications. On the other hand, the adaptive scheme is less sensitive, but has low communications load. The joint deployment of these two schemes has the potential to increase the accuracy and precision of the whole protection system while reducing the overall communications cost. We also propose a preemptive switching algorithm for the microgrid relays so that they can seamlessly handover from DMP to AMP depending on the bit error rate (BER) of the communication link. Some illustrative results are provided based on a WiMAX network to justify the proposed hybrid protection scheme.

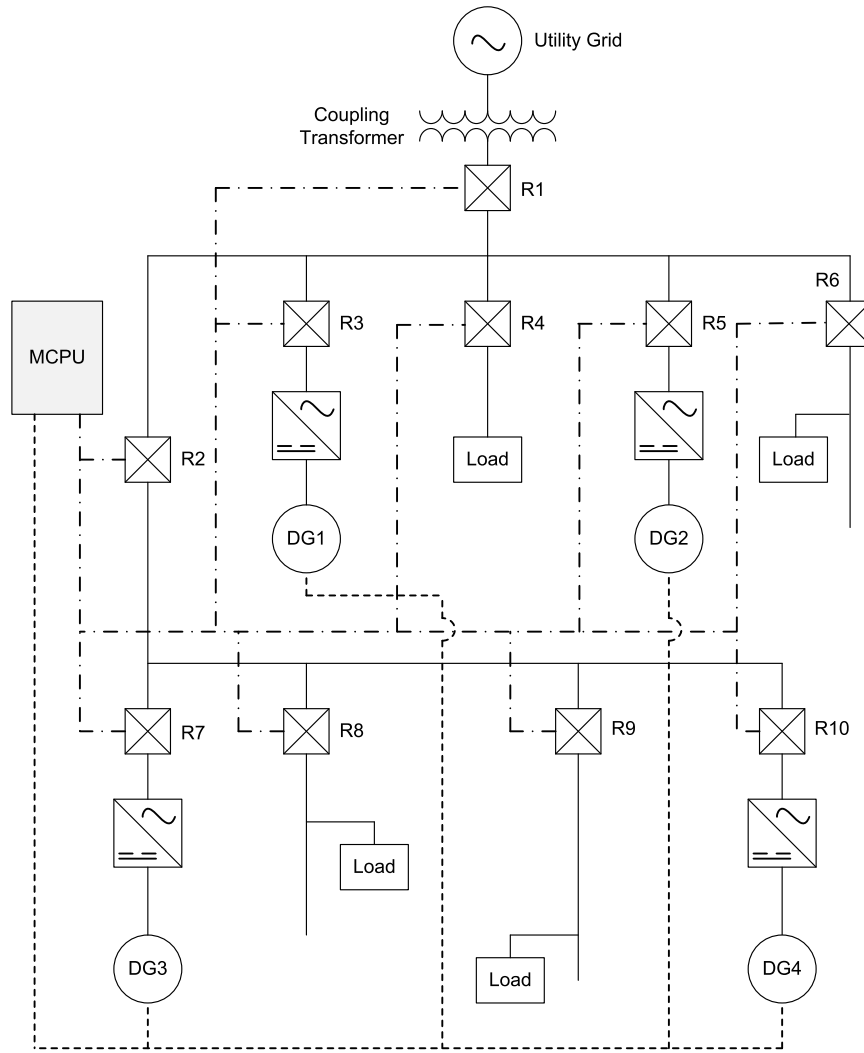


Figure 7.11: An adaptive microgrid protection system [110].

7.3.1 Adaptive Microgrid Protection Scheme

Considering the unconventional fault behaviour of the small-scale rotating machines and the non-linear fault behaviour of the inverter-interfaced DGs, an AMP system has been developed in [111, 110]. An illustrative example of such a system is shown in Figure 7.11. Each relay communicates with the MCPU through a unicast message as soon as it operates either due to a fault, or a change of topology in the system, for example, a change of direction in the current flow due to connection/disconnection of a DG. Based on this information, the MCPU re-calculates the fault current of all the relays and updates them about their new fault-current setting through a broadcast message.

For any relay x in the network, the fault current is calculated using the following formula [110]:

$$I_{fault} = (I_{faultGRID} \times OperatingMode) + \sum_{i=1}^M (k_i \times I_{faultDG_i} \times Status_{DG_i}), \quad (7.14)$$

where M is the total number of DGs in the microgrid, k_i is the impact factor of the i^{th} DG on the fault current of the relay x , $I_{faultDG_i}$ is the maximum fault current contribution of the i^{th} DG, and the ‘operating mode’ is a status bit ensuring that the fault contribution of the utility grid is considered only when the microgrid is connected to the grid; its value is 0 when the microgrid is islanded. Similarly, the fault contribution of a DG which is not in operation will be annulled by its status bit $Status_{DG_i}$. Thus, the above equation considers the fault contribution of grid and the active DGs on that particular relay.

Once the fault currents are estimated, and the operating currents are reported to relays, the relays operate individually without any further information exchange. In other words, communication is only required when there is an event in the microgrid and the fault currents need to be re-estimated. Hence, the scheme is less dependent on the communications network.

7.3.2 Differential Microgrid Protection Scheme

In Section 3.3, we have described a traditional differential protection scheme. However, in a microgrid environment, the level and direction of the fault current may change rapidly due to the dynamic behaviour of the loads and the DGs. Hence, some adaptations are required for the differential protection scheme to work effectively in this environment. In [48], the use of MCPU is proposed which is responsible for comparing the measured currents and detection of a fault, as well as issuing of the subsequent trip signal. Figure 7.12 illustrates the use of the DMP scheme in the same microgrid environment as in Figure 7.11.

As shown in the figure, DMP can be deployed over a transformer or a larger protection zone within the microgrid. However, here the key challenge for a DMP scheme is the increased communications load, as all the relays within the microgrids needs to exchange

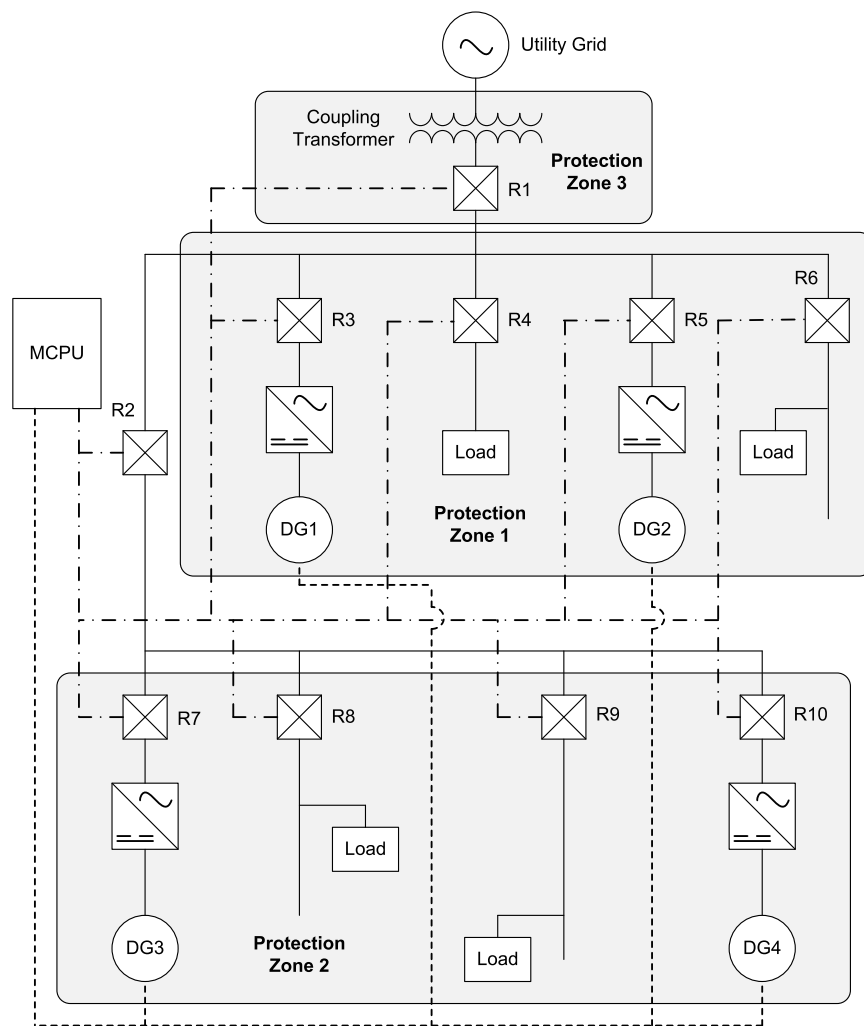


Figure 7.12: A differential microgrid protection system [48].

their current information in real time.

7.3.3 Key Communications Issues

Here, we discuss two key communications issues that highly influence the choice and performance of a microgrid protection scheme.

i) Communications Load: As mentioned earlier, the DMP scheme requires continuous communication among the peer relays. Thus, it has a fixed communications load which is given by

$$C_{DMP} = 2M \frac{L_{DMP}}{\Delta t}, \quad (7.15)$$

where L_{DMP} is the size (in bits) of a differential current measurement packet, Δt is the current measurement interval, and M is the number of relays. The factor 2 in (7.15) represents the bidirectional communications load (i.e. UL and DL) within a cellular WiMAX network.

In the case of the AMP scheme, two types of message exchanges are typically required: i) the multicast ‘status update’ message from the MCPU to the relays; and ii) the unicast ‘status change’ message from the relays to the MCPU (unicast). In addition, the MCPU can send unsolicited ‘status update’ messages to preemptively change the fault current settings of the relays; for example, during a scheduled maintenance. Since the information content of these two messages are almost the same, a fixed message size L_{AMP} can be assumed for the AMP scheme. In contrast, since each ‘status change’ message is associated with a ‘status update’ message, and the number of unsolicited ‘status update’ messages are typically low, we can readily assume an average number of microgrid events per second (denoted as n) for the AMP scheme. Additionally, in cellular wireless networks such as WiMAX, a multicast message requires the same amount overhead as a unicast message.

Thus, the overall communications load for the AMP scheme can be expressed as

$$C_{AMP} = 2MnL_{AMP}. \quad (7.16)$$

Assuming $L_{DMP} \approx L_{AMP}$ and comparing (7.15) and (7.16), we can write

$$\frac{C_{DMP}}{C_{AMP}} \approx \frac{1}{n\Delta t}. \quad (7.17)$$

Typically, the DMP scheme has a measurement interval Δt of 40ms for a distribution microgrid, which yields 25 measurement packets per second. On the other hand, the number of microgrid events n is very low compared to the amount of information exchange in the DMP scheme. Thus, from (7.17), we can clearly see that the communications load is much higher in the DMP scheme than the AMP scheme. This is further evident in Figure 3.12 that plots the data and signalling usage for various number of differential relays in a typical WiMAX network.

ii) Reliability: Reliability of the communications network is another important aspect that needs to be considered while selecting a microgrid protection scheme. In the case of a wireless network, this is even more important as the communications media is often affected by frequency-selective fading and random noise. In such a case, the SNR of the system should lie within a reasonable limit so that the BER does not fall below a certain threshold. Otherwise, packet loss will occur, which will further degrade the performance of the protection scheme during a fault event. Note that the BER of a wireless link depends on the MCS level used. The higher the MCS level, the more the data rate, the higher the SNR requirement. Figure 7.13 shows the BER vs. SNR curve for a typical WiMAX network at different MCS levels.

A typical WiMAX network can scale down its MCS level based on the SNR status of the system using the adaptive modulation and coding (AMC) technique. For a lossy communication channel, the system will first try to combat BER by reducing its MCS level. However, if the SNR does not improve after the lowest MCS level, packet-loss will happen. Although there are several retransmission schemes available in WiMAX, such as ARQ and HARQ, they require time and additional radio resources to operate. Moreover, it is quite challenging to allow retransmission within a small delay budget of a DMP scheme. Hence, under a multi-terminal DMP scheme, a failure in the communication link (due to

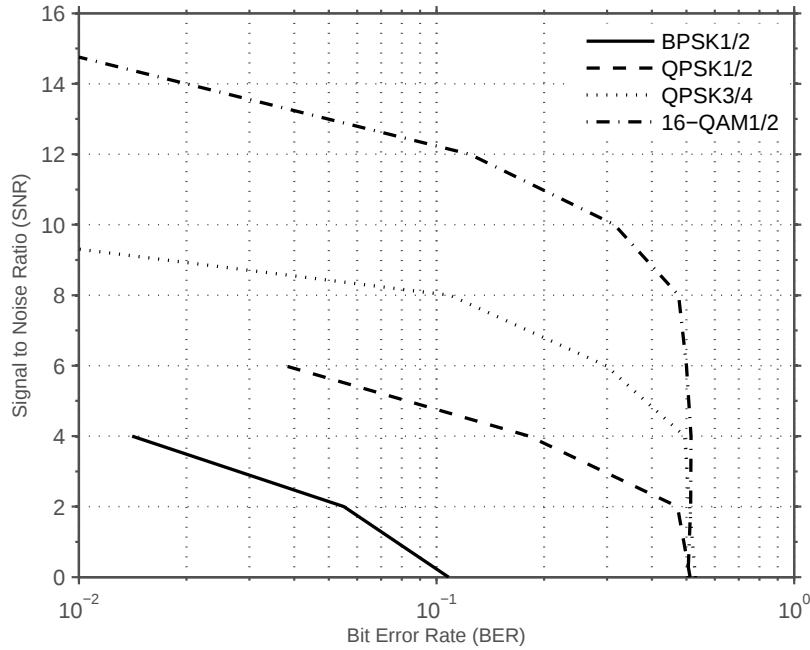


Figure 7.13: BER vs. SNR curve under different MCS levels in a typical WiMAX network.

packet-loss) for one of the relays will cause failure of the whole protection system.

In contrast, the AMP scheme is modular, as a link failure will only affect the individual relay. Moreover, if there is no configuration change prior to the fault occurrence, the link failure will not affect the performance of the protection scheme, as the affected relay will operate based on the as the ‘last-setting-valid’ method. Thus, the AMP scheme is less prone to communications failure than the DMP scheme.

7.3.4 The Hybrid Protection Scheme

From our earlier discussions, we can conclude that AMP acts rapidly, is more reliable, but is not very sensitive as it requires fault current estimation. In contrast, DMP is very sensitive in detecting the faults, but depends highly on continuous communication, and hence, is less reliable in larger systems. Based on these observations, a HMP scheme is proposed that exploits the full benefits of these two schemes and minimises their drawbacks. The key features of the HMP scheme are:

- 1) Differential protection with MCPU support could be used in important nodes such as

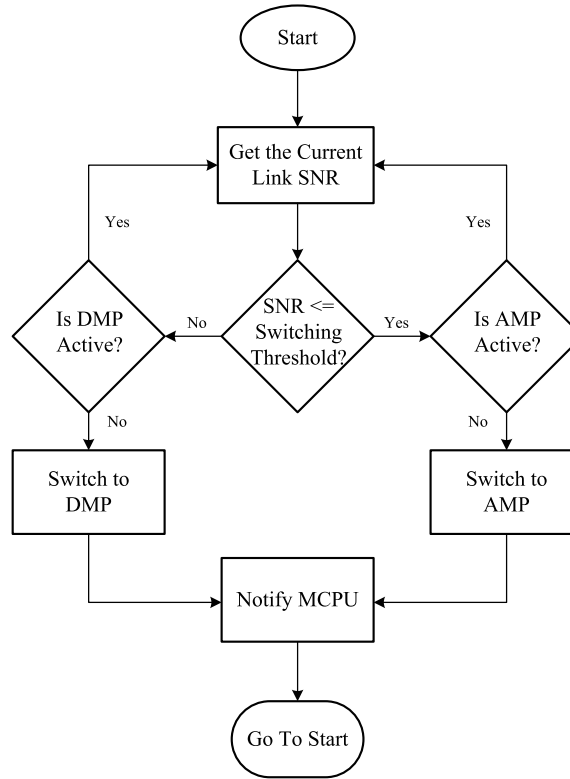


Figure 7.14: Dynamic switching between DMP and AMP schemes under the proposed HMP scheme.

transformers or in high-risk areas, where the expected fault currents are relatively higher and/or sensitive loads exist;

2) Adaptive protection could be utilised in the rest of the system as a roll-back strategy (in case of a communication failure), where differential protection is implemented.

To prevent any physical packet-loss, the HMP might be configured in such a way that it would send a switching request when the communications link operates at the lowest MCS level (i.e. QPSK with $\frac{1}{2}$ rate FEC for WiMAX system). This will ensure that the protection system will have a reasonable time to switch from DMP to AMP without losing any packets. Such a switching algorithm is illustrated in Figure 7.14.

As shown in Figure 7.14, the HMP application in the protective relays continuously observes the SNR level of its communication link with the BS (or MCPU). While operating under the DMP scheme, if the SNR level is equal to or less than a certain switching level, it switches to AMP. In contrast, if the SNR level improves, it switches back to DMP.

Table 7.2: Message Latency (in ms) of the AMP scheme in the WiMAX network.

No. of Microgrids	statusChange (UL)		statusUpdate (DL)	
	Mean	Max.	Mean	max.
2	24.39	33.83	5.90	5.90
4	24.67	34.11	5.90	5.90
6	23.88	34.08	5.90	5.90
8	23.85	34.23	5.90	5.90
10	24.25	34.29	6.10	6.10

7.3.5 Performance Analysis

In Section 3.3, we conducted a comprehensive performance analysis of the traditional DMP scheme. In this subsection, we examine the performance of the AMP scheme over a WiMAX network. Since the communications load of the AMP scheme is very small compared to the DMP scheme, we concentrate on its delay performance only. A simulation model is developed using OPNET Modeller 16.0. The WiMAX simulation parameters are same as listed in Table 3.3. However, to keep the communications latency of the AMP traffic minimum, we employ the DRA mechanism proposed in Chapter 5. Under this mechanism, the protection traffic from the microgrid relays is assigned a dedicated logical random access channel and given the highest priority in the BS scheduler to provide expedited bandwidth grant.

The baseline scenario is comprised of a single microgrid system similar to the one shown in Figure 7.11 that contains a single MCPU with 15 relays (7 bus relays and 8 standalone relays connected with the loads and the DGs). Each relay is assumed to send one protection message with a payload size of 100 bytes within the simulation duration of one hour. We further vary the number of microgrids to examine their effects on the overall message latencies in the WiMAX network. The corresponding results are listed in Table 7.2.

From the results, it can be seen that the WiMAX network, along with the proposed DRA scheme, provides almost fixed latency for both uplink and downlink components of the protection messages. Also, the total latency is below the delay bound of 40 ms for the DMP scheme. Thus, we can safely conclude that if AMP is used alongside DMP, the

WiMAX network will be able to save a considerable amount of radio resources without compromising the delay performance of the protection traffic.

7.4 Chapter Summary

In this chapter, we have presented a number of application-specific optimisations for some of the key Smart Grid M2M applications over a WiMAX network. The EV load management scheme proposed in this chapter has a very low communications load, which depends on the available power of the system rather than the number of connected EVs. This helps the system to maintain an almost fixed communications overhead, no matter how many EVs are connected to the system at any given time. Note that while the scheme is demonstrated for the EV charging applications only, it can be used in other demand management applications, such as HVAC (heating, ventilation and air conditioning) control and management of the DERs. On the other hand, the proposed predictive synchrophasor communications scheme and the HMP scheme ensure an optimum resource utilisation in the WiMAX network by increasing their bandwidth needs only when required. Hence, these optimisations are quite advantageous for a multi-service Smart Grid communications network, where the saved radio resources can be used to serve other applications.

Chapter 8

Conclusions

It is well acknowledged within the industry and the research community that a robust communications network is a fundamental building block of the next generation of electricity networks, i.e. the so called ‘Smart Grid’. However, there is still a lot of debate on which technology is the most appropriate candidate for this and how it should be adapted to the various requirements of Smart Grid applications and devices. Although wireless communication technologies have a clear leverage in this arena but there are a number of uncertainties, for example, network performance in a M2M communications environment, suitability for a vast number of emerging applications, and scalability to current and future communications needs. As a prospective Smart Grid communications technology, the IEEE 802.16-based WiMAX networks are also subjected to these uncertainties and challenges.

The goal of this research was to identify the above challenges and suggest appropriate measures, so that utility operators can receive maximum benefits from an IEEE 802.16/WiMAX-based Smart Grid communications network. In particular, this thesis has validated the notion that the random access plane is the key bottleneck for supporting M2M communications over a WiMAX network. Moreover, it has proposed a number of solutions to overcome this problem, including several other enhancements, such as the HetNet architecture in Chapter 6 and the application-specific optimisations in Chapter 7.

This chapter is devoted to a summary of the main contributions reported in this thesis and the associated key results, followed by a discussion on prospective research directions for

future work.

8.1 Summary of Contributions

One of the significant contributions of this thesis is an in-depth and technology agnostic review of the application characteristics and traffic requirements of the major Smart Grid applications. The review has two key components. First, it develops an understanding of the Smart Grid communications architecture and identifies its key components based on the well-regarded IEEE 2030-2011 standard for Smart Grid interoperability. Second, it investigates the characteristics and traffic requirements of the key Smart Grid applications, and highlights some key challenges in the context of a Smart Grid communications network. The key outcome from these discussions is that a potential Smart Grid communications technology needs to support a wide range of traffic sources with significantly varying QoS requirements. Therefore, its key challenge is to efficiently allocate radio resources among these various classes of traffic so that their end-to-end QoS requirements are met.

Based on this review, Chapter 2 developed a number of traffic models and conducted simulation studies on the performance of the three the key traffic sources of the Smart Grid, namely, smart meters/AMI, synchrophasors, and protective relays, over a conventional WiMAX network. The AMI study shows that a typical WiMAX BS offers a moderate coverage range for a suburban AMI/Smart Grid network with predominantly indoor SSs. However, the coverage can be greatly improved if the outdoor SSs are used due to their improved antenna height/gain and immunity from the building penetration losses. This was further addressed in Chapter 6, where a WiMAX-WLAN HetNet architecture was proposed to improve network coverage and utilisation. The study further reveals that the existing random access mechanism of WiMAX can significantly degrade the performance of the bursty AMI applications, especially when the access load is high. The performance is even worse for event-based traffic, such as outage alarms from the smart meters and the protection signals from the relays. On the other hand, the study on synchrophasors in Chapter 3 shows that the synchrophasor traffic can be well-supported by the conventional

UGS service of WiMAX; moreover, using advanced WiMAX features such as persistent scheduling and ROHC, their radio resource usage can be significantly reduced. To further optimise their radio resource usage, a number of application-specific optimisations were conducted in Chapter 7.

Inspired by the findings in Chapter 3, a comprehensive analytical model was developed in Chapter 4 to evaluate the performance of the CDMA-based random access procedure in WiMAX. The analysis reveals that the random access channels perform optimally when the number of devices contending simultaneously does not exceed a particular threshold. Otherwise, the channel becomes unstable in terms of access success rate and access delay. Moreover, through Monte-Carlo simulation, it shows that the BS is able to detect only a limited number of codes in the presence of MAI, random noise and frequency-selective fading in the multipath wireless channel. Thus, it validates the notion that the random access plane is the key bottlenecks for supporting M2M communications over a WiMAX network.

Chapter 5 has been devoted to addressing the above problem with a cross-layer approach. First, it has proposed an enhanced random access scheme where the fixed/low-mobility M2M devices pre-equalise their BR codes using the estimated frequency response of their slowly-varying channels. Consequently, the BS can detect a large number of codes as their mutual orthogonality remains preserved. Mathematical analysis is conducted to determine the theoretical performance limit of the code detector. The analysis reveals that the default pseudo-random code matrix specified in the IEEE 802.16 standard is not quite effective for detecting a large number of codes under the proposed scheme. As a remedy, it is argued that a Hadamard code matrix can significantly increase its code detection performance. Moreover, a DRA strategy is proposed to provide QoS-aware access service to various M2M devices. The theoretical performance of the proposed scheme is validated by simulation results under both of the two code matrices. The proposed scheme is fully compliant with the existing WiMAX specifications, except it requires a dedicated BR channel when both M2M and conventional applications need to be supported. Such a requirement is reasonable considering the volume of the M2M devices per BS and has already been

provisioned in the IEEE 802.16p amendment. Moreover, it has proposed an adaptive RRM framework based on the above DRA concept, which adaptively allocates the ranging resources based on load and priority of the application. The performance of the proposed scheme has been demonstrated using both generic traffic profiles and the pilot protection application studied in Chapter 3.

To improve the capacity and coverage issues identified in Chapter 3, Chapter 6 argued that a heterogeneous deployment of a WiMAX network with one or more short/medium range wireless technologies, such as IEEE 802.11 based WLAN or IEEE 802.15.4-based ZigBee, can efficiently address these challenges, thereby offering the ultimate wireless solution for a Smart Grid communications network. Such a multi-tier network is fully compliant with the communications architecture of the Smart Grid reviewed in Chapter 2. Moreover, it provides physical separation between the NAN/FAN and the WAN of an end-to-end Smart Grid network, which improves security and encourages the development of new distributed control and energy management applications in the NAN/FAN domain. The proof of concept was validated by a WiMAX-WLAN HetNet architecture, where the dual radio WLAN APs are embedded within the coverage umbrella of a WiMAX network.

Finally, Chapter 7 has carried out a number of application-specific optimisations to support the EV load management and synchrophasor-based protection and control schemes more efficiently over a WiMAX-based Smart Grid communications network. The EV load management scheme proposed in this chapter has a very low communications load, which depends on the available power of the system rather than the number of connected EVs. This helps the system to maintain an almost fixed communications overhead, no matter how many EVs are connected to the system at any given time. This highly efficient scheme can also be used in other demand management applications, such as HVAC control and management of the DERs. In addition, the predictive synchrophasor communications scheme and the HMP scheme presented in this chapter ensure an optimum resource utilisation in the WiMAX network by increasing their bandwidth needs only when required. Hence, these optimisations are quite advantageous for a multi-service Smart Grid communications network where the saved radio resources can be used to serve other applications.

8.2 Future Research Directions

Given that this research is one of the first investigations devoted to the applicability of a 4G wireless network within a Smart Grid paradigm, a number of potential areas could be extended for future studies, discussed as follows.

1. The random access performance analysis and enhancement carried out in Chapter 4 and Chapter 5 can be applied to other prospective Smart Grid technologies, such as the 3GPP-based LTE, which follows a similar random access procedure. The DRA concept developed in Chapter 5 can be extended and adapted to other random access-based wireless solutions such as IEEE 802.11-based WLAN networks. Moreover, the pre-equalisation based code transmission technique can also be used to optimise the initial ranging performance of the fixed/low-mobility M2M devices, which was outside the scope of this research.
2. The applicability of HetNet architecture in Chapter 6 needs further investigation and enhancements, especially where interference from a neighbouring home WiFi/ZigBee device can affect the performance of the overall network. Note that while WiMAX has an end-to-end QoS framework, the WLAN provide limited QoS support at the MAC layer based on the IEEE 802.11e standard. Hence, an important challenge for the HetNet architecture is to maintain end-to-end QoS for each traffic class. An adaptive RRM framework can help such a network architecture to optimise its performance in terms of QoS fulfillment and resource utilisation. Further, more studies are required to see how other secondary technologies of the HetNet architecture can perform in such an environment, for instance, the IEEE 802.15.4-based ZigBee and IEEE 802.11p-based Mesh WLAN.
3. Many of the application and traffic models developed in this thesis are technology agnostic. For example, the pilot protection scheme, the LCDP scheme, and the synchrophasor data transfer model. Further analysis and enhancement of the potential Smart Grid communication solutions can be carried out using these models. For example, how a pilot protection application performs in a ZigBee/WLAN based mesh

network and how its performance can be enhanced/optimised.

4. The EV load management scheme described in Chapter 7 needs further enhancement and optimisation to be used in a real life scenario, where more sophisticated energy scheduling algorithms are required. The concept of ‘energy bursts’ can be applied to other demand management applications such as HVAC, DERs and storage devices to optimise their energy scheduling and communications performance.
5. The adaptive synchrophasor reporting scheme in Chapter 7 can be developed further to incorporate both transient and steady-state voltage stabilisation issues of the power network. The concept can also be extended to other non-PMU applications, such as high resolution SCADA monitoring and remote surveillance of the Smart Grid assets.
6. The hybrid protection scheme in Chapter 7 needs further work to enable coherent decision making to switch from DMP and AMP and vice-versa, since it requires the MCPU to consider the performance of the underlying communications network as a decision metric.

Appendix A

Discrete Event Simulation

A.1 The OPNET Modeller

All the simulation models used in thesis are developed using the well-known simulation package OPNET Modeller 16.0. It is widely used for network research and development, modelling and simulation by defense organisations, research communities, and leading network equipment manufacturers.

The OPNET Modeller platform provides a comprehensive development environment supporting the modelling of communication networks and distributed systems using finite state machines (FSM). Both the behaviour and performance of modelled systems can be analysed by performing discrete event simulation (DES), where the system is modelled as a sequence of states. Each state is associated with a particular event that occurs at a particular instant in time. The model then evolves through these events over time, based on the behaviour of the its components and their interactions.

The OPNET Modeller adopts a hierarchical modelling concept that consists of three domains: viz. network domain, node domain, and process domain. The network domain at the top level defines the network topology, consisting of communication entities known as nodes and links to allow communications between nodes. The node domain at the second level comprises of a set of modules that describes the node's functionality and connections,

which allow information to flow between modules. The modules in the nodes are implemented using process models of the process domain at the lowest level. The process models provide behavioural modelling for programmable modules using Proto-C based on a combination of state transition diagrams, a library of high level commands known as kernel procedures, and general facilities of the C/C++ programming language.

A.2 Simulation Models

The OPNET Modeller provides a specialised WiMAX model that accurately emulates the operation of a WiMAX network by implementing the OFDMA physical layer and MAC layer features such as modulation and coding, QoS scheduling and BR. To conduct the simulation studies in this thesis, a number of custom models were developed by modifying the built-in WiMAX BS and SS node models and augmenting them with a customised application module. Example of such models include the AMI smart meter model (see Figure A.1), the synchrophasor models, the differential and pilot protection models, and the HetNet architecture models (see Figure A.2). Moreover, a combined energy and communications model was developed for the EV load management technique. Besides, the existing WiMAX BS and SS process models were modified to extend to include a number of additional WiMAX features, such as synchronised UGS grant allocation, the proposed adaptive synchrophasor reporting, and the proposed adaptive RRM scheme based on DRA mechanism.

For example, let us consider the EVCS and EVCC process models of the WiMAX-based EV charging application model in Section 7.1. The models are shown in Figure A.3. The FSM in the EVCS application module has five different states: *Init*, *Idle*, *Request* (initial charging request), *Response* (continue message) and *Update* (SoC update) as shown in Figure A.3(a). While the request state sends the initial charging request message, the response state is responsible for processing the energy allocations from the EVCC and the update state sends the SoC update after consuming each energy packet. The *Init* state initialises the FSM and the *Idle* state represents the waiting state of the process. Similarly, the FSM

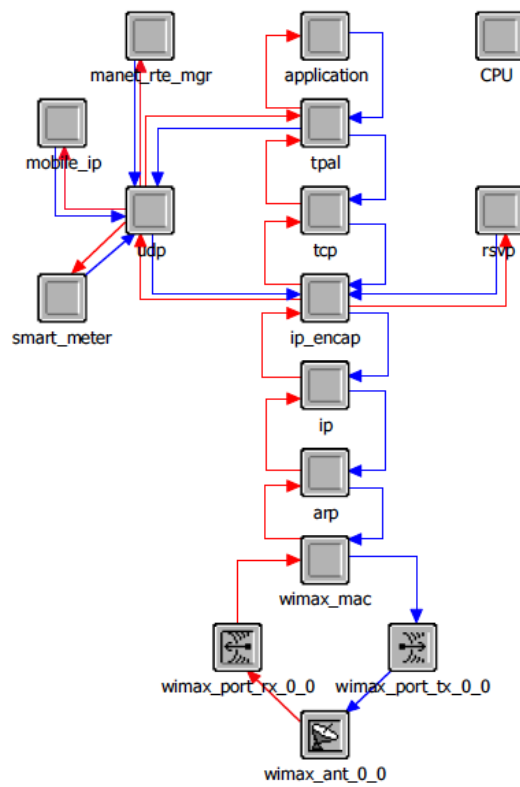


Figure A.1: The WiMAX-enabled smart meter node model.

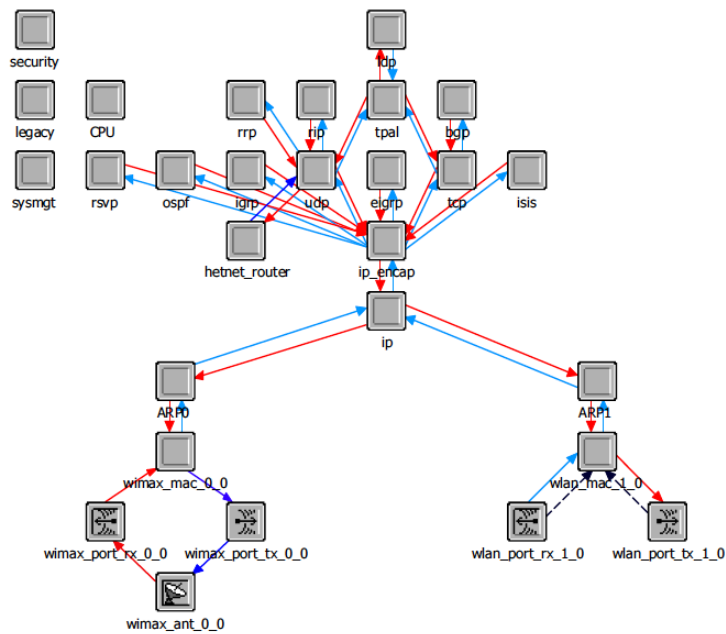


Figure A.2: The WWR node model.

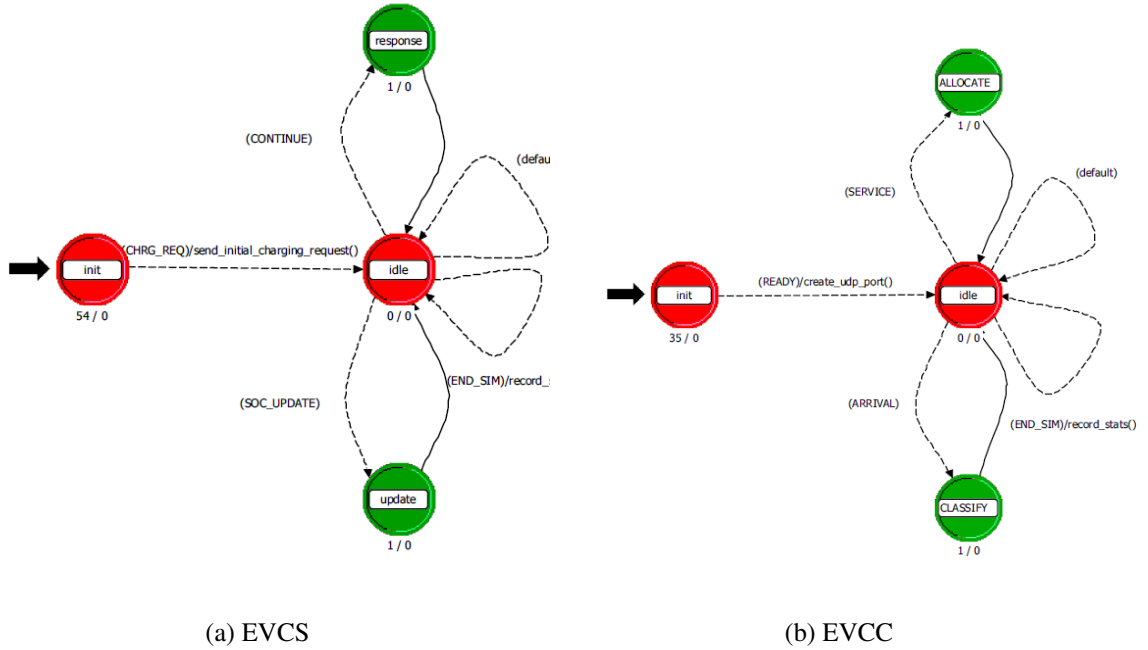


Figure A.3: Process models for the WiMAX-based EV load management application (in Section 7.1).

in EVCC application module has four states – *Init*, *Idle*, *Classify* and *Allocate* as shown in Figure A.3(b). While the *Classify* state performs request classification, the *Allocate* state performs the energy scheduling.

For backbone communication, IPv4 was used for the simulation models. Each WiMAX BS is assumed to be connected by an access service network gateway (ASN-GW) router, which also acts as an IP gateway routing the data packets from the WiMAX BS to the backbone IP network. In practice, the ASN-GW may be either co-located with the BS or it may be a stand-alone server in the backbone network. The MDMS for the smart meters, the PDC for the synchrophasors, and the MCPU for the microgrids are modelled as a remote server located in the utility's core network. For example, Figure A.4 depicts the WiMAX-based synchrophasor communications network where the PMUs are sending their measurements to their destination PDC through the WiMAX BS, ASN-GW, and the IP backbone network.

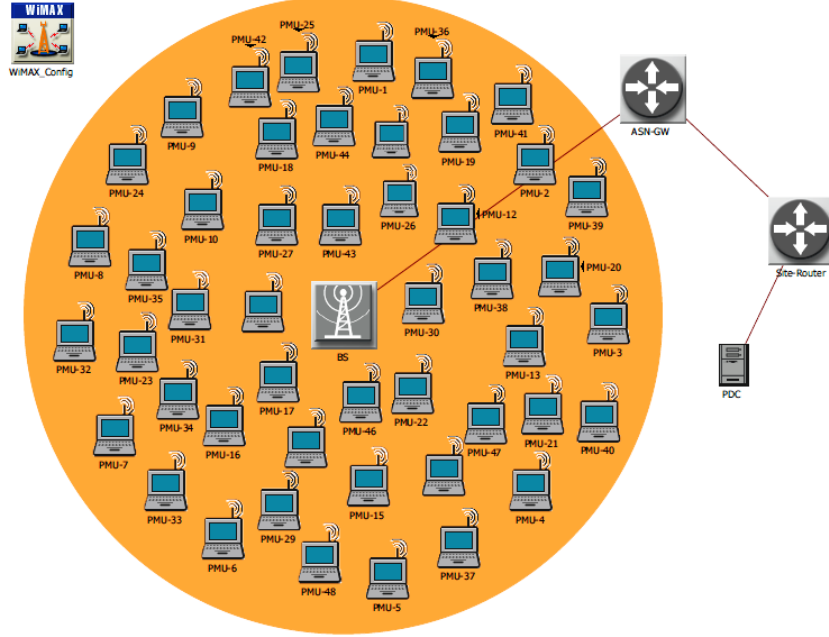


Figure A.4: WiMAX-based synchrophasor network model in OPNET.

A.3 Statistical Validity of the Results

A simulation result is considered statistically significant if it is unlikely to have occurred by pure chance. In general, system models that include stochastic behaviour have results that are dependent on the initial seeding of the random number generator. Since a particular random seed selection can potentially result in an anomalous or non-representative behaviour, it is important for each model configuration to be run with several random number seeds to determine a typical, representative behaviour. The basic principle of statistical validity applied here is that if a typical behaviour exists and many independent trials are performed, it is likely that a significant majority of these trials will fall within a close range of the representative behaviour. Therefore, all results presented throughout this thesis are computed from multiple simulation runs in which each seed value of random number generator is different.

Appendix B

Multicarrier Transmission with OFDM

In this appendix, we describe the OFDM modulation and demodulation process in detail. The basic idea of OFDM in particular, and multicarrier modulation in general, is to transmit a block of digital modulation symbols by modulating them on a large set of narrowband subcarriers (in frequency). On each subcarrier, only a fraction of the total bandwidth is available, which prolongs the symbol duration of an OFDM signal. The prolonged duration of an OFDM symbol provides enhanced immunity against Inter-symbol interference (ISI) due to multipath signal distortion. Moreover, introduction of a so-called cyclic-prefix (CP) can completely eliminated ISI from an OFDM symbol (explained in Section B.3.2). Another advantage of OFDM is that the process of modulating the source symbols onto the different subcarriers is performed by fast and computationally efficient FFT operations. Figure B.1 shows a block diagram of OFDM system outlining the OFDM modulation and demodulation process.

As shown in the Figure, the OFDM modulation is carried out digitally, after which the discrete-time OFDM signal is converted to an analog signal using a pulse-shaping filter. The analog baseband signal is then up-converted to passband frequency before continuous-time transmission across the wireless channel. During transmission, the signal is corrupted by additive noise and multipath channel distortions. At the receiver, the distorted signal is first down-converted to the baseband level and then passed through a matched-filter to improve the received SNR. The filtered signal is then sampled and fed to the OFDM demod-

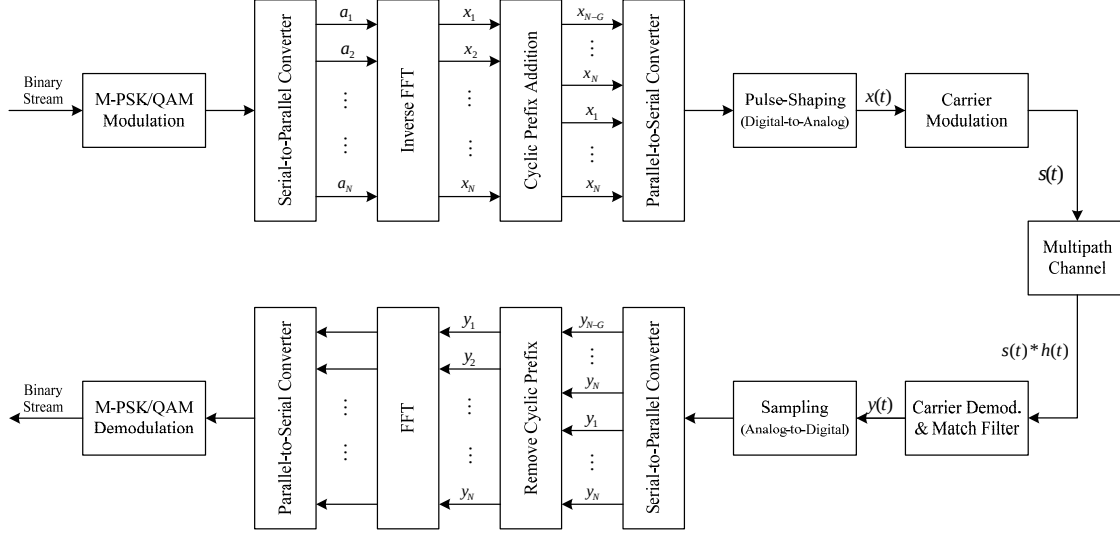


Figure B.1: Block diagram of an OFDM system.

ulator to recover the source symbols. Detailed description of these processes are provided in the subsequent sections.

B.1 OFDM Modulation

B.1.1 Multiplexing with OFDM

Consider an OFDM system that transmits $a_k, k = 1, 2, \dots, N$ source modulation symbols, each having a symbol duration T_d . The source symbols are phase/amplitude modulated using digital modulation techniques such as M-PSK or QAM. Based on the modulation scheme used, a_k can be a real or a complex number. The OFDM system uses N parallel orthogonal subcarriers to transmit these N consecutive symbols over an OFDM symbol duration $T_s = N \cdot T_d$. However, instead of the sending them as a pulse-train, the source symbols are first converted into N parallel symbols, i.e. $\mathbf{a} = [a_1, a_2, \dots, a_N]^T$ and then fed into a N point IFFT (Inverse Fast Fourier Transform) processor and converted to a new set of symbols $\mathbf{x} = [x_1, x_2, \dots, x_N]^T$.¹

Although the new symbols are transmitted as a pulse train of T_d duration, each of them

¹**Notation:** Vectors are recognised as boldface lowercase letters, e.g. $\mathbf{a}$, $\mathbf{x}$ and are always considered as columns by default. Matrices are recognised as boldface capitalised letters, e.g. $\mathbf{W}$, $\mathbf{H}$.

carriers the frequency contents of all the original source symbols. Therefore, the receiver can recover $\mathbf{a}$ given that sufficient information about $\mathbf{x}$ is available. In order to maintain orthogonality over the N subcarriers, and thereby to reduce Inter-Carrier Interference (ICI), the OFDM system packs them with a frequency spacing $\Delta f = \frac{1}{T_s}$ i.e. the k^{th} subcarrier is located at the frequency

$$f_k = k \cdot \Delta f = \frac{k}{T_s} = \frac{k}{NT_d}, \text{ for } \forall k = 1, 2, \dots, N. \quad (\text{B.1})$$

Thus, the elements of the N -point IFFT of $\mathbf{a}$ is defined as

$$\begin{aligned} x_k &= \sum_{n=1}^N a_n e^{j2\pi f_k \cdot nT_d} \\ &= \sum_{n=1}^N a_n e^{j2\pi \cdot \frac{k}{T_s} \cdot n \cdot \frac{T_s}{N}} \\ &= \sum_{n=1}^N a_n e^{j2\pi k \frac{n}{N}}, \text{ for } \forall k = 1, 2, \dots, N \end{aligned} \quad (\text{B.2})$$

The above IFFT operation can be more conveniently represented in matrix form as

$$\mathbf{x} = \mathbf{W}^H \cdot \mathbf{a}, \quad (\text{B.3})$$

where $\mathbf{W}$ is the $N \times N$ point FFT matrix with elements $[W]_{k,n} = e^{-j2\pi k \frac{n}{N}}$, for $\forall k, n \in \{1, 2, \dots, N\}$, such that

$$\mathbf{W} = \begin{bmatrix} 1 & 1 & 1 & \dots & 1 \\ 1 & W & W^2 & \dots & W^{N-1} \\ 1 & W^2 & W^4 & \dots & W^{2(N-1)} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & W^{N-1} & W^{2(N-1)} & \dots & W^{(N-1)^2} \end{bmatrix}.$$

Note that $\mathbf{W}^H$ represents the IFFT matrix, which is the hermitian transpose of $\mathbf{W}$ and ob-

tained by elements $[W]_{k,n}^*$.

From (B.2), we see that each individual source symbol a_k is distributed into N narrowband frequency channels of bandwidth $\Delta f = 1/NT_d$, instead of the total available bandwidth $1/T_d$. As a result, the duration of a time-domain OFDM symbol is much higher compared to the original source symbol duration (since, $T_s = N \cdot T_d$). In turn, this significantly reduces the ISI among the source symbols as the duration of the multipath delay spread τ_{max} is much smaller than the OFDM symbol duration, i.e. $T_s \gg T_d > \tau_{max}$.

However, the symbols transmitted in the early part of the OFDM symbol is affected by the delayed copies of the previous sample, causing ISI between two successive OFDM symbols. Hence, to completely eliminate ISI, the last part of an OFDM symbol is augmented as a CP with the beginning of the symbol. The duration of the CP is chosen in such a way that it is at least equal or more than the channel delay spread, i.e. $T_{cp} \geq \tau_{max}$ ². The discrete length of CP is therefore given by $G \geq \frac{\tau_{max}N}{T_s}$. Thus, the discrete-time OFDM signal in (B.3) becomes

$$\hat{\mathbf{x}} = [\mathbf{x}_{cp}, \mathbf{x}]^T = [x_{N-G}, \dots, x_N, x_1, \dots, x_N]^T, \quad (\text{B.4})$$

where, $\mathbf{x}_{cp} = [x_{N-G}, x_{N-G-1}, \dots, x_N]$.

Note that, the use of CP extends the length of the OFDM signal from T_s to $T = T_s + T_{cp}$.

B.1.2 Pulse Shaping and Transmission

Once the discrete-time OFDM signal $\mathbf{s}$ is determined, the corresponding analog baseband signal is constructed from pulse-shaping of the symbols, i.e. each discrete OFDM sample is multiplied by a suitable pulse $p(t)$ of duration $[0, T_d]$. The continuous pulse-shaped signal is given by

$$x(t) = \sum_{n=-G}^N x_n p(t - nT_d), \quad 0 \leq t \leq T \quad (\text{B.5})$$

²A typical WiMAX system with an FFT size of 512 has $\Delta f = 10.94$ KHz, $T_s = 91.4 \mu s$, $T_{cp} = 1/8$, and $T = 11.4 \mu s \approx 64$ samples.

In terms of source modulation symbols, $x(t)$ can be written as

$$x(t) = \sum_{k=1}^N a_k \sum_{n=-G}^N p(t - nT_d) e^{j2\pi kn/N} \quad (\text{B.6})$$

Equation (B.6) shows that each of the source modulation symbols is superimposed on N number of subcarriers over the interval $[0, T]$, thereby, transforming it into a multicarrier baseband signal. Note that to prevent out-of-band power emissions from the subcarriers along the edge of the OFDM symbol, some of the subcarriers are used as guard subcarriers, which do not contain any modulated symbols; thus, they allow the OFDM signal to decay naturally.

Let, the transfer function of $p(t)$ is $p(f)$. To minimise ISI between two successive pulses, the side-lobes of $p(f)$ must decay quickly without the expense of a large amount of excess bandwidth. To meet this criteria, in most practical implementations, pulse-shaping is carried out by a Raised-cosine filter (RCF) [112]. In time-domain,

$$p(t) = \text{sinc}\left(\frac{t}{T_d}\right) \frac{\cos(\pi\beta t/T)}{1 - 4\beta^2 t^2/T^2} \quad (\text{B.7})$$

The transfer function of such a filter is given by

$$p(f) = \begin{cases} T_d, & 0 \leq f \leq \frac{1-\beta}{2T_d} \\ T_d \cos^2\left[\frac{\pi T_d}{2\beta}(|f| - \frac{1-\beta}{2T_d})\right], & \frac{1-\beta}{2T_d} \leq |f| \leq \frac{1+\beta}{2T_d} \\ 0, & \text{Otherwise} \end{cases} \quad (\text{B.8})$$

where β denotes the roll-off factor for the RCF i.e. the excess bandwidth occupied beyond the Nyquist bandwidth of $1/2T_d$.

Before final transmission, the baseband signal $x(t)$ is superimposed on a carrier frequency f_c to create a band-pass signal $s(t)$, such that

$$s(t) = \Re \left\{ x(t) e^{j2\pi f_c t} \right\}. \quad (\text{B.9})$$

B.2 Multipath Propagation Channel

For the sake of simplicity, we ignore the channel effect on the passband signal $s(t)$, rather we focus on the baseband signal $x(t)$. Consider the impulse response of an L path channel, represented by a discrete number of Dirac delta functions

$$h(\tau, t) = \sum_{l=1}^L \alpha_l(t) \delta(\tau - \tau_l(t)), \quad (\text{B.10})$$

where α_l represents channel gain and τ_l represents channel delay for an arbitrary path $l \in 1, 2, \dots, L$ that varies with time t .

Assuming the channel is time-invariant over the OFDM symbol duration T , we can write

$$h(t) = \sum_{l=1}^L \alpha_l \delta(t - \tau_l), \quad 0 \leq t \leq T. \quad (\text{B.11})$$

Well-known multipath channel models exist that specifies the number of paths (also known as channel taps) and quantifies the delay and relative power for each of these paths. For example, the SUI-3 and 4 channel model specified in Table 4.1. Typically, a multipath channel is modelled by a FIR filter with L -tap coefficients, representing each distinct path of the channel. Each filter-tap is modelled to follow a particular distribution with the mean value described in their respective channel models. The Rayleigh fading model is widely used to model the effect of multipath propagation environment, where no dominant Line-of-Sight (LOS) path exists. If there is a dominant LOS, the Rician fading model is typically used.

According to the Rayleigh fading model, the channel impulse response can be reasonably approximated by a Gaussian random process with zero mean and phase evenly distributed between 0 and 2π radians. The envelope (i.e. magnitude) of the channel response varies randomly according to a Rayleigh distribution parameter R , which is defined as the radial component of the sum of two uncorrelated Gaussian random variables. For example, $R = \sqrt{X^2 + Y^2}$, where $X \sim \mathcal{N}(0, \sigma^2)$ and $Y \sim \mathcal{N}(0, \sigma^2)$ are independent Gaussian random variables with zero mean and a variance σ^2 .

Moreover, the variation in the value of a channel tap from one sample to another depends on the Doppler frequency, which in turn depends on the speed of the transmitter; the higher the velocity of the transmitter, the higher the Doppler frequency, and the greater the variations in the channel. The Doppler frequency is defined as

$$f_d = \frac{v \cos(\theta)}{\lambda}, \quad (\text{B.12})$$

where v is the velocity of the transmitter, θ is the angle between the direction of arrival of the signal and the direction of motion, and λ is the wavelength of the signal.

B.3 OFDM Demodulation

B.3.1 Match Filtering and Sampling

After down-conversion, the received signal $y(t)$ is given by the convolution of the transmitted baseband signal $x(t)$ with the channel impulse response $h(t)$ plus some additive Gaussian noise $z(t)$, i.e.

$$\begin{aligned} y(t) &= x(t) * h(t) + z(t) \\ &= \int_0^t x(\tau) h(t - \tau) d\tau + z(t) \quad 0 \leq t \leq T. \end{aligned} \quad (\text{B.13})$$

Before sampling $y(t)$, it is passed through a matched filter to improve the SNR [112]. The matched filter has an impulse response $g(\tau)$, which is a mirror image of the transmitted pulse $p(t)$ with a time-lag τ , such that $g(\tau) = p(t - \tau)$, $0 \leq t \leq T_d$. Finally, $r(t)$ is sampled at an integer multiple of T_d to obtain the discrete-time signal $\mathbf{r}$.

B.3.2 DE-Multiplexing and Recovery

The discrete sequence after sampling can be expressed as

$$y(k) = \sum_{i=1}^L h_i x(k-i) + z(k), \quad k = N-G, \dots, N, 1, \dots, N, \quad (\text{B.14})$$

where $z(k)$ is the additive white Gaussian noise (AWGN) component with zero mean and a variance. More elaborately,

$$\begin{bmatrix} y_{N-G} \\ \vdots \\ y_N \\ y_1 \\ \vdots \\ y_N \end{bmatrix} = \begin{bmatrix} h_1 & 0 & 0 & 0 & 0 & 0 & 0 \\ h_2 & h_1 & 0 & \vdots & 0 & 0 & 0 \\ \vdots & h_2 & \ddots & 0 & \vdots & 0 & 0 \\ h_L & \vdots & \ddots & h_1 & 0 & \vdots & 0 \\ 0 & h_L & \vdots & h_2 & h_1 & 0 & \vdots \\ \vdots & \vdots & \ddots & \vdots & \ddots & \ddots & 0 \\ 0 & 0 & \dots & h_L & \dots & h_2 & h_1 \end{bmatrix} \begin{bmatrix} x_{N-G} \\ \vdots \\ x_N \\ x_1 \\ \vdots \\ x_N \end{bmatrix} + \begin{bmatrix} z_{N-G} \\ \vdots \\ z_N \\ z_1 \\ \vdots \\ z_N \end{bmatrix}. \quad (\text{B.15})$$

The linear systems of equations in (B.15) can be compactly expressed as

$$\hat{\mathbf{y}} = \hat{\mathbf{H}} \cdot \hat{\mathbf{x}} + \hat{\mathbf{z}} \quad (\text{B.16})$$

Note that the time-domain channel matrix $\hat{\mathbf{H}}$ in (B.16) is a Toeplitz matrix, where the diagonal components of the lower-triangle are same. Since ISI is only present in the first L samples of the signal, they are removed before OFDM demodulation. Thus, after removing the CP, the sequence becomes

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ \vdots \\ Y_N \end{bmatrix} = \begin{bmatrix} h_1 & 0 & 0 & 0 & h_L & \cdots & h_1 \\ h_2 & h_1 & 0 & \vdots & 0 & 0 & h_2 \\ \vdots & h_2 & \ddots & 0 & \vdots & 0 & 0 \\ h_L & \vdots & \ddots & h_1 & 0 & \vdots & h_L \\ 0 & h_L & \vdots & h_2 & h_1 & 0 & \vdots \\ \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & 0 \\ 0 & 0 & \cdots & h_L & h_{L-1} & \cdots & h_1 \end{bmatrix} \times \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ \vdots \\ x_N \end{bmatrix} \begin{bmatrix} z_1 \\ z_2 \\ \vdots \\ \vdots \\ z_N \end{bmatrix}. \quad (\text{B.17})$$

In compact form:

$$\mathbf{y} = \tilde{\mathbf{H}} \cdot \mathbf{x} + \mathbf{z}. \quad (\text{B.18})$$

Note that $\tilde{\mathbf{H}}$ is a circulant matrix, where the row vector is rotated one element to the right relative to the preceding row. Thus, the use of CP transforms the linear convolution operation in (B.15) to circular convolution.

Now, the source symbols can be extracted by performing FFT operation on $\mathbf{y}$. The resulting signal is given by

$$\begin{aligned} \mathbf{b} &= \text{FFT} \{ \mathbf{y} \} \\ &= \mathbf{W} \cdot \tilde{\mathbf{H}} \cdot \mathbf{x} + \mathbf{W} \cdot \mathbf{z} \\ &= \mathbf{W} \cdot \tilde{\mathbf{H}} \cdot \mathbf{W}^H \cdot \mathbf{a} + \bar{\mathbf{z}} \\ &= \mathbf{H} \cdot \mathbf{a} + \bar{\mathbf{z}}, \end{aligned} \quad (\text{B.19})$$

where $\mathbf{H}$ is a diagonal matrix derived from the circulant matrix $\tilde{\mathbf{H}}$, i.e. $\mathbf{H} = \mathbf{W} \cdot \tilde{\mathbf{H}} \cdot \mathbf{W}^H$ after the FFT operation. The diagonal elements of $\mathbf{H}$ contains the frequency response coefficients of the channel, i.e.

$$\mathbf{H} = \begin{bmatrix} H_1 & 0 & \cdots & 0 \\ 0 & H_2 & 0 & \cdots \\ \vdots & \ddots & \ddots & \ddots \\ 0 & 0 & \cdots & H_N \end{bmatrix}, \quad (\text{B.20})$$

where

$$[H]_k = \sum_{n=1}^N h_n e^{-j2\pi kn/N}, \text{ for } \forall k = 1, 2, \dots, N. \quad (\text{B.21})$$

From (B.20) and (B.21), we see that the overall OFDM transmission channel can be visualised as N parallel flat-fading Gaussian channels, which have the following relationship:

$$b_k = a_k \cdot H_k + \bar{z}_k, \text{ for } \forall k = 1, 2, \dots, N, \quad (\text{B.22})$$

where $b_k; k = 1, 2, \dots, N$ is the received modulation symbols.

Note that $h_n = 0$ if $n > L$. Therefore, $[H]_k$ represents the complex-valued coefficients of the L -tap FIR filter, i.e.

$$\begin{bmatrix} H_1 \\ H_2 \\ \vdots \\ H_N \end{bmatrix} = \begin{bmatrix} 1 & 1 & 1 & \cdots & 1 \\ 1 & W & W^2 & \cdots & W^{N-1} \\ 1 & W^2 & W^4 & \cdots & W^{2(N-1)} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & W^{N-1} & W^{2(N-1)} & \cdots & W^{(N-1)^2} \end{bmatrix} \cdot \times \begin{bmatrix} h_1 \\ h_2 \\ \vdots \\ h_L \\ \vdots \\ 0 \end{bmatrix}. \quad (\text{B.23})$$

Appendix C

Summation Distribution of Lambda

In this appendix, we shall derive the mathematical expectation and variance of the variable $\mu = \sum_{k=1}^K \tau_k \lambda_k$, where $\tau_1, \tau_2 \dots \tau_K$ are real numbers, independent from λ_k . We decompose $\lambda_k = \lambda_{kr} + i\lambda_{ki}$. Let the covariance matrix of $[\lambda_{kr}, \lambda_{ki}]^T$ be $\mathcal{V}(\lambda_{kr}, \lambda_{ki})$. Since all λ_k are independent and identically distributed, the covariance matrix of $[\mu_r, \mu_i]^T$ can be expressed as

$$\mathcal{V}(\mu_r, \mu_i) = \sum_{k=1}^K S_k \mathcal{V}(\lambda_{kr}, \lambda_{ki}) S_k' \quad (\text{C.1})$$

$$\text{with, } S_k = \begin{bmatrix} \tau_k & 0 \\ 0 & \tau_k \end{bmatrix} \quad (\text{C.2})$$

We shall then apply the central limit theorem on μ to show that its covariance matrix will be S , and the distribution will be close to Gaussian. We also provide numerical validation of this claim. In the following, we shall discuss about the density function of λ . Next we derive the mathematical expectation and variance of λ . Finally, we produce density function of μ .

C.1 Probability Density of λ

Consider (5.7). For simplicity, we consider a single user, i.e. $L = 1$. Let $u_k = (1 - \alpha)h_k + e_k + \eta_k$ and $v_k = h_k + e_k$, then $\lambda_k = \frac{u_k}{v_k} = \lambda_{kr} + i\lambda_{ki}$. Since all λ_k are independent and identically distributed, we drop the subscript k . Define:

$$\Sigma = \begin{bmatrix} \mathbb{E}(u^*u) & \mathbb{E}(u^*v) \\ \mathbb{E}(uv^*) & \mathbb{E}(v^*v) \end{bmatrix} = \begin{bmatrix} \sigma_u & \rho\sigma_u\sigma_v \\ \rho\sigma_u\sigma_v & \sigma_v \end{bmatrix}, \quad (\text{C.3})$$

where

$$\begin{aligned} \sigma_u^2 &= |1 - \alpha|^2 \sigma_h^2 + \sigma_e^2 + \sigma_\eta^2 \\ \sigma_v^2 &= \sigma_h^2 + \sigma_e^2 \\ \rho &= (1 - \alpha) \sigma_h^2 + \sigma_e^2. \end{aligned} \quad (\text{C.4})$$

The variable λ has the probability density function (PDF) [90]:

$$f(\lambda_r, \lambda_i) = \frac{(1 - |\rho|^2) \sigma_u^2 \sigma_v^2}{\pi} \left(\sigma_v^2 (\lambda_r^2 + \lambda_i^2) + \sigma_u^2 - 2\rho_r \lambda_r \sigma_u \sigma_v + 2\rho_i \lambda_i \sigma_u \sigma_v \right)^{-2}. \quad (\text{C.5})$$

Note that,

$$\sigma_v^2 (\lambda_r^2 + \lambda_i^2) + \sigma_u^2 - 2\rho_r \lambda_r \sigma_u \sigma_v + 2\rho_i \lambda_i \sigma_u \sigma_v = (\sigma_v \lambda_r - \rho_r \sigma_u)^2 + (\sigma_v \lambda_i + \rho_i \sigma_u)^2 + (1 - \rho_r^2 - \rho_i^2) \sigma_u^2. \quad (\text{C.6})$$

C.2 Mathematical Variance of λ

The covariance matrix of $[\lambda_r, \lambda_i]^T$ be

$$\mathcal{V}(\lambda_r, \lambda_i) = \begin{bmatrix} \sigma_r^2 & \sigma_{ri} \\ \sigma_{ri} & \sigma_i^2 \end{bmatrix}, \quad (\text{C.7})$$

where $\sigma_r^2 = \mathbb{E}(\lambda_r^2) - (\mathbb{E}(\lambda_r))^2$, $\sigma_{ri} = \mathbb{E}(\lambda_r \lambda_i) - \mathbb{E}(\lambda_r)\mathbb{E}(\lambda_i)$; and $\mathbb{E}(x)$ denotes the mathematical expectation of x . $\mathbb{E}(\lambda_r)$ can be calculated as

$$\mathbb{E}(\lambda_r) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \lambda_r f(\lambda_r, \lambda_i) d\lambda_r d\lambda_i. \quad (\text{C.8})$$

We evaluate the integral in polar coordinate. Let:

$$\begin{aligned} \gamma &= \sqrt{\frac{(1 - \rho_r^2 - \rho_i^2)\sigma_u^2}{\sigma_v^2}}, \\ \lambda_i &= r \sin \theta - \frac{\rho_i \sigma_u}{\sigma_v}, \quad \lambda_r = r \cos \theta + \frac{\rho_r \sigma_u}{\sigma_v}; \\ \hat{\alpha} &= \frac{\rho_r \sigma_u}{\sigma_v}, \quad \beta = \frac{-\rho_i \sigma_u}{\sigma_v}, \quad \Gamma = \frac{(1 - |\rho|^2)\sigma_u^2}{\pi \sigma_v^2} \end{aligned} \quad (\text{C.9})$$

Then by combining (C.5) and (C.9), we have

$$\begin{aligned} \mathbb{E}(\lambda_r) &= \Gamma \int_0^{2\pi} \int_0^{\infty} \left(\frac{r^2 \cos \theta}{(r^2 + \gamma^2)^2} + \frac{r \hat{\alpha}}{(r^2 + \gamma^2)^2} \right) dr d\theta \\ &= \Gamma \int_0^{2\pi} \left[\cos \theta \left(\frac{-r}{2(r^2 + \gamma^2)} + \frac{1}{2\gamma} \tan^{-1}(r/\gamma) \right) - \frac{\hat{\alpha}}{2(r^2 + \gamma^2)} \right]_0^{\infty} d\theta \\ &= \Gamma \int_0^{2\pi} \left[-\left(\frac{\pi}{4\gamma} \right) \cos \theta + \frac{\hat{\alpha}}{2\gamma^2} \right] d\theta \\ &= \frac{\Gamma \pi \hat{\alpha}}{\gamma^2} = \frac{(1 - |\rho|^2) \rho_r \sigma_u^3}{\sigma_v \sigma_u^2 (1 - \rho_r^2 - \rho_i^2)} \\ &= \frac{(1 - |\rho|^2) \rho_r \sigma_u}{\sigma_v (1 - \rho_r^2 - \rho_i^2)} \\ &= \hat{\alpha}. \end{aligned} \quad (\text{C.10})$$

Next we compute $\mathbb{E}(\lambda_r^2)$ by

$$\mathbb{E}(\lambda_r^2) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \lambda_r^2 f(\lambda_r, \lambda_i) d\lambda_r d\lambda_i. \quad (\text{C.11})$$

Using the similar conversion in (C.9), we have

$$\begin{aligned}
\mathbb{E}(\lambda_r^2) &= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^3 \cos^2 \theta}{(r^2 + \gamma^2)^2} + \frac{2\hat{\alpha} r^2 \cos \theta}{(r^2 + \gamma^2)^2} + \frac{r \hat{\alpha}^2}{(r^2 + \gamma^2)^2} \right) dr d\theta \\
&= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^3 \cos^2 \theta}{(r^2 + \gamma^2)^2} \right) dr d\theta + \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{2\hat{\alpha} r^2 \cos \theta}{(r^2 + \gamma^2)^2} + \frac{r \hat{\alpha}^2}{(r^2 + \gamma^2)^2} \right) dr d\theta \\
&= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^3 \cos^2 \theta}{(r^2 + \gamma^2)^2} \right) dr d\theta + \Gamma \int_0^{2\pi} \left[- \left(\frac{\pi \hat{\alpha}}{2c} \right) \cos \theta + \frac{\hat{\alpha}^2}{2\gamma^2} \right] d\theta \\
&= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^3 \cos^2 \theta}{(r^2 + \gamma^2)^2} \right) dr d\theta + \frac{\Gamma \pi \hat{\alpha}^2}{\gamma^2} \\
&= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^3 \cos^2 \theta}{(r^2 + \gamma^2)^2} \right) dr d\theta + \hat{\alpha}^2
\end{aligned} \tag{C.12}$$

There is no closed form solution of the first integral term. In fact, the solution become unbounded when $r \rightarrow \infty$. However, it can be bounded by some fixed value of $r < \infty$. In the integration, we consider that both (λ_r, λ_i) can take values between $[-\infty, \infty]$. However in practice, σ_h is bounded and $\sigma_h > \sigma_e, \sigma_\eta$. Hence, there is a very low probability of (λ_r, λ_i) taking a very large magnitude. An approach to obtain a reliable bound of (λ_r, λ_i) would be using Monte Carlo simulations. Under this approach, one simulates a large number of realisations of the variable λ according to (5.7); then, these realisations can be used to compute an estimate of the PDF of $f(\lambda_r, \lambda_i)$. Note that the accuracy of the distribution function increases with an increase in the number of simulations performed. Thus, one can plot the distribution function with respect to (λ_r, λ_i) and the bound of (λ_r, λ_i) can be estimated from the boundary of the distribution function within which 99.9% energy resides.

Finally we compute,

$$\mathbb{E}(\lambda_r \lambda_i) = \int_{-\infty}^\infty \int_{-\infty}^\infty \lambda_r \lambda_i f(\lambda_r, \lambda_i) d\lambda_r d\lambda_i. \tag{C.13}$$

In polar coordinate,

$$\begin{aligned}
 \mathbb{E}(\lambda_r \lambda_i) &= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^3 \cos \theta \sin \theta}{(r^2 + \gamma^2)^2} + \frac{r^2(\beta \cos \theta + \hat{\alpha} \sin \theta)}{(r^2 + \gamma^2)^2} + \frac{r \hat{\alpha} \beta}{(r^2 + \gamma^2)^2} \right) dr d\theta \\
 &= \Gamma \int_0^{2\pi} \int_0^\infty \left(\frac{r^2(\beta \cos \theta + \hat{\alpha} \sin \theta)}{(r^2 + \gamma^2)^2} + \frac{r \hat{\alpha} \beta}{(r^2 + \gamma^2)^2} \right) dr d\theta \\
 &= \frac{\Gamma \pi \hat{\alpha} \beta}{\gamma^2} \\
 &= \beta \hat{\alpha}.
 \end{aligned} \tag{C.14}$$

C.3 Probability Density of μ

The PDF of the variable μ can be computed by using the following two steps:

1. Compute characteristic function of μ , and
2. The inverse Fourier transform of the characteristic function will give the PDF of μ .

The characteristic function of μ is by definition the function:

$$\Phi(\omega, \varpi) = \prod_{k=1}^K \mathbb{E}(e^{-i\tau_k(\omega \lambda_{kr} + \varpi \lambda_{ki})}), \tag{C.15}$$

where $\mathbb{E}(x)$ denotes the mathematical expectation of x . Now using (C.5), we can write:

$$\begin{aligned}
 &\mathbb{E}(e^{-i\tau(\omega \lambda_r + \varpi \lambda_i)}) \\
 &= \frac{(1 - |\rho|^2) \sigma_u^2 \sigma_v^2}{\pi} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \frac{e^{-i\tau(\omega \lambda_r + \varpi \lambda_i)}}{\{(\sigma_v \lambda_r - \rho_r \sigma_u)^2 + (\sigma_v \lambda_i + \rho_i \sigma_u)^2 + (1 - \rho_r^2 - \rho_i^2) \sigma_u^2\}^2} d\lambda_r d\lambda_i
 \end{aligned} \tag{C.16}$$

In polar coordinates, by using (C.9), we get

$$\begin{aligned}
 \mathbb{E}(e^{-i\tau(\omega \lambda_r + \varpi \lambda_i)}) &= \Gamma \int_0^{2\pi} \int_0^\infty r \frac{e^{-i(\tau \omega \{\alpha + r \cos(\theta)\} + \tau \varpi \{\beta + r \sin(\theta)\})}}{\{r^2 + \gamma^2\}^2} dr d\theta \\
 &= e^{-i(\omega \alpha + \varpi \beta)} \int_0^{2\pi} \int_0^\infty r \frac{e^{-ir(\tau \omega \cos(\theta) + \tau \varpi \sin(\theta))}}{\{r^2 + \gamma^2\}^2} dr d\theta.
 \end{aligned} \tag{C.17}$$

Write $\eta(\omega, \varpi, \theta) := r\omega \cos(\theta) + r\varpi \sin(\theta)$, and

Define

$$J(\omega, \varpi, \theta) := \int_0^\infty \frac{r e^{-ir \eta(\omega, \varpi, \theta)}}{\{r^2 + \gamma^2\}^2} dr. \quad (\text{C.18})$$

Then

$$\mathbb{E}(e^{-i\tau(\omega\lambda_r + \varpi\lambda_i)}) = \int_0^{2\pi} J(\omega, \varpi, \theta) d\theta, \quad (\text{C.19})$$

and J is essentially the Fourier transform of

$$\frac{r}{(r^2 + \gamma^2)^2}, \quad r \geq 0,$$

evaluated at $\eta(\omega, \varpi, \theta)$.

Finally, the probability density function of $[\mu_r, \mu_i]$ can be expressed in terms of $\Phi(\omega, \varpi)$ as

$$\hat{f}(\mu_r, \mu_i) = \int_{-\infty}^\infty \int_{-\infty}^\infty \Phi(\omega, \varpi) e^{i(\omega\mu_r + \varpi\mu_i)} d\omega d\varpi \quad (\text{C.20})$$

Note that, there is no closed form expression for the Fourier transform in (C.18). Hence, we need to compute (C.19) and (C.20) numerically in discrete domain.

C.4 Simulation

In this section we provide numerical validation of our claim in Section 5.1.1. By using the central limit theorem, we show that the distribution of μ can be approximated by a Gaussian distribution. We further show that the variance of μ , computed using (C.1) and the procedure in Section C.3, are similar.

C.4.1 The Distribution of μ

Let a two dimensional random vector $[x \ y]^T$ has multivariate normal distribution. The mean and variance of the vector are

$$\mathbf{v} = \begin{bmatrix} v_x \\ v_y \end{bmatrix}, \quad \mathcal{V}(x, y) = \begin{bmatrix} \sigma_x^2 & \delta \sigma_x \sigma_y \\ \delta \sigma_x \sigma_y & \sigma_y^2 \end{bmatrix}, \quad (\text{C.21})$$

where δ is the correlation between x and y . The PDF of the vector is

$$\tilde{f}(x, y) = \frac{1}{2\pi v_x v_y \sqrt{1 - \delta^2}} \exp \left(\frac{-1}{2(1 - \delta^2)} \left[\frac{(x - v_x)^2}{\sigma_x^2} + \frac{(y - v_y)^2}{\sigma_y^2} - \frac{2\delta(x - v_x)(y - v_y)}{\sigma_x \sigma_y} \right] \right). \quad (\text{C.22})$$

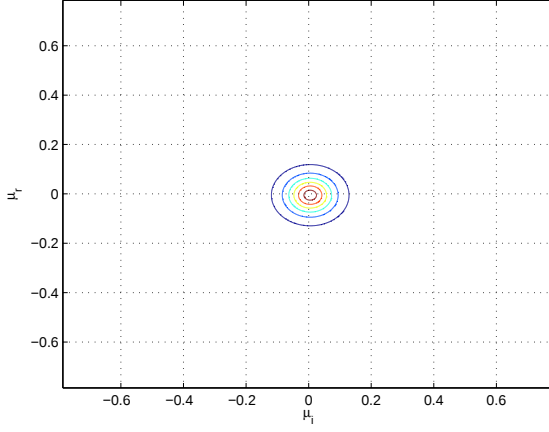
We aim to check how closely the probability distribution $\hat{f}(\mu_r, \mu_i)$, evaluated in (C.20), is matched with the Gaussian distribution $\tilde{f}(\mu_r, \mu_i)$. We evaluate $\hat{f}(\mu_r, \mu_i)$ in (C.20) for different values of μ_r and μ_i . We then apply a nonlinear regression to approximate the values of coefficients $v_x, v_y, \sigma_x, \sigma_y$ and δ in (C.22), such that the mean squared error (MSE) between $\hat{f}(\mu_r, \mu_i)$ and $\tilde{f}(\mu_r, \mu_i)$ is minimised. The MSE can be used as an indicator of fitness, i.e, a smaller MSE indicates that $\hat{f}(\mu_r, \mu_i)$ is closely matched with Gaussian distribution. In particular, we use the Levenberg-Marquardt algorithm or nonlinear least squares to compute the fit [113].

We apply the central limit theorem to validate the distribution. We have $\mu = \sum_{k=1}^K \tau_k \lambda_k$ and set $\tau_k \in [+1, -1]$. Starting from a smaller K , we allow K to take larger value. We found that the fitness increasing with increasing the value of K . Figure C.1 shows the PDF of $\mu = \sum_{k=1}^K \tau_k \lambda_k$ for different values of K . As can be seen from the Figures C.1(a)-(c), the fitness between $\hat{f}(\mu_r, \mu_i)$ and $\tilde{f}(\mu_r, \mu_i)$ increases with increasing the value of K . We take $\tau_k \in [+1, -1]$, hence the mean of $[\mu_r, \mu_i]$ remain close to zero in Figures C.1(a)-(c).

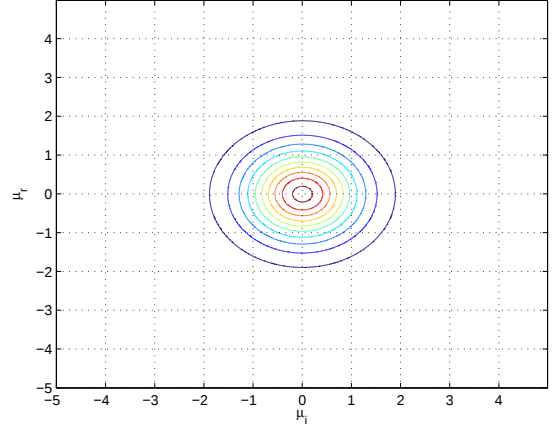
To illustrate how the mean changes with K , we take $\tau_k = 1$ in Figure C.1(d). In Figure C.2, we demonstrate how the MSE between $\hat{f}(\mu_r, \mu_i)$ and standard Gaussian distribution in (C.22) decreases with increasing the value of K .

C.4.2 Comparing Calculated Value and Simulated Value

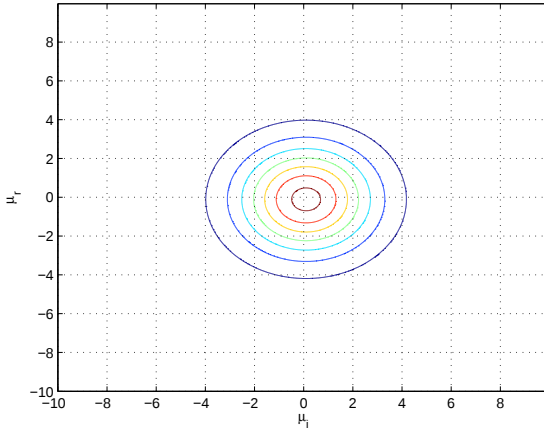
Figure C.3 demonstrates how the experimental mean and variance (by using simulations similar in Figures C.1(a)-(c)) are matched with the calculated expected value and variance evaluated in (C.1). To observe the evaluation of the variance of μ_r in Figure C.3(a), we fix a value of K and set $\tau_k \in [+1, -1]$. We then compute $\mathcal{V}(\mu_r, \mu_i)$ using (C.1) and (C.7). Again for the fixed K , we evaluate $\hat{f}(\mu_r, \mu_i)$ in (C.20) for different values of μ_r and μ_i and fit it with the Gaussian distribution in (C.22). The Gaussian variance will be called the experimental variance.



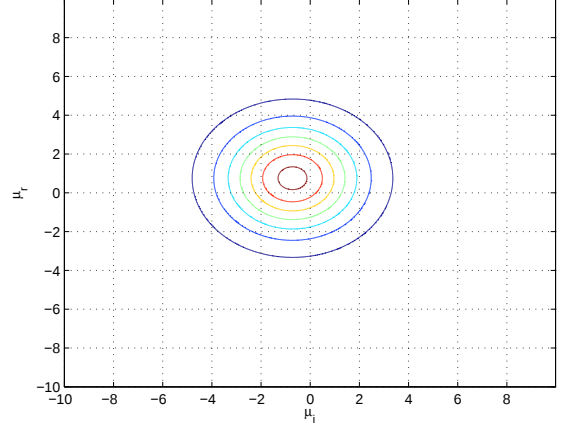
(a) $K = 1, \tau_k \in +1, -1; \sigma_{\mu_r} = 0.0077, \text{MSE} = 9.0477e - 1.$



(b) $K = 35, \tau_k \in +1, -1; \sigma_{\mu_r} = 0.9064, \text{MSE} = 2.9530e - 5.$



(c) $K = 144, \tau_k \in +1, -1; \sigma_{\mu_r} = 2.1269, \text{MSE} = 7.0171e - 7.$



(d) $K = 144, \tau_k = 1; \sigma_{\mu_r} = 2.1272, \text{MSE} = 7.0540e - 007, \text{mean}(\mu_r) = 0.7542.$

Figure C.1: Plot of PDF using $\hat{f}(\mu_r, \mu_i)$ in (C.20) and compute its fitness with $\tilde{f}(x, y)$ in (C.22). Distribution parameters are $\sigma_h = 1, \sigma_e = 0.05, \sigma_\eta = 0.05$. The value of δ is very small.

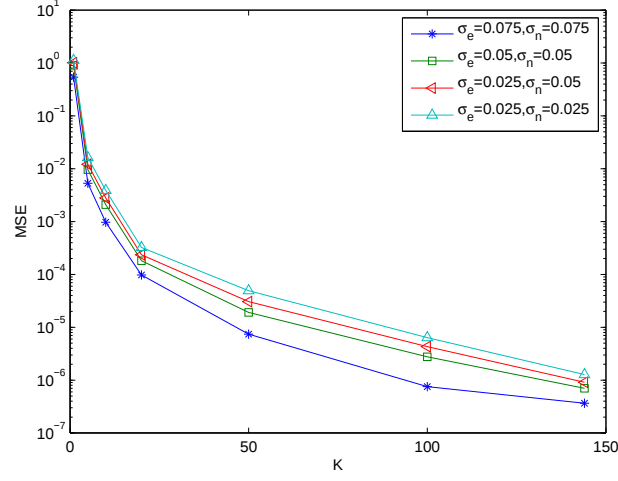
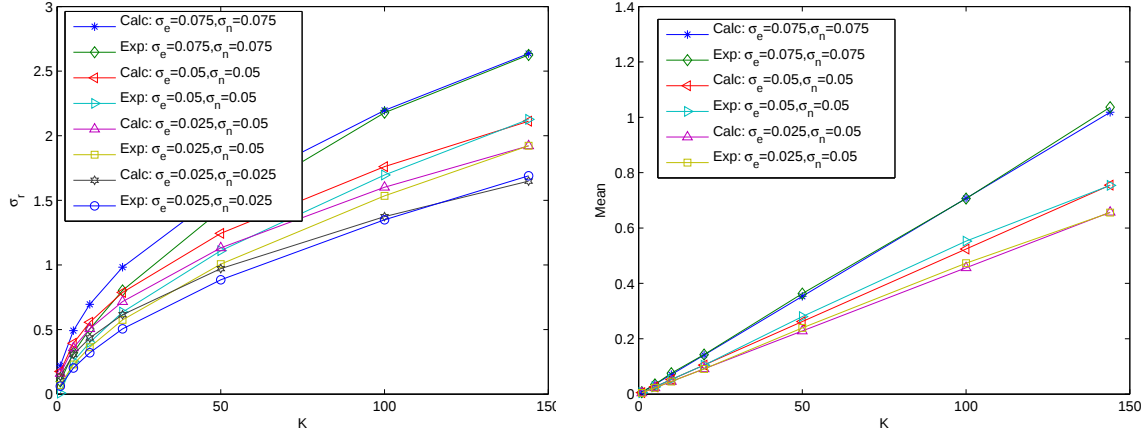


Figure C.2: Fitness update with K . Distribution parameters $\sigma_h = 1$.



(a) Evaluation of calculated variance and experimental variance of μ_r

(b) Evaluation of calculated mean and experimental mean of μ_r

Figure C.3: Comparison of calculated value and experimental value of different coefficients. Distribution parameter $\sigma_h = 1$.

Appendix D

Voltage Instability Prediction Algorithm

In this appendix, we describe the details of the voltage instability prediction algorithm used in Section 7.2 to enable adaptive synchrophasor communications.

D.1 The Prediction Algorithm

By combining (7.9) and (7.10) we have

$$P_t = \frac{Q_t X_t}{R_t} + \frac{V_t E_t}{R_t} \cos \theta_t - \frac{V_t^2}{R_t}. \quad (\text{D.1})$$

By assuming Q_t, X_t, E_t, θ_t remain constant over a short time interval $t_{(\ell-n)}$ to t_ℓ (with $\ell > n$), we can write

$$P_t = \alpha_1(\ell) + \alpha_2(\ell)V_t + \alpha_3(\ell)V_t^2, \quad t \in [t_{\ell-n}, t_\ell], \quad (\text{D.2})$$

where $\alpha_1(\ell) = \frac{Q_t X_t}{R_t}, \alpha_2(\ell) = \frac{E_t}{R_t} \cos \theta_t, \alpha_3(\ell) = -1/R_t$ remains constant in the fixed time interval. This assumption is common in some state-of-art techniques [105, 106]. We can

calculate the values of $\alpha_1(\ell), \alpha_2(\ell), \alpha_3(\ell)$ by using the past n time sampled data as

$$\begin{bmatrix} 1 & V_{t_\ell} & V_{t_\ell}^2 \\ 1 & V_{t_{\ell-1}} & V_{t_{\ell-1}}^2 \\ \vdots & \vdots & \vdots \\ 1 & V_{t_{\ell-n+1}} & V_{t_{\ell-n+1}}^2 \end{bmatrix} \begin{bmatrix} \alpha_1(\ell) \\ \alpha_2(\ell) \\ \alpha_3(\ell) \end{bmatrix} = \begin{bmatrix} P_{t_\ell} \\ P_{t_{\ell-1}} \\ \vdots \\ P_{t_{\ell-n+1}} \end{bmatrix}, \quad (\text{D.3})$$

i.e.,

$$A\alpha = \mathbf{y},$$

and solving an optimisation like

$$\min_{\alpha} \|\mathbf{y} - A_v\alpha\|_2, \quad (\text{D.4})$$

where, $\mathbf{y} = [P_{t_\ell}, P_{t_{\ell-1}}, \dots, P_{t_{\ell-n+1}}]^T$, and $\mathbf{x}^T$ indicates transpose of vector $\mathbf{x}$.

The equation (D.2) will construct a parabola. If the parabola can follow the P-V characteristics curve (see Figure 7.6), then the maximum load transfer $P_{\max}$ can be approximated by using the focus of the parabola, i.e.

$$P_{\max}^\ell = \frac{\alpha_2(\ell)^2 + 4\alpha_3(\ell)\alpha_1(\ell)}{4\alpha_3(\ell)}. \quad (\text{D.5})$$

By combining (7.13) and (D.2), it can be shown that the ZIP load curve will intersect the parabola load curve at:

$$V_i = \frac{-(\alpha_2(\ell) - k_t a_1) \pm \Gamma}{2(\alpha_3(\ell) - k_t a_2)}, \quad (\text{D.6})$$

where $\Gamma = \sqrt{(\alpha_2(\ell) - k_t a_1)^2 - 4(\alpha_3(\ell) - k_t a_2)(\alpha_1(\ell) - k_t P_c)}$.

At critical loading k_* , the value of $\Gamma = 0$. Hence,

$$k_* = \frac{-q \pm \sqrt{q^2 - 4(a_1^2 - 4P_c a_2)(\alpha_2(\ell)^2 - 4\alpha_1(\ell)\alpha_3(\ell))}}{2(a_1^2 - 4P_c a_2)},$$

where, $q = 4\alpha_3(\ell)P_c + 4\alpha_1(\ell)a_2 - 2\alpha_2(\ell)^2 a_1$.

The prediction $P_{\max}^\ell$ and k_* will be accurate if the values of Q_t, X_t, E_t, θ_t will remain constant throughout the time t_ℓ to $t = C$. However, it true only for small values of n . Further, the voltage stabilizer of modern power system keeps the voltage profile virtually fixed [114], until it reaches closer to the instability point. Hence, it is difficult to predict the P-V characteristics of a system by using sampled data, when $T - t_\ell$ is not small. Therefore, developing an prediction algorithm by only observing the P-V characteristics may not be reliable. To overcome the difficulty, we also estimate the P- θ characteristics of the system. Let us consider (7.9). By assuming E_t, X_t, V_t remain constant over a short interval $t_{(\ell-n)}$ to t_ℓ , we can write

$$P_t = \beta_1(\ell) + \beta_2(\ell) \cos \theta_t + \beta_3(\ell) \sin \theta_t, \quad t \in [t_{\ell-n}, t_\ell]. \quad (\text{D.7})$$

By using the approaches (D.3) and (D.4), we can estimate β , as well as the P- θ characteristics curve. The maximum of the curve in (D.7) (say $\hat{P}_{\max}^\ell$) will be an estimate of $P_{\max}$.

Let at time t_ℓ , the estimation error $\|P_{\max} - \hat{P}_{\max}^\ell\|$ is smaller than $\|P_{\max} - P_{\max}^\ell\|$, i.e. the P- θ curve approximate the $P_{\max}$ more accurately than P-V curve¹. However, we have to use P-V curve and (D.2), (D.6) to detect voltage instability point V_C . Hence, we need to improve the prediction accuracy of the P-V curve by using (D.2). In this work, we do this by moving $P_{\max}^\ell$ closer to $\hat{P}_{\max}^\ell$, i.e. we redefine the optimisation (D.4)

$$\begin{aligned} & \min_{\alpha} \|\mathbf{y} - A_v \alpha\|_2 \\ \text{sub. to } & \frac{\alpha_2(\ell)^2 + 4\alpha_3(\ell)\alpha_1(\ell)}{4\alpha_3(\ell)} = \hat{P}_{\max}^\ell \end{aligned} \quad (\text{D.8})$$

The optimisation cannot be solved directly. Hence, we modify the optimisation in the following way. If $P_{\max}^\ell > \hat{P}_{\max}^\ell$ then

$$\begin{aligned} & \min_{\alpha} \|\mathbf{y} - A_v \alpha\|_2 \\ \text{sub. to } & \frac{\alpha_2(\ell)^2 + 4\alpha_3(\ell)\alpha_1(\ell)}{4\alpha_3(\ell)} \leq \hat{P}_{\max}^\ell \end{aligned} \quad (\text{D.9})$$

¹A procedure of identifying prediction accuracy will be described in Section D.2.

Table D.1: Prediction Procedure

-
1. Calculate the parameters of (D.11) for all load buses.
 2. If $\Delta V^s > \xi$ then
 3. Initialise the adaptive PMU reporting scheme (Subsection 7.2.2).
 4. Calculate the values of α and β in (D.2), (D.7) respectively, by using (D.4). Compute $P_{\max}^\ell$ and $\hat{P}_{\max}^\ell$ using (D.5) and (D.7).
 5. If $|P_{\max}^\ell - \hat{P}_{\max}^\ell| > \varepsilon$, then
 6. Compare prediction accuracy of $P_{\max}$ by $P_{\max}^\ell$ and $\hat{P}_{\max}^\ell$ using Section D.2.
 7. If the prediction accuracy by $\hat{P}_{\max}^\ell$ is higher, then
 8. Update the value of α in (D.2) by using the optimisation in (D.8).
 9. End If
 10. End if
 11. Compute the magnitude of voltage V_{pm}^ℓ corresponding to the maximum power transfer $P_{\max}^\ell$ by using (D.2). Estimate unstable voltage V_C^ℓ by using (D.6):

$$V_C^\ell = \frac{-(\alpha_2(\ell) - k_* a_1)}{2(\alpha_3(\ell) - k_* a_2)}$$
 12. Approximate the times t_{pm}, t_s required to reach the voltage of bus s to V_{pm}^ℓ and V_C^ℓ respectively.

$$V_{t_\ell}^s - V_{pm}^\ell = \Delta V_{t_\ell}^s t_{pm} + \frac{1}{2} \delta (\Delta V_{t_\ell}^s) t_{pm}^2$$

$$V_{t_\ell}^s - V_C^\ell = \Delta V_{t_\ell}^s t_s + \frac{1}{2} \delta (\Delta V_{t_\ell}^s) t_s^2.$$
 13. If $V_{t_\ell}^s > V_{pm}^\ell$, then
 15. Set $t_*^s = t_s$ and $V_*^s = V_C^\ell$;
 14. Else t_*^s will be estimated using:

$$V_{t_\ell}^s - V_*^s = \Delta V_{t_s}^s t_*^s + \frac{1}{2} \delta (\Delta V_{t_s}^s) (t_*^s)^2.$$

The value of t_*^s approximates the time required to move the voltage of bus s i.e. $V_{t_\ell}^s$ to the point of instability i.e. V_C .
 15. End If
 18. If $t_* < \eta$ then
 19. Take proper control action to manage load on bus s to avoid possible voltage instability. The value of η will be determined by the latency of communication and control system.
 20. End If
 21. End If

However, if $\{P\}_{\max}^\ell < \hat{P}_{\max}^\ell$ then the optimisation becomes difficult and we approximate the solution in the following way

$$\begin{aligned} & \min_{\alpha, z} z + \tau \|\mathbf{y} - A_v \alpha\|_2 \\ & \text{sub. to, } 4(\hat{P}_{\max}^\ell - \alpha_1(\ell)) = z; \alpha_2(\ell)^2 - \alpha_3(\ell)z \leq 0, \end{aligned} \quad (\text{D.10})$$

where $\tau > 0$ is a tuning factor (see Section D.4). Both (D.9) and (D.10) are convex and has only four variables.

D.2 Comparison of Prediction Accuracy

Let the system will transfer $P_{\max}$ at $t = T$ sec. If the system P-V curve follows the expression in (D.2) approximately, one may expect that the prediction accuracy of $P_{\max}$ will increase when $|t_\ell - T|$ becomes smaller, where $|\cdot|$ indicates absolute value. Let the prediction of $P_{\max}$ be $P_{\max}^\ell$ and prediction error $e_\ell = |P_{\max}^\ell - P_{\max}|$. Consider a time sequence $\{t_{\ell_i}\}_{i=1}^m$ such that $|t_{\ell_1} - T| > |t_{\ell_2} - T| \cdots > |t_{\ell_m} - T|$. Then we may expect $e_{\ell_1} > e_{\ell_2} > \cdots e_{\ell_m}$. Let another different equation is also applied to estimate $P_{\max}$ and its estimation error is indicated by $\hat{e}_{\ell_i} = |\hat{P}_{\max}^{\ell_i} - P_{\max}|$. Let $\hat{P}_{\max}^{\ell_i}$ is closer to $P_{\max}$ than $P_{\max}^{\ell_i}$, i.e. $e_{\ell_i} > \hat{e}_{\ell_i}$. Since all equations are trying to converge to $P_{\max}$ with increasing t_ℓ , hence, the change of $P_{\max}^\ell$ with time i.e., $|P_{\max}^{\ell_{i-1}} - P_{\max}^{\ell_i}|$ will be larger than $|\hat{P}_{\max}^{\ell_{i-1}} - \hat{P}_{\max}^{\ell_i}|$. Let us observe a few time sequences of $P_{\max}^{\ell_i}$ and $\hat{P}_{\max}^{\ell_i}$. If $\hat{P}_{\max}^{\ell_i}$ is closer approximation of $\{P\}_{\max}$, we can expect: $|P_{\max}^{\ell_{i-1}} - P_{\max}^{\ell_i}| < |\hat{P}_{\max}^{\ell_{i-1}} - \hat{P}_{\max}^{\ell_i}|$; for, $i = 2, 3, \cdots m$. Hence, by comparing the magnitudes of $|P_{\max}^{\ell_{i-1}} - P_{\max}^{\ell_i}|$ and $|\hat{P}_{\max}^{\ell_{i-1}} - \hat{P}_{\max}^{\ell_i}|$, we can predict which equation approximate the $P_{\max}$ closely.

In our proposed approach, two different equations are trying to approximate two different curves. Equation (D.2) approximates P-V curve, and (D.7) approximates P- θ curve. However, the value of maximum (i.e. $P_{\max}$) is same for the two curves. Hence, we can utilise the above observation to compare the prediction accuracy of $P_{\max}$ by an equation.

D.3 Voltage Instability Predication Procedure

Let the current time instant be t_ℓ . We assume that we have the measurements of voltage magnitude and angle of all load buses at time instants $\{t_r\}_{r=1}^\ell$. Let at time t_ℓ , the voltage magnitude and angle of a load bus s be $V_{t_\ell}^s$ and $\theta_{t_\ell}^s$, respectively. We can calculate the changes of voltage magnitude, angle and the voltage decline rate [115] as

$$\Delta V_{t_\ell}^s = \frac{V_{t_{\ell_m}}^s - V_{t_\ell}^s}{M}; \Delta \theta_{t_\ell}^s = \frac{\theta_{t_{\ell_m}}^s - \theta_{t_\ell}^s}{M};$$

$$\delta(\Delta V_{t_\ell}^s) = \frac{\Delta V_{t_{\ell_m}}^s - \Delta V_{t_\ell}^s}{M}; \quad (\text{D.11})$$

where $t_\ell > t_{\ell_m}$ and $t_\ell - t_{\ell_m} = M$ sec. At t_ℓ , we run the procedure in Table D.1.

D.4 Tuning Parameter τ

The value of τ in (D.10) can be tuned by using the historical load data of the power system. Let the historical data shows that the incident of Step 2 of Table-D.1 is satisfied for S number of time instants $\{\kappa_i\}_{i=1}^S$ for different buses. Select a time instant κ_i . The historical data can help us to predict the actual value of $P_{\max}$ and V_{pm} after the incident at κ_i . Now use the value of $P_{\max}$ in (D.10) instead of $\hat{P}_{\max}^\ell$. Tune the value of τ in the range $[10^2, 10^6]$ for which the constraint in (D.8) is closely satisfied i.e., $\|\frac{\alpha_2(\ell)^2 + 4\alpha_3(\ell)\alpha_1(\ell)}{4\alpha_3(\ell)} - P_{\max}\| < 1^{-3}$. Let the value of τ is indicated by $\hat{\tau}_i$. The procedure is repeated for all times $\{\kappa_i\}_{i=1}^S$. The final value of $\tau = \frac{1}{S} \sum \hat{\tau}_i$.

Bibliography

- [1] H. Farhangi, “The path of the smart grid,” *Power and Energy Magazine, IEEE*, vol. 8, no. 1, pp. 18–28, 2010.
- [2] “The Smart Grid: An Introduction,” U.S. Department of Energy (DOE), Tech. Rep., 2010.
- [3] “Clean Energy Australia Report 2012,” Clean Energy Council, Australia, Tech. Rep., 2013.
- [4] “Energy market arrangements for electric and natural gas vehicles,” Australian Energy Market Commission, Tech. Rep., Jan. 2012.
- [5] W. Wang, Y. Xu, and M. Khanna, “A survey on the communication architectures in smart grid,” *Computer Networks*, vol. 55, no. 15, pp. 3604–3629, 2011.
- [6] “WiMAX applications for utilities,” *WiMAX Forum*, Oct. 2008.
- [7] “WiGRID : A Wireless Broadband Solution for Utilities,” *WiMAX Forum*, 2013.
- [8] “Machine-to-Machine (M2M) communications study report,” IEEE 802.16 Broadband Wireless Access Working Group, Tech. Rep., 2010.
- [9] “IEEE standard for air interface for broadband wireless access systems - Amendment 1: enhancements to support machine-to-machine applications,” *IEEE Std 802.16p-2012 (Amendment to IEEE Std 802.16-2012)*, pp. 1–82, 2012.
- [10] S. Karnouskos, “Cyber-physical systems in the smartgrid,” in *Industrial Informatics (INDIN), 2011 9th IEEE International Conference on*. IEEE, 2011, pp. 20–23.

- [11] R. H. Khan and J. Y. Khan, "A comprehensive review of the application characteristics and traffic requirements of a smart grid communications network," *Computer Networks*, 2012.
- [12] "IEEE guide for smart grid interoperability of energy technology and information technology operation with the electric power system (EPS), end-use applications, and loads," *IEEE Std 2030-2011*, pp. 1–126, 2011.
- [13] "American national standard for utility industry end device data tables," *ANSI C12.19-2008*, March 2009.
- [14] "IEEE standard for utility industry metering communication protocol application layer (End device data tables)," *IEEE Std 1377-2012 (Revision of IEEE Std 1377-1997)*, pp. 1–576, 2012.
- [15] "Application integration at electric utilities system interfaces for distribution management, part 9: interfaces for meter reading and control," *IEC 61968-9 Std*, Sept. 2009.
- [16] "Guidelines for assessing wireless standards for Smart Grid applications," NIST Priority Action Plan 2 (PAP2), Tech. Rep., 2010.
- [17] G. R. Newsham and B. G. Bowker, "The effect of utility time-varying pricing and load control strategies on residential summer peak electricity use: A review," *Energy Policy*, vol. 38, no. 7, pp. 3289–3296, 2010.
- [18] P. Palensky and D. Dietrich, "Demand side management: demand response, intelligent energy systems, and smart loads," *Industrial Informatics, IEEE Transactions on*, vol. 7, no. 3, pp. 381–388, 2011.
- [19] "Applicability of M2M architecture to Smart Grid networks," ETSI, Tech. Rep. Draft TR 102 935, Version: 0.1.3, 2010.
- [20] S. Shao, M. Pipattanasomporn, and S. Rahman, "Grid integration of electric vehicles and demand response with customer choice," *Smart Grid, IEEE Transactions on*, vol. 3, no. 1, pp. 543–550, 2012.

- [21] K. Clement-Nyns, E. Haesen, and J. Driesen, "The impact of charging plug-in hybrid electric vehicles on a residential distribution grid," *Power Systems, IEEE Transactions on*, vol. 25, no. 1, pp. 371–380, 2010.
- [22] W. Kempton and J. Tomić, "Vehicle-to-grid power implementation: From stabilizing the grid to supporting large-scale renewable energy," *Journal of Power Sources*, vol. 144, no. 1, pp. 280–294, 2005.
- [23] A. Brooks, E. Lu, D. Reicher, C. Spirakis, and B. Wehl, "Demand dispatch," *Power and Energy Magazine, IEEE*, vol. 8, no. 3, pp. 20–29, 2010.
- [24] K. Valentine, W. G. Temple, and K. M. Zhang, "Intelligent electric vehicle charging: Rethinking the valley-fill," *Journal of Power Sources*, vol. 196, no. 24, pp. 10 717–10 726, 2011.
- [25] R. H. Khan, S. Stüdli, and J. Y. Khan, "A network controlled load management scheme for domestic charging of Electric Vehicles," in *Australasian Universities Power Engineering Conference, AUPEC 2013*, 2013, pp. 1–6.
- [26] K. Turitsyn, N. Sinitsyn, S. Backhaus, and M. Chertkov, "Robust broadcast-communication control of electric vehicle charging," in *Smart Grid Communications (SmartGridComm), 2010 First IEEE International Conference on*. IEEE, 2010, pp. 203–207.
- [27] M. Erol-Kantarci, J. H. Sarker, and H. T. Mouftah, "Analysis of plug-in hybrid electrical vehicle admission control in the smart grid," in *Computer Aided Modeling and Design of Communication Links and Networks (CAMAD), 2011 IEEE 16th International Workshop on*. IEEE, 2011, pp. 56–60.
- [28] "Communication networks and systems in substations - Part 5: communication requirements for functions and device models." *IEC Std 61850-5*, 2003.
- [29] "IEEE standard for electric power systems communications - Distributed Network Protocol (DNP3)," *IEEE Std 1815-2010*, pp. 1–775, July 2010.

- [30] A. Apostolov, "Distributed protection, control and recording in IEC 61850 based substation automation systems," in *Developments in Power System Protection, 2004. Eighth IEE International Conference on*, vol. 2. IET, 2004, pp. 647–651.
- [31] V. Terzija, G. Valverde, D. Cai, P. Regulski, V. Madani, J. Fitch, S. Skok, M. M. Begovic, and A. Phadke, "Wide-area monitoring, protection, and control of future electric power networks," *Proceedings of the IEEE*, vol. 99, no. 1, pp. 80–93, 2011.
- [32] "Using spread spectrum radio communication for power system protection relaying applications," *IEEE Power System Relaying Committee, Working Group H2*, 2005.
- [33] M. Yalla, M. Adamiak, A. Apostolov, J. Beatty, S. Borlase, J. Bright, J. Burger, S. Dickson, G. Gresco, W. Hartman *et al.*, "Application of peer-to-peer communication for protective relaying," *Power Delivery, IEEE Transactions on*, vol. 17, no. 2, pp. 446–451, 2002.
- [34] I. Schweitzer, E.O., K. Behrendt, T. Lee, and D. A. Tziouvaras, "Digital communications for power system protection: security, availability, and speed," in *Developments in Power System Protection, 2001, Seventh International Conference on (IEE)*, 2001, pp. 94–97.
- [35] R. Hunt, M. Adamiak, A. King, and S. McCreery, "Application of digital radio for distribution pilot protection," in *Rural Electric Power Conference, 2008 IEEE*, 2008, pp. A3–A3–8.
- [36] "Communication networks and systems for power utility automation," *IEC Std 61850-90-5*, 2012.
- [37] "Protective relay applications using the Smart Grid communication infrastructure," IEEE Power System Relaying Committee (PSRC), Working Group H2, Tech. Rep., 2012.
- [38] A. Apostolov, "IEC 61850 communications based transmission line protection," in *Developments in Power Systems Protection, 2012. DPSP 2012. 11th International Conference on*. IET, 2012, pp. 1–6.

- [39] M. P. Sanders, "Object modeling for pilot channel equipment in iec61850 based devices," in *Protective Relay Engineers, 2010 63rd Annual Conference for.* IEEE, 2010, pp. 1–12.
- [40] A. Smit, "Method and system for programming and implementing automated fault isolation and restoration using sequential logic," Oct. 18 2012, uS Patent 20120265360 A1.
- [41] A. G. Phadke and J. S. Thorp, *Synchronized phasor measurements and their applications.* Springer, 2008.
- [42] "Real-Time Application of Synchrophasors for Improving Reliability," North American Electric Reliability Corporation (NERC), Tech. Rep., October, 2010.
- [43] M. Wache and D. Murray, "Application of synchrophasor measurements for distribution networks," in *Power and Energy Society General Meeting, 2011 IEEE.* IEEE, 2011, pp. 1–4.
- [44] M. Hadley, J. McBride, T. Edgar, L. O'Neil, and J. Johnson, "Securing wide area measurement systems," *US Department of Energy*, 2007.
- [45] "IEEE standard for synchrophasor data transfer for power systems," *IEEE C37.118.2-2011*, pp. 1–53, 2011.
- [46] B. Kroposki, R. Lasseter, T. Ise, S. Morozumi, S. Papatlianassiou, and N. Hatziairgyriou, "Making microgrids work," *Power and Energy Magazine, IEEE*, vol. 6, no. 3, pp. 40–53, 2008.
- [47] T. S. Ustun, C. Ozansoy, and A. Zayegh, "Distributed Energy Resources (DER) object modeling with IEC 61850-7-420," in *Universities Power Engineering Conference (AUPEC), 2011 21st Australasian.* IEEE, 2011, pp. 1–6.
- [48] T. S. Ustun, C. Ozansoy, and Zayegh, "Differential protection of microgrids with central protection unit support," in *TENCON Spring Conference, 2013 IEEE.* IEEE, 2013, pp. 15–19.

- [49] M. Dewadasa, A. Ghosh, and G. Ledwich, "Protection of microgrids using differential relays," in *Universities Power Engineering Conference (AUPEC), 2011 21st Australasian*. IEEE, 2011, pp. 1–6.
- [50] S. Ward and T. Erwin, "Current differential line protection setting considerations," *RFL Electronics Inc*, 2005.
- [51] "Transmission line protection principles," *GE Digital Energy*, 2007, available: <http://www.gedigitalenergy.com/>.
- [52] D. Soldani, M. Li, and R. Cuny, "QoS and QoE management in UMTS cellular networks," 2006.
- [53] "OPNET Modeler 16.0," *Riverbed Technologies, Inc*.
- [54] R. Khan and J. Khan, "A heterogeneous WiMAX-WLAN network for AMI communications in the Smart Grid," in *Smart Grid Communications (SmartGridComm), 2012 IEEE Third International Conference on*, Nov 2012, pp. 710–715.
- [55] Khan, R.H. and Khan, J.Y., "Wide area PMU communication over a WiMAX network in the Smart Grid," in *Smart Grid Communications (SmartGridComm), 2012 IEEE Third International Conference on*, Nov 2012, pp. 187–192.
- [56] R. Khan, T. Ustun, and J. Khan, "Differential protection of microgrids over a WiMAX network," in *Smart Grid Communications (SmartGridComm), 2013 IEEE International Conference on*, Oct 2013, pp. 732–737.
- [57] R. Khan, J. Brown, and J. Khan, "Pilot protection schemes over a multi-service WiMAX network in the Smart Grid," in *Communications Workshops (ICC), 2013 IEEE International Conference on*, June 2013, pp. 994–999.
- [58] "IEEE standard for local and metropolitan area networks - Part 16: air interface for broadband wireless access systems," *IEEE Std 802.16-2009*, pp. 1–2080, 2009.
- [59] "WiMAX System Evaluation Methodology," *WiMAX Forum*, May 2008.

- [60] V. Erceg, L. J. Greenstein, S. Y. Tjandra, S. R. Parkoff, A. Gupta, B. Kulic, A. A. Julius, and R. Bianchi, "An empirically based path loss model for wireless channels in suburban environments," *Selected Areas in Communications, IEEE Journal on*, vol. 17, no. 7, pp. 1205–1211, 1999.
- [61] J. Cha, "Evaluation guideline for comparison of network entry solutions," *IEEE 802.16.p Working Group*, 2011.
- [62] R. Golla, W. Bell, M. Lin, N. Himayat, S. Talwar, and K. Johnsson, "Power outage alarm for smart metering applications," *IEEE 802.16p Task Group*, Sept., 2010.
- [63] M. Adamiak, B. Kasztenny, and W. Premerlani, "Synchrophasors: definition, measurement, and application," *Proceedings of the 59th Annual Georgia Tech Protective Relaying, Atlanta, GA*, 2005.
- [64] F. H. Fitzek, G. Schulte, E. Piri, J. Pinola, M. D. Katz, J. Huusko, K. Pentikousis, and P. Seeling, "Robust Header Compression for WiMAX Femto Cells," *WiMAX Evolution*, p. 185, 2009.
- [65] J. L. Blackburn and T. J. Domin, *Protective relaying: principles and applications*. CRC press, 2006.
- [66] *Distribution Automation Handbook*. ABB Group, 2011.
- [67] J. Krinock, M. Singh, M. Paff, A. Lonkar, L. Fung, and C. Lee, "Comments on OFDMA ranging scheme described in IEEE 802.16 ab-01/01r1," *document IEEE*, vol. 802, 2001.
- [68] R. Khan and k. Mahata, "Random access performance of a WiMAX network for M2M communications in the Smart Grid," in *Smart Grid Communications (Smart-GridComm), 2014 IEEE International Conference on*, June 2014, pp. 994–999.
- [69] A. Vinel, Y. Zhang, M. Lott, and A. Tiurlikov, "Performance analysis of the random access in IEEE 802.16," in *Personal, Indoor and Mobile Radio Communications*,

2005. *PIMRC 2005. IEEE 16th International Symposium on*, vol. 3, 2005, pp. 1596–1600.
- [70] G. Bianchi, “Performance analysis of the IEEE 802.11 distributed coordination function,” *Selected Areas in Communications, IEEE Journal on*, vol. 18, no. 3, pp. 535–547, 2000.
- [71] J. He, K. Guild, K. Yang, and H.-H. Chen, “Modeling contention based bandwidth request scheme for IEEE 802.16 networks,” *Communications Letters, IEEE*, vol. 11, no. 8, pp. 689–700, 2007.
- [72] S. Perera and H. Sirisena, “Analysis of contention based access for best effort traffic on fixed WiMAX,” in *Broadband Communications, Networks and Systems, 2007. BROADNETS 2007. Fourth International Conference on*, 2007, pp. 552–558.
- [73] Y. Fallah, F. Agharebparast, M. Minhas, H. Alnuweiri, and V. Leung, “Analytical modeling of contention-based bandwidth request mechanism in IEEE 802.16 wireless networks,” *Vehicular Technology, IEEE Transactions on*, vol. 57, no. 5, pp. 3094–3107, 2008.
- [74] Q. Ni, L. Hu, A. Vinel, Y. Xiao, and M. Hadjinicolaou, “Performance analysis of contention based bandwidth request mechanisms in WiMAX networks,” *Systems Journal, IEEE*, vol. 4, no. 4, pp. 477–486, 2010.
- [75] G. Giambene and S. Hadzic-Puzovic, “Nonsaturated performance analysis for WiMAX broadcast polling access,” *Vehicular Technology, IEEE Transactions on*, vol. 62, no. 1, pp. 306–325, 2013.
- [76] C. So-In, R. Jain, and A.-K. Tamimi, “Scheduling in IEEE 802.16e mobile WiMAX networks: Key issues and a survey,” *Selected Areas in Communications, IEEE Journal on*, vol. 27, no. 2, pp. 156–171, 2009.
- [77] J.-B. Seo, N.-S. Lee, N.-H. Park, and C.-H. Cho, “Performance analysis of IEEE802.16d random access protocol,” in *Communications, 2006. ICC’06. IEEE International Conference on*, vol. 12. IEEE, 2006, pp. 5528–5533.

- [78] J.-B. Seo and V. Leung, "Design and analysis of backoff algorithms for random access channels in UMTS-LTE and IEEE 802.16 Systems," *Vehicular Technology, IEEE Transactions on*, vol. 60, no. 8, pp. 3975–3989, 2011.
- [79] D. Chuck, K.-Y. Chen, and J. Chang, "A comprehensive analysis of bandwidth request mechanisms in IEEE 802.16 networks," *Vehicular Technology, IEEE Transactions on*, vol. 59, no. 4, pp. 2046–2056, 2010.
- [80] "IEEE 802.16p Machine to Machine (M2M) Evaluation Methodology Document," *IEEE 802.16 Broadband Wireless Access Working Group*, 2011.
- [81] N. Himayat, S. Talwar, K. Johnsson, S. Andreev, O. Galinina, and A. Turlikov, "IEEE 802.16p performance requirements for network entry by large number of devices," *IEEE 802.16p Working Group*, 2010.
- [82] H. Wu, C. Zhu, R. J. La, X. Liu, and Y. Zhang, "FASA: Accelerated S-ALOHA using access history for event-driven M2M communications," *Networking, IEEE/ACM Transactions on*, vol. PP, no. 99, pp. 1–1, 2013.
- [83] Y. Liu, C. Yuen, J. Chen, and X. Cao, "A scalable Hybrid MAC protocol for massive M2M networks," in *Wireless Communications and Networking Conference (WCNC), 2013 IEEE*, 2013, pp. 250–255.
- [84] J. B. Andersen, J. O. Nielsen, G. F. Pedersen, G. Bauch, and G. Dietl, "Doppler spectrum from moving scatterers in a random environment," *Wireless Communications, IEEE Transactions on*, vol. 8, no. 6, pp. 3270–3277, 2009.
- [85] R. Khan and k. Mahata, "An adaptive RRM scheme for Smart Grid M2M applications over a WiMAX network," in *Communication Systems, Networks, and Digital Signal Processing (CSNDSP), 2014 IEEE/IET International Symposium on*, June 2014, pp. 994–999.
- [86] M. Hyder, R. Khan, and K. Mahata, "An enhanced random access mechanism for Smart Grid M2M communications in WiMAX networks," in *Smart Grid Communi-*

- cations (SmartGridComm), 2014 IEEE International Conference on*, June 2014, pp. 994–999.
- [87] R. H Khan, K. Mahata, and J. Y Khan, “A Novel Scheduling Service for Distribution Pilot Protection Over a WiMAX Network,” *Recent Patents on Telecommunication*, vol. 2, no. 1, pp. 26–40, 2013.
- [88] D. G. Jeong and M.-J. Kim, “Effects of channel estimation error in MC-CDMA/TDD systems,” in *Vehicular Technology Conference (VTC), 2000 IEEE 51st*, vol. 3, 2000, pp. 1773–1777.
- [89] K. Fazel and S. Kaiser, *Multi-carrier and spread spectrum systems: from OFDM and MC-CDMA to LTE and WiMAX*. Wiley. com, 2008.
- [90] R. Baxley, B. Walkenhorst, and G. Acosta-Marum, “Complex Gaussian ratio distribution with applications for error rate calculation in fading channels with imperfect CSI,” in *Global Telecommunications Conference (GLOBECOM 2010), 2010 IEEE*, Dec 2010, pp. 1–5.
- [91] R. H. Khan and J. Y. Khan, “A heterogeneous WiMAX-WLAN network for AMI communications in the Smart Grid,” in *Smart Grid Communications (SmartGridComm), 2012 IEEE Third International Conference on*. IEEE, 2012, pp. 710–715.
- [92] “IEEE standard for local and metropolitan area networks - Part 16: air interface for broadband wireless access systems - Amendment 1: multihop relay specification,” *IEEE Std 802.16j-2009 (Amendment to IEEE Std 802.16-2009)*, pp. 1–290, June 2009.
- [93] D. Niyato and E. Hossain, “Integration of IEEE 802.11 WLANs with IEEE 802.16-based multihop infrastructure mesh/relay networks: A game-theoretic approach to radio resource management,” *Network, IEEE*, vol. 21, no. 3, pp. 6–14, 2007.
- [94] L. Berlemann, C. Hoymann, G. Hiertz, and B. Walke, “Unlicensed operation of IEEE 802.16: Coexistence with 802.11 (a) in shared frequency bands,” in *Personal*,

- Indoor and Mobile Radio Communications, 2006 IEEE 17th International Symposium on.* IEEE, 2006, pp. 1–5.
- [95] L. Berlemann, C. Hoymann, G. R. Hiertz, and S. Mangold, “Coexistence and interworking of IEEE 802.16 and IEEE 802.11 (e),” in *Vehicular Technology Conference, 2006. VTC 2006-Spring. IEEE 63rd*, vol. 1. IEEE, 2006, pp. 27–31.
- [96] H.-T. Lin, Y.-Y. Lin, W.-R. Chang, and R.-S. Cheng, “An integrated WiMAX/WiFi architecture with QoS consistency over broadband wireless networks,” in *Consumer Communications and Networking Conference, 2009. CCNC 2009. 6th IEEE.* IEEE, 2009, pp. 1–7.
- [97] I. Corporation, “Intel® WiMAX/WLAN Link 5350 and Intel® WiMAX/WLAN Link 5150.”
- [98] “Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications - Amendment 8: Medium Access Control (MAC) Quality of Service Enhancements,” *IEEE Std 802.11e-2005 (Amendment to IEEE Std 802.11, 1999 Edition (Reaff 2003))*, pp. 1–212, Nov 2005.
- [99] T. Sakurai and H. L. Vu, “MAC access delay of IEEE 802.11 DCF,” *Wireless Communications, IEEE Transactions on*, vol. 6, no. 5, pp. 1702–1710, 2007.
- [100] “Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications: High-Speed Physical Layer in the 5 GHz Band,” *IEEE Std 802.11a-1999*, pp. i–, 1999.
- [101] M. M. Hyder, R. H. Khan, K. Mahata, and J. Y. Khan, “A predictive protection scheme based on adaptive synchrophasor communications,” in *Smart Grid Communications (SmartGridComm), 2013 IEEE International Conference on.* IEEE, 2013, pp. 762–767.
- [102] N. Leemput, J. Van Roy, F. Geth, P. Tant, and J. Driesen, “Comparative analysis of coordination strategies for electric vehicles,” *status: published*, 2011.

- [103] S. Deilami, A. S. Masoum, P. S. Moses, and M. A. Masoum, "Real-time coordination of plug-in electric vehicle charging in smart grids to minimize power losses and improve voltage profile," *Smart Grid, IEEE Transactions on*, vol. 2, no. 3, pp. 456–467, 2011.
- [104] "Use of synchrophasor measurements in protective relaying," IEEE Power System Relaying Committee (PSRC), Working Group C14, Tech. Rep. Draft 6.4, Aug. 2012.
- [105] B. Milosevic and M. Begovic, "Voltage-stability protection and control using a wide-area network of phasor measurements," *Power Systems, IEEE Transactions on*, vol. 18, no. 1, pp. 121–127, Feb. 2003.
- [106] Y. Gong, N. Schulz, and A. Guzman, "Synchrophasor-based real-time voltage stability index," in *Power Systems Conference and Exposition, 2006. PSCE '06. 2006 IEEE PES*, 29 2006–Nov. 1, pp. 1029–1036.
- [107] W. Price, C. Taylor, and G. Rogers, "Standard load models for power flow and dynamic performance simulation," *IEEE Transactions on power systems*, vol. 10, no. CONF-940702–, 1995.
- [108] "WiMAX Network Architecture - Stage 2: Architecture Tenets, Reference Model and Reference Points - Base Specification," *WiMAX Forum*, March 2013.
- [109] S. Chakrabarti and E. Kyriakides, "Optimal placement of phasor measurement units for power system observability," *Power Systems, IEEE Transactions on*, vol. 23, no. 3, pp. 1433–1440, Aug. 2008.
- [110] T. S. Ustun, C. Ozansoy, and A. Zayegh, "A microgrid protection system with central protection unit and extensive communication," in *Environment and Electrical Engineering (EEEIC), 2011 10th International Conference on*. IEEE, 2011, pp. 1–4.
- [111] T. S. Ustun, C. Ozansoy, and Zayegh, "Modeling of a centralized microgrid protection system and distributed energy resources according to IEC 61850-7-420," *Power Systems, IEEE Transactions on*, vol. 27, no. 3, pp. 1560–1567, 2012.

- [112] R. G. Gallager, *Principles of digital communication*. Cambridge University Press Cambridge, UK:, 2008, vol. 1.
- [113] G. Seber and C. Wild, “Nonlinear regression. 2003.”
- [114] N. Niglye, F. S. Peritore, R. D. Soper, C. Anderson, R. Moxley, and A. Guzmán, “Considerations for the application of synchrophasors to predict voltage instability,” in *Power Systems Conference: Advanced Metering, Protection, Control, Communication, and Distributed Resources, 2006. PS’06*. IEEE, 2006, pp. 169–172.
- [115] R. Sodhi, S. Srivastava, and S. Singh, “A simple scheme for wide area detection of impending voltage instability,” *Smart Grid, IEEE Transactions on*, vol. 3, no. 2, pp. 818–827, June, 2012.