



End to end architecture and mechanisms for mobile and wireless communications in the Internet

Lei Zhang

► To cite this version:

Lei Zhang. End to end architecture and mechanisms for mobile and wireless communications in the Internet. Computer Science [cs]. Institut National Polytechnique de Toulouse - INPT, 2009. English. NNT: . tel-00435868v1

HAL Id: tel-00435868

<https://theses.hal.science/tel-00435868v1>

Submitted on 25 Nov 2009 (v1), last revised 9 Jan 2024 (v2)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



THÈSE

En vue de l'obtention du

DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par *Institut National Polytechnique de Toulouse (INPT)*

Discipline ou spécialité : *Informatique, Réseaux*

Présentée et soutenue par *ZHANG Lei*
Le 5 Octobre 2009

Titre : *End to end architecture and mechanisms for mobile and wireless communications in the Internet*

JURY

Andrzej DUDA, Professeur, Ensimag

Fethi FILALI, Associate Professor, EURECOM

André-Luc BEYLOT, Professeur, Enseeiht

Prosper CHEMAUIL, Directeur du programme de recherche, réseaux et système, Orange Labs

Francine Krief, Professeur, ENSEIRB

Michel DIAZ, Directeur de recherche, LAAS-CNRS

Patrick SENAC, Professeur, ISAE

Ecole doctorale : *Mathématiques, Informatique, Télécommunications de Toulouse (M.I.T.T.)*

Unité de recherche : *Laboratoire d'Analyse et d'Architecture des Systèmes (LAAS-CNRS)*

Directeur(s) de Thèse : *M. DIAZ et P. SENAC*

Rapporteurs : *A. DUDA et F. FILALI*



Université
de Toulouse

Thèse

Année 2009

Préparée au Laboratoire d'Analyse et d'Architecture des Systèmes du CNRS

**En vue de l'obtention du Doctorat de l'Institut National Polytechnique de
Toulouse, Université de Toulouse**

**Ecole doctorale : Mathématiques, Informatique, Télécommunications de
Toulouse (M.I.T.T.)**

Par **ZHANG Lei**

**End to end architecture and mechanisms for mobile
and wireless communications in the Internet**

Soutenue le 5 Octobre 2009

Acknowledgment

In the first place, I would like to record my gratitude to Patrick Sénac for his 3 years' supervision in the context of my PhD study, as well as giving me extraordinary opportunities to participate several projects throughout the research work. This thesis would not have been possible without his support and guidance. His truly scientist intuition has made him as a constant oasis of ideas and passions in science, which led me to avoid redundant tortuousness, inspired me to find out innovative solutions to different research issues, and exceptionally enrich my growth as a researcher wants to be. Thanks to him, research life became smooth and rewarding for me. Apart from his endless helps on my research work, I am so grateful that he gave me many precious suggestions on the philosophy of life, the art of expression, etc. He is definitely one of the most important persons to me in my research career and I am indebted to him more than he knows.

I would like to specially thank Michel Diaz, for his supervision and advices during my PhD study. His insights and special view on the research has triggered and nourished my intellectual maturity. I benefited from his valuable advices after each science discussion with him. His kindness and his smiles will be engraved in my memory forever.

I gratefully acknowledge Emmanuel Lochin who gave me endless help on my research issues, working together on several experimental and demonstration testbeds, using his precious time to read my previously submitted papers and giving his critical comments. I am grateful for his involvement with his originality to my research framework and I hope to keep up the collaboration in the future.

I would like to show my gratitude to Aruna Seneviratne and Roksana Boreli for giving me a valuable opportunity to work in National ICT center of Australia for 3 months and resulted in fruitful cooperation results. Thank all NICTA's colleagues for their welcome and help during my staying in Australia.

Of course, I will never forget all my colleagues in the department DMIA of ISAE, who made it a convivial place to work. In particular, I would like to thank Pierre Ugo Tournoux for his friendship and help, thank Jérôme Lacan and Tanguy Pérennou for their help and cooperation work in some research issues.

Last but definitely not least, my deepest gratitude goes to my parents for their unflagging love and support throughout my life; this dissertation is simply impossible without them. I have no suitable word that can fully describe the everlasting love from my mother, Tingting Zhang; I remember her constant support when I encountered difficulties. And I am indebted to my father, Junyu Zhang, for his care and love. He spares no effort to provide the best possible environment for me to grow up.

Contents

1	Introduction	5
1.1	Access technologies to Internet	5
1.2	Limits of the wireless network evolution	6
1.2.1	Limits of the transport layer	6
1.2.2	Limits of the network layer	6
1.3	The cross layer conception	7
1.4	The E2E (End to End) conception	7
1.5	Organization of the thesis	7
1.5.1	Research issues	7
1.5.2	Thesis organization	8
2	Background and state of the art	9
2.1	Wireless networks	9
2.1.1	WPAN	10
2.1.2	WMAN	10
2.1.3	Cellular networks	11
2.2	Mobility Management	12
2.2.1	Overview of Mobility Management	12
2.2.2	Network layer solutions	15
2.2.3	Transport layer solutions	20
2.2.4	Application layer solutions	22
2.2.5	Cross layer solutions	23
2.2.6	Macro mobility solution	24
2.2.7	Micro mobility solution	25
2.2.8	802.21 based mobility management	25
2.2.9	Open issues of mobility management	26
2.2.10	Conclusion of section 2.2	27
2.3	WLAN Issues	27
2.3.1	Background of 802.11	27
2.3.2	Open issues and related work for 802.11	31
2.3.3	Conclusion of section 2.3	33
2.4	Conclusion of chapter 2	33
3	End to end architecture for mobility management	35
3.1	Introduction	35
3.2	Mobility management architecture	35
3.3	Mobile connection analytical model	36
3.4	Location/Address Management	40
3.4.1	Dynamical DNS	40
3.4.2	HIP Rendezvous server (RVS) mechanism	41
3.5	Continuous connection support	45
3.6	Handover Management	49
3.6.1	Signal analysis function	49
3.6.2	Mobility prediction	53
3.6.3	The Handover Procedure	53

3.7	Conclusion of chapter 3	64
4	WLAN performance enhancements	65
4.1	Three main issues of 802.11	66
4.1.1	Performance Anomaly of 802.11	66
4.1.2	MAC layer overflow	67
4.1.3	Fairness issues	71
4.2	MAC Layer rate modeling	72
4.2.1	Fundamental modeling parameters	72
4.2.2	Analytical model of calculating maximum bandwidth delivered by 802.11 MAC layer	73
4.3	Rate adaptation	79
4.3.1	IP layer shaper	80
4.3.2	Transport rate limiter	81
4.3.3	Simulation and validation of rate adaptation mechanisms	81
4.3.4	Conclusion of the section 4.3	85
4.4	Rate equalization mechanisms	85
4.4.1	Unfairness issue between UL and DL flows	85
4.4.2	Unfairness issue between ACK clocked and Non-ACK clocked flows	86
4.4.3	Simulation and validation of rate-equalization mechanisms	87
4.4.4	Conclusion of the section 4.4	89
4.5	Adaptive FEC based mechanism	89
4.5.1	Motivation of this work	90
4.5.2	Cross layer based erasure code	91
4.5.3	Defining the critical redundancy ratio	93
4.5.4	Validation of our proposal	96
4.5.5	Conclusion of section 4.5	98
4.6	Conclusion of chapter 4	99
5	Enhancements of WiMax modulation scheme efficiency	101
5.1	Motivation	101
5.2	FEC-based modulation adaptation	102
5.2.1	Conception view of our proposal	102
5.2.2	Mapping function	103
5.2.3	Modulation and coding scheme decision algorithm	104
5.2.4	Global goodput gain estimation	106
5.3	Simulation results	108
5.3.1	Scenario I:Uniform movement	109
5.3.2	Scenario II:Stop and start movement	111
5.4	Conclusion of chapter 5	111
6	Implementation of an end to end architecture for improving mobility management and wireless communications	113
6.1	Seamless streaming module	114
6.1.1	Java Media Framework (JMF)	114
6.1.2	Extension work to JMF	114
6.1.3	Extension work to VLC	115
6.2	Geo-Location module	115
6.3	Energy Control module	117
6.4	MAC layer Optimization module	117
6.5	3MOI Interface	118
6.6	Conclusion of chapter 6	120

7	Conclusion and perspective	121
7.1	Summary	121
7.2	Perspective	122
8	Appendix	125
8.1	Estimate of t'_{jam} in UDP mode	125
8.2	Estimate of t''_{jam} in the TFRC greedy mode	125
8.3	Estimate of t_{jam} in the rate sparing mode for the non-frequent feedback based protocols	125
8.4	Estimation of t_{jam} in the TCP/UDP coexisting case	126
9	List of publications	127

Chapter I

1 Introduction

Wireless network technologies have been rapidly developed to meet the increasing demand for mobile services with higher data rate and enhanced multimedia applications. Two approaches have been addressed by the current research. The first approach focuses on the mobility management which allows assigning and controlling the wireless connections and offers enhanced wireless services to end systems. Location and handover management are two core issues in the context of mobility management. The other approach focuses on the wireless communication techniques (i.e. wireless access methods, physical feature competence, etc.) which enhance the lower layer performances such as hyper-speed transmission rate, interference avoidance, transmission efficiency improvement. Keeping in mind the above two approaches, two main contributions have been addressed in the frame of the thesis: 1) End to end architecture for the mobility management to support seamless communications over the Internet. 2) MAC-layer performance enhancements to improve the transmission efficiency in the context of WLAN and WiMax.

This chapter introduces several issues related to our thesis: the access networks to Internet, the limits of the current wireless network evolution as well as the “Cross layer” and “End to end” conceptions, and the organization of the thesis will be shown finally.

1.1 Access technologies to Internet

The Internet, which was developed in the wired epoch, ignores the features of wireless compatibility. In the frame of this thesis, one of the main addressed contribution is, based on our mobility management architecture, how to support a continuous and seamless commutation over the Internet.

Internet was originally designed in the context of wired network, 802.3 remains the basic and most popular access method, it offers large band (from Mega to Giga bps) to support multiple Internet applications. Nowadays, the increasing use of portable devices in daily life and the dramatic and continuous growth of the “mobile and wireless Internet” greatly increase the needs of offering higher capacity, higher reliability, and more advanced multimedia services to wireless mobile users, the emergence of WiFi technology known as 802.11x is pushing a new revolution on personal wireless communication. It makes possible to deliver pervasive low cost broadband Internet access to mobile users. Specially, the next generation 802.11n wireless network is being standardized by adding several powerful technologies like MIMO [5] and 40MHz operation [6] to the physical layer which allows supporting much higher transmission rates and wider signal coverage. Another technology, originally designed for the next generation mobile broadband networks is WiMax, Worldwide Interoperability for Microwave Access [7] defined in the context of IEEE 802.16 standard, it is designed to deliver wireless broadband bitrates, with Quality of Service (QoS) guarantees for different traffic classes, robust security, and mobility. Currently, the extended version of WiMax, the IEEE standardization of 802.16m (also called

WiMax II) is in progress. This emerging standard will incorporate the most promising communication techniques to offer 100 Mbps data rates to mobile systems and 1 Gbps to fixed end systems. Meanwhile, one of the main competitors: 3G as well as its evolution LTE [8] offer a flexible way for the mobile systems to access Internet. Compared to WiMax which requires the significant changes to the current communication system, LTE inherits the 3G architecture and supports a seamless transition from 3G to the next wireless communication generation, this explains the success and the rapid deployment of LTE.

1.2 Limits of the wireless network evolution

Besides the transmission competences that rely on the wireless communication techniques, the bottleneck which constrains the development of wireless network is the lack of efficient architectures and management mechanisms especially at transport and network layers.

1.2.1 Limits of the transport layer

The rapid growth of Internet is forcing an evolution of one of its most significant technology enablers, the TCP/IP protocol suite. TCP has proved to be remarkably resilient over the years, but as the capacity of the network increases, the ability of TCP to continue to deliver higher data rates over wireless networks that spans the globe is becoming a significant concern. Some of the recent work focuses on the revised TCP flow-control algorithms that still share the network fairly with other concurrent TCP sessions, yet can ramp up to multi gigabit-per-second data-transfer rates and sustain those rates over extended periods [1]. The latest TCP stack from Microsoft in Vista uses dynamic tuning of the Receive window, and a larger inflation factor of the send window in congestion avoidance where there is a large bandwidth delay product, and improved loss-recovery algorithms that are particularly useful in wireless environments. Linux now includes an implementation of Binary Increase Congestion control (BIC), which undertakes a binomial search to reestablish a sustainable send rate. Both of these approaches can improve the performance of TCP, particularly when sending the TCP session over long distances and trying to maintain high transfer speeds. However, TCP today as deployed on the Internet is much the same as TCP of a decade ago, with only a couple of notable exceptions. How to make the TCP compatible with wireless network is a big challenge in the current research and remains unsolved.

Otherwise, two new transport protocols, the Datagram Congestion Control Protocol (DCCP) [2] and the Stream Control Transmission Protocol (SCTP) [3] has been developed, which can be regarded as refinements of TCP to cover flow control for datagram streams in the case of DCCP and flow control over multiple reliable streams in the case of SCTP. However, in a world of transport-aware middleware that is the Internet today, the level of capability to actually deploy these new protocols in the public Internet is marginal at most.

1.2.2 Limits of the network layer

Network layer provides globally-usable addresses and enables IP routing. To summarize the current wireless network at network layer, a mobile node connecting to a foreign network with the purpose of acting as a client, accessing services on the Internet will requires local support of a DHCP service. When requiring full access to the home network, a virtual private network (VPN) [9] can be used. To manage the combination of moving nodes and reachability from other nodes, the Mobile IP (MIP) is proposed [10]. MIP solves the problem with the dual meaning of the IP address as an endpoint identifier and

a location identifier. However, it brings several drawbacks such as triangle routes, lack of notification mechanisms before the mobility disconnection, etc. Current research of the network layer is still focusing on the location management and the optimization of routing.

1.3 The cross layer conception

The traditional way of designing wireless network architecture, was to identify each module/layer and then assign them separately their roles or requirements. Since each layer being treated individually, this approach has in many ways caused the research communities to split into different groups, where each group focus their resources on solving "their" own problem in the best possible way. This approach has been proved to have several drawbacks due to the lack of interaction information. The idea behind the cross layer design is to combine the resources available in different layers, and create a network architecture which can be highly adaptive and QoS-efficient by sharing state information between the lower and upper layers. In the context of the wireless network, in order to fully optimize wireless transmission, both the challenges from the physical medium and the QoS-demands from the high layers have to be taken into account. Rate, power and coding at the physical layer can be adapted to meet the requirements of the applications given the current channel and network conditions. Besides, lower layers information can also be considered as a hint for mobility awareness and to predict the mobility behaviors, which play an important role in handover issues to support seamless communications. Our proposed architecture for the mobility management is based on the cross layer and allows information sharing and transparency across the different layers, which is proved as an efficient approach to satisfy the mobility requirements.

1.4 The E2E (End to End) conception

The end-to-end concept and network intelligence at the end-hosts reaches back to packet-switching networks in the 1970's, cf. CYCLADES. A presentation entitled End-to-end arguments in system design [11] argued that reliable systems tend to require end-to-end processing to operate correctly, in addition to any processing in the intermediate system. Nowadays, the end-to-end principle remains one of the central design principles of the Internet and is implemented in the design of the underlying methods and protocols in the Internet Protocol Suite. It is also used in other distributed systems. The principle states that, whenever possible, communications protocol operations should be defined to occur at the end-points of a communications system, or as close as possible to the resource being controlled. The main advantage of the E2E, which is also adopted by our mobility management architecture, is that such an approach allows integrating efficient and intelligent mechanisms at the end systems and doesn't require the modifications to the intermediate system, which make the deployment easier and much more flexible.

1.5 Organization of the thesis

This section gives a brief description on our research issues and provides a roadmap in the context of the thesis.

1.5.1 Research issues

This thesis addresses two main contributions.

Firstly, we focus on our end-to-end architecture for the efficient mobility management; our architecture is based on cross layer design and regroups several existing and novel

mechanisms to support seamless communications over the Internet. Location management and Handover management are two main issues addressed in this part. Our validation middleware which is implemented in Java demonstrates the efficiency of the proposed architecture.

The second contribution of our thesis focuses on the WLAN and WiMax MAC-layer performance enhancements. We identified three intrinsic problems of WLAN that are related to the CSMA/CA access method, namely 1) Performance degradation due to the lack of flow control between the MAC and upper layer resulting in potential MAC buffer overflow; 2) Unfair bandwidth share issues; 3) 802.11 Performance Anomaly. Different from the traditional mechanisms at lower layers, our cross layer based solutions don't require changes to the 802.11 standard and become easier to be deployed. Meanwhile, we extended some of our proposals in the context of WiMax to optimize the transmission efficiency.

The thesis work has resulted in 8 publications which is listed at the end of this thesis report.

1.5.2 Thesis organization

The thesis consists in 7 chapters. Chapter 2 focuses on the state of the art that gives a detailed description on the background and the related work of our research in the context of the thesis. Chapter 3 shows our end to end architecture for the mobility management. Chapter 4 and 5 show our proposals for the enhancements of the WLAN and WiMax MAC layer performances. Chapter 6 shows the implementation of our end to end architecture for improving mobility management and wireless communications. Conclusion and perspective work will be discussed in Chapter 7.

Chapter II

2 Background and state of the art

Wireless networks have had a significant impact on our daily life. One of the most challenging issues focuses on the management of the mobility behaviors over the Internet. Nowadays, several Internet architectures as well as a number of novel mechanisms involved at different ISO layers have been proposed to address this problem. However, the proposed solutions always sway on the tradeoff between the mobility management efficiency and the integration of the additional infrastructures which are costly and break the Internet protocol stack.

Moreover, while the core networks are growing much more mature and robust, the wireless access networks that are also called “last mile” still remain as a bottleneck for bandwidth, end-to-end delay, stability and general compatibility. Therefore, how to efficiently manage the wireless resource and enhance the low layer performance is becoming a challenging topic.

In this section, we will give an introduction to the background and the related work of the two main issues related to our thesis, the mobility management and wireless access networks performance enhancements. We will define solutions for these issues in the next chapters by proposing protocols and mechanisms that push as far as possible the end to end approach.

2.1 Wireless networks

The vision of wireless communications supporting information exchange between mobile users or devices is the communications frontier. This vision will allow multimedia communication from anywhere in the world using a small handheld device or laptop. Nowadays, cellular systems have experienced exponential growth over the last decade, cellular phones have become a critical business tool and part of daily life, and they are rapidly supplanting antiquated wire-line systems across the world. In addition, many new applications including wireless sensor networks, automated highways and factories, smart homes and appliances, and remote telemedicine are emerging from research ideas to concrete systems. The explosive growth of wireless systems coupled with the proliferation of laptop and palmtop computers suggests a bright future for wireless networks, both as stand-alone systems and as part of the larger networking infrastructure.

Basically, wireless networks can be categorized into 4 types according to their coverage areas, working functions as well as their customer features: WPAN (Wireless Personal Area Network), WLAN (Wireless Local Area Network), WMAN (Wireless Metropolitan area networks), and Cellular networks. The different physical layer properties decide their different application domains. Details of the WLAN background and its related work will be discussed in section 2.3.

2.1.1 WPAN

Wireless Personal Area Network (WPAN) is a type of wireless network that interconnects devices within a relatively small area. Bluetooth [16] is a representative technology which allows interconnecting devices within reach of a person. Nowadays, another typical WPAN technology, ZigBee [17] becomes a hot topic which targets at radio-frequency applications that require a low data rate, long battery life, and secure networking. The low cost allows this technology to be widely deployed in wireless control and monitoring field and the mesh networking provides high reliability and larger range. Home Automation is one of its typical applications.

2.1.2 WMAN

Wireless Metropolitan Area Networks (WMAN) offer wireless connections which are optimized for a larger geographical area ranging from several blocks of buildings to entire cities. The typical technology of WMAN in IEEE context is called Worldwide Interoperability for Microwave Access (WiMax)[20]. WiMax is a wireless solution that offers broadband data transmission service over metropolitan areas. It was originally designed as a wireless alternative to cable and DSL for "last mile" broadband access. The current WiMAX incarnation, Mobile WiMAX that provides mobility connectivity, is based upon IEEE Standard 802.16e-2005, approved in December 2005. It is a supplement to the IEEE Standard 802.16-2004 that addresses only fixed systems. Some of the WiMax salient features that deserve highlighting are as following:

OFDM-based physical layer: The WiMAX physical layer is based on orthogonal frequency division multiplexing, a scheme that supplies good resistance to multipath, and allows WiMAX to operate in NLOS conditions. OFDM [21] is now widely considered as the method of choice for mitigating multipath for broadband wireless.

Very high peak data rates: WiMAX supports very high peak data rates which can be as high as 74Mbps on a 20MHz wide spectrum. More typically, using a 10MHz spectrum operating using TDD scheme with a 3:1 downlink-to-uplink ratio, the peak PHY data rate reaches 25Mbps and 6.7Mbps for the downlink and the uplink, respectively. These peak rates are achieved when using 64 QAM modulation under good signal conditions, and higher peak rates may be achieved using multiple antennas and spatial multiplexing.

Adaptive modulation and coding (AMC): WiMAX supports different modulation and forward error correction (FEC) coding schemes [22] and allows the scheme to be changed on a per user and per frame basis, based on channel conditions. AMC is an effective mechanism to maximize throughput in dynamic link channel. The adaptation algorithm allows the use of the highest modulation and coding scheme that can be supported by the signal-to-noise and interference ratio at the receiver side; therefore each user is provided with the highest possible data rate that can be supported in their respective links.

Support for TDD and FDD: IEEE 802.16-2004 and IEEE 802.16e-2005 supports both TDD time division duplexing and FDD frequency division duplexing. TDD is favored by a majority of implementations because of its advantages: (1) flexibility in choosing uplink-to-downlink data rate ratios, (2) ability to exploit channel reciprocity, (3) ability to implement in non-paired spectrum. (4) less complex transceiver design. All the initial WiMAX profiles are based on TDD, except for two fixed WiMAX profiles in 3.5GHz.

QoS support: The WiMAX MAC layer has a connection-oriented architecture that is designed to support a variety of applications, including voice and multimedia services. The system offers the support for constant bit rate, variable bit rate, real-time, and non-real-time traffic flows as well as best-effort data traffic. WiMAX MAC is designed to support a large number of users, with multiple connections per terminal, each with its own QoS requirement.

Robust security: WiMAX supports strong encryption, using Advanced Encryption Standard (AES) [23], and has a robust privacy and key-management protocol. The system also offers a very flexible authentication architecture based on Extensible Authentication Protocol (EAP) [24], which allows for a variety of user credentials, including username/password, digital certificates, and smart cards.

Support for mobility: The mobile WiMAX supports secure seamless handovers for delay-tolerant full-mobility applications, such as VoIP. The system also supports power-saving mechanisms that extend the battery life of handheld subscriber devices. Physical-layer enhancements, such as more frequent channel estimation, uplink sub-channelization, and power control, are also specified in support of mobile applications.

IP-based architecture: The WiMAX Forum has defined a reference network architecture that is based on an all-IP platform. All end-to-end services are delivered over an IP architecture relying on IP-based protocols for end-to-end transport, QoS, session management, security, and mobility. Reliance on IP allows WiMAX to ride the declining cost-curves of IP processing, facilitate easy convergence with other networks, and exploit the rich ecosystem for application development that exists for IP.

2.1.3 Cellular networks

Cellular telephone networks, have experienced explosive growth in the past two decades. Today millions of people around the world use cellular phones. Cellular phones allow a person to make or receive a call from almost anywhere. Likewise, a person is allowed to continue the phone conversation while on the move. Cellular communications is supported by an infrastructure called a cellular network, which integrates cellular phones into the public switched telephone network. The cellular telephone network has gone through three generations. The first generation of cellular networks is analog in nature. To accommodate more cellular phone subscribers, digital TDMA (time division multiple access) and CDMA (code division multiple access) technologies are used in the second generation (2G) to increase the network capacity. With digital technologies, digitized voice can be coded and encrypted. Therefore, the 2G cellular network is also more secure. General Packet Radio Service (GPRS) is a packet oriented Mobile Data Service available to users of the 2G cellular communication systems; It's often described as "2.5G", that is, a technology between the second and third generations of mobile telephony. A GPRS connection is established by reference to its Access Point Name (APN). The APN delivers the services such as Wireless Application Protocol (WAP) access, Short Message Service (SMS), Multimedia Messaging Service (MMS), and for Internet communication services such as email and World Wide Web access.

The third generation (3G) integrates cellular phones into the Internet world by providing high speed packet-switching data transmission in addition to circuit-switching voice transmission. IMT-2000 is a term that represents several incompatible 3G standards lumped together under one banner. The hope of IMT-2000 is that phones using these different standards will be able to move seamlessly between all networks, thus providing global roaming. WCDMA is one of the 3G cell phone technologies, which is also being developed into a 4G technology. Currently, the most common form of UMTS uses WCDMA as the underlying air interface. WCDMA is considered as a wideband spread-spectrum mobile air interface that utilizes the direct-sequence spread spectrum method of asynchronous code division multiple-access to achieve higher speeds and support more users. In the other hand, CDMA2000 is a hybrid 2.5G / 3G technology of mobile telecommunications standards that use CDMA, a multiple access scheme for digital radio, to send voice, data, and signaling data (such as a dialed telephone number) between mobile phones and cell sites. CDMA2000 is considered a 2.5G technology in 1xRTT and a 3G technology

in EVDO. The CDMA2000 standards (CDMA2000 1xRTT, CDMA2000 EV-DO, and CDMA2000 EV-DV) are approved radio interfaces for the ITU's IMT-2000 standard and a direct successor to 2G-CDMA, IS-95 (cdmaOne).

4G (also known as Beyond 3G), an abbreviation for Fourth-Generation, is a term used to describe the next complete evolution in wireless communications. A 4G system will aim to provide a comprehensive IP solution where voice, data and streamed multimedia can be given to users on an "Anytime, Anywhere" basis, and at higher data rates than previous generations. Nowadays, there is no formal definition for what 4G is; however, there are certain objectives that are projected for 4G. These objectives include: that 4G will be a fully IP-based integrated system. 4G will be capable of providing between 100 Mbit/s and 1 Gbit/s speeds both indoors and outdoors with premium quality and high security. Currently, the technologies considered to be "pre-4G" include WiMax, WiBro [25], iBurst [26], and 3GPP LTE (Long Term Evolution) [27]. One of the first technology really fulfilling the 4G requirements as set by the ITU-R will be LTE which is currently being standardized by 3GPP. LTE will be an evolution of the 3GPP Long Term Evolution. Higher data rates are for instance achieved by the aggregation of multiple LTE carriers that are currently limited to 20MHz bandwidth.

2.2 Mobility Management

With the development of multimedia communication and personal mobile devices, an ubiquitous broadband network supporting mobility is the goal of future wireless network. Different complementary wireless technologies are able to cooperate together to empower mobile users to use the best available access network that suits their requirements. The integration of different networks entails several research challenges due to the following heterogeneities:

- Different wireless access technologies: Wireless communication systems will include many heterogeneous networks using different radio technologies. These networks may have overlapping coverage areas and different cell sizes, ranging from a few square meters to hundreds of square kilometers.
- Different network architecture and mechanisms: Wireless communication systems have different network architectures and protocols for application, transport flow congestion and error control, routing, mobility management, etc
- Different services requirements: Different mobile users and application services require different services ranging from low data rate and non real-time applications to high-speed real-time multimedia applications offered by various access networks.

The above intrinsic technology heterogeneities require a common infrastructure to interconnect multiple access networks and guarantee the QoS of application services. For interoperation of different communication protocols, an adaptive protocol suite is required that will adapt itself to the characteristics of the underlying networks and provide optimal performance across a variety of wireless environments. Furthermore, adaptive terminals in conjunction with "smart" base stations will support multiple air interfaces and allow users to seamlessly switch among different access technologies. One important component of the adaptive protocol suite is the integration of mobility management schemes.

2.2.1 Overview of Mobility Management

This section firstly gives a brief introduction on the classification of mobility management in terms of different architecture, and then we will show several mobility management requirements as well as the related techniques.

2.2.1.1 Fundamental functions and classification of mobility management schemes

In the current wireless systems, mobility management contains two main components: location management and handover management. The former concerns how to locate a mobile node, track its movement, and update the location information, while the latter focuses mostly on the control of the change of a mobile node's access point during active data transmission.

Location management enables the system to track the locations of mobile nodes in real time. It includes two major tasks, which are location registration and location update, to enable the network to discover the current attachment point of the mobile user. The mobile node is required to update its logical address in the location server once address is changed due to the migration across different subnets. The object of the location management is to make the mobile node's location always be transparent to other nodes across the Internet.

Handover management is the process by which a mobile node keeps its connection active when it moves from one access point to another. The handoff process can be intra or inter-system. Intra-system handoff is the handoff in homogeneous networks. Inter-system handover refers to the handover between different technologies and network domains. Handoff management research concerns issues such as: efficient and expedient packet processing; minimizing the signaling load on the network; optimizing the route for each connection; efficient bandwidth reassignment; evaluating existing methods for standardization; and refining quality of service for wireless connections. Handoff process is normally divided into three stages namely: BS Discovery; Handoff Decision; Handoff Execution.

BS Discovery: When a station decides to perform a handoff, it needs to find candidate access points. In order to discover the possible base stations (BS) to which it may switch, a mobile node performs a link layer procedure called "scan". In 802.11, there are two methods of scan, passive and active. In active scanning, the mobile node listens for beacon frames issued by the BS at regular intervals. The beacon includes information such as SSID, supported rates and security parameters. The mobile node can also obtain the same information by using active scanning. Active scanning (or probing) consists on issuing probe request frames to which BS will respond with probe response frames including information similar to the information included in beacon frames.

Handoff Decision: In the Handoff Decision stage, it determines which network can be used to connect to, based on different parameters such as quality of signal, access network capabilities, user preferences & policy, handoff signaling delay, jitter, delay, packet loss rate and maximum bit rate. However, most of the current proposals choose the quality of signal as the main critical condition to choose the best BS due to the efficiency and simplicity.

Handoff Execution: In the Handoff Execution stage, lower layer and upper layer handover are respectively executed to establish connection to the new subnet. This execution stage involves mechanisms based on several layers such as the MAC, network and transport layers. Efficient cross layer interactions that aims to reduce handover latency and deliver seamless communications is still an open issue that will be addressed in this thesis.

In the point of view on the classification of the handover management, it can normally be categorized as Soft / Hard handover and Horizontal / Vertical handover.

Soft handover refers: While performing handoff, the terminal may connect to multiple BS's simultaneously and use some form of signaling diversity to combine the multiple signals. Soft handover is able to re-establish the connection before break the previous one. Thus, it is also called "Make-before-break Handover". Contrarily, if the terminal stays connected to only one BS at a time, clearing the connection with the former BS

immediately before establishing a connection with the target BS, then the process is referred to as hard handoff, which is also called “Make-after-break Handover”

In the other hand, vertical handoff usually refers to handover between different types of networks (i.e. WLAN, cellular networks). Vertical handoffs between WLAN and 3G cellular networks have attracted a great deal of attention in all the research areas of the 4G wireless network, due to the benefit of utilizing the higher bandwidth and lower cost of WLAN as well as better mobility support and larger coverage of 3G. Vertical handoffs among a range of wired and wireless access technologies including WiMAX can be achieved using Media independent handover which is standardized as IEEE 802.21. Otherwise, horizontal handoff refers to handover between the same type of network, however, since horizontal handoff doesn’t allow the simultaneous work of multiple different wireless cards on the same mobile terminal, thus, the reduction of handover delay remains a big challenge.

Meanwhile, from an OSI model point of view, mobility management can be classified into the following categories, which will be discussed in the following sections:

- **Network layer solutions**
- **Transport layer solutions**
- **Application layer solutions**
- **Cross-layer based solutions**

These mobility management solutions will be respectively detailed in section **2.2.2-2.2.5**

Finally, in the point of view of the network domain that the mobility decision spans, two kinds of mobility management can be classified as well:

1) Macro mobility management, i.e. mobile node’s movements between different domains, to which inter-domain mobility management schemes can be employed, acting as a global mobility solution with the advantages of flexibility, robustness, and scalability.

2) Micro mobility management, i.e. mobile node’s movements inside a domain, to which intra-domain mobility management solutions are suitable, focusing mainly on a fast, efficient, seamless mobility support within a restricted coverage.

It is worthy to note that the concepts micro and macro mobility based on the definition of domain are possibly recursive: a movement may be micro in one domain whereas macro in another. These two types of mobility management will be respectively detailed in section **2.2.6** and **2.2.7**

2.2.1.2 mobility management requirements and the related techniques

In order to design or select an efficient mobility management scheme, several requirements on performance and scalability should be carefully taken into account:

1) Fast handover: the handover operations should be quick enough in order to ensure that the mobile node can receive data packets at its new location within a reasonable time interval and as a result to reduce the end to end packet delay as much as possible.

2) Seamless handover: the handover algorithm should minimize the packet loss rate to near-zero which, together with fast handover, is sometimes referred to as smooth handover.

3) Signaling traffic overhead: the control data load, e.g. the number of signaling packets or the number of accesses to the related databases, should be lowered to within an acceptable range.

4) Routing efficiency: the routing paths between the communication nodes to the mobile nodes should be optimized to exclude redundant transfer or bypass path as e.g. triangle routing.

5) Quality of Service (QoS): the mobility management scheme should support the establishment of new QoS reservation in order to deliver a variety of traffic, while minimizing the disruptive effect during the establishment.

6) Fast security process: the mobility scheme should support different levels of security requirements such as data encryption and user authentication, while limiting the traffic and time of security process e.g. key exchange.

7) Special support required: it is better for a new mobility mechanism to require minimal special changes on the components, e.g. mobile node, router, communication media, networks, other communication nodes, etc.

Otherwise, there are many distinct but complementary techniques especially for mobility management to achieve the performance and scalability requirements listed above, including:

- Buffering and forwarding, to cache packets by the old attachment point during the MN in handoff procedure, and then forward to the new attachment point after the processing of MN's handoff.
- Movement detection and prediction, to detect and predict the movement of mobile host between different access points so that the future visited network is able to prepare in advance and packets can be delivered to the new location during handoff.
- Handoff control, to adopt different mechanisms for the handoff control, e.g. layer two or layer three triggered handoff, hard or soft handoff, mobile-controlled or network-controlled handoff.
- Paging area, to support continuously reachable with low overhead on location update registration through location registration limited to the paging area.

2.2.2 Network layer solutions

Mobile IP [28] is a typical mobility-enabling protocol for the global Internet. It integrates three new functional entities: home agent (HA), foreign agent (FA), and mobile node/host (MN / MH).

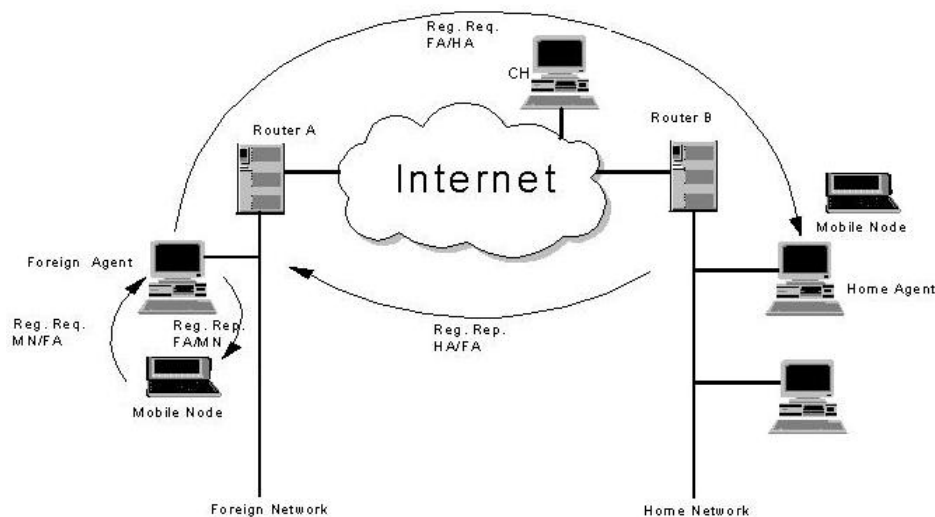


Fig 2.1 Mobile IP scenario

Fig 2.1 represents a typical scenario of Mobile IP, the MN is able to send and receive packets to/from a CH (corresponding host) when the MN is in the home network. However if the MN is away from its home network, it needs an agent to work on behalf of it. That agent is called Home Agent; this agent must be able to communicate with the MN all the time that it is “on-line”, independently of the current position of the MN. In order to make

the HA aware of location of the MN in real time when the MN is away from home, MN should communicate to the HA its current point of attachment and care-of address (CoA), which can be obtained by several ways. One typical way is that the MN is assigned a CoA by the currently attached agent, which is called Foreign Agent (FA). The MN sends registration request to HA via FA, then registration reply is sent back to the MN (through the FA) with allowing or denying information. If the registration is approved, the HA will work as a proxy of the MN. When MN's home network receives packets destined to the MN, HA intercepts those packets (using Proxy ARP), encapsulates them, and sends them to the care-of address, which is one of the addresses of the FA. The FA will then de-capsulate those packets and forward them to the MN. Encapsulation is the method used by the HA to deliver information to the MN putting an extra IP header on top of the packet and tunneling that packet to the MN.

When the MN moves to another network, it should send a new registration request to HA through the new FA, MN's care-of address is then updated in HA. Some extensions of Mobile IP allows MN to have several care-of addresses. In this case, HA will send the tunnels the packets to all the care-of addresses. This is particularly useful when the MN is at the edges of cells on a wireless environment with frequent mobility.

Mobile IP supports mobility management using the following procedures:

- Agent discovery: An MN is able to detect whether it has moved into a new subnet by periodically receiving unsolicited Agent Advertisement messages broadcasted from each FA. If an MN cannot receive the agent advertisement from its original agent for a specified time, it will assume to have moved into a new area. In this case the MN may passively wait for another agent advertisement. When an MN needs to get a new CoA but does not want to wait for the periodic agent advertisement, it can also actively broadcast an agent solicitation to solicit a new mobility agent in order to get a new CoA.

- CoA obtaining: A CoA can be of two different types: FA's IP address or located CoA (CCoA). A CoA can just be the FA's IP address, in which multiple MNs share a same CoA. By this mode the limited IP address space can be saved. Another type of CoA (CCoA) is a temporary IP address that is assigned to one of the MN's network interfaces by some external assignment mechanism such as Dynamic Host Control Protocol (DHCP). For a CCoA, all the natural FA functions have to be performed by the MN itself.

- Registration: When an MN discovers it is in a foreign network, it obtains a new care-of address (CoA). A registration process will be initialized to inform the HA of the new IP address. First the MN (when CCoA is used) or the FA sends to the HA a registration request which is known as Binding Update (BU). A BU can be treated as a triplet containing the MN's HA, CoA, and registration lifetime. Once the HA receives the BU, it updates the corresponding entries in its routing table and approves the request by sending a registration reply back to the MN. On the other hand, the HA may send a binding request to the MN to get the MN's current binding information if the latest BU expires. The MN then responds to the binding request with its new BU and waits for a binding acknowledgment from the HA.

- Routing and tunneling: Packets sent by a correspondent host to an MN are intercepted by the HA. The HA encapsulates the packets and tunnels them to the MN's CoA. With an FA CoA, the encapsulated packets reach the FA serving the MN, which decapsulates the packets and forward them to the MN.

A number of extended work from Mobile IP have been proposed in current researches.

Mobile IPv6 (MIPv6) [29] was developed by the Internet Engineering Task Force (IETF) to maintain connectivity while users roam through IPv6 networks. Without this mobile IP based proposals, packets sent to a mobile node cannot be received if the MN does not use its home access link (roams outside its network). While the mobile roams, it must change its IP address each time it crosses a new network. However, all connections

including the transport layer will be lost since the mobile node's IP address changes each time it moves or changes networks. Indeed, all transport protocols in the TCP/IP family define their connections with the host's and correspondent IP addresses, hence the loss of the contracted connection when one of the addresses changes. IPv6 mobility management is an important aspect of global mobility since it is envisioned that in the near future most of the Internet will be populated by IPv6 mobile nodes [30, 31].

Hierarchical Mobile IPv6 (HMIPv6) [32] complements MIPv6 by facilitating local mobility management. HMIPv6 aims to minimize global signalling and to provide improved local mobility management by introducing a hierarchical architecture. Hierarchical mobility management tends to reduce signalling overhead among the mobile node, its correspondent nodes, and its home agent. Indeed, by decomposing the network in several domains managed by a Mobility Anchor Point (MAP), a mobile node does not need to inform its correspondent nodes when it moves or roams within the same domain. Moreover, by using MAPs, a network is likely to improve MIPv6 performances in terms of handover latency.

Fast Handovers for Mobile IP is a protocol used to reduce communication interruption time during L3 handovers [34]. Simultaneous Bindings [35] is complementary to FMIPv6 and aims at reducing the number of packets lost by a mobile node during the handover procedure. To reduce packet losses, a mobile node's traffic is bi-casted or n-casted to all locations. This procedure eliminates any ambiguity regarding the moment where the traffic should be rerouted toward the mobile node's new location after a fast handover and enables the protocol to decouple L2 and L3 handovers. Moreover, Simultaneous Bindings reduce service interruption time in cases where messages are exchanged in ping pong style.

Fast Handovers for Mobile IPv6 (FMIPv6) [33] is a protocol in IPv6 family that enables a mobile node to configure a new CoA before it changes subnet. Thus, the MN is able to use the CoA right after a connection with the new Access Router is established. The goal of FMIPv6 is to minimize handover latency since an MN can neither send nor receive packets until the handover completes. The principle is to create a new CoA before losing the connection with the old AR. Therefore, when the mobile node finally connects to the new AR, it can continue its communications using a known address. If the early address registration failed, the MN can always resort to the traditional handover provided by MIPv6. The creation of the new CoA before the mobile node's migration requires the protocol to anticipate the movement. This anticipation or prediction is based on messages exchanged by the physical layer or simply information gathered on layer 2 like signal strength, for example. The goal here is to initiate a layer 3 handover (L3 handover) before the handover at Layer 2 (L2 handover) completes.

Cellular IP (CIP) [36] is proposed to offer local mobility and handoff support for frequently moving nodes. It supports fast handoff and paging in CIP access networks. For mobility between different CIP networks, it can cooperate with MIP to provide wide-area mobility support. Different wireless access networks are connected to the Internet through a gateway (GW), which handles the mobility within one domain. Packets destined to a mobile host (MH) reach the GW first. Then the GW forwards the packets to the MH using the host-specific routing path. The design of CIP is based on four fundamental principles:

- Distributed caches are used to store location information of MHs.
- Location information of an active MH is updated by regular IP datagrams originated by itself. For an idle MH, this is achieved by the use of dummy packets that are sent by the idle host at regular intervals.
- Location information is stored as soft states.

- Location management for idle MHs is separated from location management of MHs that are actively transmitting/receiving data.

CIP uses respectively distributed paging cache and distributed routing cache for location management and routing. Distributed paging cache coarsely maintains the position of the idle MHs for efficient paging, whereas the routing cache maintains the position of an active MH up to subnet level accuracy. When an MH performs handoff, the routing states in the routing cache are dynamically updated. The handoff process of CIP is automatic and transparent to the upper layers. When the strength of the beacon signals from the serving BS is lower than that of a neighboring BS, the MH initiates a handoff. The first packet that travels to the GW through the new BS configures a new path through the new BS. This results in two parallel paths from the GW to the MH: one through the old BS and one through the new BS. If the MH is capable of listening to both BSs at the same time, the handoff is soft; otherwise, the handoff is hard. The path through the old BS will be active for duration equal to the timeout of route caches. After timeout, the entries corresponding to the MH in the nodes that belong only to the old path are deleted. Thereafter, only the new path exists between the GW and the MH.

HAWAII [37] is a domain-based approach to mobility support. When an MH is in its home domain, packets destined to the MH are routed using typical IP routing. When the MH is in a foreign domain, packets for the MH are intercepted by its HA first. The HA tunnels the packets to the domain root router (*all issues related to mobility management within one domain are handled by a gateway called a domain root router*) that serves the MH. The domain root router routes the packets to the MH using the host-based routing entries. When the MH moves between different subnets of the same domain, only the route from the domain root router to the BS serving the MH is modified, and the remaining path remains the same. Thus, during an intra-domain handoff, the global signaling message load and handoff latency are reduced. To establish and maintain a dynamic path to the MH, HAWAII uses three types of messages: powerup, path refresh, and path update. The path setup messages after powerup establish the host-specific path from the domain root router to the MH by creating host-specific forwarding entries in the routers along the path. When the MH is in its home domain, once the host-specific forwarding entries are created in the routers along the path from the domain root router to the MH, the powerup procedure is complete. When the MH is in a foreign domain, it registers its CoA with its HA upon receipt of the acknowledgment from the domain root router in reply to the path setup message. Once the host-specific forwarding entries are created for an MH, they remain active for a time period. The MH periodically sends path refresh messages to its current BS before timeout occurs. In response to the path refresh messages, the BS sends aggregate hop-by-hop refresh messages to the next-hop router towards the domain root router. The path update messages are used to maintain end-to-end connectivity when an MH moves from one BS to another within the same domain. HAWAII also supports IP paging. It uses IP multicasting to page idle MHs when packets destined to an MH arrive at the domain root router and no recent routing information is available.

Besides the IETF proposition, the mobility management issues are also studied and investigated in the scientific community.

Routing scheme for macro mobility handover: Vivaldi and al. [38] use multicasting to reduce latency and packets lost during a handover. HMIPv6 reduces latency and signalization overhead during handover. However, in the case of macro-mobility or MAP inter-domain handover, the protocol is not able to guarantee QoS requirements for delay sensitive or real time flows.

Pointer forwarding MIPv6 mobility management: Chu and Weng [39] solve the handover latency problem and packet loss problems by adding pointer forwarding

capabilities to MIPv6. In large-coverage networks, signalization traffic and handover latency generated by MIPv6 registrations can be high, especially when long distances separate the mobile node's home networks from its visited network.

Mobile node extension employing buffering function: Omae et al. [40] introduce a model that improves MIPv6 handover performance by involving buffering functions into the mobile node. Mobile IP cannot perform a handover without interrupting the connections existing between the mobile node and its HA. In handover mode, packets sent to the mobile node are lost and must be retransmitted to its new address. Introducing of HMIPv6 allows minimizing service interruption time since it takes less time to update a MAP compared to a distant HA.

HiMIPv6+: Lee et al. [41] discuss the handoffs issues that include horizontal and vertical handoffs. They present a scheme for integrating wireless local area network and wide area access networks, and propose a micro-mobility management method called HiMIPv6+. They also propose a quality-of-service-based vertical handoff scheme and algorithm that consider wireless network transport capacity and user service requirement.

LCoA and SHMIP: Lai and Chiu [42] analyze the handoff delay and find that the duplication address detection (DAD) time represents a large portion of handoff delay. They assign a unique on-link care-of address (LCoA) to each mobile node and switch between one-layer IPv6 and two-layer IPv6 addressing. By this way, they propose a Stealth-time HMIP (SHMIP) which allows reducing the effect of the DAD time on the handoff delay and, thus, reduce the handoff time significantly. To further reduce packet losses, they also adopt pre-handoff notification to request previous MAP to buffer packets for the mobile node.

Registration scheme for reducing handover delay: Kaffle et al. [43] present a new scheme for fast binding update with correspondent nodes to expedite the handover operation of mobile nodes in Mobile IPv6 protocol supporting networks. The scheme is an extension of the Mobile IPv6 correspondent registration process that modifies some messages exchanged between a mobile node and a correspondent node. This scheme reduces the handover latency while providing the same level of security assurances for revealing the location information to the correspondent node as in the current Mobile IPv6 protocol.

A fast neighbor discovery and DAD scheme: Byungjoo et al. [44] present an efficient movement detection procedure which quickly sends Router Solicitation (RS) and Router Advertisement (RA) message without Random Delay using Stored RA messages with unicast. Such rapid process determines the uniqueness of a new CoA using the modified neighbor cache of an access router. In particular, they focus on the delay optimization of movement detection and DAD for fast handover in Mobile IPv6 Networks.

Predictive HO by removing the router discovery time: An et al. [45] propose an enhanced handover mechanism with new additional primitives and parameters to the media-independent handover (MIH) services defined in the IEEE 802.21. The proposed scheme can reduce handover latency for mobile IPv6 (MIPv6) by removing the router discovery time. Moreover, when the proposed mechanism is applied to the FMIPv6, they can increase the probability that the FMIPv6 can be performed in predictive mode by reducing the handover initiation time, thereby they can reduce the expected handover latency in the FMIPv6.

This section gives a brief description on several mechanisms for mobility management which are mainly based on the Mobile IP architecture. However, such approach requires additional infrastructures and intermediate systems in the networks and makes the deployment more complicated.

2.2.3 Transport layer solutions

While Mobile IP is a network layer scheme which makes mobility transparent to upper layers by increasing the burden and responsibility of the Internet infrastructure, transport layer schemes are based on an end-to-end approach to mobility that attempt to keep the Internet infrastructure unchanged by allowing the end hosts to take care of mobility. In this part, we discuss a number of transport layer schemes that have been proposed in the literature.

Msocks [46]: Maltz et al. propose TCP Splice to split a TCP connection at a proxy by dividing the host-to-host communication into host-proxy and proxy-host communications. MSOCKS [47] uses TCP Splice for connection migration and supports multiple IP addresses for multiple interfaces. When an MH disconnects itself from a subnet during handoff, it obtains a new IP address from the new subnet using DHCP, and establishes a new connection with the proxy using its second interface. The communication between proxy and CN, however, remains unchanged. The data flow between MH and CN thus continues, with the CN being unaware of the mobility. Location management is done through the proxy that is always aware of the location of the MH; this limits the mobility within the coverage of the proxy.

SIGMA [48]: Seamless IP diversity based Generalized Mobility Architecture (SIGMA) is a complete mobility management scheme implemented at the transport layer, and can be used with any transport protocol that supports IP diversity. SIGMA supports IP diversity-based soft handoff. As an MH moves into the overlapping region of two neighboring subnets, it obtains a new IP address from the new subnet while still having the old one as its primary address. When the received signal at the MH from the old subnet goes below a certain threshold, the MH changes its primary address to the new one. When it leaves the overlapping area, it releases the old address and continues communicating with the new address thus achieving a smooth handoff across subnets. Location management in SIGMA is done using DNS as almost every Internet connection starts with a name lookup. Whenever an MH changes its address, the DNS entry is updated so that subsequent requests can be served with the new IP address.

Migrate TCP [49] is a transparent mobility management scheme which is based on connection migration using Migrate TCP [49], and uses DNS for location management. In Migrate TCP, when an MH initiates a connection with a CN, the end nodes exchange a token to identify the particular connection. A hard handoff takes place when the MH reestablishes a previously established connection using the token, followed by migration of the connection. Similar to SIGMA, this scheme proposes to use DNS for location management.

Freeze-TCP [50] is a connection migration scheme that lets the MH 'freeze' or stop an existing TCP connection during handoff by advertising a zero window size to the CN, and unfreezes the connection after handoff. This scheme reduces packet losses during handoff at the cost of higher delay. Although it provides transparency to applications, Freeze TCP requires changes to the transport layer at the end nodes. Freeze-TCP only deals with connection migration, but does not consider handoff or location management. It can be employed with some other schemes like Migrate to implement a complete mobility management scheme.

Radial Reception Control Protocol (R2CP) [51] is based on Reception Control Protocol (RCP), a TCP clone in its general behavior but moves the congestion control and reliability issues from sender to receiver on the assumption that the MH is the receiver and should be responsible for the network parameters. R2CP has some added features over RCP like the support of accessing heterogeneous wireless connections and IP diversity that enables a soft handoff and bandwidth aggregation using multiple interfaces. A location management scheme might be integrated with R2CP to deploy a complete scheme.

Mobile Multimedia Streaming Protocol (MMSP) [52] supports transparent soft handoff through IP diversity and uses bi-casting (duplicate flow simultaneously) to prevent losses during the handoff period. This protocol uses Forward Error Correction (FEC) and fragmentation to mitigate wireless errors and does not include location management.

Indirect TCP (I-TCP) [53] is a mobility solution that requires a gateway along the communication path between CN and the MH to enable mobility. In this scheme, a TCP connection between CN and gateway and an I-TCP connection between the gateway and MH is established to provide CN to MH communication. The TCP portion remains unchanged during the lifetime of the communication and remains unaware of the mobility of MH. In the I-TCP portion, when the MH moves from one subnet to another one, a new connection between MH and the gateway is established and the old one is replaced by the new one. The transport layer of the MH needs to be modified but applications enjoy a transparent view of the mobility at both ends. I-TCP does not support IP diversity and soft handoff. Location management is not included in this scheme.

Mobile TCP (M-TCP) [54], an enhanced version of ITCP, is implemented at MH which works like a link layer one hop protocol that connects to the gateway via wireless. The gateway maintains a regular TCP connection with the CN and redirects all packets coming from CN to MH. This redirection is unnoticed by both the MH and CN. The enhancement of M-TCP over I-TCP is in requiring less complexity in the wireless part of the connection. Similar to I-TCP, M-TCP does not support IP diversity or location management but ensures transparency to applications.

Mobile UDP (M-UDP) [55] is an implementation of UDP protocol with mobility support similar to I-TCP and M-TCP. Like M-TCP, M-UDP uses a gateway to split the connections between MH and CN to ensure one unbroken gateway to CN connection and continuously changing MN to gateway connection. This also does not include IP diversity or location management.

The Bay Area Research Wireless Access Network (BARWAN) [56] is a solution to heterogeneous wireless overlay network. It has a gateway centric architecture on an assumption that the wireless networks are built around the gateways. Diverse overlapping networks are integrated through software that operates between the MH and the network. This allows the MH to move among multiple wireless networks, whenever MH moves out of a lower coverage network (e.g. WLAN), moves into a higher coverage network (e.g. WWAN) and MH changes its connection from lower to higher one. This scheme supports IP diversity for the MH hence enables seamless handoff across different networks. BARWAN requires the application to be aware of mobility as the decision to make a handoff is taken by the application. This scheme does not specify a location manager.

TCP Redirection (TCP-R) [57] is a connection migration scheme that keeps active TCP connections during handoff by updating end-to-end address pairs. Whenever MH gets a new IP address, TCP-R updates the address at CN and the already existent connection continues with the new address. TCP-R does not implement connection timeout to support long disconnection. Transport layer at both end-systems needs modification for this support, yet it gives application transparency. Similar to Migrate, TCP-R proposes to use DNS as location manager. Combined with a handoff management scheme, this scheme might be deployed as a complete mobility scheme.

Mobile SCTP (mSCTP) [58] supports IP diversity and soft handoff. The handoff is similar to the one of SIGMA. mSCTP is able to maintain application transparency but it does not support location management. There are several other protocols and schemes that have the potential to support (e.g. MTCP [59] is a TCP based transport protocol that supports live connections to seamlessly migrate servers) or improve performance of

mobility (e.g. pTCP [60] supports IP diversity and can achieve bandwidth aggregation of wireless networks through multiple interfaces) schemes. However, they have not been proposed as mobility schemes.

This section presents several mechanisms at transport layer to support the mobility management, however, most of these proposals only focus on the transport layer and lacks a cooperation work between the different layers. Moreover, the handover duration issues haven't been carefully investigated in this framework.

2.2.4 Application layer solutions

The Session Initiation Protocol (SIP) [61] allows two or more participants to establish a session consisting of multiple media streams. The media streams can be audio, video or any other Internet-based communication mechanism, such as distributed games, shared applications, shared text editors and white boards that have been demonstrated in practice. The media streams for a single user can be distributed across a set of devices, e.g., specialized audio and video network appliances in addition to a workstation. The protocol is standardized by the IETF and is being implemented by a number of vendors, primarily for Internet telephony. SIP may be extended to provide presence, event notification and instant messaging services and accommodate features and services such as call control services, mobility, interoperability with existing telephony systems, and more.

Entities in SIP are user agents, proxy servers and redirect servers. A user is addressed using an email-like address “user@hos”, where “user” is a user name or phone number and “host” is a domain name or numerical address. Responses to methods indicate success or failure, distinguished by status codes, 1xx (100 to 199) for progress updates, 2xx for success, 3xx for redirection, and higher numbers for failure. Each new SIP transaction has a unique call identifier, which identifies the session. If the session needs to be modified, e.g. for adding another media, the same call identifier is used as in the initial request, in order to indicate that this is a modification of an existing session. The SIP user agent has two basic functions: Listening for incoming SIP messages, and sending SIP messages upon user actions or incoming messages. The SIP user agent typically also starts appropriate applications according to the session that has been established. The SIP proxy server relays SIP messages, so that it is possible to use a domain name to find a user, rather than knowing the IP address or name of the host. A SIP proxy can thereby also be used to hide the location of the user. A redirect server returns the location of the host rather than relaying the SIP message. This makes it possible to build highly scalable servers, since it only has to send back a response with the correct location, instead of participating in the whole transaction which is the case for the SIP proxy. Both the redirect and proxy server accepts registrations from users, in which the current location of the user is given. The location can be stored either locally at the SIP server, or in a dedicated location server (more about the location server further below). Deployment of SIP servers enables personal mobility, since a user can register with the server independently of location, and thus be found even if the user is changing location or communication device. SIP requests and responses are generally sent using UDP, although TCP is also supported

SIP supports personal mobility, i.e., a user can be found independent of location and network device (PC, laptop, IP phone, etc.). The step from personal mobility to IP mobility support is basically the roaming frequency, and that a user can change location (IP address) during a traffic flow. Therefore, in order to support IP mobility, we need to add the ability to move while a session is active. It is assumed that the mobile host belongs to a home network, on which there is a SIP server (in this example, a SIP redirect server), which receives registrations from the mobile host each time it changes location. This is similar to home agent registration in Mobile IP. Note that the mobile host does not need

to have a statically allocated IP address on the home network. When the correspondent host sends an INVITE to the mobile host to establish the session, the redirect server has current information about the mobile host's location and redirects the INVITE there. If the mobile host moves during a session, it must send a new INVITE to the correspondent host using the same call identifier as in the original call setup and tells the correspondent host where it wants to receive future SIP messages. To redirect the data traffic flow, it also indicates the new address in the message field, where it specifies transport address.

Nowadays, numerous studies have focused on the mobility supported SIP, which is considered as one of the main approaches for mobility management at the application layer.

Application-Layer Mobility Using SIP [62] describes how SIP can provide terminal, personal, session and service mobility. It also describes when MIP should be preferred for terminal mobility.

Mobility Support using SIP [63] proposes mobility support at the application-layer for real-time communication. They propose to use the HA for location lookup or to let SIP redirect the invitation to the home address of the MH and let the HA forward id to the MH.

Multilayer Mobility Management for All-IP Networks: Pure SIP vs. Hybrid SIP/Mobile IP [64] evaluates the protocols ability to handle real-time sessions. Through evaluation they show the SIP superiority in real-time situations while MIP is preferred in non real-time situations. The study investigates possibilities of combining the two protocols into a hybrid solution.

In [65], authors introduce application layer techniques to achieve fast handoff for real-time RTP/UDP based multimedia traffic in a SIP signaling environment. These techniques are based on standard SIP components such as user agent and proxy which usually participate to set up and tear down the multimedia sessions between the mobiles. Unlike network layer techniques, this mechanism does not have to depend upon any additional components such as home agent and foreign agent.

Otherwise, integrated Mobile IP and SIP approach for advanced location management [66] expresses the benefits of an integrated solution to a hybrid solution to avoid redundant signaling. A Mobility Server is proposed to handle all mobility related functionality (e.g. all SIP server functionality, IP address distribution, AAA and forwarding agent).

2.2.5 Cross layer solutions

Whether the traditional layered based protocols and models are efficient for the mobility management are still questionable and remains an open topic in the current research community since they can't give a "global" view of the communication context. Nowadays, more and more researches tend to establish mobility management architectures based on a cross layer system that cooperates with the different layers to enhance the management efficiency. For example, mobility related context can be exchanged between layers to make the mobility adaptation more efficient. Link-layer parameters like signal strength, handoff initiation and completion events can be used to make network-layer handover decisions more efficient. Link-layer parameters like packet loss and bit-error-rate could also be used to fine tune the performance at the transport layer (e.g. retransmission) and for adoption at the application-layer (e.g. media encoding and transmission rate). The network-layer notifications and mobility information can be used to enhance the application-layer movement detection (e.g. new IP address). Even if the network interface automatically receives a new IP address the application is not informed and it has to perform repeated polling of the interface to detect mobility. Otherwise, mobile user's profiles (i.e. personal routing) can also be taken into consideration to enhance the efficiency of the mobility management

architecture. Cross-layer information exchange also supports service differentiation based on higher layer requirements. If a device has multiple wireless network interfaces, the selection of the appropriate interface to use could be based on user needs or multimedia application requirements.

Several studies have addressed the cross layer based mobility management as follows: The “Interlayer Signaling Pipe” [67] is proposed to store the cross layer information in the Wireless Extension Headers of IPv6 packets. This method makes use of IP data packets as in-band message carriers with no need to use a dedicated message protocol.

In [68], authors propose a mechanism that allows propagating information across layers by using ICMP (Internet Control Message Protocol) [69] messages. Since a message could be generated from any layer and then terminates at a higher layer, cross-layer signaling is carried out through these selected “holes”, rather than the “pipe” presented in [67]. Thus, this method appears more flexible and efficient. However, ICMP messages encapsulated by IP packets have to pass by the network layer even if the signaling is only desired between the link layer and the application layer. Furthermore, only upward ICMP messages were reported.

In [70], authors propose that the cross-layer information is abstracted from each related layer respectively and stored in separate profiles within a MH. Other interested layer(s) can then select profile(s) to fetch desired information. This method is flexible since profile formats can be tailored to specific layers, which in turn can access to the information directly. However, it is not suitable for time-stringent tasks.

In [71], channel and link information from the physical and link layers are gathered, abstracted and managed by third parties, the distributed Wireless Channel Information (WCI) servers. As a network service, it is complementary to the former schemes within a MH although some communication overheads would be incurred and interfaces have to be defined among the MH, the WCI server and application servers.

The Mobile people architecture [72] introduces the concept of personal routing between people that could be hosted on different devices. They focus on personal adoption and media conversion to handle cell phones that are turned off, or PCs on faraway desktops.

Cross-Layer Signaling Interactions [73] express the benefits of cross-layer information for enhancement in mobility management. It also identifies each layer’s contribution in various mobility support tasks.

Mode switching and QoS issues in software radio [74] describes how a mobile device can be reconfigured to satisfy application requirements. The authors propose switching between different air-interface standards according to QoS requirements.

Simple but effective cross-layer networking system for mobile ad hoc networks [75] explains how adaptive routing can be achieved based on low-layer information.

Wireless Channel-Aware Ad Hoc Cross-Layer Protocol with Multi-Route Path Selection Diversity [76] proposes a route discovery and route optimization mechanism based on channel information.

Cross-layer based architecture allows information from different layers potentially on different peers to be shared and transparent, especially the low layer information results in the mobile context awareness which can be used to enhances handover decisions, transport performance and media adoption, as a consequence to improve the mobility management efficiency.

2.2.6 Macro mobility solution

Macro-mobility deals with the MH’s movement across different domains. There are two main IP mobility management protocols proposed by the IETF for enabling macro mobility: MobileIPv6 that has been presented in the previous sections and the Host Identi-

tification Protocol (HIP) [100]. The Host Identity Protocol (HIP) provides a method of separating the end-point identifier and locator roles of IP addresses. It introduces a new Host Identity (HI) name space, based on public keys. The public keys are typically, but not necessarily, self-generated. HIP is being specified in the IETF's HIP working group. An IRTF HIP research group looks at the broader impacts of HIP.

MobileIPv6 and HIP protocols both maintain a fixed proxy (Home Agent / Rendezvous Server), a host that is aware of the current location and address of a node. This architecture enables permanent reachability even with mobile nodes. MobileIPv6 and HIP also offer an address change notification mechanism to preserve established transport sessions in the presence of macro mobility. For both of them, Internet drafts are proposed which describe extensions to enable multi-homing [109-111]. Note however that these two IP mobility management protocols and their extensions for multi-homing require that a node explicitly knows the access networks over which its packets are forwarded to the Internet. To deal with macro mobility, a moving node updates its address to topologically fit to the access network relaying its packets and notifies its fixed proxy as well as its communication peers about its address change.

2.2.7 Micro mobility solution

Micro mobility solutions are presented for the intra-domain mobility management to implement a fast and seamless handoff and minimized control traffic overhead. A movement of mobile node in a foreign network domain need not inform the MN's HA of the new attachment. The micro mobility protocols ensure that the packets arriving at the mobility server (gateway) can be correctly forwarded to the appropriate access point that the MN currently attaches.

There are three typical proposals for micro mobility solution: HMIP, Cellular IP [36], and HAWAII [37]. Table 1 shows a simple comparison of the three proposals. Note that none of these suggestions are trying to replace the Mobile IP. Instead they are enhancements to the basic Mobile IP with the micro mobility management capability. There are still many other typical micro mobility proposals existing, e.g. the Intra-Domain Mobility Management Protocol (IDMP) [112], Edge Mobility Architecture [113], Hierarchical Mobile IPv6 (HMIPv6) [32], etc.

	Hierarchical MIP	Cellular IP	HAWAII
OSI Layer	"L3.5"	L3	L3
Nodes Involved	FAs	all CIP nodes	all routers
Mobile Host ID	home addr	home addr	c/o addr
Intermediate Nodes	L3 routers	L2 switches	L2 switches
Means of Update	signalling msg	data pkt	signalling msg
Paging	explicit	implicit	explicit
Tunnelling	yes	no	no
L2 Triggered Handoff	no	optional	optional
MIP Messaging	yes	no	yes

Table 2.1 Comparison of Cellular IP, Hawaii and Hierarchical Mobile IP

2.2.8 802.21 based mobility management

802.21 [121] is an IEEE emerging standard which enables seamless handover between networks of the same type as well as the Media independent handover (MIH) between different network types (i.e. cellular, GPRS, Bluetooth, 802.11 and 802.16 networks). 802.21 defines an infrastructure intended to improve mobile devices' handover decisions

based on mainly lower-layer information from both mobile devices and the access network (i.e. neighboring access networks and other link layer information as well as information about a limited set of available higher-layer services like QoS and Internet connectivity). The infrastructure is used for the cases of handover optimization, particularly between heterogeneous networks. Meanwhile, context-aware services framework is essential in the context of IEEE 802.21 to support a three-tier architecture consisting of a sensor, a convergence, and a service applications.

Generally speaking, the 802.21 framework consists of the following elements:

1) An architecture that enables transparent service continuity while a mobile node (MN) switches between heterogeneous link-layer technologies. The architecture relies on the identification of a mobility-management protocol stack within the network elements that support the handover. The description of the architecture does not address implementation details and does not provide indication of preferred implementations of the IEEE 802.21 standard. The architecture presents MIH Reference models for different link-layer technologies.

2) A set of handover-enabling functions within the mobility-management protocol stacks of the network elements and the creation therein of a new entity called the MIH Function. A media independent Service Access Point (called the MIH_SAP) and associated primitives are defined to provide MIH users with access to the services of the MIH Function, listed below:

- The Media Independent Event service detects events and delivers triggers from both local and remote interfaces.
- The Media Independent Command service provides a set of commands for the MIH users to control handover-relevant link states.
- The Media Independent Information service provides the information model and an information repository to make more effective handover decisions. The mobile terminal obtains information from the repository using its current network point(s) of attachment.

Based on the IEEE 802.21, the handover part of our proposed mobility management architecture has been designed to satisfy the requirements of seamless communication during frequent handovers, thanks to the intelligent modules involved such as *Link-layer signal analysis module*, *mobility prediction module*, *context awareness* and *best access point selection module*.

2.2.9 Open issues of mobility management

A rapidly growing interest has been focused on the next generation wireless networks that support global roaming across multiple wireless and mobile networks, there are principally three main issues that should be carefully taken into consideration regarding the mobility management.

1) Firstly, a QoS aware mobility architecture should be taken into consideration to achieve the end-to-end QoS guarantee for mobile nodes' handover between heterogeneous networks. This architecture should take account of low layer signal information as well as the available QoS offered by the candidate access networks to support a smooth handover on the upper layer across heterogeneous networks with different QoS supports. It refers to a problem of optimal choice of the access networks, or how to be best connected. If a mobile user is offered connectivity from more than one technology at any one time, one has to efficiently decide which access network is most optimal for the services involved in the mobile terminals according to different parameters such as available QoS, cost, degree of mobility, etc. A handover algorithm should both determine which network to connect to as well as when to perform a handover between the different networks. Ideally, the handover algorithm would assure that the best overall wireless link is chosen.

2) The second issue focuses on fast, seamless horizontal/vertical handovers (IP micro-mobility), as well as guaranteed quality of service (QoS), security and accounting. Real-time multimedia applications in the future will require fast/seamless handovers for smooth operation. Meanwhile, high variations in the network Quality of Service (QoS) leads to significant variations of the multimedia transmission quality. The result could sometimes be unacceptable to the users. Avoiding such symptom requires choosing an adaptive encoding framework for multimedia transmission. The network should signal QoS variations to allow the application to be aware in real time of the network conditions. A “cross layer” interaction is supposed to help ensure personalized adaptation of the multimedia presentation.

3) The third issue is to find an efficient mobility management architecture, which requires minimal changes to the current Internet stack. Nowadays, Mobile IP is developed to support the mobility management which is compatible with the all IP-based network, however, this approach requires significantly additional infrastructures and intermediate system which complicate the networks as well as the its management. How to integrate the efficient and intelligent mechanisms at the end systems to support efficient mobility management, which make the deployment easier and much more flexible, is a big challenge in the context of wireless networks.

2.2.10 Conclusion of section 2.2

In this section, we introduce the backgrounds and related work of mobility management architectures and mechanisms. We give an analysis on the currently proposed approaches that are involved at different layers and scheme patterns; we conclude that the efficient mobility management entails all the layers’ participation in a highly co-operative way. Cross-layer interactions play a crucial role to enhance the performance of the mobility management.

2.3 WLAN Issues

Wireless Local Area Network (WLAN) is a wireless alternative to a Local Area Network (LAN) that uses wireless radio to transmit data in a small area. Wireless LANs are standardized under the IEEE 802.11 series. As the importance of mobility and nomadic user profiles has increased, WLANs gained attention especially in home, office, and campus environments. Low infrastructure cost, ease of deployment and support for nomadic communication are among the strengths of WLANs. Deployment without cabling and ease of adding a new user to the network decrease the deployment cost of a WLAN dramatically.

2.3.1 Background of 802.11

The 802.11 family have defined full-scale standards which covers different aspects of 802.11 issues, such as the physical transmission techniques, wireless access method, security, connection issues, etc. This section will give a description on 802.11 family.

2.3.1.1 802.11 Physical layer

Various PHY layers are available in the IEEE 802.11 family, The 802.11a standard uses an OFDM based air interface and operates in the 5 GHz band with a maximum net data rate of 54 Mbit/s. Since the 2.4 GHz band is heavily used to the point of being crowded, using the relatively un-used 5 GHz band gives 802.11a a significant advantage. However, this high carrier frequency also brings a disadvantage of short overall range and 802.11a

signals cannot penetrate as far as the other 802.11 schemes (i.e. 802.11b/g) because they are absorbed more readily by walls and other solid objects in their path due to their smaller wavelength.

802.11b products appeared on the market in early 2000. Since 802.11b uses a new PHY layer, High Rate DSSS (HR-DSSS) based on DSSS, a relative bigger range as well as substantial price reductions led to the rapid acceptance of 802.11b as the definitive wireless LAN technology. However, 802.11b devices suffer interference from other products operating in the 2.4 GHz band, which includes microwave ovens, Bluetooth devices, baby monitors and cordless telephones.

802.11g [14,15] was rapidly adopted by consumers starting in January 2003, well before ratification, due to the desire for higher data rates, and reductions in manufacturing costs by using various methods (e.g., Atheros SuperG [18]). By summer 2003, most dual-band 802.11a/b products became dual-band/tri-mode, supporting a and b/g in a single mobile adapter card or access point. 802.11g works also in the 2.4 GHz band, but uses the same OFDM based transmission scheme as 802.11a. It operates at a maximum physical layer bit rate of 54 Mbit/s exclusive of forward error correction codes, or about 19 Mbit/s average throughputs.

802.11n [19] is a proposed amendment which improves upon the previous 802.11 standards by adding multiple-input multiple-output (MIMO) [5] and many other newer features such as 40MHz operation [6] and supports a maximum transmission rate of 580 Mbit/s. The workgroup is not expected to finalize the amendment until December 2009. Enterprises, however, have already begun migrating to 802.11n networks based on Draft 2 of the 802.11n proposal. A common strategy for many businesses is to set up 802.11n networks to support existing 802.11b and 802.11g client devices and while gradually moving to 802.11n clients as part of new equipment purchases.

Most of the IEEE 802.11 PHY layers work in the 2.4 GHz frequency band (2.414–2.484 GHz) with 14 (includes US) distinct channels. The availability of these 14 channels varies from country to country. IEEE 802.11 does not have a fixed channel bandwidth but the standard dictates several rules about signaling such as the center frequency of these channels must be at least 5 MHz apart from each other and the power levels of the signals for nearby frequencies cannot exceed certain thresholds. A typical APs signal does not extend more than 22 MHz from center frequency of the selected frequency. As a result only three of the 14 channels do not overlap. IEEE 802.11a, on the other hand, works in the 5 GHz (5.15–5.825 GHz) frequency band with a fixed channel center frequency of 5MHz. The number of channels varies from 36 to 161 depending on the frequency band. There are 12 non-overlapping channels (with center frequencies 20 MHz apart from each other) in the frequency band used by IEEE 802.11a in the US and 19 non-overlapping channels in Europe.

A typical WLAN access point (AP) uses one omni-directional antenna and has an average range of 70 m indoors and 200 m outdoors. This range (especially indoors range) is greatly affected by the obstacles between the AP and the mobile node, link condition, and the security measures used in the WLAN. Using directional antennas, directed peer-to-peer (P2P) WLAN links can be established within range of a few kilometers. Working at a higher frequency band, IEEE 802.11a networks suffer more from increased range and attenuation compared to IEEE 802.11b/g networks. In [77], it is shown that using sectorized antennas instead of omni-directional antennas greatly increases the aggregate WLAN data rate in a given area with a factor of two or three.

2.3.1.2 802.11 MAC layer

The object of 802.11 MAC layer is to coordinate the use of the physical medium by controlling the transmission of packets from multiple stations, where a station is either a host or an AP. This is accomplished by using the random access protocol CSMA/CA [117] protocol.

The IEEE 802.11 MAC uses a Distributed Coordination Function (DCF) for media access. DCF is an implementation of CSMA/CA protocol, which offers a best effort type of service, each mobile station checks whether the medium is idle before transmission, immediate transmission is available if the medium idles for longer than DIFS (Distributed Inter Frame Space). Contrarily, the mobile station waits as access deferral for DIFS and retries after an exponential backoff delay if the medium is busy. The receiving station will check the CRC of the received packet and send acknowledgment packet (ACK). Receipt of the acknowledgment will indicate to the transmitter that no collision occurred. If the sender does not receive the acknowledgment, it will retransmit the fragment until it gets acknowledged or thrown away after a given number of retransmissions.

In order to avoid the hidden node problem, Virtual Carrier Sense [122] has been proposed as an option to cooperate with DCF to reduce the probability of two stations colliding because they cannot hear each other. A station that wishes to transmit a packet first needs to 'sense' the channel. Following this, it generates a random backoff timer chosen uniformly from the window $[0, w-1]$, where w is the contention window. Initially w is set to CW_{min} (16 for 802.11b). After the backoff timer expires, the node sends a short Request To Send (RTS) message to the intended receiver of data. If this message is received properly by the receiver and if it is able to receive any transmission, it responds back with a short Clear To Send (CTS) message. A node cannot receive any transmission if some other node in its vicinity has already reserved the channel for packet reception or transmission. Both RTS and CTS messages carry the duration information for which the channel is going to be occupied by the proposed data transmission. Upon hearing RTS and CTS, all other nodes update their Network Allocation Vectors (NAVs) with the information about the duration for which the channel is going to be busy and defer their transmissions and receptions to avoid collisions. The CTS message is followed by the DATA transmission, and if the data frame is received successfully, a MAC level ACK is returned to the sender. In case of the packet losses, the data frame is repeatedly retransmitted in the absence of ACKs till a threshold number of retransmissions are carried out. Once the retransmissions exceed the threshold, the transmission is assumed to be unsuccessful. After an unsuccessful transmission attempt, the sender follows a binary exponential backoff (BEB) and doubles its contention window size to reduce the channel contention between nodes. The contention window is not incremented further if it already equals CW_{max} (256 for 802.11b). After every successful transmission, the contention window is reset back to CW_{min} . RTS, CTS, DATA, and ACK are separated by a time spacing of Short Inter Frame Space (SIFS). Each successful transmission follows the procedure of RTS-CTS-DATA-ACK. A node may choose to disable the virtual carrier sense (RTS, CTS) to reduce its overhead when the probability of existence of hidden nodes is known to be small.

Otherwise, 802.11 MAC layer offers link management entities, 802.11 defines several messages which are transmitted as link level frames. The frames that are exchanged between different nodes of an 802.11b network are classified into three categories:

- 1) Management frames: Management frames mainly refer to the frames that are used for security purpose, affiliation activity of any node to a particular cell and access-point such like Authentication, Association frames. Beacon frame which carries important management information like time-stamp, supported rates, traffic indication maps etc.

- 2) Control frames: Control frames refer to those which allows controlling the data

transmission, such like RTS, CTS, ACK, PS-Poll (power save poll) CF-Poll, CF-ACK, CF-END (Contention free channel information in PCF [123]).

3) Data frame: Data frame refers to that carries the upper-layer payload.

2.3.1.3 802.11 QoS support

802.11 standard was originally designed with best effort traffic in mind. The contention based access method allows mobile node to compete and capture the channel resources for data transmission, but it does not guarantee any QoS to the connections. Based on DCF, 802.11e [12] defines the Enhanced Distributed Coordination Function (EDCF)[118] which supports a priority based best effort service and significantly enhances QoS support in WiFi. In EDCF, frames corresponding to different traffic categories are transmitted through different backoff instances. The scheduling of frames for every traffic category is done the same way as in DCF. The different priority is achieved by setting different probabilities for different categories for competing the channel. The probability is changed by varying the so-called Arbitration Inter Frame Space (AIFS) which is the listening interval for channel contention. For each traffic category, the value of AIFS determines the priority. If the AIFS values are low, the listen interval required for channel contention is lower and hence the probability of winning the channel contention is higher. For compatibility with legacy DCF, AIFS should be at least equal to DIFS. The backoff procedure in case of collisions is similar to that in DCF whose contention window is always doubled. The contention window in EDCF is expanded by a predetermined persistence factor (PF). A single node can have up to eight traffic categories. These different categories are realized as eight different virtual nodes with varying parameters (i.e. AIFS, CW, and PF). These parameters are responsible for determining the priority of each traffic category. If the backoff counters of multiple traffic categories reach zero at the same time, a virtual collision happens within the same physical node but different traffic categories, then the EDCF scheduler inside the node allows the data transmissions for the traffic category with the higher priority.

Otherwise, other mechanisms [79, 80] in the literature have been proposed to further improve the performance of IEEE 802.11e. It has been shown that using adaptive contention window, Inter-Frame Space (IFS) parameters based on the link condition can enhance the overall QoS performance.

2.3.1.4 Other 802.11 standards

IEEE 802.11d is an amendment to the 802.11 specification that adds support for "additional regulatory domains". This support includes the addition of a country information element to beacons, probe requests, and probe responses.

IEEE 802.11f or Inter-Access Point Protocol is a recommendation that describes an optional extension to 802.11 that provides wireless access-point communications among multi-vendor systems.

IEEE 802.11h addresses the Spectrum and Transmit Power Management Extensions; it specifies two mechanisms to IEEE 802.11, Dynamic Frequency Selection (DFS) and Transmit Power Control (TPC) [125]. With DFS, the AP detects other networks operating at the same frequency band in the same region and changes the operating frequency of the WLAN to prevent collision. TPC is used to keep the signal level below a threshold if a satellite signal is available in nearby channels. This mechanism can also be used to improve link condition by changing the working frequency to a clearer channel and also to reduce power consumption.

IEEE 802.11i [13] specifies security mechanisms for 802.11 networks. The draft standard came out in June 2004, and supersedes the previous security specification, Wired

Equivalent Privacy (WEP), which was shown to have severe security weaknesses. Wi-Fi Protected Access (WPA) had previously been introduced by the Wi-Fi Alliance as an intermediate solution to WEP insecurities. WPA implemented a subset of 802.11i. The Wi-Fi Alliance refers to their approved, interoperable implementation of the full 802.11i as WPA2, also called RSN (Robust Security Network). 802.11i makes use of the Advanced Encryption Standard (AES) block cipher. The 802.11i architecture contains the following components: 802.1X for authentication (entailing the use of EAP and an authentication server), RSN for keeping track of associations, and AES-based CCMP to provide confidentiality, integrity and origin authentication, which enhance the security performance.

IEEE 802.11k enables management of the air interface between multiple APs. This standard specifically addresses environments with many APs. An AP can order a mobile station to make a site survey and report the results. Then, the AP takes admission control decisions based on this information. Normally, a mobile station searching for an AP connects to the AP with the strongest signal. Using IEEE 802.11k, a congested AP may reject a new mobile station and force it to connect to a less congested AP within the STAs range.

IEEE 802.11p (Access for Vehicular Environments - WAVE) supports data exchange between high-speed vehicles and between the vehicles and the roadside infrastructure for the purpose of ITS (Intelligent Transportation) [124]. Using this protocol, vehicles send information about their traffic parameters (speed, distance from other vehicles, etc.) to nearby vehicles. Thus, each vehicle knows the current traffic status and acts accordingly. IEEE 802.11p is planned to work at the 5.9 GHz frequency band, which is not compatible with IEEE 802.11a/b/g/n.

IEEE 802.11r specifies fast Basic Service Set (BSS) transitions between access points by redefining the security key negotiation protocol, allowing both the negotiation and requests for wireless resources, it allows that the mobile nodes keep track of the nearby APs and communicates with them before making the actual handover, and then decrease delay during the handover procedure.

IEEE 802.11s is a universal solution which allows forming a mesh topology between APs, and remedying interoperability problems between different mesh support mechanisms. IEEE 802.11s has broadcast, multicast, and uni-cast support and is expected to include multiple routing algorithms between APs.

IEEE 802.11u contains requirements in the areas of enrollment, network selection, emergency call support, emergency alert notification, user traffic segmentation, and service advertisement.

2.3.2 Open issues and related work for 802.11

Apart from the existing family of 802.11 standards, there are still some open issues which have been addressed by many contributions in the current scientific community.

2.3.2.1 MAC layer throughput

Although the transmission rate of 802.11 has been significantly increasing since its original version by integrating several promising technologies (ex: transmission rate of 802.11g reaches 54Mbps), the real MAC throughput is much less than that rate especially when RTS/CTS are taken into consideration. In [81], the authors show that the maximum MAC layer throughput is limited to 75Mbps which does not depend on the PHY layer data rate. IEEE 802.11e and IEEE 802.11n introduce new mechanisms to increase the MAC layer throughput; However, as proved in [81], it's difficult to significantly improve

the MAC layer throughput without modifying the current 802.11 standard and MAC layer access method.

2.3.2.2 Security issues

Security is another important issue in 802.11 context, the initial available encryption method (WEP: Wired Equivalent Privacy) [116] has been identified to have big security problems. The open system and the shared key authentication are proved not to work very well to defend attacks. WiFi alliance then developed a new encryption method named WiFi Protected Access (WPA) [120] to improve the security performance. IEEE 802.11i addressed this issue and incorporated IEEE 802.1X authentication method which is used in all IEEE 802 family standards. The users can authenticate their identities by a RADIUS or a Diameter server [78] with this method. Many studies are currently focusing more efficient security mechanism without scarifying the MAC throughput.

2.3.2.3 Reliability and performance

Like other wireless networks, data can be lost during the transmission when signal fades or interference occurs. FEC [83] and ARQ [84] are introduced to ameliorate the transmission reliability. [85, 86] present a new approach which is based on a precise analysis of the error process and allows predicting the PER (Packet Error Rate) according to the convolution code parameters as well as the SNR over the AWGN channel [119], they give an analysis of the evaluation of PER in Wireless Networks. In [87], authors propose a model to predict BER of time varying 802.11 wireless channel with a neural network system. In [88], the authors present an analysis of PER of WLAN under the Interference of other wireless networks such as IEEE 802.15.4. On the other side, some different rate adaptation mechanisms such as ARF (Auto Rate Feedback) [89], CARE (Collision-Aware Rate Adaptation) [90], RRAA (Robust Rate Adaptation Algorithm) [91] have been introduced to balance the tradeoff between error rate and bandwidth and choose an optimal modulation according to the quality of link channel. Higher modulations are applied to mobile nodes to offer high throughputs when signal is optimal, on the contrary, lower modulations should be applied to guarantee the reliability of transmission when nodes experience poor signal periods. However, all of these work focus on the packet losses due to the signal fading and interference, none of them has take into account the losses at the 802.11 MAC layer due to the lack of interaction between the higher and WLAN MAC layer. Indeed, the higher layer is unaware of the MAC layer information, then the sending rate from higher layer can surpass the maximum bandwidth supported by the 802.11 MAC layer, as a consequence to have packets lost at the MAC buffer. This problem was firstly addressed in [82].

Performance of TCP/IP over wireless networks is a hot topic as well. TCP behavior is addressed under different wireless scenarios. The primary thrust of this research is on the behavior of TCP in response to the error conditions in wireless networks. TCP, which is known to be a very stable and robust protocol on wired networks, does not perform as well on wireless networks. Some of the serious issues with TCP are that of performance degradation over wireless links. The primary reason for performance degradation is the assumption by TCP that all losses in the network are due to congestion [96] while wireless losses are mainly due to bad channel conditions. Distinguishing the natures of losses remains a big challenge to optimize the communication efficiency in the context of wireless networks. Meanwhile, it is evident that all layers (not only the high transport layer) in the protocol stack should adapt to the variations in the wireless link appropriately in order to optimize the network performance. This should be done while considering the adaptive strategies at other layers and the cross layer based architecture.

2.3.2.4 Fairness issues

Fairness Issues is another crucial problem in the 802.11 context. Numerous work have focused on the unfairness issues of TCP connections in WLAN as it is a widely accepted and implemented protocol over the Internet today. Shugong Xu et al. [94] analyze the behavior of TCP protocol in multi-hop 802.11 networks. Using simulations they show that TCP suffers from instability and unfairness problem in WLAN networks. The instability causes the throughput of available wireless network to fluctuate because of interactions between different nodes carrying TCP-data and TCP-ACK traffic. The unfairness problem leads to indefinitely long timeouts causing multiple retransmissions and connection breakups. Koksai et al. [95] describe a scenario where TCP performance degrades because of short term unfairness exhibited by MAC protocols. Because of short term unfairness the TCP acknowledgments fail to reach the sender in a timely fashion. This bursty traffic results into ACK compression which aggravates the “burstiness” of the stream. The bursty traffic has many disadvantages like packet loss in response to bursty traffic and throughput loss because of idle links during two consecutive bursts. In [92], authors also give an analysis on the short-term fairness for TCP Flows and analyze the unfairness problems caused by channel unavailability. In [93], authors lead an in depth analysis on the TCP fairness over Wireless LAN, they conclude that the buffer size at the base station plays an important role in the observed unfairness. Even in a scenario that considers only TCP connections, the unfairness in TCP throughput ratio between upstream and downstream flows can be as large as 800. Otherwise, the contention avoidance procedures implemented at the 802.11MAC layer of access points can slower the rate of returned ACK packets, which degrades the performance of ACK clocked protocols such as TCP and results in the unfairness problems between the ACK clocked (i.e. TCP) and Non- ACK clocked (i.e. UDP) connections .Because of this TCP intrinsic potential of strong unfairness in the context of wireless access networks, TCP appears un-adapted to WLANs.

2.3.3 Conclusion of section 2.3

This subsection gives a description on the background of 802.11 technique, we specifies the functions and performance of MAC and physical layer as well as an introduction of 802.11 family. We also mention the related studies of 802.11 which address several open issues such as fairness issues, transmission performance and reliability, security.

2.4 Conclusion of chapter 2

After a brief introduction on the wireless networks, this section focuses on the backgrounds and the related work of mobility management and WLAN performance issues, which are the two principal issues that will be addressed in the thesis. The next section will introduce a novel mobility management architecture, which is based on the IEEE 802.21 framework, to support efficient handover and location management and allows offering seamless communication in the mobility context. Then in chapter 4, 5 and 6, we will focus on several intrinsic problems related to the access method CSMA/CA in the context of 802.11 as well as the modulation schemes of 802.16 networks, we will see that our cross layer based solutions are able to significantly enhance the wireless MAC performance and transmission efficiency.

Chapter III

3 End to end architecture for mobility management

This section introduces an end to end architecture for mobility management which is not dedicated to a specific protocol or technology. The aim of the various mechanisms and protocols introduced in this chapter is to offer efficient mobility management and preserves efficiently session continuity in the context of mobile and wireless networks.

3.1 Introduction

The proliferation of laptops, cellular phones, and other mobile computing platforms connected to the Internet has triggered numerous research work into mobile networking. The increasingly dense set of wireless access networks that can be potentially accessed by mobile users open the door to an era of pervasive computing. However, the puzzle of wireless access networks that tends to become the natural access networks to the Internet pushes legacy “wire-oriented” communication architectures to their limit. Indeed, there is a critical gap between the increasingly used stream centric multimedia applications and the incapacity of legacy communication stacks to insure the continuity of these multimedia sessions for mobile users.

Nowadays, a great number of studies about the mobility management focus on the IP layer that does not take into account of the “end-to-end” and “cross-layer based” conceptions which allows simplifying the deployment and enhancing the performance of the mobility management due to the information transparency among different ISO layers. This chapter proposes an end to end communication architecture which aims to fill the gap between the application layer continuity needs and the discontinuity of the communication service inherent to the physical layer of wireless mobile networks. Different from the other proposed mobility management models, our end to end architecture aims to introduce as few changes as possible in the current Internet protocol architecture, meanwhile, our architecture binds a set of novel mechanisms that can be dynamically inserted in new generation reconfigurable protocols such as defined in [97] to support efficient mobility management. The proposed contribution efficiently satisfies mobility requirements such as efficient location management, fast handover, and continuous connection support.

3.2 Mobility management architecture

In this emerging pervasive computing era, computer networking will be carried on by a variety of mobile end-systems which can migrate anytime across different wireless subnets while keeping seamlessly their communication sessions. Nowadays, multimedia continuous streams (e.g. video or audio streams) take a bigger and bigger part of the information flows accessed or exchanged by Internet users. This feature enforces the requirement of offering seamless communication to mobile users. In today’s Internet, there is no widely

disseminated, available and used communication architecture to efficiently and with a reduced infrastructure cost address the whole scope of mobility management issues.

By performance of the mobility management, we mainly address the following issues: 1) Continuous connection support: an established connection should be suspended instead of being cut off during the migration of the mobile nodes, and the continuous communication is available as soon as the host gets reconnected. 2) Fast handover: minimize the duration of handover to support a seamless communication. 3) Efficient location management: the current mobile node's network address should be accessible any time at a light cost. 4) Seamless communication support application as well as the related mechanisms.

Our mobility management architecture will then be presented in details with the following 4 principal parts:

- Mobile connection analytical model: we establish an analytical model, based on a Markovian model, to estimate the stationary probability of simultaneous connectivity for the two corresponding mobile terminals in ideal conditions, which is considered as an optimal utility that our efficient architecture for mobility management should try to approach as close as possible but will never be able to reach.
- Location/Address Management (LM): Two options for mobility location management have been studied and integrated in our architecture, Dynamic DNS [98] and HIP Rendezvous server (RVS) mechanism [99] which is an extension work of Host Identity Protocol [100]. According to our experimental tests, Dynamic DNS is more adapted to infrequently moving nodes, while RVS fits better with frequently moving nodes.
- Continuous connection support: It has been modeled and designed at the transport layer, based on the location/address management, to support continuous connection and guarantee the mobile nodes to connect as soon as possible once the communication is available.
- Handover Management: Handover Management is the key part of our cross layer based architecture, which comprises 1) Signal analysis function that allows an intelligent collection and analysis of the signal information (SNR, Rx power, MAC layer) from lower layers, these information is used to predict the signal evolution and estimate an optimal moment to start the handover procedure. 2) Implementation of mobility prediction in our architecture allows greatly improving the performance of the handover and transmission quality thanks to the several pre-reservation based services. 3) Handover procedure allows minimizing the delay based on signal intelligent analysis and pre-reservation based services.

Following the modeling and performance analysis of these mechanisms, we have developed a Java-based middleware interface that integrates the proposed mechanisms involved in our mobility management architecture, which will be presented in chapter 6.

3.3 Mobile connection analytical model

In this section, we introduce a Markovian analysis of the notion of mobile connection. This analytical model gives the limit of the mobile connection utility at the physical layer and ignores handover processing overhead. This modeling gives the higher bound of the probability to be able to communicate for two corresponding mobile nodes. In practice, this limit can never be reached but can be approached thanks to our proposed architecture.

Differently from a classic transport connection, we define a mobile connection by a stochastic tuple: $(@S(t); , @D(t))$, where $@S(t)$ and $@D(t)$ are the Source Host Name and the Destination Host Name that are continuous functions of time and dynamically evolve following the mobility behavior of the two end systems. The dynamic association between host names and their current network address can be simply managed by DDNS or the Rendezvous server (RVS to be introduced in section 3.4.2). The above definition of a mobile connection leads to define a stochastic model that enables their stochastic performance analysis.

Based on a Markovian model, our analysis helps to estimate the stationary probability of simultaneous connectivity for the two corresponding mobile terminals (MN1, MN2) in ideal conditions (i.e. lossless channel and perfect seamless handover). This estimated probability gives the higher bound of the mobile connection utility and can also be useful for mobility prediction purpose.

We denote C_i , the available communication state for MN_i ; and NC_i , the non-available communication state for MN_i

We suppose, when MN1 is already in the state $C1$, at a moment “t” given:

- The probability of staying in the communication state ($C1$) is $p1(t)$.
- The probability of passing to the non-available communication state ($NC1$) is $1-p1(t)$.

When MN1 is already in the state $NC1$, at a moment “t” given:

- The probability of staying in the state $NC1$ is $q1(t)$.
- The probability of passing to the state $C1$ is $1-q1(t)$.

When MN2 is already in the communication state $C2$, at a moment “t” given:

- The probability of staying in the state $C2$ is $p2(t)$.
- The probability of passing to the non-available communication state $NC2$ is $1-p2(t)$.

When mobile 2 is already in the state $NC2$, at a moment “t” given:

- The probability of staying in the state $NC2$ is $q2(t)$.
- The probability of passing to the state $C2$ is $1-q2(t)$.

Fig 3.1 represents the transition diagram of 4 different states. ($C1, C2$), ($C1, NC2$), ($NC1, C2$), ($NC1, NC2$). (We replaced $pi(t), qi(t)$ by pi, qi to simplify the expression).

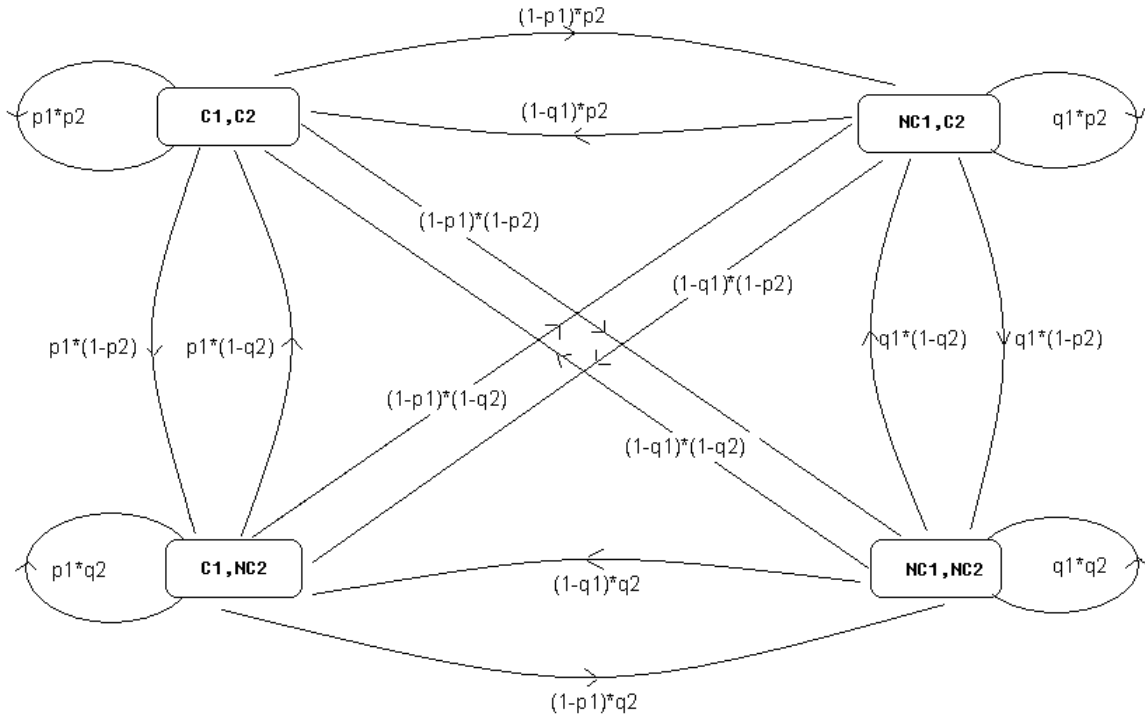


Fig 3.1 Transition diagram

We suppose that each of the two communicating mobile nodes MN1 and MN2 follows a model of mobility defined by a two-state continuous Markov chain and we have: $pi(t) = \exp(-\lambda i * t)$; $qi(t) = \exp(-\mu i * t)$, where λi is the rate of the exponential distribution that defines for MNi its probability to be able to communicate (so $\frac{1}{\lambda i}$ represents average period available for communication), and μi is the rate of the exponential distribution that gives for MNi its probability to be disconnected from an access network (so $\frac{1}{\mu i}$ represents the average period NOT available for communication). The global behavior of the two mobile nodes can be analytically modeled by composing these two continuous Markov chains.

We introduce a matrix A which represents the probabilities of passing from one state to another one.

Matrix A:

	(C1,C2)	(NC1,C2)	(C1,NC2)	(NC1,NC2)
(C1,C2)	$\exp((- \lambda 1 + \lambda 2) * t)$	$(1 - \exp(- \lambda 1 * t)) * \exp(- \lambda 2 * t)$	$(1 - \exp(- \lambda 2 * t)) * \exp(- \lambda 1 * t)$	$(1 - \exp(- \lambda 1 * t)) * (1 - \exp(- \lambda 2 * t))$
(NC1,C2)	$(1 - \exp(- \mu 1 * t)) * \exp(- \lambda 2 * t)$	$\exp(- \mu 1 * t) * \exp(- \lambda 2 * t)$	$(1 - \exp(- \mu 1 * t)) * (1 - \exp(- \lambda 2 * t))$	$(1 - \exp(- \lambda 2 * t)) * \exp(- \mu 1 * t)$
(C1,NC2)	$(1 - \exp(- \mu 2 * t)) * \exp(- \lambda 1 * t)$	$(1 - \exp(- \mu 2 * t)) * (1 - \exp(- \lambda 1 * t))$	$\exp(- \mu 2 * t) * \exp(- \lambda 1 * t)$	$(1 - \exp(- \lambda 1 * t)) * \exp(- \mu 2 * t)$
(NC1,NC2)	$(1 - \exp(- \mu 2 * t)) * (1 - \exp(- \mu 1 * t))$	$(1 - \exp(- \mu 2 * t)) * \exp(- \mu 1 * t)$	$(1 - \exp(- \mu 1 * t)) * \exp(- \mu 2 * t)$	$\exp(- \mu 2 * t) * \exp(- \mu 1 * t)$

We suppose matrix A* is the matrix derived from matrix A when t=0, we have matrix A*:

	(C1,C2)	(NC1,C2)	(C1,NC2)	(NC1,NC2)
(C1,C2)	$-(\lambda 1 + \lambda 2)$	$\lambda 2$	$\lambda 1$	0
(NC1,C2)	$\mu 1$	$-(\mu 1 + \lambda 2)$	0	$\lambda 2$
(C1,NC2)	$\mu 2$	0	$-(\lambda 1 + \mu 2)$	$\lambda 1$
(NC1,NC2)	0	$\mu 2$	$\mu 1$	$-(\mu 1 + \mu 2)$

We suppose that Xi represents the probability in different states. X1 for state (C1,C2); X2 for state (NC1,C2); X3 for state (C1,NC2); X4 for state (NC1,NC2). We have $X A^* = 0$, Then the following 4 equations are induced:

$$-(\lambda 1 + \lambda 2) * X1 + \mu 1 * X2 + \mu 2 * X3 = 0 \quad (1)$$

$$\lambda 1 * X1 - (\mu 1 + \lambda 2) * X2 + \mu 2 * X4 = 0 \quad (2)$$

$$\lambda 2 * X1 - (\lambda 1 + \lambda 2) * X3 + \mu 1 * X4 = 0 \quad (3)$$

$$\lambda 2 * X2 + \lambda 1 * X3 + -(\mu 1 + \mu 2) * X4 = 0 \quad (4)$$

And we have:

$$X1 + X2 + X3 + X4 = 1 \quad (5)$$

Then from the resulting composed chain, the stationary probability of simultaneous connectivity for the two MHs (X1) can be derived and is given by:

$$X1 = \frac{\mu_1 \mu_2}{\lambda_1 \lambda_2 + \lambda_1 \mu_2 + \lambda_2 \mu_1 + \mu_1 \mu_2} \quad (6)$$

Fig 3.2 represents the probability of the simultaneous connectivity of the two mobile hosts when varying λ_i and μ_i .

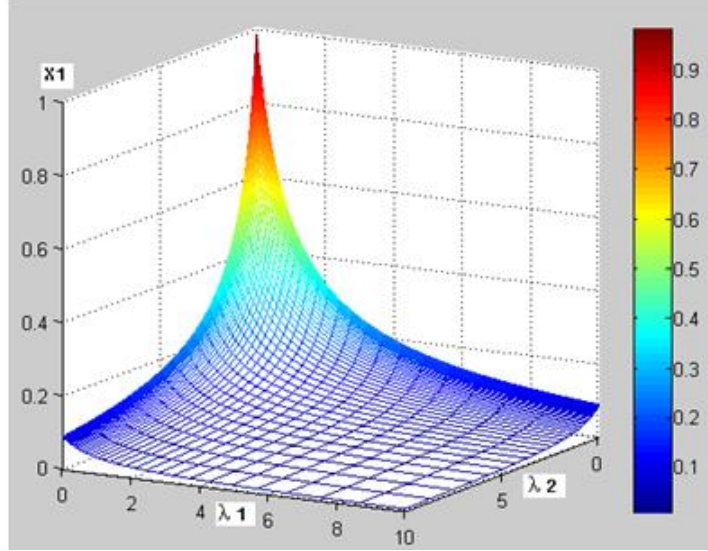


Fig 3.2-a probability of the simultaneous connectivity according to λ

If we suppose μ_1 and μ_2 (average frequencies NOT available for communication) are constants and equal to 1. We verified that, X1: the probability of the state (C1,C2) becomes smaller when the time available for communication decreases (the frequency available for communication λ_1, λ_2 increase).

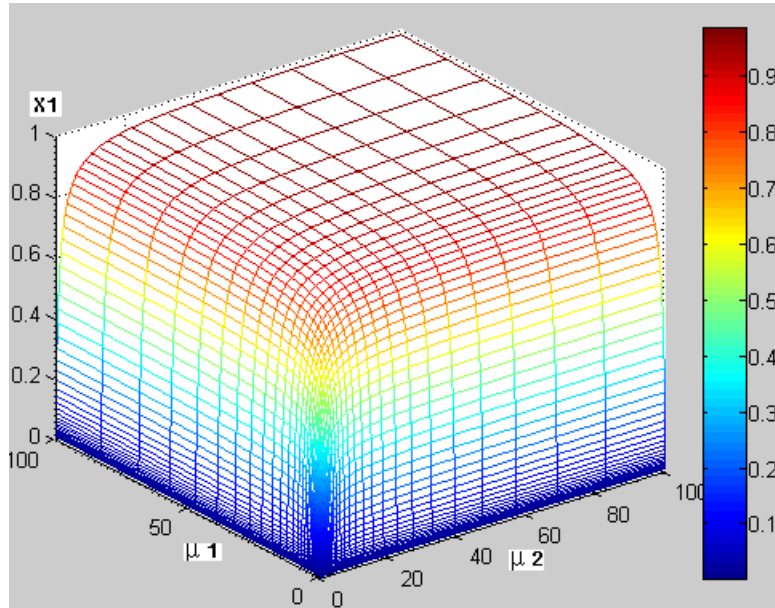


Fig 3.2-b probability of the simultaneous connectivity according to μ

We suppose λ_1 and λ_2 (average frequencies available for communication) are constants and equal to 1. We verified that, X1: the probability of the state (C1,C2) becomes smaller

when the non-available communication period increases (the frequency of non-available for communication μ_1, μ_2 decrease).

This section shows an analytical model that allows estimating the utility of a mobile connection as a function of the mobility behavior of the two mobile nodes. It is worth noting that the utility given by such a Markovian analysis considers the availability of the mobile connection only at the physical layer and does not take into account handover and higher protocols overhead; This is an optimal utility which is considered as an object that our mobility management architecture should try to approach as close as possible but will never be able to reach.

3.4 Location/Address Management

In the context of mobility scenarios as two communicating hosts can potentially move anytime, a fixed point is necessary to conserve mobile nodes' updated location information (i.e. updated dynamic IP address). The Dynamic DNS and the HIP-Rendezvous-server [99] are the two mechanisms we have experimented on which we leveraged for our mobility location management. They are used for supporting a consistent reactivation of the communication in the dynamical context of mobile nodes' migration. Dynamic DNS is more adapted to infrequent migration mobile nodes due to the inefficiency when query load increases, while RVS fits better with frequent migration mobile nodes.

3.4.1 Dynamical DNS

Dynamic DNS allows domain names held by a name server to be updated dynamically. It allows mobile nodes' logical name (i.e. node's name + its original domain name) to be mapped with a varying dynamic IP address. Therefore, this makes it possible for other mobile nodes on the Internet to establish (or reestablish) connections to current mobile node without needing to track mobile node IP address themselves.

Dynamic DNS enables a user to automate the discovery and registration of client's public IP addresses. The client program is executed on a device in the private network. It connects to the service provider's systems and allows those systems to link the discovered public IP address of the home network with a hostname in the domain name system. The hostname is registered within a domain owned by the provider or the customer's own domain name. These services can function by a number of mechanisms such as HTTP service, which is available even in restrictive environments. Dynamic DNS allows the end user to run a fully functional Internet server, despite changing IP addresses.

Security is an important issue in DDNS where TSIG [101] is largely implemented in DDNS systems; TSIG is short for transaction signature and is a crypto-graphical signature that the server can check. If the signature is correct, the server knows that the update either came from the authorized client or from someone who has stolen the secret signing key. TSIG uses a mechanism called HMAC-MD5 [102] to authenticate the sender and message content of the updates. HMAC is a mechanism for message authentication to be used in combination with a cryptographic hash routine MD5.

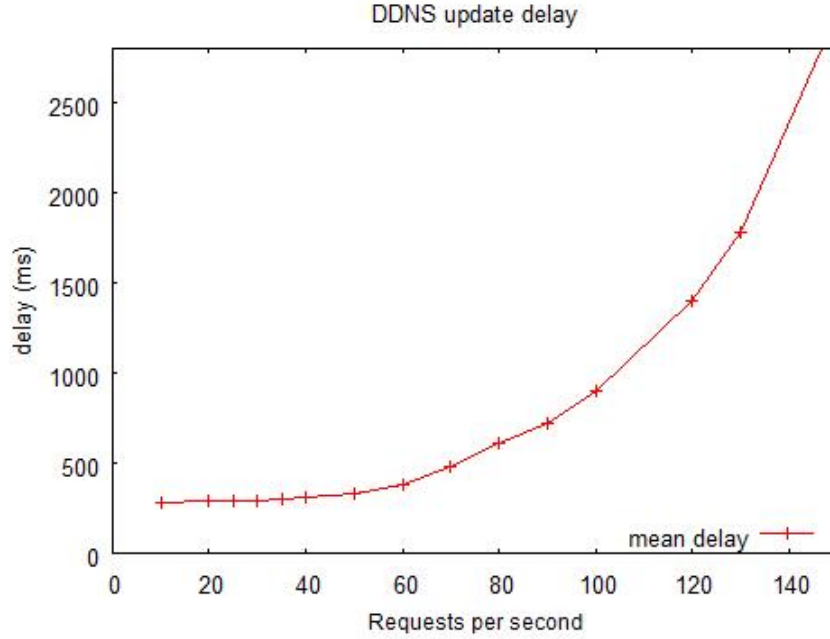


Fig 3.3 DDNS update delay in terms of request frequency

In our architecture, each mobile node has its own registered DNS server which offers the DDNS service. As soon as a mobile node acquires a new IP address in a new access network, it updates its IP address record in its own DNS server using TSIG. We experiment in a scenario where mobile nodes send very frequently the IP address update request and measure the update response delay. We use Simple DNS Plus as the DNS server [103] which includes TSIG generator for security purpose. We coded in Java to capture the response packets and calculate the response delay. Fig 3.3 shows the DDNS update latency in terms of the load of the DDNS server (i.e. N : the number of DDNS registration requests per second). We find that the latency increases rapidly when N is around 60, and the latency increases rapidly towards an incompatible value when N surpasses 120. This performance measurement shows that DDNS overhead is not negligible and that it is not an optimal choice when mobile nodes move frequently. Otherwise, another shortage which can't be ignored is the security issue caused by the sharing of a secret key. However, we considered DDNS as a configurable option in our architecture due to that it doesn't require any other additional infrastructure and doesn't change the current architecture of the Internet.

3.4.2 HIP Rendezvous server (RVS) mechanism

HIP RendezVous Server, which is used for efficient mobility location management purpose, is an extension work of HIP. This subsection gives a description of the HIP, RVS as well as several experimental results.

3.4.2.1 Host Identity Protocol (HIP) and HIP-enabled mobility

In real life, if you have to prove your identity and the asking person is unsure, you show your ID-card. Respectively, if you are asked to give your address, you will give the street address providing your (home) location. If this analogy is used in the Internet, the host identity and location information must be separated from each other. The Host Identity Protocol provides one possible solution for decoupling the location from the identity.

HIP has been developed in the framework the IETF IP working group for a few years. It comes from the need to communicate everywhere anytime security. For this reason,

it was necessary to distinguish between topological locators and identifiers, whose roles are both played by IP address nowadays. Thanks to this new approach, IP addresses act only as locators while host identities are the identifiers themselves. This solution, though, requires adding a new layer, the Identifiers layer, between the transport layer and the IP layer. One of the issues completely defined in HIP is that the Host Identity (HI) is the public key from a public/private key pair. This key can be represented by the Host Identity Tag (HIT), a 128-bit hash of the HI, and has to be globally unique in the whole Internet universe. Another representation of the HI is the Local Scope Identity (LSI) which is 32-bits size and can only be used for local purposes.

The cryptographic nature of the Host Identifiers is the security cornerstone of the new architecture. Each end-point generates exactly one public key pair. The public key of the key pair acts as the Host Identifier. The end-point is supposed to keep the corresponding private key secret and not disclose it to anybody. The use of the public key as the name allows a node to directly check, via an end-to-end authentication procedure, that a party is actually entitled to use its name. Compared to solutions where names and cryptographic keys are separate, the key-oriented naming does not require any external infrastructure to authenticate identity. In other words, no explicit Public Key Infrastructure (PKI) is needed. Since the identity is represented by the public key itself, and since any proper public key authentication protocol can be used to check that a party indeed possesses the private key corresponding to a public key, a proper authentication protocol suffices to verify that the peer indeed is entitled to the name. But how can we trust that the authenticated identity we are communicating with is the one we wanted to? If the mapping between a known domain name and an identity (performed in the Domain Name System - DNS) was hacked, then we could be actually contacting with an unknown site, although security, after the authentication protocol. For this reason, when it is claimed that no explicit PKI is needed, it is as long as a secure protocol is used for the requests to the DNS (e.g., DNSSEC [104]).

Several HIP-enabled mobility mechanisms have been addressed by the current research, compared to the traditional technique like Mobile IPv6, several significant advantages of HIP-enabled mobility are worthy of underlining. Mobile IPv6 with route optimization allows a correspondent host to directly route packets to the mobile host's visited address, rather than through the home network, to improve on latency, robustness, and reduce home network congestion. It does so by maintaining a "binding cache" that matches a mobile node's presently visited address with the permanent home address. This mechanism is an optimization of the basic Mobile IPv6 technique of "reverse tunneling" to the home agent, which must be used whenever a binding has not been established with the correspondent host. While Mobile IPv6 with route optimization uses a Binding Update exchange to notify a peer of an address change, HIP uses a Readdress packet. Mobile IPv6 may be used with or without IPsec, while HIP is tightly integrated with IPsec, although a non-IPsec mode may be possible. Other key differences between HIP-enabled mobility and Mobile IPv6 with route optimization are the following: 1) HIP does not use the concept of a home network, while Mobile IPv6 requires that initial packet exchanges between a mobile host and a correspondent host flow through the home network even if route optimization is subsequently invoked. In HIP, the location of a mobile node is obtained directly from DNS or other directory services, rather than from a home agent; 2) HIP does not incur additional per-packet overhead for carrying Mobile IPv6 Home Address or Routing Header options; 3) Mobile IPv6 can generalize to include subnet mobility (mobility of a router and its attached subnet), while HIP is purely a host-based approach; and 4) HIP inherently secures the readdressing process, while Mobile IPv6 must rely on additional mechanisms.

3.4.2.2 HIP Rendezvous server (RVS)

The HIP protocol employs an infrastructure, the HIP RendezVous Server (RVS), which is used for efficient mobility location management purpose. The rendezvous mechanism is needed if both of the nodes happen to change their address at the same time; it is designed to support double-jump scenarios - simultaneous host-mobility. A RVS functions as a fix point in a network and it keeps track of host mobile nodes.

The clients of an RVS are nodes that use the HIP Registration Protocol [18] to register their HIT->IP address mappings with the RVS. Essentially, the clients of an RVS become reachable at the RVS' IP addresses. Peers can initiate a HIP base exchange with the IP address of the RVS, which will relay this initial communication so that the base exchange may successfully complete. As describes in Fig 3.4-a, an Initiator does not know the IP address of a Responder (mobile node), the Initiator can send the initial message (I1) containing the Responder's HI to a RVS with known IP address. The RVS then relays the I1 message to the Responder, then responder returns reply packet R1 to Initiator to begin establishing the connection.

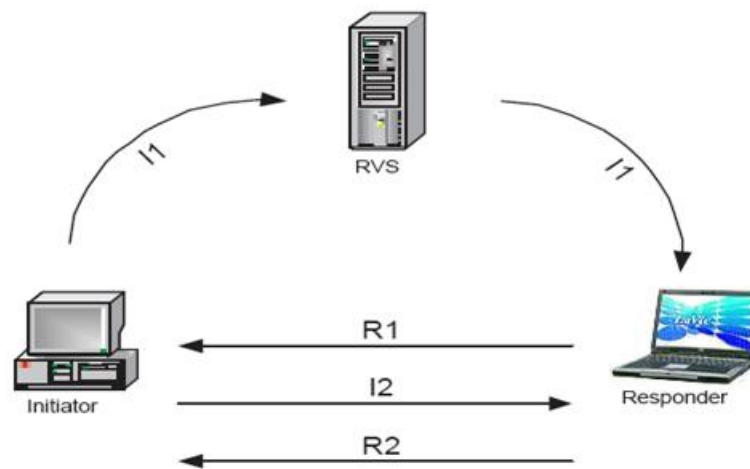


Fig 3.4-a Connection establishment via RVS

But how do we know the IP of the RVS where the peer is registered? For this reason the addition of a new Resource Record (RR) is necessary in the DNS (called IPRVS). As the Fig 3.4-b shows, the complete procedure of establishing the connection includes: 1) the registration of the mobile peer's IP address in its RVS. 2) Registration of the mobile peer's domain name and its RVS IP address in the DNS (FQDNr to HIr, HITr and FQDNrvs mapping). 3) Then, when the mobile node wants to establish the connection with peer node, the host sends the peer's domain name to DNS server (FQDNr). 4) The peer's RVS domain name (FQDNrvs) and the HI/HIT (HIr, HITr) of the peer node will returned to the mobile node by DNS server. 5) With this information, mobile node performs another lookup requesting the IP address of the RVS (IPrvs) with RVS domain name (FQDNrvs) provided in the first lookup. 6) Finally, the mobile node send initial message to RVS server of peer node to start the connection establishment. RVS server relay the initial packet to peer node. When the peer node receives this message, it can then reply directly to initial node without further assistance from RVS because the packet contains the peer's source address.

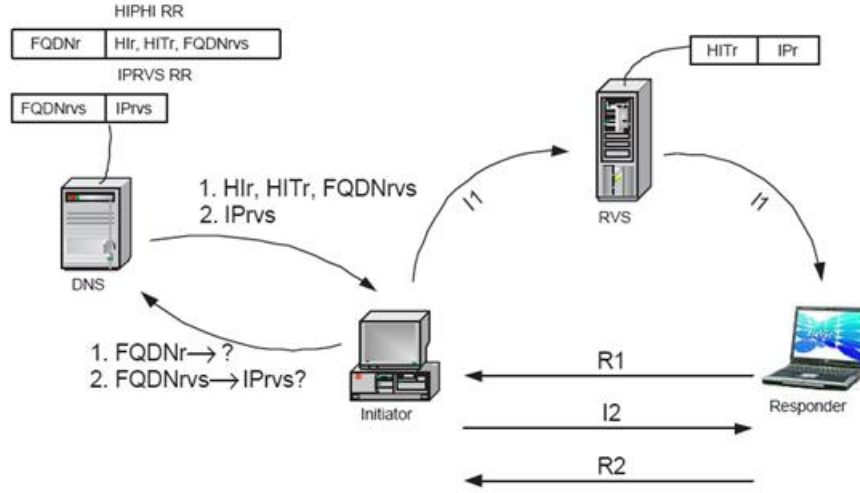


Fig 3.4-b Complete procedure of connection establishing via RVS

3.4.2.2 Experiments

We used the tool OpenHIP [105], which is a free, open-source implementation of the Host Identity Protocol (HIP) in our experiments. OpenHIP is being developed within the Internet Engineering Task Force (IETF) and the Internet Research Task Force (IRTF) to study and experiment with HIP and related protocols. The latest version (V6.0) came out in May, 2009. The source of OpenHIP is in code C, they are compiled and installed under our testbeds based on Ubuntu 8.04 LTS.

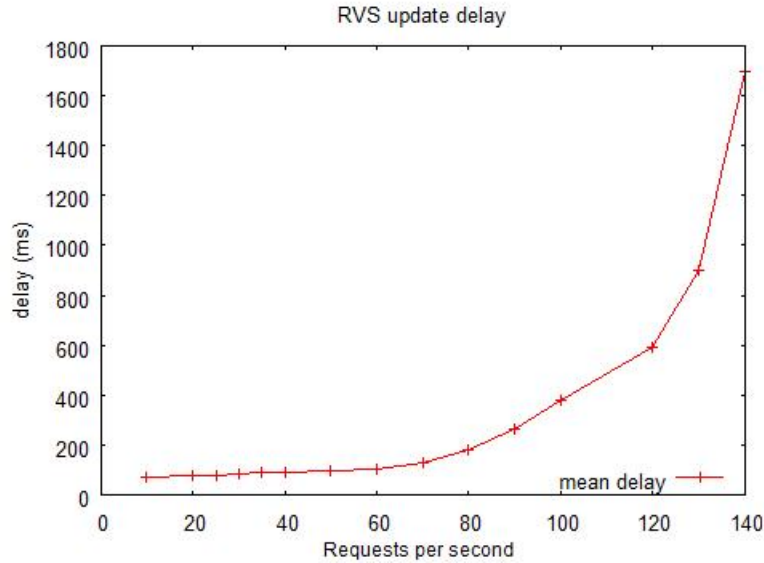


Fig 3.5 RVS update delay in terms of request frequency

In order to measure the registration delay of IP addresses on a RVS, we set a scenario where two mobile nodes are trying to establish the connection via RVS. We suppose that a great number of mobile nodes with high mobility frequency require register their dynamic IP addresses in their RVS, one RVS allows serving for number of mobile nodes.

Several XML configuration of OpenHIP have been carefully configured. The corresponding node's name should be added by ".hip" which tells the hip daemon to use HIP for this domain name. We coded in Java to capture the HIP packets to calculate the response delay.

Fig 3.5 shows the registration of IP latency in function of the load of the RVS server (N : the number of registration requests per second). We find that the latency starts to increase rapidly when N is around 120. Compared to the DDNS, it's observed that IP address update latency is much less than that of DDNS and RVS tolerates a much more frequent update requests.

This section shows two location management based network services involved in our architecture, DDNS and HIP-RVS server, that deliver supports for the end to end management of node mobility and transparent connections in the Internet. Our experiments results show that, compared to DDNS, HIP-RVS performances much better in frequent mobility scenarios. However, the DDNS doesn't require any other additional infrastructure and doesn't change the current architecture of the Internet, which is much easier to be deployed.

3.5 Continuous connection support

HIP protocol supports the separation of identifiers from locations and allows the user not to be aware of the mobility, it offers a solution for mobile users to be easily connected in a dynamic mobile context. However, it can't manage the transport layer directly (on the issues such like congestion, losses, delay, etc.) to support a continuous connection. Based on the location management presented above, we establish a modeling work which allows the mobile nodes to get connected as soon as possible when the communication is available.

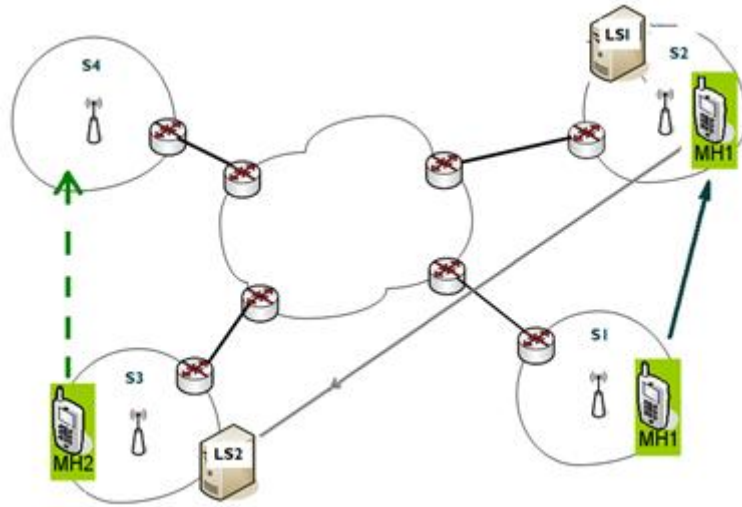


Fig 3.6-a Mobility scenario

In fact, an established network connection can be suspended instead of being cut off if one or both of the mobile nodes migrate(s) and lose its on-going communication. [49] has introduced the notion of connection migration as an extended TCP option to address the continuity of a transport layer connection when only one of the two transport peers moves. In our architecture, with the help of an efficient mobility location service and the data link prediction mechanism to be introduced in section 3.6.1.1, we have extended the notion of connection migration (also called mobile connection) to encompass more general mobility scenarios where the two communicating peers can simultaneously migrate. Besides, for a transport protocol to be able to manage efficiently, it has to satisfy a minimum set of features among which connection management and message numbering are at the first rank. The robust and efficient management of mobile connections that supports the simultaneous migration of the transport peers raises several issues. Indeed, such complex

mobility scenarios expose a protocol for mobile connection management to subtle potential cases of deadlock or address inconsistency.

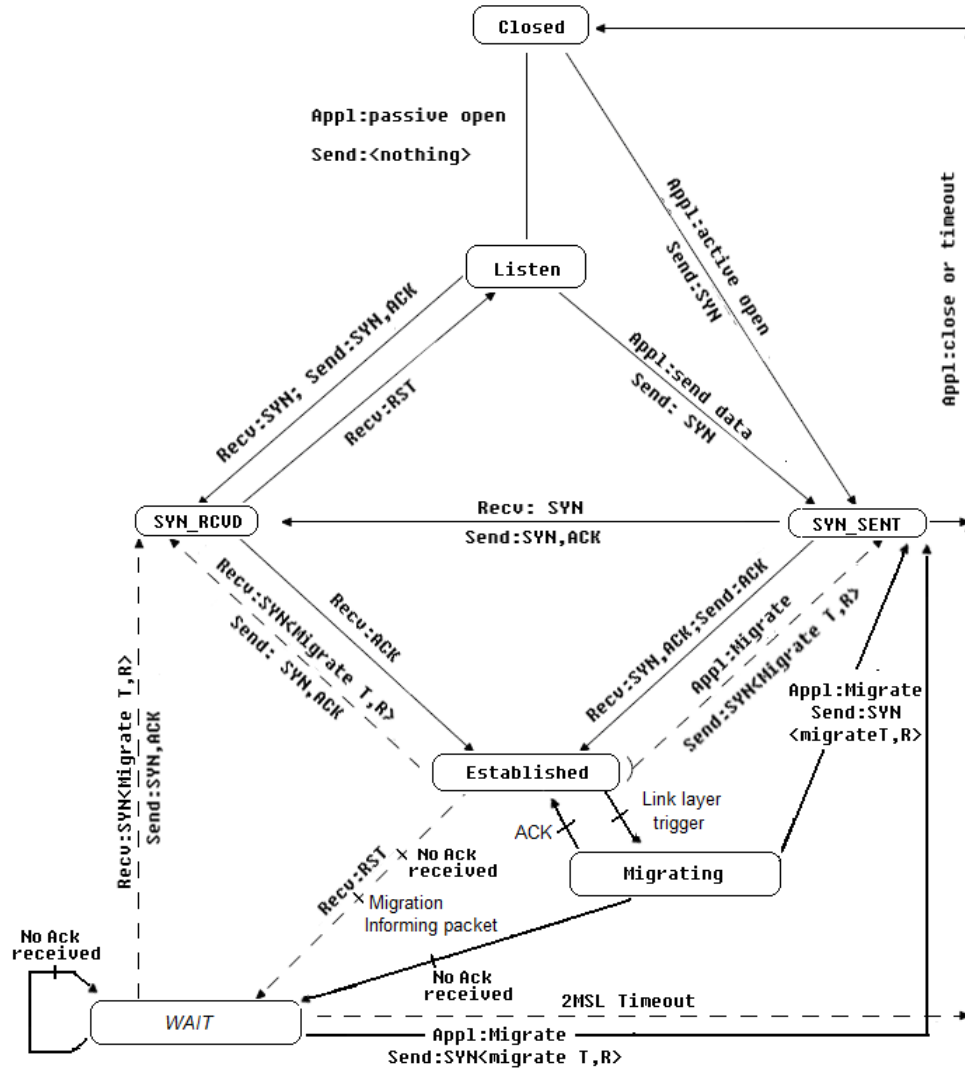


Fig 3.6-b State graph of continuous connection support

The next description gives a big picture of the proposed continuous connection support issue. Fig 3.6-a shows a mobility scenario where the MH1 and MH2 can migrate in a generic way (i.e. two corresponding mobile hosts can migrate simultaneously).

We suppose MH1 and MH2 have initially established a connection. When MH1 starts to move away from subnet-1(S1), based on our data link prediction mechanism (to be presented in section 3.6.1.1), MH1 estimates an optimal moment to send “migration-informing” packets to MH2 and start the pre-DHCP processing (to be presented in section 3.6.3.3.3, There is a greater interest in getting rid of DHCP latency during handover in order to deliver seamless communication). MH1 then get disconnected to the current network and enters into “*Migrating*” state. However, both of them can conserve their current connection state even if the communication between the two nodes is cut off. MH2 then enters into a “*WAIT*” state, and they note down both the sequence number of the last packet they respectively received. When MH1 arrives to its destination access network S2, it activates the pre-reserved IP address as its new IP address and saves its new IP address at once in its own location server LS1, where LS (Location Server) refers to a DDNS or a HIP-RVS Server. In our mobile connection management protocol, a mobile node contacts its corresponding node as soon as it gets re-connected. MH1 then verifies MH2’s current IP address in MH2’s LS because MH2 might also migrate simultaneously

during MH1's migration. Then MH1 pro-actively sends a packet to MH2 to reactive the connection. However, if MH2 is migrating at this moment or the IP information in LS2 has not been updated on time, MH1 will enter into a “**WAIT**” state and wait being woken up by MH2 until MH2 finishes its migration and gets ready for communication. When the communication is recovered, they resume their communication from the last packet they respectively received. Fig 3.6-b shows, take the protocol TCP for example, the extended state graph which represents our continuous connection support mechanism.

A set of mobility scenarios have been modeled and validated in the TURTLE formal language [114], which is an integration of Real Time Lotos and UML. TURTLE is a UML profile dedicated to the modelling and formal validation of real-time systems.

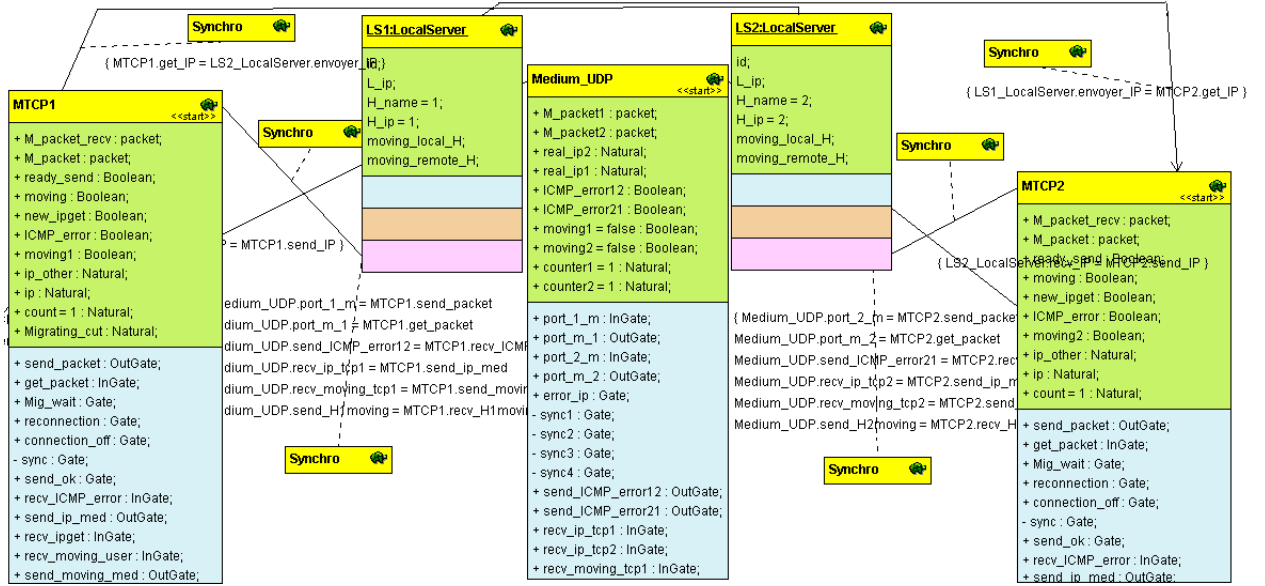


Fig 3.7-a Class diagram

Fig 3.7-a,b,c shows one example of the TURTLE modeling work as well as the simulation results. Fig 3.7-a represents the TURTLE class diagrams which is programmed to describe the interfaces of classes (i.e. Location server, mobility profiles of mobile nodes, etc.), and the relation between the different classes. Interfaces of classes include regular attributes (boolean and natural types) in green, and gates in blue, which are the only way for classes to communicate with each others. Fig 3.7-b represents the modeling work of the mobility scenarios, we suppose that, based on the scenario presented in Fig 3.6-a, the two mobile nodes respectively migrate during the time intervals [150,250], [400,500] and [260,360], [620,650].

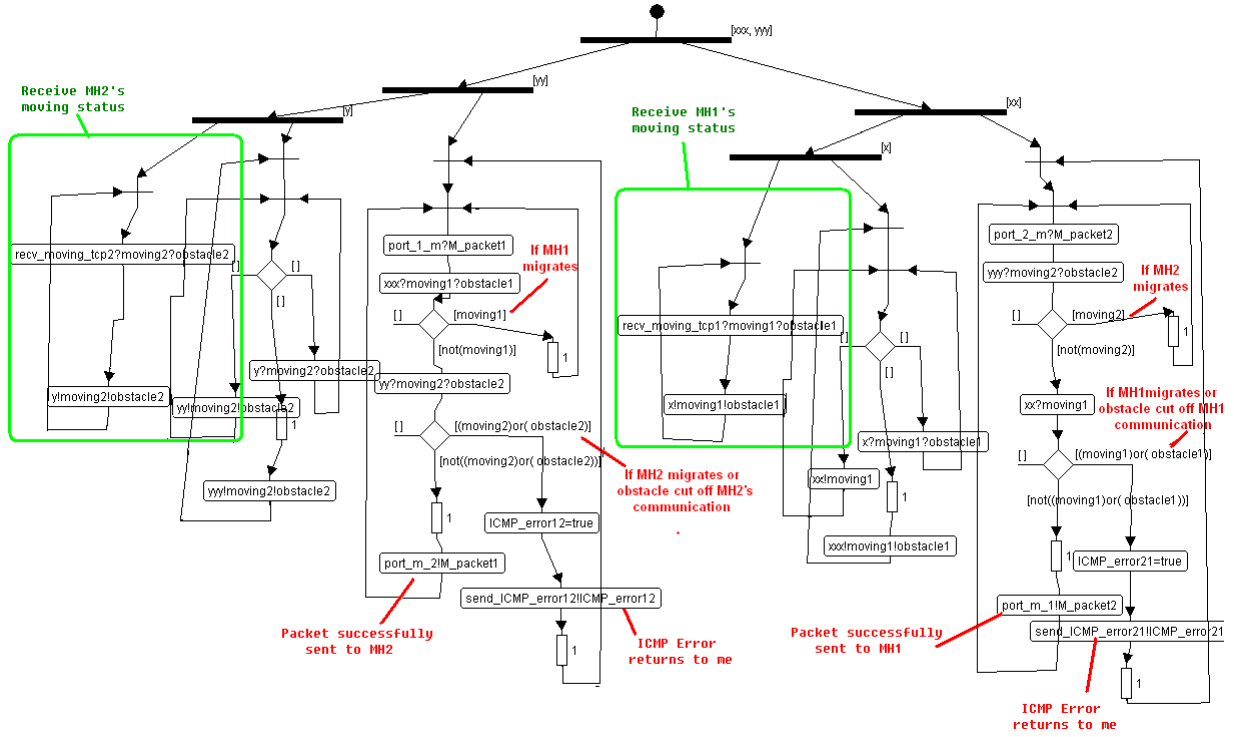


Fig 3.7-b TURTLE modeling work

Fig 3.7-c shows the simulation results, Lines B represents the moment the nodes enter into the **WAIT** state. Lines A represent the time when mobile nodes are able to communicate. The simulation and validation analysis have shown that the mobile nodes can always be “connected as soon as the communication is available”. The formal modelling of our communication architecture allowed us to exhibit initial design choice that induced potential deadlocks and inconsistencies and to finally insure the liveness and consistency of our proposal.

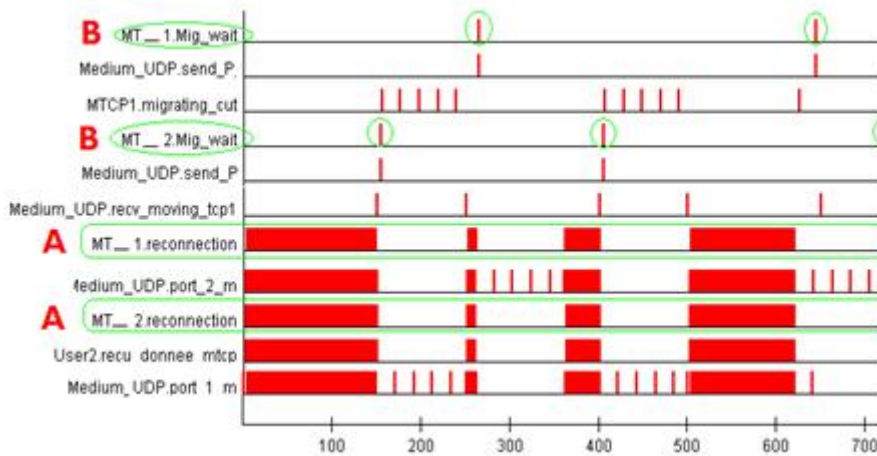


Fig 3.7-c TURTLE simulation results

This section shows, based on the location management, a continuous connection that has been modeled and designed at the transport layer to guarantee the mobile nodes to connect as soon as the communication is available. Although this model allows efficiently offer a continuous communication in a dynamical wireless context, however, it doesn't take into consideration the optimization of the handover delay, which might not be compatible with the real-time multimedia applications, therefore the handover delay optimization will be further studied (section 3.6) in our proposed handover procedure.

3.6 Handover Management

Handover Management is the key part of our cross layer based architecture, which allows minimizing the losses and delay during the handovers. Our proposed Handover Management is based on the cross layer interactions between different layers, and it comprises three principal parts:

Signal analysis function (section 3.6.1), which is based on several intelligent monitoring mechanisms, to efficiently estimate and predict the signal evolutions and mobility behaviors. This part is important for the handover trigger decision which will be presented in section 3.6.3.2.

We integrate the mobility prediction mechanism (section 3.6.2) that is extended from the study [106] in our architecture to support pre-reservation services such like IP address pre-reserved before the handover.

Handover procedure (section 3.6.3) is the core part of the handover management. Firstly, we present the best access point selection policy involved in our architecture and the handover trigger decision which is based on our intelligent signal analysis function. Then, we mainly introduce two proposed optimization mechanisms : pre-DHCP that allows avoiding the layer3 handover delay, and a RTT estimation based mechanism to further minimize the handover delay and support a seamless communications.

3.6.1 Signal analysis function

The signal analysis function allows an intelligent collection and analysis of the signal information (SNR, Rx power, MAC layer parameters, etc.) from lower layers. It is composed by 2 principal parts: 1) Signal Evolution Prediction : this function integrates an intelligent signal filtration process to synthesize the collected signal information and estimates the future evolution of signals. Cooperating with the handover decision function (see section 3.6.3.2), it helps to estimate an optimal moment to start the handover procedure. 2) Dynamic monitoring interval mechanism: this function manages the optimal interval to collect the signal information from the associated and surrounding visible BS or APs. For instance, the frequency of signal information collection can become higher when MN moves to the edge of the BS or AP, since the collected information is more useful in the critical points for the procedure of handover management.

3.6.1.1 Signal evolution prediction

Signal evolution prediction is based on data link layer information, it makes our mobile node aware of its relative position and speed in its current wireless access network. The proposed location estimation technique makes it possible to predict when a mobile node will get disconnected from its current wireless access network according to the signal strength evolution. A calibration based formula is initially defined by a set of experiments to represent the relationship between SNR (Signal Noise Ratio in dB) and the relative distance (in meter) for wireless network (Note that a long relative distance can represent a number of obstacles). Basically, the mobile node measures SNR periodically every t seconds, with a low pass filter integrated in the system to smooth the SNR variation, the evolution of SNR can be translated into the evolution of the relative distance between node and the base station, and a relative speed can be estimated in real time. If we define a threshold which supports the critical wireless communication, our mechanism allows the mobile node being aware of its relative position in base station coverage and estimating the time when the mobile node is about to get disconnected according to its current relative positive speed.

In a practical view, for every wireless access point, a calibration based on a formula such as $SNR(dB) = A - B * \log_{10}(\text{distance})$ can be initially done in order to establish the relationship between SNR (Signal Noise Ratio in dB) and the distance (in meter) from the access point of the current wireless access network (coefficients A and B vary according to the frequency of emitting signal). The basic idea of this estimation mechanism is to measure periodically the SNR (i.e. every n second(s)) so that we can calculate the relative speed V of the mobile node as:

$$V = [10^{(A-SNR[i])/B} - 10^{(A-SNR[i-1])/B}] / n \quad (7)$$

where $SNR[i]$ is the current measurement and $SNR[i-1]$ is the previous one. $V > 0$ means the mobile node is moving away from the AP, $V < 0$ means the mobile node is moving towards the AP. Note that this equivalent speed does not represent the real speed of the mobile node specially in the indoor case. For example, a sudden large speed variation corresponds to a brutal fall of the SNR, which can be caused by the appearance of obstacles between the mobile node and the access point. Therefore, the rough estimator must be enhanced with filtering techniques that aim to suppress the outlier.

If we define $SNR_threshold$ as the critical threshold under which wireless communication cannot be supported anymore between the MN and the AP: then we can estimate the relative time (T) when the mobile node will get disconnected according to its current relative positive speed.

$$T = \frac{10^{(A-SNR_threshold)/B} - 10^{(A-SNR[i])/B}}{V} = n * \frac{10^{(A-SNR_threshold)/B} - 10^{(A-SNR[i])/B}}{10^{(A-SNR[i])/B} - 10^{(A-SNR[i-1])/B}} \quad (8)$$

However, the reality is somewhat different from this theoretical analysis. According to our experimental studies for an indoor case with obstacles wrapped up by Fig 3.8 (where 3 APs are located at positions 90, 210 and 380 respectively), we find that the measured SNR (in red line) is not always stably decreasing while the distance between the MN and the AP increases (that is, we observed a quite large variance of the measured SNR around the estimated one).

Therefore, these experimental results lead us to apply a low pass filter to the measured SNR in order to smooth its variation. This loss pass filter is based on an exponential moving average of the processed SNR given by the following formula:

$$SNR[i]^* = K1 * SNR[i] + (1 - K1) * SNR[i-1]^* \quad (0 < K1 < 1.0) \quad (9)$$

The so resulting $SNR[i]^*$ estimator is represented by the green line in Fig 3.8. Note that the choice of $K1$ has a significant influence on the performance of our mechanism, we identified an optimal value of $K1=0.6$ in our experimental tests.

The above work allows estimating the future signal variation as well as mobile nodes' mobility evolution, and then deducing when a mobile node is about to migrate out of the current wireless coverage. Based on this prediction result, the mobile node is able to inform its peer mobile node the disconnection message in advance before its migration, which is of a great help on the mobility status awareness for the mobile nodes.

While with the traditional methods, a mobile node cannot be aware of its peer node's migration at once in the context of mobility. We have assessed various predictors based on traditional ICMP error packets or ACK message timeouts. In practice, these mechanisms induce feedback loops with a magnitude of several RTTs that can potentially entail lengthy handovers and discontinuities on multimedia streams. Moreover these network or transport layer mechanisms induce a waste of energy for a "blind" node that continues sending packets in vain while its corresponding node is disconnected or in conditions where the signal gets too weak to support the communication.

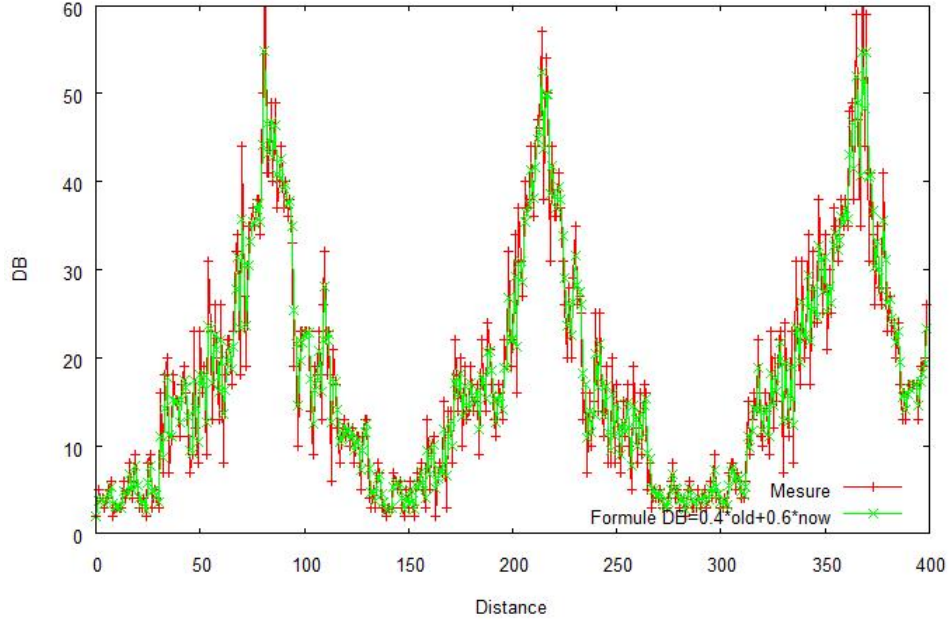


Fig 3.8 measured and resulting SNR

In our proposed approach, a mobile node (MN1) can be immediately informed as soon as its corresponding node (MN2) is about to migrate. Such a “handover in progress” message avoids MN1 sending blindly while MN2 is out of connection. According to its auto-estimated relative speed and its estimation of the RTT between the two host, MN2 can automatically find a threshold (in terms of SNR) from which it starts to send informing packets to MN1 in order to assure with a given probability that MN1 will be informed before MN2’s disconnection.

This signal evolution prediction function has been implemented in our demonstration tool which will be presented in chapter 6. In our experimental tests, we used Dlink 614+ as wireless routers. We supposed a 95dB constant noise level and an 8dB critical communication threshold. We set T, the monitoring interval, to 2 seconds. We experimentally found that the coefficients A and B, associated to the formula which describes the SNR (dB) in terms of distance for the Dlink 614+ wireless routers, are respectively A=60 and B=20.5, that is:

$$SNR(dB) = 60 - 20.5 * \log_{10}(distance) \quad (10)$$

Fig 3.9 represents the approximate SNR formula and the experimental measurements of SNR in terms of distance between mobile node and access point.

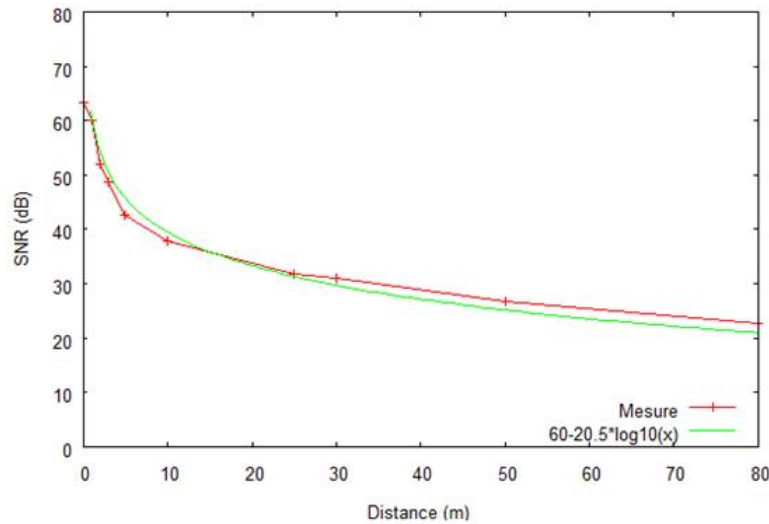


Fig 3.9 SNR in terms of distance

Furthermore, this mechanism can also cooperate with handover decision function to estimate an optimal moment to process Pre-DHCP mechanism (to be introduced in section 3.6.3.3.3), which allows significantly minimizing the handover delay and support seamless communications.

3.6.1.2 Dynamic Monitoring Interval function

SNR monitoring is integrated in our cross layer based architecture to offer low layer information to upper layers for handover management (i.e. new AP discovery, disconnection prediction as described in last section, etc). In our original experiments, we forced mobile nodes to monitor SNR and low layer information periodically every T seconds. We found that such approach is quite energy costly for an energy limited mobile terminal. Indeed, when a mobile node is close to its attached AP where it can receive a good signal, it doesn't make sense to measure frequently the SNR to be aware of the disconnection information, or to discover the new AP and prepare for the handover procedure. So in this case, the frequency of signal monitoring should be lower. In contrast, when the mobile node is moving to the edge of the coverage where the SNR is poor, the frequency of signal monitoring should increase to be aware of the disconnection information more precisely (i.e. predict the disconnection information) and discover new AP as soon as possible.

- ***Signal Monitoring Frequency in the attached AP coverage***

We denote $F_{\min}$, the minimal signal monitoring frequency which corresponds to the case when mobile node is very close to the AP and receives a most optimal SNR ($SNR_{\max}$) in the coverage. We also denote, $F_{\max}$, the maximum signal monitoring frequency that the mobile node supports. We have the signal monitoring frequency F :

$$F = F_{\min} + (F_{\max} - F_{\min}) * (1 - SNR_{\text{curr}} / SNR_{\max})$$

Where SNR_{curr} is the current monitored SNR. We calculate the monitoring period $T_1 = 1/F$.

- ***Signal Monitoring Interval to discover new AP***

We denote $SNR_{\text{new_ap}}$, the new AP's SNR monitored by the MN (the new AP represents the AP the mobile node will handover to; if multiple neighbor APs exist, the new AP with the highest SNR is taken into consideration). SNR_{th} is the threshold to support critical communication. $T_{\min}$ is the minimal monitoring period and $T_{\max}$ is the maximum monitoring period. We have the following algorithm to calculate the monitoring period T_2

$$\begin{aligned} & \text{if } (SNR_{\text{new_ap}} > SNR_{\text{th}}) \\ & \quad T_2 = \max(T_{\min}, T_{\max} + (T_{\max} - T_{\min}) * (1 - SNR_{\text{new_ap}} / SNR_{\text{th}})); \\ & \text{else} \\ & \quad T_2 = T_{\max}; \end{aligned}$$

Then, we have the Signal Monitoring period $T = \min(T_1, T_2)$.

Calculation of T_2 can be disabled when the mobile nodes are not willing to process the handover. In our experimental tests, we set $F_{\min} = 0.2$, $F_{\max} = 1$; $T_{\min} = 1s$, $T_{\max} = 5s$, and $SNR_{\text{th}} = 8dB$, $SNR_{\max} = 65dB$ for the access router Dlink 614+. This Dynamic Monitoring Interval function is proved to more efficient for the handover management and the wireless context-awareness, furthermore, if we suppose an uniform mobility model (walking speed) across the wireless LAN 802.11b zones, compared to the default mode where signal monitoring is done periodically every 2 seconds, 36% (average in indoor environment) and 58% (average in open air environment) of the energy for signal monitoring can be saved according to the experimental results.

3.6.2 Mobility prediction

Context awareness is essential in the evolution of today's wireless networks; one of its important objects is the mobility prediction, which leads to a more efficient planning and management of the wireless network and pre-reservation based services (e.g. adaptive buffering, network resource pre-reserving and DHCP IP address pre-reserving).

Nowadays, mobility prediction algorithms are mainly based on historical records such as aggregate mobility and handoffs history and locations. They can be based on the mobile nodes' profiles, habits, speed, location, or direction and incorporate with geographic maps with identifiable landmark objects. In our architecture, mobility prediction is based on a K-order Markov chain, concretely speaking, the mobility prediction is based on the records of the K last visited attached access zones. In [106], authors have given a detailed analysis on the Markov based mobility prediction for different order K. It's observed in the experimental tests that a higher Markov order which starts from K=2 results in poorer estimation. The reason is that, although the predictors use more information in the prediction process, they are also more likely to encounter a context that has not been seen before, and thus be unable to make a prediction. A missing prediction is not a correct prediction, and these unpredicted moves bring down the accuracy of the higher-order predictors. They conclude that a 2-order Markov based mobility prediction models is the most optimal estimator, which is also implemented in our architecture.

Furthermore, the Geo-Location function (GL) which is to be presented in chapter 6 allows enhancing the precision of the mobility prediction, especially in the cases where multiple WLAN zones overlap together.

In our experimental tests, we set 8 access points uniformly in the corridor, offices and class hall on the first floor of the building F in campus Jolimont of ISAE. We extended the work from the study of [106], a four-dimension database which represents the mapping between the historical probability (P) of entering into a wireless zone (Z3) and the last two visited zones (Z2 and Z1) has been integrated in our implementation. In our prototype which is to be presented in chapter 6, we found that, with the integration of Geo-Location function, the prediction precision significantly improves from 73% to 93%, which will do a great of help on the pre-reservation based services such as Pre-DHCP that is to be introduced in the next section.

3.6.3 The Handover Procedure

Handover is one of the principal parts in our mobility management architecture, which allows supporting multimedia seamless communication. The handover procedure is composed by 3 sub-functions: **1) Best AP selection function:** this function allows choosing an optimal AP according to the detection of environment signal distribution. **2) Handover trigger function:** this sub-function determines the optimal moment to start handover procedure based on the Signal Evolution Prediction function. **3) Handover delay optimization:** this functions comprises two delay optimization levels: Level 1 focuses on minimizing HO delay with an in-advance address reservation mechanism (Pre-DHCP) which pre-acquire an IP address of the next subnet via the current AR. Level 2 is based on an intelligent RTT estimation based mechanism to further reduce the handover delay and enhance the handover efficiency.

3.6.3.1 Best access point selection

If multiple neighbor APs are available to offer services to the moving mobile node, the node should select the best AP which allows offering a stable and guaranteed QoS according to the link layer parameters and the applications required by the mobile node. Many

studies have been investigated to find an efficient way of select best APs according to several parameters such like SNR, admission control issues, QoS charges of potential APs. However, all of these proposals require additional infrastructure or data loads, and since the principal of our thesis does not focus on the admission control issues, therefore the best AP is selected only according to the strength and variation of monitored signal in our experimental tests, that is, the AP with better SNR is chosen, if multiple APs exist with the same monitored SNR, the one with less signal variation is selected as candidate AP.

Our future work will further focus on a QoS aware mobility architecture, without entailing additional infrastructures, which not only takes into consideration the low layer signal information, but also the available QoS offered by the candidate access networks to support an intelligent AP selection.

3.6.3.2 Handover trigger decision

In principle, handovers can be categorized as imperative and alternative handovers according to initiation reasons. Handovers due to low link quality are imperative; because both the handover decision and execution have to be done rapidly in order to maintain on-going connections. Primarily, the signal-to-noise ratio (SNR) monitored from the attached access point and neighboring access points are used for handover decisions. In the other hand, handovers that are sometimes used to provide a mobile user with better performance or to meet a particular preference, can be considered as alternative handovers, for example, a mobile node can alternatively associate with another AP with better SNR which offers a higher bandwidth, lower link errors and enhanced QoS.

This handover decision function is able to find an optimal moment from which the mobile node starts the pre-reserved service for handover procedure such as pre-DHCP which is processed via the current access router to acquire new IP addresses for the potential future wireless access networks to visit. According to the number of potential candidate APs, it cooperates with the Signal Evolution Prediction function to decide an optimum moment to start Pre-DHCP processes and guarantee that these DHCP pre-reservations will be successfully finished before the signal gets too poor or the current node gets disconnected. Concretely, we denote $K * T1$, the overall reserved duration of DHCP pre-reservations procedure in millisecond, where $T1$ represents the delay dedicated for the Pre-DHCP procedure and K is a coefficient which is bigger than 1. Then, according to the equations in section 3.6.1.1 which calculate the time that a mobile node will get disconnected, we can estimate an optimal moment ($K * T1$ milliseconds before the current node's disconnection) to start the Pre-DHCP procedure, in order to guarantee the Pre-DHCP procedure should be finished before the node's disconnection from the current AP.

For the alternative handover, if an AP with optimal SNR is monitored, it can be considered as a potential candidate AP for mobile node to visit in order to enhance the transmission efficiency, in this case, Pre-DHCP procedure can also be triggered to acquire a temporary new IP addresses form this AP with short lease duration.

3.6.3.3 Handover delay optimization

This section will firstly give a brief introduction on the handover latency in the context of WLAN, then we will focus on two proposed mechanisms: Pre-DHCP, which is based on a pre-reservation service, to avoid the layer3 handover delay; and a RTT estimation based protocol to further minimize the handover delay and support a seamless communications.

3.6.3.3.1 Handover delay

The handoff procedure latency can be divided into 4 parts: AP scanning delay, Authentication Delay, Re-association Delay, Network layer delay (i.e. DHCP procedure).

- AP scanning delay: When a mobile node is moving away from the AP and the SNR drops under a communication threshold. The mobile node should process a AP scanning procedure to discovery new available AP to associate. Scanning can be accomplished either in passive or active mode. In passive scan mode, the mobile node listens to the wireless medium for beacon frames. Beacon frames provide a combination of timing and advertising information to the mobile nodes. Using the information obtained from beacon frames the mobile node can select to join an AP. During this scanning mode the mobile node listens to each channel of the physical medium to try and locate an AP. Active scanning involves transmission of probe request frames by a mobile node on each channel (one by one) and processing of the received probe responses from the APs. Normally, there are 11 channels that should be scanned, if no response has been received by minChannelTime for the current scanned channel (This value is device dependent [115]), then the mobile node scans next channel. The number of “minChannelTime timeout” can be up to 11 in “bad luck”, which can be dramatically increasing the handover latency.

- Authentication delay: When an AP is selected after the scanning procedure, the mobile node attempts to authenticate to the AP. The network access server and AAA Server use the Authenticator to demonstrate that both know the same secret password. Authentication consists of two or four consecutive frames depending on the authentication method used by the AP.

- Association Delay: Upon successful authentication process, the mobile node sends a association request frame to the AP and receives a association response frame and completes the link layer handoff.

- Network layer delay: When the mobile node enters into new access point coverage, the mobile node will request a new IP address and the new network configuration information from the DHCP server. The DHCP latency is considered as one of the main delay in the handover procedure, which reaches up to 2.6 seconds in our experimental tests.

Table 3.1 shows, in the context of IEEE 802.11, the different layer latency during the handover procedure.

3.6.3.3.2 Mobility prediction based selective scanning

The AP scanning procedure can be quite long if up to 11 channels are scanned with “min-ChannelTime timeout” before finding an available AP. The main reason is due to mobile node’s blindness to the adjacent APs’ channel information. The solution integrated in our architecture makes the scanning procedure more pertinent to the adjacent available APs, that is, during the layer 2 handover, the mobile node can only scan the channels corresponding to the APs that the mobile node will potential visit (neighboring APs), which is based on the mobility prediction and significantly reduces the MAC layer handover latency.

We suppose each AP has a table which records the channels information of the neighboring APs. The mobility prediction mechanism offers the mobile node a list of the APs that includes the potential targets the node may process the handover. The mobile node can request the channel information for these potential APs via the attached AP; therefore, the mobile node can only scan these channels pertinently during the handover instead of 11 channels, which is a significant improvement on the handover performance for some real time application such as VoIP.

Latency Budget

March 2004

doc.: IEEE 802.11-04/0377r1

Layer	Item	IPv4 Best Case (ms)	IPv4 Worst Case (ms)	IPv6 Best Case (ms)	IPv6 Worst Case (ms)
L2	802.11 scan (passive)	0 (cached)	1 sec (wait for Beacon)	0 (cached)	1 sec (wait for Beacon)
L2	802.11 scan (active)	20	300	20	300
L2	802.11 assoc	20	80	20	80
L2	802.1X authentication (full)	750	1200	750	1200
L2	802.1X Fast resume	150	300	150	300
L2	Fast handoff (4-way handshake only)	10	80	10	80
L3	DHCPv4 (6to4 scenario only)	200	500	0	0
L3	IPv4 DAD	0 (DNA)	3000	0	0
L3	Initial RS/RA	0	0	5	10
L3	Wait for more RAs	0	0	0	1500
L3	IPv6 DAD	0	0	0 (Optimistic DAD)	1000
L3	MN-HA BU	0	200	0	200
L3	MN-CN BU	100	200	100	200
L4	TCP adjustment	0	Varies	0	Varies

Table 3.1 802.11 handover latency budget

3.6.3.3.3 First Optimization Level: Pre-DHCP

According to our experimental tests, we found that the delay to acquire an IP address is often significantly longer for a wireless node than for a wired one, and is the most important latency observed during the handover procedure. The main reason for this relatively high DHCP latency is that when the mobile node successfully associates with an AP, it may still be at the edge of the area covered by the AP. The strength of the signal may be very weak, so some link-layer frames corresponding to the DHCP request or its response may be lost. Fig 3.10 summarize the results of our measurements (for a DLink 614+ wireless router) of the mean DHCP processing latency (ms) in function of the SNR (dB). These measurements show that DHCP latency can be important at the edge of a wireless access network when SNR is poor.

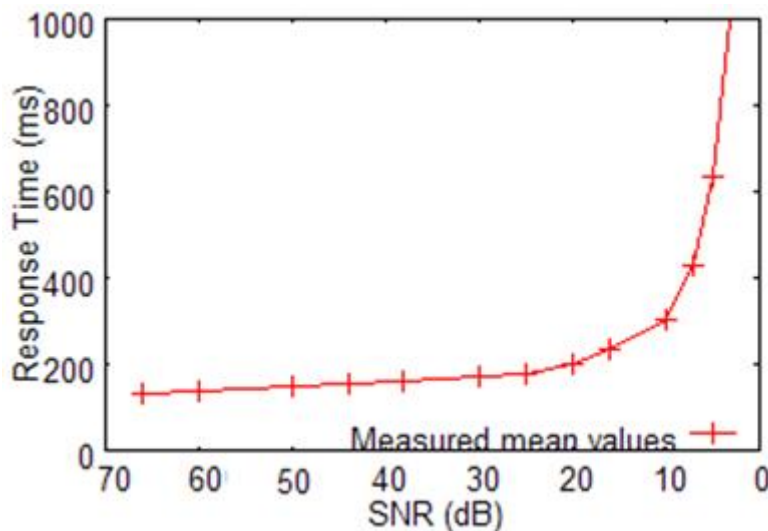


Fig 3.10 DHCP delay in terms of SNR

Therefore, there is a great interest in processing the DHCP procedure before the handover, when the mobile node is still in its previous subnet, instead of doing it when the MN arrives at the edge of its new access network where the signal is weak. However, the moment when this pre-reservation of DHCP addresses in the previous access network starts is crucial, it should be triggered neither too early (i.e. IP addresses could be reserved in vain) nor too late (i.e. signal gets too weak or no enough time to finish the DHCP processing before getting disconnected). The Signal Evolution Prediction function (section 3.6.1.1) can help to find an optimal moment from which the mobile node starts its pre-DHCP processing via the current access router to acquire new IP addresses for the potential future wireless access networks to visit. This mechanism can guarantee that these DHCP pre-reservations will be successfully finished before the signal gets too poor or the current node gets disconnected.

DHCP is a protocol that allows assigning a new IP address and offering related network parameters for the node, it comprises four steps: DHCPDISCOVER, DHCPOFFER, DHCPREQUEST, DHCPACK. DHCPDISCOVER message is firstly broadcast by a DHCP client; DHCP server then responds a DHCPOFFER message including the candidate IP address from an address pool as well as a list of DHCP configuration parameters; the client confirms the correctness of the previously allocated IP address by sending a DHCPREQUEST message; Finally, a DHCPACK message is sent from server to confirm the offered options and validate the assigned IP address.

However, since DHCP acquirement messages are broadcast by the client from the source network, DHCP server that situates in destination subnet is not able to assign an IP address for the client. DHCP Relay Agent is developed to solve this problem, it helps to direct DHCP messages between the DHCP server and the DHCP client in different subnets. Normally, the Relay Agent is implemented on the router between the DHCP client and DHCP server subnets. When DHCPDISCOVER message is captured by DHCP Relay Agent, it sets the Relay IP Address field (also known as GIADDR: the Gateway IP Address field) with the IP address of the interface on which this message is received, and changes the source IP address of the DHCPDISCOVER message to the gateway IP address. Then, the DHCP Relay Agent sends this DHCPDISCOVER message as a unicasted IP packet to the destination DHCP server. When responding to the DHCP client's request for an IP address, the DHCP server uses the Relay IP Address field in the following 2 ways: 1) The Relay IP Address and the subnet masks of the server's configured scopes are compared through a logical AND comparison to find a scope whose network ID matches the network ID of the Relay IP Address. When a match is found, the DHCP server allocates an IP address from that scope. 2) When sending the offer back to the client, the DHCP server sends the DHCPOFFER message to the Relay IP Address as the destination IP address. Once received by the DHCP Relay Agent, the Relay IP Address is used to determine the interface to which the DHCPOFFER message is to be forwarded. It then forwards the DHCPOFFER message to the client using the offered IP address as the destination IP address and the client's MAC address as the destination MAC address. Similarly, the relay agent allows relaying the DHCPREQUEST and DHCPACK packets to finalize the DHCP procedure.

Our proposal Pre-DHCP is based on DHCP relay agent technique, the principle of our solution is that the mobile node pretends acting as a role of relay agent and acquires a new IP address of the neighboring subnets. The mobile node is firstly set in "relay agent mode". It should configure several parameters in DHCPDISCOVER packet; it set the hop number field to N ($N > 0$ to go across the routers), it will "relay" an IP address request for itself; it fills in its MAC address and sets the relay gateway address (GIADDR field) to its current local address. We add an extended option in the current DHCP protocol; we denote it

“Pre_DHCP” option. The DHCPDISCOVER packet is then unicast to the neighboring Access Router (with DHCP server). For the modification on the side of DHCP server in the neighboring AR, once the DHCP server receives the DHCP packet with this option, DHCP server will either disable the “AND” comparison operation or change the subnet mask in order to assign an IP address which belongs to the subnet where the MN is going to. The DHCP server in neighboring AR then generates DHCPOFFER packet including the assigned new IP address for the mobile node (corresponding the MAC address in DHCPDISCOVER packet), and the DHCPOFFER packet (with new IP address and network parameters) is sent back to GIADDR address which is also the mobile node’s current IP address. Similarly, after the exchange of DHCPREQUEST and DHCPACK packets, the IP address is reserved for the mobile node. When the mobile node enters into the neighboring subnet, it just needs to process the lower layer handover and active the reserved IP address to start the upper layer communication, which avoids the layer-3 handover duration. However, we should pay attention on 3 possible aspects in order to make it work efficiently: 1) some related security mechanisms are required (i.e. secret key encryption and exchange) to guarantee that the DHCP server assign an IP address to legal mobile nodes. 2) Since the DHCPDISCOVER packet should be unicast to the neighboring AR(s) to pre-reserve the IP address, the mobility prediction is a crucial part for the mobile node to precisely estimate the future AR that the mobile is about to visit. 3) The pre-reserved IP address should be reserved as in soft states associated with short lease duration in order to avoid wasting IP address resources if mobility prediction gives error estimation and IP addresses are reserved for vain. In the context of our mobility management architecture, we haven’t taken into consider the first aspect, the security issues.

3.6.3.3.4 Validation of Pre-DHCP

In our first experimental tests, we use Dlink-614+ as access routers with DHCP server integrated. Fig 3.11 represents our scenario setting, all the ARs are attached to AR1, there network ID are set respectively to 192.168.N.0, N=1,2,3,4. The mobile node is initially in the coverage of AR4, AR1,2,3,5 are identified as the potential ARs that the mobile node is about to visit. The mobile node process the Pre-DHCP procedure to pre-reserve IP addresses for the four candidate subnets before getting out of the AR4 coverage. The object of our experiments is to verify the feasibility of our proposal and measure the procedure duration of Pre-DHCP according to the number of candidate subnets from which the mobile node demands pre-reserve IP addresses simultaneously. Our experimental tests are based on the java package JDHCP [107], we made some modification to fit our proposal, we record the timestamps of DHCPDISCOVER and DHCPACK packets to measure the duration of Pre-DHCP procedure.

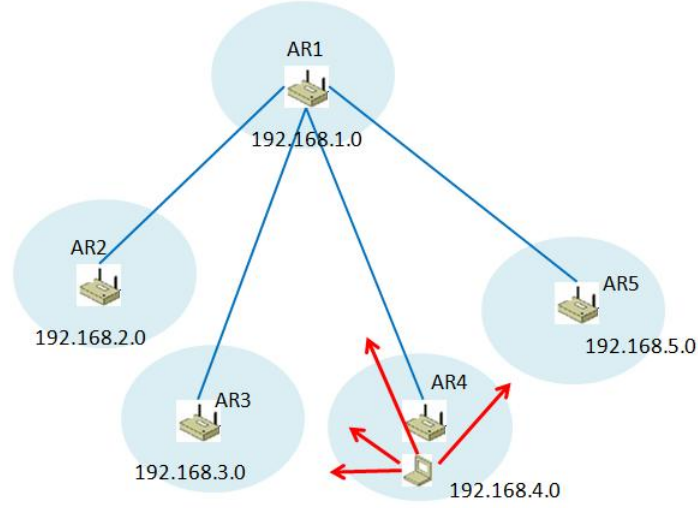


Fig 3.11 Pre-DHCP tests scenario

Fig 3.12 represents the Pre-DHCP latency according to the number of candidate subnets from which the mobile node demands pre-reserve IP addresses simultaneously. We don't find much difference if number of candidate subnets changes, that is because the DHCP packets transmission durations are much lower than processing delay.

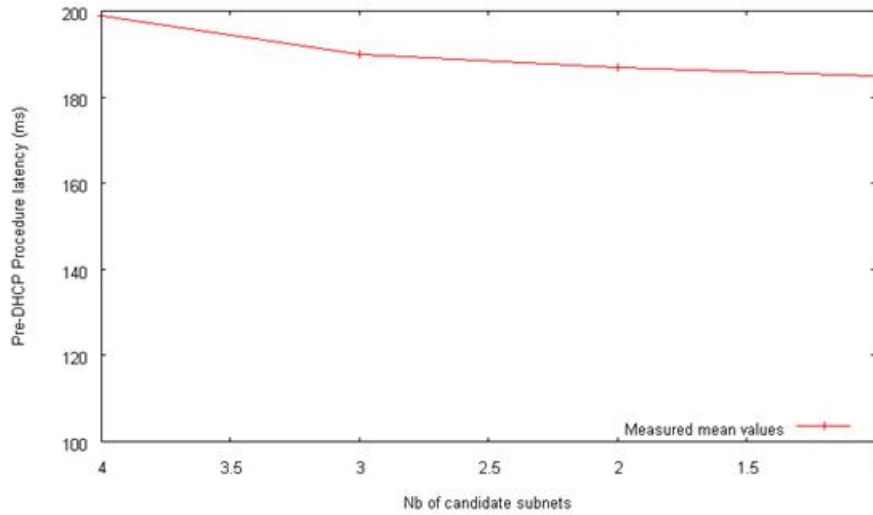


Fig 3.12 Pre-DHCP latency

However, according to experimental tests, we find surprisingly that even if we don't have any modifications on the DHCP server integrated in several AR models such as DLink 614+, the DHCP server in the neighboring AR can still return the IP address of the neighboring subnet, the reason is that it doesn't have a IP address pool corresponding to the IP addresses of previous subnet (not in DHCP agent mode), in this case, after the "AND" comparison operation, DHCP server in the neighboring AR can't find a IP address scope whose network ID matches the previous subnet, and then it assigns an IP address of the neighboring subnet for the mobile node.

3.6.3.3.5 Second Optimization Level: RTT estimation based mechanism for Handover delay optimization

Pre-DHCP and selective scanning introduced in the previous section allow avoiding the layer3 delay and reducing the layer2 delay during the handover, in this section, based on

the two proposed techniques, we integrate an intelligent RTT estimation based mechanism to optimize the handover procedure, and further minimize the handover duration and enhance the handover efficiency.

The principal object is to make the handover delay as close as possible to the critical limits that is the layer2 handover delay which comprises the selective AP scanning, authentication and association delay, these delays can't be reduced any more for a given AP products without modifying its low layer driver.

In default case, when a mobile node finishes the layer2 and layer3 handover procedure, it can't receive the data from its corresponding node at once, it should firstly inform its corresponding node that the handover is over and it's available to continue the data transmission, then the corresponding node starts sending packets to the mobile node, this "classical" procedure induces a RTT delay and is not neglectable for the real time application.

Our proposal focuses on this RTT delay avoidance and minimizes the handover delay to layer2 latency. Based on the RTT delay estimation, our proposal allows the corresponding node to reactive the transmission and start the data transmission in advance before the informing packet (inform that the mobile node finished the handover) is received by corresponding node. Concretely, the corresponding node re-activates the transmission and starts to send the data to mobile node's new IP address (new IP address has already assigned via Pre-DHCP and been informed to the corresponding node) at a estimated moment in order that the data arrives at the mobile node side as soon as the mobile node finishes layer2 handover and activates the reserved IP address for receiving data. Therefore in mobile node side, the handover duration reaches the critical limits: the layer2 delay.

Fig 3.13 represents the sequence procedure of our proposal. We denote t , the layer2 handover delay which comprises the selective scanning, authentication and association delay; t_1 , the delay of packet sent from corresponding node to the mobile node in the previous subnet before the handover; t_2 represents the delay of packet sent from corresponding node to the mobile node after the handover procedure in new subnet. t is measured in advance and can be considered as a constant; t_1 and t_2 can be considered approximately to the half of the RTT in previous and next subnet. When the mobile node is in the previous subnet before the handover, t_1 can be calculated by the mobile node and t_2 can be predicted by an analysis of the different delays of the two access networks. Normally, the delay of the different natures of access networks (i.e. WLAN, 3G, etc.) have different orders that have been pre-measured.

When a mobile node moves to the edge of WLAN coverage and the monitored SNR (based on the low pass filter, see section 3.6.1.1) drops under a defined threshold, the mobile node starts the handover procedure, it firstly sends a informing packet (I1, format detailed in Fig 3.15) to its corresponding node including the pre-reserved new IP address in the next subnet as well as the value of t , t_1 and t_2 in millisecond, this packet is considered as management packet which is sent with lowest modulation and with large persistence of MAC level retransmission to guarantee the reachability to corresponding node with a high probability. When the corresponding node receives the informing packet, it returns another management ACK packet (A1, see Fig 3.15) to permit mobile node's handover, then the corresponding node enters into "sleep" state for a duration of T' :

$$T' = \max(0, t + t_1 - t_2) \quad (11)$$

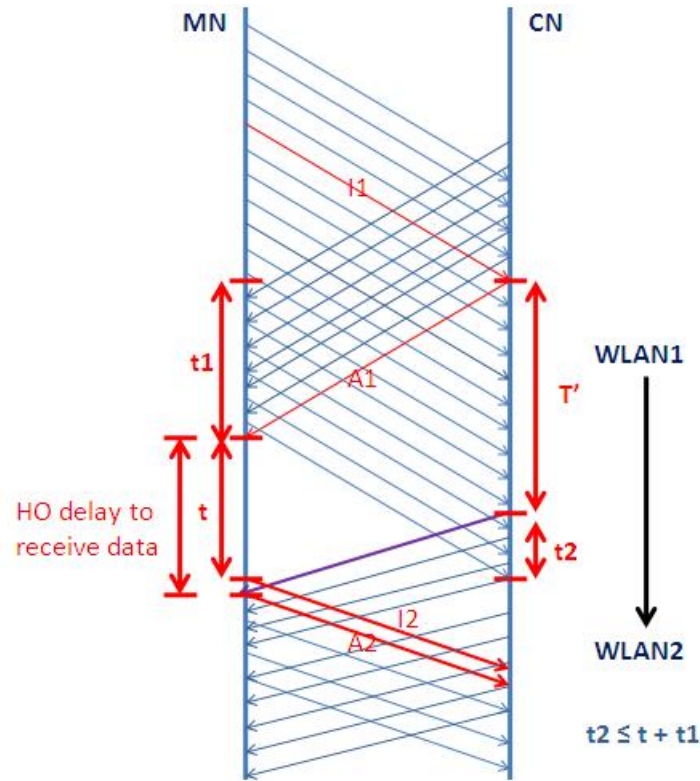


Fig 3.13 Sequence procedure of handover in case $t_2 \leq t_1 + t$

The mobile node starts the handover procedure as soon as the ACK packet is received. When the layer2 handover (corresponding to the delay t) is finished and reserved IP address is active, it send a new “informing” packet (I2, see Fig 3.15) to the corresponding node to inform that the handover is over, and meanwhile, it can start to receive the packets from corresponding node at once in the optimal case if ($t_2 < t + t_1$). When the mobile node receives the first packet from corresponding node after the handover, then it will also return an ACK packet (A2, see Fig 3.15) to its corresponding node that the transmission is re-active. If the t_2 in the new WLAN subnet is very long ($t_2 > (t + t_1)$) (see Fig 3.14), then CN doesn’t stop sending data even during MN’s handover (T is set to 0, destination IP address changes after sending A1 packet), and the mobile node has to wait ($t_2 - (t + t_1)$) ms to start receiving data after the layer-2 handover.

We argue that for handovers across WLANs, transmission delay (or RTT) from the corresponding node to the mobile node in neighboring WLAN subnets are similar and T' is set to t , the handover latency of layer2 delay that is indispensable in handover procedure.

In case of the losses of the handover management packets, the handover procedure switches to the default mode automatically, mobile node is forced to start layer 2 handover after a timeout of four times of RTT ($8 \cdot t_1$) in the previous subnet in our proposal. Corresponding node stops data transmission if no packet arrives during a timeout T_{out} and it can be re-active when a new informing packet (i.e. I2 packet) is received indicating that the mobile node has finished the handover and available for data transmission.

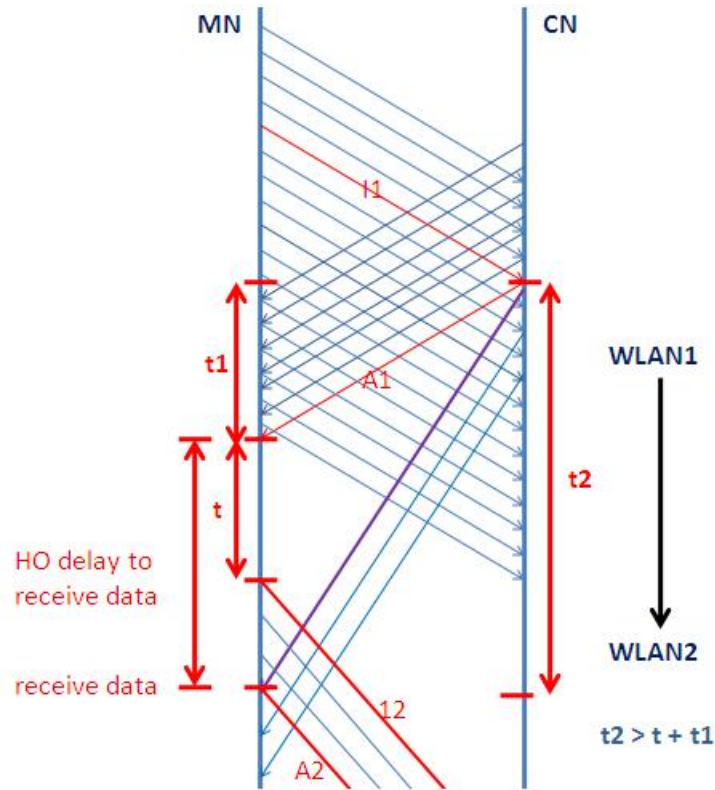


Fig 3.14 Sequence procedure of handover in case $t_2 > t + t_1$

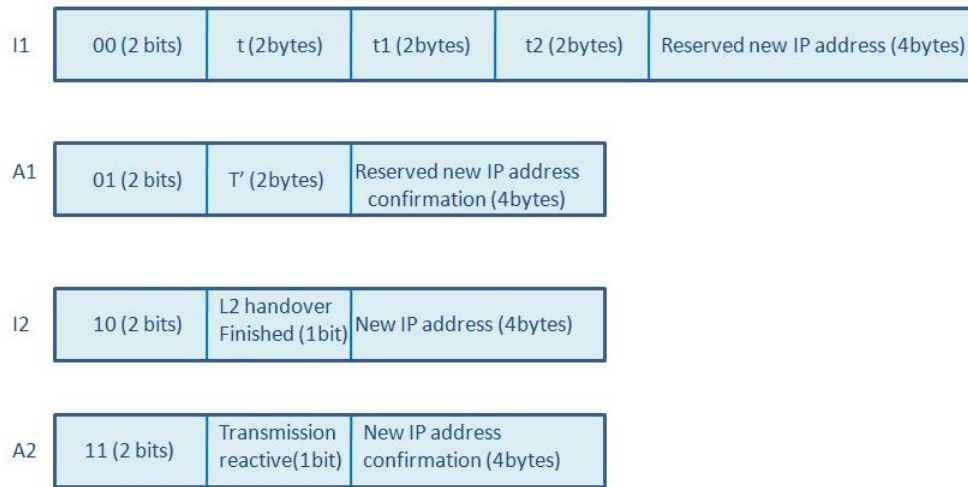


Fig 3.15 Packets format

3.6.3.3.6 Validation of RTT estimation based mechanism

In our experimental test, we use DLink 614+ as access router, our test bed is based on Ubuntu 7.05, we extend the JMF (Java Media Framework) to support our end to end multimedia-transmission based handover procedure. The object of our experiments is to measure the handover delay and compare them to the layer2 average handover latency (AP scanning, authentication and association delay) that is measured in advance. Our proposal allows the handover duration to get as close as possible to the layer2 delay. In our tests, the RTT are the same for different subnets, $t_1 = t_2$, and they are measured to be 92ms in average; the layer2 handover delay is measured to be 418ms by analysis of the channel scanning and association response packets (channel scanning takes most part of the time). The mobile node migrates from AR3 to AR4 with walk speed. We record, at

the IP layer, the timestamps of the last received packet before the handover and the first arriving packet after the handover from corresponding node, which is a JMF video server. Fig 3.16-a represents our scenario setting that a mobile node is moving from AR3 to the overlapped AR4.

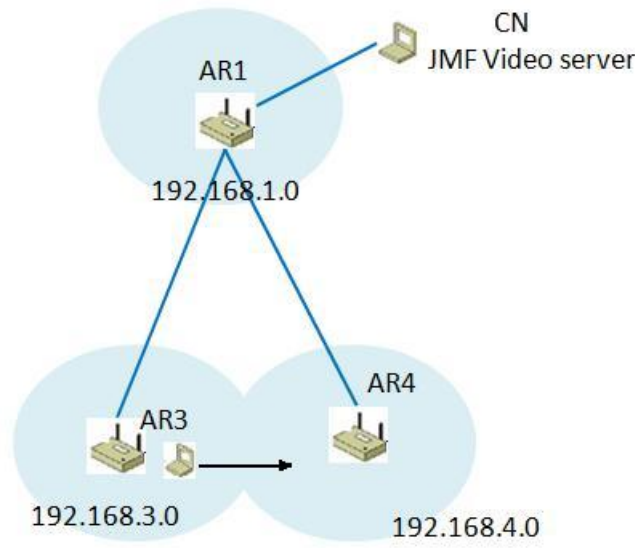


Fig 3.16-a Handover experiments scenario

Fig 3.16-b represents the handover latency with our proposal in different tests as well as the layer2 handover delay. The main reason for the handover delays not to reach the layer2 latency is that the t_2 measured in advance varies more or less during the experimental tests.

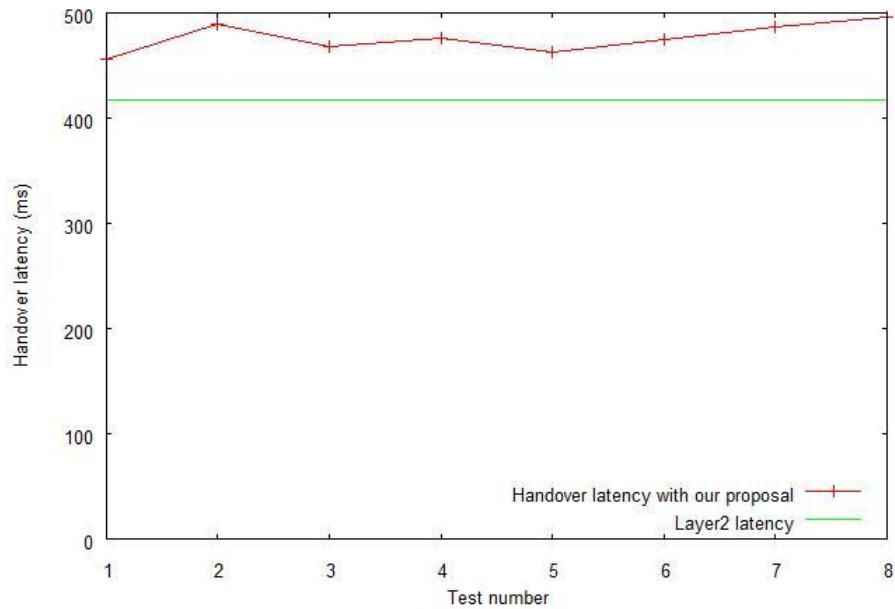


Fig 3.16-b Handover latency from experiments

We showed that with this approach that the handover delays are very close to the layer2 handover delay, which is the critical latency that cannot be reduced any more for a given AR product without modifying the MAC layer driver.

3.7 Conclusion of chapter 3

In this chapter, we introduce our end-to-end mobility architecture that relies on a bundle of cross layer based mechanisms (existing and novel mechanisms), which are dynamically configurable and allows significantly enhancing the performance of the mobility management. Meanwhile, due to the adoption of the “end-to-end” paradigm, our architecture entails as few changes as possible in the current Internet protocol architecture. Our architecture focuses on two principal parts, which are location management and handover management. The cooperation of different functions results in information sharing and transparency, and greatly improves the management efficiency, which have been validated by a set of our experiment results. We will show in chapter 6, based on the proposed mobility management architecture, our middleware and implementations which demonstrates its capability of supporting context-aware mobility management and efficiently satisfies dynamic mobility requirements.

Chapter IV

4 WLAN performance enhancements

The deployment of the wireless networks is accelerated by the emergence of portable terminals in business and house environments. The real potential of broadband wireless networks lies on the mobility, the success of Wi-Fi network with IEEE 802.11x technology makes it possible to access broadband with low cost. Specially, the next generation 802.11n wireless network is being standardized by adding several powerful technologies like MIMO [131] and 40MHz operation [132] to the physical layer which support much higher transmission rates and wider signal coverage. One of the core features of 802.11, the contention based CSMA/CA access method which guarantees equity for long term channel access, has never be changed since its first original version. It entails a round robin scheduling for the access to the medium, which simplifies the access method and aims to offer an easy and fair way to share the bandwidth resource by the mobile hosts attached to the same access point. However, much work have shown that this medium access method can degrade performances on the wireless transmission [133-135]. In this chapter, we will tackle three of these intrinsic problems associated to the 802.11 access method: 1) 802.11 Performance Anomaly; 2) Performance degradation due to MAC layer overflow; 3) Fairness degradation. Then we will propose cross layer based solutions to efficiently solve these problems and significantly improve the wireless transmission efficiency using the end to end and cross layer based solutions which do not entail any changes to the standard of MAC layer.

The 802.11 performance anomaly was first raised in [133], where it was proved that mobile node with the lowest transmission rate impact on the performance of every mobile node in the coverage of the same access point, and that the 802.11 access method which guarantees the equal sending opportunity for each mobile node is the origin of such syndrome. Simulation and experimental results in [133] have shown that this anomaly can dramatically reduce the network transmission efficiency if one or several mobile nodes degrade their transmission rate due to the signal fading and interferences.

Meanwhile, the UpLink bandwidth resulted from this anomaly is principally restricted by the contention based MAC layer, when the sending rate from transport layer becomes higher than the offered rate, packets can be lost in MAC buffers, which degrades the quality of transmission special for the reliable connections.

And finally, fairness degradation is one of the most challenging issues in the context of 802.11. We identify two main unfairness problems in this chapter: 1) Unfairness issue between Uplink(UL) and DownLink(DL) flows; 2) Unfairness issue between ACK clocked (i.e. TCP) and Non-ACK clocked (i.e. UDP) flows. Since the access point, AP, is considered as a normal contention-based mobile node, it has the same opportunity of sending packets (to all the download mobile nodes) as any of the sending mobile nodes, so the aggregated bandwidth of the UL mobile nodes can be much higher than of the DL nodes. In the other hand, we identified that, in the scenarios when ACK and Non-ACK clocked UL flows coexist, the contention avoidance procedures implemented at the 802.11

MAC layer of access points can slower the rate of returned ACK packets to the uploading MNs in the coverage of AP, and as a result can slower the sending rates of the ACK clocked connections.

Furthermore, in the current research community on wireless networks, one of the obstacle facing *the client and network load balancing function* and *real-time radio resource management* is the lack of an accurate and efficient bandwidth estimation technique for WLANs. Few work [127-130] have investigated in this issue, and unfortunately, all of them consider best effort flows only, none of them takes into account the whole complexity of the communication context. We propose a generic model which allows calculating 802.11 MAC-layer available rate while taking into consideration different wireless scenarios' parameters such as mobile node's dynamical transmission status, flows profiles and the different protocols which can be possibly involved (i.e. UDP, TCP, TFRC, ACK clocked, Non-ACK clocked flows).

In this chapter, we propose several novel cross-layer based mechanisms which allow efficiently alleviating and solving the above three main issues of 802.11. These performance improvement mechanisms are based on an analytical model of the rate delivered by the WLAN MAC layer. This analytical model opens the door to fruitful cross layer interactions between the lower and higher layers of the WLAN architecture, and can be used to apply efficient rate & congestion control mechanisms to suppress the previously introduced issues, the analytical model has been validated from intensive simulations and real experiments and programmed in a python based tool to calculate the WLAN MAC supported bandwidth in terms of different wireless scenario parameters. Furthermore, we introduce a upper layer FEC based mechanism that allows efficiently alleviating the 802.11 performance anomaly and improve the transmission goodput. Our end-to-end based mechanisms respect the current 802.11 access method and induces no change to the standard, which makes our proposals easy to deploy.

4.1 Three main issues of 802.11

This section identifies the three main hidden issues of the CSMA/CA access method of 802.11 that hinder WLAN performances.

4.1.1 Performance Anomaly of 802.11

The access method of the IEEE 802.11 standard defines the Distributed Coordination Function (DCF) that uses CSMA/CA to allow mobile nodes to contend for accessing the wireless media, which is wildly implemented in the current 802.11 products. DCF offers a best effort type of service, each mobile station checks whether the medium is idle before transmission, immediate transmission is available if the medium idles for longer than DIFS (Distributed Inter Frame Space). If the medium is busy, the mobile station waits as access deferral for DIFS and retries after an exponential backoff delay. This contention-based access method offers equal sending opportunity for each mobile node in long term channel access and allows multiple stations to interact without central control. However, this access method is the root cause of the so called 802.11 Performance Anomaly: namely the performance of the mobile nodes in the coverage of an access point (AP) is degraded when one or several mobile nodes use lower transmission rates than others. The main reason is that the mobile node can occupy the channel for a long time because of its lower transmission rate, and as a result penalizes other hosts that use higher transmission rate. This syndrome was identified and analyzed in [133]. The useful UL bandwidth for each mobile node in the same AP coverage is proved to be identical and strongly dependent on the number of competing hosts and their respective transmission rates.

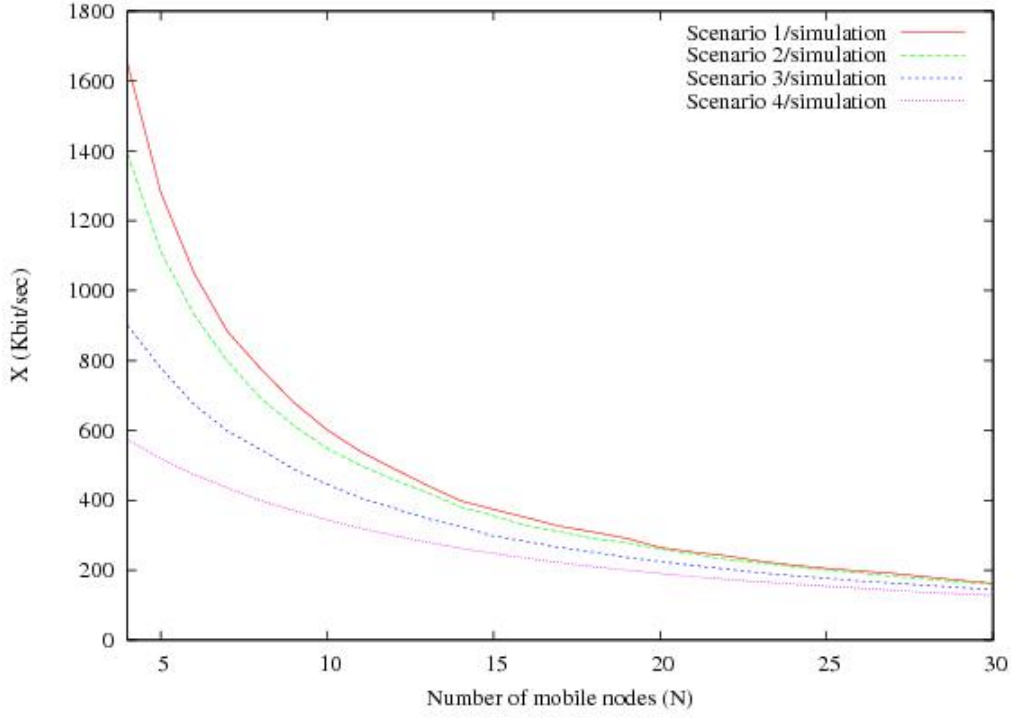


Fig 4.1 Rate in terms of number of mobile nodes

In order to illustrate this phenomenon, Fig 4.1 shows simulation results that we have done with OPNET [139]. The results show that the evolution of the UDP sending rate for each mobile node (each MN is the source of a UDP flow only) in the AP coverage according to number of the uploading mobile nodes ($N = [4, 30]$) in four different scenarios. In the first scenario, all the mobile nodes have a transmission rate of 11Mbit/s. Then, for each other three scenarios, only one among N mobile nodes has respectively a transmission rate of 5.5Mbit/s, 2Mbit/s and 1Mbit/s. The simulation results clearly show that the slower nodes may considerably limit the throughput of other mobile nodes roughly to a much lower level, and such anomaly alleviates when the number of nodes increases.

Several contributions [136-138, 140, 141] attempted to solve this anomaly issue from different point of views, such as packet aggregation and fragmentation solutions, CW size adjustment solutions, etc. However, all of these studies require either a complicated scheduler or deep modification to contention windows at the MAC layer, which is not compatible to the current 802.11 standard and makes difficult their wide deployment. Our proposal which is detailed in section 4.5 introduces a new approach that integrates a flexible adaptive FEC (also known as erasure code) mechanism at the higher layer to keep a high transmission rate for the mobile nodes as long as possible while optimizing their goodput. As a consequence our contribution alleviates the performance anomaly and delivers a higher global goodput.

4.1.2 MAC layer overflow

The above analysis shows the maximum UL bandwidth supported by the MAC layer of each mobile node in an AP coverage strongly depends on the number of competing nodes and their transmission rates, the upper layers are unaware of this information because of the lack of cross-layer interactions. Therefore, when the sending rate from higher layer surpass this offered rate, packets can be lost in MAC buffers, which degrades the quality of transmission and can raise several QoS issues for QoS constrained connections. A set of simulations and experiments have been done to demonstrate the discrepancy between the transport layer sending rate and the real sending rate offered by the MAC layer.

For UDP flows that offer the best effort service, our experimental testbed is composed of 3 hosts, two of them are equipped with Atheros IEEE802.11a cards (to avoid interference by others 802.11b hotspots situated in the lab). A mobile node A sends a UDP flow (packets size 1500bytes) to a remote wired node B (receiver) via an access point (AP) which is 1m distant from the mobile host. Both wireless stations run FreeBSD6.1 and use *ifinfo* tool in order to check the number of packets sent by the wireless interface. This tool returns information from the wireless card and in particular: the instantaneous length, the maximum length and the number of drops in the send queue of the wireless interface. The MAC buffer of node A is set to 50 packets. We measure the percentage of losses according to node A's sending rate. We made a similar simulation scenario with OPNET. The results of these simulations and experiments, as illustrated in Fig 4.2, show that the throughput at the transport layer can surpass the maximum bandwidth that the MAC layer can support and can lead to massive MAC buffer overflow and significant packet loss rate.

We then addressed the impact of this lack of cross layer rate control on congestion control streams such like TFRC and TCP flows. TFRC protocol [142] has been proved to be able to offer a smooth, low delay and TCP-Friendly packet transmission. We firstly developed the module of TFRC protocol in the OPNET, which has been validated by comparing numerous simulation results with the same scenarios under NS-2 simulator. Our simulation work focuses on two TFRC mobile nodes uploading data to remote servers with a transmitting rate of 5.5Mb/s where congestion occur at the MAC layer of the mobile node.

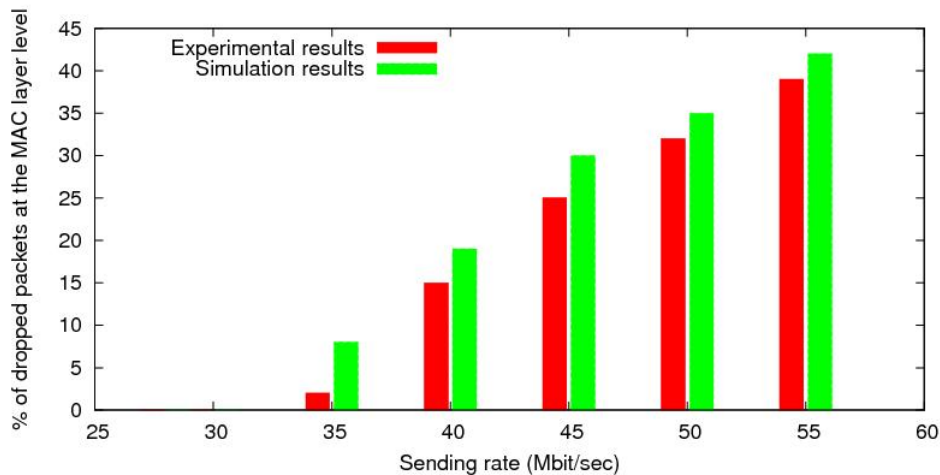


Fig 4.2 MAC layer overflow in UDP case

Fig 4.3 gives the result of the TFRC throughput and the maximum available throughput supported by the MAC layer (X_t). We can observe that TFRC obtains unstable rate variations around X_t . The transport layer throughput has a standard deviation of 115.8Kb/s, the unstable oscillations are proved to be caused by the packet losses in the MAC buffer.

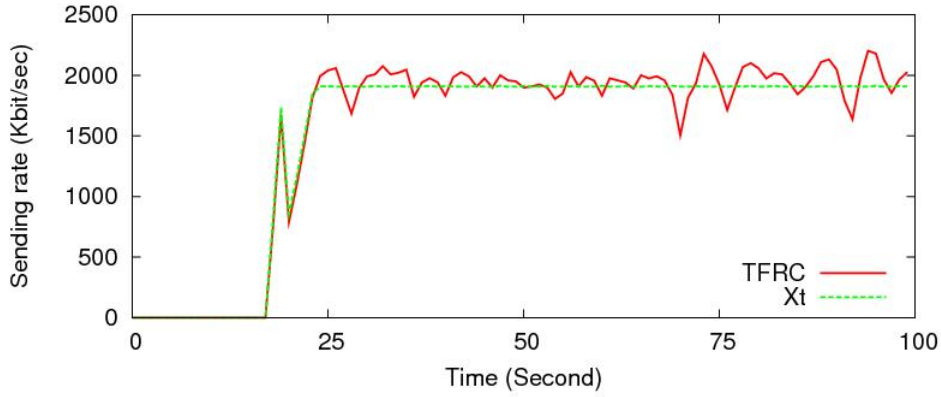


Fig 4.3 TFRC case

Furthermore, we also illustrate the impact of these losses on the behavior of a TCP flow in real conditions. We are not interested in quantifying the exact number of TCP losses as it depends of the wireless card used. Indeed, the MAC buffer characteristics might be differently sized following the card or the system in use. Thus, we propose to generate a TCP traffic from a mobile host to a wired host through a wireless LAN (topology and scenario similar to [133]). This traffic is a file transfer of 15min duration. The wireless station is connected to the base station (BSS) with 802.11b access mode and the transmit rate is set to 1Mbps. Behind the BSS, a wired host is connected at 100Mbps. In order to identify possible TCP losses which occur at the MAC layer of the mobile node, we realize a Tcpdump capture [136] at the IP level at the wireless input of the BSS. The traffic is generated from a GNU/Linux mobile node to the wired host. We use the TCP Newreno version and the TCP window size set to the maximum size (64000 bytes) as well as the window scale option. The idea of identifying MAC layer losses of TCP flow is to look at the BSS trace when a packet appears out of order following the well-known algorithm described in [137]. The principle of this algorithm is as follows: if the BSS observes a hole in the TCP sequence number followed by an out of order packet, then it means that a loss is occurred between the MN and the BSS; however if a duplicate sequence number is observed, it means that a loss is occurred between the BSS and the wired node. Indeed, a duplicate sequence number corresponds to a retransmission of a packet lost after its capture by the BSS, while a rupture in the sequence means a loss before the BSS capture. We have analyzed the TCP traces and observed periodic bunches of lost packets (every 2 minutes in our experiment), which is shown in Fig.4.4-a (section “OR” of the sequence number represents the out-of-order data). Fig.4.4-b is a zoom of an “OR” section.

The periodic character of these losses clearly proves experimentally that they are due to buffer overflow instead of channel losses (which should be very rare and should occur randomly according to the “perfect” communication experimental conditions). In this experiment, we have found 555 over 79159 packets dropped due to MAC buffer overflow (several others experiments have shown that this value oscillates between 0.7% and 0.8% with the same experimentation characteristics). We observed that an average of 90 packets (± 10) are lost during every loss period. As a matter of fact, this bunch of losses is obviously prejudicial for the performance of TCP flows. Such losses can be much more significant if the size of MAC buffer decreases or other nodes (specially greedy UDP mobile nodes) are involved in the scenario.

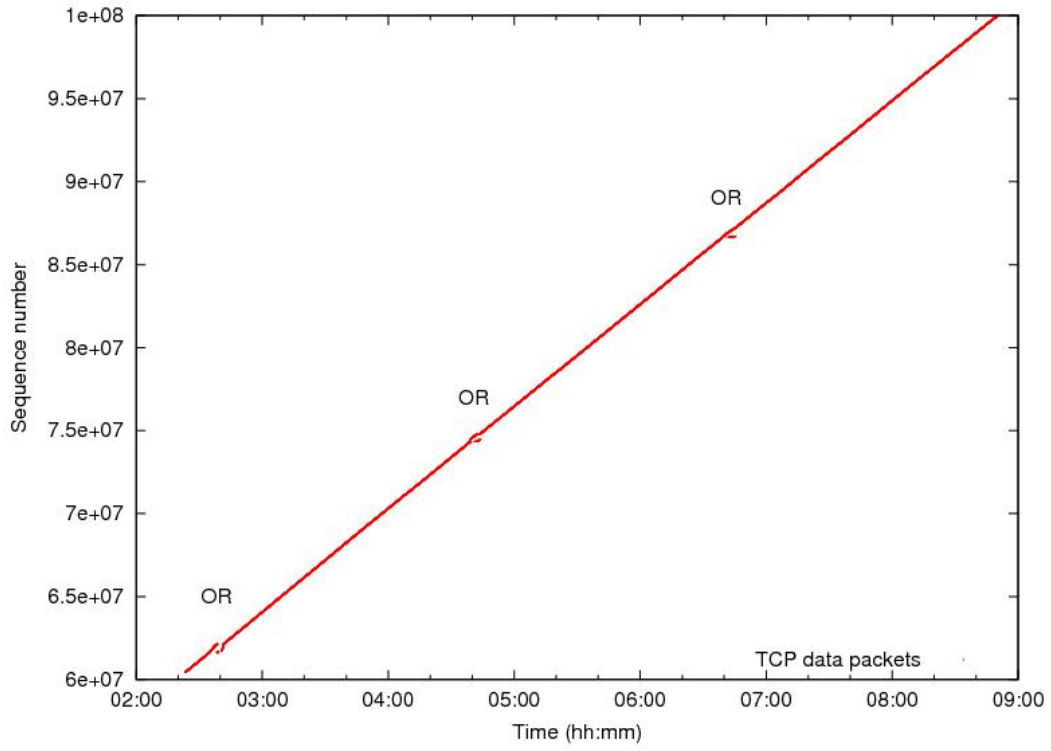


Fig 4.4-a MAC overflow for TCP scenario

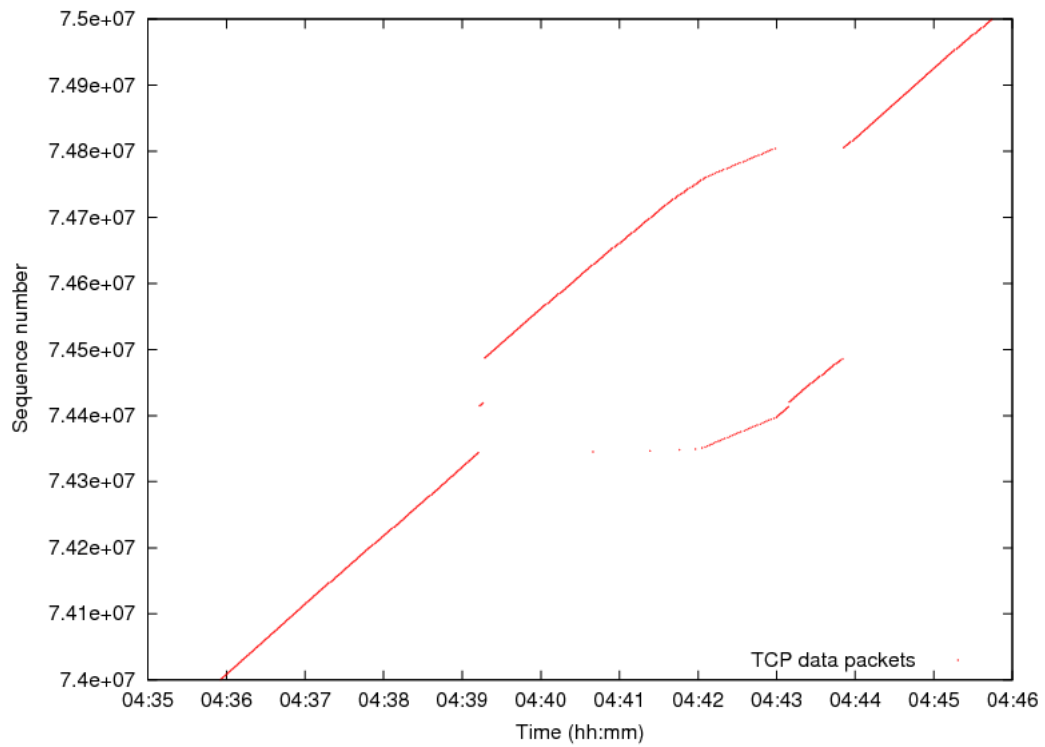


Fig 4.4-b Zoom of "OR"

We deduce from the simulation and experimental results, that coordinating the higher layer sending rate with the rate delivered by the WLAN MAC layer can entail an important loss reduction, thus lowering end to end delays and jitter variations, and as a result improving the transmission efficiency. A solution to solve this issue is to be introduced in section 4.3.

4.1.3 Fairness issues

In the context of 802.11, AP is considered as a normal contention-based mobile node, it has the same opportunity of sending frames (to all the download mobile nodes) as any of the upload mobile node in long term, which significantly degrades the aggregated downloading bandwidth. If we suppose that U mobile nodes are uploading UDP streams to remote servers via an AP, and that D UDP flows are concurrently sent by this same access point to D downloading mobile nodes, then the average throughput of each upload flow is equal to the aggregated throughput of the download flows sent by AP. In other words, instead of getting a fair throughput ratio of $\frac{D}{U+D}$, the aggregated download flows occupy a unfair throughput ratio of $\frac{1}{U+1}$. We have identified and addressed this unfairness issue between UL and DL flows in [135], our simulation illustrated in Fig 4.5, (with $U = 2, D = 2$ in 802.11b scenario, all of the nodes use a transmission rate of 11Mbps), shows that each UL mobile node occupies the bandwidth two times higher than each DL node, the aggregated download flow receive a unfair throughput ratio around 1/3.

In the other hand, let's consider the case when ACK clocked and Non-ACK clocked flows coexist. Take TCP and UDP for example, we should take into account the ACK packets frequently returned to each TCP uploading mobile node via AP, which occupy a part of DL wireless network resource. As the DL bandwidth is limited due to the above-stated unfairness problem, the contention avoidance procedures implemented at the 802.11 MAC layer of access points can slower the rate of returned ACK packets to the uploading MNs in the coverage of AP, as a result slowing the TCP sending rate. UDP sending rates are not constrained by the ACK packets, and can make full use of the available wireless network resource offered by the CSMA contention-based mechanism.

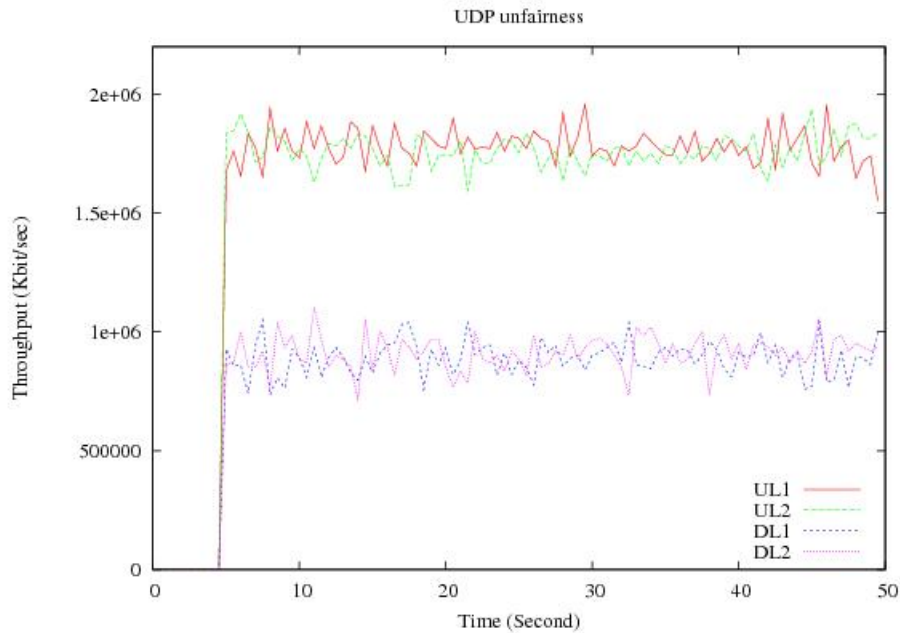


Fig 4.5 UDP unfairness

We introduce in section 4.4 a rate-equalization mechanism uniquely based on an end to end approach (i.e. with no modification of the MAC layer) to solve the two main unfairness issues in the context of 802.11.

4.2 MAC Layer rate modeling

As shown in section 4.1.2 the simulation and experimental results, following the systematic absence of flow control between the IEEE 802.11MAC layer and the upper layers, the rate at the upper layer can strongly diverge from the rate offered by the MAC layer, this mal-adjustment between these two layers can induce numerous losses at the WLAN MAC layer, and therefore degrade the quality of service delivered by the various transport layer to users and applications. In order to efficiently model and calculate the rate delivered by the 802.11 MAC layer for mobile nodes, which is a fundamental parameter of our proposed rate control method, we firstly introduce in this section an analytical model to estimate this maximum supported rate by 80211 MAC layer in generic 802.11b scenarios according to the number of competing nodes, their transmission rates as well as different applied protocols.

4.2.1 Fundamental modeling parameters

Firstly, we introduce several parameters that are involved in our 802.11 MAC layer analytical model.

4.2.1.1 Node profiles

Firstly, in order to distinguish idle nodes from those which make full use of the available bandwidth delivered by the MAC layer, we define greedy nodes, whose sending rate at upper layer can reach or surpass the maximum bandwidth that the MAC layer can support. On the contrary, we define rate sparing nodes, whose throughputs are limited by their application layer (e.g. a VoIP node that requires a relatively low bandwidth) or by congestion over the network. In a contention based MAC layer, rate sparing nodes can alleviate the contention level in MAC layer and give more sending opportunities to other greedy nodes.

4.2.1.2 ACK effect

ACK effect happens when ACK clocked (TCP) and Non-ACK clocked (UDP) flows co-exist. We suppose there are several TCP and UDP based uploading mobile nodes in the 802.11 access point (AP) coverage. In the context of 802.11, AP is considered as a normal contention based mobile node, it has the same opportunity of sending packets (specially ACK packets to all the TCP based sending mobile nodes) as any of the UDP uploading mobile node in long term. Therefore, the bandwidth available for forwarding ACK packets by the AP is, roughly speaking, conversely proportional to the number of UDP based uploading nodes. In the context of TCP, we denote K : the average number of acknowledged TCP segments (on average with TCP Newreno $K=2$) in the steady state (i.e. in congestion avoidance phase). When the number of TCP uploading mobile nodes increases and surpass K , the contention avoidance procedures implemented at the 802.11 MAC layer of the AP can slower the rate of returned ACK packets to the TCP uploading nodes. As a consequence, the reduced pace of ACK packets slows the TCP sending rate. Therefore, each TCP uploading mobile node doesn't make full use of the available wireless network resource offered by the CSMA contention-based mechanism, this part of TCP unused network resources are therefore captured by the Non-ACK clocked flows (UDP connections) which are not limited by the return pace of ACK packets. This analysis is detailed in section 4.2.2.2 and shows that the ratio of upload rate between each TCP and UDP node tends to $\frac{K}{N_t}$ if $N_t > K$ and tends to 1 if $N_t \leq K$ (packet sizes of the TCP and UDP are identical). This unfair phenomena between the ACK clocked and Non-ACK clocked connects is called "ACK effect" in this chapter.

4.2.1.3 Overall transmission duration of the 802.11 MAC layer data frame

We suppose that N mobile nodes transmit data frames to remote hosts through 802.11 Access Point (AP). Four groups of mobile nodes is classified among the N mobile nodes according to their respective transmission rates, N_i ($i = 1, 2, 3, 4$) mobile nodes use a transmission rate of R_i , where $R_1 = 11\text{Mb/s}$, $R_2 = 5.5\text{Mb/s}$, $R_3 = 2\text{Mb/s}$, $R_4 = 1\text{Mb/s}$. We denote S_m : the MAC-layer frame length in bits. T_i^{tr} represents the duration to transmit a data frame with a certain transmission rate R_i :

$$T_i^{tr} = \frac{S_m}{R_i} \quad R_i \in (1, 2, 5.5, 11\text{Mbps}) \quad (12)$$

We denote that T_i^{ov} represents a constant overhead which comprises DIFS, SIFS, two times of the PLCP preamble and the header transmission time as well as the MAC acknowledgment transmission time t^{ack} .

$$T_i^{ov} = DIFS + SIFS + 2 * t_i^{pr} + t^{ack} \quad (13)$$

In 802.11b, $DIFS = 50\mu s$, $SIFS = 10\mu s$, $t^{ack} = \frac{112}{R_i}$, $t_i^{pr} = 192\mu s$, when the mobile node use the transmission rate of 1Mb/s and $t_i^{pr} = 96\mu s$ when the mobile node uses the other transmission rates.

In [133], the authors give an approximation to express the average duration of backoff process $T^{cont}(N)$ as function of N : the total competing mobile nodes in the 802.11 coverage and CW_{min} : the minimum contention window size. In 802.11b, $CW_{min} = 31$ and $SLOT = 20\mu s$.

$$T^{cont}(N) = \frac{SLOT * (1 + Pc(N)) * CW_{min}}{4 * N} \quad (14)$$

Where $Pc(N)$ is the proportion of collisions experienced for each packet successfully acknowledged at the MAC level.

$$Pc(N) = 1 - \left(1 - \frac{1}{CW_{min}}\right)^{(N-1)} \quad (15)$$

With the above equations, we can calculate T_i , the average overall duration to transmit one data frame for each node in the group i :

$$T_i = T_i^{tr} + T_i^{ov} + T^{cont}(N) \quad (16)$$

4.2.2 Analytical model of calculating maximum bandwidth delivered by 802.11 MAC layer

This section focuses on an analytical model which allows accurately estimating the available bandwidth delivered by the WLAN MAC layer. The proposed model pushes further the approach proposed in [133] by considering both the mobile nodes' different transmission rate profiles and the specificities of the different type of protocols that compete for this bandwidth.

4.2.2.1 Maximum bandwidth supported by 802.11 MAC layer for the Non-ACK clocked mobile node

In this section, we will consider the case where only non-ACK clocked streams compete for the available bandwidth, we will presents our modeling work which allows accurately calculating the maximum bandwidth (X_m) delivered to mobile nodes by the MAC layer for generic 802.11b scenarios. Following the definition in section 4.2.1.3, N mobile nodes

in 802.11b coverage are classified into four groups N_i ($i = 1, 2, 3, 4$) according to their transmission rates. For each group N_i , we also suppose that there are K_i rate sparing nodes among N_i mobile nodes using the transmission rate R_i . The average throughput of these K_i nodes limited by their application or the congestion in the network is X_i^j ($j = 1, 2, \dots, K_i$). Then each of the $(N - \sum_{i=1}^4 K_i)$ greedy nodes fully uses the maximum throughput X_m delivered by the MAC layer. Meanwhile, AP is considered as a normal contention based mobile node and X_{AP} represents the aggregated throughput from AP to the N mobile nodes in the AP coverage (i.e. TFRC feedback packets or downloaded data). We can calculate the aggregated bandwidth X between all the mobile nodes and the access point is given by:

$$X = \sum_{i=1}^4 \sum_{j=1}^{K_i} X_i^j + (N - \sum_{i=1}^4 K_i) * X_m + X_{AP} \quad (17)$$

We also define the proportion of the throughput used by each rate sparing nodes in group i : $P_i^j = \frac{X_i^j}{X}$ ($j = 1, 2, \dots, K_i$), the proportion of the aggregated throughput used by the AP: $P_{AP} = \frac{X_{AP}}{X}$, as well as the proportion of the throughput for each of the $(N - \sum_{i=1}^4 K_i)$ greedy mobile nodes: $P_b = \frac{X_m}{X}$.

Since the contention based access method of 802.11 allows all the greedy mobile nodes to share fairly the radio channel and have the same opportunity of sending packets, theoretically, once a data frame is sent out by any of the greedy mobile nodes, it should wait an average period of time to send another data frame, this average time T between the two successive packet emissions comprises the following 4 parts:

(1) the time required for sending one packet by each of the greedy node with different transmission rate: $\sum_{i=1}^4 T_i * (N_i - K_i)$;

(2) the time required for sending packets by the sparing nodes with different transmission rate. According to the above analysis on the rate proportion, every time a packet is sent out by a greedy node, there should be $\frac{\sum_{i=1}^4 \sum_{n=1}^{K_i} P_i^n}{P_b}$ packets sent by all the rate sparing nodes;

(3) the time required for sending packets (i.e. TFRC feedback packets or downloading data) from the AP to N mobile nodes. Similarly, every time a packet is sent out by a greedy node, there should be $\frac{P_{AP}}{P_b}$ packets sent by AP;

(4) the time spent in collisions (T_{col}), since AP is considered as a normal competing mobile node when it sends packets, we have a total of equivalent $(N+1)$ contention mobile nodes in the coverage, we have:

$$T_{col} = Pc(N+1) * t_{jam} * (N+1) \quad (18)$$

t_{jam} represents the average time spent in collision for each node in case of collision, whose calculation is given in appendix 8.3.

The average time T is represented by:

$$T = \sum_{i=1}^4 T_i * (N_i - K_i) + \frac{\sum_{i=1}^4 \sum_{n=1}^{K_i} P_i^n * T_i}{P_b} + T_{AP} * \frac{P_{AP}}{P_b} + T_{col} \quad (19)$$

T_{AP} and T_i can be calculated with equation (16). S_m is defined as the length of MAC layer packet in bits, with the expression of the average time between the two successive packet emissions T . Then, we can calculate the maximum throughput supported by the MAC layer for greedy nodes X_m :

$$X_m = \frac{S_m}{T} \quad (20)$$

If we suppose S_t , the length of transport layer packet in bits, then we can calculate the maximum available throughput at the transport layer:

$$X_t = \frac{S_t}{T} \quad (21)$$

This section has defined an analytical model to calculate the maximum bandwidth delivered by 802.11 MAC layer scenarios (when non ACK clocked stream are considered only) while taking account of mobile nodes' different transmission rate profiles. Next subsections specify the calculation of the rate X_m in several scenarios where specific protocols are implemented.

4.2.2.1.1 UDP case

Let's suppose that only UDP upload flows exist in our scenario and that all the nodes are in greedy mode, so we have $P_{AP} = 0, K_i = 0, P_i^n = 0$. We can simplify the equation (19) to:

$$T = \sum_{i=1}^4 T_i * N_i + T'_{col} \quad (22)$$

Since AP is no longer considered as a transmission mobile node (no download packets are sent from AP to N mobile nodes), we have:

$$T'_{col} = Pc(N) * t'_{jam} * N \quad (23)$$

The calculation of t'_{jam} is given in appendix 8.1.

Thus, the maximum bandwidth capacity supported by mobile nodes' MAC layer X_m (and X_t , maximum available bandwidth at the transport layer) in UDP mode can be calculate with the above equations.

Our analytical model has been validated by a set of simulations under OPNET. Fig 4.1 in section 4.1.1 shows the simulation results under OPNET the evolution of UDP sending rate supported by the MAC layer in function of the number of uploading mobile nodes in different scenarios (In the first scenario, all the mobile nodes have a transmission rate of 11Mbit/s. Then, for each other three scenarios, one among N mobile nodes has respectively a transmission rate of 5.5Mbit/s, 2Mbit/s and 1Mbit/s). Fig 4.6 compares this calculated rate issued from our analytical model to the simulation results, we find that our analytical results are very close to the simulation results.

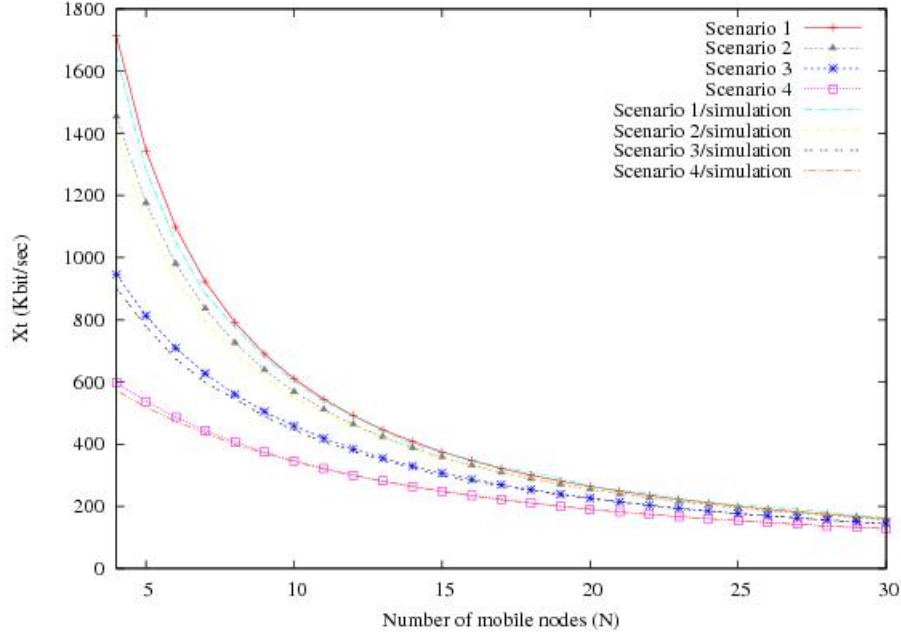


Fig 4.6 Rate in function of the number of mobile nodes

4.2.2.1.2 TFRC case

In the TFRC scenario, let's consider TFRC streams only, we suppose the uplink TFRC mobile nodes are in greedy mode, we have $K_i = 0$, $P_i^n = 0$. We can simplify the equation (19) to:

$$T = \sum_{i=1}^4 T_i * N_i + T_{AP} * \frac{P_{AP}}{P_b} + T_{col}'' \quad (24)$$

Where T_{col}'' is the time spent in collisions:

$$T_{col}'' = Pc(N + 1) * t_{jam}'' * (N + 1) \quad (25)$$

The calculation of the average time spent in collision t_{jam}'' for each node in TFRC case is given in appendix 8.2. With the above proportion analysis, we have:

$$\frac{P_{AP}}{P_b} = \frac{X_{AP}}{X_m} \quad (26)$$

Then the maximum bandwidth capacity supported by mobile nodes' MAC layer X_m in TFRC scenario can be calculated with the above equations. If we suppose there are only TFRC upload flows in our scenario and the TFRC receiver returns a feedback packet every RTT, we can calculate the bandwidth X_{AP} (corresponding to the TFRC feedback packets) from the AP to the N mobile nodes as follows:

$$X_{ap} = \frac{N}{RTT} * S_f \quad (27)$$

Where, S_f is the size of TFRC feedback packet (816 bits) and RTT is the averaged round trip time. If we suppose that X remains the same for UDP and TFRC scenarios where only upload flows exist (with the same distribution of N_i). We can approximately estimate the maximum available throughput at the transport layer $X_t^{UDP} = \frac{X}{N}$ for UDP case and $X_t^{TFRC} = \frac{(X - X_{AP})}{N}$ for TFRC case. The relative error between the two available bandwidths is:

$$E = \frac{X_t^{UDP} - X_t^{TFRC}}{X_t^{UDP}} = \frac{X_{AP}}{X} = \frac{S_f * N}{RTT * X} \quad (28)$$

Since TFRC is not a ACK clocked protocol and its feedback packets take a small part of bandwidth resource among the aggregated bandwidth capacity. This relative error E is normally much less than 5% in the context of 802.11b, as a consequence, we can integrate the UDP based model into the TFRC scenarios to simplify the calculation when the relative error E is little.

4.2.2.2 Maximum throughput limited by contention based MAC layer for UDP nodes in scenarios where TCP and UDP flows coexist

Due to the “ACK effect”, when the number of TCP uploading mobile nodes surpasses K ($K=2$ is the most current use case), we have previously underlined a bandwidth share unfairness issue between UDP and TCP flows, and in this case, MAC buffer overflow is not observed for the TCP based uploading nodes during our simulation and experiments due to their self clocking property and the resulting implicit rate control at the TCP transport layer. Therefore, this section will mainly focus on the estimation of the 802.11 MAC-layer available rate for UDP based nodes while taking into consideration the coexistence of TCP and UDP flows uploaded by the mobile nodes.

For the sake of simplicity and without loss of generality, we consider N_t TCP uploading mobile nodes and N_u UDP “greedy” sending mobile nodes in the coverage of an 802.11b AP. N_t^i and N_u^i ($i=1; 2; 3; 4$) represent respectively the number of TCP and UDP mobile nodes using a transmission rate of R_i , $R_i \in (1, 2, 5.5, 11Mbps)$. Firstly, We suppose that each TCP uploading node has the same number of TCP flows. We denote T_{ta}^i : the overall duration for transmitting an ACK packet from AP to a TCP based uploading mobile node which belongs to the group N_t^i . When $N_t > K$, the average interval (T) between the two successive emissions of UDP packets from the same UDP uploading node comprises the following 4 parts:

1. T1: The time required for sending packets by the TCP uploading mobile nodes that are limited by the “ACK effect”: $\sum_{i=1}^4 (N_t^i * T_t^i * \frac{K}{N_t})$, where T_t^i represents the duration of sending a TCP packet from the mobile nodes in group N_t^i . This formula can be explained as follows: since AP is considered as a normal contention based mobile node, it has the same opportunity of sending ACK packets as any of the UDP uploading mobile node. In the stationary state in which the AP and the UDP mobile nodes are continuously ready to send packets, we consider that between two UDP packet transmitted by the same mobile nodes, one ACK packet can be forwarded to TCP uploading node by the AP, and then a right to emit K TCP packets is offered to the TCP nodes. We can consider, according to the law of large numbers, that these availability of transmitting K packets is shared by the N_t TCP uploading nodes and each TCP node is able to transmit the equivalent of $\frac{K}{N_t}$ packets.
2. T2: The time required by each UDP mobile node to send out one packet: $\sum_{i=1}^4 (N_u^i * T_u^i)$, T_u^i represents the duration of sending a UDP packet from the node in group N_u^i . Note that they are not limited by “ACK effect” to send packets.
3. T3: The time required to send one ACK packet by AP to one TCP uploading node, $\frac{\sum_{i=1}^4 (T_{ta}^i * N_t^i)}{N_t}$. The ACK packet can be sent to any node of the N_t TCP uploading nodes.

4. T4: The time spent in collisions (T_{col}), $T_{col} = Pc(N_u + K + 1) * t_{jam} * (N_u + K + 1)$. In the case when $N_t > K$, N_t TCP uplink nodes are equivalent to K competing nodes that make full use of the bandwidth delivered by the MAC layer. Furthermore, the AP is considered as a normal contention node, and therefore we have an “equivalent” $(N_u + K + 1)$ greedy competing mobile nodes in total.

Then we have:

$$T = \sum_{i=1}^4 (N_t^i * \frac{K}{N_t} * T_t^i) + \sum_{i=1}^4 (N_u^i * T_u^i) + \sum_{i=1}^4 (T_{ta}^i * \frac{N_t^i}{N_t}) + T_{col} \quad (29)$$

where t_{jam} represents the average time spent in collision for TCP-UDP coexisting flows, whose calculation is given in appendix 8.4.

In the case where $N_t \leq K$, TCP mobile nodes are no longer limited by the “ACK effect” and are able to make full use of the bandwidth delivered by the 802.11MAC layer. Thus, we set $T1 = \sum_{i=1}^4 (N_t^i * T_t^i)$ in the above equation

$$T3 = \frac{N_t}{K} * \frac{\sum_{i=1}^4 (T_{ta}^i * N_t^i)}{N_t} = \frac{\sum_{i=1}^4 (T_{ta}^i * N_t^i)}{K}$$

and

$$T_{col} = Pc(N_u + N_t + \frac{N_t}{K}) * t_{jam} * (N_u + N_t + \frac{N_t}{K}).$$

With the expression of T : the average interval between two successive emission UDP packets from the same UDP node, we can calculate the maximum throughput X_m supported by the MAC layer for each of the UDP flow based node:

$$X_m = \frac{S}{T} \quad (30)$$

With S : the length of MAC layer packet in bits.

The maximum available throughput at the transport layer is:

$$X_t = \frac{S_t}{T} = X_m * \frac{S_t}{S} \quad (31)$$

With S_t : the length of transport layer packet in bits.

Note, the that these value also represent the maximum bandwidth supported by the 802.11 MAC layer for TCP based nodes when $N_t \leq K$. When $N_t > K$, the available throughput at the transport layer for TCP node equals to $X_t * \frac{K}{N_t}$.

This analytical model has been validated by a set of simulations under OPNET. Fig 4.7-4.9 show the analytical and simulation results of X_m according to the number of TCP nodes ($N = [2,15]$) and the number of the UDP nodes ($N = [2,10]$) in three different scenarios when two mobile nodes (one UDP and one TCP based node) have transmission rate of 11Mbps and all the other mobile nodes have respectively transmission rate of 5.5Mbps, 2Mbps, 1Mbps (respectively corresponding to scenario 1,2 and 3).

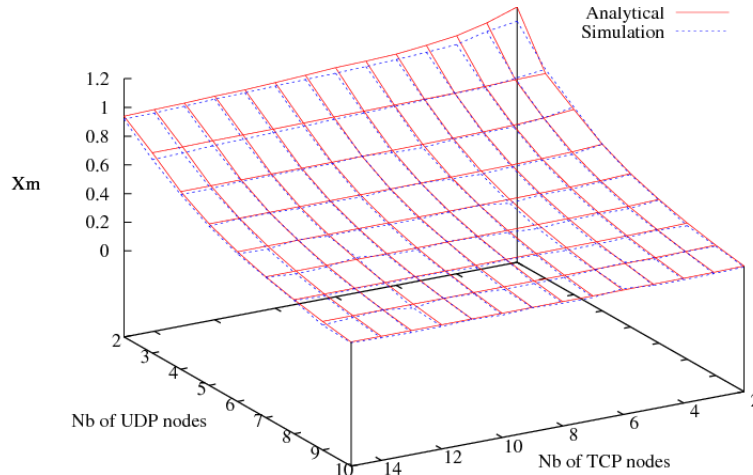


Fig 4.7 Scenario 1

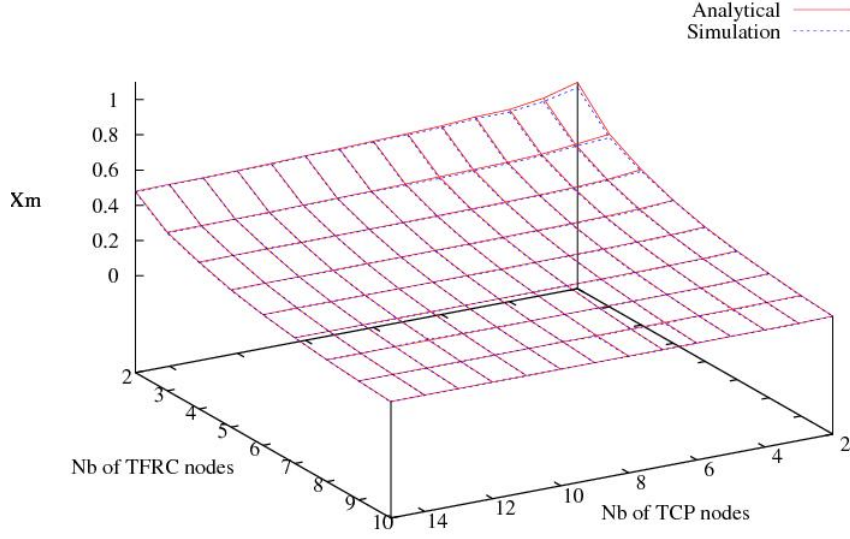


Fig 4.8 Scenario 2

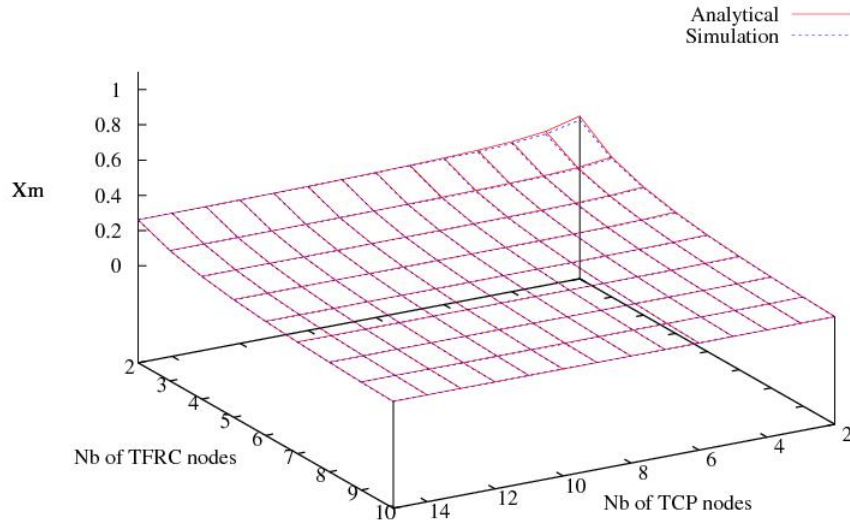


Fig 4.9 Scenario 3

Based on our analytical model, we developed a graphical interface to calculate the WLAN MAC supported bandwidth in terms of different wireless scenario settings, which will be shown in chapter 6.

4.3 Rate adaptation

In the previous section, we have introduced an analytical model to process the maximum available sending rate delivered by the 802.11MAC layer to each mobile node. When the sending rate from upper layer exceeds this rate, MAC layer buffer can be overflowed leading to potential massive packet losses due to the lack of rate control between the MAC and higher layers (as shown through our experiments in section 4.1.2). Such a pathological behavior is mainly due to the independence promoted by the OSI model between its different layers, this feature is at the origin of potential maladjustment between the different layers. This approach leads to the absence of rate control between the upper

layers and the MAC layer. Specially, for TCP (when $N_t \leq K$) and non ACK clocked flow such as TFRC, the classical WLAN MAC buffer size is not big enough to induce a buffer based close-feedback loop that would implicitly adjust the TCP or equation based TFRC sending rates from the varying RTT entailed by the buffering delay. Therefore, the adaptation of upper layer sending rate (i.e. TCP and TFRC sending rate) to the available rate supported by MAC layer is mainly done from the packet losses in the MAC buffer, which entails as seen previously a potentially dramatic sub-optimal reduction of the TCP and TFRC sending rates. In this section, we propose a cross layer based-approach resulting from a cooperation between the MAC and the upper layers that aims to avoid MAC layer overflow and results in the improvement of the QoS (that is loss rate, bandwidth, delay and jitter) delivered to TCP, UDP and TFRC flows. More precisely, we have experimented two levels of cross layer interactions that implement respectively flow control either at the network layer on flow aggregate and at the transport layer on a flow per node basis.

4.3.1 IP layer shaper

We propose a generic rate based flow control mechanism, implemented as an IP layer traffic shaper, which can be fruitfully applied to a multiplex of UDP, TCP and TFRC flows. It is worth noting that such an approach, when applied to the aggregate multiplex of higher transport flows, entails no change to the various transport layers. This network layer shaper is integrated at the IP layer (with a defined buffering capacity B) of each mobile node to adapt the aggregated sending rates from transport layers to the delivered rate supported by the MAC layer (X_t). As underlined by the simulation and real experimental results, this approach reduces packet losses at the MAC layer and rate variations, and as a consequence improve the QoS delivered to applications' flows. From a practical point of view, the processing of the MAC layer available rate is operated by the AR or AP, this offered rate can be calculated from dynamic parameters that can be easily monitored in real time (e.g. PLCP frame fields of received packets, MAC address, etc.) and broadcast to the mobile nodes by the AR/AP. Sending and receiving rate profiles can also be taken into account when rate sparing nodes exist.

When considering self clocked TCP flows or equation based congestion controlled TFRC flows, because the buffer size at the IP layer is much bigger than the MAC layer one and the resulting buffer-based close-feedback loop allows implicitly adjusting their sending rate according to RTT variation resulting from IP buffering, therefore this approach leads to significant losses reduction at the MAC layer.

For TFRC flows, several feedback packets can be configured to be sent back to TFRC sender during one RTT in order to better and more quickly adjust the TFRC sending rate. However, the hypothesis of our proposal is that the bottleneck of the flow is not situated inside the network but at the WLAN MAC layer (last meter network). Meanwhile, our proposal allows efficiently avoid losses at the MAC layers, therefore the loss event rate of TFRC can keep decreasing slowly until overflow at the IP buffer (in a "perfect" communication channel). A set of simulation done for TFRC flows have proved that this IP overflow is much less aggressive than the frequent overflows at the MAC layer for the traditional TFRC cases, that is mainly due to the IP buffer based close-feedback loop which implicitly adjusts the TFRC sending rate.

IP layer shaping can be also applied to the TCP flows when $N_t \leq K$. In this case, when the aggregate throughput from TCP transport layer exceeds the rate X_t , then RTT increases because of the buffering at IP layer, which implicitly impacts on the TCP sending rate. In other words, this close feedback loop allows the aggregate TCP sending rate varying around the shaped limited throughput X_t according to the varying RTT.

Compared to the constant MAC buffer size, the size of IP buffer can be parametrized and dynamically adjusted. We consider that our proposal introduces a “novel equivalent MAC buffer” that is configurable to replace static and constant MAC buffer size which is implemented in the wireless card and not modifiable. However, the trade-off have to be found between large buffer sizes which allows progressive rate self-adaptation with low rate loss but potential high delays and small buffers that can entail frequent losses at the benefit of low end to end delays.

When considering UDP flows, their sending rate is also constrained by the IP shaper to the MAC layer available rate. When the sending rate from these “not rate and congestion controlled transport flows” exceeds X_t , an exception is raised by the IP layer and propagated by an upcall to the application layer that is aware of its rate maladjustment and can adapt its behavior. There is a clear advantage of the IP shaper solution over the default blind behavior of the UDP based applications. For instance when considering a UDP based video transmission, the application can apply codec change or video resolution reduction (i.e. VLC transcode functions) to lower its sending rate, and then to avoid MAC losses and support continuous video transmission.

4.3.2 Transport rate limiter

We have addressed another cross layer flow control proposal that focuses on the implementation of a rate limiter inside the transport layer. However, this approach can only be implemented within rate-based protocols (i.e. TFRC) on a per flow basis and therefore, is less generic than the IP shaper solution. Indeed, in the context of TFRC flows, the processed TFRC equation based sending rate X_{tfrc} can be bounded by the maximum transport layer bandwidth X_t offered by the MAC layer, resulting in a transport sending rate X_{send} given by the following equation which a simple additional constraint introduced in the final step of the processing of the TFRC rate:

$$X_{tfrc} = \frac{s}{R * \text{sqrt}(2 * b * p/3) + (t_{RTO} * (3 * \text{sqrt}(3 * b * p/8) * p * (1 + 32 * p^2)))} \quad (32)$$

$$X_{send} = \min(X_{tfrc}, X_t) \quad (33)$$

Where, X_{tfrc} is the transmit rate in bytes/second; s is the packet size in bytes; R is the round trip time in seconds; p is the loss event rate, between 0 and 1.0, of the number of loss events as a fraction of the number of packets transmitted t_{RTO} is the TCP retransmission timeout value in seconds; b is the number of packets acknowledged by a single TCP acknowledgment.

However, such a transport layer should be applied on a flow per node basis. It can raise the issue of the fair share of the bandwidth if the different transport flows are produced by the end systems. Indeed, when several flows co-exist, their respective transport layer rate must be dynamically adapted according to the profiles of different TFRC flows. Therefore, such a transport layer approach increases greatly the complexity compared to the cross layer solution between the network and the MAC layer (IP shaper).

4.3.3 Simulation and validation of rate adaptation mechanisms

A set of wireless scenarios have been simulated under OPNET to validate our proposed methods. We present in this section three typical scenarios. In all these scenarios, we set the link bandwidth capacity of the access router to $C = 10\text{Mb/s}$ in order to have $C \gg \sum_1^N X_m$ with N : the number of mobile nodes. As a result, X_m (i.e. the MAC layer bandwidth delivered by the WLAN access network) is considered as the bottleneck

between the mobile node and the destination. The size of data frame (S_t) is equal to 8192bit (MPDU size $S = 8614\text{bit}$), the number of TCP segments acknowledged by each ACK : $K = 2$. In all the simulations, the traffic generation starts at $t = 15\text{sec}$. We show, in the following two scenarios, how the IP and transport shaper can be applied to the different data flows to avoid MAC layer overflows and substantially improves the quality of the transmission.

1.3.3.1 Scenario I: IP and Transport shaper for TFRC flows

In this scenario, we suppose that 4 TFRC mobile nodes are transmitting data to the remote server via the current AP, two of the mobile nodes keep using a transmission rate of 11Mb/s. The other two mobile nodes use the transmission rate of 5.5Mb/s between $t = [15\text{sec}, 80\text{sec}]$. When the captured signal gets better since they are moving towards the access router, their transmission rates turn to 11Mb/s at $t = 80\text{sec}$. During the first phrase $t = [15\text{sec}, 80\text{sec}]$, X_t is processed and equals to 1.14Mb/s. From $t = 80\text{sec}$, X_t rises to 1.48Mb/s. Since the bottleneck always situates on the MAC layer of each mobile node, the transport shaper allows the sending rate X_{send} (inserted in the TFRC protocol) to be adjusted to X_t to avoid congestion and losses at the MAC layer. Then we reproduce the same simulation scenario for the IP shaper approach, we suppose that there are three greedy TFRC connections for one of the moving mobile nodes (MN_{IP}), we constrain the sending rate from IP layer to X_t to alleviate the 802.11 MAC layer congestion.

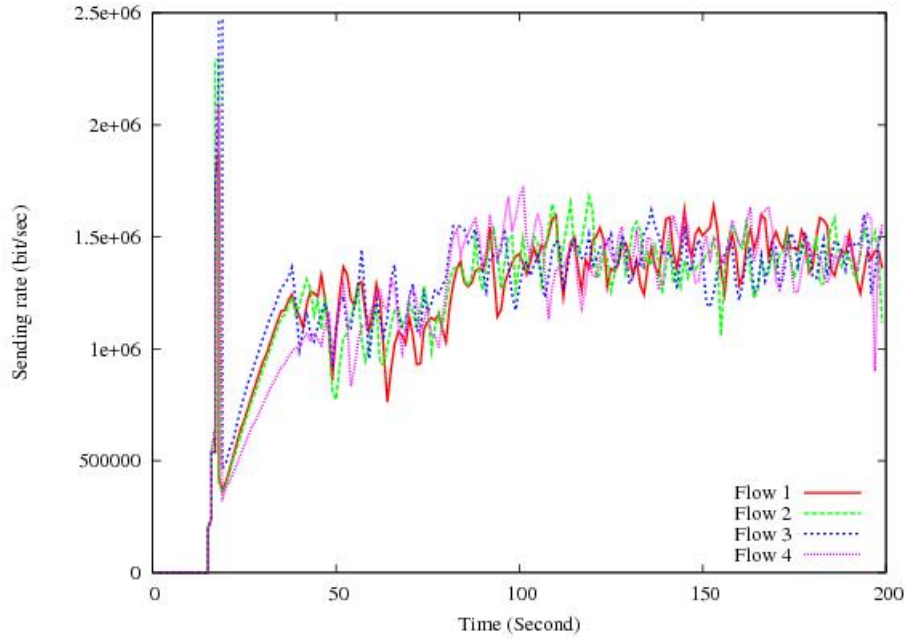


Fig 4.10 Default TFRC rate

Fig 4.10 shows the sending rate of default TFRC flows, Fig 4.11-a,b respectively show the sending rates of transport limiter based TFRC flows and the IP shaper based flow which aggregates the three TFRC sending rates of the mobile node MN_{IP} . After the slow-start phase, 1.51 packets per second on average are dropped by the MAC layer for the default TFRC flows while zero packets are dropped for Transport and IP Shaper based flows. Although oscillation of the transport sending rate can be observed when IP shaper approach is applied (as figure Fig 4.11-b illustrates), the MAC pipe is continuously fed at a constant rate of X_m . The simulation results illustrate that our proposal efficiently avoids the losses on the MAC layer and substantially improves the quality of the service delivered to application flows.

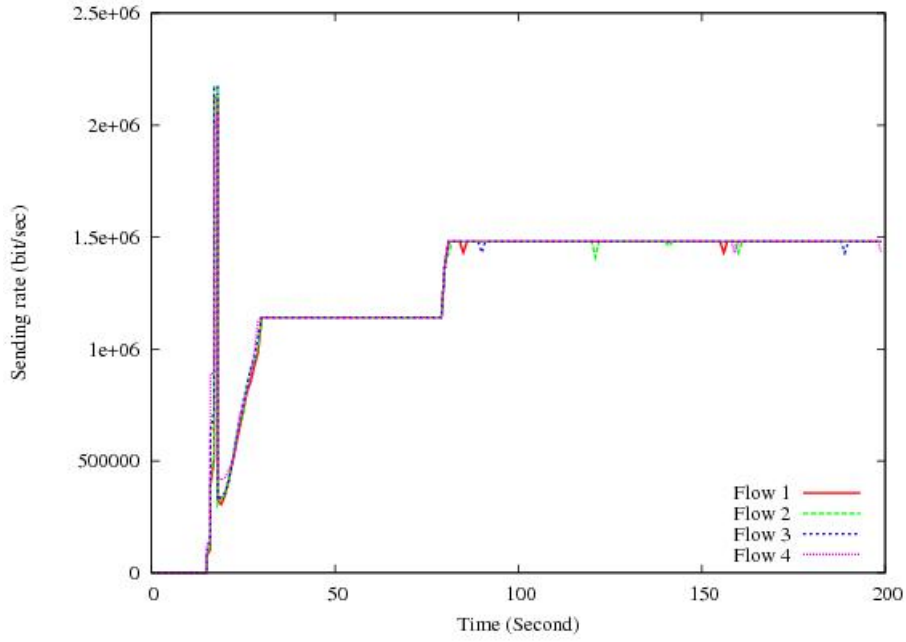


Fig 4.11-a TFRC rate transport shaper

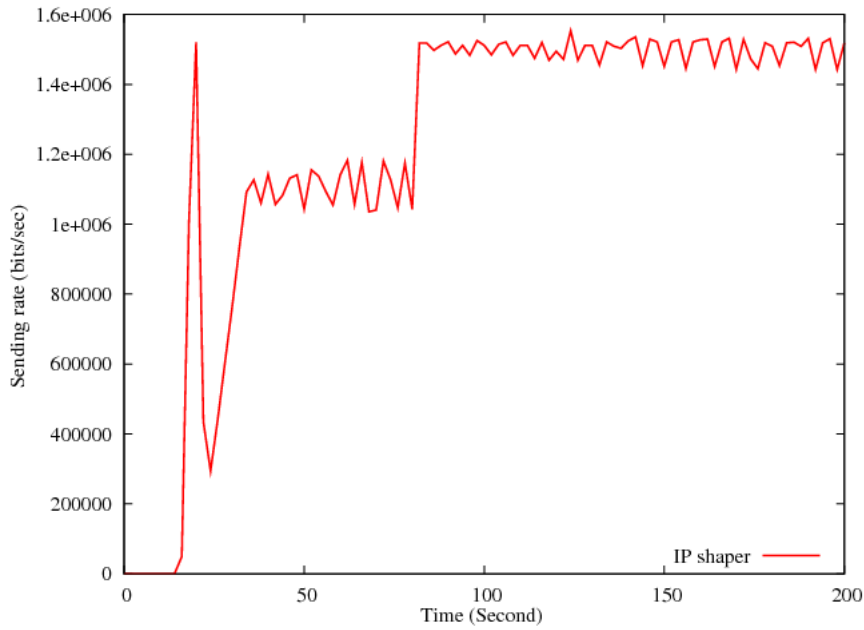


Fig 4.11-b TFRC rate IP shaper

4.3.3.2 Scenario II: UDP/TCP coexisting flows

in this scenario, we suppose that there are three TCP uploading mobile nodes (MN_{t1} , MN_{t2} , MN_{t3}) and three UDP uploading mobile nodes (MN_{u1} , MN_{u2} , MN_{u3}) in the coverage of AP. MN_{t3} and MN_{u3} always keep a transmission rate of 11Mbps. The four mobile nodes (MN_{t1} , MN_{t2} , MN_{u1} and MN_{u2}) that are moving towards the AP have an initial transmission rate of 5.5Mbps, their transmission rate rise to 11Mbps at the time of 120sec. Between the period $t = [18\text{sec}; 120\text{sec}]$, X_t is processed and equal to 839.9Kbps according to our analytical model (with $N_t^1 = 1$; $N_t^2 = 2$; $N_u^1 = 1$; $N_u^2 = 2$), then X_t rises to 991Kbps between $t = [120\text{sec}; 200\text{sec}]$ with $N_t^1 = N_u^1 = 3$. In this scenario, we constrain the UDP sending rate to X_t to avoid MAC layer congestion and losses. Fig 4.12 and Fig 4.13 represent respectively the sending rate for 3 UDP flows and the average TCP

throughput of the 3 TCP flows in the default case and when a transport layer shaper is used respectively. Fig 4.14 represents the losses at the MAC layer of each UDP node in default case, note that we don't observe any losses at the MAC layer of the UDP mobile nodes when a transport layer shaper is applied.

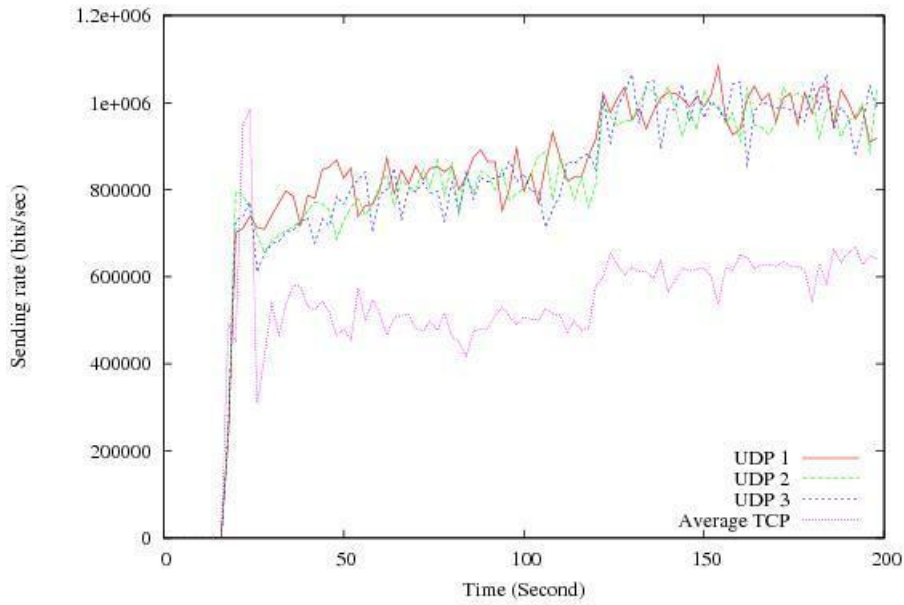


Fig 4.12 Default rate

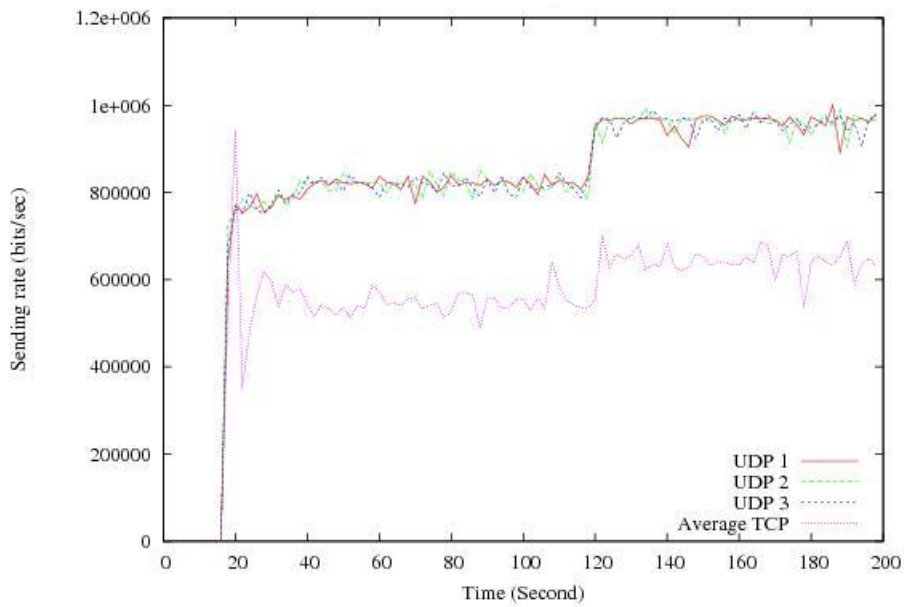


Fig 4.13 Rate with our proposal

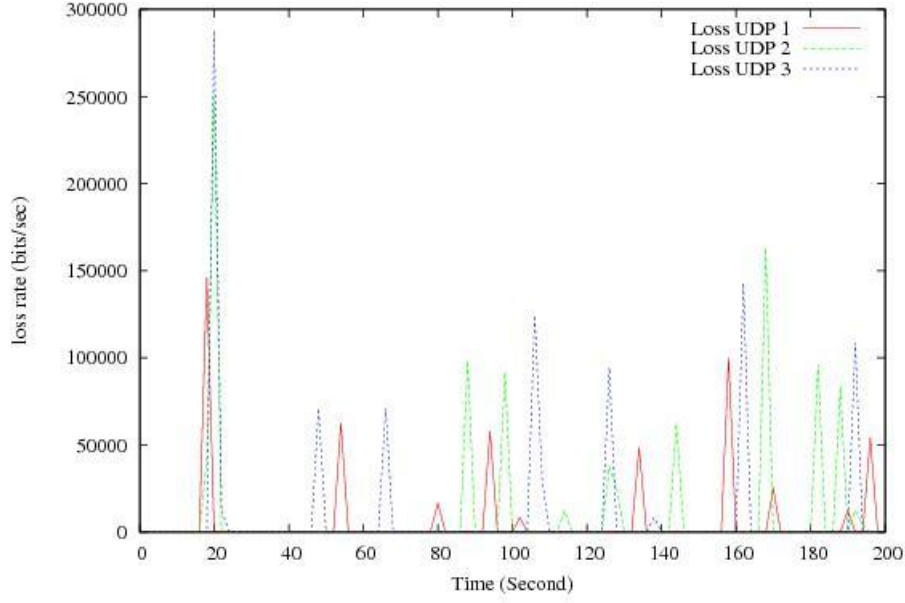


Fig 4.14 Losses at MAC layer

4.3.4 Conclusion of the section 4.3

In this section, we have presented our analytical model to accurately calculate the maximum bandwidth supported by the 802.11 MAC layer. Following a cross layer approach, we have shown how this processed bandwidth can be used at the IP or transport layer in order to introduce efficient rate control mechanisms that prevent MAC layer overflow. Simulation results show that this proposal significantly improves the quality of transmission. Furthermore, the results of our analytical model can also be inserted in the rate-equalization mechanism which is to be presented in the following section, to resolve the unfairness problems in the context of 802.11.

4.4 Rate equalization mechanisms

Among the various performance syndromes intrinsic to the CSMA/CA access method, unfairness is one of the most challenging issues. We introduce in this section an approach for addressing two main unfairness problems in the 802.11 context: 1) Unfairness between UpLink (UL) and DownLink (DL) flows. 2) Unfairness between ACK clocked and Non-ACK clocked flows. Different from other proposals that have addressed to these unfairness issues, our proposal are based on the previously introduced upper layer analytical model and entails no change to the current 802.11 standard.

4.4.1 Unfairness issue between UL and DL flows

In the context of 802.11, access point (AP) has to compete for sending packets to the downloading mobile nodes. It is considered as a normal competing node and has the same probability of sending packets as any of the uploading mobile node in long terms which is proved to be the root of the unfairness problem between the UL and DL flows. This unfairness symptom is more critical for the Non-ACK clocked flows (i.e. UDP, TFRC flows). Take TFRC flows for instance, then as previously underlined, if we suppose that U TFRC mobile nodes are uploading streams to remote servers via AP, and that D TFRC flows are concurrently sent by this same AP to D TFRC downloading mobile nodes, then the average throughput of each upload flow equals to the aggregate throughput of the

download flows sent by AP. In other words, instead of getting a fair mean throughput ratio of $\frac{D}{U+D}$, the aggregated download flows get an unfair throughput ratio of $\frac{1}{U+1}$. We introduce in this section a rate-equalization mechanism based on the modeling results issued from section 4.2 to make the downloading flows gain a fair share of the throughput delivered by the MAC layer.

The analytical model in section 4.2 allows calculating the maximum bandwidth supported by the MAC layer for any of the competing uploading mobile node ($X_m = X_u$, where X_u represents the throughput for each UL mobile node). Since the AP is considered as a normal transmission mobile node, the average sending rate of AP to downloading nodes (the aggregated downloading rate) equals to that of each UL mobile node, $X_{AP} = X_u = D * X_d$, where X_d is the average bandwidth for each downloading flows.

The aggregated bandwidth (X) exchanged between the AP and all the mobile nodes is given by:

$$X = U * X_u + D * X_d = (U + 1) * X_m \quad (34)$$

Our proposed rate equalization mechanism makes it possible re-assigning this total bandwidth X more fairly to each of the mobile nodes (download or upload mobile nodes). Indeed, each mobile node should get a bandwidth of :

$$X_{fair} = \frac{X}{U + D} = \frac{(U + 1) * X_m}{U + D} \quad (35)$$

In case of TFRC streams, by considering both the fair share constraint and the MAC rate constraint, we constraint our transport layer sending rate to:

$$X_{send} = \min(X_{tfrc}, X_t, X_{fair}) \quad (36)$$

Our proposal doesn't touch the lower layers; it limits the sending rate of each of the U uploading mobile node at the transport level to X_{fair} . As a consequence this approach allows alleviating the sending competition by the UL nodes at the MAC layer and therefore, allows AP to gain more sending opportunities, which corresponds to an improved bandwidth of $(X - U * X_{fair})$ for the DL flows. Thus, each download node can get a bandwidth of X_{fair} . Simulation results, as shown in section 4.4.3, demonstrate that such an end to end approach ensures a fair share of the bandwidth between upload and download flows. Once again note that this approach can be applied either at the transport layer or on flows' aggregation at the IP layer.

4.4.2 Unfairness issue between ACK clocked and Non-ACK clocked flows

We have mentioned in section 4.2.1.2 the "ACK effect" for ACK clocked flows such as TCP connections. When the number of TCP uploading nodes (N_t) surpasses the the number of segments acknowledged by each ACK (K), the contention avoidance procedure implemented at the 802.11 MAC layer can slower the TCP sending rate because the AP cannot gain enough sending opportunity of sending ACK packets back to the uploading mobile nodes. Of course, this effect has no influence to the Non-ACK clocked based nodes that can still make full use of the bandwidth delivered by the 802.11 MAC layer. Similarly to the analysis presented above, the principal of our proposal is to constrain the sending rates of Non-ACK clocked flows to allow the AP to gain more sending opportunity for forwarding ACK packets, therefore increasing the sending rate of the ACK clocked flows.

Let's consider a set of UDP (Non-ACK clocked) and TCP (ACK clocked) flows for instance. If $N_t \leq K$, we have seen previously that UDP and TCP flows share fairly the upload bandwidth. Therefore, we focus in the following on the case where $N_t > K$. We suppose that the average bandwidth of each greedy uploading UDP node is

R packets/sec (which corresponds to X_m processed by our analytical model) and the packet sizes S are supposed to be identical for TCP and UDP flows to simplify the analysis without lack of generality, S_{ack} denotes the size of TCP ACK packet in bits. The uplink bandwidth obtained by all the TCP uploading mobile nodes are equivalent to the bandwidth captured by K UDP mobile nodes without “ACK effect”. So the aggregated bandwidth (X) exchanged between the AP and all the mobile nodes is given by:

$$X = N_u * R * S + R * S_{ack} + R * K * S \quad (37)$$

where

$$R = \frac{X_m}{S} \quad (38)$$

The object of our proposal is to re-distribute this total network resource more fairly to all of the TCP and UDP mobile nodes. If we denote X_{fair} , the sending rate of each TCP and UDP (corresponding to R_{fair} packets/sec) node after the network resource re-distribution. We have,

$$X = (N_t + N_u) * R_{fair} * S + \frac{N_t * R_{fair}}{K} * S_{ack} \quad (39)$$

with

$$X_{fair} = R_{fair} * S \quad (40)$$

where, $\frac{N_t * R_{fair}}{K}$ represents the number of TCP ACK packets that must be sent by the AP for continuously feeding the TCP flows. Therefore the fair bandwidth share is given by:

$$X_{fair} = \frac{(N_u + K) * X_m + \frac{X_m * S_{ack}}{S}}{N_t + N_u + \frac{N_t * S_{ack}}{K * S}} \quad (41)$$

Following our cross layer approach, we propose to constrain the sending rate of each UDP node to X_{fair} in order to deliver a fair bandwidth share to the TCP nodes. Therefore, the contention based mechanism allows AP to gain more sending opportunities to forward ACK packets, which corresponds to an additional bandwidth of $(\frac{N_t * R_{fair}}{K} - R)$ packets/sec. Thus, each TCP node is no longer constrained by the “ACK effect” and shares the same network resources as each UDP node.

4.4.3 Simulation and validation of rate-equalization mechanisms

We introduce in this section two typical simulation scenarios to show how the rate-equalization mechanism allows avoiding the unfairness problems intrinsic to 802.11 WLANs. Simulation parameters are configured similarly to those defined in section 4.3.3. Scenario 1 focuses on the solution to the unfairness issue between TFRC (or UDP) UL and DL flows, scenario 2 focuses on the unfairness issue between ACK clocked (TCP) and Non-ACK clocked (TFRC, or UDP) flows.

4.4.3.1 Scenario I

In this scenario, we address fairness issues between TFRC upload and download flows in this scenario. We consider that two mobile nodes upload TFRC flows (of which the transmission rates are respectively 11Mb/s and 2 Mb/s) and three mobile nodes download TFRC flows with transmission rate of 5.5 Mb/s. Fig 4.15 shows that for the default case, each upload flow occupies much more bandwidth (average ratio of three) than each download flow. Conversely, we observe a fair share of the bandwidth between the upload and download flows when implementing our rate-equalization mechanisms (Fig 4.16).

Indeed, in this case, when applying our analytical model, since the AP is considered as a upload node we have $N = 3, N_1 = N_3 = 1$, and $N_2 = 1$ (which corresponds to the aggregated 3 download nodes). Then we get $X_{up} = X_m = 968Kbps$ and $X_{fair} = X_{up} * \frac{3}{5} = 581Kbps$. The sending rate for each of the upload mobile nodes is then constrained to X_{fair} to allow downloading nodes sharing the same bandwidth than uploading nodes

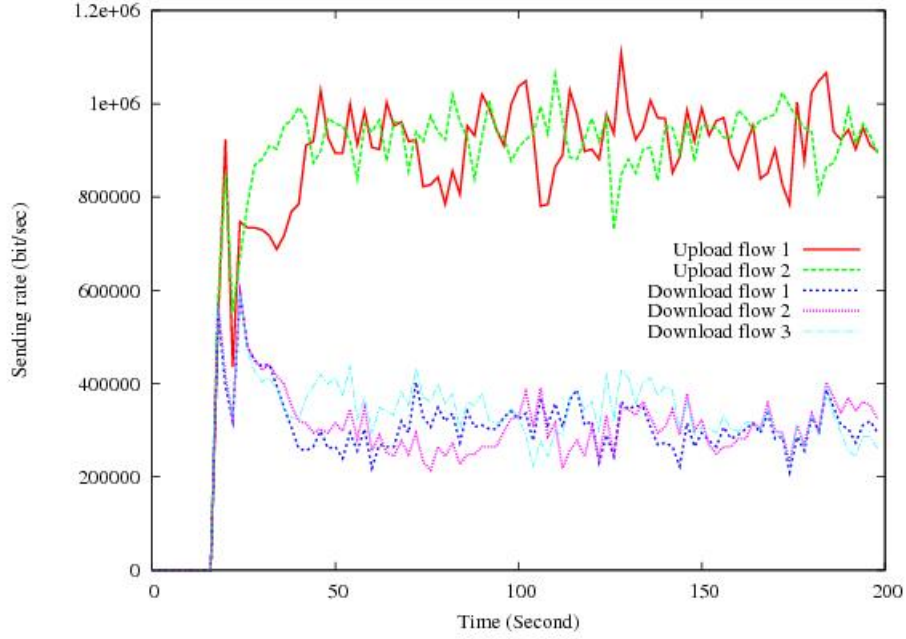


Fig 4.15 Upload and Download rates for TFRC flows in default case

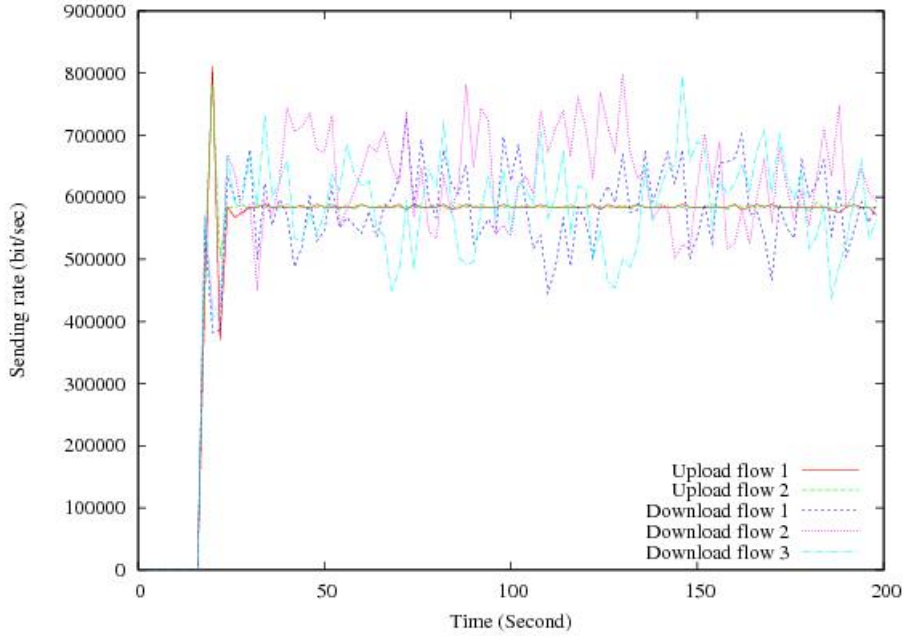


Fig 4.16 Upload and Download rates with rate equalization mechanism

4.4.3.2 Scenario II

Simulation results in section 4.3.3.2 have shown that our proposal allows suppressing losses at the MAC buffer and improving the quality of service delivered to application flows. However, we observe in Fig 4.13, each UDP connection occupies much more bandwidth compared to each TCP flow. We reproduce the same scenario setting represented in

section 4.3.3.2 while applying our rate equalization mechanism. Following our proposal, we constrain the sending rate of each UDP nodes to X_{fair} to make AP gain more opportunity of sending ACK packets, in order to increase the TCP sending rate and improve the fairness issue. According to the rate equalization model, the X_{fair} is processed and equal to 689.5Kbps between $t = [18\text{sec}; 120\text{sec}]$ and rises to 826.6Kbps between $t = [120\text{sec}; 200\text{sec}]$. We observed in Fig 4.17 that our proposal allows not only avoiding the losses at the MAC layer, but also improving the unfairness issue between coexisting TCP and UDP flows.

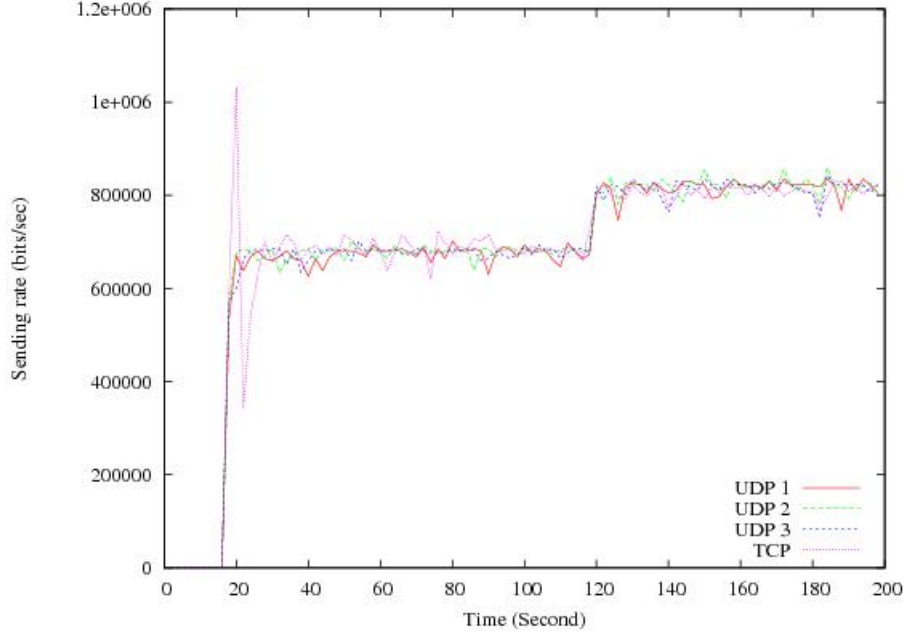


Fig 4.17 TCP and UDP flows with our proposal

4.4.4 Conclusion of the section 4.4

Based on the analytical model introduced in section 4.2, we have introduced a rate-equalization mechanisms which solves several unfairness issues that plague WLANs while avoiding modifications to the MAC layer. We apply our proposal to several specific protocols; simulation results show a significant improvement on the fairness issues.

4.5 Adaptive FEC based mechanism

In the context of 802.11, the CSMA/CA access method entails a round robin scheduling for the access to the medium by the mobile hosts attached to the same access point. As previously introduced, this long term behavior is at the origin of a syndrome called 802.11 performance anomaly caused by low rate senders lowering the performance of fast senders and inducing a strong correlation between the collective and the individual performances. Indeed, in order to optimize the individual goodput (meaning the application level throughput), the Auto Rate Feed- back (ARF) mechanism applied by the 802.11 MAC layer, dynamically adapts end-systems sending rate according to dynamic channel communication conditions. In this section, we will show that a systematic recourse to the ARF mechanism, when communication conditions worsen, is not always the best solution. We will demonstrate that cross layer cooperation between an adaptive upper layer FEC based mechanism (also known as erasure code) and the ARF mechanism can alleviate this performance anomaly and improve the overall goodput. Indeed, instead of having systematically recourse to the ARF mechanism when communication condition worsen, our

approach aims at keeping the current communication rate and increasing the reliability of the current communication context by, in first resort, applying to the packet flow a FEC mechanism of which the expansion ratio is defined according to the monitored Packet Error Rate. our solution also allows assessing the limit of this approach and makes explicit the threshold when ARF has to take over from FEC for ultimately delivering smoother and more efficient rate variations than the usual ARF mechanism alone. We show in our simulation and trace based measurements that our proposal can significantly increase the global goodput of the mobile nodes in certain communication contexts.

4.5.1 Motivation of this work

Cross-layer reliability management is definitely an open problem in wireless networks (see e.g. [145, 146]). Classical mechanisms are error correcting codes and modulation at the physical layer, erasure codes or retransmissions at the link and/or the transport layer. The choice between these possibilities must be carefully done by taking into account the properties of the channel and the interactions between the mechanisms (see e.g. [147]). From this viewpoint, a transmission using TCP/IP over 802.11b already uses a multilayer reliability system, mixing retransmissions mechanisms at the transport and at the link layer and adaptive modulation at the physical layer. The components mechanisms and the parameters of this system were tuned in order to provide the best reliability to each connection. However, an unexpected consequence of this complex system is the "anomaly" phenomenon of 802.11 [133]. Clearly, the best way to solve this anomaly is to rethink the whole reliability system. Unfortunately, due to the wide deployment of this technology, it is more realistic to not strongly modify the mechanisms of the lower layers but rather to propose new solutions for the higher layers. This is the main objective of the contribution introduced in this section. In [133], Heusse et Al. provides a performance analysis of the IEEE 802.11b networks. Through simulation and experimental work, they show that overall performances are considerably degraded when mobile hosts transmit with a lower rate than others in the same hot-spot. The contention based CSMA/CA access method is demonstrated to be the root cause of this problem known as the 802.11 performance anomaly. In [147], the authors propose a deeper analytical study of this anomaly performance and obtain similar results. Several contributions attempted to solve this anomaly issue [148-150]; however, all of them require modifications on the MAC layer, which makes difficult their wide deployment. The origin of physical layer rate variations suffered by mobile node takes its root in the Auto Rate Feedback mechanism that aims to discretely adapt mobile nodes communication rate according to the monitored channel state. The loss rate is the parameter usually monitored by wireless cards for channel state estimation purpose. Consecutive packet losses, after applying some degree of ARQ persistence, are then considered for downgrading the current selected rate in order to improve the quality of the transmission. Conversely, consecutive successful packet transmissions lead end systems to upgrade their transmission rate. Performance anomaly aside, such communication rate enslavement to the channel conditions potentially induces erratic rate changes that have a negative impact on the quality of transmission. Few contributions have tackled the performance anomaly issue with a non intrusive approach for the CSMA/CA access method. In the following of this chapter, we introduce an original cooperative scheme between the ARF mechanism at the origin of discrete communication rate variation and an upper layer Forward Error Correction mechanism. Instead of applying significant discrete rate changes when an error threshold is overcome, we give the floor first to an adaptive FEC mechanism that aims to reduce the losses observed at the current rate and coding conditions. The usual ARF mechanism takes over when the safe state of this "FEC extended channel" cannot be preserved anymore. This cooperative scheme between

FEC and ARF aims to induce a more continuous rate delivery to application flows and to reduce the 802.11 performance anomaly by preserving simultaneously individual and collective interests of mobile nodes. Furthermore, the approach promoted in this chapter entails no change in the 802.11 MAC layer if the option to enable and disable the ARQ mechanism is available.

4.5.2 Cross layer based erasure code

This section describes our proposal to alleviate the performance anomaly of 802.11 and to offer a significant gain on the aggregated goodput delivered to the mobile nodes attached to the same access point (we denote this global gain: GG) compared to the traditional scheme defined in 802.11 standard. The purpose of our contribution aims to increase the collective aggregated mobile nodes goodput while not sacrificing individual goodput. Therefore in order to estimate the individual impact of our approach we also define the individual gain (GI) of the FEC-based mobile nodes as the ratio between the goodput delivered by our approach and the one delivered by the standard.

4.5.2.1 Default rate in 802.11 standard scenario

For generality purpose, we suppose that N mobile nodes are associated to an 802.11b access point. All of them are supposed to have initially a stable transmission rate of Tr_i with $Tr_i = (11, 5.5, 2, 1)$ Mb/s for $i = (4, 3, 2, 1)$. We also suppose that N_2 mobile nodes (among the N mobile nodes) are moving away from the AP. Following the 802.11 standard, their transmission rates will be downgraded to Tr_{i-1} after 2 consecutive packets are lost over the wireless channel. The analytical model previously introduced in section 4.2 gives the maximum uploading bandwidth (R) offered by the MAC layer to each of the N mobile nodes in long term. We denote p , the Packet Error Rate that for clarity purpose is supposed to be equal for all the mobile nodes that have downgraded their rate. Normally, small values of p are observed as the adaptive modulation is efficient enough to assure a low PER. The aggregated useful uploading throughput of the N mobile nodes is given by:

$$G_a = N * R * (1 - p) \quad (42)$$

and the individual goodput for each of the mobile nodes is:

$$G_i = R * (1 - p) \quad (43)$$

These calculated uploading goodputs will be compared to the optimized goodputs resulted in the following section to process the gain offered by our proposal.

4.5.2.2 Algorithm of our proposal

In this subsection, we present the algorithm of our proposal which is based on a cooperation between the ARF mechanism at the origin of discrete communication rate variation and an upper layer adaptive FEC mechanism. The principal is illustrated in Fig 4.18.

Our algorithm is initially triggered when M consecutive packets are lost in the wireless link (i.e. $M=2$ according to the 802.11 standard). Following such a loss event, the transmission rate is downgraded in default case when standard ARF mechanism is implemented. In our proposal, instead of downgrading the rate, we maintain the current higher transmission rate (Tr_i) and adaptively enforce the communication reliability with redundancy packets added at the IP layer. The proportion of redundancy packets among the whole set of transmitted packets is called redundancy ratio and is denoted rr . The

redundancy ratio rr ($0 \leq rr < 1$) is estimated periodically every N_{pkt} packets sent by the wireless card (in our experiments, we set N_{pkt} to 50). During this N_{pkt} packets period, the MAC layer can monitor the number of packets successfully received by the AP according to the number of acknowledgments returned (N_{ack}). The redundancy ratio which is able to correct this monitored loss rate is given by:

$$rr' = \frac{N_{pkt} - N_{ack}}{N_{pkt}} \quad (44)$$

However, this processed redundancy ratio is based on a past state of the communication channel which may differ from the current one. Indeed, if we suppose that rr' is processed at instant t and directly used to build the redundancy packets, these packets will be sent out at $(t + \Delta t)$ due to buffering and processing delays at IP and MAC layers. As the channel conditions might have changed during this time interval, rr' might be no longer valid and might not be able to recover enough data at the receiver side. In order to address this issue, we introduce an estimation function ϕ that takes into account the evolution of the redundancy ratio. The rr value is estimated to predict the variation of channel condition during the interval Δt . In the experimental part of this study (section 4.5.4.2), we used a proportional estimator derived from an offline analysis of the real wireless traces used in our experiments. The evaluation of others estimator is planned in a future work.

During a transmission, losses can occur following a burst pattern. If the burst pattern corresponds to N_{max} consecutive losses (in our experiment N_{max} is set to 5), we can asset that signal conditions get weak, then the transmission rate will be immediately downgraded from Tr_i to Tr_{i-1} . We denote rr_{max} the maximum redundancy ratio. If the periodically estimated rr is bigger than rr_{max} , which means the useful throughput for the nodes implemented with FEC redundancy method is too little and is out of the tolerant range, then the transmission rate is immediately downgraded from Tr_i to Tr_{i-1} and rr is reset. Similarly to the 802.11 standard [143], the transmission rate increases from Tr_{i-1} to Tr_i when X (X can be set to 10 according to standards) consecutive packets are successfully sent to AP by a mobile node. In our proposal, rr is reset and N_{ack} is re-started to count every time when the transmission rate is changed. Fig 4.18 represents the basis of our algorithm.

In summary, when the channel conditions get worse, the mobile node, instead of downgrading its transmission rate at once, enforces dynamically the higher transmission rate with adaptive FEC in order to delay the occurrence of the performance anomaly syndrome. Such a conceptual approach raises the fundamental issues of the existence and determination of the threshold rr_{max} allowing the algorithm to switch from the ‘‘FEC extended’’ communication state to the ARF mechanism. Ideally the maximum admissible value rr_{max} , which can be used as the FEC redundancy ratio, should prevent the occurrence of the performance anomaly syndrome while preserving the goodput of the mobile nodes that enter in the FEC-extended communication state. Indeed, the proposed algorithm must prevent from replacing one anomaly by another which would lead to sacrifice significantly the individual goodput to favor the collective benefit. In the next subsection, we show that such a maximum value exists and can be processed. Therefore the adaptive FEC processed in the framework of the proposed algorithm can be applied up to rr_{max} in order to increase the collective goodput while not sacrificing significantly individual benefit.

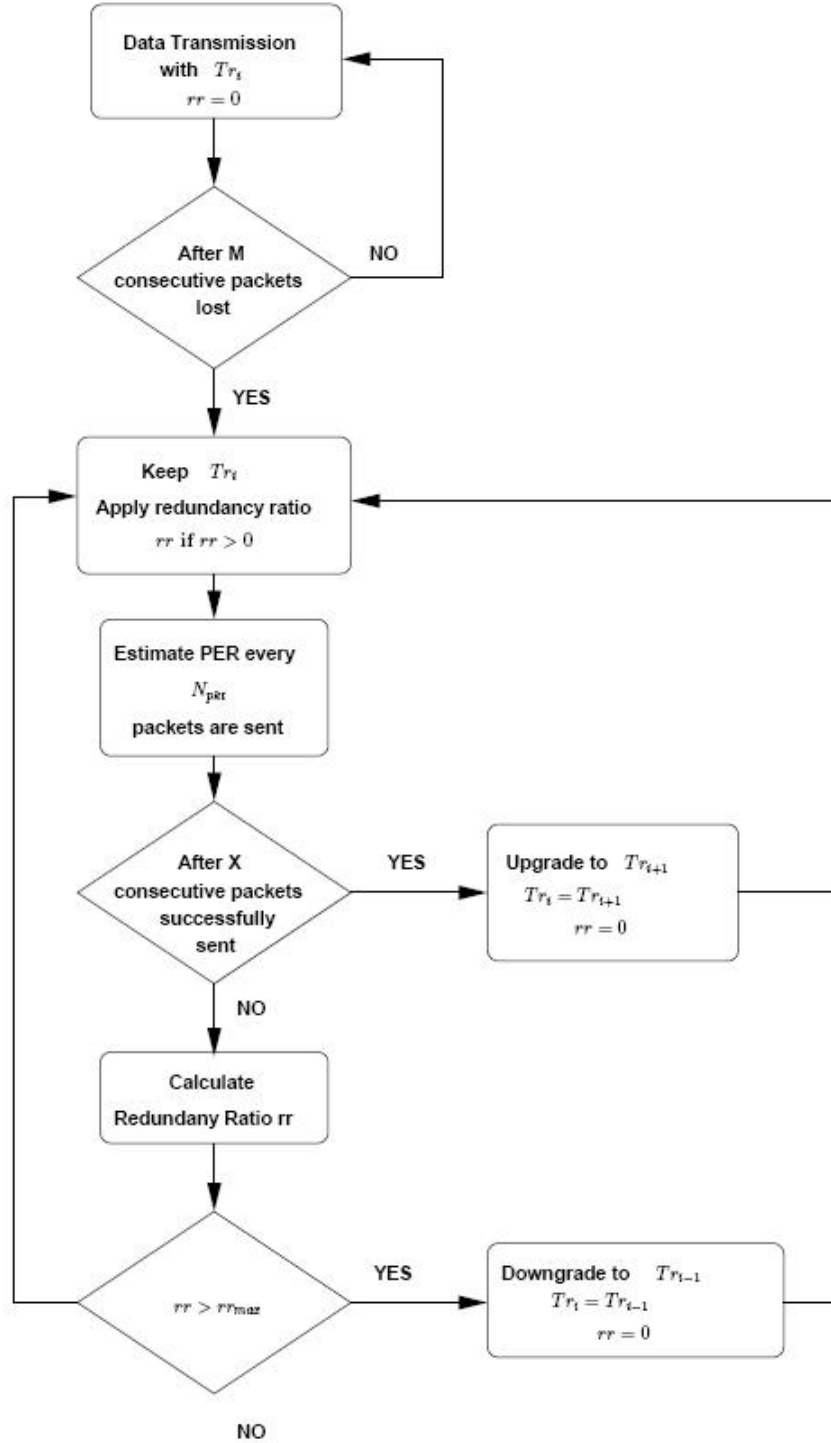


Fig 4.18 Algorithm

4.5.3 Defining the critical redundancy ratio

Our proposal allows trying to maintain the high transmission rate (Tr_i) of the N2 mobile nodes even if their signals get poor. According to the analytical model presented in section 4.2, we can calculate the maximum uploading bandwidth (R_{FEC}) supported by the MAC layer for each of the mobile nodes. In the other hand, the estimated rr cannot guarantee to recover all the useful packets. We denote p' the percentage of the packets which cannot be recovered. We denote p , the PER of channel losses for mobile nodes with good signal (i.e. the N1 nodes which have not activated the FEC mechanism), then the aggregated uploading goodput of all the N mobile nodes is:

$$G'_a = N1 * R_{FEC} * (1 - p) + N2 * R_{FEC} * (1 - rr) * (1 - p') \quad (45)$$

and the uploading goodput for each of the N2 mobile node is:

$$G'_i = R_{FEC} * (1 - rr) * (1 - p') \quad (46)$$

We suppose that the adaptive modulation is efficient enough to ensure a low PER (p), and we also suppose the applied redundancy ratio allows recovering the lost packets at the receiver side, then p and p' can be considered as negligible. In order to obtain a higher global goodput with our proposal compared to the standard, we should have:

$$N1 * R_{FEC} + N2 * R_{FEC} * (1 - rr) > N * R \quad (47)$$

When p and p' are negligible, we have:

$$rr < \frac{N * (R_{FEC} - R)}{N2 * R_{FEC}} \quad (48)$$

Therefore, rr_{GG} , the maximum redundancy ratio threshold allowing a global rate gain is given by:

$$rr_{GG} = \frac{N * (R_{FEC} - R)}{N2 * R_{FEC}} \quad (49)$$

In order to obtain a higher individual goodput for each of the N2 mobile nodes with our proposal compared to the standard, we should have:

$$R_{FEC} * (1 - rr) > R \quad (50)$$

Which means:

$$rr < 1 - \frac{R}{R_{FEC}} \quad (51)$$

Therefore, the maximum redundancy ratio threshold denoted rr_{GI} that delivers an individual gain is given by:

$$rr_{GI} = 1 - \frac{R}{R_{FEC}} \quad (52)$$

Theorem 1. Whatever the number of mobile nodes N , there is a redundancy ratio rr that satisfies simultaneously both global and individual gains (respectively denoted as GG and GI).

Proof. Let's demonstrate that the individual redundancy ratio is always lower than the global redundancy ratio, that is: $rr_{GI} \leq rr_{GG}$ (i.e. (40) $\leq$ (37)), indeed:

$$\frac{R_{FEC} - R}{R_{FEC}} \leq \frac{N * (R_{FEC} - R)}{N2 * R_{FEC}} \Leftrightarrow N2 \leq N$$

Since $N2 \leq N$ is always true, all the values in $(0, rr_{GI}]$ offers both individual and global gains.

The resulting rr_{GG} and rr_{GI} are represented as a function of $N1$, $N2$ in Fig 4.19 where $Tr_i = 11Mbps$ with adaptive FEC.

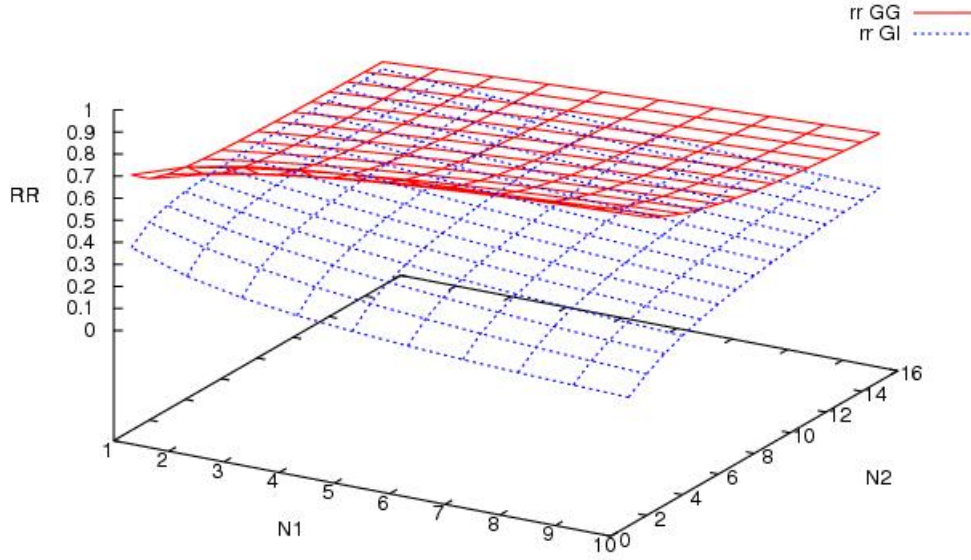
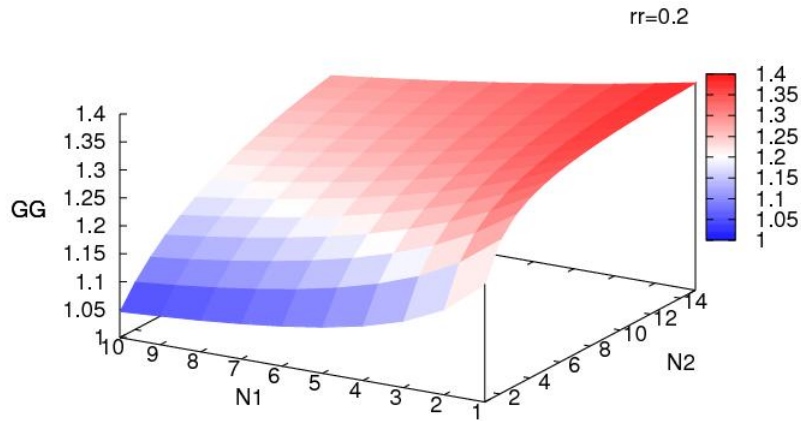


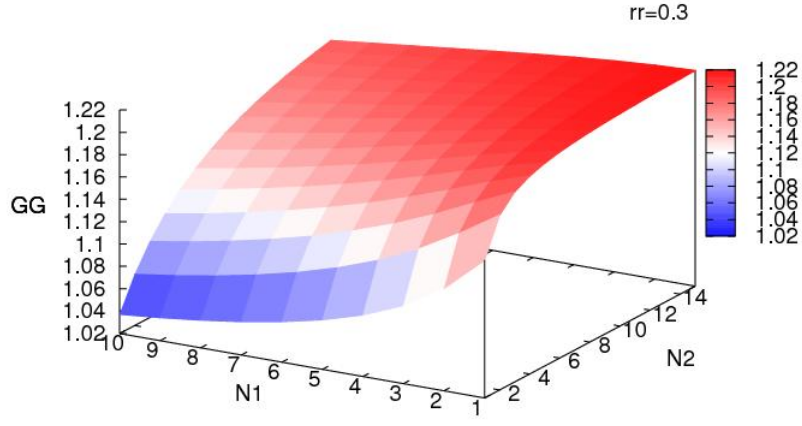
Fig 4.19 Value of RR in terms of N1 and N2

Fig 4.19 shows the existence of two distinct regions. The first one associated to a redundancy ratio rr , lower than rr_{GG} but higher than rr_{GI} , which delivers a global gain while sacrificing goodput for each of the $N2$ mobile nodes that have recourse our adaptive FEC proposal (i.e. $GG > 1$ and $GI < 1$). The second region is defined by the values of rr lower than rr_{GI} where we have both global and individual gains ($GG > 1$ and $GI > 1$). The previous theorem shows that the redundancy ratio, which can be applied to a stream transmission in order to avoid a brutal rate decrease, covers an interval that can potentially lead to diverse policies going from the stringent respect of both individual and global goodput to the respect of global goodput only. This can result for instance in priority based policies that, according to the priority given to a node, can endeavor to more or less preserve the individual goodput of the considered node.

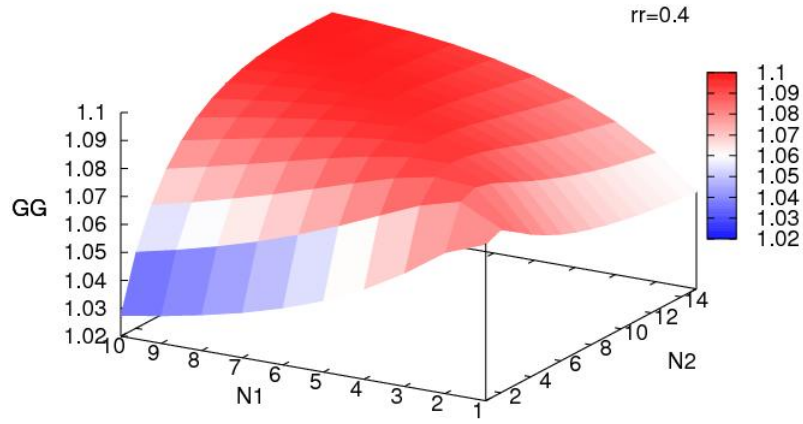
Moreover, we can observe in Fig 4.20 that a high redundancy ratio can also impact negatively on the global gain. The policy function that applies an optimal trade-off between the global and individual goodput is highly dependent of the flow's types and will be carefully studied in our future work.



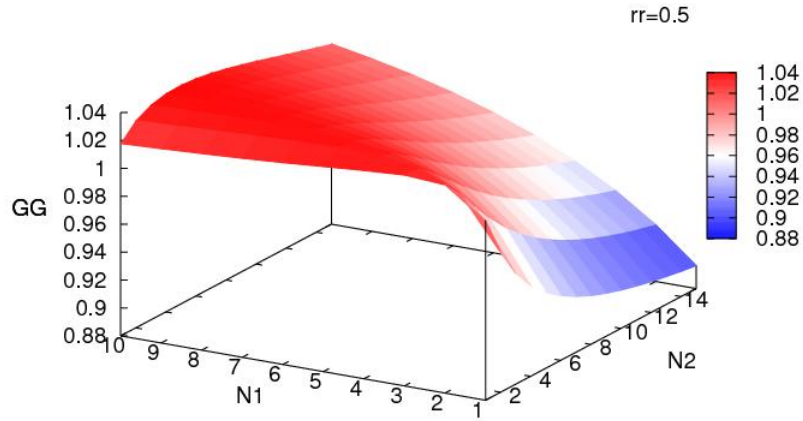
(a)



(b)



(c)



(d)

Fig 4.20 Global gain with a given rr as a function of $N1$ and $N2$

4.5.4 Validation of our proposal

This section focus on the validation of our adaptive FEC based proposal through some experimental results.

4.5.4.1 Proof of concept

We have experimentally measured the PER (Packet Error Rate) as a function of SNR for different transmission rates in the context of 802.11b in order to assess the compatibility of the analytical model with real wireless signal behavior (Fig 4.21). These measurements show that there is a clear covering of the curves of contiguous rates for a significant interval of SNR values. These results fit with the ones given in [144]. Therefore, these covering intervals give room for applying, at a constant given rate and up to a maximum redundancy ratio, the proposed adaptive FEC enhanced communication scheme which enhances flows goodput compared to a brutal rate downgrade.

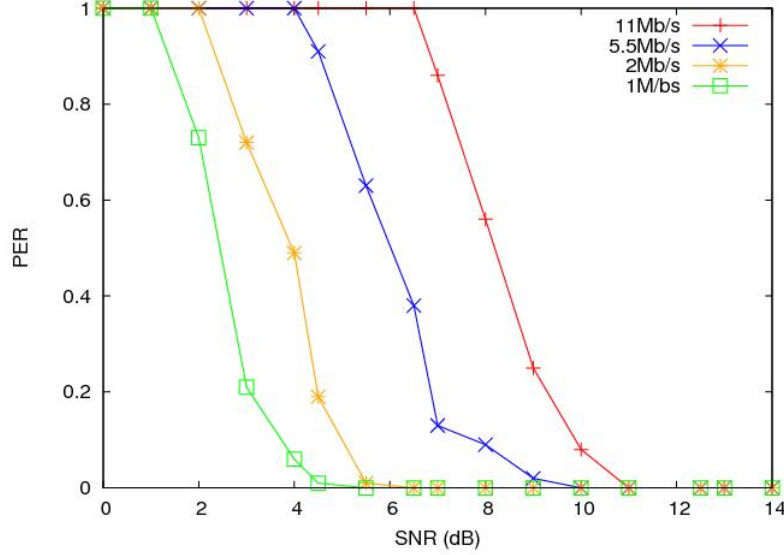


Fig 4.21 PER as a function of SNR

The above Fig 4.21 shows that, when the signal (SNR) degrades, the FEC implemented mobile node can maintain the current transmission rate ($Tr_i = 11\text{Mb/s}$ in the illustrated example) until the estimated redundancy ratio reaches rr_{max} which is defined as the critical threshold (represented here by the PER value of around 0.25). The transmission rate is downgraded after this maximum threshold. Such method allows efficiently “delay” the applying of the lower modulation and alleviate the 802.11 anomaly performance.

Furthermore, instead of applying significant discrete rate changes resulted from transmission rate degradation, our proposal gives the floor first to an adaptive FEC mechanism that aims to recover the losses observed at the current rate and coding conditions. As a consequence, such an approach replaces the traditional discrete rate evolution by a more continuous and optimal one, and offers continuous goodput evolution when SNR degrades.

4.5.4.2 Experimental scenario

This subsection focuses on the wireless traces used to validate the proposal. The traces consist in a tcpdump output of UDP data and redundancy packets lost during a communication between 4 mobile hosts and one AP. During the experiment, three mobile nodes have a good signal coverage and use a transmission rate $Tr_i = 11\text{Mbps}$ ($N1 = 3$). The fourth mobile node is initially in a poor signal coverage and then moves towards the good signal coverage zone for around 65 seconds and returns to its initial point ($N2 = 1$). During its move, the mobile node’s transmission rate remains at 11Mb/s when the estimated rr is less than rr_{max} (we set $rr_{max} = 0.35$ in our experimentation). When the signal continues to degrade and the estimated rr is out of tolerated range ($rr > rr_{max}$), then the transmission rate of the moving node N2 has to be downgraded to 5.5Mb/s. In

Fig 4.22, curves (C1) and (C2) represent respectively the goodput of the moving node and the global goodput of the 4 mobile nodes with the 802.11 standard behavior (in this case, the transmission is directly degraded to 5.5Mb/s when signal is poor).

We can see in this figure that our proposal maintains a high transmission rate of 11Mb/s by applying FEC redundancy packets at the IP layer. The redundancy ratio rr is estimated periodically and applied to the IP layer packet when rr is lower than $rr_{max} = 0.35$. In this experimental test, we simply used a proportional estimator for the function ϕ previously introduced. Indeed, this function is defined as follows: $rr = k * rr'$, where k results from an offline analysis of the loss evolution processed on our traces. We set k to 1.45 during our experiment to guarantee that 97.2% of useful data packets can be successfully received at the receiver side. According to this approach, the estimated rr ranges from 0.2 to 0.31. In Fig 4.22, C3 and C4 represent respectively the individual goodput of the moving node and the global goodput of the 4 nodes when our FEC adaptive approach is used. Compared to the standard, although the goodput of the FEC based mobile node slightly decreases ($GI = 0.93$), we can observe a significant global gain of $GG = 1.12$.

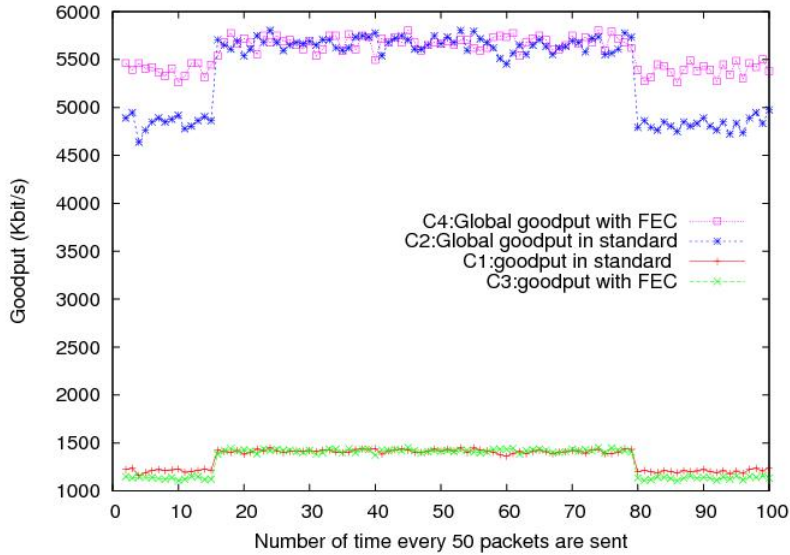


Fig 4.22 Validation result

4.5.5 Conclusion of section 4.5

In this section, we introduce an original cooperative scheme between the ARF mechanism at the origin of discrete communication rate variation and an upper layer Forward Error Correction mechanism. Instead of applying significant discrete rate changes when an error threshold is reached, we give the floor first to an adaptive FEC mechanism that aims to recover the losses observed at the current rate and coding conditions. The usual ARF mechanism takes over when the safe state of this “FEC extended channel” cannot be preserved anymore. This cooperative scheme between FEC and ARF aims to induce higher goodput and a more continuous rate delivery to application flows and to reduce the 802.11 performance anomaly by preserving simultaneously individual and collective interests of mobile nodes. Furthermore, the approach promoted in this chapter complies with the end to end approach followed throughout this thesis, and entails no change in the 802.11 MAC layer if the option to enable and disable the ARQ mechanism is available. Our future work will focus on the definition of efficient algorithms to predict the redundancy ratio and the definition of several joint rate-FEC adaptation policies that take into consideration the QoS constraints and optimality criteria.

4.6 Conclusion of chapter 4

In this chapter, based on an analytical model of the rate delivered by the WLAN MAC layer, we have proposed three novel cross-layer based mechanisms which allow solving three symptomatic issues of 802.11. Such analytical model results from a cross layer interaction between the lower and higher layers of the WLAN architecture, and can be used to apply efficient rate & congestion control and error control mechanisms to suppress the MAC layer overflows and enhance the unfairness issues. Furthermore, we have also introduced a upper layer FEC based mechanism that allows efficiently alleviating the 802.11 performance anomaly and improve the transmission goodput. Different from the other proposals, our end-to-end based mechanisms respect the current 802.11 access method and induces no change to the standard, which makes our proposals easy to deploy.

Chapter V

5 Enhancements of WiMax modulation scheme efficiency

The emergence of portable terminals in daily life and the dramatic and continuous growth of the “mobile and wireless Internet” greatly increase the needs of offering higher capacity, higher reliability, and more advanced multimedia services to wireless mobile users. The IEEE 802.16, as a complementary access network to IEEE 802.11, offers several features that can potentially respond to this acute needs in global wireless access for mobile end systems. The IEEE802.16 standard proposes an adaptive modulation scheme which allows WiMax nodes to communicate from various modulation coding according to the link quality. However, the standard does not define a detailed link adaptation algorithm and currently, the most largely used modulation adaptation technique is based on a channel quality lookup table. We argue that this method is not able to make the best adaptation decisions and delivers a sub-optimal goodput in numerous communication contexts. In this chapter, following the same approach as indicated in the previous chapters, we introduce a novel cross layer-based modulation adaptation mechanism which combines the use of adaptive erasure code with the PHY layer information to significantly improve the goodput and transmission efficiency. Simulation results show that our proposal adapts more efficiently to real environments and achieves a significant gain on the goodput delivered to mobile nodes.

5.1 Motivation

The IEEE 802.16 standard defines within its scope four PHY layers according to the various modes of system operation (i.e. NLOS and LOS environments, license and license-exempt frequency bands). Specially, the IEEE 802.16e [151] is designed to incorporate the current competitive technologies in communications and digital signal processing to deliver broadband Internet experience to nomadic or mobile users over a wide or metropolitan area. Link adaptation plays an essential role for the efficient use of the radio resources in WiMax networks. IEEE 802.16 defines a radio link control (RLC) framework to enable the implementation of PHY layer adaptation schemes. This framework includes the definitions of the signal indicator message and signal flow for link adaptation. Currently, the largely accepted adaptive modulation mechanism is based on a channel quality lookup table, from which the modulation and coding scheme is chosen according to the monitored channel quality such as SNR, SINR or CINR value. However, this method is not always able to adapt efficiently to different terrain types (SUI_i models) [152] or when different FEC codings are integrated [153]. Therefore, while being simple, such a lookup table based approach entails sub-optimal modulation/coding choices that lower the useful throughput delivered to end systems and impact negatively on the transmission efficiency of WiMax networks. In response to this issue, [154] proposes a dynamic threshold link adaptation algorithm for different channel conditions. [153] proposes a cross-layer link adaptation algorithm based on the QoS demand and channel condition. [155, 156] intro-

duce a mean carrier-to-interference-noise ratio (CINR) based link rate adaptation scheme, where the modulation mode is upgraded or downgraded when the updated mean CINR is out of the required range of the current PHY mode. However, all of these work either significantly modify the standard or are not able to efficiently adapt to real varying channel condition. In this chapter, we propose a cross layer-based modulation adaptation mechanism which incorporates upper layer (i.e. IP layer) adaptive erasure code with the PHY layer information to significantly improve the goodput and transmission efficiency.

5.2 FEC-based modulation adaptation

There are two main approaches in the current commercial market of WiMax deployment: 1) QoS supported deployment according to the standard. 2) Pre-wimax deployment which assigns the mobile nodes certain time slots and frequencies according to their service layer agreement. At lower layers, both of these solutions implement an adaptive modulation scheme which is classically based on channel quality lookup tables. In approach 1, five QoS classes have been defined in the standard; specially, the throughput of UGS (Unsolicited Grant Service) connection is guaranteed and does not depend on the link quality, which means the BS has to assign more time slots or frequencies to the mobile nodes when they use lower modulation when the signal is worsening. This approach can decrease the network capacity offered to other nodes attached to the considered base station. This behavior entails a syndrome similar to the well known performance anomaly issue observed in the context of 802.11 [157] and addressed in the previous chapter. In approach 2, the base station assigns up to a maximum number of time slots or frequencies to the mobile nodes for data transmission. Since mobile WiMax is a TDD based technology; there is a strong interest in maximizing the goodput for given assigned time slots and frequencies. According to the WiMax standard, transmission reliability can be optionally insured from an ARQ mechanism applied at the MAC layer; however, such mechanism potentially impacts negatively on the end to end delay and reduces the transmission throughput [158]. In order to obtain a good trade-off between reliability, delay and throughput, we propose to set the ARQ retry timeout to zero (timer not scheduled) to avoid retransmission [161] and still get feedback about the channel state from ACK/NACK messages. Based on a cross layer approach, this low layer feedback information feeds a high layer adaptive FEC mechanism that aims to improve channel reliability while delivering high goodput and low end to end delay. This high layer FEC mechanism is based on Maximum Distance Separable Convolutional Codes [164] which is compatible with the implementation in [163]. Such an approach can be applied both to Up-Link (UL) and Down-Link (DL) flows. In this chapter, the analysis focuses on the UL flows.

5.2.1 Conception view of our proposal

Our proposal focuses on a novel modulation decision mechanism which maximizes the goodput of mobile nodes and improves the transmission efficiency. In other words, the approach promoted in this chapter aims to replace the default staircase behavior, resulting from the use of a SNR lookup table, by a more “continuous” one that delivers a higher goodput and lowers end to end delay. The improvement is achieved by applying a high level (i.e. network layer) adaptive FEC mechanism which allows correcting packet errors according to the monitoring of packet error rate and the instantaneous SNR at the BS side (for UL flows) and MN side (for DL flows). Our approach is based on a self-learning decision engine that defines a dynamic mapping function which allows the mobile nodes, by considering the current SNR and packet error rate, to select the modulation/coding scheme that delivers the best goodput. Following this decision mechanism described in the next section, the optimal modulation/coding scheme corresponding to the best goodput

is then dynamically selected. Compared to the default modulation adaptation mechanism, our proposal takes into account not only the signal to noise ratio, but also the instantaneous packet error rate corresponding to the dynamically varying communication conditions. We will show in the following sections that our high layer adaptive FEC proposal makes it possible to delay the recourse to a lower modulation when signal degrades and delivers a better goodput. Moreover such an approach avoids systematically resorting to a brutal discrete rate reduction entailed by a sudden modulation change.

5.2.2 Mapping function

In our proposal, the ARQ mechanism is enabled in order to receive feedback of the channel state from ACK/NACK messages. According to the standard, Service Data Units (SDU) are partitioned into fixed-length ARQ blocks (ARQ-BLOCK-SIZE, of which the size is limited to 2040 bytes for fixed WiMax [161] and 1024 bytes for mobile WiMax [151]). At the sender side, the ARQ Block Error Rate (A-BLER) can be directly derived from the returned ACK/NACK information (represented by selective ACK bitmaps) for a continuously monitored sliding window of N transmitted ARQ blocks. Moreover, in order to have an easier and more efficient error control mechanism, we consider that the IP layer fragments its Network-PDU to a maximum size of ARQ-BLOCK-SIZE (1024 bytes in our proposal). That is, each MAC-SDU can be associated by the MAC layer to an ARQ block which is in turn encapsulated into a MAC-PDU, an approach also suggested in [160]. Therefore, as a result, the packet error rate (PER) is equivalent to the ARQ Block Error Rate and can be calculated by a simple analysis of the returned ACK/NACK information.

In order to define an efficient FEC adaptation policy that aims to dynamically supply to channel losses, we dynamically update a mapping function defined for the current communication context of the mobile node. This continuously evolving mapping function dynamically maintains, according to the mobile nodes communication experience, the packet error rates (p_{SNR}^i) when a given modulation/coding scheme i is used and a SNR value (defined with a precision step of 0.1dB in our proposal) is observed. For different SUI channels, this mapping database is initially populated with the packet error rate values resulted either from standard IEEE channel models or from simulations. This initial mapping database is dynamically enriched by the here after introduced self-learning mechanism that takes into consideration the dynamical channel conditions experienced by the mobile node, so introducing adaptability in the channel modeling structure.

This mapping function maintains an adaptive knowledge database that results from the communication experience of the mobile node. :

$$p_{SNR}^i = f(modulatio_coding_i, SNR) \quad (53)$$

From a practical point of view, we suppose that the mobile nodes monitor the packet error rate and are aware of the current Signal to Noise Ratio SNR_{cur} (at the BS side for UL approach) which are either periodically sent as feedback by BS or derived from the SNR monitored by mobile node and the transmission power parameters [162]. For the DL approach, SNR_{cur} is the monitored Signal to Noise Ratio that can result from the Channel Quality Indicator field. The SNR variations can be smoothed periodically with an exponential moving average:

$$SNR = (1 - k) * SNR + k * SNR_{cur} \text{ where } k \in [0; 1] \quad (54)$$

The resulting smoothed SNR is approximated with a precision of 0.1dB. The instantaneous packet error rate, p_{cur} can be periodically processed by analyzing the ACK/NACK

information (represented by selective ACK bitmaps) for a sliding window of N instantaneous transmitted ARQ blocks. Similar to the SNR process, the packet error rate p_{SNR}^i corresponding to the estimated SNR and current modulation is smoothed and updated periodically by an exponential moving average :

$$p_{SNR}^i = a * p_{cur} + (1 - a) * p_{SNR}^i \text{ where } a \in [0; 1] \quad (55)$$

This periodic processing of the SNR and the related packet error rates make it possible to establish a dynamic database that contains the packet error rate p_{SNR}^i according to a given modulation/coding scheme i and a given SNR range. This mapping database allows the mobile node to estimate the goodput it could get if it uses a higher or lower modulation/coding scheme (i.e. when using neighbor modulations).

5.2.3 Modulation and coding scheme decision algorithm

In reaction to an observed current packet error rate, p_{cur} , in order to insure packet errors correction, we propose to apply instantaneously an adaptive FEC mechanism of which the FEC ratio is a function of the current packet error rate, $f_{cur} = g(p_{cur})$ to Network-SDUs. If we define a group of Network-SDU as a contiguous set of w Network-SDUs, then the number of redundancy packets (with the same packet size of 1024 bytes) for each group is given by:

$$m = \text{ceiling}(w * \frac{f_{cur}}{1 - f_{cur}}) \quad (56)$$

The FEC ratio $f_{cur} = g(p_{cur}) = p_{cur}$ is in a first step based on the hypothesis of an uniform packet error distribution that makes possible with such a FEC redundancy the full recovery of the lost packets. The definition of more elaborate FEC ratio function adapted to non uniform and variable packet error distributions is planned in future work that will address different channel models.

In the context of UDP based transmission, within time slots t and subcarriers s assigned by BS, we denote respectively R_c : the current UL transmission rate of a mobile node, R_l : the UL transmission rate delivered if the mobile node uses the next lower modulation/coding scheme, R_h : the UL transmission rate if the mobile node uses the next higher modulation/coding scheme. Then, the UL goodput within the assigned time slots (useful transmission rate successfully received by the BS during t) is given by:

$$G_c = R_c * (1 - f_{cur}) = R_c * (1 - p_{cur}) \quad (57)$$

The object of our contribution is to choose an optimal FEC based modulation to probe the higher communication rate while optimizing the goodput. This approach aims to deliver both an optimal individual and collective goodput to the whole set of mobile nodes connected to the considered base station. In order to take an efficient modulation-choice decision, we compare this resulting rate, G_c , to the goodput obtained if the mobile node uses either the next lower or higher modulations/coding schemes. Indeed, we consider the case where the MN uses respectively the next higher and lower modulations/coding schemes with the current smoothed SNR. From the dynamically elaborated database, we can get the corresponding packet error rates p_{SNR}^i (with $i = l$ or $i = h$), the percentage of losses with the current smoothed SNR and the next lower ($i = l$) or higher ($i = h$) modulations/coding schemes : $p_{SNR}^i = f(i, SNR)$.

Similar to the previous analysis, if we set the FEC ratio $f_i = p_{SNR}^i$, then the UL goodput becomes G_i if the mobile node uses the neighbor modulations/coding schemes:

$$G_i = R_i * (1 - f_i) = R_i * (1 - p_{SNR}^i) \text{ with } i = h \text{ or } i = l \quad (58)$$

The MN will compare G_c to G_l and G_h periodically, the optimal modulation and coding scheme that delivers the higher goodput is chosen and sent to the BS, the UL-MAP is then generated by the BS according to the updated modulation and burst profile. We don't need to compute the exact value of R_c , R_l and R_h , that depend not only of the modulation code but also of several other SLA based parameters, we just need to calculate the ratio between them for the comparison purpose. We first define the combined efficiency c as:

$$c = \text{coderate} * E = \text{coderate} * \log_2(M) \quad (59)$$

where coderate and M represent respectively the coding rate and the M-ary phase.

We define the different combined efficiency c_i as: c_1 (QPSK1/2), c_2 (QPSK3/4), c_3 (16QAM1/2), c_4 (16QAM3/4), c_5 (64QAM2/3), c_6 (64QAM3/4). Indeed, with the given time slots t and subcarriers s , the transmission rate depends on the combined efficiency:

$$\frac{R_m}{R_n} = \frac{c_m}{c_n} \quad (60)$$

Table 5.1 lists the so obtained different modulations ratios.

	QPSK3/4	16QAM1/2	16QAM3/4	64QAM2/3	64QAM3/4
c_i	1.5	2	3	4	4.5
c_i/c_{i-1}	1.5	4/3	3/2	4/3	9/8

Table 5.1 Different modulations ratios

The algorithm of our proposal is described below:

```

While (True){
  When (timer runs out) do
  {
    Timer initialized;
    Fetch the current SNR, calculate  $p_{cur}$  and update the database  $p_{SNR}^i$ ;
    Process FEC ratio  $f_i$  according to the current SNR and the next higher and lower
modulations;
    Process the comparison of  $G_c, G_l$  and  $G_h$ ;
    if ( $G_c \geq G_l$ ) and ( $G_c \geq G_h$ )
    {
      Inform BS that MN should keep its current modulation and coding scheme ;
      FEC ratio =  $f_{cur} = p_{cur}$ ;
    }
    else
    {
      if ( $G_c < G_l$ )
      {
        FEC ratio =  $f_l = p_{SNR}^l$ ;
        Inform BS that MN should degrade its coding scheme or modulation rate,  $c_i = c_{i-1}$ .
(UL-MAP is then scheduled by BS);}
      else
      if ( $G_c < G_h$ )
      {
        FEC ratio =  $f_h = p_{SNR}^h$ ;

```


*Inform BS that MN should upgrade its coding scheme or modulation rate, $c_i = c_{i+1}$.
 (UL-MAP is then scheduled by BS);}*
}
}
}

Our proposal allows determining an optimal modulation and coding scheme to obtain higher goodput for the mobile nodes with given time slots and subcarriers. The promoted approach can efficiently improve both the individual and collective transmission efficiency in the whole coverage of the considered WiMax base station. The following section will give an analysis on the global gain on the goodput delivered by the proposed mechanism, specially when UGS connections are managed by the base station.

5.2.4 Global goodput gain estimation

We will show in this section that the previously introduced FEC based adaptive modulation scheme offers a significant goodput gain to the whole set of connections managed by the base station. Take UGS connection for example, whose bandwidth is guaranteed by the BS. When the mobile node's signal condition worsen, the default modulation mechanism leads him to downgrade his modulation/coding scheme c_i . As a result, the BS has to allocate more time slots or subcarriers to guarantee the UL/DL bandwidth. Our proposal aims to delay this modulation/coding downgrade while guaranteeing the negotiated QoS and the transmission reliability (in other words, such an approach allows improving the default goodput). Indeed, this approach not only preserves the QoS of the considered mobile node but also "saves" time slots that can be used for the other mobile nodes in the WiMAX BS coverage, and as a result improves the global goodput of the network.

In order to make an analysis of the resulting goodput gain, let us consider N UGS connections of which the respective required UL bandwidths are B_k bps ($k=1,..N$). When considering a given SNR, their combined efficiency (bits/symbol) are respectively c_{def}^k ($k=1...N$) according to the default modulation lookup table [151]. We suppose that all the N connections use default modulations that can be improved by our FEC based mechanism, so with the Modulation Decision Algorithm presented in section 5.2.3, a more efficient modulation/coding scheme c_{new}^k is then chosen for the N connections with a FEC ratio f_k . Furthermore, we also suppose all the other UL connections (except for the UGS connections) occupy a number of data symbols s_i ($i \in [1, 6]$) which respectively corresponds to the different combined efficiency c_i ($i \in [1, 6]$).

The total number of symbols per second (N_{sym1}) used for N upload UGS connections considered in the default scheme is then given by:

$$N_{sym1} = \sum_{k=1}^N \frac{B_k}{c_{def}^k} \quad (61)$$

Since our proposal allows applying redundancy packets (with a FEC ratio of f_k) to the UGS connections, in order to guarantee the bandwidth B_k which is required by MN's applications, the bandwidth assigned (by BS) to the UGS connections becomes $\frac{B_k}{1-f_k}$ (which includes the required bandwidth B_k and the bandwidth for the redundancy packet transmission: $\frac{B_k * f_k}{1-f_k}$). However, although the assigned bandwidth increases, the assigned time slots (or subcarriers) by BS decrease thanks to the use of higher modulation and more efficient coding scheme c_{new}^k , (where $c_{new}^k > c_{def}^k$). Then, the total number of symbols per second (N_{sym2}) used for the N UGS connections with the proposed FEC based scheme is given by:

$$N_{sym2} = \sum_{k=1}^N \frac{B_k}{c_{new}^k} \quad (62)$$

Therefore, our approach allows $(N_{sym1} - N_{sym2})$ symbols to be “saved” and to be available for the data transmission of the other mobile nodes. The property $(N_{sym1} - N_{sym2}) > 0$ is guaranteed by the Modulation Decision Algorithm presented in section 5.2.3 (with $f_{cur} = 0$, since no FEC redundancy packets are applied in default case). If we suppose that these “saved” symbols are used by other nodes with an equivalent combined efficiency c_{other} bits/symbol, then this leads to a bandwidth gain B_{gain} :

$$B_{gain} = c_{other} * (N_{sym1} - N_{sym2}) \quad (63)$$

In the default case, we suppose that the efficiency of the non-UGS connections is modeled by the following mean efficiency :

$$c_{non_ugs} = \frac{\sum_{i=1}^6 si * ci}{\sum_{i=1}^6 si} \quad (64)$$

We denote S , the number of available symbols in each frame for UL data transmission:

$$S = Sym_{UL} * N_{subcarriers} - cont_{RNG} - cont_{BWR} \quad (65)$$

Then, the global UL bandwidth in default case is:

$$BP_{UL} = \sum_{k=1}^N B_k + (\frac{S}{T_f} - N_{sym1}) * c_{non_ugs} \quad (66)$$

Where Sym_{UL} represents the number of symbol time in a frame used for UL, $N_{subcarriers}$ represents the total number of subcarriers used for data transmission. $cont_{RNG}$ and $cont_{BWR}$ represent the symbols used for contention in UL part. T_f represents the frame duration.

If we suppose that the lower modulation in default case is efficient enough to guarantee a low error rate, so we can define the global gain ratio G as:

$$G = 1 + \frac{B_{gain}}{BP_{UL}} \quad (67)$$

5.2.4.1 Case study

In this section, for a given communication condition, we will estimate the global goodput gain that our approach can deliver to a mixed set of UGS and non UGS connections. Considering a SNR of 14.5dB, according to the default modulation lookup table, the N uploading UGS connections use the modulation of 16QAM1/2 ($c_{def} = 1/2 * \log_2(16) = 2bits/symbol$). Our proposal allows the mobile nodes to use a higher modulation/coding scheme 16QAM3/4 ($c_{new} = 3/4 * \log_2(16) = 3bits/symbol$) with a FEC ratio of $f = 25\%$ according to the dynamic mapping database introduced in section 5.2.2. If we suppose that the “saved” symbols are used by other nodes with an equivalent combined efficiency c_{other} which is the same as c_{non_ugs} (the equivalent combined efficiency for the non-UGS connections in default case), and that B_{UGS} represents the total bandwidth required by the N UGS uploading connections

$$B_{UGS} = \sum_{k=1}^N B_k \quad (68)$$

Then, global gain G can be calculated from the analysis in the previous section:

$$G = 1 + \frac{c_{non_ugs} * \left(\frac{B_{UGS}}{c_{def}} - \frac{B_{UGS}}{c_{new}} \right)}{B_{UGS} + \left(\frac{S}{T_f} - \frac{B_{UGS}}{c_{def}} \right) * c_{non_ugs}} \quad (69)$$

According to [151], $Sym_{UL} = 12$ for the UL part in each frame. $N_{subcarries} = 1120$ for UL data transmission. $cont_{RNG} = 96$ and $cont_{BWR} = 192$, T_f is the frame duration (5ms).

Following this analytical expression, the Fig.5.1 represents the global goodput gain in function of the percentage of UGS upload bandwidth (B_{UGS}/BP_{UL}) and the combined efficiency c_{non_ugs} .

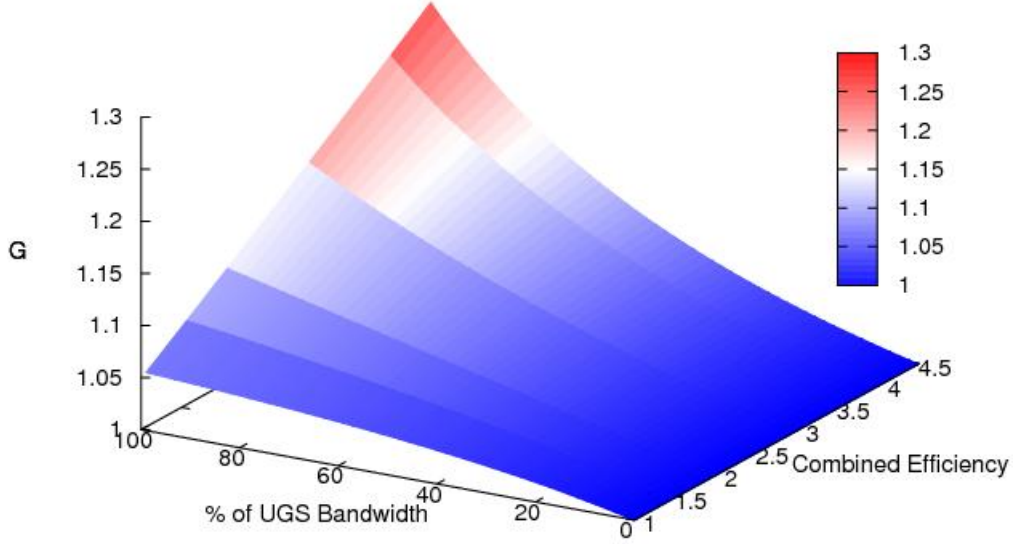


Fig.5.1 Global Gain

This section has introduced the principal of our proposal which, compared to the default behavior, allows mobile nodes to benefit from higher modulation rates while guaranteeing their transmission reliability. We have shown that this approach delivers a higher global goodput compared to the default modulation mechanism which is based on discrete and brutal modulation and rate variations.

5.3 Simulation results

This section focuses on the simulation results of our proposal as get from OPNET simulations [139]. We suppose that a mobile node is in the SUI_1 environment (Category C: Flat/light tree density) [165]. The Wireless OFDMA 20 MHz bandwidth is used in the considered scenario; the UL Sub-frame size (Sym_{UL}) is set to 5. All the available UL sub-frames and data subcarriers of the BS are assigned to the mobile node for the UL data transmission. Table 5.2 describes the simulation parameters setting.

Frequency	3.4GHz	Bandwidth	20MHz
Frame Duration	5ms	Symbol duration	102.86us
UL subframe Size	5 Symbols	Frame Preambles	1 Symbol
TTG	106us	RTG	60us
UL Data subcarriers	1120	UL subchannels	70

Table 5.2. Simulation parameters setting

5.3.1 Scenario I: Uniform movement

We suppose that the considered mobile node is moving away from the BS. As a result of this mobility scenario, the SNR monitored at the receiver side (i.e. BS) ranges from 22dB to 4dB (covering all the different modulation/coding scheme $c_i (i \in [1, 6])$).

Fig.5.2 represents the default and the FEC based goodput experienced by the mobile node. Fig.5.3 represents the gain of UL goodput with our proposal compared to the default method. Fig.5.4 represents the FEC ratio at higher layer according to the SNR value. Fig.5.5 represents the default and FEC based modulation applied (6: 64QAM3/4; 5: 64QAM2/3; 4: 16QAM3/4; 3: 16QAM1/2; 2: QPSK3/4; 1: QPSK1/2).

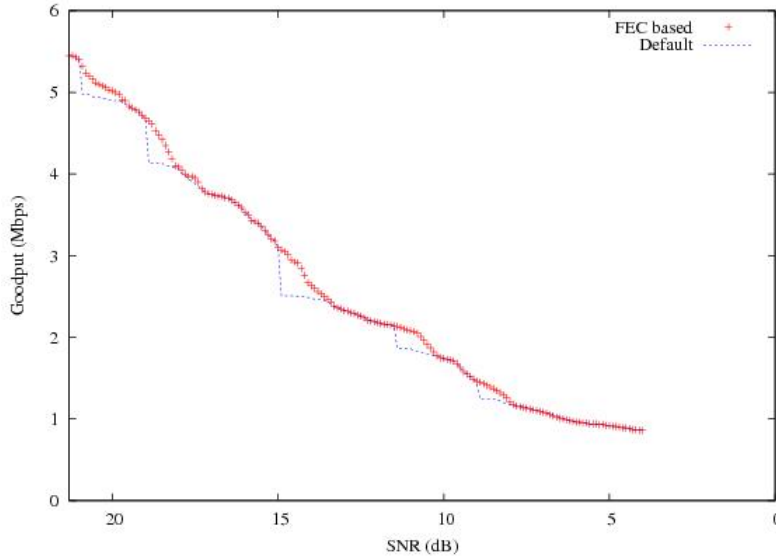


Fig 5.2 Default and FEC based UL goodput

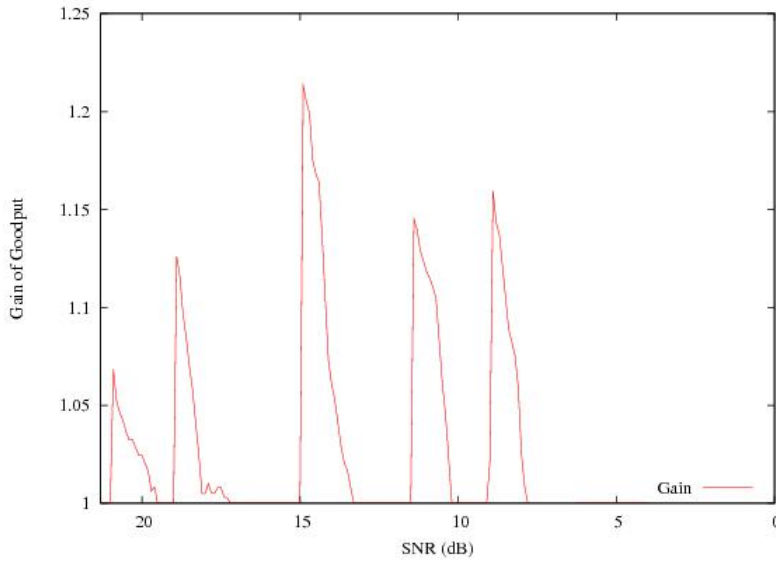


Fig.5.3 Gain of UL goodput

We observe that the goodput gain is not homogeneously spatially distributed along the radius of the WiMAX coverage. In the following, we will estimate the percentage of the WiMax coverage, approximated by a disk, where a gain can be observed. In order to estimate this ratio, we will use the model introduced in [159] for the processing of the distance between the BS and a MN in open air (D) according to the monitored SNR:

$$E = \frac{PE[dBm] + 10\log(GE)[dB] + 10\log((GR)[dB]) - SNR[dB] - N[dBm]}{20} \quad (70)$$

$$D = \frac{\lambda * 10\exp(E)}{4 * \pi} \quad (71)$$

Where PE is the emitted power, GE is the emitter antenna gain, GR is the receiver antenna gain, N is the thermal noise and λ is the wavelength. If we suppose a frequency of 3.4GHz used for the outdoor WiMax in France and a bandwidth of 20MHz, then the thermal noise is equal to $-100.97dBm$. According to a maximum allowed Effective Isotropic Radiated Power (EIRP) of 1W, the emitters are assumed to have an emission power PE of 1W. From Fig.5.3 and the above formulars, if we suppose that the available SNR for MN is between 4dB and 55dB, we can get goodput gain in the areas where the distance between BS and MN are respectively (2030-2600m), (2980-4030m), (6360-8460m), (12400-15300m), (19500-24000m). So a gain is observed for 26.64% of the whole BS coverage.

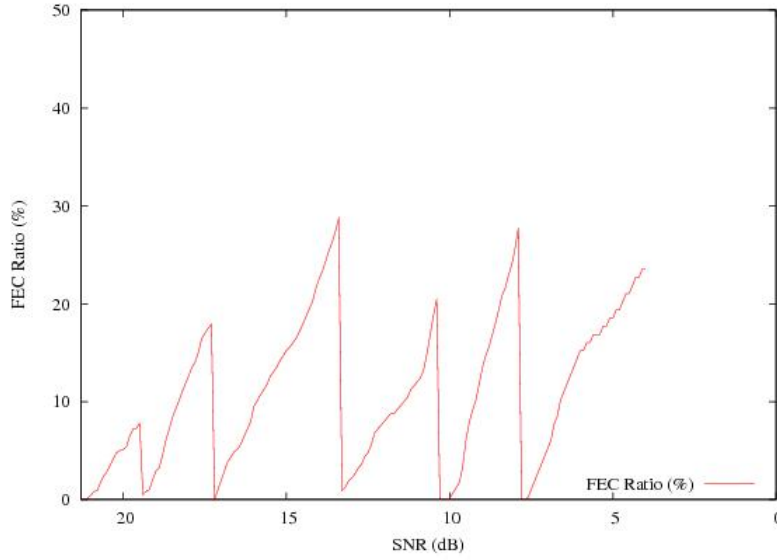


Fig.5.4 FEC ratio according to SNR

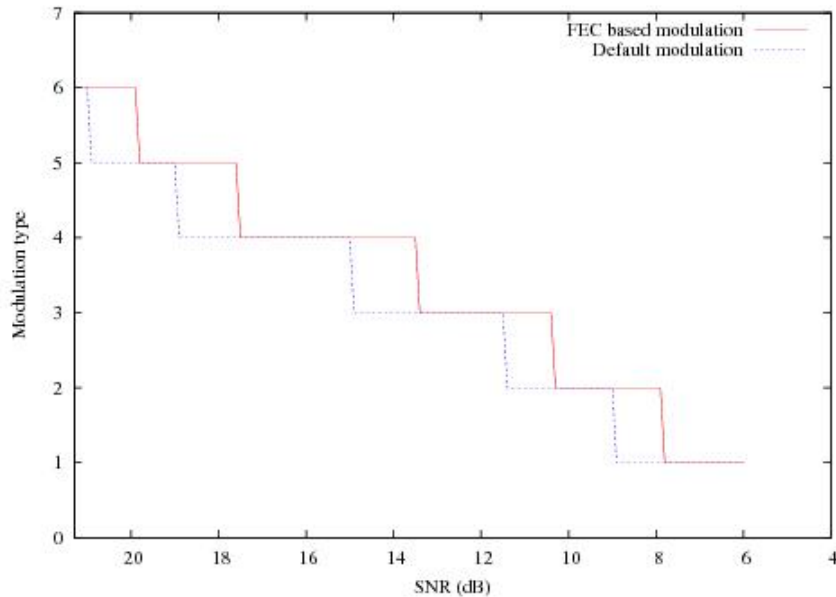


Fig.5.5 Modulation

5.3.2 Scenario II: Stop and start movement

We suppose now that a mobile node stays for 60 seconds in a place where the SNR (at the BS receiving side) is around 16dB, then it moves during 20 seconds towards a place where the SNR drops to 14.5dB, and stays there for 100seconds.

Fig.5.6 shows the SNR evolution, the default and FEC based UL Goodput. We observe that, when SNR drops to around 14.5, the FEC based mechanism offers a higher goodput (16QAM3/4 with FEC applied) compare to the default one where 16QAM1/2 is applied. The goodput gain is around 1.15.

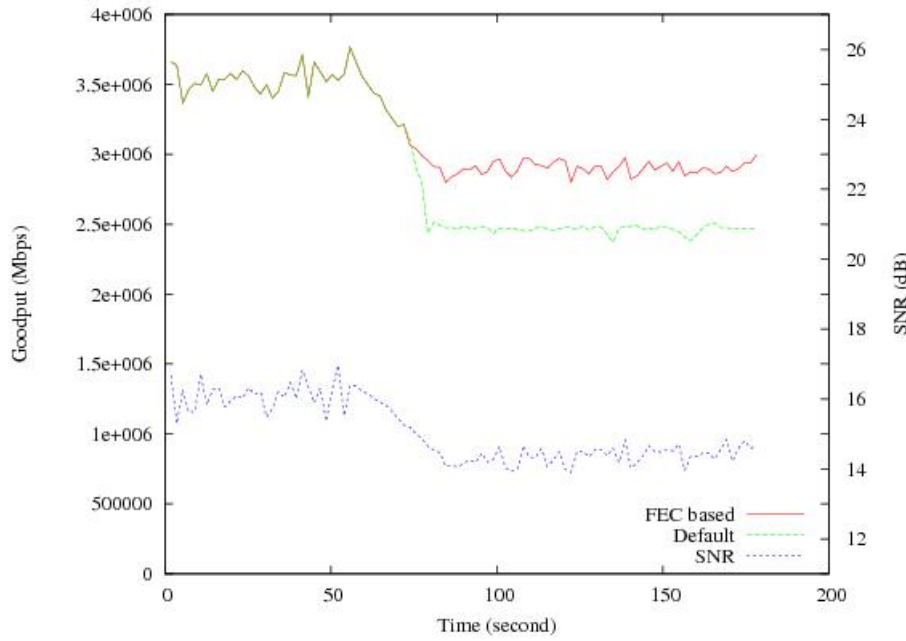


Fig.5.6 SNR evolution and Default/FEC based UL goodput

The simulation results show that our proposal always offers a higher UL goodput compared to the default method. Since the FEC mechanism is implemented in the high layer, our proposal induces no modification to the standard WiMax protocol but simply requires a simple thin interface between the MAC and network layer for the management of the mapping database and high layer adaptive erasure code. This practical consideration makes our proposal compliant with the large and fast deployment.

5.4 Conclusion of chapter 5

In this chapter we have introduced a new cross layer-based modulation adaptation mechanism which leverages on high layer (i.e. IP layer) adaptive erasure code and low layer information such as SNR and packet loss rate to significantly improve the goodput and transmission efficiency of WiMax connections. The proposed self-learning mechanism is “adaptive” to the varying channel conditions. Such an approach also results in a smoother goodput evolution rather than the default coarser evolution entailed by the brutal modulation rate changes. Indeed, instead of downgrading the transmission rate immediately when signal degrades, we propose the mobile node to keep its current high transmission rate and enforce its flows reliability with FEC redundancy packets at higher layer. The analytical study and the simulation results introduced in this chapter demonstrate a significant improvement on the goodput and transmission efficiency offered to WiMax mobile nodes.

Chapter VI

6 Implementation of an end to end architecture for improving mobility management and wireless communications

Based on the proposed mobility management architecture and mechanisms described in chapter 3-5, we have designed and implemented a Middleware for the Mobility Management Over the Internet (3MOI) [126] which is capable of supporting continuous connection; efficient physical and logical location management; efficient handover management; awareness of the dynamic mobility context. All of these features are based on the cross layer end to end approach promoted in this thesis. Fig 6.1 presents the different modules of the global structure of 3MOI, which includes: the Signal Analysis module (SA); the Energy Control module (EC); the MAC layer Optimization module (MO); the Geo-Location module (GL); the Location Management module (LM); the Mobility Prediction module (MP); the HandOver module (HO); and finally the Seamless Streaming support module(SS). The middleware has been implemented in Java under Ubuntu 7.04.

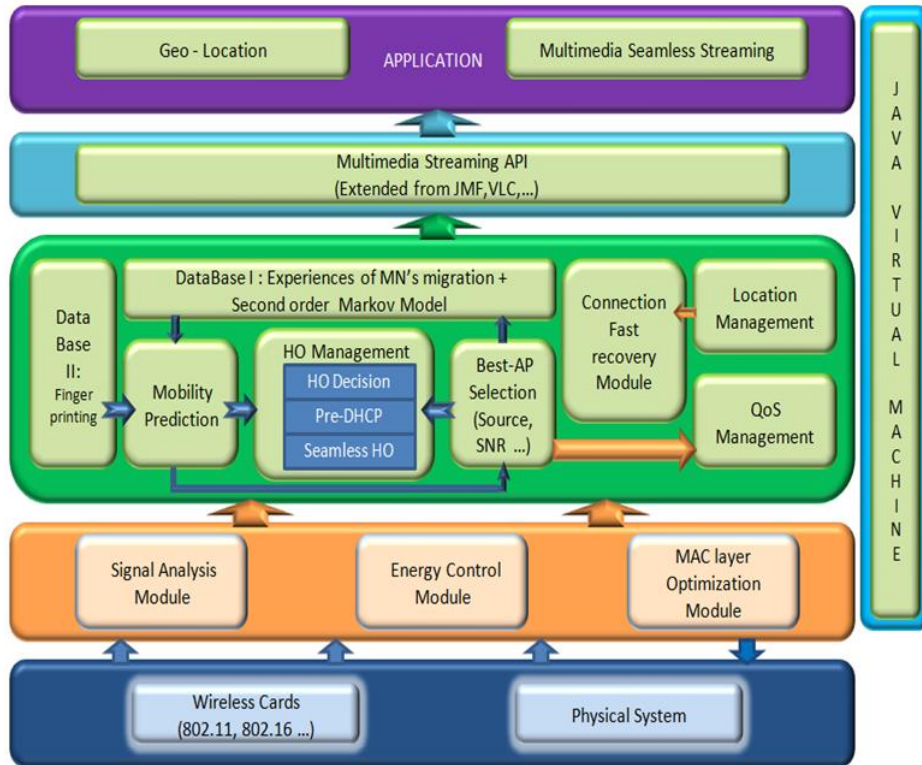


Fig 6.1 Structure of 3MOI

The 4 core modules LM, SA, MP and HO have been detailed in chapter 3. The application modules (SS, GL and EC) and MO module which will be introduced in the

following sections have been integrated in our middleware to provide: performance enhanced streaming transmission; Geo-Localization; Power optimization as well as the QoS enhancement.

6.1 Seamless streaming module

The seamless streaming module is a multimedia based extension to VLC (VideoLAN) and Java Media Framework developed by Sun Microsystems, which cooperates with the location management and handover modules to support multimedia seamless transmission.

6.1.1 Java Media Framework (JMF)

The Java Media Framework API [108] enables audio, video and other time-based media to be added to applications and applets built on Java technology. This optional package, which can capture, playback, stream, and transcode multiple media formats, extends the Java 2 Platform, Standard Edition (J2SE) for multimedia developers by providing a powerful toolkit to develop scalable, cross-platform technology.

The latest version V2.1.1 integrates an improved RTP API which provides end-to-end network transport functions suitable for applications transmitting real-time data. It also supports H.263/1998 (RFC 2429) and can inter-operate with Darwin based RTSP servers. Furthermore, it allows direct audio renderer and capturer and provides HTTP and FTP streaming support on the client side.

The reason we choose JMF as multimedia transmission tool is mainly because it offers a simple, efficient and unified architecture to synchronize and control audio, video and other time-based data within Java applications and applets.

6.1.2 Extension work to JMF

We have inserted a new mobility management module in JMF based end to end systems, which are the multimedia client and server. This module is based on the continuous connection support and optimized handover mechanism which are presented in chapter 3.

Specially, our design is integrated with other intelligent cross-layer modules such as signal analysis module and Pre-DHCP module. Indeed, this seamless streaming module is based on RTP multimedia transmission. At the client part, two ports are respectively created for the video and audio RTP transmission session. In order to improve the handover efficiency, according to the handover procedure presented in section 3.6.3.3.5, we define new formats for three management packets as shown in Fig 3.15 (I1, I2 and A2 packets in chapter 3) to respectively inform the server that the mobile client is about to *start handover procedure*; *the handover is finished*, and *transmission is re-active*. These packets are encapsulated in UDP payload in our experimental tests.

At the server side, the multimedia is transmitted to the client via unicast and the initial destination IP address and ports are set for the RTP session. In order to enhance the handover efficiency, we define the packet format of the management packet A1 (Fig 3.15) to inform the mobile client that the server is aware of the handover behavior of the mobile client and enters into sleep status for a delay of T' (please see section 3.6.3.3.5). We use the multimedia process freezing and stream re-start functions defined in JMF which support multimedia synchronization control to suspend and reactive the multimedia transmission in the server side.

6.1.3 Extension work to VLC

We developed a bandwidth-estimator based tool (M_adaptor: Multimedia bandwidth Adaptor, Fig 6.2) which cooperates with VLC to support continuous multimedia transmission.

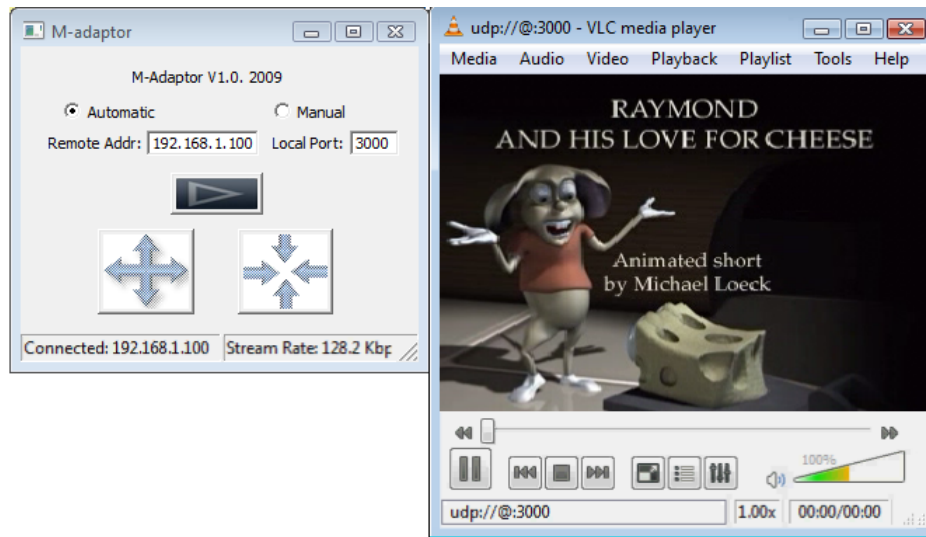


Fig 6.2 M_adaptor interface

In the context of dynamical wireless networks, the bottleneck of an end-to-end multimedia connection can vary significantly and frequently. The available bandwidth (bottleneck) for the mobile client can drop dramatically when any other mobile node downgrades the transmission rate or a slow-rate mobile node gets connected with the current access point according to the 802.11 performance anomaly. In contrast, this bottleneck can also suddenly increase if any mobile node upgrades its transmission rate or the slow-rate mobile node disconnected with the current access point. Therefore, multimedia flows which normally require a steady rate transmission can suffer in such a dynamical wireless network context. Take an example, a video transmission with a coded rate 800Kbps can no longer be decoded at the client if the bottleneck drops to 400Kbps.

The tool Multimedia bandwidth adaptor (M_adaptor) allows monitoring the available bandwidth (bottleneck) of an end-to-end media connection in real time, and informs the media server to apply an optimal codec/resolution (corresponding to the bottleneck) to the multimedia stream and therefore, to deliver continuous media streaming in the dynamical context of wireless networks. Concretely, when the bottleneck drops, the integrated intelligent algorithm analyzes the received available rate and distinguishes the case of “rate oscillation” and “bottleneck dropping”. In case of “bottleneck dropping”, it informs the server about the degraded bottleneck and demand the server to change the codec or resolution to fit the current bottleneck. Meanwhile, the integrated algorithm also allows server to send periodically “sensing data” to detect whether the bottleneck rate increases, and then adapt a better video-resolution for the media stream.

This tool also has a “Manual” mode which allows the mobile user to manually change the resolution of video in order to adapt the varying bottleneck and support a continuous media transmission.

6.2 Geo-Location module

Geo-Location module is based on the Fingerprinting technique, which cooperates with the signal analysis module that collects and analyzes the signals vectors from all the APs visible around the mobile node. The mobile node is localized by comparing the vectors

of received SNR and the database of n-uplet signal power associated with different geo-positions. This module can not only offer to mobile node the location information, but also allows improving precision of the mobility prediction, which is useful in the Pre-DHCP procedure.

Firstly, we established, by pre-measurements, a SNR database that maps the position (Position j with an interval of 1 meter) and vectors of measured SNR_i^j ($i=1,2,\dots,8$), SNR_i^j represent the monitored SNR in average from the installed 8 APs, SNR_i^j is set to 0 if no signal from APi is measured in position j. Table 6.1 represents an example of our off-line measured database.

	SNR_1	SNR_2	SNR_3	SNR_4	SNR_5	SNR_6	SNR_7	SNR_8
Position1	SNR_1^1	SNR_2^1	SNR_3^1	SNR_4^1	SNR_5^1	SNR_6^1	SNR_7^1	SNR_8^1
Position2	SNR_1^2	SNR_2^2	SNR_3^2	SNR_4^2	SNR_5^2	SNR_6^2	SNR_7^2	SNR_8^2
...								

Table 6.1 Mapping database

When the mobile nodes start Geo-location module, the mobile node collects and analyzes the signal SNR_i' ($i=1,2,\dots,8$) in real time from the APs around the mobile node, and compare them to the database represented in table 6.1, and then the position with lowest relative error is chosen, the relative error for position j is calculated as,

$$E^j = \sqrt{\sum_{i=1}^8 (SNR_i^j - SNR_i')^2} \quad (72)$$

Where SNR_i^j represents the off-line SNR from APi for position j.

In order to optimize the algorithm and lighten the calculation, we ignore and don't calculate error for the position j where $(SNR_i^j - SNR_i') > M$ for any $i \in (1,2,\dots,8)$. We set $M=15$ dB in our experiments. We don't consider the potential positions (Position j) when the difference between the measure SNR from any APi and the off-line SNR_i^j is higher than 15 dB.

The scenario of our experiments is set on the first floor of the building F in campus Jolimont of ISAE where we settled 8 APs (Fig 6.3). We firstly established an off-line database that maps between positions (zone corridor, class hall and some of the offices with 1 meter interval) and corresponding SNR_i^j which are measured with Intel PRO/Wireless 3945ABG Network card. During our location test, we walk along the corridor (in brown track), a vector of SNR_i' is scanned every 2 seconds; the green track represents our estimated positions and trajectory. According to the experimental results, we found the distance error is around 2.8 meters. We observed that distance error drops when the mobile node gets close to one or several APs.

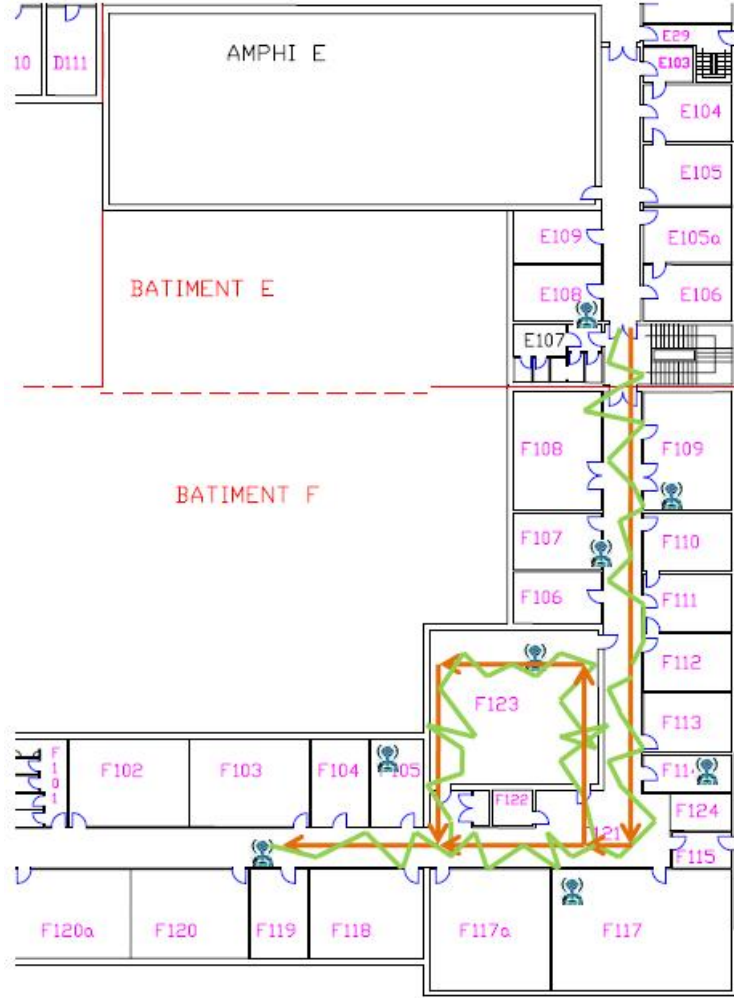


Fig 6.3 Geo-localization experiments

6.3 Energy Control module

This module collects system resources (battery, CPU utilization, etc.) in real time. In case of battery save mode or when energy is critical, it can disable some services/modules which are not indispensable and optimize the CPU utilization (i.e. disable mobility prediction and decrease the signal scanning frequency). It also allows estimating the optimal wireless network to be associated with according to the available QoS and the power level if heterogeneous access networks are available.

6.4 MAC layer Optimization module

In the context of wireless networks, the upper layers are blind to the available bandwidth resource available at the MAC layer; when the sending rate from upper layer exceeds this bandwidth, MAC layer buffer can be overflowed and lead to potential massive packet losses due to the lack of rate control between the MAC and higher layers. This discrepancy results in the degradation of the quality of transmission. In order to solve this issue, this module makes the upper layers aware of the characteristics of the MAC layer; QoS is then improved with the help of cross layer based flow control mechanism.

Based on our analytical model described in chapter 4, we developed a graphical interface to calculate the WLAN MAC supported bandwidth in terms of different wireless scenario settings. (Figure 6.4), this estimated bandwidth can then be used by the upper layers of the mobile node to efficiently avoid losses at the MAC buffer.

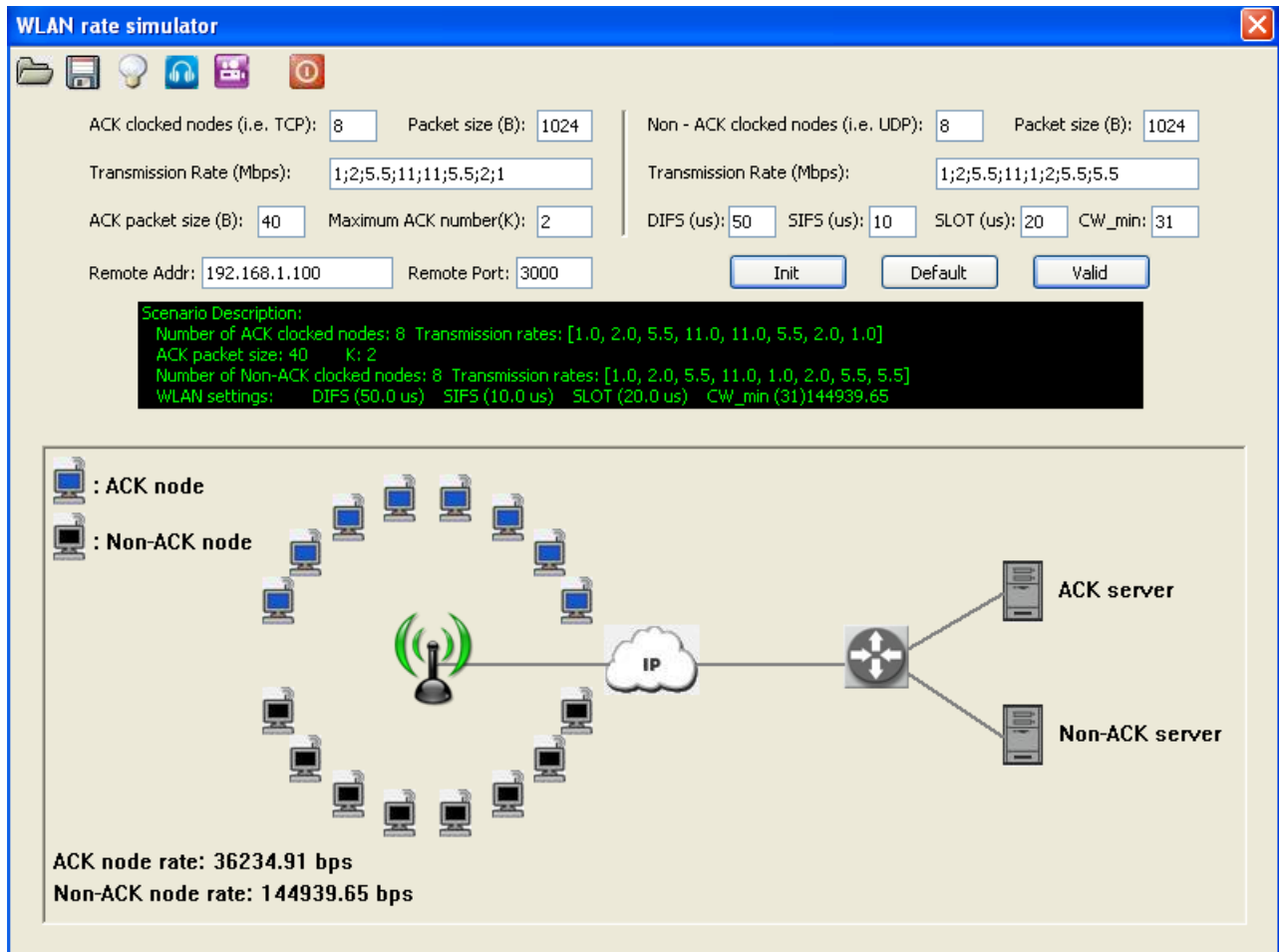


Fig 6.4 WLAN MAC bandwidth estimator

This tool has been developed to process the maximum bandwidth supported by 802.11 MAC layer for each mobile node according to the number of ACK clocked mobile nodes (i.e. TCP based nodes); number of Non-ACK clocked mobile nodes (i.e. UDP and TFRC based nodes); their respective transmission rates; data and ACK packet sizes as well as the WLAN parameter settings. The 'Default' button allows setting the default WLAN and scenario parameters. The 'Init' button allows initializing the scenario settings. The 'Valid' button allows to show description of the scenario setting, and to calculate the available rates for ACK and Non-ACK based node. It's available to open and save the scenario settings by clicking the 'Open' and 'Save' toolbar buttons. Media button allows playing a media stream from local host or Internet (by configuring the IP address and port). The 'Media_Sender' button allows sending multimedia with an optimized codec/ resolution according to the available estimated bandwidth.

6.5 3MOI Interface

The middleware that represents our global mobility management architecture has been implemented in Java. Fig 6.5 shows its user-friendly graphical interface. It comprises different modules involved in our mobility management architecture.

- Part A allows the mobile client to configure the location server IP address (i.e. DDNS, HIP RVS).
- Part B offers the access points configuration, when new access points are inserted in the experimental scenario or deleted from the experimental scenario.

- Part C allows the configuration of the mobile node. (i.e. Interface, mode of signal scanning, etc.)
- Part D offers different options for handover procedures, the mobile client is able to choose the default handover and our proposed handover procedure which cooperates with mobility prediction and Pre-DHCP. Mobile client is also able to manually connect or disconnect to any available AP by “One click”.
- Part E allows video transmission configuration and setting the peer node’s IP address (IP address of multimedia client or server) if applicable.
- Part F shows in real time the network status of the mobile client and its IP address. This feature is useful during the handover procedure to track mobile client’s status.
- Part G shows the SNR evolution for the current associated access point.



Fig 6.5 Middleware interface

- Part H lists all the APs the mobile detects and has detected (including the history of the connected APs), it integrates an intelligent self-learning mechanism to record the time corresponding to the associated and detected access points, which enriches the database for mobility prediction.
- Part I shows the information for the attached access point, which includes the measured SNR, the name, MAC address and position of the associated access point.
- Part J shows the results of the mobility prediction; a radar-like scanning interface shows, in real time, the respective possibilities of the candidate APs that the mobile node is going to visit.
- Part K represents the transmission status of data transmission (i.e. the rate of different flows).
- Part L represents the imported scenario map where the positions of different access points and mobile node is marked thanks to the Geo-Localization. The access points associated, detected or un-detected are displayed with different colors and icons.

- Part M is a JMF (or VLC) video display screen that tests the seamless multimedia transmission.

6.6 Conclusion of chapter 6

We have introduced in this chapter the implementation of our end to end mobility management architecture, based on all the contributions described in the previous chapters with additional mechanisms such as geo-localization that prefigures the intelligent high level adaptive and decision making mechanisms that will be at the core of the new generation communication architectures. The proposed communication architecture can support efficient and context-aware mobility management and satisfy new mobility requirements such as dynamical location management, fast handover, and continuous connection support.

Chapter VII

7 Conclusion and perspective

7.1 Summary

Future wireless networks will provide access to a wide range of services and enable mobile users to communicate regardless of their geographical location and their roaming characteristics. Specially, multimedia continuous streams (e.g. video or audio streams) take a bigger and bigger part of the information flows accessed or exchanged by Internet users. This feature enforces the requirement of offering seamless and broadband communication to mobile users. Therefore, reliable mobility management which aims to continuously assign and control the wireless connections of mobile nodes amongst a space of wireless access networks becomes a big challenging issue. Nowadays, In today's Internet, there is no widely disseminated, available and used communication architecture to efficiently and with a reduced infrastructure cost address the whole scope of mobility management issues. In our thesis framework, we push the end to end paradigm as far as possible for solving the following 4 issues which aim to deliver seamless and continuous communications through a space of wireless access networks: 1) Continuous connection support: an established connection should be suspended instead of being cut off during the migration of the mobile nodes, and the continuous communication is available as soon as the host gets reconnected. 2) Fast handover: minimize the duration of handover to support seamless communication. 3) Efficient location management: the current mobile node's network address should be accessible any time at a light cost. 4) Seamless communication support application as well as the related mechanisms. Different from the traditional contributions, our mobility management mechanisms and architecture take into consideration the "end-to-end" and "cross-layer based" paradigms and propose a set of novel mechanisms that can be dynamically inserted in new generation reconfigurable protocols, we offer efficient and light solutions to satisfy the mobility requirements and entail a minimum impact on the legacy protocols and Internet infrastructure.

Wireless networks which are becoming the new dominant access networks to the Internet, have introduced new lower layers quite independently of the legacy "wire-oriented" protocols that are still at the heart of the protocol stacks of the end systems. The legacy Internet protocols, still dominant at that time, have not been designed with mobility and wireless in mind. The principle of isolation and independence between layers advocated by the OSI model has its drawbacks of inadaptation between new access methods and higher-level protocols built on the assumption of a wired Internet. Concretely, we addressed in this thesis 3 principal issues: 1) the 802.11 performance anomaly, which makes the mobile node with the lowest transmission rate degrade the performances of every mobile node in the coverage of the same access point due to the 802.11 access method; 2) We identified that, due to the lack of interaction between the MAC and upper layers, the sending rate from upper layer can surpass the rate offered by the 802.11 MAC layer, and then packets can be lost in MAC buffers, which degrades the quality of transmission specially for the reliable connections. 3) Fairness degradation is one of the most challenging issues in the

context of 802.11. We addressed two main unfairness problems: Unfairness issue between Uplink(UL) and DownLink(DL) flows; Unfairness issue between ACK clocked (i.e. TCP) and Non-ACK clocked (i.e. UDP) flows. In order to resolve these discrepancies, we proposed several solutions from an original and innovative approach that, based on a “cross layer” paradigm, makes it possible to enhance transmission efficiency and solve several syndromes that plague the performances of current wireless networks. We developed an analytical model to establish a cross layer interaction between the lower and higher layers of the WLAN architecture, and can be used to apply efficient rate & congestion control mechanisms to suppress the previously introduced negative issues and enhance the unfairness issues. Meanwhile, a upper layer FEC based mechanism allows efficiently alleviating the 802.11 performance anomaly and improve the transmission goodput. Furthermore, in the context of 802.16 WiMax, we have shown that the default adaptive modulation method is not able to make the best adaptation decisions and delivers a sub-optimal goodput in numerous communication contexts. We then introduced a novel cross layer based modulation adaptation mechanism which leverages on high layer (i.e. IP layer) adaptive erasure code and low layer information such as SNR and packet loss rate. The proposed self-learning mechanism is “adaptive” to the varying channel conditions and allows to significantly improve the goodput and transmission efficiency of WiMax connections. Different from the other proposals, our end-to-end based mechanisms respect the current 802.11 access method, 802.16 framework and induces no change to the standards, which makes our proposals easy to deploy.

This thesis addresses these issues by combining analytical models, simulations and real experiments. Several softwares/tools coded in Java and Python have been developed to demonstrate the efficiency of our proposals. The resulting mechanisms have been developed and integrated into adaptive mobility management communication architecture that delivers high performing communication services to mobile wireless systems, with a focus on WiFi and WiMax access networks.

7.2 Perspective

In the framework of optimization of mobility architecture for the next generation wireless networks, our future work will focus on two principal parts: Seamless handover enhancements and MAC layer optimization of heterogeneous wireless access networks.

Seamless handover requires not only low delay, but also low/no losses during the handover procedures. In our future work, 5 approaches will be taken into consideration simultaneously:

- Optimization of link layer handover.

Minimizing the link layer handover delay without entailing significant changes to standard is a big challenge. We propose the following issues that should be carefully investigated in our future work: 1) More intelligent predictive and event handover triggers. Predictive triggers carry predictive information, including an expected time bound for the occurrence of the event and a level of confidence that the event will take place. Predictive information allows the mobile node to pre-reserve several related services in the target networks in advance to support smooth handover. The accuracy of the predictive triggers is a crucial and important issue in our future work. Specially, based on our current signal evolution prediction function, the extension work will be focused on more intelligent signal filter techniques which are compatible in different link channel models, and results in a more accurate prediction of signal evolution. 2) Efficient scanning mechanism to reduce the L2 handover delay. The object is, based on a more accurate mobility prediction and self-sniffed wireless en-

vironment, to efficiently find the channel information of the optimal target network as fast as possible.

- Optimization of upper layer handover.

In our proposed mobility management architecture, we have driven the upper layer handover delay close to zero thanks to several involved prediction and pre-reservation mechanisms. In our future work, we will integrate our work in the context of IPV6 and propose more efficient mechanisms to estimate precisely the QoS information (especially the delay) in target networks and push further the upper-layer handover delay to 0.

- Loss avoidance during the handover.

Numerous studies, which are based on the multicast routing, buffering, tunneling as well as several FEC coding mechanisms, have been proposed to avoid losses (or allow recovering lost packets) during the handover procedure. However, these proposals either require significant change to the current network standard and infrastructure or entail a long handover delay which cannot support seamless communications. Our future work will focus on balancing the tradeoff between the handover delay and losses.

- QoS guarantee.

A QoS aware mobility architecture will be addressed in the context of our future work to achieve the end-to-end QoS guarantee for mobile nodes' handover between heterogeneous networks. This architecture not only takes into consideration the low layer signal information, but also the available QoS offered by the candidate access networks (such as the loads, offered bandwidth, jitter, etc.) which supports a smooth handover on the upper layer across heterogeneous networks with different QoS supports. For example, cooperating with the rate/congestion control mechanisms, we can apply an optimal pre-reserved bandwidth to the mobile node immediately after handover instead of the slow-start procedure (for the dominant protocol TCP) in order to support a seamless communications.

- No or minimal changes entailed to the current standards and Internet stack.

Any solution at a price of entailing significant changes and additional infrastructures to the current network standard stacks is difficult to be largely adopted and deployed in practice. Keeping the above four principal aspects in mind, our future mobility architecture will cooperate with several existing standards (i.e. 802.11 framework; 802.11r which supports fast and secure handovers from one base station to another managed in a seamless manner) and seeks efficient solutions which entails no or minimal changes to the current network stack.

In the other hand, we have developed several efficient mechanisms to enhance the 802.11 and 802.16 MAC layer performances in the context of this thesis. Specially, the upper layer based FEC solution has been proved to be a big forwarding step on the optimization of the current modulation schemes in static wireless channels. However, due to the weakness of the dynamic link-layer PER (packet error rate) prediction, our proposal entails a discount on the gain of optimization in such dynamic link channels. Our current work will be extended, by integrating several intelligent prediction and self-learning mechanisms, to be compatible with different wireless channel models and result in a more significant gain on the transmission goodput. Meanwhile, we will focus on the upper layer FEC technique itself – a packet level redundancy coding. How to improve the adaptive packet level coding efficiency is also a big challenge in our future work.

Since our mobility architecture and MAC layer optimization mechanisms are “media independent” and not dedicated in some certain RATs (Radio Access Technique). We will then push further our proposals to the implementation on several different access networks (i.e. 802.16 WiMax, 3G and LTE).

8 Appendix

8.1 Estimate of t'_{jam} in UDP mode

In this section, we focus on the calculation of the average time spent in collision (t'_{jam}) to send one frame for each mobile node. When a mobile node (MN) of the N_i nodes wants to send out one frame, it risks being deferred because it exists a probability of collisions $P_c(N)$. If the collision happens, the node which causes MN deferring can be one of the other $(N_i - 1)$ mobile nodes in group N_i with a probability of $\frac{N_i-1}{N-1}$, in this case, MN has to wait minimum T_i to prepare resending; the node can also be one of the N_j nodes in group j , ($j \neq i$) with a probability of $\frac{N_j}{N-1}$, in this case, MN has to wait minimum T_j to prepare resending. So for any of the node in group N_i , the average time spent in collision to send one frame is expressed by the following equation:

$$t'_{jam} = \frac{N_i-1}{N-1} * T_i + \sum_{j=1}^4 \frac{N_j}{N-1} * T_j$$

Here, we try to find the average time spent in collision for each of the N nodes, so:

$$t'_{jam} = \frac{N1}{N} * t'_{jam} + \frac{N2}{N} * t'_{jam} + \frac{N3}{N} * t'_{jam} + \frac{N4}{N} * t'_{jam}$$

With the above equations, we have:

$$t'_{jam} = \frac{\sum_{i=1}^4 (N_i * T_i)}{N}$$

8.2 Estimate of t''_{jam} in the TFRC greedy mode

Similar to the calculation in the previous section 8.1, for a mobile node (MN) in group N_i , if a collision happens, the node which causes MN deferring can be one of the other $(N_i - 1)$ mobile nodes in group N_i with a probability of $\frac{(N_i-1)*X_m}{(N-1)*X_m+X_{AP}}$, the node can also be one of the N_j nodes in group j , ($j \neq i$) with a probability of $\frac{N_j*X_m}{(N-1)*X_m+X_{AP}}$, the node can also be the AP with a probability of $\frac{X_{AP}}{(N-1)*X_m+X_{AP}}$. So similarly, the average time spent in collision for AP and every node in group N_i to send one frame is:

$$t'_{jam} = \frac{(N_i-1)*X_m}{(N-1)*X_m+X_{AP}} * T_i + \sum_{j=1}^4 \frac{N_j*X_m}{(N-1)*X_m+X_{AP}} * T_j + \frac{X_{AP}*T_{AP}}{(N-1)*X_m+X_{AP}}$$

For the AP:

$$t'_{jam} = \sum_{j=1}^4 \frac{N_j*X_m}{(N-1)*X_m+X_{AP}} * T_j$$

The average time spent in collision for each of the $N + 1$ mobile nodes (including AP) is:

$$t''_{jam} = \sum_{j=1}^4 \frac{N_j}{N+1} * T_{jam} + \frac{1}{N+1} * T_{jam}^{AP}$$

Finally, we have approximately:

$$t''_{jam} = \frac{\sum_{i=1}^4 (N_i * X_m * T_i) + X_{AP} * T_{AP}}{N * X_m + X_{AP}}$$

8.3 Estimate of t_{jam} in the rate sparing mode for the non-frequent feedback based protocols

We denote P_i^m the proportion of throughput for each mobile node in group N_i (m is the index of the mobile nodes in each group N_i). So P_m represents the proportion of throughput for both the rate sparing nodes and the greedy nodes. For each mobile node, if its proportion of throughput increases, this node has upper hand for the contention on the MAC layer. For example, if this proportion is $q\%$, the collisions caused by this node has a probability of $q\%$. Similar to the analysis in appendix 8.1, for each node in group

Ni, the collision can be caused by the other $(N_i - 1)$ nodes or the nodes in other groups $j, (j \neq i)$. The average time spent in collision to send out one frame for AP and any of the node in group Ni can be expressed by the following equation:

$$t_{jam_i}^m = \frac{\sum_{k=1, k \neq m}^{N_i} (P_i^k * T_i) + \sum_{j=1}^4 \sum_{k=1}^{N_j} (P_j^k * T_j)}{1 - P_i^m}$$

$$t_{jam}^{AP} = \frac{\sum_{j=1}^4 \sum_{k=1}^{N_j} (P_j^k * T_j)}{1 - P_{AP}}$$

Finally, the average time spent in collision for each of the N nodes is:

$$t_{jam} = \sum_{i=1}^4 \sum_{m=1}^{N_i} (P_i^m * t_{jam_i}^m) + P_{AP} * t_{jam}^{AP}$$

8.4 Estimation of t_{jam} in the TCP/UDP coexisting case

t_{jam} represents the average time spent in collision, for a TCP mobile node in group N_t^i , if a collision happens, the mobile node which causes this node deferring can be one of the other $(N_t^i - 1)$ mobile nodes in group N_t^i , or a node in other TCP groups, or a node in UDP groups as well as the AP. The average time spent in collision for a TCP node in group N_t^i is,

$$T_{jam}^{TCPi} = \frac{(N_t^i - 1) * \frac{K}{N_t} * T_t^i + \sum_{j=1}^4 \sum_{k=1}^{N_j} (N_t^j * \frac{K}{N_t} * T_t^j) + \sum_{i=1}^4 (N_u^i * T_u^i) + \sum_{i=1}^4 (T_{ta}^i * \frac{N_t^i}{N_t})}{\frac{K}{N_t * (N_t - 1) + N_u + 1}}$$

Similarly, the average time spent in collision for a UDP node in group N_u^i is,

$$T_{jam}^{UDPi} = \frac{\sum_{i=1}^4 N_t^i * \frac{K}{N_t} * T_t^i + (N_u^i - 1) * T_u^i + \sum_{j=1}^4 \sum_{k=1}^{N_j} (N_u^j * T_u^j) + \sum_{i=1}^4 (T_{ta}^i * \frac{N_t^i}{N_t})}{K + N_u}$$

The average time spent in collision for AP is,

$$T_{jam}^{ack} = \frac{\sum_{i=1}^4 N_t^i * \frac{K}{N_t} * T_t^i + \sum_{i=1}^4 (N_u^i * T_u^i)}{K + N_u}$$

Then we can calculate the the average time spent in collision for AP and every node in group Ni to send one frame is:

$$t_{jam} = \sum_{i=1}^4 T_{jam}^{TCPi} * \frac{\frac{N_t^i * K}{N_t}}{N_u + K + 1} + \sum_{i=1}^4 T_{jam}^{UDPi} * \frac{N_u^i}{N_u + K + 1} + T_{jam}^{ack} * \frac{1}{N_u + K + 1}$$

9 List of publications

1. A Dose of Cross-Layer approach to Treat WLANs (Journal under review). Lei Zhang, Patrick Sénac, Emmanuel Lochin and Michel Diaz. IEEE Transactions on Mobile Computing.
2. Optimization of WiMax modulation scheme with a cross layer erasure code. Lei Zhang, Patrick Sénac, Roksana Boreli and Michel Diaz. IEEE International Symposium on a World of Wireless, Mobile and Multimedia Networks (IEEE WoWMoM 2009)
3. Enhancements of WLAN MAC performance. Lei Zhang, Patrick Sénac and Michel Diaz. IEEE INFOCOM 2009 Student Workshop, 19-25 Avril 2009, Rio de Janeiro, Brazil.
4. Mobile TFRC : a Congestion Control for WLANs. Lei Zhang, Patrick Sénac, Emmanuel Lochin and Michel Diaz, IEEE International Symposium on a World of Wireless, Mobile and Multimedia Networks (IEEE WoWMoM 2008)
5. A Novel Middleware for the Mobility Management Over the Internet. Lei Zhang, Patrick Sénac, Emmanuel Lochin and Michel Diaz, IEEE International Symposium on a World of Wireless, Mobile and Multimedia Networks (IEEE WoWMoM 2008)
6. Cross-layer based erasure code to reduce the 802.11 performance anomaly : when FEC meets ARF. Lei Zhang, Patrick Sénac, Emmanuel Lochin, Jérôme Lacan and Michel Diaz. ACM Mobiwac 2008. Vancouver, Canada, October 2008
7. Cross-Layer based Congestion Control for WLANs. Lei Zhang, Patrick Sénac, Emmanuel Lochin and Michel Diaz, ICST QShine, Hongkong, 2008
8. A Generic Communication Architecture for End to End Mobility Management in the Internet. Lei Zhang, Patrick Sénac and Michel Diaz, International conference on Wireless Broadband and Ultra Wideband Communication (AusWireless07), Sydney, 2007

References

- [1] Huston, G., "Gigabit TCP," IPJ, Volume 9, No. 2, June 2006.
- [2] Datagram Congestion Control Protocol (DCCP). RFC4340
- [3] Stream Control Transmission Protocol RFC2960
- [4] Tim O'Reilly (2005-09-30). "What Is Web 2.0". O'Reilly Network. Retrieved on 2006-08-06.
- [5] J. Sharony. Introduction to Wireless MIMO. Talk of the Communications Society of the IEEE Long Island Section.
- [6] J. M. Wilson. The Next Generation of Wireless LAN Emerges with 802.11n. Intel Corporation, 2004. White paper.
- [7] IEEE 802.16 WirelessMAN Standard: Myths and Facts. ieee802.org. Retrieved on 2008-03-12.
- [8] "The Long Term Evolution of 3G" on Ericsson Review, no. 2, 2005
- [9] R. Venkateswaran, Virtual private networks, IEEE Potentials, vol. 20, no. 1, pp. 11-15, Feb, 2001.
- [10] IP Mobility Support for IPv4. RFC 3344
- [11] End-to-End Arguments in System Design, Saltzer, J., Reed, D., and Clark, D.D., Second International Conference on Distributed Computing Systems (April 1981) pages 509-512, ACM Transactions on Computer Systems, 1984, Vol. 2, No. 4, November, pp. 277-288
- [12] IEEE 802.11e-2005, IEEE Standard for Local and Metropolitan Area Networks Specific Requirements – Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications MAC Enhancements for QoS, 2005.
- [13] IEEE 802.11i-2004, IEEE Standard for Local and Metropolitan Area Networks Specific Requirements – Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications Amendment 6: Medium Access Control (MAC) Security Enhancements, July 23, 2004.
- [14] Atheros Communication, Super G Maximizing Wireless Performance, www.super-g.com/atheros_superg_whitepaper.pdf, 2004.
- [15] U.S. Robotics, 802.11g Speed Acceleration How We Do It, www.usr.com/download/whitepapers/125mbps-wp.pdf 2004.
- [16] "How Bluetooth Technology Works". Bluetooth SIG. Retrieved on 2008-02-01.
- [17] "ZigBee Specification". zigbee.org. Retrieved on 2008-03-18.
- [18] Super G (wireless networking) <http://www.super-g.com/>
- [19] 802.11n: Next-Generation Wireless LAN Technology. 2006 by BROADCOM CORPORATION.

- [20] 802.16e-2005. IEEE Standard for Local and Metropolitan Area Networks. Part 16: Air Interface for Fixed and Mobile Broadband Wireless Access Systems,” February 28, 2006.
- [21] Bahai, A. R. S., Saltzberg, B. R., Ergen, M. (2004)., Multi Carrier Digital Communications: Theory and Applications of OFDM, Springer, 2004.
- [22] Wicker, Stephen B. Error Control Systems for Digital Communication and Storage. Englewood Cliffs, N.J.: Prentice-Hall, 1995
- [23] Joan Daemen and Vincent Rijmen, "The Design of Rijndael: AES - The Advanced Encryption Standard." Springer-Verlag, 2002.
- [24] RFC 3748: Extensible Authentication Protocol (EAP) (June 2004)
- [25] Soon Young Yoon, Introduction to WiBro Technology. Telecom R&D Center, Samsung Electronics Co, Ltd
- [26] Iburst ((or HC-SDMA, High Capacity Spatial Division Multiple Access), <http://www.iburst.org/>
- [27] E. Dahlman, H. Ekström, A. Furuskär, Y. Jading, J. Karlsson, M. Lundevall, and S. Parkvall, "The 3G Long-Term Evolution - Radio Interface Concepts and Performance Evaluation,"IEEE Vehicular Technology Conference (VTC) 2006 Spring, Melbourne, Australia, May 2006
- [28] C. E. Perkins, “IP Mobility Support for IPv4,” RFC 3220, Jan. 2002.
- [29] D. Johnson, C. Perkins, J. Arkko, Mobility Support in Ipv6. IETF Draft draft-ietf-mobileip-ipv6-21.txt, Fe´vrier 2003.
- [30] N. Montavont, T. Noel, Handover management for mobile nodes in IPv6 networks, IEEE Commun. Magn. 40 (8) (2002) 38–43.
- [31] S. Vaughan-Nichols, Mobile IPv6 and the future of wireless internet access, Computer 36 (2) (2003) 18–20.
- [32] H. Soliman, C. Castelluccia, K. El-Malki, L. Bellier, Hierarchical Mobile IPv6 mobility management (HMIPv6), IETF Mobile IP working Group, draft-ietf-mobileip-hmipv6-07.txt, October 2002.
- [33] R. Koodli, Fast handovers for Mobile IPv6. IETF Mobile IP Working Group, draft-ietf-mipshop-fast-mipv6-01.txt, January 2004.
- [34] P. McCann, Mobile IPv6 Fast Handovers for 802.11 Networks. IETF Mobile IP Group, 2004.
- [35] El K. Malki, H. Soliman, Simultaneous Bindings for Mobile IPv6 Fast Handovers. IETF, draft-elmalki-mobileip-bicasting-v6-03-txt, Mai 2003.
- [36] A. T. Campbell et al.,”Design, Implementation, and Evaluation of Cellular IP,” IEEE Pers. Commun., Aug. 2000, pp. 42–49.
- [37] R. Ramjee et al., “HAWAII: A Domain-Based Approach for Supporting Mobility in Wide-Area Wireless Networks,” IEEE/ACM Trans. Net., vol. 10, no. 3, June 2002, pp. 396–410.

- [38] I. Vivaldi, B. Ali, H. Habaebi, V. Prakash, A. Sali, Routing scheme for macro mobility handover in hierarchical mobile IPv6 network, in: Proceedings of the Fourth National Conference on Telecommunication Technology, 2003, pp. 88–92.
- [39] C. Chu, C. Weng, Pointer forwarding MIPv6 mobility management, in: Proceedings of the IEEE Global Telecommunications Conference GLOBECOM'02, Vol. 3, 2002, pp. 2133–2137.
- [40] K. Omae, T. Ikeda, M. Inoue, I. Okajima, N. Umeda, Mobile node extension employing buffering function to improve handoff performance, in: Proceedings of the Fifth International Symposium on Wireless Personal Multimedia Communications, Vol. 1, 2002, pp. 62–66.
- [41] C. Lee, L. Chen, M. Chen, Y. Sun, A framework of handoffs in wireless overlay networks based on mobile IPv6, *IEEE J. Sel. Areas Commun.* 23 (11) (2005) 2118–2128.
- [42] W. Lai, J. Chiu, Improving handoff performance in wireless overlay networks by switching between two-layer IPv6 and one-layer IPv6 addressing, *IEEE J. Sel. Areas Commun.* 23 (11) (2005) 2129–2137.
- [43] V. Kaffe, E. Kamioka, S. Yamada, Extended correspondent registration scheme for reducing handover delay in mobile IPv6, in: Proceedings of the Seventh International Conference on Mobile Data Management, 2006, pp. 110–110.
- [44] P. Byungjoo, L. Sunguk, H. Latchman, A fast neighbor discovery and DAD scheme for fast handover in mobile IPv6 networks, in: Proceedings of the International Conference on Systems and International Conference on Mobile Communications and Learning Technologies, 2006, ICN/ICONS/MCL 2006, pp. 201–205.
- [45] Y. An, B. Yae, K. Lee, Y. Cho, W. Jung, Reduction of handover latency using MIH services in MIPv6, in: Proceedings of the International Conference on Advanced Information Networking and Applications, 2006, Vol. 2, pp. 229–234.
- [46] D. Maltz and P. Bhagwat, TCP splicing for application layer proxy performance, IBM Research Report RC 21139, IBM T.J. Watson Research Center, March 1998.
- [47] D. A. Maltz and P. Bhagwat, MSOCKS: an architecture for transport layer mobility, *IEEE INFOCOM*, San Francisco, CA, pp. 1037–1045, March 29 - April 2, 1998.
- [48] S. Fu, L. Ma, M. Atiquzzaman, and Y. Lee, Architecture and performance of SIGMA: A seamless handover scheme for data networks, *IEEE ICC*, Seoul, South Korea, May 16–20, 2005.
- [49] A. C. Snoeren and H. Balakrishnan, An end-to-end approach to host mobility, *MOBICOM*, Boston, MA, pp. 155–166, August 6–11, 2000.
- [50] T. Goff, J. Moronski, D. S. Phatak, and V. Gupta, Freeze-TCP: a true end-to-end TCP enhancement mechanism for mobile environments, *IEEE INFOCOM*, Tel Aviv, Israel, pp. 1537–1545, March 26–30, 2000.
- [51] H. Hsieh, K. Kim, Y. Zhu, and R. Sivakumar, A receiver-centric transport protocol for mobile hosts with heterogeneous wireless interfaces. *ACM MobiCom*, San Diego, CA, pp. 1–15, September 14–19, 2003.
- [52] H. Matsuoka, T. Yoshimura, and T. Ohya, End-to-end robust IP soft handover, *IEEE ICC*, Anchorage, Alaska, pp. 532–536, May 11–15, 2003.

- [53] A. Bakre and B. R. Badrinath, I-TCP: indirect TCP for mobile hosts, IEEE International Conference on Distributed Computing Systems, Vancouver, Canada, pp. 136 143, May 30 - June 2, 1995.
- [54] Z. J. Haas and P. Agrawal, Mobile-TCP: an asymmetric transport protocol design for mobile systems, IEEE ICC, Montreal, Canada, pp. 1054 1058, June 8 - 12, 1997.
- [55] K. Brown and S. Singh, M-UDP: UDP for mobile cellular networks, Computer Communication Review, vol. 26, no. 5, pp. 60 78, October 1996.
- [56] R. H. Katz, E. A. Brewer, E. Amir, H. Balakrishnan, A. Fox, S. Gribble, T. Hodes, D. Jiang, G. T. Nguyen, V. Padmanabhan, and M. Stemm, The bay area research wireless access network (BARWAN), IEEE COMPCON, Santa Clara, CA, pp. 15 20, February 25-28, 1996.
- [57] D. Funato, K. Yasuda, and H. Tokuda, TCP-R: TCP mobility support for continuous operation, IEEE International Conference on Network Protocols, Atlanta, GA, pp. 229 236, October 28 - 31, 1997.
- [58] S.J. Koh, M. J. Chang, and M. Lee, mSCTP for soft handover in transport layer, IEEE Communications Letters, vol. 8, no. 3, pp. 189 191, March 2004. 59. F. Sultan, K. Srinivasan, D. Iyer, and L. Iftode, Migratory TCP: connection migration for service continuity in the internet, IEEE International Conference on Distributed Computing Systems, Vienna, Austria, pp. 469-470, July 2002.
- [59] F. Sultan, K. Srinivasan, D. Iyer, and L. Iftode, Migratory TCP: connection migration for service continuity in the internet, IEEE International Conference on Distributed Computing Systems, Vienna, Austria, pp. 469-470, July 2002.
- [60] H. Hsieh and R. Sivakumar, pTCP: an end-to-end transport layer protocol for striped connections, IEEE International Conference on Network Protocols, Paris, France, pp. 24 33, November 12-15, 2002.
- [61] RFC3261. SIP: Session Initiation Protocol
- [62] H. Schulzrinne and E. Wedlund, Application-layer mobility using SIP. SIGMOBILE Mob. Comput. Commun. Rev., vol. 4, pp. 47-57 2000.
- [63] E. Wedlund and H. Schulzrinne, "Mobility support using SIP," Proceedings of the 2nd ACM international workshop on Wireless mobile multimedia, 1999.
- [64] C. Politis, K. Chew, and R. Tafazolli, "Multilayer Mobility Management for All-IP Networks: Pure SIP vs. Hybrid SIP/Mobile IP," 57th IEEE Semiannual Vehicular Technology Conference VTC2003, Apr. 2003.
- [65] Ashutosh Dutta, Sunil Madhani, Wai Chen. Computer Science Department, Columbia University, New York. Fast-handoff Schemes for Application Layer Mobility Management.
- [66] Q. Wang and M. A. Abu-Rgheff, "Integrated Mobile IP and SIP Approach for Advanced Location Management," Proc. IEE 4th International Conference on 3G Mobile Communication Technologies 3G2003, June 2003.
- [67] G. Wu, Y. Bai, J. Lai and A. Ogielski, "Interaction between TCP and RLP in Wireless Internet", in Proceedings of IEEE GLOBECOM'99, Rio de Janeiro, Brazil, Dec. 1999.

- [68] P. Sudame and B. R. Badrinath, "On Providing Support for Protocol Adaptation in Mobile Wireless Networks", *Mobile Networks and Applications*, Vol. 6, No. 1, Jan. 2001, pp. 43-55.
- [69] J. Postel, "Internet Control Message Protocol", RFC 792, 1981.
- [70] K. Chen, S. H. Shan and K. Nahrstedt, "Cross-Layer Design for Data Accessibility in Mobile Ad Hoc Networks", *Wireless Personal Communications*, Vol. 21, No. 1, Apr. 2002, pp. 49-76.
- [71] B-J "J" Kim, "A Network Service Providing Wireless Channel Information for Adaptive Mobile Applications: Part I: Proposal", in *Proceedings of ICC'01*, Helsinki, Finland, June 2001.
- [72] Mobile People Architecture. *Mobile Computing and Communications Review*, Volume 1, Number 2, 1999.
- [73] Architecture Using Cross-Layer Signalling Interactions," *Proc. IEE 5th European Personal Mobile Communications Conference EPMCC'03*, Apr. 2003.
- [74] A. Aghvami, T. Le, N. Olaziregi, "Mode switching and QoS issues in software radio," *IEEE Personal Communications* 8 (5) (2001) 38-44.
- [75] Wingho Yuen, Heung-No Lee and Timothy Anderson, "Simple but effective cross-layer networking system for mobile ad hoc networks," *Proc. of IEEE PIMRC*, pp. 1952-6, Lisbon, Portugal, Sept. 2002
- [76] M. Park, J. G. Andrews, and S. Nettles, "Wireless Channel-Aware Ad Hoc Cross-Layer Protocol with Multi-Route Path Selection Diversity", *IEEE Vehicular Technology Conference*, Orlando, FL, October 2003.
- [77] A. Otyakmaz, U. Fornefeld, J. Mirkovic, D.C. Schultz, E. Weiss, Performance evaluation for IEEE 802.11g hot spot coverage using sectorized antennas, in: *15th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC '04)*, vol. 2, 2004, pp. 1460-1465.
- [78] J.C. Chen, M.C. Jiang, Y.W. Liu, Wireless LAN security and IEEE 802.11i, *IEEE Wireless Communications* 12 (1) (2005) 27-36.
- [79] Q. Ni, Performance analysis and enhancements for IEEE 802.11e wireless networks, *IEEE Network* 19 (4) (2005) 21-27.
- [80] P. Garg, R. Doshi, R. Greene, M. Baker, M. Malek, X. Cheng, Using IEEE 802.11e MAC for QoS over Wireless, *IEEE International Performance, Computing, and Communications Conference 2003 (IPCCC '06)* 2003, pp. 537-542.
- [81] Y. Xiao, Packing mechanisms for the IEEE 802.11n wireless LANs, *IEEE communications society*, in: *47th Annual IEEE Global Telecommunications Conference 2004 (Globecom'04)*, Dallas, TX, USA, 2004, pp. 3275-3279.
- [82] L. Zhang, P. Senac, E. Lochin and M. Diaz. Mobile TFRC: a congestion control for WLANs. *The IEEE International Symposium on a World of Wireless, Mobile and Multimedia Networks (WOWMOM 2008)*, 23-26 June 2008, Newport Beach CA, USA.
- [83] Wicker, Stephen B. *Error Control Systems for Digital Communication and Storage*. Englewood Cliffs, N.J.: Prentice-Hall, 1995

- [84] G. Fairhurst. RFC 3366 "Advice to link designers on link Automatic Repeat reQuest (ARQ)"
- [85] R. Khalili and K. Salamatian. A new analytic approach to evaluation of packet error rate in wireless networks. In Proceedings of the 3rd Annual Communication Networks and Services Research Conference, Washington, DC, USA, 2005. IEEE Computer Society.
- [86] Ramin Khalili, Kav'e Salamatian. EVALUATION OF PACKET ERROR RATE IN WIRELESS NETWORKS. LIP6-CNRS, Universit'e Pierre et Marie Curie.
- [87] Gowrishankar and Satyanarayana P.S. Recurrent Neural Network Based Bit Error Rate Prediction for 802.11 Wireless Local Area Network. IJCSES International Journal of Computer Sciences and Engineering Systems, Vol.2, No.3, July 2008.
- [88] Dae Gil Yoon, Soo Young Shin and Wook Hyun Kwon. Packet Error Rate Analysis of IEEE 802.11b under IEEE 802.15.4 Interference. IEEE 2006
- [89] IEEE 802.11TM WIRELESS LOCAL AREA NETWORKS STANDARD.
- [90] Jongseok Kim Seongkwan Kim Sunghyun Choi. CARA: Collision-Aware Rate Adaptation for IEEE 802.11 WLANs. Infocom 2006
- [91] Starsky H.Y. Wong, Hao Yang, Songwu Lu and Vaduvur Bharghava. Robust Rate Adaptation for 802.11 Wireless Networks. MobiCom'06, September 23–26, 2006, Los Angeles, California, USA.
- [92] M. Bottigliglio, C. Casetti, C.-F. Chiasserini and M. Meo, Shortterm Fairness of TCP Flows in 802.11b WLANs, in: Proceedings of IEEE INFOCOM'04 (Hong Kong, China, March 2004).
- [93] S. Pilosof, R. Ramjee, D. Raz, Y. Shavitt and P. Sinha, Understanding TCP fairness over Wireless LAN, in: Proceedings of IEEE INFOCOM' 03 (San Francisco, CA, March–April 2003) pp. 836–843.
- [94] Shugong Xu and Tarek Saadawi. "Does the IEEE 802.11 MAC Protocol WorkWell in MultihopWireless Ad Hoc networks?". IEEE Communications Magazine, pages 130–137, June 2001
- [95] Can Emre Koksall, Hisham Kassab, and Hari Balakrishnan. An analysis of short-term fairness in wireless media access protocols (poster). In Measurement and Modeling of Computer Systems, pages 118–119, 2000.
- [96] George Xylomenos and George C. Polyzos. "TCP Performance Issues over Wireless Links". IEEE Communications Magazine, pages 52–58, April 2001.
- [97] E. Exposito, P. Senac, Michel Diaz, Compositional Architecture Pattern for QoS-oriented Communication Mechanisms, Multimedia Modelling 2005, Melbourne, Jan. 2005.
- [98] P. Vixie, Editor, Dynamic Updates in the Domain Name System. RFC 2136
- [99] J. Laganier, L. Eggert. Host Identity Protocol (HIP) Rendezvous Extension. Draft-ietf-hip-rvs-05
- [100] R. Moskowitz, P. Nikander. Host Identity Protocol Architecture. Draft-ietf-hip-arch-03

- [101] D. Eastlake. DNS Request and Transaction Signatures. RFC 2931
- [102] RFC1321, The MD5 Message-Digest Algorithm
- [103] <http://www.simplifiedns.com/>
- [104] “DNSSEC: DNS Security Extensions. Securing the Domain Name System“, www.dnssec.net.
- [105] Openhip, www.openhip.org
- [106] Libo Song, David Kotz, Ravi Jain, and Xiaoning He, “Evaluating location predictors with extensive Wi-Fi mobility data,” in Proc. IEEE INFOCOM 2004, March 2004.
- [107] JDHCP, <http://www.dhcp.org/javadhcp/>
- [108] JMF, <http://java.sun.com/javase/technologies/desktop/media/jmf/>
- [109] T. Henderson, “End-Host Mobility and Multihoming with the Host Identity Protocol,” draft-ietf-hip-mm-04, 2006.
- [110] K. U. F. Teraoka, M. Ishiyama, “LIN6: A Solution to Multihoming and Mobility in IPv6,” draft-teraoka-multi6-lin6-00, 2003.
- [111] M. B. E. Nordmark, “Level 3 multihoming shim protocol,” draft-ietfshim6- proto-05, 2006.
- [112] A. Misra, S. Das, A. Mcauley, A. Dutta, and S. K. Das, “IDMP: an intra-domain mobility management protocol using mobility agents,” Internet Draft, draft-misra-mobileip-idmp-00, Work in Progress, Jan. 2000.
- [113] A. O’Neill, G. Tsirtsis, and S. Corson, “Edge mobility architecture,” Internet Draft, draft-oneill-ema-02, Work in Progress, July 2000.
- [114] TURTLE, <http://labsoc.comelec.enst.fr/turtle/>
- [115] Héctor Velayos, Gunnar Karlsson, Techniques to reduce the IEEE 802.11b handoff time.
- [116] Nikita Borisov, Ian Goldberg, David Wagner. Intercepting Mobile Communications: The Insecurity of 802.11. Published in the proceedings of the Seventh Annual International Conference on Mobile Computing And Networking, July 16–21, 2001.
- [117] CSMA/CA. http://en.wikipedia.org/wiki/Carrier_sense_multiple_access_with_collision_avoidance
- [118] Dajiang He, Charles Q. Shen. Simulation Study of IEEE 802.11e EDCF. Work under European IST Moby Dick project.
- [119] AWGN. http://en.wikipedia.org/wiki/Additive_white_Gaussian_noise
- [120] WPA. http://en.wikipedia.org/wiki/Wi-Fi_Protected_Access
- [121] IEEE 802.21. <http://www.ieee802.org/21/index.html>
- [122] Kaixin Xu, Mario Gerla, Sang Bae. How Effective is the IEEE 802.11 RTS/CTS Handshake in Ad Hoc Networks? Work supported in part by ONR “MINUTEMAN” project.

- [123] Daji Qiao, Sunghyun Choi, Amjad Soomro, Kang G. Shin. Energy-Efficient PCF Operation of IEEE 802.11a Wireless LAN. Infocom 2002
- [124] Intelligent Transportation. http://en.wikipedia.org/wiki/Intelligent_transportation_system
- [125] Anmol Sheth, Richard Han. An Implementation of Transmit Power Control in 802.11b Wireless Networks. Technical Report CU-CS-934-02. Department of Computer Science, University of Colorado.
- [126] A Novel Middleware for the Mobility Management Over the Internet. Lei Zhang, Patrick Sénac, Emmanuel Lochin and Michel Diaz, IEEE International Symposium on a World of Wireless, Mobile and Multimedia Networks (IEEE WoWMoM 2008)
- [127] Fethi Filali, Wimeter - a bandwidth estimation tool and its assistance to QoS Provisioning in Multiple Hot Spots WLANs, NEWCOM Technical Dissemination Day, February 17th 2007, Paris, France.
- [128] Mark Davis, A wireless traffic probe for radio resource management and QoS provisioning in IEEE 802.11 WLANs, Proceedings of the 7th ACM international symposium on Modeling, analysis and simulation of wireless and mobile systems, October 04-06, 2004, Venice, Italy
- [129] A. johnsson, B.Melander and M.Bjorkman. Bandwidth Measurement in Wireless Networks, Mediterranean Ad Hoc Networking Workshop (Med-Hoc-Net), June 2005.
- [130] {key-14}Karthik Lakshminarayanan , Venkata N. Padmanabhan , Jitendra Padhye, Bandwidth estimation in broadband access networks, Proceedings of the 4th ACM SIGCOMM conference on Internet measurement, October 25-27, 2004, Taormina, Sicily, Italy
- [131] J. Sharony. Introduction to Wireless MIMO. Talk of the Communications Society of the IEEE Long Island Section.
- [132] J. M. Wilson. The Next Generation of Wireless LAN Emerges with 802.11n. Intel Corporation white paper, 2004.
- [133] M. Heusse, F. Rousseau, G. Berger-Sabbatel, and A. Duda. Performance anomaly of 802.11b. In Proceedings of IEEE INFOCOM 2003, San Francisco, USA, Mar. 2003
- [134] S. Pilosof, R. Ramjee, D. Raz, Y. Shavitt, and P. Sinha. Understanding TCP fairness over Wireless LAN. 2003 IEEE. In Proc. of IEEE Infocom, 2003.
- [135] L. Zhang, P. Senac, E. Lochin, and M. Diaz. "Cross-Layer based Congestion Control for WLANs." In ICST QShine 2008, Hong Kong, China, July 2008
- [136] R. G. Garroppo, S. Giordano, S. Lucetti, and E. Valori. Twhtb: A transmission time based channel-aware scheduler for 802.11 systems. In Proc. of 1st Workshop on Resource Allocation in Wireless Networks (RAWNET), Riva del Garda, Italy, Apr. 2005.
- [137] R. G. Garroppo, S. Giordano, S. Lucetti, and E. Valori. The wireless hierarchical token bucket: A channel aware scheduler for 802.11 networks. In Proc. of the Sixth IEEE International Symposium on World of Wireless Mobile and Multimedia Networks (WoWMoM), Washington, DC, USA, 2005.

- [138] M. Heusse, F. Rousseau, R. Guillier, and A. Duda. Idle sense: an optimal access method for high throughput and fairness in rate diverse wireless lans. In SIGCOMM '05, pages 121–132, New York, NY, USA, 2005. ACM Press.
- [139] OPNET Modeler. www.opnet.com
- [140] L. Iannone and S. Fdida. Sdt.11b : Un schema division de temps pour eviter l'anomalie de la couche mac 802.11b. In CFIP, Colloque Francophone sur l'Ing'enerie des Protocoles, April.
- [141] TheInstitute of Electrical and Electronics Engineers (IEEE). Ieee 802.11e/d12.0. draft supplement to part 11: Wireless medium access control (mac) and physical layer (phy) specifications: Medium access control (mac) enhancements for quality of service (qos). IEEE 802.11 WG, 2004.
- [142] RFC 3448, TCP Friendly Rate Control (TFRC)
- [143] IEEE 802.11TM WIRELESS LOCAL AREA NETWORKS STANDARD.
- [144] R. Khalili and K. Salamatian. A new analytic approach to evaluation of packet error rate in wireless networks. In Proceedings of the 3rd Annual Communication Networks and Services Research Conference, Washington, DC, USA, 2005. IEEE Computer Society.
- [145] A. Bouabdallah, M. Kieffer, J. Lacan, G. Sabeva, G. Scot, C. Bazile, and P. Duhamel. Evaluation of cross-layer reliability mechanisms for satellite digital multimedia broadcast. Broadcasting, IEEE Transactions on, 53(1), Mar. 2007.
- [146] P. Frossard, C. Chen, C. Sreenan, K. Subbalakshmi, D. Wu, and Q. Zhang. Guest editorial cross-layer optimized wireless multimedia communications. Selected Areas in Communications, IEEE Journal on, 5(4), May 2007.
- [147] G. Fairhurst and L. Hood. Advice to Link Designers on Link Automatic Repeat reQuest (ARQ). In Request For Comments 3366. IETF, Aug. 2002.
- [148] E. Lopez-Aguilera, M. Heusse, Y. Grunengerger, F. Rousseau, A. Duda, and J. Casademont. An asymmetric access point for solving the unfairness problem in wlans. IEEE Transactions on Mobile Computing, Mar. 2008.
- [149] R. G. Garroppo, S. Giordano, S. Lucetti, and E. Valori. Twhtb: A transmission time based channel-aware scheduler for 802.11 systems. In Proc. of 1st Workshop on Resource Allocation in Wireless Networks (RAWNET), Riva del Garda, Italy, Apr. 2005.
- [150] R. G. Garroppo, S. Giordano, S. Lucetti, and E. Valori. The wireless hierarchical token bucket: A channel aware scheduler for 802.11 networks. In Proc. of the Sixth IEEE International Symposium on World of Wireless Mobile and Multimedia Networks (WoWMoM), Washington, DC, USA, 2005.
- [151] 802.16e-2005. IEEE Standard for Local and Metropolitan Area Networks—Part 16: Air Interface for Fixed and Mobile Broadband Wireless Access Systems,” February 28, 2006.
- [152] Mohammad Azizul Hasan. Performance Evaluation of WiMAX/IEEE 802.16 OFDM Physical Layer. Master's Thesis of HELSINKI UNIVERSITY OF TECHNOLOGY

- [153] Yu-Ting Yu and Hsi-Lu Chao. "A Cross-layer Approach of Link Adaptation in IEEE 802.16 Broadband Wireless Access Systems" 2007 IEEE
- [154] Tsz Ho CHAN, Mounir HAMDI. "A Link Adaptation Algorithm in MIMO-based WiMAX systems". JOURNAL OF COMMUNICATIONS, VOL. 2, NO. 5, AUGUST 2007.
- [155] Nai'an Liu, Dechun Sun, Wuqiang Li, and Changxing Pei, "A link adaptation solution to IEEE 802.16d WMAN," Proc. IEEE PDCAT 2006.
- [156] Juan Li and Sven-Gustav Haggman, "Performance of IEEE802.16-2004 based system in jamming environment and its improvement with link adaptation," Proc. IEEE PIMRC 2006
- [157] M. Heusse, F. Rousseau, G. Berger-Sabbatel, and A. Duda. Performance anomaly of 802.11b. In in Proc. of IEEE Infocom, 2003
- [158] Jeffrey G.Andrews, Fundamentals of Wimax, P208, P263
- [159] Chadi Tarhini, Tijani Chahed, Modeling of streaming and elastic flow integration in OFDMA-based IEEE802.16 WiMAX. Computer Communications 30 (2007) 3644–3651, 2007 Elsevier B.V.
- [160] YangXiao, WiMAX/MobileFi, P189. Auerbach Publications
- [161] Air interface for fixed broadband wireless access systems. IEEE Standard 802.16, Jun 2004.
- [162] Katja Schwieger and Gerhard Fettweis, Power and Energy Consumption for Multi-Hop Protocols: A Sensor Network Point of View. nternational Workshop in Wireless Ad-Hoc Networks, 2005
- [163] Luigi Rizzo, Effective Erasure Codes for Reliable Computer Communication Protocols. ACMComputer Communication Review Vol.27. n.2 Apr.97 pp24-36.
- [164] Joachim Rosenthal, Roxana Smarandache. Maximum Distance Separable Convolutional Codes (MDS codes)
- [165] Channel Models for Fixed Wireless Applications, IEEE 802.16 Broadband Wireless Access Working Group. IEEE 802.16.3c-01/29r4

End to end architecture and mechanisms for mobile and wireless communications in the Internet

Wireless networks, because of the potential pervasive and mobile communication services they offer, are becoming the dominant Internet access networks. However, the legacy Internet protocols, still dominant at that time, have not been designed with mobility and wireless in mind. Therefore, numerous maladjustments and “defaults of impedance” can be observed when combining wireless physical and MAC layers with the traditional upper layers. This thesis proposes several solutions for a pacific coexistence between these communication layers that have been defined and designed independently.

Reliable mobility management and Low layer performance enhancements are two main challenging issues in the context of wireless networks. Mobility management (which is mostly based on mobile IP architecture nowadays) aims to continuously assign and control the wireless connections of mobile nodes amongst a space of wireless access networks. Low layer performance enhancements mainly focus on the transmission efficiency such as higher rate, lower loss, interference avoidance.

This thesis addresses these two important issues from an original and innovative approach that, conversely to the traditional contributions, entails a minimum impact on the legacy protocols and internet infrastructure. Following the “end to end” and “cross layer” paradigms, we address and offer efficient and light solutions to fast handover, location management and continuous connection support through a space of wireless networks. Moreover, we show that such an approach makes it possible to enhance transmission efficiency and solve efficiently several syndromes that plague the performances of current wireless networks such as performance anomaly, unfairness issues and maladjustment between MAC layer and upper layers. This thesis tackles these issues by combining analytical models, simulations and real experiments. The resulting mechanisms have been developed and integrated into adaptive mobility management communication architecture that delivers high performing communication services to mobile wireless systems, with a focus on WIFI and WIMAX access networks.

Key words: Wireless networks, Mobility management, MAC layer performance, IEEE 802.11, IEEE 802.16, cross layer.



AUTEUR : ZHANG Lei

TITRE : End to end architecture and mechanisms for mobile and wireless communications in the Internet

DIRECTEUR DE THESE : M. DIAZ et P. SENAC

LIEU ET DATE DE SOUTENANCE: ISAE (Campus ENSICA),Toulouse. Le 5 Octobre 2009

RESUME EN FRANÇAIS :

La gestion performante de la mobilité et l'amélioration des performances des couches basses sont deux enjeux fondamentaux dans le contexte des réseaux sans fil. Cette thèse apporte des solutions originales et innovantes qui visent à répondre à ces deux problématiques empêchant à ce jour d'offrir des possibilités de communication performantes et sans couture aux usagers mobiles accédant à l'Internet via des réseaux d'accès locaux sans fil (WLAN). Ces solutions se distinguent en particulier par l'impact minimum qu'elles ont sur les protocoles standards de l'Internet (niveaux transport et réseau) ou de l'IEEE (niveaux physique et liaison de données).

S'inscrivant dans les paradigmes de "bout en bout" et "cross-layer", notre architecture permet d'offrir des solutions efficaces pour la gestion de la mobilité : gestion de la localisation et des handover en particulier. En outre, nous montrons que notre approche permet également d'améliorer l'efficacité des transmissions ainsi que de résoudre efficacement plusieurs syndromes identifiés au sein de 802.11 tels que les anomalies de performance, l'iniquité entre les flux et l'absence de contrôle de débit entre la couche MAC et les couches supérieures. Cette thèse résout ces problèmes en combinant des modèles analytiques, des simulations et de réelles expérimentations. Ces mécanismes adaptatifs ont été développés et intégrés dans une architecture de communication qui fournit des services de communication à haute performance pour réseaux sans fils tels que WIFI et WIMAX.

MOTS-CLES : La gestion de la mobilité, IEEE802.11, IEEE802.16, couche MAC, cross layer

DISCIPLINE : Informatique, Réseaux

LAAS-CNRS (Laboratoire d'Analyse et d'Architecture des Systèmes) :

7 avenue du Colonel Roche - 31077 Toulouse Cedex 4
Tél. : (33) 05 61 33 62 00 - Fax : (33) 05 61 55 35 77