



Performance evaluation of mobile wireless networks

Ahmad Al Hanbali

► To cite this version:

Ahmad Al Hanbali. Performance evaluation of mobile wireless networks. Networking and Internet Architecture [cs.NI]. Université Nice Sophia Antipolis, 2006. English. NNT: . tel-00327868

HAL Id: tel-00327868

<https://theses.hal.science/tel-00327868>

Submitted on 9 Oct 2008

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.

Université de Nice - Sophia Antipolis – UFR Sciences
École Doctorale STIC

THÈSE

Présentée pour obtenir le titre de :
Docteur en Sciences de l'Université de Nice - Sophia Antipolis

Spécialité : INFORMATIQUE

par

Ahmad AL HANBALI

Équipe d'accueil :
MAESTRO – INRIA Sophia Antipolis

Evaluation des performances des réseaux sans-fil mobiles

Soutenue publiquement à l'INRIA le 20 Novembre 2006 devant le jury composé de :

Président :	Ernst	BIERSACK	Institut Eurécom
Directeurs :	Eitan	ALTMAN	INRIA
	Philippe	NAIN	INRIA
Rapporteurs :	Ger	KOOLE	Université Libre d'Amsterdam
	John C.S.	LUI	Université Chinoise de Hong Kong
Examineurs :	Hossam	AFIFI	Institut National des Télécommunications
	Walid	DABBOUS	INRIA

Evaluation des performances des réseaux sans-fil mobiles

Ahmad AL HANBALI

Titre de la thèse en anglais :

Performance evaluation of mobile wireless networks

À ma famille et mon épouse Hajar

Acknowledgments

This work would not have been possible without the help of many people whom I would like to thank here.

First of all, I must express my deepest gratitude to my advisors Eitan Altman and Philippe Nain. Their help, feedback, and support were absolutely crucial for the successful completion of this thesis. I am considering myself extremely fortunate to being their student.

I gratefully thank Ger Koole and John C.S. Lui for being my thesis reviewers, and for making the trip to attend the defense. My thanks go also to the other members of my thesis committee: Ernst Biersack who presided the committee, Hossam Afifi, and Walid Dabbous.

I am also deeply grateful to Arzad Alam Kherani a past postdoc of Maestro team. Our collaboration was so enriching and thrilling.

This phd. thesis was carried out within the MAESTRO team at INRIA Sophia Antipolis, such a great environment to work. I wish to thank all the permanent researchers of the team. In particular Philippe Nain, Eitan Altman, Kostantin Avrachenkov, Sara Alouf, and Alain Jean-Marie. Sara and Kostantin are soo helpful. I am thankful to many past members of MAESTRO for their support especially during the beginning of my thesis: Robin Groenevelt, Urtzi Ayesta, Florence Clevenot-Perronnin, Balakrishna Prahbu, Naceur Malouch, Rachid Elazouzi, and Daniele Miorandi. Also, I am thankful to all the current members of MAESTRO for the great working atmosphere and the time shared: Nicolas Bonneau, Mohammad Ibrahim, Dinesh Kumar, Natalia Osipova, and Giovanni Neglia. A special thanks to our wonderful assistant Ephie Deriche and for my kind officemate Adulhalim Dandoush.

Love and thanks to my parents, Mohammad and Mouna, who devoted their life to raise me and make me a successful man. I am also grateful to my wife Hajar her love change my life to the best. Finally many thanks to my sister Nisrine and my brothers Abdelrahman, Jihad, and Houssam for their support.

Ahmad Al Hanbali
Novembre 2006, INRIA Sophia Antipolis

Table of Contents

Acknowledgments	iii
1 Introduction	1
1.1 The main challenges in ad hoc networks	2
1.2 Thesis contributions and organization	7
2 A Survey of TCP over Ad Hoc Networks	9
2.1 Introduction	10
2.2 TCP performance over ad hoc networks	11
2.2.1 TCP performance over MANETs	11
2.2.2 TCP performance over SANETs	13
2.2.3 Summary	15
2.3 Proposals to improve TCP performance in ad hoc networks	15
2.3.1 Proposals to distinguish between losses of route failures and congestion	16
2.3.2 Proposals to reduce route failures	21
2.3.3 Proposals to reduce the wireless channel contention	25
2.3.4 Proposals to improve TCP fairness	27
2.4 General comparison and discussion	28
2.5 Conclusion	29
2.A Overview of TCP versions	29
3 Impact of Mobility on the throughput of Relaying in Ad Hoc Networks	33
3.1 Introduction	34
3.2 The system model	35
3.3 Queueing model for the relay buffer	36
3.3.1 Single source, destination, and relay nodes	36
3.3.2 Multiple source, destination, and relay nodes	40
3.4 Comparison of mobility models	43
3.5 Throughput in Random Waypoint and Random Direction models	45
3.5.1 One dimension	46
3.5.2 Two dimensions	48

3.6	Numerical results	50
3.6.1	Validation of Theorem 1	50
3.6.2	Validation of Theorem 2 and Section 3.5	51
3.6.3	Validation of Section 3.3.2	53
3.6.4	Validation of two-hop route probability	56
3.7	Conclusions	57
4	Impact of Mobility on the Relay Buffer Behavior	59
4.1	Introduction	60
4.2	The system model	60
4.3	Relay buffer behavior	62
4.3.1	One-dimensional Random Walk	63
4.3.2	Two-dimensional RD model	64
4.4	Validation and numerical results	67
4.4.1	Validation of Section 4.3.1	67
4.4.2	Validation of Section 4.3.2	67
4.4.3	Comparison between $E[\tilde{B}]$ and $E[B]$	69
4.5	Conclusions	70
5	Performance of Multicopy Two-Hop Relay Routing with Limited Packet Lifetime	71
5.1	Introduction	72
5.2	The system model	72
5.3	Markovian analysis	75
5.3.1	Delivery delay	77
5.3.2	Number of copies in the system at delivery time	77
5.3.3	Total number of copies generated by the source before delivery time	78
5.4	Packet round trip time	80
5.4.1	The n^{th} -order moment of the packet round trip time	82
5.5	Asymptotic analysis	82
5.5.1	Delivery delay	83
5.5.2	Expected number of copies at delivery instant	85
5.5.3	Expected number of copies generated	85
5.6	Validation and numerical results	85
5.6.1	Model validation	86
5.6.2	Comparison of two-hop relay and epidemic routing protocols	86
5.7	Concluding remarks	88
5.A	Proof of Lemma 2	88
5.B	Distribution of delivery delay	91
6	Performance Evaluation of a Class of Two-Hop Relay Protocols	93
6.1	Introduction	94
6.2	The system model	95
6.3	K-copies MTR protocol	95

6.3.1	Delivery delay	98
6.3.2	Number of copies at delivery instant	98
6.3.3	Number of transmissions before absorption	98
6.3.4	Minimizing the consumed energy	99
6.4	K-transmissions MTR protocol	99
6.4.1	Expected delivery delay	103
6.4.2	Distribution of the number of copies generated	103
6.4.3	Numerical evaluation	104
6.4.4	Asymptotic analysis	105
6.5	Two-hop relay with Erasure Coding	106
6.5.1	Computation of $H(s, z)$	108
6.5.2	Erasure coding vs simple multicopy scheme	108
6.5.3	Asymptotic analysis	109
6.6	Concluding remarks	109
7	Impact of Constant TTL and Arbitrary Inter-meeting on MTR	113
7.1	Introduction	114
7.2	Impact of constant TTL	114
7.3	Impact of arbitrary inter-meeting time distribution	117
7.4	Concluding remarks	121
8	Conclusion	123
8.1	Summary of the results	123
8.2	Future work	124
A	IEEE 802.11 standard	125
B	Routing Protocols in Ad Hoc Networks	129
B.1	Proactive routing protocols	129
B.2	Reactive routing protocols	130
C	Thesis Summary in French, “Présentation des Travaux de Thèse”	133
C.1	Introduction	133
C.2	Contributions et organisation	134
C.2.1	Première partie	134
C.2.2	Deuxième partie	135
C.3	Conclusion	142
C.4	Perspectives	143
	List of Acronyms	145
	Bibliography	154

Chapter 1

Introduction

For the last twenty years, mobile communications have experienced an explosive growth. In particular, one area of mobile communication networks, the ad hoc networks, has attracted significant attention due to its challenging research problems.

Ad hoc networks are complex distributed systems that consist of wireless mobile or static nodes that can freely and dynamically self-organize. In this way they form arbitrary and temporary “ad hoc” network topologies, allowing devices to seamlessly interconnect in areas with no pre-existing infrastructure. Depending on whether the wireless nodes are mobile or static, an ad hoc network is called a Mobile Ad hoc Network (MANET) or a Static Ad hoc Network (SANET). Throughout this thesis, the term ad hoc network will represent both MANET and SANET.

The ad hoc networks have spun off many new research areas, such as mesh-based mobile networks and sensor networks. What characterizes all of those wireless networks is that during the exchange of information between the source and destination nodes, there exists a set of links which constitutes an end-to-end route over which the packets are transmitted, called the multi-hop route. However, in some application scenarios, the multi-hop route suffers from frequent failures due to mobility, low node density, or some wireless channel conditions. This causes many challenges to the protocols of ad hoc networks.

On the other hand, the wireless local area network (WLAN) is a single hop wireless network with a small area of coverage. But, WLAN standards may support two modes of operation: an infrastructure based mode and an ad hoc mode. In the infrastructure based mode, the communications between two nodes or between a node and the Internet are done through the Access Point (AP). On the contrary, the ad hoc mode is a peer-to-peer mode where each node may act as a router for the others. The latter mode is used in ad hoc networks to allow multi-hop communications between nodes, and for this reason some of WLAN standards are called the enabling technology of ad hoc networks. Among those

standards we could find the de-facto WLAN standard IEEE 802.11 [IEE]. For details about this standard see Appendix A.

In the following, we will overview the main challenges in ad hoc networks, and then report the contributions of this thesis to some of these challenges related to mobility.

1.1 The main challenges in ad hoc networks

In ad hoc networks, the protocols face many new challenges. This is due to several reasons specific to these networks: lossy channels, hidden and exposed stations, path asymmetry, network partitions, route failures, and power constraints.

Lossy channels

The main causes of errors in wireless channel are the following:

- Signal attenuation: This is due to a decrease in the intensity of the electromagnetic energy at the receiver (e.g., due to the distance), which leads to low signal-to-noise ratio (SNR).
- Doppler shift: This is due to the relative velocities of the transmitter and the receiver. Doppler shift causes frequency shifts in the arriving signal, thereby complicating the successful reception of the signal.
- Multipath fading: Electromagnetic waves reflecting off objects or diffracting around objects can result in the signal traveling over multiple paths from the transmitter to the receiver. Multipath propagation can lead to fluctuations in the amplitude, phase, and geographical angle of the signal received at a receiver.

In order to increase the success of transmissions, link layer protocols implement the following techniques: Automatic Repeat reQuest (ARQ), or Forward Error Correction (FEC), or both. For example, IEEE 802.11 implements ARQ, so when a transmitter detects an error, it will retransmit the frame. Error detection is timer based. Bluetooth implements both ARQ and FEC on some synchronous and asynchronous connections [blu].

Note that packets transmitted over a fading channel may cause the routing protocol to incorrectly conclude that there is a new one-hop neighbor. This one-hop neighbor could provide a shorter route to even more distant nodes. Unfortunately, this new shorter route is usually unreliable. As an example in [CJWK02], Chin et al. deploy DSDV [PW94] and AODV [PBRD03] ad hoc routing protocols in a real network. They find that neither of these protocols can provide a stable multi-hop route because of the physical channel behaviors, and in particular fading. For more details about routing protocols see Appendix B.

Hidden and Exposed stations

In ad hoc networks, stations may rely on a physical carrier-sensing mechanism to determine an idle channel, such as in the IEEE 802.11 DCF function. This sensing mechanism does

not solve completely the *hidden station* and the *exposed station* problems [TK75]. Before explaining these problems, we need to clarify the “transmission range” term. The transmission range is the range, with respect to the transmitting station, within which a transmitted packet can be successfully received.

A typical hidden terminal situation is depicted in Figure 1.1. Stations A and C have a frame to transmit to station B. Station A cannot detect C’s transmission because it is outside the transmission range of C. Station C (resp. A) is therefore “hidden” to station A (resp. C). Since A and C transmission areas are not disjoint, there will be packet collisions at B. These collisions make the transmission from A and C toward B problematic. To alleviate the hidden station problem, virtual carrier sensing has been introduced [IEE, BDSZ94]. It is based on a two-way handshaking that precedes data transmission. Specifically, the source station transmits a short control frame, called Request-To-Send (RTS), to the destination station. Upon receiving the RTS frame, the destination station replies by a Clear-To-Send (CTS) frame, indicating that it is ready to receive the data frame. Both RTS and CTS frames contain the total duration of the data transmission. All stations receiving either RTS or CTS will keep silent during the data transmission period (e.g., station C in Figure 1.1).

However, as pointed out in [FZL⁺03, XMB03], the hidden station problem may persist in IEEE 802.11 ad hoc networks even with the use of the RTS/CTS handshake. This is due to the fact that the power needed for interrupting a packet reception is much lower than that of delivering a packet successfully. In other words, the node’s transmission range is smaller than the sensing node range. For more details, see the model of the physical layer implemented in NS-2 and the Glomosim simulators [ns2, glo].

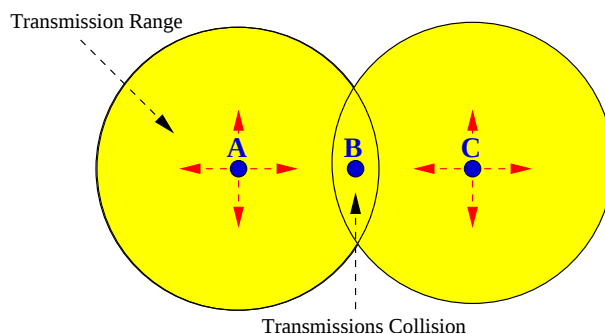


Figure 1.1: Hidden terminal problem: Packets sent to B by A and C will collide at B.

The exposed station problem results from a situation where a transmission has to be delayed because of the transmission between two other stations within the sender’s transmission range. In Figure 1.2, we show a typical scenario where the exposed terminal problem occurs. Let us assume that A and C are within B’s transmission range, and A is outside C’s transmission range. Let us also assume that B is transmitting to A, and C has a frame to be transmitted to D. According to the carrier sense mechanism, C senses a busy channel because of B’s transmission. Therefore, station C will refrain from transmitting to D, although this transmission would not cause interference at A. The exposed station problem may thus result in a reduction of channel utilization.

It is worth noting that hidden terminal and exposed terminal problems are correlated with the transmission range. By increasing the transmission range, the hidden terminal problem occurs less frequently. On the other hand, the exposed terminal problem becomes more important as the transmission range identifies the area affected by a single transmission.

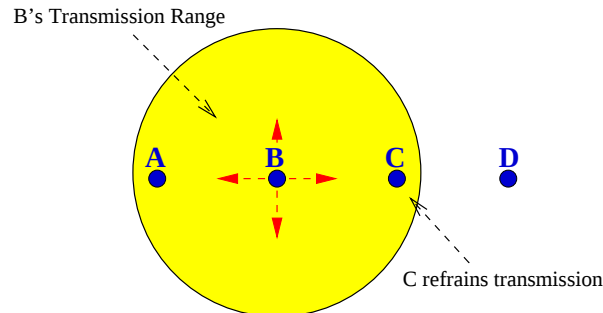


Figure 1.2: Exposed terminal problem: Because of B's transmission C refrains transmission to D.

Path asymmetry

Path asymmetry in ad hoc networks may appear in several forms like bandwidth asymmetry, loss rate asymmetry, and route asymmetry.

Bandwidth asymmetry: Satellite networks suffer from high bandwidth asymmetry, resulting from various engineering tradeoffs (such as power, mass, and volume), as well as the fact that for space scientific missions, most of the data originates at the satellite and flows to the earth. The return link is not used, in general, for data transferring. For example, in broadcast satellite networks the ratio of the bandwidth of the satellite-earth link over the bandwidth of the earth-satellite link is about 1000 [DMT96]. On the other hand in ad hoc networks, the degree of bandwidth asymmetry is not very high. For example, the bandwidth ratio lies between 2 and 54 in ad hoc networks that implement the IEEE 802.11 version g protocol [IEE]. The asymmetry results from the use of different transmission rates. Because of these different transmission rates, even symmetric source destination paths may suffer from bandwidth asymmetry.

Loss rate asymmetry: This type of asymmetry takes place when the backward path is significantly more lossy than the forward path. In ad hoc networks, this asymmetry is due to the fact that packet losses depend on local constraints that can vary from place to place. Note that loss rate asymmetry may produce bandwidth asymmetry. For example, in multi-rate IEEE 802.11 protocol versions, senders may use the Auto-Rate-Fallback (ARF) algorithm for transmission rate selection [KM97]. With ARF, senders attempt to use higher transmission rates after consecutive transmission successes, and revert to lower rates after failures. So, as the loss rate increases the sender will keep using low transmission rates.

Route asymmetry: Unlike the previous two forms of asymmetry, where the forward

path and the backward path can be the same, route asymmetry implies that distinct paths are used for data packets and acknowledgment packets like in TCP protocol [rfc81]. This asymmetry may be due to the routing protocol used. Route asymmetry increases routing overheads and packet losses in case of high degree of mobility¹. Because when nodes move, using distinct forward and reverse routes increases the probability of route failures. However, this is not the case in static networks or networks that have low degree of mobility. So, it is up to the routing protocols to select symmetric paths when such routes are available in the case of ad hoc networks of high mobility.

Network partition

An ad hoc network can be represented by a simple graph G . Nodes are the “vertices”. A successful transmission between two nodes is an undirected “edge”. Network partition happens when G is disconnected. The main reason of this disconnection in MANETs is node mobility. Another factor that can lead to network partition is energy constrained operation of nodes. An example of network partition is illustrated in Figure 1.3. In this figure, A is the source node and F is the destination node, the dashed lines are the links between nodes. When node D moves away from node C it results in a partition of the network into two separate components. Clearly, F cannot receive the data packets of A. If the disconnectivity persists for a long duration the network performance (throughput, delay, loss rate) will degrade; Grossglauser and Tse [GT02] propose to take advantage of the nodes mobility to relay packets between nodes in order to increase the network throughput in the case of network partition. This proposal is based on the store-carry-and-forward paradigm and it works as follows. When there is no route between the source node and the destination node, the source node transmits its packets to one of its neighboring nodes called relay node. Then the relay node will store these packets in its buffer until the time it will encounter the destination, which justifies the name of this proposal “two-hop relay protocol”. Chapters 3--7 evaluate the performance (throughput, delay, energy consumption) of the two-hop relay protocol.

Another implication of network partition occurs when TCP is running between nodes. For example in Figure 1.3, let node A be the TCP sender and node F be the TCP sink. If the disconnectivity persists for a duration greater than the retransmission timeout (RTO) of the TCP sender, this sender will trigger the *exponential backoff* algorithm [PA00]. This algorithm consists of doubling the RTO whenever the timeout expires. Originally, TCP does not have indication about the exact time of network reconnection. This lack of indication may lead to long idle periods during which the network is connected again, but TCP is still in the backoff state. In Chapter 2, we survey the main proposals of the literature to overcome this problem.

¹A network with routes of short lifetime with respect to the session transfer time is an example of a network of high degree of mobility

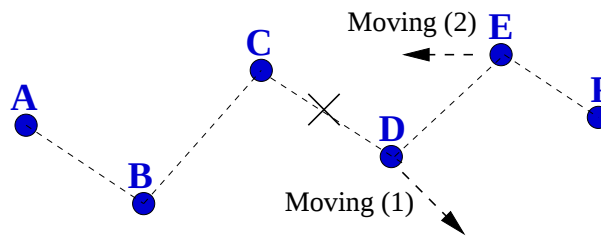


Figure 1.3: Network partition scenario: When D is moving away from C. The network is reconnected when E is moving toward C.

Routing failures

In wired networks route failures occur very rarely. In MANETs they are frequent events. The main cause of route failures is node mobility. Another factor that can lead to route failures is the link failures due to the contention on the wireless channel, which is the main cause of performance degradation in SANETs. The route re-establishment duration after a route failure in ad hoc networks depends on the underlying routing protocol, the mobility pattern of mobile nodes, and the traffic characteristics. As already discussed, if a TCP sender does not have indications on the route re-establishment event, the throughput and session delay will degrade because of the long idle time. We note that the two-hop relay protocol is used also in the scenarios of frequent route failures. See Chapters 3--7.

In addition routing protocols in ad hoc networks that rely on broadcast signaling messages, like the Hello messages, to detect neighbors reachability may suffer from the “communication gray zones” problem. In these zones, data messages cannot be exchanged although broadcast Hello messages and control frames indicate that neighbors are reachable. So on sending data messages, routing protocols will experience routing failures. [LNT02] Lundgren et al. have conducted experiments and subsequently concluded that this problem arises from the heterogeneous transmission rates, the absence of acknowledgment for broadcast packets, the small packet size of Hello messages, and the fluctuation of wireless links.

Power constraints

Because batteries carried by each mobile node have a limited power supply, the processing power is limited. This is a major issue in ad hoc networks as each node is acting as an end system and as a router at the same time, with the implication that additional energy is required to forward and relay packets. Ad hoc network protocols must use this scarce power resource in an “efficient” manner. Here, efficiency refers to minimizing the number of unnecessary retransmissions at both the transport and link layers. In general, in ad hoc networks there are two correlated power problems: the first one is “power saving” that aims at reducing the power consumption; the second one is “power control” that aims at adjusting the transmission power of mobile nodes. Power saving strategies have been investigated at several levels of a mobile device including the physical layer transmissions, the operating

systems, and the applications [JSAC01]. When evaluating the performance of the two-hop relay protocol in Chapters 5 and 6, we will be interested in the power consumption and how we can minimize it subject to a delay constraint.

After reporting the main challenges in ad hoc networks, we will proceed in the next section with the thesis contributions and organization.

1.2 Thesis contributions and organization

This thesis focuses on the impact of mobility on the performance of MANETs protocols, especially TCP and the two-hop relay protocol. The chapters of the thesis can be classified into two parts. The first part deals with TCP, and the second part deals with the two-hop relay protocol.

The first part consists of Chapter 2. This part examines the performance of TCP over ad hoc networks with a special interest on MANETs, and it surveys the main proposals that aim at improving the TCP performance. The main conclusion is that mobility is the main cause of TCP performance degradation in MANETs as it induces frequent route failures and network partitions. Originally, TCP can not differentiate between losses due to route failure and network congestion. To overcome this problem, most of the TCP proposals over MANETs suggest to notify the TCP sender when the routing layer detects a route failure. Upon receiving this notification, TCP sender enters a *freezing* state, in which TCP stops sending data packets and freezes all of its variables. This approach motivates us to look at the performance gain that can be achieved when the mobility is exploited to relay packets between nodes during periods of disconnection. This may be done using the two-hop relay protocol.

The second part of the thesis consists of Chapters 3--7. This part evaluates the performance of the two-hop relay protocol over MANETs by proposing and developing several stochastic models. More precisely these chapters are organized as follows:

- Chapter 3 evaluates the impact of nodes mobility on the relay throughput of the two-hop relay protocol through a queueing analysis based on the G/G/1 queue. The relay throughput is defined as the maximum rate at which a node can relay data from the source to the destination. The main findings are: **(1)** The relay throughput depends on the random mobility of the nodes only via their stationary spatial distributions. **(2)** The random mobility models that result in a uniform stationary spatial distribution of nodes achieve the lowest relay throughput. **(3)** We report the stability condition of the Relay Buffer (RB) of a relay node and some simple approximation of the relay throughput for a number of random mobility models such as the Random Waypoint [BHPC04], the Random Direction [PNL05], and the Random Walk [Fel68] models.
- Chapter 4 studies the behavior of the relay buffer as a function of the nodes mobility using the queueing model of Chapter 3. More precisely, we find an approximation of the mean buffer size in heavy traffic case using the Kingman bound. The result is

that the mean buffer size depends on both the expectation and the variance of the duration of time when two nodes are in contact, called contact time. This analysis is done for the Random Direction and Random Walk models in the one-dimensional and two-dimensional scenarios.

- Chapter 5 proposes a Markovian analysis for the multicopy two-hop relay (MTR) protocol where the source of a packet may generate multiple packet copies. The source disseminates these copies to different relay nodes. This is done under the assumption that copies at relay nodes have a limited time to live (TTL). MTR reduces the delivery delay of the packet at the expense of an increase in the energy consumed to deliver the packet. Both closed-form expressions and asymptotic results (when the number of nodes is large) are reported for the distribution of the delivery delay, the distribution of the number of copies at the delivery time, and the expected total number of copies generated by the source. Finally, we compare MTR with the epidemic routing, a scheme that is closed to MTR and is shown to reduce the delivery delay.
- Chapter 6 evaluates the performance of a class of two-hop relay protocols called K-copies and K-transmissions MTR that aim at limiting the energy consumption by limiting to K either the maximum number of copies in the network or the total number of copies generated. These performance results are used to find the optimal K that minimizes the energy consumption subject to a delay constraint. Finally, we study an approach that reduces the variance of the delivery delay over the two-hop relay protocol called Erasure Coding and we compare it with the simple multicopy scheme like in MTR.
- Chapter 7 studies the impact of considering constant TTL and arbitrary distribution of nodes inter-meeting times on MTR protocol. In this case a non-Markovian analysis is done. Both closed-form expressions and approximations of the delivery delay are derived. Further, we show analytically that constant TTL yields stochastically smaller delivery delay than exponential TTLs, and hyper-exponential inter-meeting times yield stochastically larger delivery delay than exponential inter-meeting times.
- Chapter 8 concludes the thesis and gives some research directions.

Notations

A word on the notation: throughout $\mathbf{1}_A$ will designate the indicator function of any event A ($\mathbf{1}_A = 1$ if A is true and 0 otherwise), rv refers to random variable, CDF refers to cumulative distribution function, $A \stackrel{st}{=} B$ denotes that rvs A and B are equal in distribution, $a := b$ refers that a equal to b by definition, $\text{diag}(a_1, \dots, a_n)$ is a n -by- n diagonal matrix with (i, i) -entry a_i .

Chapter 2

A Survey of TCP over Ad Hoc Networks

The Transmission Control Protocol (TCP) was designed to provide reliable end-to-end delivery of data over unreliable networks. In practice, most TCP deployments have been carefully designed in the context of wired networks. Ignoring the properties of wireless ad hoc Networks can lead to TCP implementations with poor performance. In order to adapt TCP to the ad hoc environment, improvements have been proposed in the literature to help TCP to differentiate between the different types of losses. Indeed, in mobile or static ad hoc networks losses are not always due to network congestion, as it is mostly the case in wired networks. In this Chapter, we present an overview of this issue and a detailed discussion of the major factors involved. In particular, we show how TCP can be affected by mobility and lower layers protocols. In addition, we survey the main proposals which aim at adapting TCP to mobile and static ad hoc environments.

Note: The materials in this chapter have appeared in [HAN04a, HAN04b].

2.1 Introduction

Ad hoc Networks are complex distributed systems that consist of wireless mobile or static nodes that can freely and dynamically self-organize. In this way they form arbitrary, and temporary “Ad hoc” networks topologies, allowing devices to seamlessly interconnect in areas with no pre-existing infrastructure. Recently, the introduction of new protocols such as Bluetooth [blu], IEEE 802.11 [IEE] and Hyperlan [hyp] are making possible the deployment of ad hoc networks for commercial purposes. As a result, considerable research efforts have been put on this new challenging wireless environment. For simplicity, in the following we will use the term MANETs instead of mobile ad hoc networks, and SANETs instead of static ad hoc networks. Also we note that the term ad hoc networks will represent both mobile ad hoc networks (MANETs) and static ad hoc networks (SANETs).

TCP (Transmission Control Protocol) [rfc81] was designed to provide reliable end-to-end delivery of data over unreliable networks. In theory, TCP should be independent of the technology of the underlying infrastructure. In particular, TCP should not care whether the Internet Protocol (IP) is running over wired or wireless connections. In practice, it does matter because most TCP deployments have been carefully designed based on assumptions that are specific to wired networks. Ignoring the properties of wireless transmission can lead to TCP implementations with poor performance.

In ad hoc networks, the principal problem of TCP lies in performing congestion control in case of losses that are not induced by network congestion. Since bit error rates are very low in wired networks, nearly all TCP versions nowadays assume that packet losses are due to congestion. Consequently, when a packet is detected to be lost, either by timeout or by multiple duplicated ACKs, TCP slows down the sending rate by adjusting its congestion window. Unfortunately, wireless networks suffer from several types of losses that are not related to congestion, making TCP not adapted to this environment. Numerous enhancements and optimizations have been proposed over the last few years to improve TCP performance over one-hop wireless (not necessarily ad hoc) networks. These improvements include infrastructure based WLANs [AHM03, BPSK97, BSAK95, BB97], mobile cellular networking environments [BS97, BSK95], and satellite networks [DMT96, HK99]. Ad hoc networks inherit several features of these networks, in particular high bit error rates and path asymmetry, and add new problems that come from mobility and multi-hop communications, such as network partitions, route failures, and hidden (or exposed) terminals. We note that the following TCP versions: Tahoe, Reno, Newreno, and Vegas perform differently in ad hoc networks [XS02]. However, all these versions suffer from the same problem of inability to distinguish between packet losses due to congestion from losses due to the specific features of ad hoc networks. For more details about TCP versions see Appendix 2.A, and for a survey about TCP versions we refer to [BAD00].

By examining the TCP’s performance studies over MANETs and SANETs, we identified the following four major problems: **(i)** TCP is unable to distinguish between losses due to route failures and network congestion. **(ii)** TCP suffers from frequent route failures. **(iii)** The contention on the wireless channel. **(iv)** TCP unfairness. We note that the first two problems are the main causes of TCP performance degradation in MANETs, and the other

two problems are the main causes of TCP performance degradation in SANETs. Based on these four problems, the proposals that aim to improve TCP performance over ad hoc networks are regrouped in four sets. We classify the proposals of a set in two categories: cross layer proposals and layered proposals. In cross layer proposals, TCP and its underlying protocols work jointly. For example, considerable improvements are possible when TCP can differentiate between packet losses due to congestion that should activate the congestion control algorithm, and losses due to the specific features of MANETs. In order to do that, some proposals suggest that when the routing layer detects a route failure, it notifies the TCP sender about a routing failure [CRVP98a, HV02, LS01, WZ02, KTC01]. Upon receiving this notification, TCP sender enters a *freezing* state. In this state, TCP stops sending data packets, and it freezes all its variables to their current value, such as the congestion window and the retransmission timer. After route re-establishment, TCP sender goes back to the normal state. In layered proposals, the problems of TCP are addressed at one of the OSI layers. For example in [AJ03], Altman and Jimenez use adaptive TCP delayed ACK to reduce the contention on wireless channel in SANETs. In [FZL⁺03], Fu et al. propose two link layer techniques, called Link RED and adaptive pacing, to improve TCP performance over SANETs.

The rest of the chapter is organized as follows: In Section 2.2 we survey papers on TCP performance over mobile and static ad hoc networks. In Section 2.3 we review the main proposals for improving TCP performance. Section 2.4 compares and discusses these proposals. Section 2.5 concludes the paper and gives research directions.

2.2 TCP performance over ad hoc networks

In this section, first we report the studies of TCP performance over MANETs. Second, we report the studies of TCP performance over SANETs. At the end, we summarize the major problems of TCP over static and mobile ad hoc networks. In passing, we point out that routing protocols and link layer protocols in ad hoc networks are beyond the scope of this survey. The interested readers are referred to Appendix A and B for details.

2.2.1 TCP performance over MANETs

In [HV02] Monks et al. investigate the impact of mobility on TCP throughput in MANETs. In their simulation scenarios, nodes move according to the Random Waypoint model with 0s pause time [BHPC04]. The speed of the node was uniformly distributed in $[0.9v - 1.1v]$ for some mean speed v . At the routing layer, the authors use DSR [JMH03]. They report that when the mean speed increases from 2m/s to 10m/s the throughput drops sharply. But, when it increases from 10m/s to 30m/s the throughput drops slightly. Also they remark that, for a given mean speed, certain mobility patterns achieve throughput close to 0, although the other mobility patterns are able to achieve high throughput. By analyzing the simulations trace of patterns of low throughput, they found that the TCP sender's routing protocol is unable to quickly recognize and purge stale routes from its cache, which

results in repeated routing failures and TCP retransmission timeouts. For patterns of high throughput they found that most of the time the TCP sender and receiver are close to each other. By examining the mobility patterns, the authors observe that as the sender and receiver move closer to each other, DSR can maintain a valid route. This is done by shortening the existing route before a routing failure occurs. However, as the sender and receiver move away from each other, DSR waits until a failure occurs to lengthen a route. The route failure induces up to a TCP-window worth of packet losses [APSS04], and subsequent route discovery latency often results in repeated TCP timeout. To prevent TCP invocation of congestion control that deteriorates TCP throughput in case of losses induced by mobility, the authors suggest to use the explicit link failure notification (ELFN) technique. An overview of this technique is discussed in Section 2.3.1.1.

In [APSS04], in addition to the problem highlighted in [HV02] that is “TCP treats losses induced by route failures as signs of network congestion”, Anantharaman et al. identify a set of factors that contribute also to the degradation of TCP throughput in the presence of mobility. These factors are: MAC failure detection and route computation latencies. The MAC failure detection latency is defined as the amount of time spent before the MAC concludes a link failure. They found that in the case of the IEEE 802.11 protocol, when the load is light (one TCP connection) this latency is small and independent of the speed of the nodes. However in case of high load, the value of this latency is magnified and becomes a function of the nodes speed. The route computation latency is defined as the time taken to recompute the route after a link failure. They found that, as for MAC failure detection latency, the route computation latency increases with the load and becomes a function of the nodes speed in the high load case. Also, the authors identify another problem, called MAC packet arrival, that is related to routing protocols. In fact, when a link failure is detected, the link failure is sent to the routing agent of the packet that triggered the detection. If other sources are using the same link in the path to their destinations, the node that detects the link failure has to wait until it receives a packet from these sources before they are informed of the link failure. This also contributes to the delay after which a source realizes that a path is broken.

In [DB01], Dyer and Boppana report simulation results on the performance of TCP Reno over three different routing protocols (AODV [PBRD03], DSR [JMH03], ADV [BK01]). It is found that ADV performs well under a variety of mobility patterns and topologies. Furthermore, they propose an heuristic technique called fixed RTO to improve the performance of on-demand routing protocols (AODV and DSR). An overview of this technique is discussed in Section 2.3.1.2. Also in [LXG03], Lim et al. show that TCP’s performance degrades when the multipath routing protocol SMR [LG01] is used. Multipath routing effects TCP by two factors: The first one is inaccuracy of the average RTT measurement that leads to more premature timeouts. The second one is out-of-order packet delivery via different paths, which triggers duplicated ACK, which in turn triggers TCP congestion control.

2.2.2 TCP performance over SANETs

In [AJ03, FZL⁺03, XS01, GTB99] the authors report simulation results on TCP throughput in a static linear multi-hop chain, where IEEE 802.11 protocol is used. In Figure 2.1, we display a multi-hop chain of N nodes. It is expected that, as the number of hops increases, the spatial reuse will also increase. However, simulation results indicate that TCP throughput decreases “rapidly” up to a point as the number of hops increases. It is argued that this decrease is due to the hidden terminals problem, which increases frames collisions. After a repeated¹ transmission failure the MAC layer will react by two actions. First, the MAC will drop the head-of-line frame destined to the next hop, we note that this type of drops is known also as drops due to contention on wireless channel. Second, the MAC will notify the upper layer about a link failure. When the routing protocol of a source node detects a routing failure, it will initiate a route re-establishment process. In general the route re-establishment duration is greater than the retransmission timer of the TCP agent; hence the TCP agent will enter the backoff procedure and will set its congestion window to 1. Also, as TCP sender’s does not have indications on the route re-establishment event, TCP will suffer from a long idle time. During this time, the network may be connected again, but TCP is still in the backoff state.

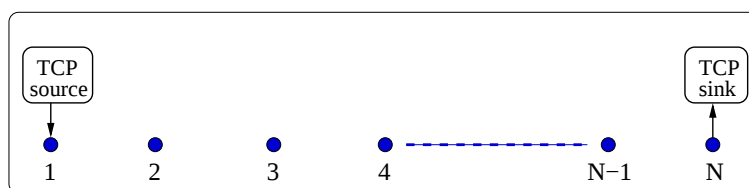


Figure 2.1: Multi-hop chain topology.

In [XS02], Xu and Saadawi study the performance of TCP Tahoe, Reno, New Reno, Sack and Vegas² over the multi-hop chain topology shown in Figure 2.1, in the case where the IEEE 802.11 protocol is used. It is shown that TCP Vegas delivers the better performance and does not suffer from instability. By tuning the sender TCP’s maximum window size (advertised window) to approximately four packets, all TCP versions perform similarly. Furthermore, the authors investigate the performance of these TCP versions using the delayed-ACK option, as defined in RFC 1122. According to the mentioned RFC, the TCP sink will send one TCP ACK packet for every two TCP packets received. This option will reduce the contention on the wireless channel, because the data and Ack packets share the same wireless channel. Simulating the multihop chain with the delayed-ACK option, they report that an improvement of 15% to 32%.

In [XS01, TG99], the authors study how the TCP connections share the bandwidth of the channel in the context of Ad hoc network. They report some unfair bandwidth sharing using the actual MAC 802.11 in a multi-hop communication environment. Moreover in

¹On repeated transmission failure, 802.11 MAC is allowed to retransmit short frames seven times and long frames four times.

²For details about TCP versions see Appendix 2.A

[XBLG02], Xu and al. show that also in scenarios where TCP crosses wireless ad hoc and wired networks, the TCP unfairness problem persists. Xu et al. in [XGQS03] report that RED [FJ93] did not solve TCP's unfairness in ad hoc networks. The reason is that congestion does not happen in a single node, but in an entire area involving multiple nodes. So, the local packet queue at any single node cannot completely reflect the network congestion state. For this reason, they define a new distributed queue that contains all packets whose transmissions will affect the node transmission in addition to its own packets. An overview of this proposal is given in Section 2.3.4.1. In [ACG03], Anastasi et al. investigate the performance of IEEE 802.11 ad hoc network by means of an experimental study. Their findings match those obtained by simulations, namely, TCP connections may experience significant unfairness. They mention several aspects that are usually neglected in simulation studies of the IEEE 802.11b protocol. For instance, since the control frames (RTS, CTS, ACK) and data frame may be transmitted at different rates, this produces different transmission ranges and carrier sensing ranges in the network.

In [CXN03], Chen et al. study the impact of the TCP's congestion window limit (CWL) on TCP throughput in SANETs. In fact, they relate the problem of setting TCP's optimal CWL to identifying the bandwidth-delay product (BDP) of multi-hop paths. As they argue that TCP's congestion window should be less than the BDP. They prove that regardless of the MAC layer being used, the value of the BDP in bytes at multi-hop routes cannot exceed the value of the Round-Trip Hop-Count (RTHC) times the size of data packets in bytes of these multi-hop routes. This is done by assuming similar bottleneck bandwidths along the forward and reverse route. In the case of IEEE 802.11 MAC layer protocol, when a node is transmitting, all other nodes that are inside its interference range or that receive the RTS or CTS frames will remain silent till the transmission ends. The limit on CWL is illustrated in the scenario of Figure 2.1; the distance between adjacent nodes equal to the node transmission range, which is taken to be $250m$, and node interference range equals to $500m$. For that scenario, the authors report that the BDP is less one $1/4$ of the RTHC, as only 4-hops nodes away can transmission of data packets be done concurrently without collisions. Using simulation of the previous scenario, they found that when the CWL exceed one $1/5$ of the RTHC, TCP throughput decreases substantially. So based on this tighter bound, the authors propose an algorithm to adjust TCP CWL according to the RTHC of the routes used. Using their algorithm they report an improvement in TCP performance up to 16% even in mobile scenarios that contain multiple TCP connections. In [FZL⁺03], Fu et al report that given a specific network topology and flow patterns, there exists an optimal TCP's window size W^* . Using W^* TCP achieves the best throughput via improved spatial reuse. But, unfortunately, in practice, TCP operates at an average window size that is much larger than W^* . This leads to increased packet loss due to the contention on the wireless channel. To help TCP operating around W^* , they propose two techniques: link RED and adaptive pacing at the link layer. By coupling these two techniques they show a 30% improvement in the performance. An overview of these techniques is given in Section 2.3.3.3.

2.2.3 Summary

We conclude that the cause of TCP performance degradation in MANETs are due to two major problems. The first problem is, TCP is unable to distinguish between losses due to route failures and network congestion. The second problem is, TCP suffers from frequent route failures. However in SANETs, TCP faces two major problems. The first problem is the contention on the wireless channel. The second problem is TCP unfairness. Basing on these four problems, the TCP proposals of the next section will be regrouped into four sets.

2.3 Proposals to improve TCP performance in ad hoc networks

In this section we present the various proposals which have been made in the literature to improve the performance of TCP in ad hoc networks. We regroup these proposals to four sets according to the four problems identified in Section 2.2.3. These four problems are: **(1)** TCP is unable to distinguish between losses due to route failures and those due to network congestion, **(2)** frequent route failures, **(3)** contention on wireless channel, and **(4)** TCP unfairness. We note that problems (1) and (2) are the main causes of TCP performance degradation in MANETs. However, problems (3) and (4) are the main causes of TCP performance degradation in SANETs. Figure 2.2 shows a general classification of each set of proposals. We classify the proposals that belong to the same set to two types: cross layer proposals, and layered proposals. The cross layer proposals rely on interactions between two layers of the Open System Interconnection (OSI) architecture. These proposals were motivated by the fact that “providing lower layer informations to upper layer should help the upper layer to perform better”. Thus, depending on between which two OSI layers there will be informations exchange, cross layer proposals can be further classified to four types: TCP and network, TCP and link, TCP and physical, and network and physical. Layered proposals rely on adapting OSI layers independently of other layers. Thus, depending on which layer is involved, layered proposals can be further classified to three types: TCP layer, network layer, and link layer proposals.

In general, cross layer solutions report better performance than layered solution. But layered solutions respect the concept of designing protocols in isolation, thus they are considered as long term solutions. So, to choose between cross layer and layered solutions we have to answer first what is prior for us, performance optimization or architecture. Performance optimization can lead to short term gain, while architecture is usually based on longer term considerations. In addition, cross layer solutions are more complex to implement and to design than layered solutions. The reason is that the implementation of cross layer solutions require at least modifications of two OSI layers, and their design requires that the system is considered in its entirety. For more discussions about cross layer design in ad hoc networks we refer the reader to [KK05].

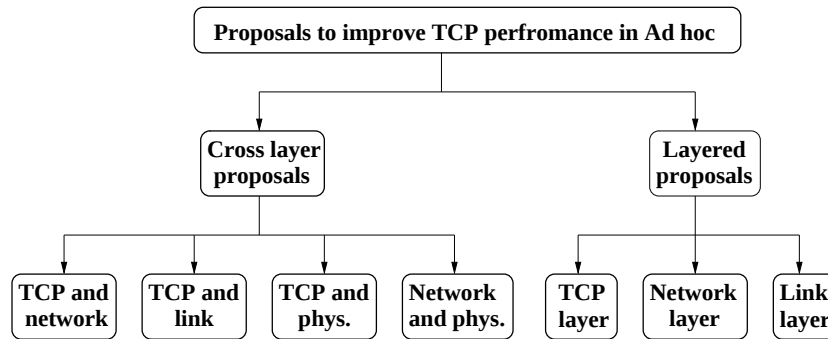


Figure 2.2: Classification of proposals to improve TCP performance in ad hoc networks.

2.3.1 Proposals to distinguish between losses of route failures and congestion

The proposals that address the problem of TCP inability to distinguish between losses due to route failures and network congestion in MANETs can fall into two categories: TCP and network cross layer proposals, and TCP layer proposals.

2.3.1.1 TCP and network cross layer proposals

TCP-F: TCP Feedback [CRVP98b] is a feedback based approach to handle route failures in MANETs. This approach allows the TCP sender to distinguish between losses due to routes failure and those due to network congestion. This is done as follows. When the routing agent of a node detects the disruption of a route, it explicitly sends a Route Failure Notification (RFN) packet to the source. On receiving the RFN, the source goes into a snooze state. A TCP sender in snooze state will stop sending packets, and will freeze all its variables, such as timers and congestion window size. The TCP sender remains in this snooze state until it is notified of the restoration of the route through Route Re-establishment Notification (RRN) packet. On receiving the RRN, the TCP sender will leave the snooze state and will resume transmission based on the previous sender window and timeout values. To avoid blocking in the snooze state, the TCP sender, on receiving RFN, triggers a route failure timer. When this timer expires the congestion control algorithm is invoked normally.

The authors report an improvement by using TCP-F over TCP. The simulation scenario is basic and is not based on an ad hoc network. Instead, they emulate the behavior of an ad hoc network from the viewpoint of a transport layer.

ELFN-based technique: Explicit Link Failure Notification technique [HV02] is similar to TCP-F. However in contrast to TCP-F, the evaluation of the proposal is based on a real interaction between TCP and the routing protocol. This interaction aims to inform the TCP agent about route failures when they occur. The authors use an ELFN message, which is piggy-backed on the route failure message sent by the routing protocol to the

sender. The ELFN message is like a “host unreachable” Internet Control Message Protocol (ICMP) message, which contains the sender receiver addresses and ports, as well as TCP packet’s sequence number. On receiving the ELFN message, the source responds by disabling its retransmission timers and enters a “standby” mode. During the standby period, the TCP sender probes the network to check if the route is restored. If the acknowledgment of the probe packet is received, the TCP sender leaves the standby mode, resumes its retransmission timers, and continues the normal operations.

In the mentioned reference, the authors study the effect of varying the time interval between probe packets. Also, they evaluate the impact of the RTO and the Congestion Window (CW) upon restoration of the route. They find that a probe interval of 2 sec. performs the best, and they suggest to make this interval a function of the RTT instead of giving it a fixed value. For the RTO and CW values upon route restoration, they find that using the prior values before route failure performs better than initializing CW to 1 packet and/or RTO to 6 sec., the latter value being the initial default value of RTO in TCP Reno and New Reno versions.

This technique provides significant enhancements over standard TCP, but further evaluations are still needed. For instance, different routing protocols should be considered other than the reactive protocol DSR considered in [HV02], especially proactive protocols such as OLSR [CJ03]; In [APSS04], Anantharaman et al. report that in case of high load ³ ELFN performs worse than standard TCP, because ELFN is based on probing the network periodically to detect route re-establishment. Also in [MSB00], Monks et al. find that even in case of light load ELFN performs worse than standard TCP by 5% in case of static Ad hoc networks.

ATCP: Ad hoc TCP [LS01] utilizes network layer feedback too. In addition to the route failures, ATCP tries to deal with the problem of high Bit Error Rate (BER). The TCP sender can be put into persist state, congestion control state or retransmit state. A layer called ATCP is inserted between the TCP and IP layers of the TCP source nodes. ATCP listens to the network state information provided by ECN (Explicit Congestion Notification) messages [RFB01] and by ICMP “Destination Unreachable” message; then ATCP puts TCP agent into the appropriate state. On receiving a “Destination Unreachable” message, TCP agent enters a persist state. The TCP agent during this state is frozen and no packets are sent until a new route is found by probing the network. The ECN is used as a mechanism to explicitly notify the sender about network congestion along the route being used. Upon reception of ECN, TCP congestion control is invoked normally without waiting for a timeout event. To detect packet losses due to channel errors, ATCP monitors the received ACKs. When ATCP sees that three duplicate ACKs have been received, it does not forward the third duplicate ACK but puts TCP in the persist state and quickly retransmits the lost packet from TCP’s buffer. After receiving the next ACK, ATCP will resume TCP to the normal state. Note that ATCP allows interoperability with TCP sources or destinations that do not implement ATCP.

ATCP was implemented in a testbed and evaluated under different scenarios, such as

³25 TCP connections

congestion, lossy links, partition, and packet reordering. In all cases the transfer time of a given file using ATCP yielded better performance than TCP. However, the used scenario was somewhat special, since neither wireless links nor ad hoc routing protocols were considered. In fact, the authors used an experimental testbed consisting of five PCs equipped with Ethernet cards. With these PCs, the authors formed a four hop network.

In addition to routes failure, ATCP tries to deal with the problem of high BER, network congestion, and packet reorder. This advantages make ATCP a more robust proposal for TCP in MANETS. But some assumptions such as ECN-capable node as well as sender node being always reachable might be somehow hard to meet in a mobile ad hoc context. Also, the probing mechanism used to detect routes re-establishment generates problems in the case of high load, like in the ELFN proposal.

TCP-BuS: TCP Buffering capability and Sequence information [KTC01], like previous proposals, uses the network feedback in order to detect route failure events and to take convenient reaction to this event. The novel scheme in this proposal is the introduction of *buffering capability* in mobile nodes. The authors select the source-initiated on-demand ABR [Toh97] (Associativity-Based Routing) routing protocol. The following enhancements are proposed:

- Explicit notification: two control messages are used to notify the source about the route failure and the route re-establishment. These messages are called Explicit Route Disconnection Notification (ERDN) and Explicit Route Successful Notification (ERSN). On receiving the ERDN from the node that detected the route failure, called the Pivoting Node (PN), the source stops sending. And similarly after route re-establishment by the PN using a Localized Query (LQ), the PN will transmit the ERSN to the source. On receiving the ERSN, the source resumes data transmission.
- Extending timeout values: during the Route ReConstruction (RRC) phase, packets along the path from the source to the PN are buffered. To avoid timeout events during the RRC phase, the retransmission timer value for buffered packets is doubled.
- Selective retransmission request: as the retransmission timer value is doubled, the lost packet along the path from the source to the PN are not retransmitted until the adjusted retransmission timer expires. To overcome this, an indication is made to the source so that it can retransmit these lost packet selectively.
- Avoiding unnecessary requests for fast retransmission: when the route is restored, the destination notifies the source about the lost packets along the path from the PN to the destination. On receiving this notification, the source simply retransmits these lost packets. But the packets buffered along the path from the source to the PN may arrive at the destination earlier than the retransmitted packets. So the destination will reply by duplicate ACK. These unnecessary request packets for fast retransmission are avoided.
- Reliable retransmission of control message: in order to guarantee the correctness of TCP-BuS operation, they propose to transmit reliably the routing control messages

ERDN and ERSN. The reliable transmission is done by overhearing the channel after transmitting the control messages. If a node has sent a control message but did not overhear this message relayed during a timeout, it will conclude that the control message is lost and it will retransmit this message.

This proposal introduces many new techniques for the improvement of TCP. The novel contributions of this paper are the buffering techniques and the reliable transmission of control messages. In their evaluation, the authors found that TCP-BuS outperforms the standard TCP and the TCP-F under different conditions. The evaluation is based only on the ABR routing protocol and other routing protocols should be taken into account. Also, TCP-BuS did not take into account that pivoting node may fail to establish a new partial route to the destination. In this case, what will happen to the packets buffered at intermediate nodes is not handled by the authors.

2.3.1.2 TCP layer proposals

Fixed RTO: This technique [DB01] is a sender-based technique that does not rely on feedback from the network. In fact, the authors employ a heuristic to distinguish between route failures and congestion. When two timeouts expire in sequence, which corresponds to the situation where the missing ACK is not received before the second RTO expires, the sender concludes that a route failure event has occurred. The unacknowledged packet is retransmitted but the RTO is not doubled a second time. This is in contrast with the standard TCP, where an “exponential” backoff algorithm is used. The RTO remains fixed until the route is re-established and the retransmitted packet is acknowledged.

In [DB01] Dyer et al. evaluate this proposal by considering different routing protocols as well as the TCP selective and the delayed acknowledgment options. They report that significant enhancements are achieved when using fixed-RTO with on-demand routing protocols. Nevertheless, as stated by the authors themselves, this proposal is restricted to wireless networks only, a serious limitation since interoperation with wired networks is clearly necessary. Also, the supposition that two consecutive timeouts are the exclusive results of route failures need more analysis, especially in cases of congestion.

TCP DOOR: TCP Detection of Out-of-Order and Response (DOOR) is an end-to-end approach [WZ02]. This approach, which does not require the cooperation of intermediate nodes, is based on out-of-order (OOO) delivery events. OOO events are interpreted as an indication of route failure. The detection of OOO events is accomplished either by the means of a sender-based or a receiver-based mechanism. The sender-based mechanism uses the non-decreasing property of the ACKs sequence number to detect the OOO events. In case of duplicate ACK packets, these ACKs will have the same sequence number, so that the sender needs additional information to detect OOO event. This information is a one byte option added to ACKs called ACK Duplication Sequence Number (ADSN). The ADSN is incremented and transmitted with each duplicate ACK. However, the receiver-based mechanism needs an additional two bytes TCP option to detect OOO events, called TCP Packet Sequence Number (TPSN). The TPSN is incremented and transmitted with

each TCP packets including the retransmitted packets. If the receiver detects an OOO event, it should notify the sender by setting a specific option bit, called OOO bit, in the ACK packet header.

Once the TCP sender knows about an OOO event, it takes the following two response actions: temporarily disabling congestion control, and instant recovery during congestion avoidance. In the former action, the TCP sender disables the congestion algorithm for a specific time period (T_1). In the latter action, if the congestion control algorithm was invoked during the past time period (T_2), the TCP sender should recover immediately to the state before the invocation of the congestion control. In fact, the authors make the time periods T_1 and T_2 function of the RTT.

In the simulation study presented in [WZ02], different scenarios are considered by combining all mechanisms and actions mentioned above. Their results show that the sender and receiver-based mechanisms behave similarly. So, they recommend the use of the sender detection mechanism as it does not require notifications from the sender to the receiver. Regarding the actions, temporarily disabling congestion control and instant recovery during congestion avoidance, to be taken upon an OOO event detection, they have found that both of them lead to significant improvement. In general, TCP DOOR improves TCP performance up to 50%. Nevertheless, the supposition that OOO events are the exclusive results of route failure deserves much more analysis. Actually, multipath routing protocols such as TORA [Per01] may produce OOO events that are not related to route failures.

2.3.1.3 Comparison

Six proposals have been presented. These proposals address the problem of TCP's inability to distinguish between losses due to route failures and network congestion. The authors of these proposals proceed with a TCP and network cross layer solution, like TCP-F, ELFN, ATCP, and TCP-BuS, or a TCP layered solution, like Fixed RTO, and TCP-DOOR. The four TCP and network cross layer proposals, TCP-F, ELFN, ATCP, and TCP-BuS, are based on an explicit notification from the network layer to detect the route failures. But, they differ in how to detect route re-establishments. TCP-F and TCP-BuS use the explicit notification from the network layer. However, ELFN and ATCP use a probing mechanism. Comparing with the explicit notification, the probing mechanism is easy to implement. But what is the optimal value of the probing interval? And what is the implication of this mechanism in case of high load? Especially we saw that in the case of high load, the ELFN proposal performs worst than the standard TCP. Between the four cross layer proposals, ATCP emerges as a robust proposal as it supports mechanisms to alleviate the negative impact of high BER and packets out-of-order on TCP performance. But This proposal need to revise some assumptions such as ECN-capable node. Table 2.1 compares the key features of the four TCP and network cross layer proposals. The TCP layer proposals, Fixed RTO and TCP DOOR, distinguish between packet losses induced by routes failures and congestion by employing end-to-end TCP layer approach. In Fixed RTO, this is done by considering two consecutive timeouts as a sign of route failures. In TCP-DOOR, when the source or the sink receives out-of-order packets, it considers that a route failure has

occurred. The main advantage of these two proposals is that they do not require explicit notification from routing layer, neither do they require other nodes cooperation to detect route failures. Comparing these two proposals, TCP DOOR performs better than fixed RTO, but at the cost of more modifications.

	TCP-F	ELFN	ATCP	TCP-BuS
High BER packet loss	Not handled	Not handled	Handled	Not handled
Route failures (RF) detection	RFN pkt freezes TCP	ELFN pkt freezes TCP	ICMP “destination unreachable” freezes TCP sender	ERDN pkt freeze TCP
Route Reconst. (RR)	RRN pkt resumes TCP	Probing mechanism	Probing mechanism	ERSN pkt resumes TCP
Packet reordering	Not handled	Not handled	Handled	Not handled
CW and RTO after RR	Old CW and RTO	Old CW and RTO	Reset for each route	Old CW and new RTO
Reliable trx of control messages	Not handled	Not handled	Not handled	Handled
Evaluation	Emulation no routing	Simulation	Experimental no routing	Simulation

Table 2.1: Comparison of TCP and network cross layer proposals to distinguish between losses due to route failures and congestion.

2.3.2 Proposals to reduce route failures

The proposals that address the problem of frequent route failures in MANETs can be classified as follows to three categories: TCP and network cross layer proposals, network and physical cross layer proposals, and network layer proposals.

2.3.2.1 TCP and network cross layer proposal

Split TCP: TCP connections that have a large number of hops suffer from frequent route failures due to mobility. To improve the throughput of these connections and to resolve the unfairness problem, the Split TCP scheme was introduced to split long TCP connections into shorter localized segments [KKFT02] – see Figure 2.3. The interfacing node between two localized segments is called a proxy. The routing agent decides if its node has the role of proxy according to the inter-proxy distance parameter. The proxy intercepts TCP packets, buffers them, acknowledges their receipt to the source (or previous proxy) by sending a local acknowledgment (LACK). Also, a proxy is responsible for delivering the packets, at an appropriate rate, to the next local segment. Upon the receipt of a LACK (from the next proxy

or from the final destination), a proxy will purge the packet from its buffer. To ensure the source to destination reliability, an ACK is sent by the destination to the source similarly to the standard TCP. In fact, this scheme splits also the transport layer functionalities into those end-to-end reliability and congestion control. This is done by using two transmission windows at the source which are the congestion window and the end-to-end window. The congestion window is a sub-window of the end-to-end window. While the congestion window changes in accordance with the rate of arrival of LACKs from the next proxy, the end-to-end window will change in accordance with the rate of arrival of the end-to-end ACKs from the destination. At each proxy, there would be a congestion window that would govern the rate of sending between proxies.

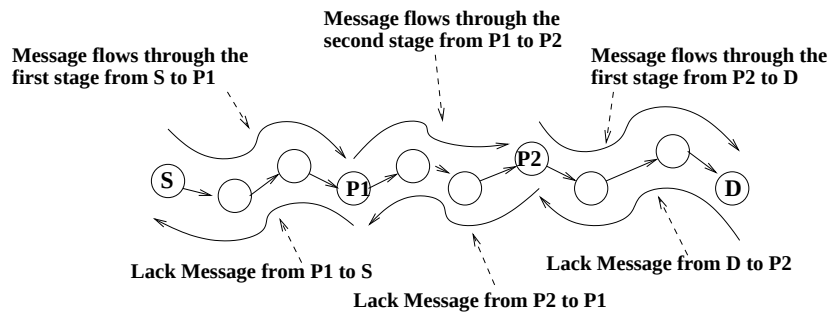


Figure 2.3: TCP Split.

Simulation results indicate that an inter-proxy distance of between 3 and 5 has a good impact on both throughput and fairness. The authors report that an improvement up to 30% can be achieved in the total throughput by using Split TCP. The drawbacks are large buffers and network overhead. Also, this proposal makes the role of proxy nodes more complex, as for each TCP session they have to control packet delivery to succeeding proxies.

2.3.2.2 Network and physical cross layer proposals

Preemptive routing in ad hoc networks: As we report in Section 1.1, in MANETs TCP may suffer from long idle periods induced by frequent route failures. This proposal [GAGPK01] addresses this problem by reducing the number of route failures. In addition, it also reduces the route reconstruction latency. These are achieved by switching to a new route when a link of the current route is expected to fail in the future (see below). This technique is coupled with the on-demand routing protocol AODV and DSR. The detection mechanism of failure is power based. More specifically, when an intermediate node along a route detects that the signal power of a packet received from its upstream node drops below a given, called preemptive threshold, this intermediate node will detect a routing failure. For example in Figure 2.4, when node 4 senses that the signal power of a packet received

from node 5 drops below the preemptive threshold, it will detect the routing failure event. On detecting this event, node 4 will notify the source of the route, node 1. On receiving this notification, the source's routing agent proactively looks up a new route. When the new route is available, the routing agent switches to this new route. The value of the preemptive threshold appears to be critical. Indeed, in the case of a low threshold value, there will not be sufficient time to discover an alternative path before the route fails. Also, in the case of high value, the warning message will be generated too early. To overcome the fluctuations of the received signal power due to channel fading and multipath effects, that may trigger a preemptive route warning and cause unnecessary route request floods, the authors use a repeated short message probing to verify the correctness of the warning message. For example in Figure 2.4, when the signal power of the packet sent from node 5 and received at node 4 drops under the preemptive threshold, node 4 sends a ping message to node 5. On receiving this message, node 5 replies with a pong message. On receiving the pong message, node 4 checks the signal power of the pong message to verify that the link 4 – 5 is going to fail. The ping-pong handshake is repeated several times.

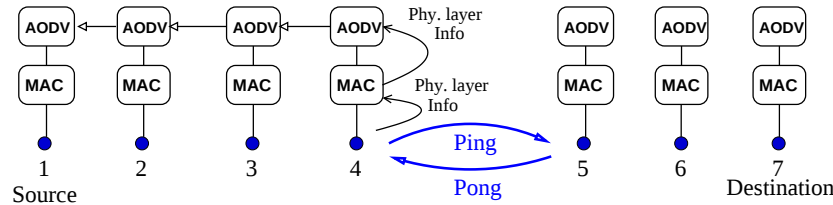


Figure 2.4: Network and physical cross layer.

Using simulations the authors show that their scheme yields a reduction of the number of route failures and decreases the latency by 30%. It should be noted that this scheme is “packet receipt event-driven” and that failures cannot be detected if no packets are transmitted.

Signal strength based link management in ad hoc network: This algorithm [KKT03] is similar to the previous one. However, in this algorithm each node keeps a record of the received signal strengths of 1-hop neighboring nodes. Using these records, the routing protocol predicts link break event in the immediate future⁴; this prediction is called Proactive Link Management. On detecting this event, the source's routing agent is notified by a Going Down message – see Figure 2.4. On receiving this message the source's routing agent stops sending packets, and initiates a route discovery procedure. The novelty of this proposal is the Reactive Link Management mechanism. This mechanism increases the transmission power to re-establish a broken link. Reactive and Proactive Link Management mechanisms can be coupled in the following way: on predicting that a link is going to be down, the node's routing agent notifies the source to stop sending, and this node increases its transmitting

⁴Immediate future means after 0.1sec.

power to handle packets in transit that use this link.

The authors use simulations to show that their scheme yields 45% improvement in TCP performance. Note that only light load scenarios are considered.

2.3.2.3 Network layer proposal

Backup path routing: This proposal [LXG03] aims to improve the path availability of TCP connections using multipath routing. The authors found that the original multipath routing deteriorates TCP performance. This is due to inaccuracy in average RTT measurement and out-of-order packet delivery. Thus, they introduce a new variation of multipath routing, called backup path routing. The backup path routing proposal maintains several paths from source to destination, but it only uses one path at any time. When the current path breaks, it can quickly switch to an alternative path. Using simulations, the authors observed that maintaining one primary path and one alternative path for each destination reports the best TCP performance. For the selection criteria of paths, the authors consider two schemes. The first scheme consists of selecting the shortest-hop path as primary and shortest-delay path as alternative. The second scheme consists of selecting the shortest-delay path as primary and the maximally disjoint path as alternative. The alternative maximally disjoint path is the path which has the fewest overlapped intermediate nodes with the primary path. Comparing the two selection schemes, they found that the first scheme outperforms the second one. The reason is that routes tend to be longer in number of hops using the second scheme. Comparing TCP performance over the DSR routing protocol with backup routing, the authors report an improvement in TCP throughput of up to 30% with a reduction in routing overheads. These results are based on various mobility and traffic load scenarios. However, the authors did not evaluate the scheme where the shortest path is selected as primary and the maximal disjoint as alternative.

2.3.2.4 Comparison

Four proposals have been presented. These proposals address the problem of frequent route failures that induce long idle times. The authors of these proposals proceed with a TCP and network cross layer solution, like Split TCP, or a network and physical cross layer solution, like Preemptive routing, and Signal strength based link, or a network layer solution, like Backup routing. The main cause of route failures in MANETs is node mobility. Split TCP is based on splitting long TCP connections, in terms of hops, into short segments to decrease the number of routing failures. However, this generates more overhead. In Preemptive routing and Signal strength based link management, the problem is addressed by predicting link failures and initiating a route reconstruction before the current route breaks. Signal strength based link management proposal uses a more robust approach to predict failures than Preemptive routing. Backup routing improves the TCP path availability by storing an alternative path. This path is used when a routing failure is detected. Backup routing

reports an improvement in TCP throughput of up to 30% with a reduction in routing overhead. But, further evaluations are needed especially for the route selection criteria.

2.3.3 Proposals to reduce the wireless channel contention

The proposals that address the problem of the contention on the wireless channel in SANETs can be classified to three categories: TCP layer proposals, network layer proposals, and link layer proposals.

2.3.3.1 TCP layer proposals

Dynamic delayed Ack: This approach [AJ03] aims to reduce the contention on the wireless channel, by decreasing the number of TCP ACKs transmitted by the sink. It is a modification of the delayed ACK option (RFC 1122) that has a fixed coefficient $d = 2$. In fact, d represents the number of TCP packets that the TCP sink should receive before it acknowledges these packets. In this approach, the value of d is not fixed and it varies dynamically with the sequence number of the TCP packet. For this reason, the authors define three thresholds $l1$, $l2$, and $l3$ such that $d = 1$ for packets with sequence number N smaller than $l1$, $d = 2$ for packets with $l1 \leq N \leq l2$, $d = 3$ for $l2 \leq N \leq l3$ and $d = 4$ for $l3 \leq N$. In their simulations, they study the packet loss rate, throughput, and session delay of TCP New Reno, in the case of short and persistent TCP sessions on a static multihop chain. They show that their proposal, with $l1 = 2$, $l2 = 5$, and $l3 = 9$, outperforms the standard TCP as well as the delayed ACK option for a fixed coefficient $d = 2, 3, 4$. They suggest that better performance could be obtained by making d a function of the sender's congestion window instead of a function of the sequence number.

2.3.3.2 Network layer proposals

COPAS: COntention-based PAtH Selection proposal [CDA03] addresses the TCP performance drop problem due to the contention on the wireless channel. It implements two techniques: the first one is disjoint forward and reverse routes, which consists of selecting disjoint routes for TCP data and TCP ACK packets. The second one is dynamic contention-balancing, which consists of dynamically update disjoint routes. Once the contention of a route exceeds a certain threshold, called backoff threshold, a new and less contented route is selected to replace the high contended route. In these proposals, the contention on the wireless channel is measured as a function of the number of times that a node has backed off during each interval of time. Also at any time a route is broken, in addition to initiating a route re-establishment procedure, COPAS redirects TCP packets using the second alternate route. Comparing COPAS and DSR, the authors found that COPAS outperforms DSR in term of TCP throughput and routing overheads. The improvement of TCP throughput is up to 90%. However, the use of COPAS, as reported by the authors, is limited to static networks or networks with low mobility. Because, as nodes move faster using a disjoint

forward and reverse routes increases the probability of route failures experienced by TCP connections. So, this may induce more routing overheads and more packet losses.

2.3.3.3 Link layer proposals

Link RED: Link Random Early Detection (RED) [FZL⁺03] aims to reduce the contention on the wireless channel. This is done by monitoring the average number of retransmissions (*avg*) at the link layer. When *avg* number becomes greater than a given threshold, the probability of dropping/marking is computed according to the RED algorithm [FJ93]. Since it marks packets, Link RED can be coupled with ECN in order to notify the TCP sender about the congestion level [RFB01]. However, instead of notifying the TCP sender about the congestion level, the authors increase the backoff time at the MAC layer.

Adaptive pacing: The goal of this proposal [FZL⁺03] is to improve spatial channel reuse. In the current IEEE 802.11 protocol, a node is constrained from contending for the channel by a random backoff period, plus a single packet transmission time that is announced by a RTS or CTS frame. However, the exposed receiver problem persists due to the lack of coordination between nodes that are two hops away from each other. Adaptive pacing solves this problem by increasing the backoff period by an additional packet transmission time. This proposal works together with Link RED as follows. Adaptive pacing is enabled by Link RED. When a node finds its average number of transmission retries to be less than a threshold, it calculates its backoff time as usual. When the average number of retries goes beyond this threshold, adaptive pacing is enabled and the backoff period is increased by a duration equal to the transmission time of the previous packet. Working together, Link RED and Adaptive pacing have reported an improvement in TCP throughput as well as the fairness between multiple TCP sessions. However, in this proposal the additional backoff time is based on the packet size. So the existence of different data packets sizes in the network should be inspected.

2.3.3.4 Comparison

Three proposals have been presented. These proposals address the problem of contention on the wireless channel, which is the main cause of TCP performance degradation in SANETs. The authors of proposals proceed with a TCP layer proposal, like Dynamic delayed ACK, or with a network layer proposal, like COPAS, or with a link layer proposals, like LRED/Adaptive pacing. Dynamic delayed ACK is a simple approach that aims to reduce the contention on the wireless channel by decreasing the number of TCP ACKs transmitted by the sink. However, COPAS attacks this problem by using disjoint forward and reverse routes, and dynamic update of disjoint routes is based on contention level. COPAS reports an improvement in TCP throughput up to 90%. LRED and Adaptive pacing attack the contention problem from its origin at the link layer. This is done by detecting the contention and increasing the backoff time at the link layer before packets transmission.

However, these two proposals is based on the packet size to compute the backoff time, and this may generate problems in case of networks that supports multiple packet sizes.

2.3.4 Proposals to improve TCP fairness

The proposals that address the problem of TCP unfairness in SANETs can be classified to one category link layer proposals.

2.3.4.1 Link layer proposals

Non work-conserving scheduling: The goal of this proposal [YSY03] is to improve fairness among TCP flows crossing wireless ad hoc and wired networks. The authors adopt a “non work-conserving scheduling” policy for ad hoc networks instead of the “work-conserving scheduling”. This is done as follows. The link layer queue⁵ sets a timer whenever it sends a data packet to the MAC. The queue only outputs another packet to the MAC when the timer expires. The duration of the timer is updated according to the queue output rate value. Specifically, the duration of the timer is a sum of three parts D_1 , D_2 , and D_3 . D_1 is equal to the data packet length divided by the bandwidth of the channel. D_2 is a delay, the value of which is decided by the recent queue output rate. D_3 is a random value uniformly distributed between 0 and D_2 . D_3 is used to avoid synchronization problems and to reduce collisions. The queue calculates the output rate by counting the number of bytes, C , it output in every fixed interval T . To decide on the value of D_2 , the authors set three thresholds X , Y , and Z ($X < Y < Z$) for C , and four delay values D_{21} , D_{22} , D_{23} , and D_{24} ($D_{21} < D_{22} < D_{23} < D_{24}$) for D_2 , as shown in (2.1). The heuristic behind (2.1) is to penalize greedy nodes with high output rate by increasing their queuing delay D_2 and to favor nodes with small output rates. By means of simulations, the authors report that their scheme greatly improves the fairness among TCP connections at the cost of moderate total throughput degradation.

$$\left\{ \begin{array}{ll} D_2 = D_{21} & \text{for } C \leq X \\ D_2 = D_{22} & \text{for } X < C \leq Y \\ D_2 = D_{23} & \text{for } Y < C \leq Z \\ D_2 = D_{24} & \text{for } C > Z \\ 0 \leq D_{21} < D_{22} < D_{23} < D_{24} & \end{array} \right. \quad (2.1)$$

Neighborhood RED: This proposal [XGQS03] aims to enhance TCP fairness in MANETs. Unlike in wired networks, the authors show that RED does not solve TCP’s unfairness in MANETs, because the congestion does not happen in a single node, but in an entire area involving multiple nodes. The local packets queue at any single node cannot completely reflect the network congestion state. For this reason, the authors define a new distributed queue, called neighborhood queue. At a node, the neighborhood queue should contain all

⁵In NS-2, Link layer queue is called InterFace queue (IFq).

packets whose transmissions will affect its own transmission in addition to its packets⁶. Since it is difficult to get information about all these packets without introducing significant communication overhead, which may need 2-hop information exchange, a simplified node neighborhood queue is introduced. It aggregates the node's local queue, and the upstream and downstream queues of its 1-hop neighbors. Now, the RED algorithm is based on the average queue size of the neighborhood queue. The authors use a distributed algorithm to compute the average queue size. In this algorithm, the time is slotted. During each time slot, the idle period of the channel is measured. Using this measurement, a node estimates the channel utilization and the average neighborhood queue size. The accuracy of the estimation is controlled by the duration of the slots. Using simulations of SANETs, they verify the effectiveness of their proposal and the fairness improvement of TCP.

2.3.4.2 Comparison

Two proposals have been presented. These proposals address the problem of TCP unfairness SANETs. To deal with the TCP unfairness problem, the authors of proposals proceed with a link layer proposals, like Non work-conserving scheduling, and Neighborhood RED. In Non work-conserving scheduling proposal, this is done through penalizing greedy nodes of high output rate, by increasing their queuing delays at the link layer. However, in Neighborhood RED it is up to TCP to regulate its transmission rate when it senses packets drop. The RED algorithm in this proposal is applied to a distributed queue, called neighborhood queue, that does not contain only the node local queue but also the upstream and downstream queues of its 1-hops neighbor. Neighborhood RED is better than Non work-conserving as it does not cause degradation in total throughput.

2.4 General comparison and discussion

It is difficult to compare the different proposals that aim to improve TCP's performance in MANETs. Some of them have been designed to solve different problems. By examining all these proposals, we identify four major problems: the first one is "TCP is unable to distinguish between losses due to route failures and network congestion", the second one is "frequent route failures", the third one is "wireless channel contention", and the fourth one is "TCP unfairness". We note that the first two problems are the main causes of TCP performance degradation in MANETs. However, the second two problems are the main causes of TCP performance degradation in SANETs. To solve these problems, the authors of the proposals proceed with a cross layer solution or a layered solution. In terms of complexity, cross layer solutions are more complex to implement and to design than layered solutions. The reason is that the implementation of cross layer solutions requires at least modification of two OSI layers. Also, as they break the concept of designing protocols to be independent of other layer protocols, their design requires that we consider the system

⁶Affect means : prevent a node from transmitting because of channel capturing, or to induce transmission failure to carrier sensing MAC protocols.

in its entirety. In the following, we compare the proposals that share a common problem, and Table 2.2 contains a detailed comparison of these proposals. The comparison is based on five features: the first one is the solution type that can be: cross layer or layered, the second one is the degree of complexity, the third one is the network type, the fourth one is TCP connections load and type, and the fifth one is the evaluation basis. For the degree of complexity, we identify three degrees: high, medium, and low. Layered solutions are considered to be medium or low complexity, whereas cross layer solutions are considered to be medium or high complexity.

2.5 Conclusion

We have presented the state-of-the-art of TCP over static and mobile ad hoc networks (SANETs and MANETs). The principal problem of TCP in this MANETs environment is clearly its inability to distinguish between losses induced by network congestion and other types of losses. TCP assumes that losses are always due to network congestion. But while this assumption in most cases is valid in wired networks, it is not true in MANETs. In MANETs, there are indeed several types of losses, including losses caused by routing failures, by network partitions, and by high bit error rates. Performing congestion control in these cases, like TCP does, yields poor performance. In static multi-hop ad hoc networks the principal problem of TCP is the contention on wireless channel that induces routes failures and losses. In order to solve these problems, several proposals have been made in the literature. We classified these proposals as layered and cross layer proposals. In cross layer proposals, TCP and the underlying protocols cooperate to improve Ad hoc network performance. For example, TCP and network cross layer proposals use explicit notification to inform TCP about route failures, which requires cooperation from intermediate nodes like in TCP-F, TCP-ELFN, ATCP, TCP-BuS. In layered proposals, one OSI layer is adapted. For example, TCP layer proposals require only the cooperation of the sender and receiver, like in TCP-DOOR and Fixed-RTO. However, cross layer proposals yield higher improvement than layered ones.

The frequent route failures and route re-establishments in MANET environment introduce a new challenge to TCP congestion control algorithm, and leads us to investigate how we can take advantage of the mobility of the nodes in MANETs. This will be done in the second parts of thesis using the two-hop relay protocol.

2.A Overview of TCP versions

TCP is a window-based acknowledgement-clocked flow control protocol. It uses additive-increase multiplicative-decrease strategy for changing its window as a function of network conditions. Packets of a TCP connection are sent with increasing consecutive sequence numbers. In the simplest operation of TCP, at each arrival of a packet at the destination, an ACK is sent back to the source with the information of the next sequence number that is expected. Thus if all packets up to packet $n - 1$ have reached the destination, then the last arrival will trigger an ACK with sequence number n . If a packet n is lost in the network

Problem	Proposal	Solution type	Complexity degree	Network type	TCP load /type	Evaluat.
TCP is unable to distinguish RFs and congestion losses	ELFN [HV02]	TCP-net. cross layer	medium	random mobile	one persistent	simulation no routing
	ATCP [LS01]	TCP-net. cross layer	high	random mobile	one persistent	experiment. no routing
	TCP-BuS [KTC01]	TCP-net. cross layer	high	random mobile	multi. persistent	simulation no routing
	TCP-F [CRVP98b]	TCP-net. cross layer	medium	random mobile	one persistent	emulation no routing
	TCP-DOOR [WZ02]	TCP layer	low	random mobile	one persistent	simulation
	TCP-RTO [DB01]	TCP layer	low	random mobile	one persistent	simulation
Frequent route failures	Split TCP [KKFT02]	TCP-net. cross layer	medium	random mobile	one persistent	simulation
	Preemptive routing [GAGPK01]	Net.-phy. cross layer	medium	random mobile	no TCP load	simulation
	Signal strength based link [KKT03]	Net.-phy.	high	random mobile	two persistent	simulation
	BR [LXG03]	Network layer	low	random mobile	one persistent	simulation
Contention on wireless channel	Dynamic delayed ACK [AJ03]	TCP layer	low	static chain	one persist. and short	simulation
	COPAS [CDA03]	Network layer	medium	static random	multi. persistent	simulation
	LRED/Adaptive pacing [FZL ⁺ 03]	Link layer	low	static	multi. persistent	simulation
TCP unfairness	NWC [YSY03]	Link layer	low	static chain	multi. persistent	simulation
	Neighborhood RED [XGQS03]	Link layer	low	static/mobile	multi. persistent	simulation

Table 2.2: General comparison of proposals.

and packet $n + i$, $i = 1, 2, 3$, arrives at the destination, then each of these packets will trigger an Acknowledgment indicating that the destination is expecting packet n . These are called duplicated ACKs. In the absence of losses, starting from one packet or from a larger value, the window is increased exponentially by one packet every non-duplicate ACK until the source estimate of network capacity is reached. This is the Slow Start (SS) phase, and the capacity estimate is called the slow start threshold (ssthresh). In most versions of TCP (Tahoe, Reno and New Reno) once ssthresh is reached, the source switches to a slower increase in the window by one packet for every window's worth of ACKs. This phase is called Congestion Avoidance (CA) phase. The window increase is interrupted when a loss is detected. Two mechanisms are available for the detection of losses: the expiration of a retransmission timer (timeout), or the receipt of three duplicate ACKs (the latter is called the fast retransmit (FRXT) phase). The source then sets its estimation of the capacity to half the current window; this action is due to the fact that when these TCP versions have been developed, losses were indication of congestion as TCP was then deployed only over wireline networks. Tahoe [Jac98], the first version of TCP to implement congestion control, at this point sets the window to one packet and enter the slow start phase to reach the new ssthresh. Slow starting after every loss detection deteriorates the performance given the low bandwidth utilization during SS. When the loss is detected via timeout a more drastic reaction is taken as a more drastic congestion is understood to occur, since the ACK stream has stopped. In the FRXT case, ACKs still arrive at the source, and losses are recovered without SS. This is the behavior of the newer versions of TCP (Reno, NewReno, SACK, etc.) that call a Fast Recovery (FRCV) algorithm to retransmit the losses while maintaining enough packets in the network to preserve the ACK clock. Once the losses are recovered, this algorithm ends and normal CA is called. If FRCV fails (to recover the losses), the ACK stream stops, a timeout occurs, and the source resorts to SS as with Tahoe. Among the TCP versions that use the FRCV, the difference is in the estimation of the number of packets in the flight during FRCV. Reno [FF96] considers every duplicate ACK a signal that a packet has left the network. The problem of Reno is that it leaves FRCV when an ACK for the first loss window is received. This prohibits the source from detecting the other losses with FRXT. A long timeout is required to detect the other losses. NewReno [Heo96] has been proposed to overcome this problem. The idea is to stay in FRCV until all the losses in the same window are recovered. Another problem of Reno and newreno is that they rely on ACKs to estimate the number of packets in flight. ACKs can be lost on the return path, which results in an underestimation of the number of packets that have left the network. More information is needed to estimate more precisely the number of packets in the pipe. This information is provided by the selective ACK (SACK) [AGS99], a TCP option containing the three blocks of contiguous data most recently received at the destination. Finally, we mention TCP Vegas that aims to decouple congestion detection from losses. In TCP Vegas, the RTT of the connection and the window size are used to compute the number of packets in the network buffers. The window is decreased when this number exceeds a certain threshold and is increased when it is below some threshold. In other words, Vegas also uses delay as a congestion indication and then reacts to reduce its

throughput.

Chapter 3

Impact of Mobility on the throughput of Relaying in Ad Hoc Networks

Considered is a mobile ad hoc network consisting of three types of nodes: source, destination, and relay nodes. All the nodes are moving over a bounded region with possibly different mobility patterns. This chapter defines and studies the notion of *relay throughput*, i.e., the maximum rate at which a node can relay data from the source to the destination. We show that the relay throughput depends on the node mobility pattern only via its (stationary) node position distribution and that a node mobility pattern that results in a uniform steady-state distribution for all nodes achieves the lowest relay throughput. Random Waypoint and Random Direction mobility models in both one and in two dimensions are studied and approximate simple expressions for the relay throughput are provided.

Note: Part of the results in this chapter have appeared in [HKG⁺05, HKG⁺06].

3.1 Introduction

In a Mobile ad hoc Network (MANET), since there is no fixed infrastructure and nodes are mobile, links between nodes are set up and turned down dynamically. A link between two nodes is up when these nodes are inside one another's communication range, and a link is down otherwise. The establishment of a route from a source node to a destination node requires the simultaneous availability of a number of links that are all up, one originating at the source node and another one ending at the destination nodes. Indeed, MANETs often experience route failures and network disconnectivity, especially when the nodes are moving frequently and the network is sparse. This is a typical scenario of the Delay Tolerant Network (DTN) [dtn], which is a subclass of MANET.

Grossglauser and Tse [GT02] have observed that mobility in MANETs can be used to increase the average network throughput. Their idea was to look at the diversity gain achieved by using the mobile nodes as relays. Their relay mechanism, called *two-hop relay* protocol, is simple: if there is no route between the source node and the destination node, the source node transmits its packets to all neighboring nodes (called *relay* nodes) that it meets for delivery to the destination. A relay node is only allowed to send a packet to its destination node, and it is not allowed to send the packet to another relay node, thereby justifying the name of this protocol. It was then shown in [GMPS04] that a bounded delay can be guaranteed under this relaying mechanism. The aim of these studies (see also [GK00]) is the scaling property of the throughput or delay as the number of nodes in the network becomes large. Our interest in the present chapter is in the performance of the above mentioned relaying mechanism in a network consisting of a fixed finite number of nodes.

It is important to mention that most of the studies of scaling laws of delay or throughput in wireless ad hoc networks assume a uniform spatial distribution of nodes, which is the case, for example, when the nodes perform a symmetric Random Walk over the region of interest [GMPS04, GT02]. This chapter studies the impact of the node mobility pattern on the throughput of the two-hop relay routing of [GT02]. The interest is in the *maximum relay throughput* of a mobile node, i.e., the maximum that a node can contribute as a relay to the communication between two other nodes. The relaying of data for other nodes requires a relay node to allocate its own resources. In particular, a relay node has to keep the data to be relayed in its buffer. Hence, the study of the buffer behavior of a relay node forms an important topic of research, this issue will be studied in details in Chapter 4. The present chapter will focus on the impact of the nodes mobility on the maximum relay throughput and the stability of the *relay buffer* (RB).

The point of departure is a simple observation which relates the evolution of a relay node buffer to the evolution of the workload process in a G/G/1 queueing system. The service requirements and inter-arrival times in this queueing system are determined by the characteristics of the mobility pattern of the nodes.

The main findings are the following:

1. The relay throughput depends only on the stationary distribution of the nodes position. Hence, any two different mobility patterns that have the same stationary distribution will achieve the same relay throughput.

2. This *relay throughput* achieved is the lowest when the steady-state distribution of the nodes location are uniformly distributed.

An important point that needs to be emphasized is that, unlike [GMPS04, GT02, GK00] which study the system performance when the number of nodes is large, this chapter interests in a relay node performance while it is involved in relaying data between two particular nodes.

The rest of the chapter is organized as follows: Section 3.2 describes the system model considered. In Section 3.3 a queueing model is developed for the relay buffer (RB). Section 3.4 studies the impact of mobility models on the relaying throughput, and in Section 3.5 the expressions are found for the relay throughput for the Random Waypoint (RWP) and the Random Direction (RD) models in both one and in two dimensions. Section 4.4 reports the numerical results on the stability, relay throughput, probability of a 2-hop route. Section 3.7 concludes the chapter.

3.2 The system model

To study the maximum rate at which a node can relay data, this chapter starts by considering the scenario where three nodes move in a two-dimensional bounded region. One of these nodes is the source of packets, one is the destination, and the third one is the relaying node. The mobility patterns of the three nodes are independent and may be different from each other; this is in contrast with [GMPS04, GT02] where the authors assume that the mobility pattern of the nodes is such that the steady-state distribution of the location of all the nodes is uniform over the region of interest. In fact, [GMPS04] assumes that nodes perform random walks (there are other mobility models which also result in a uniform stationary distribution, e.g., the Random Direction (RD) model [PNL05]). As mentioned earlier, this chapter focuses in the maximum relay throughput of a relay node. As a starting point restricts yourselves to the case where there is only one relay node. At a later stage this assumption will be relaxed. It is assumed that a node detects its one-hop neighbor(s) by sending periodically Hello messages. However, to detect the two-hop neighbor(s) nodes exchange the addresses of their neighbor(s).

The model is the following:

1. The three nodes move independently of each other according to a (possibly node-dependent) mobility model inside a bounded 2-dimensional region.
2. The source node transmits a packet only once (either to the relay or to the destination node). Thus, the source node does not keep a copy of the packet once it has been sent. When the source node transmits a packet to the destination node (when their locations permit such a transmission), the source node transmits packets that it has not transmitted before.

3. The source node has always data to send to the destination node. This is a standard assumption, also made in [GMPS04, GT02, GK00], because we are interested in the maximum relay throughput of the relay node.
4. When the relay node comes within the transmission range of the source node (we will also say that nodes are *in contact* in this case), and if the destination node is outside the transmission range of the source and of the relay node, then the relay node accrues packets to be relayed to the destination node at a constant rate r_s . These packets are buffered in a dedicated buffer at the relay node called the relay buffer (RB). [We could allow for a stochastic nature of traffic generated by the source by assuming that r_s is an independent stochastic process. However, such a study is out of the scope of this chapter.]
5. When the destination node comes within the transmission range of the relay node, and if the destination and the relay node are outside transmission range of the source node, then the relay node sends the relay packets (if any packets in its RB) to the destination node at a constant rate r_d .
6. If the relay node is within transmission range of *both* the source node and the destination node, then the relay node does not contribute to *relaying*. In this case there is either a direct communication between the source and destination or there is a two-hop route via the relay node so that the relay node acts as a forwarding node and not as a relay. Further, we assume that forwarding data has much more priority than relay data. Since the source is assumed to have always data to send, the relay will not get the chance to transmit the relay data from its RB in this case.
7. We assume that a data packet can not be relayed more than once to reach its destination (i.e., a relay node may only send packets to the destination and not to another relay node).

The objective is to study the properties of the relay buffer (stability, throughput). To this end, a queueing model will be developed to give many insights into the system behavior.

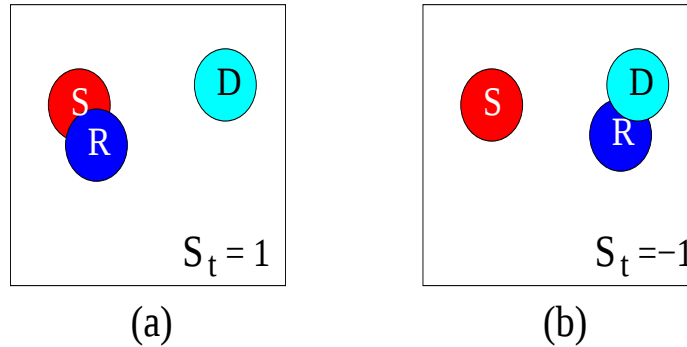
3.3 Queueing model for the relay buffer

After addressing the case where there are only three mobile nodes Section 3.3.1 investigates the situation of an arbitrary number of source/destination/relay nodes (cf. Section 3.3.2).

3.3.1 Single source, destination, and relay nodes

The state of the relay node at time t is represented by the random variable (rv) $S_t \in \{-1, 0, 1\}$ where:

- $S_t = 1$ if at time t the relay node is neighbor (i.e., within transmission range) of the source, and if the destination is a neighbor neither of the source nor of the relay node. See Figure 3.1.a. In other words, when $S_t = 1$, the source node sends relay packets to the relay node at time t ;
- $S_t = -1$ if at time t the relay node is a neighbor of the destination, and if the source is a neighbor neither of the destination nor of the relay node. See Figure 3.1.b. When $S_t = -1$ the relay node delivers relay packets (if any) to the destination;
- $S_t = 0$ otherwise.


 Figure 3.1: The process $\{S_t\}$.

Mobiles have finite speeds. Assume that the relay node may only enter state 1 (resp. -1) from state 0: if $S_{t-} \neq S_t$ then necessarily $S_t = 0$ if $S_{t-} = 1$ or $S_{t-} = -1$, where

$$t- = \lim_{\epsilon \rightarrow 0, \epsilon > 0} t - \epsilon.$$

Denote by B_t the RB occupancy at time t . The rv B_t evolves as follows:

- it increases at rate r_s if $S_t = 1$. This is because when $S_t = 1$, the relay node receives data to be relayed from the source node at rate r_s ;
- it decreases at rate r_d if $S_t = -1$ and if the RB is non-empty. This is because if $S_t = -1$, and if there is any data to be relayed, then the relay node sends data to the destination node at rate r_d .
- it remains unchanged in all other cases, for more details see system model of Section 3.2 especially items 4 and 6.

Let $\{Z_n\}_n$ ($Z_1 < Z_2 < \dots$) denote the consecutive jump times of the process $\{S_t, t \geq 0\}$. An instance of the evolution of S_t and B_t as a function of t is displayed in Figure 3.2.

The evolution of the discrete indexed process $\{S_{Z_k}, k \geq 1\}$ consists of sequences of 1, 0 and -1 . Without loss of generality assume that the process $\{S_t, t \geq 0\}$ is a right-continuous process.

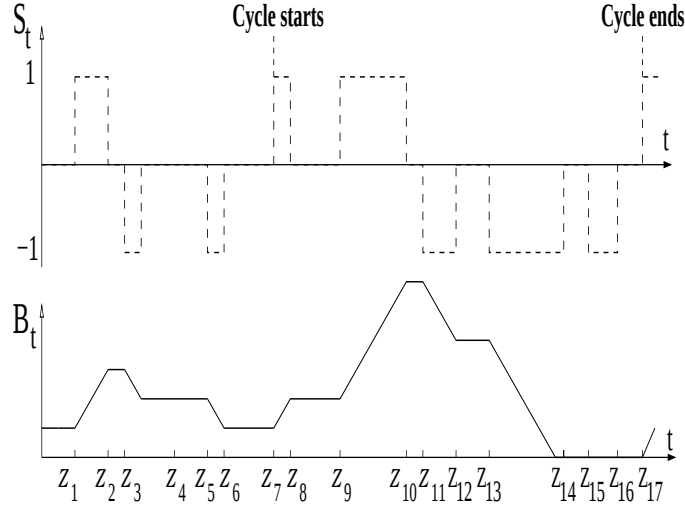


Figure 3.2: Evolution of $\{S_t\}_t$ and relay node buffer occupancy.

Define a *cycle* as the interval of time that starts at $t = Z_k$, for some k with $S_t = 1$, and (necessarily) $S_{Z_{k-1}} = 0$ and $S_{Z_{k-2}} = -1$, and ends at the smallest time $t + \tau$ such that $S_{t+\tau} = 1$ and $S_{t+s} = -1$ for some $s < \tau$. In Figure 3.2, the time-interval $[Z_7, Z_{17})$ constitutes a cycle. Let C_n denotes the duration of the n^{th} cycle. Note that there is no restriction on the number of times the relay node becomes neighbor of the source node or of the destination node during a cycle. Also, note that the end time of the n^{th} cycle is also the beginning time of the $(n+1)^{th}$ cycle.

Let W_n be the time at which the n^{th} cycle begins. Let

$$\sigma_n := \int_{t=W_n}^{W_{n+1}} \mathbf{1}_{\{S_t=1\}} dt \quad (3.1)$$

be the amount of time spent by the relay node in state 1 during the n^{th} cycle. Similarly, let

$$\alpha_n := \int_{t=W_n}^{W_{n+1}} \mathbf{1}_{\{S_t=-1\}} dt \quad (3.2)$$

be the amount of time spent by the relay node in state -1 during the n^{th} cycle. Observe that during the amount of time σ_n , the RB increases at rate r_s (source to relay node transmission rate), and it decreases at rate r_d during the amount of time α_n . Let $\tilde{B}_n$ be the RB occupancy at the beginning of the n^{th} cycle, i.e., $\tilde{B}_n = B_{W_n}$. Clearly,

$$\tilde{B}_{n+1} = [\tilde{B}_n + r_s \sigma_n - r_d \alpha_n]^+ \quad (3.3)$$

where $[x]^+ := \max(x, 0)$. In other words, $\tilde{B}_{n+1}$ can be interpreted as the workload seen by the $(n+1)^{st}$ arrival in a G/G/1 queue, where $r_s \sigma_n$ is the service requirement of the n^{th}

customer, and $r_d\alpha_n$ is the inter-arrival time between the n^{th} and the $(n+1)^{st}$ customer. This interpretation will be used next.

Assumption A: Throughout Section 3.3.1 assume that the sequence $\{C_n, \sigma_n, \alpha_n\}_n$ is stationary and ergodic, with $0 < E[C_n] < \infty$, $0 < E[\sigma_n] < \infty$ and $0 < E[\alpha_n] < \infty$.

Clearly, the statistical properties of the random variables C_n , σ_n , and α_n will depend on the node mobility patterns. Hence, the study will be restricted to the class of mobility models under which the stationarity and ergodicity assumptions hold for the sequence $\{C_n, \sigma_n, \alpha_n\}_n$.

Definition: The long-term fraction of time the RB receives data is

$$\pi_s := \lim_{t \rightarrow \infty} \frac{1}{t} \int_{u=0}^t \mathbf{1}_{\{S_u=1\}} du, \quad (3.4)$$

and the long-term fraction of time that the destination node is the neighbor of only the relay node (i.e., the fraction of time that the RB may drain off) is

$$\pi_d := \lim_{t \rightarrow \infty} \frac{1}{t} \int_{u=0}^t \mathbf{1}_{\{S_u=-1\}} du. \quad (3.5)$$

It is easy to show that these limits exist under Assumption A. Moreover, Assumption A implies that (see [BB87, Part I, Sec. 4.3]),

$$\pi_s = \lim_{t \rightarrow \infty} P(S_t = 1) = \frac{E[\sigma_n]}{E[C_n]} \quad (3.6)$$

and

$$\pi_d = \lim_{t \rightarrow \infty} P(S_t = -1) = \frac{E[\alpha_n]}{E[C_n]}. \quad (3.7)$$

Theorem 1 *If $r_s E[\sigma_n] < r_d E[\alpha_n]$ then $\tilde{B}_n$ converges in probability to a proper and finite rv $\tilde{B}$ (i.e., $\lim_n P(\tilde{B}_n < x) = P(\tilde{B} < x) \forall x \geq 0$). $r_s E[\sigma_n] > r_d E[\alpha_n]$, then $\tilde{B}_n$ converges to $+\infty$ P - a.s.* $\square$

Proof. Follows from the relation to the G/G/1 queue made above and [Loy62]. $\blacksquare$

In the case when $r_s E[\sigma_n] < r_d E[\alpha_n]$, the RB will be said to be stable, however when $r_s E[\sigma_n] \geq r_d E[\alpha_n]$ the RB will be said to be unstable. Thus, for this reason the condition $r_s E[\sigma_n] < r_d E[\alpha_n]$ will be called the stability condition of Theorem 1.

Remark 1 *In terms of π_s and π_d , the stability condition of Theorem 1 writes*

$$r_s \pi_s < r_d \pi_d.$$

Remark 2 *If all nodes have the same mobility pattern, then clearly $\pi_s = \pi_d$, since the relay node is equally likely to be within the transmission range of the source and of the destination. Therefore, by Remark 1, the stability condition is*

$$r_s < r_d.$$

Theorem 2 *If $r_s\pi_s < r_d\pi_d$, then the relay throughput T_r , defined as the stationary output rate of RB of the relay node, is given by*

$$T_r = r_s\pi_s. \quad \square$$

Proof. In steady-state the RB can be thought of as a standard G/G/1 queue so that the output rate is the same as the input rate and is given by $r_s\pi_s$. ■

Remark 3 *The relay throughput T_r only depends on r_s and the stationary distribution of the node mobility pattern. In particular, two different mobility patterns with the same stationary distribution (for the location of the nodes) will yield the same relay throughput.*

It is clear from Theorem 2 and Remark 1 that π_s and π_d play an important role in determining the stability and the throughput of the RB. Much of the rest of this chapter will be devoted to the study of these quantities.

3.3.2 Multiple source, destination, and relay nodes

Assume now that there are K source nodes, M destination nodes and N relay nodes, all with the same transmission range R (the latter will be assumed throughout). The source and destination nodes are stationary (not moving). The relay nodes move independently of each other inside a connected area A according to a common mobility model. The distance between any two source nodes, and between any two destination nodes, is assumed to be greater than $2R$. This implies that a relay node can not receive (resp. transmit) data from (resp. to) two or more source (resp. destination) nodes at the same time.

Furthermore, assume that the routing protocol generates routes of length no more than h hops, i.e., the lifetime of a packet in number of hops is not greater than h . The distance between any source and any destination node is set to be greater than hR . Therefore, there does not exist a direct route from any source to any destination node, which implies that packets of a source have to use the mobile relay nodes in order to be transferred to their destinations.

The RB of a relay node is composed of M queues; one for each of the M destinations. The system behaves as follows:

1. When there are i relay nodes inside the transmission range of source node k , where $i \in \{1, \dots, N\}$ and $k \in \{1, \dots, K\}$, then the source transmits to the i relay nodes the packets addressed to destination node $m \in \{1, \dots, M\}$ with probability P_k^m in a round-robin scenario, where $\sum_{m=1}^M P_k^m = 1$. Note that we assume that there is only one copy of the packet. Thus queue m of the relay node accrues packets at a fixed rate $r_{S_k} P_k^m / i$, where r_{S_k} is the transmission rate of source node k .

2. When the relay node receives a packet from a source that is destined to destination node m , it buffers this packet in its queue of index m .
3. When there are j relay nodes with non-empty queue m inside the transmission range of destination m , these relay nodes share the channel bandwidth fairly. More precisely, queue m of these j relay nodes drains off at a fixed rate r_{D_m}/j , where r_{D_m} is the transmission rate of a relay node to the destination node m . The service discipline in queue m of the relay node is FIFO.

Let $f(x)$, $x \in A$, be the stationary node location probability density. The stationarity of the node location is under time shift. Denote by x_{S_k} and x_{D_m} the fixed location in A of source $k \in \{1, \dots, K\}$ and destination $m \in \{1, \dots, M\}$, respectively.

Hence, the probability that a relay node is the neighbor of a node located in $x \in A$ is

$$\pi(x) = \int_{\{y \in A: d(x, y) \leq R\}} f(y) dy, \quad (3.8)$$

where $d(u, v)$ is the Euclidean distance between vectors u and v .

By conditioning on the number of nodes within range of source node k , one can find that the input rate at queue m of each relay node is

$$\begin{aligned} \tau_{S_k}^m &= \sum_{i=1}^N \frac{1 P_k^m r_{S_k}}{i} \binom{N-1}{i-1} \pi(x_{S_k})^i \overline{\pi(x_{S_k})}^{N-i} \\ &= \frac{P_k^m r_{S_k}}{N} \sum_{i=1}^N \binom{N}{i} \pi(x_{S_k})^i \overline{\pi(x_{S_k})}^{N-i} \\ &= P_k^m r_{S_k} \frac{1 - (1 - \pi(x_{S_k}))^N}{N}, \end{aligned} \quad (3.9)$$

where $\bar{a} := 1 - a$.

The overall long-term arrival rate to queue m of a relay node from all of the sources is

$$\tau_S^m := \sum_{k=1}^K \tau_{S_k}^m = \frac{1}{N} \sum_{k=1}^K P_k^m r_{S_k} \left[1 - (1 - \pi(x_{S_k}))^N \right]. \quad (3.10)$$

The exact derivation of τ_D^m , the long-term service rate of queue m at a relay node, is intractable since it depends on the (stationary distribution of the) location of the other relay nodes with respect to the destination D_m , and on whether or not queue m at each relay node is empty or not and located within transmission range of D_m . More precisely, if i relay nodes are within the transmission range of destination D_m , and if queue m in each of these relay nodes is non-empty, then the service rate in queue m at each of the i relay nodes is r_{D_m}/i . The above reasoning indicates that r_{D_m}/N is the minimum instantaneous service

rate at each queue m . This yields the following lower bound—called $\hat{\tau}_D^m$ —on the long-term service rate of queue m :

$$\begin{aligned}\hat{\tau}_D^m &= \sum_{i=1}^N \frac{r_{D_m}}{i} \binom{N-1}{i-1} \pi(x_{D_m})^i \overline{\pi(x_{D_m})}^{N-i} \\ &= r_{D_m} \frac{1 - (1 - \pi(x_{D_m}))^N}{N}.\end{aligned}\quad (3.11)$$

As a result, a sufficient condition for the stability of queue m at each relay node reads

$$\tau_S^m < \hat{\tau}_D^m. \quad (3.12)$$

If queue m at a relay node is stable, then the relay throughput T_r^m at this queue is equal to its long-term arrival rate, i.e.,

$$T_r^m = \tau_S^m = \sum_{k=1}^K \tau_{S_k}^m = \frac{1}{N} \sum_{k=1}^K P_k^m r_{S_k} \left[1 - (1 - \pi(x_{S_k}))^N \right]. \quad (3.13)$$

The network throughput, T , is the sum of the relay throughputs at all the M queues of all the N relay nodes, namely

$$T = \sum_{n=1}^N \sum_{m=1}^M T_r^m = \sum_{k=1}^K r_{S_k} \left[1 - (1 - \pi(x_{S_k}))^N \right]. \quad (3.14)$$

Observe that $1 - (1 - \pi(x_{S_k}))^N$ is nothing than the probability that there is at least one relay node inside the transmission range of the source node k .

Now consider the situation where all the nodes are moving including the sources and the destinations. Since an exact calculation of the throughput of queue m at a relay node is very difficult, an approximation will be derived for this quantity. This approximation is based on the assumption that routes cannot exceed two hops. Assume that all nodes move independently of each other with the same mobility pattern, and that they have the same transmission range. Let p_1 be the probability that two nodes are within transmission range of one another. Let p_2 be the probability that three nodes constitute a two-hop route. Then, under the above simplifying assumption, the probability that a relay node is in contact with the source k and that there is no route to the destination through this relay node is $p_1 - p_2$. On the other hand, the probability that a relay node is not in contact with the source k is $1 - p_1$. Hence,

$$\sum_{i=1}^N \frac{P_k^m r_{S_k}}{i} \binom{N-1}{i-1} (p_1 - p_2)^i (1 - p_1)^{N+1-i} = P_k^m r_{S_k} \frac{(1 - p_1) ((1 - p_2)^N - (1 - p_1)^N)}{N} \quad (3.15)$$

is the contribution of source node k to the long-term arrival rate in queue m at any relay node. Therefore, the overall long-term input rate at queue m at any relay node can be approximated by summing up the r.h.s. of the above identity over all the values of k . This gives

$$\tau_S^m \approx \frac{(1-p_1)((1-p_2)^N - (1-p_1)^N)}{N} \sum_{k=1}^K P_k^m r_{S_k}. \quad (3.16)$$

When $P_k^m = 1/M$ (that is, there is a uniform probability that the source k sends packets to destination m) and when the transmission rates of all sources are equal to r_S , then (3.16) becomes

$$\tau_S^m \approx r_S \frac{(1-p_1)((1-p_2)^N - (1-p_1)^N)}{MN}. \quad (3.17)$$

The next section investigates the impact of the mobility pattern on the relay throughput. We will show that the throughput is minimized when in steady-state the nodes are uniformly distributed over the area.

3.4 Comparison of mobility models

Consider the scenario where nodes move independently of each other according to same mobility pattern. Assume that the nodes location distribution is stationary. The nodes position can take values in a discrete set X with cardinality $\#X = G$. Let $G(x), x \in X$ denote the set of all points in the transmission range of a node located at x . We assume that there is complete symmetry, so that for all $x, y \in X$ if $x \in G(y)$ then $y \in G(x)$. This can be assumed when there is no *boundary effect*, for example, as is the case of motion over a torus or over a circle (representing, respectively, motion over a plane or line with wrap around).

Let P be the probability measure over X that represents the stationary node location distribution. As the cardinality of X is equal to G , P can be represented as a G -dimensional (column) vector. The uniform stationary node location over X , called U , is a G -dimensional vector whose entries are all equal to $\frac{1}{G}$. Let $e_x, x \in X$, denote a probability measure over X which gives all mass to position x , i.e., e_x is an G -dimensional vector whose entries are all equal to 0 except for the x^{th} components which is equal to 1.

For any stationary node location distribution P over X , let $g(P)$ denote the probability that two nodes are neighbor of each other. Let H denote the neighborhood matrix, i.e., $H_{x,y} = 1$ if $y \in G(x)$ and $H_{x,y} = 0$ otherwise. Note H is a symmetric matrix. In terms of P_x (resp. P_y), the probability that a node is at location x (resp. y) in the stationary regime $g(P)$ writes

$$g(P) = \sum_{x \in X} P_x \sum_{y \in G(x)} P_y = \sum_{x \in X} P_x \sum_{y \in X} H_{x,y} P_y = P^T H P$$

where P^T is the transpose of P and we use the fact that the locations of the nodes are independent. Since in this chapter we assumed that the nodes move independently of each other, see Section 3.2.

Theorem 3 *A uniform distribution of nodes over the region of interest achieves the minimum probability of contact between any two nodes.* $\square$

Proof: Consider any P of the form

$$P = U + \delta e_x - \delta e_y, \quad (3.18)$$

for some $0 < \delta < 1$ and $x, y \in X$, $x \neq y$. Then

$$\begin{aligned} g(P) &= P^T H P \\ &= (U + \delta e_x - \delta e_y)^T H (U + \delta e_x - \delta e_y) \\ &= g(U) + \delta^2 (e_x - e_y)^T H (e_x - e_y) + \delta U^T H (e_x - e_y) + \delta (e_x - e_y)^T H U \\ &= g(U) + \delta^2 (e_x^T H e_x + e_y^T H e_y - e_x^T H e_y - e_y^T H e_x) + 2\delta (e_x - e_y)^T H U, \end{aligned} \quad (3.19)$$

where we have used the fact that for all $x \in X$ $e_x^T H e_x = 1$ as $x \in G(x)$. Since H is a symmetric matrix, we have $P^T H Q = Q^T H P$ for all P, Q probability measures on X . Also, it is easy to see that $e_y^T H e_x = 1$ if $x \in G(y)$ and is 0 otherwise. Hence we get

$$\begin{aligned} g(P) &= g(U) + 2\delta^2 (1 - \mathbf{1}_{\{x \in G(y)\}}) + 2\delta (e_x - e_y)^T H U \\ &= g(U) + 2\delta^2 \mathbf{1}_{\{x \notin G(y)\}} + \frac{2\delta (\#G(x) - \#G(y))}{G}, \end{aligned}$$

where in the last expression we have used the, easy to observe, fact that $e_x^T H U = \frac{\#G(x)}{G}$. Hence, since $\#G(x) = \#G(y)$ for all $x, y \in X$, it is seen that $g(P) - g(U) = 2\delta^2 \mathbf{1}_{\{x \notin G(y)\}} \geq 0$. Which implies that

$$U \in \underset{P=U+\delta e_x-\delta e_y}{\operatorname{argmin}} g(P). \quad (3.20)$$

Now, any other probability distribution over the set X is a point in G -dimensional canonical simplex. The uniform distribution is at the centroid of this simplex and any other distribution P , when viewed as an G -dimensional vector (a point in the simplex), can be written as

$$P = U + \epsilon, \quad (3.21)$$

where U is the uniform distribution and ϵ is an G -dimensional vector whose entries are in the interval $[-\frac{1}{G}, \frac{G-1}{G}]$ and the entries sum to zero. Clearly, any such ϵ can be written as a (possibly non-unique) finite sum

$$\epsilon = \sum_{x \in I(P)} (e_x - e_{y(x)}) \delta_x, \quad (3.22)$$

where $I(P) \subset X$ is some index set, $y(x) \in X$, and $\delta_x > 0$, $x \in I(P)$. This is because e_x forms a basis for the G -dimensional space and because P is a probability vector with $\sum_{x \in X} \epsilon_x = 0$.

Recall that if the stationary node distribution is P , we can write $g(P)$ as $g(P) = P^T H P$, where H is an $G \times G$ symmetric matrix indicating the neighborhood relation. We have already shown that when $P = U$, the uniform distribution, the directional derivative of $P^T H P$ is positive along any direction of the form $(e_x - e_y)$ where e_x is G -dimensional vector with all except the x^{th} entry equal to zero. We now use continuity of the derivative of $g(U)$ to conclude that its directional derivative along any direction is positive. Hence $U \in \underset{P}{\operatorname{argmin}} g(P)$. $\blacksquare$

The above result does not imply that the *relay throughput* achieves its minimum under the uniform stationary node distribution. This is because the relay throughput under distribution P , denoted $T_r(P) = r_s \pi_s(P)$ and with $\pi_s(P) = \lim_{t \rightarrow \infty} P(S_t = 1)$ under the probability measure P , is

$$T_r(P) = r_s \left(g(P) - \sum_{x \in X} P_x \sum_{y \in G(x)} P_y \sum_{z \in G(x) \cup G(y)} P_z \right), \quad (3.23)$$

and it can be easily seen that for any $x \in X$, $\pi_s(e_x) = 0$. Since $\pi_s(\cdot)$ is a probability, this implies that $P = e_x$ achieves minimum of $\pi_s(\cdot)$. However, it is reasonable to assume that the uniform distribution is a local minimum for $\pi_s(\cdot)$ because the second term in expression for $\pi_s(\cdot)$ above is of smaller order as compared to the first term.

Observe that if the source node and the destination node are fixed, and if they are far apart (so that a two-hop route between them via a relay node is not possible), then a uniform distribution of relay node achieves the minimum relay throughput.

The next section finds expressions for the relay throughput in the case nodes move according to the Random Waypoint (RWP) and the Random Direction (RD) models.

3.5 Throughput in Random Waypoint and Random Direction models

In this section, we compute the relay throughput in the case where (i) the relay node moves along a finite interval according either to the RWP or to the RD model, the source and destination nodes being stationary (Section 3.5.1.1), (ii) all nodes move independently of each other, with the same mobility model (RD or RWP), either along a finite interval (Section 3.5.1) or inside a square (Section 3.5.2).

Let now define how the RWP and RD models work respectively: (a) In the RWP model [BMJ⁺98] each node is assigned an initial location in a given area (typically a square) and travels at a constant speed to a destination chosen randomly in this area. The speed is chosen randomly in $(v_{\min}, v_{\max})$, independently of the initial location and destination.

After reaching the destination the node may pause for a random time, after which a new destination and speed are chosen, independently of previous speeds, destinations, and pause times. (b) In the RD model [PNL05] each node is assigned an initial direction, speed and travel time. The node then travels in that direction at the given speed and for the given duration. When the travel time has expired, the node may pause for a random time, after which a new direction, speed and travel time are chosen at random, independently of all previous directions, speeds and travel times. When a node reaches a boundary it is either reflected or the area wraps around so that the node reappears on the other side.

Coming back to our problem, in Theorem 2 we have shown that the relay throughput T_r is given by $T_r = r_s \pi_s$, where r_s is the transmission rate of the source to the relay node (r_s is a given parameter), and π_s is the stationary probability that the source is sending packets to the relay node (see (3.4) in Section 3.3). In the following, we will compute π_s for each case mentioned above. This will be carried out under the assumption that all nodes have the same transmission range R .

3.5.1 One dimension

For the RWP mobility model over the interval $[0, L]$, the stationary probability density function of a node location is [BHPC04]

$$f(x) = \frac{6(L-x)x}{L^3}, \quad x \in [0, L]. \quad (3.24)$$

The stationary probability density function under the RD mobility model is uniform [PNL05], i.e.,

$$f(x) = \frac{1}{L}, \quad x \in [0, L]. \quad (3.25)$$

3.5.1.1 Only relay node is mobile

Assume that the source and the destination nodes are fixed in $[0, L]$, and that the relay node moves along this interval according to either the RD or the RWP mobility model.

From Remark 1 the stability condition is given by $r_s \pi_s < r_d \pi_d$, where these quantities are defined in Section 3.3. Let compute π_s and π_d for either mobility model (recall that r_s and r_d are given parameters). Hence

$$\pi_s = \int_{x=(s-R)^+}^{(s+R) \wedge L} f(x) dx, \text{ and } \pi_d = \int_{x=(d-R)^+}^{(d+R) \wedge L} f(x) dx, \quad (3.26)$$

where $f(\cdot)$ is the stationary node location distribution, and $a \wedge b = \min(a, b)$. Thus, the stability condition reads

$$r_s \int_{x=(s-R)^+}^{(s+R) \wedge L} f(x) dx < r_d \int_{x=(d-R)^+}^{(d+R) \wedge L} f(x) dx. \quad (3.27)$$

Consider now the relay throughput. In the stable case it is given by (see Theorem 2)

$$r_s \int_{x=(s-R)^+}^{(s+R)\wedge L} f(x) dx. \quad (3.28)$$

In the particular case where the relay node moves according to the RD mobility model, the stability condition is (use (3.27)) with $f(x)$ given in (3.25))

$$r_s[(s+R) \wedge L - (s-R)^+] < r_d[(d+R) \wedge L - (d-R)^+], \quad (3.29)$$

and the relay throughput, T_{RD}^f , achieved is (cf. (3.25) and (3.28))

$$T_{RD}^f = r_s \frac{((s+R) \wedge L) - (s-R)^+}{L}. \quad (3.30)$$

For $R < s < L - R$ and $R < d < L - R$, $T_{RD}^f = r_s \frac{2R}{L}$ and $\pi_s = \pi_d$, regardless of the position of the source node and of the destination node. In this case, the stability condition reduces to $r_s < r_d$.

When the relay moves according to the RWP mobility model, then the stability condition is (use (3.27) with $f(x)$ given in (3.24))

$$r_s[2(A-B)(3-(A^2+AB+B^2))] < r_d[2(C-D)(3-(C^2+CD+D^2))],$$

with

$$A := (s+R) \wedge L, \quad C := (d+R) \wedge L, \quad B := (s-R)^+, \quad D := (d-R)^+,$$

and the relay throughput, T_{RW}^f , is given by

$$T_{RW}^f = r_s \frac{2(A-B)(3-(A^2+AB+B^2))}{L^3}. \quad (3.31)$$

3.5.1.2 All nodes are mobile

In this section, the source, destination and relay nodes are all mobile, and move along $[0, L]$ according to the same mobility model: the RWP or the RD model. First, observe from Remark 2, that in this case the stability condition is given by $r_s < r_d$.

Let us now compute the throughput $T_r = r_s \pi_s$ for each mobility model. We have

$$\begin{aligned} \pi_s = & \int_0^L f(x) \int_{(x-R)^+}^{(x+R)\wedge L} f(y) dy dx - \int_0^L f(x) \left[\int_{(x-R)^+}^{(x+R)\wedge L} f(y) dy \right]^2 dx \\ & - \int_0^L f(x) \int_{(x-R)^+}^{(x+R)\wedge L} f(y) \int_{(y-R)^+}^{(y+R)\wedge L} f(z) dz dy dx, \end{aligned} \quad (3.32)$$

where $f(x)$ is given either by (3.24) or by (3.25), depending on the mobility model in use. Note that the second term of the r.h.s of (3.32) is the probability that the relay node and the destination are in contact with the source, and the third term is the probability that the source and destination form a two hop routes via the relay node. It is easy to compute π_s in explicit form for both functions $f(x)$. We will instead provide compact approximation formulas, since the exact ones are lengthy. Approximate π_s by the first term in the r.h.s. of (3.32). Note that this term is the probability that two nodes are neighbors. This approximation is justified by the fact that, for each function $f(x)$ in (3.24) and (3.25), the second and the third term in the r.h.s. of (3.32) are much smaller than the first term, when the ratio $\rho := R/L$ is small with respect to 1. This yields the following approximate throughputs:

$$T_{RW} \approx r_s \frac{\rho(12 - 20\rho^2 + 15\rho^3 - 2\rho^5)}{5} \quad (3.33)$$

for the RWP mobility model, and

$$T_{RD} \approx r_s \rho(2 - \rho), \quad (3.34)$$

for the RD mobility model. Observe that these formulas depend on R and L only through their ratio, and that $T_{RD} < T_{RW}$ for all $\rho < 1$.

We conclude this section by considering the case where $r_s = r_d := r$. In this case, the relay buffer is not stable. To handle this situation, it is proposed in [GT02] to use a probability of relaying, p_r , which is close to 1, so that when the relay node enters the neighborhood of the source node, the source node transmits data to be relayed with probability $p_r < 1$, and does not transmit with the complementary probability. Note that this scheme ensures stability and gives near maximum throughput as well, since p_r is close to 1.

3.5.2 Two dimensions

In this section nodes are moving in a square. This section starts by computing the relay throughput; then it finds the probability that these nodes form a two-hop route.

3.5.2.1 Three nodes moving

Nodes move independently of each other inside a square of side length L . They move according to the same mobility model, the RD or the RWP model.

Similarly to Section 3.5.1.2 the stability condition is $r_s < r_d$. The probability that two nodes are neighbors is

$$\pi(f) = \int_{x \in [0, L]^2} f(x) \int_{x_1: |x - x_1| \leq R} f(x_1) dx_1 dx \approx \pi R^2 \int_{x \in [0, L]^2} f^2(x) dx, \quad (3.35)$$

where the continuity of $f(\cdot)$ is used, and the assumption that R is negligible with respect to L . This approximation is in agreement with Theorem 3, which states that the minimum

probability of contact is achieved by the uniform distribution, since the latter integral is the L_2 -norm of $f(\cdot)$, which is minimized when $f(\cdot)$ is the uniform distribution.

When the RB is stable, the relay throughput is approximated by $r_s \pi(f)$, where r_s is the source transmission range. It can be shown, numerically, that for the RWP mobility model over a square $\pi(f) \approx 1.36\pi(R/L)^2$ [Gro05], this implies that the relay throughput, T_{RW}^{2d} , is approximated by

$$T_{RW}^{2d} \approx 1.36\pi r_s \rho^2, \quad (3.36)$$

where $\rho = R/L$. Under the RD mobility we find that $\pi(f) \approx \pi\rho^2$. Hence, the relay throughput, T_{RD}^{2d} , is approximated by

$$T_{RD}^{2d} \approx \pi r_s \rho^2. \quad (3.37)$$

Note that $T_{RD}^{2d} < T_{RW}^{2d}$.

3.5.2.2 Probability of a two-hop route

The throughput of data between a pair of nodes when there exists a route between them, called *forwarding throughput*, is a function of the route length in hops. A first step to derive the forwarding throughput is to compute the distribution of the route length in hops. This section computes the probability of a two-hop route between two nodes s and d , assuming there are N other nodes. All the nodes are mobile inside a square $A = [0, L]^2$ and moving (independently) according to RWP or to RD model.

Let x_s (resp. x_d) represents the position of node s (resp. d). There is a two-hop route between node s and d , if the node d is inside the annulus, C , of center x_s and of interior and exterior radius equal to R and $2R$, and there is at least an intermediate node inside, D_I , the intersection of the two disks of radius R centered around x_s and x_d , see Figure 3.3. Hence, the probability of the two-hop route between nodes s and d when $R \ll L$ reads

$$\begin{aligned} P^N(2) &\approx \int_A f(x_s) \int_C f(x_d) \left[1 - \overline{\left(\int_{D_I} f(x) dx \right)^N} \right] dx_d dx_s \\ &\approx 2\pi \sum_{i=1}^N (-1)^{i+1} \binom{N}{i} u(i) v_f(i) \left(\frac{R}{L} \right)^{2(i+1)}, \end{aligned} \quad (3.38)$$

where $u(i) = \int_1^2 r A(r)^i dr$, and $v_f(i) = \int_{[0,1]^2} f(x)^{i+2} dx$. Here $A(r)$ is the area of D_I when $R = 1$ and the distance between nodes s and d is equal to r . This gives

$$A(r) = 2 \arccos\left(\frac{r}{2}\right) - \frac{r}{2} \sqrt{4 - r^2}. \quad (3.39)$$

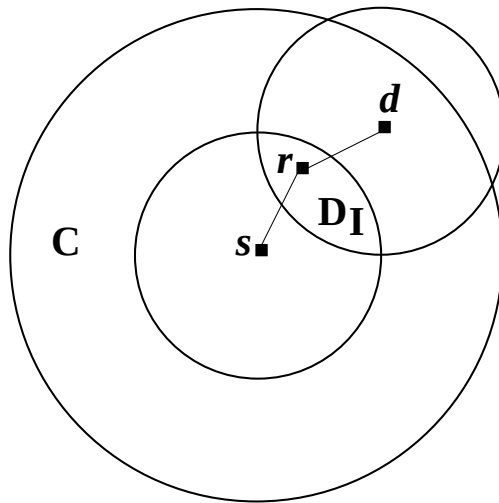


Figure 3.3: Two-hop route between two nodes s and d .

Observe that $P^N(2)$ is a function of $(R/L)^2$. Note that $u(i)$ is independent of the mobility model of the nodes, and that in the case RD mobility $v_f(i) = 1$ for $i \geq 1$. The numerical values of $u(i)$ and $v_f(i)$ for some typical values of i will be listed in Table 3.1 in Section 4.4.

Until now we have looked at the effect of mobility patterns on the “average” throughput and showed that throughput depends only on the stationary node distribution. We then showed that minimum throughput is achieved when the stationary distribution of node position is uniform. In the next section, we will show the numerical results of the queueing model and the relay throughput approximations.

3.6 Numerical results

In this section we present numerical results to validate results in Theorem 1 (stability issues), Theorem 2 (throughput depends only on stationary distribution), Section 3.5 (throughputs obtained by RWP and RD Models), the probability of two-hop route. Numerical results of RWP and RD models are obtained using the NS-2 code of the random trip model [BV05]. Throughout this section, we will assume that the transmission range of the nodes is constant and is equal to R .

3.6.1 Validation of Theorem 1

Consider the scenario of three nodes: a source, a destination, and a relay node. Nodes are moving according to a symmetric Random Walk over a circle. It follows from Theorem 1 that the relay node buffer occupancy is stable *iff* the ratio $p = \frac{r_s}{r_d} < 1$. Figure 3.4 plots the evolution of relay node buffer with time for different values of p .

It is evident from the figure that when $p = 1.0$, the buffer occupancy process is unstable. While for the case $p = 0.9 < 1.0$, this process is stable. Similar results were

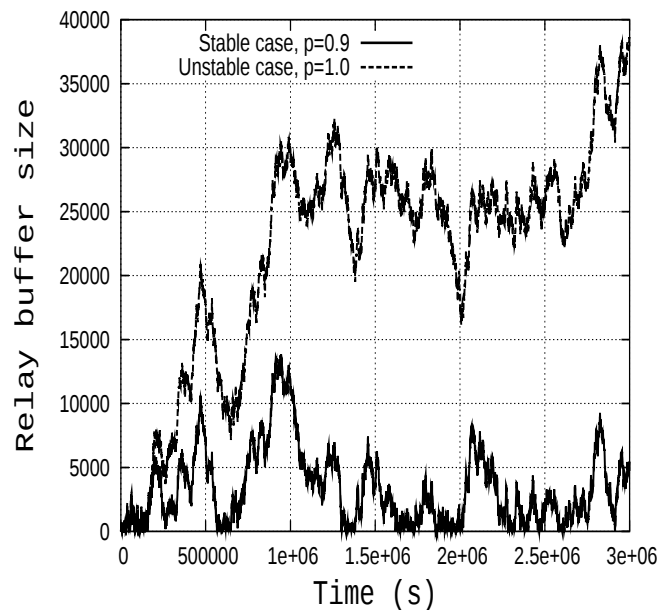


Figure 3.4: Time-evolution of relay node buffer for Random Walk model over a circle for different values of ratio, $p = \frac{r_s}{r_d}$.

obtained even for $p \approx 1.0$ with $p < 1.0$ but are not shown here.

3.6.2 Validation of Theorem 2 and Section 3.5

Theorem 2 states that the relay throughput depends only on the stationary node distribution. Section 3.5 provides the value of relay throughput under RD and RWP mobility models. To validate both of these results, we ran simulations to find the relay throughput for case of the RD and the RWP mobility models with different parameters.

To illustrate that the throughput depends only on the stationary distribution of the node position we look at the scenario where three nodes move over a line of length $L = 4$ kilometers according to the RD model. Assume that the time between two consecutive decision instants (travel time) is fixed and equal to 15 seconds and the distribution of speed was chosen to be i) Uniform over some interval, ii) Exponentially distributed, and iii) fixed. Note that the case where speed is fixed corresponds to the Random Walk. Since the stationary node location distribution is same for all the three choice of speed distribution, Theorem 2 implies that the relay throughput will be identical. The numerical results plotted in Figure 3.5 are in accordance with this result. Also evident is the fact that relay throughput $T_r = \pi_s$, for $r_s = 1$.

In Figure 3.6 we plot values of π_s when the three nodes move according to the RWP model over a square of side-lengths L for different values of transmission range, R . We keep the speed of the mobiles fixed. The plot shows that π_s (and hence T_r) is a function of $\rho = \frac{R}{L}$ alone. The numerical values also support the result of Section 3.5.2 where for the RWP model in square, the throughput is approximately $1.36\pi\rho^2$.

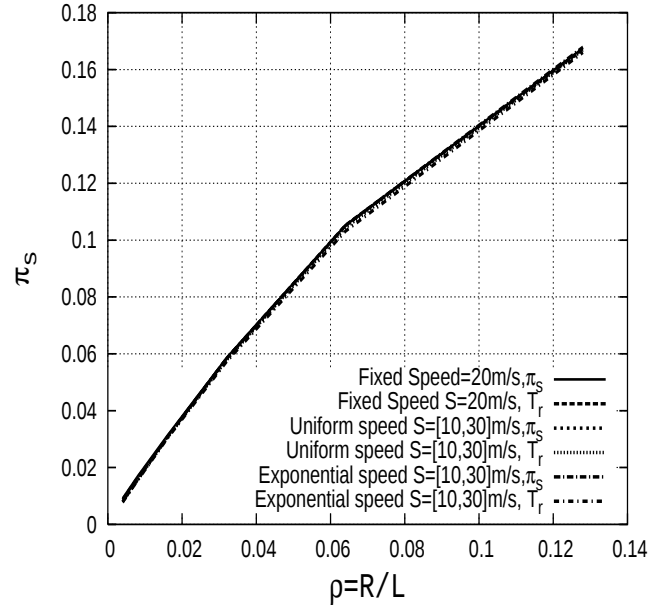


Figure 3.5: Values of π_s obtained for the RD model over a segment of length L . Various distributions for the speed were taken.

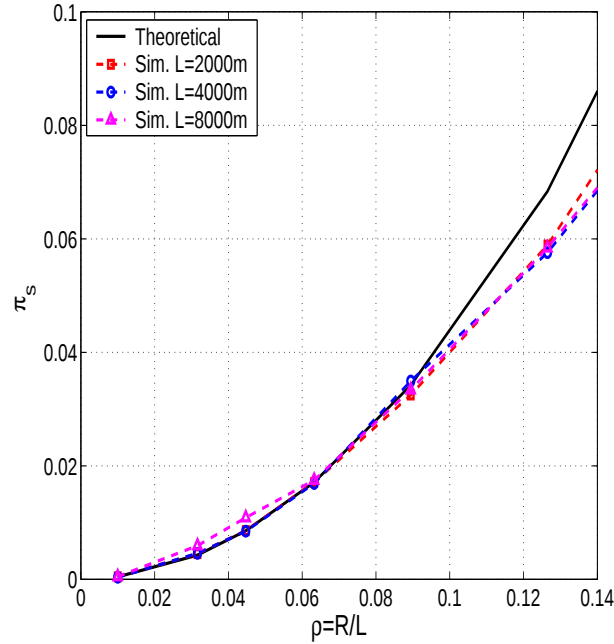


Figure 3.6: Simulation results showing π_s for RWP model over a square of side-lengths L for different values of transmission range, R . Also shown are corresponding values from Section 3.5.2.

Similarly, the values of the throughput from theory and simulations provide a good match for all of the scenario studied in Section 3.5. Because of the space restriction, we did not include these numerical results due to space constraints.

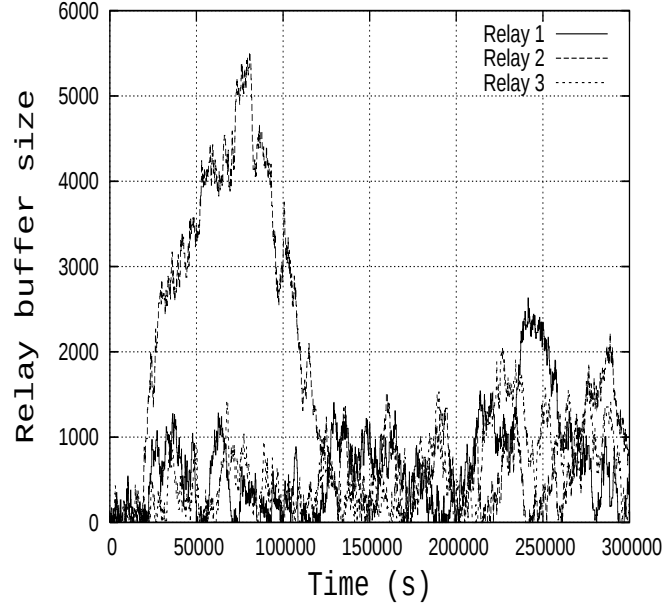


Figure 3.7: Time-evolution of relay buffer for three relay nodes moving according to RWP model inside a square of side-length $4000m$, with fixed and symmetric source and destination node w.r.t. to square center. Here $r_s = 0.9r_d$ (the stable case).

3.6.3 Validation of Section 3.3.2

Section 3.3.2 studies multiple relay nodes with fixed source nodes and destination nodes. It reports stability condition and derives the value of the relay throughput and the network throughput. To validate the stability condition, we take the scenario of one source node, one destination node, and 3 relay nodes move according to the RWP model inside a square of side-length $L = 4000m$. The source and destination nodes are fixed and they are symmetric according to the center of the square, and the separated distance between them is of $2000m$. The stable case is shown in Figure 3.7 where $r_s = 0.9r_d$, and the unstable case is shown in Figure 3.8 where $r_s = 1.1r_d$. The relay throughput and the network throughput as a function of the number of the relay nodes are shown in, respectively, Figures 3.9 and 3.10 for different value of R/L . In this scenario, the probability that the relay node is neighbor of the source of location $(1000, 1000)$ is equal to 0.0485, 0.0275, and 0.0126 for R/L equal to 0.1, 0.075, and 0.05 respectively.

We validate the approximation of Section 3.3.2 of the case where all source and destination nodes are moving (c.f, Equation 3.16). We consider a scenario of N relay nodes and K source nodes and M destination nodes. All the nodes move according to the RWP model within a square region of side-length L equal to 4000. We assume that the lifetime of the packets in hops is equal to 4. In Figure 3.11, we show the approximation of Equation 3.16 as well as the simulation result of, τ_g^m , the long-term arrival rate to the queue m of the relay node n from the source nodes. We observe that the above approximation is accurate for $R/L \leq 0.05$, and the relative error between the approximation and simulation is less than 5%, for $R/L = 0.05$.

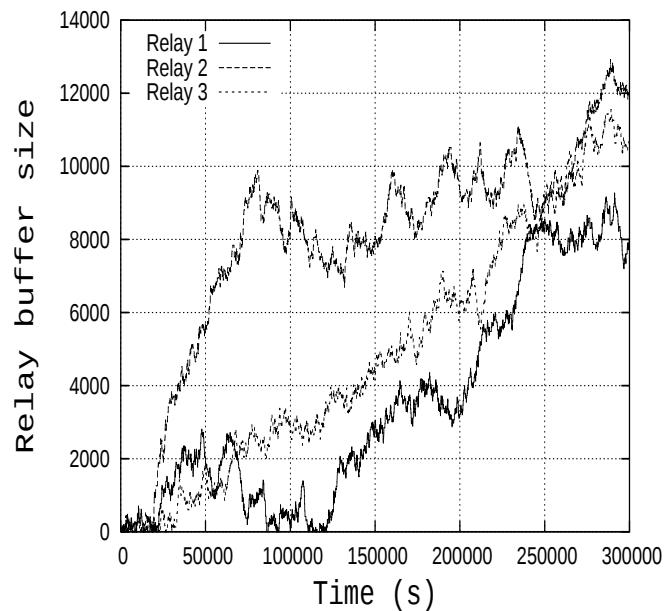


Figure 3.8: Time-evolution of relay buffer for relay nodes moving according to RWP model inside a square of side-length $4000m$, with fixed and symmetric source and destination nodes w.r.t. to square center. Here $r_s = 1.1r_d$ (the unstable case).

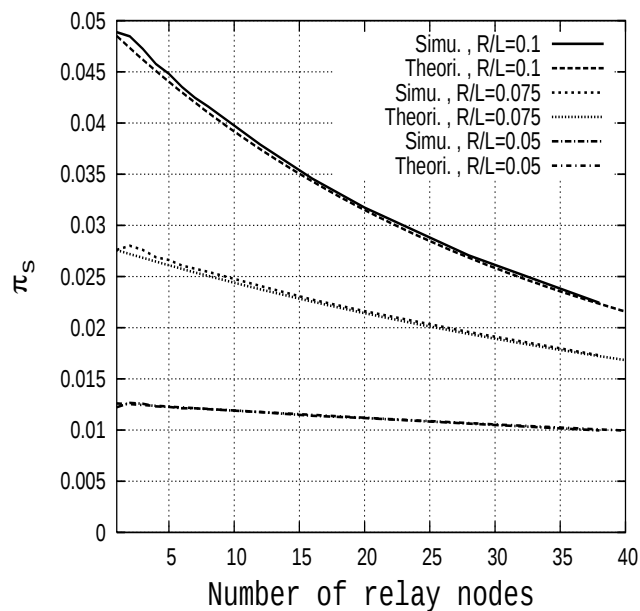


Figure 3.9: Relay throughput as function of number of relay nodes in a square region of side-length $4000m$ with fixed source and fixed destination nodes.

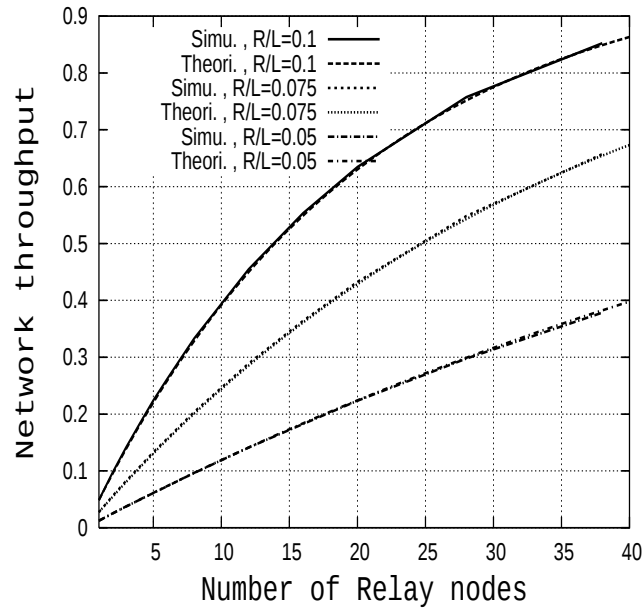


Figure 3.10: Network throughput as function of number of nodes in square region of side-length $4000m$ with fixed source and fixed destination nodes.

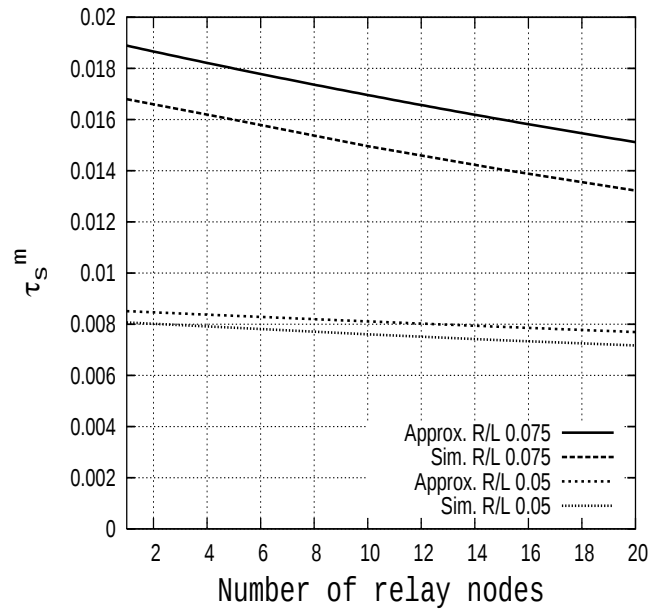


Figure 3.11: Long-term arrival rate at queue m of relay node n from source node k as function of number of relay nodes N , where $K = 4$, $M = 5$, $P_k^m = 1/M$, and $r_{S_k} = 1$.

3.6.4 Validation of two-hop route probability

We have $N + 2$ nodes moving inside a square of side length $L = 4000$ according to the RWP model. We validate the approximation formula for the probability of a two hops route for different values of N (cf., section 3.5.2.2). In Table 3.1, we list the numerical values of $u(i)$ as function of i , and of $v_{fw}(i)$, the values of $v_f(i)$ for the RWP model, for details see Equation 3.38. In Figure 3.12, we show the results of the simulation and the approximation for $R/L \in \{0.025, 0.0375, 0.05\}$. We observe that for $R/L \leq 0.05$ the approximation is accurate.

i	$u(i)$	$v_{fw}(i)$	i	$u(i)$	$v_{fw}(i)$
1	0.649	2.156	8	0.4513	170.9
2	0.47	3.664	9	0.494	341.7
3	0.406	6.541	10	0.547	688.8
4	0.384	12.08	20	2.159	$9.93 \cdot 10^5$
5	0.383	22.87	40	66.23	$3.68 \cdot 10^{12}$
6	0.395	44.1	60	2703.08	$3.80 \cdot 10^{18}$
7	0.418	86.3	120	$3.09 \cdot 10^8$	$7.09 \cdot 10^{38}$

Table 3.1: The numerical values of the $u(i)$ and $v_{fw}(i)$.

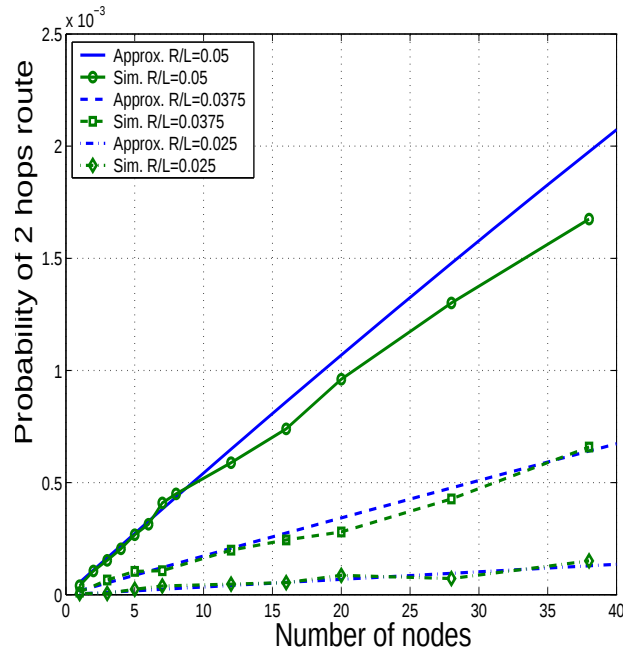


Figure 3.12: Probability of two-hop route as function of number of nodes. All nodes move according to RWP model inside square of side length $4000m$.

3.7 Conclusions

This chapter studied the performance of relaying in mobile ad hoc networks by developing a queueing model. The parameters of the queueing model depend on the node mobility pattern.

Our main findings are that (under the assumptions placed on our model) the relay throughput only depends on the stationary node location distribution, and that uniform stationary distribution of nodes results in the smallest relay throughput. Approximate throughput formulas have been derived for both the RWP and the RD mobility models; these formulas have been found to be in agreement with simulation results.

The next chapter will study the behavior of the relay buffer function of the mobility of the nodes.

Chapter 4

Impact of Mobility on the Relay Buffer Behavior

Considered is a mobile ad hoc network consisting of three types of nodes: source, destination, and relay nodes. This chapter studies the behavior of the relay buffer as a function of the node mobility. We show that the expected relay buffer size depends on both the expected and the variance of the nodes contact time. Random Walk and Random Direction mobility models in both one and in two dimensions are studied and approximate simple expressions of the expected relay buffer size are provided.

Note: Part of the results in this chapter have appeared in [HKG⁺06, HKG⁺05].

4.1 Introduction

In this chapter, we consider the two-hop relay protocol [GT02]. The network consists of three type of nodes, source, destination, and relay nodes. The objective is to study the behavior of the relay buffer (RB) as a function of the nodes mobility models. To this end, we find the expected RB size, in the heavy traffic case, embedded at certain instants of time. This expectation is called the event average. Note that the expected RB size in the heavy traffic case serves also as an upper bound of the expected RB size [Kle76]. Further, we show numerically that under the mobility models considered the event average converges toward the time average of the RB as the load of the relay buffer tends to one (heavy traffic case). This will be done for three different mobility models: Random Walk [Fel68], Random Direction (RD) [PNL05], and Random Waypoint (RWP) [BHPC04].

The point of departure is the observation made in Chapter 3 Section 3.3 which relates the evolution of a relay node buffer to the evolution of the workload process in a $G/G/1$ queueing system. The expected waiting time of this $G/G/1$ queue will be approximated by the corresponding quantity in the $GI/G/1$, i.e. the inter-arrival times are independent of the service requirements.

The rest of this chapter is organized as follows: Section 4.2 summarizes the system model that is based on Chapter 3 Section 3.3. In section 4.3 we study the RB behavior for the Random Walk and Random Direction models. Section 4.4 reports the numerical results on the contact time distribution, the expected RB size, and the convergence of the time average of the RB size. Section 4.5 concludes this chapter.

4.2 The system model

As mentioned earlier, this chapter focuses in the relay buffer (RB) behavior and shows the dependency of the relay node buffer behavior on the mobility model. We consider the two-hop relay protocol introduced in Chapter 3 Section 3.3. The following is a summary of the model:

1. There are three nodes consisting of one source, one destination, and one relay node.
2. The source node has always data to send to the destination node.
3. When the relay node comes within the transmission range of the source node, and if the destination node is outside the transmission range of the source and of the relay node, then the relay node accrues packets to be relayed to the destination node at a constant rate r_s .
4. When the destination node comes within the transmission range of the relay node, and if the destination and the relay node are outside transmission range of the source node, then the relay node sends the relay packets (if any) to the destination node at a constant rate r_d .

The state of the relay node at time t is represented by the rv $S_t \in \{-1, 0, 1\}$ where:

- $S_t = 1$ if at time t the relay node is in contact with the source, and if the destination is in contact neither of the source nor of the relay node;
- $S_t = -1$ if at time t the relay node is in contact of the destination, and if the source is in contact neither of the destination nor of the relay node;
- $S_t = 0$ otherwise.

Assume that the relay node may only enter state 1 (resp. -1) from state 0: if $S_{t-} \neq S_t$ then necessarily $S_t = 0$ if $S_{t-} = 1$ or $S_{t-} = -1$.

Let B_t be the RB occupancy at time t . Based on the definition of S_t , B_t increases at rate r_s if $S_t = 1$, B_t decreases at rate r_d if $S_t = -1$ and if the RB is non-empty, and B_t remains unchanged in all other cases.

Let $\{Z_n\}_n$ ($Z_1 < Z_2 < \dots$) denote the consecutive jump times of the process $\{S_t, t \geq 0\}$. Define a *cycle* as the interval of time that starts at $t = Z_k$, for some k with $S_t = 1$, and (necessarily) $S_{t-} = 0$ and $S_{Z_{k-2}} = -1$, and ends at the smallest time $t + \tau$ such that $S_{t+\tau} = 1$ and $S_{t+s} = -1$ for some $s < \tau$. Let C_n denotes the duration of the n^{th} cycle. Let W_n denote time at which the n^{th} cycle begins. Let

$$\sigma_n \triangleq \int_{t=W_n}^{W_{n+1}} \mathbf{1}_{\{S_t=1\}} dt \quad (4.1)$$

be the amount of time spent by the relay node in state 1 during the n^{th} cycle. Similarly, let

$$\alpha_n \triangleq \int_{t=W_n}^{W_{n+1}} \mathbf{1}_{\{S_t=-1\}} dt \quad (4.2)$$

be the amount of time spent by the relay node in state -1 during the n^{th} cycle. Let $\tilde{B}_n$ be the RB occupancy at the beginning of the n^{th} cycle, i.e. $\tilde{B}_n = B_{W_n}$. Clearly,

$$\tilde{B}_{n+1} = [\tilde{B}_n + r_s \sigma_n - r_d \alpha_n]^+ \quad (4.3)$$

where $[x]^+ = \max(x, 0)$. In other words, $\tilde{B}_{n+1}$ can be interpreted as the workload seen by the $(n+1)^{\text{st}}$ arrival in a $G/G/1$ queue, where $r_d \alpha_n$ is the inter-arrival time between the n^{th} and the $(n+1)^{\text{st}}$ customer, and $r_s \sigma_n$ is the service requirement of the n^{th} customer. According to Chapter 3 Theorem 1, the $G/G/1$ is stable iff $r_s E[\sigma_n] < r_d E[\alpha_n]$ and then $\tilde{B}_n$ converges in probability to a proper and finite rv $\tilde{B}$. The distribution of $\tilde{B}$ in the $G/G/1$ queue will be approximated by the corresponding quantity in a $GI/G/1$, i.e., by assuming that α_n and σ_n are independent. In the $GI/G/1$ queue, it is known that $E[\tilde{B}]$ verifies the following inequality called the Kingman upper bound [Kle76, P. 29]

$$E[\tilde{B}] \leq \frac{r_d^2 \text{Var}(\alpha_n) + r_s^2 \text{Var}(\sigma_n)}{2(r_d E[\alpha_n] - r_s E[\sigma_n])}, \quad (4.4)$$

and furthermore in the *heavy traffic* case, i.e., when $r_s E[\sigma_n] \approx r_d E[\alpha_n]$ with $r_s E[\sigma_n] < r_d E[\alpha_n]$, the stationary waiting time is exponentially distributed with mean [Kle76, P. 29]

$$E[\tilde{B}] \approx \frac{r_d^2 \text{Var}(\alpha_n) + r_s^2 \text{Var}(\sigma_n)}{2(r_d E[\alpha_n] - r_s E[\sigma_n])}. \quad (4.5)$$

This approximation will be used in the following to estimate the expected RB size.

4.3 Relay buffer behavior

In this section, the impact of the mobility model on the relay buffer occupancy is studied. Assume that the mobility models under consideration have stationary node location distributions. The plan is to view this system as a $GI/G/1$ queue in heavy traffic and then to look at the effect of mobility patterns on the relay buffer occupancy. It is known from heavy-traffic analysis of a $GI/G/1$ queue [Kle76] that the tail behavior (the large deviation exponent) of the buffer occupancy is determined by the variance of the service and inter-arrival times.

Moreover, it is also to be understood that the effective arrival process to the RB in the queueing model (introduced in Section 4.2) is not the contact time between the relay and source nodes, i.e.,

$$\int_{u=Z_n}^{Z_{n+1}} \mathbf{1}_{\{S(u)=1\}} du, \quad (4.6)$$

but is composed of many (random number of) such contact times since

$$\sigma_n = \int_{u=W_n}^{W_{n+1}} \mathbf{1}_{\{S(u)=1\}} du. \quad (4.7)$$

Clearly, a larger relay buffer occupancy would imply that the amount of time required to deliver all the packets would be composed of many contact periods between the relay node and the destination, hence there can be several inter-visits between the relay node and the destination required to deliver the packets. This implies that we can not study the delay incurred by the nodes by considering only one inter-visit time (or the meeting time) or only one contact time. This shows that the buffer behavior (hence the delays) will depend on *both* contact times and the inter-visit times.

This section is devoted to such a study for the following cases where: first the relay node performs a Random walk and the source and destination are fixed, second the source, the destination, and the relay node move inside a square according to the RD model. This section is meant for illustration of the above phenomenon.

requirement of a customer in the $G/G/1$ queue of Section 4.3 is $r_s\sigma$, where

$$\sigma = \sum_{j=1}^L A_j. \quad (4.9)$$

In the following, A denotes a generic rv with the same distribution as A_j .

Using results from random walk literature (for example [Fel68, Chap. 14, Sec. 4]), we find that

$$E[A] = 2R\frac{\mu}{V}, \quad (4.10)$$

$$Var[A] = \left(\frac{\mu}{V}\right)^2 \cdot \frac{4R}{3}(2R+1)(R+1), \quad (4.11)$$

$$p = 1 - \frac{1}{w}, \quad (4.12)$$

$$E[L] = w, \quad (4.13)$$

$$Var[L] = w(w-1). \quad (4.14)$$

Since L is independent of A , we get

$$E[\sigma] = E[A]E[L] = 2wR\frac{\mu}{V}, \quad (4.15)$$

$$\begin{aligned} Var[\sigma] &= Var[A]E[L] + (E[A])^2Var[L], \\ &= 4wR\left(\frac{\mu}{V}\right)^2\left(wR + \frac{1}{3}(2R^2 + 1)\right). \end{aligned} \quad (4.16)$$

Hence in the scenario under study and when $r_s \approx r_d$ with $r_s < r_d$ the mean RB occupancy in (4.5) reads

$$E[\tilde{B}] \approx \frac{r_d^2 Var(\alpha_n) + r_s^2 Var(\sigma_n)}{2(r_d E[\alpha_n] - r_s E[\sigma_n])} = \frac{(r_s^2 + r_d^2) Var(\sigma)}{2E[\sigma](r_d - r_s)} = \frac{\mu(r_s^2 + r_d^2)}{V(r_d - r_s)}\left(Rw + \frac{1}{3}(2R^2 + 1)\right), \quad (4.17)$$

where we have used the fact that α_n and σ_n are identically distributed.

4.3.2 Two-dimensional RD model

We consider the following scenario of three nodes: the source, the destination, and the relay node. All these nodes are moving inside a square of side-length L . They move according to the RD model with wraparound and with constant speed V , and constant travel time T [PNL05].

Based on the queueing model of Section 4.2, the RB accumulates data at rate r_s when the process $S_t = 1$. This means that the relay node is inside the transmission range of the source, in *contact* with the source, and the destination is outside the transmission range of the source and of the relay node. In one hand, note that when $R \ll L$ the probability that

the source and the relay node are in contact is of order $(R/L)^2$, see Section 3.5.2.1. On the other hand, the probability that the relay node and the destination are in contact with the source is of order $(R/L)^4$, the same order as the probability that the three nodes form a two hop route, see Section 3.5.2.2. This observation motivates us to approximate the duration of time where $S_t = 1$ by the duration of time where the source and the relay node are in contact regardless of the position of the destination.

Let $A_j, j \geq 1$, be independent and identically distributed random variables representing the *contact times* between the source and the relay node. Recall that in Section 4.2 a cycle contains at least one contact time between source and relay node and one contact time between relay node and destination. Thus let L_d denotes the rv that represents the number of contact times between the source and relay node inside a cycle. The service requirement in the queueing model of Section 4.2 is $r_s \sigma$, where

$$\sigma = \sum_{j=1}^{L_d} A_j. \quad (4.18)$$

As we saw in Section 4.3.1, the mean RB size at cycle instants in heavy-traffic depends on the first two moments of L_d and A_j . Thus in the rest of this section, first we will derive the probability distribution of L_d , and next we will derive the probability distribution of A_j .

For $i \geq 1$, let I_i denote the inter-meeting time between the source and the relay node, i.e. the time passes between two consecutive of the two nodes contact. Let F_r denotes the residual time, at a cycle starts, that the relay node and the destination that are not in contact to become in contact. We will assume that I_i , A_i , and F_r are mutually independent. For $k \geq 1$, let $X_k = \sum_{i=1}^k (A_i + I_i)$ and $X_0 \triangleq 0$, so the probability that $L_d \geq k$ given that $L_d \geq 1$ reads

$$P(L_d \geq k) = \frac{P(X_{k-1} + A_k \leq F_r)}{P(A_1 \leq F_r)}. \quad (4.19)$$

For $R \ll L$, the inter-meeting times between any pair of nodes, I_i , are iid rv of exponential distribution with parameter $\lambda = \frac{8RV}{\pi L^2}$ [GNK05]. Then, basing on the memoryless property of the exponential distribution, it follows that F_r , the residual time, is exponentially distributed with the same parameter λ . Thus, (4.19) gives

$$\begin{aligned} P(L_d \geq k) &= 2^{1-k} \frac{\sum_{i=0}^{+\infty} \left(\frac{(-\lambda)^i}{i!} E[(A_1 + \dots + A_k)^i] \right)}{\sum_{i=0}^{+\infty} \frac{(-\lambda)^i}{i!} E[A_1^i]} \\ &\approx \beta^{k-1}, \end{aligned} \quad (4.20)$$

where $\beta = \frac{1}{2} \exp(-\lambda E[A])$. The above approximation is based on the fact that for $R \ll L$ the variation of A_j is small comparing with X_k and F_r . This is equivalent to approximate the distribution of L_d by the geometrical distribution of parameter β . In the rest of this section, we will derive the probability distribution of A_j .

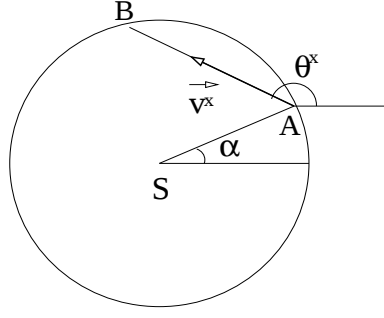


Figure 4.2: Contact distance between 2 mobile nodes.

Let A denotes the rv with the same distribution of A_j . Assume that during the contact of two nodes there is no change of node direction. This is valid when T , the node travel time, is sufficiently large comparing with the contact time between these nodes. By conditioning on the direction of the source, θ_s , and of the relay node, θ_r , we derive the relative speed $v^* = |\vec{V}(\theta_s) - \vec{V}(\theta_r)|$ and the relative angle $\theta^* = \theta_s - \theta_r$ between these two nodes. Given that the relay node has crossed the transmission range of the source, the crossing distance, AB , will depend on the crossing angle α , see Figure 4.2. Conditioning on α , A is computed by dividing AB by the relative speed v^* . We note that in the case of the RD model that θ_s , θ_r , and α are uniform distributed in $[0, 2\pi]$. The CDF of A is computed by unconditioning on θ_s , θ_r , and α , over the constraint that $|\theta_s - \theta_r| > \delta$, where δ is a threshold to avoid very large contact time. Let $\theta_1 = \frac{\theta_s + \theta_r}{2}$, and $\theta_2 = \frac{\theta_s - \theta_r}{2}$, thus the CDF of the normalized contact time $\tilde{A} = A \frac{V}{R}$ is the following

$$\begin{aligned}
 P(\tilde{A} \leq x) = & \frac{1}{C} \left[\int_{\delta}^{2\pi} \int_0^{\theta_s - \delta} \int_0^{\frac{\pi}{2}} d\alpha d\theta_s d\theta_r + \int_0^{2\pi - \delta} \int_{\theta_s + \delta}^{2\pi} \int_0^{\frac{\pi}{2}} d\alpha d\theta_s d\theta_r + \right. \\
 & \left. \int_{\delta}^{2\pi} \int_0^{\theta_s - \delta} \int_{\frac{\pi}{2}}^{2\pi} d\alpha d\theta_s d\theta_r + \int_0^{2\pi - \delta} \int_{\theta_s + \delta}^{2\pi} \int_{\frac{\pi}{2}}^{2\pi} d\alpha d\theta_s d\theta_r \right] \quad (4.21)
 \end{aligned}$$

$\theta_1 - \pi \leq \alpha \leq \theta_1, \sin(\theta_1 - \alpha) \leq x \cdot \sin(\theta_2)$ $0 \leq \alpha \leq \theta_1 - \pi, \sin(\theta_1 - \alpha) \geq x \cdot \sin(\theta_2)$
 $\theta_1 + \pi \leq \alpha \leq 2\pi, \sin(\theta_1 - \alpha) \leq x \cdot \sin(\theta_2)$ $\theta_1 \leq \alpha \leq \theta_1 + \pi, \sin(\theta_1 - \alpha) \geq x \cdot \sin(\theta_2)$

where C is a normalization constant. We observe that the CDF of the normalized contact time is independent of R and V , and that it is easy to compute numerically the moments of $\tilde{A}$. For example for $\delta = 0.1\pi$, we have that $C = 0.3298$, $E[A] = 1.1778 \frac{R}{V}$, and $E[A^2] = 2.3365 \left(\frac{R}{V}\right)^2$.

Now return back to the mean RB occupancy assume that L_d is independent of A , the mean and variance of σ can be derived. So, in heavy-traffic the mean RB size of (4.5) reads

$$E[\tilde{B}] \approx \left(0.403 + 0.5889 \frac{\beta}{1 - \beta} \right) \frac{r_s^2 + r_d^2}{r_d - r_s} \frac{R}{V}, \quad (4.22)$$

for $r_s \approx r_d$ with $r_s < r_d$. We observe that $E[\tilde{B}]$ is function of R , and V only through the ratio R/V .

4.4 Validation and numerical results

In this section we present the numerical results to validate the approximation of expected relay buffer behavior size as suggested in Section 4.3, of the convergence of the mean RB size at cycle time to the time average of the RB size. Numerical results of the Random Direction model are obtained using the NS-2 code of the random trip model [BV05]. Throughout this section, we will assume that the transmission range of the nodes is constant and is equal to R .

4.4.1 Validation of Section 4.3.1

We consider the relay node buffer occupancy in the scenario in Section 4.3. The third line of Table 4.1 reports the percentage of the relative error of the relay buffer occupancy found in (4.17) and the corresponding simulated value $E[B_{sim}]$, for different values of the parameters R and w with $\mu = 50$ meters and $r_d\mu/V = 200$ data units. Parameters R and w are chosen so that the circumference of the circle is equal to 3000 meters (i.e. $(4R + 2w)\mu = 3000$ meters). In the simulation the relay node buffer is sampled at the beginning of each cycle (as defined in Section 4.2). Throughout these experiments $\frac{r_s}{r_d} = 0.95$, so as to reflect the heavy-traffic scenario under which (4.17) was established. We observe that the relative error between Equation (4.17) and the simulation is very small.

R	11	9	7	5
w	8	12	16	20
$\frac{ E[B_{sim}] - E[B] }{E[B]}$ (%)	2	1	1	2

Table 4.1: Validation of Equation (4.17) (for $\mu = 50$ meters and $r_d\mu/V = 200$ data units).

4.4.2 Validation of Section 4.3.2

We consider the scenario where the nodes move according to the RD model inside a square of side length $L = 4000m$. The nodes' speed is constant and is equal to V . The node travel time, T , is constant and is greater than R/V . First, we validate the approximation of, $\tilde{A}$, the normalized contact time, and next the approximation of the mean buffer size $E[\tilde{B}]$ at cycle time. In Figure 4.3, we compare the normalized contact time distribution of the simulations for different values of R/V and T with the approximation of (4.21) for $\delta = 0.1\pi$. We observe that similarly as the model predict that the CDF of $\tilde{A}$ is almost independent of R/V , of R/L , and of T . We note that same observation is also valid in the case of the RWP model, see Figure 4.4. In Table 4.2, we show the relative error between the mean buffer size of (4.22) and the simulation as a function of R and V and for $\frac{r_s}{r_d} = 0.99$ (i.e. the heavy traffic case). We observe that the model is accurate for $R/V < 10$ and $R \ll L$.

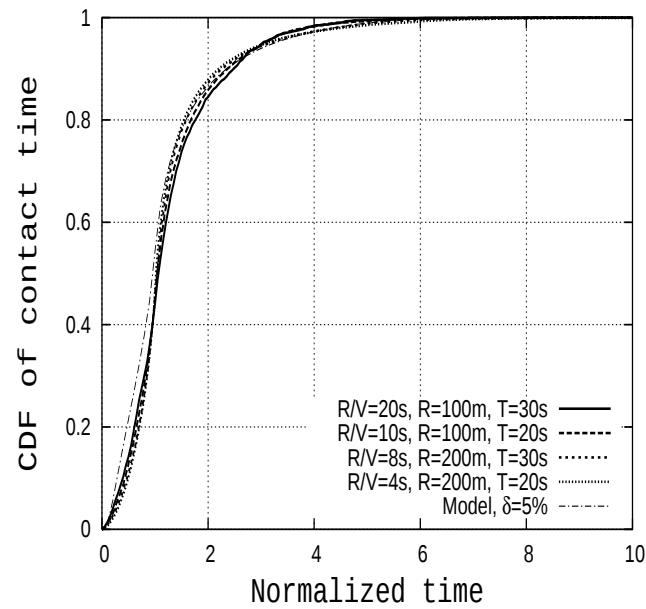


Figure 4.3: Normalized contact time distribution for different values of R/V and travel time T of the simulation and of (4.21) for the RD model.

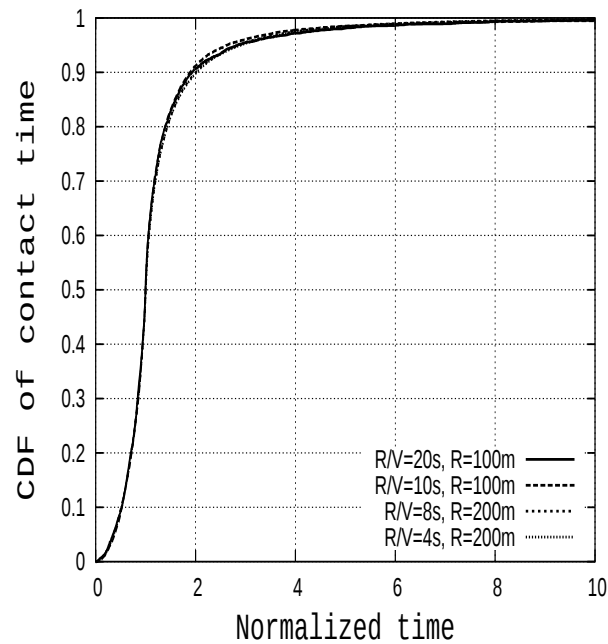


Figure 4.4: Normalized contact time distribution of simulation for Random Waypoint model.

R (m)	75		100	
V (m/s)	7.5	15	10	20
$\frac{ E[B_{sim}] - E[\tilde{B}] }{E[\tilde{B}]}$	0.14	0.03	0.14	0.05

 Table 4.2: Validation of Equation (4.22) (for $r_s/r_d = 0.99$ and $r_d = 10$).

4.4.3 Comparison between $E[\tilde{B}]$ and $E[B]$

In this section we study the relative difference between, $E[B]$, the RB time average and, $E[\tilde{B}]$, the RB average at the cycle start time (event average). We consider the scenario where three nodes move according to the RWP model or to the RD model inside a square of side-length $L = 1000m$. For both cases, the nodes speed is uniformly distributed in the interval $[5m/s, 15m/s]$. Tables 4.3 and 4.4 show the relative difference defined as follows $\frac{E[B] - E[\tilde{B}]}{E[B]}$ as a function of R , the node's transmission range, and the ratio $\frac{r_s}{r_d}$. Note that the time average $E[B]$ is computed by sampling the RB size every 0.1s. The simulation time is considered to be equal to 10^9s . We observe that for different values of the R as the $\frac{r_s}{r_d}$ increases the relative difference decreases and when $\frac{r_s}{r_d} = 0.99$ the relative difference is almost negligible. This means that the event average $E[\tilde{B}]$ is converging to the time average $E[B]$ and these two quantities are almost equal in heavy-traffic. Note that this observation is true for the two mobility model considered.

R (m)	r_s/r_d			
	0.25	0.5	0.75	0.99
50	0.75	0.41	0.17	0.006
100	0.75	0.42	0.18	0.004
200	0.73	0.42	0.18	0.003

 Table 4.3: Relative difference between the mean buffer size at t and the mean buffer size at cycle start time for the RWP mobility model.

R (m)	r_s/r_d			
	0.25	0.5	0.75	0.99
50	0.75	0.42	0.17	0.0015
100	0.74	0.42	0.18	0.004
200	0.73	0.42	0.18	0.005

 Table 4.4: The relative difference between the mean buffer size at t and the mean buffer size at cycle start time for the RD model.

4.5 Conclusions

We have studied the behavior of the relay buffer of the two-hop relay routing in mobile ad hoc networks by developing a queueing model. The parameters of the queueing model depend on the node mobility pattern.

The main finding is that the expected relay buffer size depends on the expectation and the variance of the nodes contact time. Such analysis is done for the one dimensional random walk over a circle, and for the two RD mobility inside a square.

In the next chapter we will study the delay energy tradeoff of the two-hop relay protocol with multiple packet copies that will be referred to as (MTR) protocol.

Chapter 5

Performance of Multicopy Two-Hop Relay Routing with Limited Packet Lifetime

Considered is a mobile ad hoc network consisting of three types of nodes (source, destination and relay nodes) and using the two-hop relay routing protocol. Packets at relay nodes are assumed to have a limited lifetime in the network. Both closed-form expressions, and asymptotic results when the number of nodes is large, are provided for the packet delivery delay and the energy needed to transmit a packet from the source to its destination. Our model is validated through simulations for two mobility models (random waypoint and random direction mobility models), numerical results for the two-hop relay protocols are reported, and the performance of the two-hop routing and of the epidemic routing protocols are compared.

Note: Part of the results in this chapter have appeared in [HNA06b, HNA06c, HNA06a].

5.1 Introduction

The objective of this chapter is to evaluate the performance of a variant of the two-hop relay protocol that aims at reducing the delivery delay of the packets from the source node to the destination node. This will be done by allowing the source node to make several copies of a packet and to transmit these copies to different relay nodes (one copy per each relay node) at the instants of time when it meets them [SH05, ZNKT06, GNK05]. Much of the rest of this chapter is devoted to study this variant that will be called Multicopy Two-hop Relay (MTR) protocol. The interest is on the performance metrics that bears on the packet delivery delay, packet round trip time, and the overhead induced by the MTR protocol. This will be done under the assumption that, unlike in [GNK05, SH05], packets at relay nodes have a limited lifetime in the network.

Another relay protocol closely related to MTR protocol is the so-called *epidemic routing* protocol [SH03, ZNKT06]. This protocol is identical to MTR, except that in the epidemic routing a relay node is allowed to transmit a packet to any node that it meets, including another relay node. Epidemic routing decreases the delivery delay of packets at the cost of increasing the energy consumed by the network. The performance of the MTR protocol and the epidemic routing protocol will also be compared in this chapter.

The rest of this chapter is organized as follows: Section 5.2 gives a careful description of the MTR protocol, sets the modeling assumptions, and defines the performance metrics of interest (delivery delay, packet round trip time, overhead in terms of the number of copies of a packet). In Section 5.3 we develop a Markovian analysis that yields closed-form expressions for these performance metrics except for the packet round trip time where its n^{th} order-moment will be reported in Section 5.4. Section 5.5 presents an asymptotic analysis of the performance metrics as the number of nodes is large; this analysis uses a mean-field approximation. Validation of the model, and comparison of the performance of the MTR protocol and the epidemic routing protocol are given in Section 5.6. Section 5.7 concludes this chapter.

5.2 The system model

We consider the model introduced in [GNK05]. In this model the characteristics of MANETs are captured through a single parameter, $1/\lambda$, representing the expected inter-meeting time between any pair of nodes. More precisely, there are $N + 1$ nodes consisting of: one source node, one destination node, and $N - 1$ relay nodes. Two nodes may only communicate at certain points in time, called meeting times. The time that elapses between two consecutive meeting times of a given pair of nodes is called the inter-meeting time. In [GNK05] it is assumed that inter-meeting times are mutually independent and identically distributed (iid) random variables (rvs), with an exponential distribution with intensity $\lambda > 0$.

Throughout this chapter we address the scenario where the source node wants to send a single packet to the destination node. To this end the source may use the relay nodes according to the MTR protocol, as explained below.

The MTR protocol works as follows. The source node keeps sending a copy of the packet to all nodes that it meets and that do not have a copy, including the destination node, until the destination node has received a copy of the packet. Transmissions between two nodes are assumed to be instantaneous. This corresponds to the situation where the transfer time of a packet between two nodes is negligible with respect to their inter-meeting time. The way the source node is notified that the destination node has received the packet, either directly from it or from a relay node, will be studied in Section 5.4 by introducing VACCINE the anti-packet dissemination mechanism [SH05]. A relay node that possesses a copy of the packet is only allowed to send it to the destination node.

In addition to the model in [GNK05] we assume throughout this chapter that each *copy* of the packet has a *Time-To-Live* (TTL). When the TTL of a copy expires then the copy is destroyed. TTLs are assumed to be iid rvs with an exponential distribution with rate $\mu > 0$. The packet to be sent by the source has no TTL associated with it, so that the source is always able to send a copy to another node. (If the packet at the source has a TTL then there is a non-zero probability that the destination node will never receive the packet. This scenario is not considered in this chapter.)

We assume that the source is ready to transmit the packet to the destination at time $t = 0$. The (packet) delivery time (or delivery delay), T_d , is the first time after $t = 0$ when the destination node receives the packet (or a copy of the packet). In addition to T_d , we will evaluate C_d , the number of copies in the system at the delivery time, and G_d , the total number of copies generated by the source before the delivery time (Section 5.3). The latter is related to the overhead induced by the MTR protocol and, in particular, to the total energy needed to deliver the packet to the destination. Note that these three performance metrics are independent of the anti-packet dissemination mechanism VACCINE since the latter is activated at the delivery time of the packet. Considering VACCINE, the packet round-trip-time, T_r , that defines the first time after $t = 0$ when the source receives the anti-packet is also evaluated, more details will be provided in Section 5.4.

What are the results obtained in our simple setting (single packet and instantaneous transmission times) that could shed light on the performance of the MTR protocol in more realistic contexts (multiple packets, non-zero transmission times, limited relay storage capacity, etc.)? First, note that the packet delivery delay obtained in our setting gives a lower-bound, as a consequence of the instantaneous transmission time. Second, the protocol overhead, measured in terms of the total number of copies per-packet generated, gives an upper-bound. This is so because in the realistic context the source will not systematically be able to transmit a packet to a relay node that it encounters.

In this chapter, we will use the theory of the finite-state continuous-time absorbing Markov chain to compute the above performance metrics. Lemma 1 summarizes some results of this theory used in the following.

Lemma 1 *Consider a finite-state, continuous-time, Markov chain $\{M(t), t \geq 0\}$, of states*

space $\zeta = \{1, \dots, m, m+1, \dots, m+s\}$ and infinitesimal generator matrix, $\mathbf{G}$, of form

$$\mathbf{G} = \left(\begin{array}{c|c} \mathbf{U} & \mathbf{V} \\ \hline \mathbf{0}_m & \mathbf{0}_s \end{array} \right),$$

where $\mathbf{U}$ is a m -by- m matrix, $\mathbf{V}$ is a m -by- s matrix, $\mathbf{0}_m$ is a s -by- m matrix of all entries equal to 0, $\mathbf{0}_s$ is a s -by- s matrix of all entries equal to 0. The states $\{m+1, \dots, m+s\}$ are the absorbing states. Then,

- (a) The states $\{1, \dots, m\}$ are all transient if and only if the matrix $\mathbf{U}$ is non-singular.
- (b) The cumulative probability distribution, $F(\cdot)$, of the time until the absorption in one of the absorption states $\{m+1, \dots, m+s\}$, given that $M(0) = i$ for $1 \leq i \leq m$, reads

$$F(t) = 1 - \alpha_i \exp(\mathbf{U}t)e, \quad t \geq 0, \quad (5.1)$$

where α_i is the m -dimensional row vector whose all components equal to 0 except the i th one that it is equal to 1, e is the m -dimensional column vector whose all components are equal to 1, and

$$\exp(\mathbf{U}t) := \sum_{i=0}^{\infty} \frac{(\mathbf{U}t)^i}{i!},$$

with $(\mathbf{U}t)^0 = I$ the m -by- m identity matrix. Given that $M(0) = i$, the n^{th} order-moment of time until absorption reads

$$\mu_i^n = (-1)^n n! (\alpha_i \mathbf{U}^{-n} e), \quad i \geq 0. \quad (5.2)$$

- (c) The (i,j) -entry of $(-\mathbf{U}^{-1})$, for $1 \leq i, j \leq m$, is the expected amount of time spent in the transient state j , given that $M(0) = i$.

- (d) Let Δ denote the m -by- m diagonal matrix of diagonal entries equal to those of $\mathbf{U}$. The (i,j) -entry of $(\mathbf{U}^{-1}\Delta)$, for $1 \leq i, j \leq m$, gives the expected number of visits to state j , given that $M(0) = i$.

- (e) The (i,j) -entry of $(-\mathbf{U}^{-1}\mathbf{V})$, for $1 \leq i \leq m$ and $m+1 \leq j \leq m+s$, is the probability that absorption occurs in the state j , given that $M(0) = i$. $\diamond$

Proof: Assertions (a) and (b) are proved in [Neu81, Lemma 2.2.1 and 2.2.2]. Assertion (c) is the consequence of [NM81, Corollary 2]. To prove Assertion (d) and (e), consider the finite-state, discrete-time, absorbing Markov chain embedded just before the jump times of the Markov chain $\{M(t), t \leq 0\}$. The one-step transition probability matrix of the embedded Markov chain reads

$$\left(\begin{array}{cc} \mathbf{I} - \Delta^{-1}\mathbf{U} & -\Delta^{-1}\mathbf{V} \\ \mathbf{0}_m & \mathbf{I}_s \end{array} \right),$$

where $\mathbf{I}_s$ is the s -by- s identity matrix. Theorems [GS97, Theorems 11.4 and 11.6] applied on the embedded Markov chain gives assertion (d) and (e) respectively. $\square$

Coming back to our problem, in the following section we provide a Markovian Analysis that gives closed-form expressions of the distribution of T_d and C_d , and expected value of G_d .

5.3 Markovian analysis

The state of the system is represented by the random variable $I(t) \in \{1, 2, \dots, N, a\}$, where $I(t) \in \{1, 2, \dots, N\}$ gives the number of copies when the packet has not been delivered to the destination (i.e., for $0 \leq t < T_d$) and $I(t) = a$ for $t \geq T_d$. Under the assumptions made in Section 5.2, $\{I(t), t \geq 0\}$ is an absorbing, finite-state, continuous-time Markov chain, with absorbing state a (referred to as $\mathbf{M}$ from now on). Let $\mathbf{P} = [p(i, j)]$ be the infinitesimal generator matrix of $\mathbf{M}$. From the transition rate diagram of $\mathbf{M}$ in Figure 5.1 we find

$$\begin{aligned} p(i, i+1) &= (N-i)\lambda, \quad i = 1, \dots, N-1, \\ p(i, i-1) &= (i-1)\mu, \quad i = 2, \dots, N, \\ p(i, a) &= i\lambda, \quad i = 1, \dots, N, \\ p(i, i) &= -(N\lambda + (i-1)\mu), \quad i = 1, \dots, N, \\ p(i, j) &= 0, \text{ otherwise,} \end{aligned}$$

The infinitesimal generator matrix $\mathbf{P}$ of the Markov chain $\mathbf{M}$ can be written as

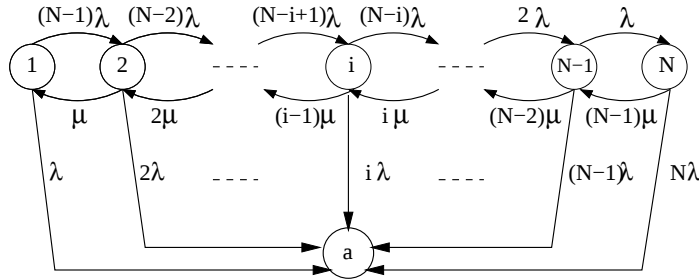


Figure 5.1: Transition rate diagram of $\mathbf{M}$.

$$\mathbf{P} = \left(\begin{array}{c|c} \mathbf{Q} & \mathbf{R} \\ \hline \mathbf{0} & 0 \end{array} \right), \quad (5.3)$$

where $\mathbf{Q} = [p(i, j)]_{1 \leq i, j \leq N}$, $\mathbf{R} = (p(1, a), \dots, p(N, a))^T$, and $\mathbf{0}$ is the N -dimensional row vector whose all components are equal to 0.

Let $\mathbf{Q}^* := -\mathbf{Q}^{-1}$. The following Lemma shows that $\mathbf{Q}^*$ exists and finds the closed-form expression of its (i, j) -entry.

Lemma 2 *The matrix $\mathbf{Q}$ has N distinct, real, and strictly negative eigenvalues equal to $\mu z_1, \dots, \mu z_N$ where z_k is given by*

$$z_k := \frac{-N(2\rho + 1) + 1 - (N + 1 - 2k)\sqrt{4\rho + 1}}{2}, \quad k = 1, 2, \dots, N \quad (5.4)$$

with $\rho := \lambda/\mu$. Therefore the matrix $\mathbf{Q}^* = -\mathbf{Q}^{-1}$ exists, and its (i, j) -entry is given by

$$q^*(i, j) := -\frac{1}{\mu \binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi_i^k \Psi_j^k}{z_k \Psi^k \tau^2 (\Psi^k)^T}, \quad (5.5)$$

with $\Psi^k := (\Psi_1^k, \dots, \Psi_N^k)$ where

$$\Psi_i^k := (-1)^{i-1} x_2^{N-i} \sum_{l=l_0}^{l_1} \binom{k-1}{l} \binom{N-k}{i-1-l} \left(\frac{x_1}{x_2}\right)^{k-1-l}, \quad (5.6)$$

$$x_1 := \frac{-1 - \sqrt{1 + 4\rho}}{2\rho}, \quad x_2 := \frac{-1 + \sqrt{1 + 4\rho}}{2\rho},$$

and where $l_0 := \max(0, i-1-N+k)$, $l_1 := \min(i-1, k-1)$, and $\tau := \text{diag}(\tau_1, \dots, \tau_N)$, with

$$\tau_i := \left(\binom{N-1}{i-1} \rho^{i-1} \right)^{-1/2}. \quad (5.7)$$

◇

Proof: See Appendix 5.A. □

Lemma 2 implies that the matrix $\mathbf{Q}$ is non-singular. Thus, by Lemma 1.a the states $\{1, 2, \dots, N\}$ of $\mathbf{M}$ are all transient. Further, Lemma 1.c gives that the (i, j) -entry of $\mathbf{Q}^*$, $q^*(i, j)$, is the expected time spent in state j before absorption given that $I(0) = i$; On the other hand, let n_{ij} be the number of visits to state j before absorption given that $I(0) = i$, and let T_j be the sojourn time in state j . Then $m(i, j) := E[n_{ij}]$ and $q^*(i, j)$ verify the following identity

$$m(i, j) = \frac{q^*(i, j)}{E[T_j]} = q^*(i, j)(N\lambda + \mu(j-1)), \quad j = 1, 2, \dots, N, \quad (5.8)$$

where $E[T_j] = 1/(N\lambda + \mu(j-1))$ for $j = 1, 2, \dots, N$, $q^*(i, j)$ is given in closed-form in (5.5), and where the first equality follows from Lemmas 1.c and 1.d.

We are now in the position to compute the expected delivery delay, the distribution of the number of copies at the delivery instant, and the expected number of copies generated by the source.

5.3.1 Delivery delay

In this section we first determine $E_i[T_d]$, the expected delivery delay given that $I(0) = i \in \{1, 2, \dots, N\}$, from which the expected delivery delay $\bar{T}_d = E_1[T_d]$ will follow.

$E_i[T_d]$ is the expected time before absorption starting from the transient state i . Recall that $q^*(i, j)$ gives the expected time spent in state j before absorption given that $I(0) = i$. Then,

$$E_i[T_d] = \sum_{j=1}^N q^*(i, j). \quad (5.9)$$

Note that (5.9) can be derived directly from Lemma 2.b. Plugging the value of $q^*(i, j)$ in Lemma 2 in the latter equation gives

$$E_i[T_d] = -\frac{1}{\mu \binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi^k \mathbf{1}^T}{z_k \Psi^k \tau^2 (\Psi^k)^T} \Psi_i^k, \quad (5.10)$$

where $\mathbf{1}^T$ is the N -dimensional column vector whose all components are equal to 1. Quantities Ψ_i^k , τ and z_k are defined in Lemma 2. Note that these quantities are only dependent on ρ and N .

More generally, the tail probability distribution of T_d , starting from $I(0) = i$ is given by

$$P_i(T_d \geq t) = \frac{1}{\binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi^k \mathbf{1}^T}{\Psi^k \tau^2 (\Psi^k)^T} \Psi_i^k e^{z_k \mu t}. \quad (5.11)$$

The proof of this result is shown in Appendix 5.B. Observe that $P_i(T_d \geq t)$ is a weighted sum (mixture) of exponentials with parameters that depend on ρ , μ , and N .

5.3.2 Number of copies in the system at delivery time

Let $P_i[C_d = j]$ be the probability that the number of copies in the network at the delivery time is j , given there are i copies in the network at time $t = 0$. We assume without loss of generality that the Markov chain $\mathbf{M}$ is left-continuous so that $P_i[C_d = j] = P[I(T_d-) = j]$ (by convention $I(t-)$ is the state of the process $\mathbf{M}$ just before time t). In words, $P_i[C_d = j]$ is the probability that the last visited state before absorption is j , given that the initial state is i .

If we split the absorbing state a into N absorbing states $a_1, \dots, a_N$, as shown in Figure 5.2, we will not affect the dynamics of the original Markov chain before absorption. This means that the matrix $\mathbf{Q}$ in (5.3) of the original absorbing Markov chain is the same as its corresponding matrix of the modified absorbing Markov chain. Clearly, $P_i[C_d = j]$ is now equal to the probability that the modified chain is absorbed in state a_j . Let b_{i,a_j} denote

this probability. From Lemma 1.e we know that

$$b_{i,a_j} = \sum_{k=1}^N q^*(i, k) r(k, a_j),$$

where $r(k, a_j)$ is the infinitesimal transition rate from state k to the absorbing state a_j in the modified Markov chain. Clearly (see Figure 5.2) $r(k, a_j) = j\lambda$ if $k = j$ and $r(k, a_j) = 0$ if $k \neq j$. Therefore,

$$P_i[C_d = j] = q^*(i, j) r(j, a_j) \quad (5.12)$$

The n^{th} order-moment of C_d is equal to (with $\mathbf{J}_n := (1, \dots, l^n, \dots, N^n)$)

$$E_i[C_d^n] = -\frac{1}{\binom{N-1}{i-1} \rho^{i-2}} \sum_{k=1}^N \frac{\Psi_i^k}{z_k \Psi^k \tau^2 (\Psi^k)^T} \Psi^k \mathbf{J}_{n+1}^T. \quad (5.13)$$

Coming back to the original problem, the probability distribution of the number of copies at delivery time is given by $P_1[C_d = j]$, and the n^{th} order-moment is given by $E_1[C_d^n]$.

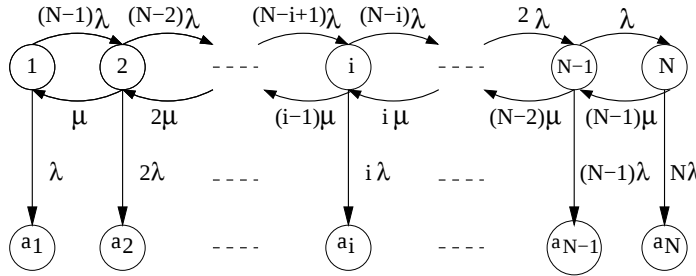


Figure 5.2: The modified absorbing Markov chain with N absorbing states.

5.3.3 Total number of copies generated by the source before delivery time

The objective is to find $\overline{G_d}$, the expected number of copies generated by the source before the delivery time (or equivalently, before absorption). Let $G_d^{i,j}$ be the number of copies generated by the source before absorption given that the chain starts in state i and that state j is the last state visited before absorption (i.e., $I(T_d-) = j$). Introduce $J^{i,j}(k, k+1)$ (resp. $J^{i,j}(k+1, k)$) the number of transitions from state k (resp. state $k+1$) to state $k+1$ (resp. state k) given that $I(0) = i$ and $I(T_d-) = j$. It is easy to see that for $k = 1, 2, \dots, N-1$ that

$$J^{i,j}(k, k+1) = J^{i,j}(k+1, k) + \mathbf{1}_{\{i \leq k < j\}} - \mathbf{1}_{\{j \leq k < i\}}, \quad (5.14)$$

A copy of the packet is generated by the source each time there is a transition from state k to state $k + 1$ for all states $k = 1, 2, \dots, N - 1$. Hence,

$$G_d^{i,j} = \sum_{k=1}^{N-1} J^{i,j}(k, k+1).$$

On the other hand, $n_{i,k}^j$, the total number of visits to state k given that $I(0) = i$ and $I(T_d-) = j$, satisfies the relation

$$\sum_{k=1}^N n_{i,k}^j = \sum_{k=1}^{N-1} J^{i,j}(k, k+1) + \sum_{k=1}^{N-1} J^{i,j}(k+1, k) + 1.$$

From (5.14) we find that

$$\sum_{k=1}^{N-1} J^{i,j}(k+1, k) = \sum_{k=1}^{N-1} J^{i,j}(k, k+1) + i - j.$$

Combining the three last identities gives

$$G_d^{i,j} = \frac{1}{2} \left[\sum_{k=1}^N n_{i,k}^j + j - i - 1 \right].$$

The expected number of copies given that $I(0) = i$, denoted by $\overline{G_d^i}$, is given by (Hint: remove the conditioning on $C_d = j$)

$$\overline{G_d^i} = \frac{1}{2} \left[\sum_{k=1}^N m(i, k) + E_i[C_d] - i - 1 \right] \quad (5.15)$$

where $m(i, k)$ is given in (5.8) and $E_i[C_d]$ is given in (5.13) (with $n = 1$). Finally, $\overline{G_d} = \overline{G_d^1}$. Note that the probability distribution of G_d can be computed by defining a two-dimensional and continuous-time absorbing Markov chain with state (i, c) , where $i \in \{1, \dots, N\}$ represents the number of copies in the network, and $c \in \mathbf{N}$ denotes the total number of copies generated by the source. Thus, $P[G_d = l]$ is the probability that the absorption occurs at one of the following transient states $\{(i, l) : 1 \leq i \leq \min(l+1, N)\}$.

Consider only the energy consumption due to packet transmission and decoding. Let p_t be the energy needed at the sender to transmit a packet to another node and let p_r be the energy needed at the receiver to decode a packet. The energy consumed by the source before the packet is delivered to the destination $P_s = p_t \overline{G_d}$, since the source needs to generate on the average $\overline{G_d}$ copies of the packet before one copy reaches the destination. The energy consumed by all nodes before the delivery time is given by $P_d = (p_t + p_r) \overline{G_d}$.

In the next section we will introduce a two-dimensional absorbing Markov chain that gives the n^{th} order-moment of the packet round-trip-time (RTT).

5.4 Packet round trip time

This section finds the n^{th} order-moment of the packet round trip time, T_r , which defines the first time, after time $t = 0$, that the source receives the acknowledgment (Ack) of the packet. This occurs either when the source meets directly the destination or when the source meets a relay node that carries the Ack. The first time the Ack is generated is when the destination receives the packet for the first time. In our case the Ack is referred to as the *anti-packet*, since on receiving the Ack a node deletes the corresponding copy (if any). The node that carries the antipacket disseminates it to all the nodes in the network including the ones that do not carry a copy of the packet. This anti-packet dissemination mechanism is called VACCINE, for more details see [SH05]. For simplicity, in the following model it is assumed that the destination does not contribute in the dissemination of the anti-packet.

Let $\{(I(t), L(t)), t \geq 0\}$ denote a two-dimensional finite-state continuous-time process with state (i, l) , $i \in \{1, \dots, N\}$ and $l \in \{0, \dots, N - i\}$, if the number of copies is i and the number of relay nodes that carry the anti-packet is l in the network at $t < T_r$, and $(I(t), L(t)) = a$ when $t \geq T_r$ (a is the absorbing state). Under the above assumption made $\{(I(t), L(t)), t \geq 0\}$ is an absorbing, two-dimensional, finite-state, continuous-time Markov chain (referred as $\mathbf{M}_r$). The infinitesimal generator matrix $\mathbf{P}_r = [p_r((i, l), \cdot)]$ of $\mathbf{M}_r$ gives

$$\begin{aligned}
 p_r((i, l), (i + 1, l)) &= \lambda(N - i - l), \quad i = 1, \dots, N - 1, \quad l = 0, \dots, N - i, \\
 p_r((i, l), (i, l + 1)) &= \lambda l(N - i - l), \quad i = 1, \dots, N - 1, \quad l = 0, \dots, N - i, \\
 p_r((i, l), (i - 1, l)) &= \mu(i - 1), \quad i = 2, \dots, N, \quad l = 0, \dots, N - i, \\
 p_r((i, l), (i - 1, l + 1)) &= \lambda(l + 1)(i - 1), \quad i = 2, \dots, N, \quad l = 0, \dots, N - i, \\
 p_r((i, l), (i, l)) &= -[\lambda(N - l)(l + 1) + \mu(i - 1)], \quad i = 1, \dots, N - 1 \\
 &\quad, \quad l = 0, \dots, N - i, \\
 p_r((i, l), a) &= \lambda(l + 1), \quad i = 1, \dots, N, \quad l = 0, \dots, N - i \\
 p_r((i, l), (j, d)) &= 0, \quad \text{otherwise}
 \end{aligned}$$

Let $\mathbf{SC}_l$, for $l = 0 \dots N - 1$, denote the sub-chain of $\mathbf{M}_r$ of state space $\{(i, l) : i = 1 \dots N - l\}$. The infinitesimal generator matrix $\mathbf{P}_r$ of the Markov chain $\mathbf{M}_r$ can be written as

$$\mathbf{P}_r = \left(\begin{array}{cccc|c} \mathbf{Q}_0 & \mathbf{\Lambda}_0 & & & \mathbf{A}_0 \\ & \ddots & \ddots & & \\ & & \mathbf{Q}_l & \mathbf{\Lambda}_l & \mathbf{A}_l \\ & & & \ddots & \\ & & & & \mathbf{Q}_{N-2} & \mathbf{\Lambda}_{N-2} & \mathbf{A}_{N-2} \\ & & & & & \mathbf{Q}_{N-1} & \mathbf{A}_{N-1} \\ \hline & & & & & & 0 \end{array} \right), \quad (5.16)$$

where $\mathbf{Q}_l = [p_r((i, l), \cdot)]_{1 \leq i \leq N-l}$, for $l = 0 \dots N - 1$, is the tridiagonal matrix that represents the infinitesimal generator matrix of the sub-chain $\mathbf{SC}_l$, $\mathbf{\Lambda}_l$ for $l = 0 \dots N - 2$ is the $(N -$

l)-by- $(N - l - 1)$ bidiagonal matrix that represents the transition rates from the states of $\mathbf{SC}_l$ to the state of $\mathbf{SC}_{l+1}$, and where $\mathbf{A}_l = [p_r((i, l), a)]_{1 \leq i \leq N-l}$, for $l = 0 \cdots N - 1$, is the vector that represents the transition rates from $\mathbf{SC}_l$ to a . We refer as $[\mathbf{Q}]$ the block bidiagonal matrix composed of $\mathbf{Q}_l$ and $\mathbf{A}_l$. Defines $\mathbf{Q}_r^* = -[\mathbf{Q}]^{-1}$. Thus, $\mathbf{Q}_r^*$ is an upper block triangular matrix of form

$$\mathbf{Q}_r^* = \begin{pmatrix} -\mathbf{Q}_0^{-1} & \mathbf{U}_{0,1} & \cdots & \mathbf{U}_{0,N-1} \\ & \ddots & \ddots & \vdots \\ & & -\mathbf{Q}_l^{-1} & \mathbf{U}_{l,N-1} \\ & & & \ddots \\ & & & -\mathbf{Q}_{N-2}^{-1} & \mathbf{U}_{N-2,N-1} \\ & & & & -\mathbf{Q}_{N-1}^{-1} \end{pmatrix}, \quad (5.17)$$

where

$$\mathbf{U}_{l,j} = (-1)^{j-l+1} \left(\prod_{k=l}^{j-1} \mathbf{Q}_k^{-1} \mathbf{A}_k \right) \mathbf{Q}_j^{-1}, \quad 0 \leq l, j \leq N - 1, \quad l < j. \quad (5.18)$$

Thus in order to find $\mathbf{Q}_r^*$ in closed-form one needs to find $\mathbf{Q}_l^{-1}$ in closed-form and use (5.18). In the following, we will find closed-form of $\mathbf{Q}_l^{-1}$.

The matrix $\mathbf{Q}_l$ can be decomposed as

$$1/\mu \mathbf{Q}_l = \mathbf{A}_{N-l} - \rho l(N-l) \mathbf{I}, \quad (5.19)$$

where $\mathbf{I}$ the identity matrix, and $\mathbf{A}_{N-l}$ is the $(N-l)$ -by- $(N-l)$ matrix given in (5.39) in Appendix 5.A with N replaced by $N-l$. In Appendix 5.A for all integer $N > 0$, the inverse of the matrix $\mathbf{A}_N$, and its N distinct and strictly negative eigenvalues and their corresponding left/right eigenvectors have been computed in explicit form.

Let w_k , $\mathbf{\Omega}_k$, and $\mathbf{\Xi}_k$ denote the k th eigenvalue and its associated left/right eigenvector of $\mathbf{A}_{N-l}$ for $k = 1 \cdots N-l$. From (5.19), we see that the eigenvalues of $\mathbf{Q}_l$ are equal to $\mu w_k - \lambda l(N-l)$ for $1 \leq k \leq N-l$, and are all distinct and strictly negative. Further, the left/right eigenvectors of $\mathbf{Q}_l$ are equal to left/right eigenvectors $\mathbf{\Omega}_k/\mathbf{\Xi}_k$ of $\mathbf{A}_{N-l}$. From Lemma 2, the (i,j) -entry of $-\mathbf{Q}_l^{-1}$ reads

$$q_l^*(i, j) = -\frac{1}{\mu \binom{N-l-1}{i-1} \rho^{i-1}} \sum_{k=1}^{N-l} \frac{\Omega_i^k (\mathbf{\Omega}^k \mathbf{1}^T)}{\mathbf{\Omega}^k \delta^2 (\mathbf{\Omega}^k)^T} \frac{1}{w_k - \rho l(N-l)}. \quad (5.20)$$

where $\mathbf{1}^T$ is the N -dimensional column vector whose all components are equal to 1. Quantities w_k , Ω_i^k , and δ are z_k , Ψ_i^k , and τ that are defined in (5.4), (5.6), and (5.7) (with N replaced by $N-l$). We showed that $\mathbf{Q}_l^{-1}$ exists in closed-form, and from (5.18) it follows by induction argument that the matrix $\mathbf{Q}_r^*$ exists. Hence, all the states of $\mathbf{M}_r$ are transient and a is the only absorption state (Lemma 1.a).

We are now in a position to derive the n^{th} order-moment of the packet round trip time.

5.4.1 The n^{th} -order moment of the packet round trip time

The time to absorption is the first time that the source receives the antipacket (packet acknowledgment). Let $E[T_r]$ denote the expected time to absorption given that $(I(0), L(0)) = (1, 0)$. Lemma 1.c gives that

$$E[T_r] = \mathbf{Q}_{\mathbf{r}(1, \cdot)}^* \mathbf{e}, \quad (5.21)$$

where $\mathbf{Q}_{\mathbf{r}(1, \cdot)}^*$ is the first row of $\mathbf{Q}_{\mathbf{r}}^*$ where its entries are given in (5.18) and (5.20), and $\mathbf{e}$ is a $(N^2 - N)$ -dimensional vector of all entries equal to 1. More general the n^{th} order-moment of T_r reads

$$E[T_r^n] = (-1)^n n! (\mathbf{Q}_{\mathbf{r}}^*)_{(1, \cdot)}^n \mathbf{e}, \quad (5.22)$$

where $(\mathbf{Q}_{\mathbf{r}}^*)_{(1, \cdot)}^n$ is the first row of $(\mathbf{Q}_{\mathbf{r}}^*)^n$ (see Lemma 1.b).

5.5 Asymptotic analysis

In this section we derive asymptotic results for the expected delivery delay, the expected number of copies at delivery instant, and the expected total number of copies generated at delivery time of the MTR protocol when the number of nodes N is large. Deriving these asymptotic results from the explicit formulas in (5.10) and (5.13), respectively, is not easy (in the more simpler case when there are no timeouts getting asymptotics from the explicit results were already quite involved [GNK05, Appendix A]).

We shall instead follow a mean field approach to find approximations of these asymptotics. The same approach was used in [SH03] and in [ZNKT06] to derive asymptotic results for epidemic models.

The mean field approximation says that $X(t)$ (resp. $G(t)$), the *expected* number of copies (resp. the expected number of copies generated by the source) in the network at time t , before absorption, when N is large, can be approximated by the solution of the following 1st-order differential equations (see [Kur70] for the general theory)

$$\dot{X}(t) = \lambda(N - X(t)) - \mu(X(t) - 1), \quad t > 0. \quad (5.23)$$

$$\dot{G}(t) = \lambda(N - X(t)), \quad t > 0. \quad (5.24)$$

Equation (5.23) simply reflects the fact that at time t $X(t)$ increases with the rate $\lambda(N - X(t))$ and decreases with the rate $\mu(X(t) - 1)$ and (5.24) gives that at time t $G(t)$ increases with the rate $\lambda(N - X(t))$. We need to complement these equations with another equation whose the solution approximates $D(t) := P(T_d < t)$, the probability distribution of the delivery delay. It was found in [SH03] that

$$\dot{D}(t) = \lambda X(t)(1 - D(t)), \quad t > 0. \quad (5.25)$$

Solving (5.23), (5.24), and (5.25) with the initial conditions $X(0) = x_0$ ($x_0 = 1$ in our model), $G(0) = 0$, and $D(0) = 0$ yields

$$X(t) = \frac{N\lambda + \mu}{\lambda + \mu} + \left(x_0 - \frac{N\lambda + \mu}{\lambda + \mu}\right)e^{-(\lambda + \mu)t} \quad (5.26)$$

$$G(t) = \lambda Nt - f_N(t), \quad D(t) = 1 - e^{-f_N(t)} \quad (5.27)$$

where $f_N(t) := \frac{\lambda}{\lambda + \mu} \left[(N\lambda + \mu)t + \left(x_0 - \frac{N\lambda + \mu}{\lambda + \mu}\right)(1 - e^{-(\lambda + \mu)t}) \right]$. It can be checked that $D(0) = 0$, $\lim_{t \rightarrow \infty} D(t) = 1$ and $t \rightarrow D(t)$ is nondecreasing, so that $D(t)$ is indeed a probability distribution of a proper rv. As expected from the very definition of $X(t)$, we note that $X(\infty) = (N\lambda + \mu)/(\lambda + \mu)$ is the expected stationary number of customers in a finite-state birth and death process, with birth rate (resp. death rate) $\lambda(N - i)$ (resp. $\mu(i - 1)$) in state $i \in \{1, 2, \dots, N\}$.

5.5.1 Delivery delay

By definition, $E[T_d] = \int_0^{+\infty} P(T_d > t) dt$, so that from (5.27) $E[T_d]$ can be approximated by

$$E[T_d] \approx \int_0^{+\infty} e^{-f_N(t)} dt \quad (5.28)$$

when N is large. When N is large it is easily seen that the dominant contribution of $e^{-f_N(t)}$ to the above integral comes from small values of t since $f_N(t)$ is a nondecreasing function of N . Hence, $e^{-f_N(t)}$ can be approximated by $e^{-f_N''(0)t^2/2}$ since $f_N(0) = 0$ and since $f_N'(0) = \lambda x_0$ does not depend on N , with $f_N''(0) = \lambda(N\lambda + \mu - (\lambda + \mu)x_0)$. For $0 \leq x_0 < X(\infty)$ this gives the 1st-order asymptotics

$$E[T_d] \approx \sqrt{\frac{\pi}{2\lambda(N\lambda + \mu - (\lambda + \mu)x_0)}} \approx \frac{1}{\lambda} \sqrt{\frac{\pi}{2N}} \quad (5.29)$$

for $N \rightarrow \infty$. The 2nd-order asymptotics of $E[T_d]$ can be obtained by expanding $f_N(t)$ in Taylor series at the order three in the vicinity of $t = 0$. We find

$$\begin{aligned} E[T_d] &\approx \int_0^{+\infty} e^{-\frac{f_N''(0)}{2!}t^2} \left(1 - \frac{f_N^{(3)}(0)}{3!}t^3\right) dt \\ &= \sqrt{\frac{\pi}{2\lambda(N\lambda + \mu - (\lambda + \mu)x_0)}} + \frac{(\lambda + \mu)(N\lambda + \mu - (\lambda + \mu)x_0)}{3\lambda^3(N - 1)^2} \end{aligned} \quad (5.30)$$

for $N \rightarrow \infty$.

Figure 5.3 displays the 1st-order and 2nd-order asymptotics of $E[T_d]$, given in (5.29) and in (5.30), respectively, as a function of N , and compare them with the exact value

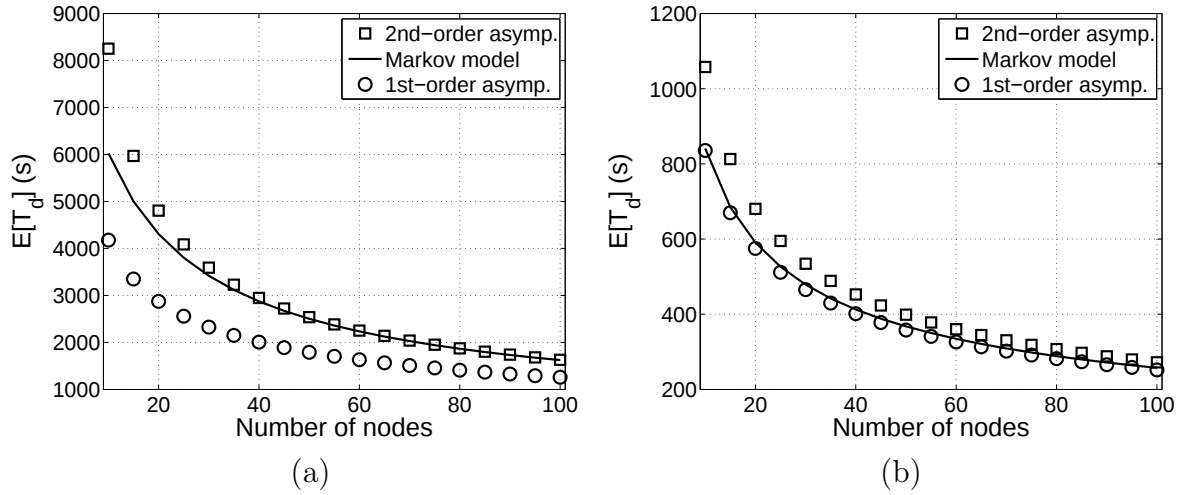


Figure 5.3: Comparing asymptotics for the expected delivery delay to the exact result ($\mu=0.001$: (a) $\lambda=0.0001$, (b) $\lambda=0.0005$.)

obtained in (5.10). We observe that, as N increases, both asymptotics converge to the exact result.

Asymptotics, as N is large, for any order moment of T_d can be derived using a similar approach as shown in the following lemma.

Lemma 3 *When N is large, the n^{th} order-moment of T_d verifies the following identity*

$$\frac{E[T_d^n]}{E[T_d]} \approx \frac{n!}{(\lambda x_0)^{n-1}}. \quad (5.31)$$

◇

Proof: the n^{th} order-moment of T_d is equal to

$$\begin{aligned} E[T_d^n] &= n \int_0^{+\infty} t^{n-1} (1 - D(t)) dt \\ &= n \int_0^{+\infty} t^{n-1} e^{-f_N(t)} dt \\ &= \lambda \int_0^{+\infty} t^n f'_N(t) e^{-f_N(t)} dt \\ &= \lambda X(\infty) \int_0^{+\infty} t^n e^{-f_N(t)} dt + \lambda(x_0 - X(\infty)) \int_0^{+\infty} t^n e^{-f_N(t) - (\lambda + \mu)t} dt, \end{aligned} \quad (5.32)$$

where we integrate by parts and we used (5.27). When N is large the integral

$$\int_0^{+\infty} t^n e^{-f_N(t) - (\lambda + \mu)t} dt \approx \int_0^{+\infty} t^n e^{-f_N(t)} dt.$$

So the n^{th} order-moment of T_d becomes equal to

$$E[T_d^n] \approx \lambda x_0 \int_0^{+\infty} t^n e^{-f_N(t)} dt = \frac{\lambda x_0}{n+1} E[T_d^{n+1}], \quad (5.33)$$

which by recurrence concludes the proof of this Lemma. $\square$

5.5.2 Expected number of copies at delivery instant

When N is large, $E[C_d]$, the mean number of copies at the delivery time T_d , is approximated by $\int_0^{+\infty} X(t) dD(t)$. With the use of (5.26) an integration by part gives

$$E[C_d] \approx x_0 + (N\lambda + \mu - (\lambda + \mu)x_0) \int_0^{\infty} e^{-f_N(t) - (\lambda + \mu)t} dt \quad (5.34)$$

for $N \rightarrow \infty$. By using again the property that the dominant contribution of $e^{-f_N(t) - (\lambda + \mu)t}$ to the above integral comes from small values of t . we may approximate $e^{-f_N(t) - (\lambda + \mu)t}$ by $e^{-f_N''(0)t^2/2}$. Hence,

$$E[C_d] \approx x_0 + \sqrt{\frac{\pi}{2\lambda}} \sqrt{N\lambda + \mu - (\lambda + \mu)x_0} \approx \sqrt{\frac{\pi N}{2}} \quad (5.35)$$

for $N \rightarrow \infty$.

5.5.3 Expected number of copies generated

When N is large, $E[G_d]$, the mean number of copies generated by the source before delivery time T_d , is approximated by $\int_0^{+\infty} G(t) dD(t)$. With the use of (5.27) an integration by part gives

$$E[G_d] \approx \lambda N \int_0^{+\infty} e^{-f_N(t)} dt - 1 \approx \sqrt{\frac{\pi N}{2}} \quad (N \rightarrow \infty). \quad (5.36)$$

From these asymptotic results of $E[T_d]$ and $E[G_d]$ we derive the Little-like formula $E[G_d] \approx \lambda N E[T_d]$ ($N \uparrow \infty$) relating the total expected number of copies generated by the protocol (protocol overhead) to the expected delivery time.

5.6 Validation and numerical results

In this section we first validate the Markov model introduced in Section 5.3 by comparing its performance (expected delivery delay) to that obtained by simulations, for two different mobility models (Random Waypoint (RWP) and Random Direction (RD) models). Simulation results of RWP and RD are obtained using the NS-2 code of the random trip model [BV05]. We then compare the expected delivery delay and the energy consumption induced by the MTR protocol and the epidemic protocol.

5.6.1 Model validation

We have simulated the MTR protocol with exponential timeouts for both the RWP and the RD mobility models. In the RWP model [BMJ⁺98] each node is assigned an initial location in a given area (typically a square) and travels at a constant speed to a destination chosen randomly in this area. The speed is chosen randomly in $(v_{\min}, v_{\max})$, independently of the initial location and destination. After reaching the destination the node may pause for a random time, after which a new destination and speed are chosen, independently of previous speeds, destinations, and pause times. In the RD model [Bet01] each node is assigned an initial direction, speed and travel time. The node then travels in that direction at the given speed and for the given duration. When the travel time has expired, the node may pause for a random time, after which a new direction, speed and travel time are chosen at random, independently of all previous directions, speeds and travel times. When a node reaches a boundary it is either reflected or the area wraps around so that the node reappears on the other side. In both mobility models nodes move independently of each other.

In our simulation settings, for both the RWP and the RD models the area is a square of side-length $L = 2000m$, the speed is constant and equals to $V = 10m/sec.$, there is no pause time, and the transmission range R is constant and the same for all nodes. In addition, in the RD model the travel time is constant and equals to $30sec.$ and the nodes reflect on reaching the boundaries. It has been experimentally observed in [GNK05] that whenever $R \ll L$ then the node inter-meeting time is exponentially distributed with rate $\lambda = 10.94 \frac{RV}{\pi L^2}$ for RWP and $\lambda = 8 \frac{RV}{\pi L^2}$ for RD.

For different values of the ratio N (resp. R/L), Table 5.1 (resp. Table 5.2) reports the expected delivery delay obtained from the Markovian model in (5.10) and by simulations for the RWP and the RD, and give relative errors.

Tables 5.1 and 5.2 show that, for both mobility models, the Markovian model is accurate for N and R/L relatively small. Observe that more accurate results are reported in the case of RD (relative error of 6% when the ratio R/L is less than or equal to 1.25% for 20 nodes – see Table 5.2-B).

N	10	20	30	40	100		10	20	30	40	100
$E_m[T_d]$ (s)	4344	3154	2596	2257	1436		6116	4416	3622	3141	1987
$E_{sim}[T_d]$ (s)	4093	2851	2237	1839	1068		6208	4141	3297	2867	1512
$ 1 - \frac{E_{sim}[T_d]}{E_m[T_d]} $ (%)	6	9	14	18	26		1	6	9	9	24

(A)
(B)

Table 5.1: Expected delivery delay calculated from (5.10) and by simulations ($\mu = 0.0001$, $R/L = 0.5\%$: (A) RWP model, (B) RD model).

5.6.2 Comparison of two-hop relay and epidemic routing protocols

In this section, we compare the expected delivery delay, $E[T_d]$, and the expected number of packets transmitted, $E[G_d]$, as a function of μ , the timeout intensity, for the two-hop relay

R/L (%)	1.25	1	0.5	0.1	1.25	1	0.5	0.1
$E_m[T_d]$ (s)	1216	1529	3154	20102	1678	2116	4416	30264
$E_{sim}[T_d]$ (s)	945	1245	2851	20861	1596	1988	4176	31651
$ 1 - \frac{E_{sim}[T_d]}{E_m[T_d]} $ (%)	22	18	10	4	6	6	5	4

(A)

(B)

Table 5.2: Expected delivery delay calculated from (5.10) and by simulations ($\mu = 0.0001$, $N = 20$: (A) RWP model, (B) RD model).

and the epidemic routing protocols. The absorbing Markov chain modeling the epidemic routing protocol is the same as the absorbing Markov chain in Section 5.3, except that the birth rate in state i is now equal to $\lambda i(N-i)$, since in the epidemic routing protocol all nodes are allowed to generate copies of the packet. The death rate (resp. absorption rate) in state i is unchanged and equal to $\mu(i-1)$ (resp. λi). The computation of the expected delivery delay and the expected number of packets transmitted for the epidemic routing protocol is therefore similar to that carried out for the MTR protocol, except that the fundamental matrix for the epidemic routing protocol cannot be computed in explicit form. This matrix was obtained numerically.

As expected, we observe that the epidemic routing protocol induces a smaller expected delivery delay than the MTR protocol, but at the expense of a much more important overhead in terms of the number of copies generated Figures 5.4 and 5.5. We also point out that the conclusions drawn from the results in Figure 5.5 apply for the energy consumptions P_s and P_d in the case where the energy to transmit (resp. decode) a packet is constant, since we have shown in Section 5.3.3 that in this case P_s and P_d are both linear functions of $E[G_d]$.

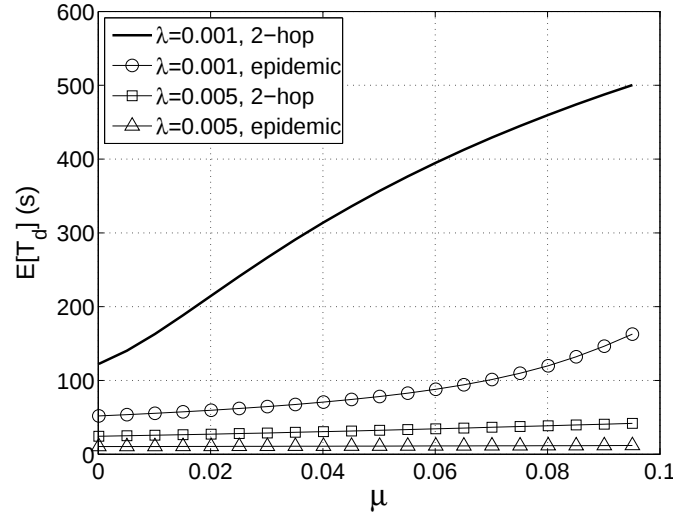


Figure 5.4: Expected delivery delay for two-hop relay and epidemic routing protocols as a function of μ ($N = 100$).

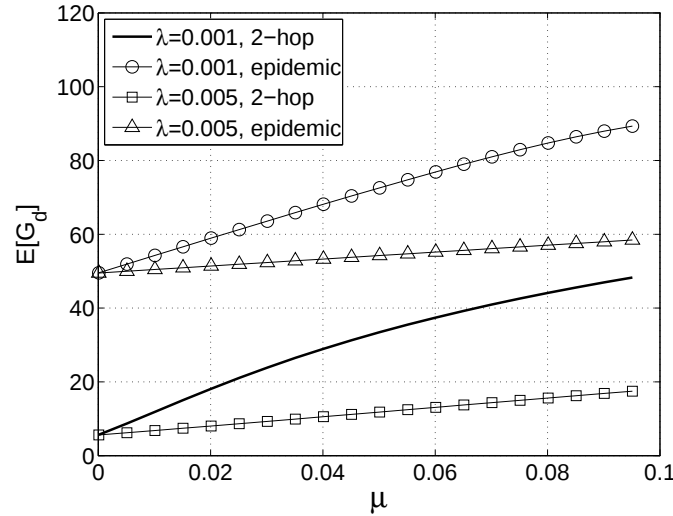


Figure 5.5: Expected number of packet transmitted for two-hop relay and epidemic routing protocols as a function of μ ($N = 100$).

5.7 Concluding remarks

In this chapter, we have evaluated the main performance metrics of the multicopy two-hop relay protocol (MTR) under the assumption that packets in relay nodes have a limited lifetime. Closed-form expressions have been derived for the probability distribution of the packet delivery delay, the distribution of the number of copies in the system at the delivery instant, and the overall expected number of copies generated by the source at the delivery instant, and the moments of the round trip time of the packet. We have observed that the number of packet generated is directly related to the energy needed to transmit the packet to the destination node, in the case when the energy needed to transmit a packet between two nodes and the energy needed to decode a packet are constant.

This chapter aims towards developing *simple* analytical models for quantifying the performance of relay protocols for MANETs and, in particular, for better understanding the delay-energy tradeoff of this class of protocols. In the following chapter will study a class of variants of the MTR protocol that aim at limiting the energy consumption.

5.A Proof of Lemma 2

Lemma 4 *The matrix $\mathbf{Q}$ has N distinct, real, and strictly negative eigenvalues $\mu z_1, \dots, \mu z_N$ where z_k is given by*

$$z_k = \frac{-N(2\rho + 1) + 1 - (N + 1 - 2k)\sqrt{4\rho + 1}}{2},$$

for $k = 1, 2, \dots, N$. Therefore the matrix $\mathbf{Q}^* = -\mathbf{Q}^{-1}$ exists, and its (i, j) -entry is given by

$$q^*(i, j) = -\frac{1}{\mu \binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi_i^k \Psi_j^k}{z_k \Psi^k \tau^2 (\Psi^k)^T}, \quad (5.37)$$

with $\Psi^k = (\Psi_1^k, \dots, \Psi_N^k)$ where

$$\begin{aligned} \Psi_i^k &= (-1)^{i-1} x_2^{N-i} \sum_{l=l_0}^{l_1} \binom{k-1}{l} \binom{N-k}{i-1-l} \left(\frac{x_1}{x_2}\right)^{k-1-l}, \\ x_1 &= \frac{-1 - \sqrt{1+4\rho}}{2\rho}, \quad x_2 = \frac{-1 + \sqrt{1+4\rho}}{2\rho}, \end{aligned}$$

and where $l_0 = \max(0, i-1-N+k)$, $l_1 = \min(i-1, k-1)$, and $\tau = \text{diag}(\tau_1, \dots, \tau_N)$, with $\tau_i = \left(\binom{N-1}{i-1} \rho^{i-1}\right)^{-1/2}$. $\diamond$

Proof. To simplify the computation of $\mathbf{Q}^*$ we introduce the matrix $\mathbf{A}_N$ defined as

$$\mathbf{A}_N = \mathbf{B}\mathbf{Q}, \quad (5.38)$$

where $\mathbf{B} = \text{diag}(b(1), \dots, b(N))$ with $b(i) = 1/\mu$. Matrices $\mathbf{Q}^*$ and $\mathbf{A}_N$ are related through the simple identity

$$\mathbf{Q}^* = -\mathbf{A}_N^{-1}\mathbf{B}. \quad (5.39)$$

In the following we will compute $\mathbf{A}_N^{-1}$. We will follow the approach developed in [AMS82]. We first compute the eigenvalues and left/right eigenvectors of $\mathbf{A}_N$.

Eigenvalues of $\mathbf{A}_N$.

Let z be some eigenvalue of $\mathbf{A}_N$ and let $\Psi = (\Psi_1, \dots, \Psi_N)$ be the associated left eigenvector. That is, $\Psi \mathbf{A}_N = z\Psi$, or equivalently,

$$\rho(N - (i-1))\Psi_{i-1} - (\rho N + i - 1 + z)\Psi_i + i\Psi_{i+1} = 0 \quad (5.40)$$

for $i = 1, \dots, N$, with $\Psi_0 = \Psi_{N+1} = 0$ by convention. Let $\psi(x) = \sum_{j=1}^N \Psi_j x^j$ denote the generating function of Ψ . Multiplying (5.40) by x^i and then summing over i yields

$$\frac{\psi'(x)}{\psi(x)} = \frac{\rho N x - (\rho N - 1 + z) - 1/x}{\rho x^2 + x - 1}. \quad (5.41)$$

Let the zeros of $x^2 + x/\rho - 1/\rho$ be $x_1 = \frac{-1 - \sqrt{1+4\rho}}{2\rho}$ and $x_2 = \frac{-1 + \sqrt{1+4\rho}}{2\rho}$. The unique solution of (5.41) such that $\Psi_N = 1$ is

$$\psi(x) = x(x_1 - x)^{c_1} (x_2 - x)^{c_2} \quad (5.42)$$

$$c_1 := \frac{x_1^2 \rho N - x_1(\rho N - 1 + z) - 1}{\rho x_1(x_1 - x_2)}, \quad c_2 := \frac{-x_2^2 \rho N + x_2(\rho N - 1 + z) + 1}{\rho x_2(x_1 - x_2)}.$$

It is easily seen that $c_1 + c_2 = N - 1$ (Hint: use $x_1 x_2 = -1/\rho$), so that (5.42) also writes

$$\psi(x) = x(x_1 - x)^{c_1} (x_2 - x)^{N-1-c_1}. \quad (5.43)$$

Because $\psi(x)$ is a polynomial of degree N , we observe from (5.43) that necessarily c_1 is an integer lying in the set $\{0, 1, \dots, N-1\}$ since x_1 and x_2 are always distinct.

The equations $c_1 = k - 1$ for $k = 1, \dots, N$ give the following N eigenvalues of $\mathbf{A}_N$:

$$z_k = \frac{-N(2\rho + 1) + 1 - (N + 1 - 2k)\sqrt{4\rho + 1}}{2}, \quad (5.44)$$

for $k = 1, \dots, N$. All eigenvalues of $\mathbf{A}_N$ are distinct (obvious from (5.44)). Furthermore, z_k increases as k increases, and it is easily seen that $z_N < 0$ for $\rho > 0$. Thus, $z_k < 0$ for all $k = 1, \dots, N$.

Left eigenvectors of $\mathbf{A}_N$:

Recall that $\Psi^k = (\Psi_1^k, \dots, \Psi_N^k)$ is the left eigenvector associated with the eigenvalue z_k of $\mathbf{A}_N$. The i th component Ψ_i^k of the eigenvector Ψ^k is the coefficient of x^i in the polynomial $x(x_1 - x)^{k-1} (x_2 - x)^{N-k}$ that is

$$\Psi_i^k = (-1)^{i-1} x_2^{N-i} \sum_{l=l_0}^{l_1} \binom{k-1}{l} \binom{N-k}{i-1-l} \left(\frac{x_1}{x_2}\right)^{k-1-l}.$$

where $l_0 = \max(0, i-1-N+k)$ and $l_1 = \min(i-1, k-1)$.

Right eigenvectors of $\mathbf{A}_N$:

Recall that $\Phi^k = (\Phi_1^k, \dots, \Phi_N^k)^T$ is the right eigenvector associated with the eigenvalue z_k , for $k = 1, \dots, N$. We proceed like in [AMS82, Section 2.4], that is we look for a diagonal matrix $\tau = \text{diag}(\tau_1, \dots, \tau_N)$ such that

$$\tau^{-1} \mathbf{A}_N \tau = (\tau^{-1} \mathbf{A}_N \tau)^T. \quad (5.45)$$

It is easily found that (Hint: solve $\tau_i^2 / \tau_{i+1}^2 = \rho(N-i)/i$ for $i = 1, \dots, N$ with $\tau_1 = 1$)

$$\tau_i = \left(\binom{N}{i} \frac{i}{N} \rho^{i-1} \right)^{-1/2}, \quad i = 1, \dots, N$$

satisfy (5.45). The identity $\Psi^k \mathbf{A}_N = z_k \Psi^k$ implies that $\Psi^k \tau (\tau^{-1} \mathbf{A}_N \tau) = z_k \Psi^k \tau$. Therefore, $\Psi^k \tau$ is a left eigenvector of the matrix $\tau^{-1} \mathbf{A}_N \tau$ associated with the eigenvalue z_k . Since the matrix $\tau^{-1} \mathbf{A}_N \tau$ is symmetric, it has identical left and right eigenvectors. Hence,

$$(\tau^{-1} \mathbf{A}_N \tau)(\Psi^k \tau)^T = z_k (\Psi^k \tau)^T$$

which gives that $\mathbf{A}_N \tau^2 (\Psi^k)^T = z_k \tau^2 (\Psi^k)^T$. This shows that $\alpha_k \tau^2 (\Psi^k)^T$ is a right eigenvector associated with the eigenvalue z_k for any constant $\alpha_k \neq 0$.

Without loss of generality we select the constants $\alpha_1, \dots, \alpha_N$ so that $\Psi^k \Phi^k = 1$ for every $k = 1, \dots, N$. Hence,

$$\alpha_k = \frac{1}{\Psi^k \tau^2 (\Psi^k)^T}$$

for $k = 1, \dots, N$. Finally, $\Phi^k = \tau^2 (\Psi^k)^T / \Psi^k \tau^2 (\Psi^k)^T$, or equivalently

$$\Phi_i^k = \frac{1}{\Psi^k \tau^2 (\Psi^k)^T} \left(\binom{N}{i} \frac{i}{N} \rho^{i-1} \right)^{-1} \Psi_i^k, \quad i = 1, \dots, N. \quad (5.46)$$

The proof is concluded by noting that

$$\mathbf{A}_N^{-1} = \mathbf{F} \text{diag} (1/z_1, \dots, 1/z_N) \mathbf{F}^{-1},$$

where the j^{th} right eigenvector of $\mathbf{A}_N$, Φ^j , is the j^{th} column of the matrix $\mathbf{F}$, and the i^{th} left eigenvector, Ψ^i , is the i^{th} row of the matrix $\mathbf{F}^{-1}$. Hence, the (i,j) -entry of $\mathbf{A}_N^{-1}$ is given by

$$\hat{a}(i, j) = \sum_{k=1}^N \frac{\Phi_i^k \Psi_j^k}{z_k} = \frac{1}{\binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi_i^k \Psi_j^k}{z_k \Psi^k \tau^2 (\Psi^k)^T}, \quad (5.47)$$

by using (5.46). (5.47) together with $q^*(i, j) = -\hat{a}(i, j)/\mu$ (coming from (5.39)) gives (5.37). $\square$

Remark 4 Replacing x by 1 in (5.41) implies the following relation between the eigenvalue z_k and its corresponding left eigenvector Ψ^k

$$\sum_{j=1}^N j \Psi_j^k = -\frac{z_k}{\rho} \sum_{j=1}^N \Psi_j^k, \quad 1 \leq k \leq N. \quad (5.48)$$

5.B Distribution of delivery delay

The delivery delay, T_d , given that there are i copies in the network at time 0 is the time to absorption of the Markov chain $\mathbf{M}$ of Figure 5.1. From Lemma 1.b, the tail distribution of the time to absorption of $\mathbf{M}$ starting from $I(0) = i$ gives

$$P_i(T_d > t) = \alpha_i \exp(\mathbf{Q}t) e, \quad t \geq 0, \quad (5.49)$$

with $\mathbf{Q}$ the transition rate between transient states of $\mathbf{M}$ defined in (5.3), α_i is the N -dimensional row vector whose all components equal to 0 except the i th one that it is equal to 1, e is the N -dimensional column vector whose all components are equal to 1, and

$$\exp(\mathbf{Q}t) := \sum_{i=0}^{\infty} \frac{(\mathbf{U}t)^i}{i!}, \quad (\mathbf{U}t)^0 = \mathbf{I}.$$

It is shown in Appendix 5.A that the matrix $\mathbf{A}_N = \mathbf{1}/\mu\mathbf{Q}$ is invertible and diagonalizable, namely, there exists an invertible matrix $\mathbf{F}$ such that

$$\mathbf{A}_N = \mathbf{F} \text{diag} (z_1, \dots, z_N) \mathbf{F}^{-1}. \quad (5.50)$$

where $z_1, \dots, z_N$ are the eigenvalues of $\mathbf{A}_N$, the j^{th} right eigenvector of $\mathbf{A}_N$, Φ^j , is the j^{th} column of the matrix $\mathbf{F}$, and the left eigenvector, Ψ^i , is the i^{th} row of the matrix $\mathbf{F}^{-1}$. Hence,

$$\mathbf{Q} = \mathbf{F} \text{diag} (\mu z_1, \dots, \mu z_N) \mathbf{F}^{-1}, \quad (5.51)$$

and

$$\exp(\mathbf{Q}t) = \mathbf{F} \text{diag} (e^{\mu z_1 t}, \dots, e^{\mu z_N t}) \mathbf{F}^{-1}. \quad (5.52)$$

Plugging (5.52) into (5.49) gives that

$$P_i(T_d \geq t) = \sum_{k=1}^N \sum_{j=1}^N \Psi_j^k \Phi_i^k e^{z_k \mu t}, \quad (5.53)$$

and (5.46) into the last equation implies that

$$P_i(T_d \geq t) = \frac{1}{\binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi^k \mathbf{1}^T}{\Psi^k \tau^2 (\Psi^k)^T} \Psi_i^k e^{z_k \mu t}, \quad (5.54)$$

where $\mathbf{1}^T$ is the column vector of dimension N whose all components are equal to 1. Since $z_k < 0$ for $1 \leq k \leq N$ and $\rho > 0$, thus $\lim_{t \rightarrow +\infty} P_i(T_d > t) = 0$. Further, $P_i(T_d > 0) = 1$ (Hint: $\Psi^i \Phi^j = 0$ for $i \neq j$, and $\Psi^k \Phi^k = 1$).

Chapter 6

Performance Evaluation of a Class of Two-Hop Relay Protocols

In this chapter, we evaluate the performance of a class of two-hop relay protocols for MANETs via simple models. The focus is on the multicopy two-hop relay (MTR) protocol, where the source may generate multiple copies of a packet, and on the two-hop relay protocol with erasure coding, where a piece of information is fragmented into n blocks. Performance metrics of interest are the time to deliver a single packet to its destination, the number of copies of the packet at delivery instant, and the total number of copies that the source generates; the latter number will be larger when TTLs are associated with the copies of a packet, a situation that we address. We also investigate separately the cases where the number of copies of a packet currently in the network and the number of copies generated are limited so as to limit the energy consumption. Finally, we characterize the delivery delay in the two-hop relay protocol with erasure coding and compare this scheme with the multicopy scheme.

Note: Part of the results in this chapter are under review of INFOCOM 2007.

6.1 Introduction

This chapter evaluates two extensions of the multicopy two-hop relay (MTR) protocol that was introduced in Chapter 5. The first extension called K-copies MTR limits the number of copies of a packet in the network. However, the second extension called K-transmissions MTR limits the number of the copies generated in the network. In addition, the performance of the two-hop relay protocol with *erasure coding* is evaluated and compared with the simple multicopy mechanism. Let now introduce these extensions.

K-copies MTR. In the K-copies MTR, the source node limits the number of packet copies in the network to K . Thus, if the source meets a relay node before meeting the destination and the number of packet copies in the network is less than K , then it sends a copy of the packet to this relay node; this relay node will in turn send the packet to the destination when it comes close to it. Note that in the two-hop relay protocol the source is the only node that generates copies and further it sets the TTLs of the copies. Thus, before the delivery of the packet to the destination the source knows exactly the current number of copies in the network. Basing on this information the source decides to transmit a new copy to a relay node that it encounters if the number of copies is less than K . We note that the limitation of the number of the packet copies in the network of the K-copies MTR comes at the expense of an increase in the packet delivery delay. In Section 6.3.4, we find the optimal value of K , the maximum number of packet copies in the network, that minimizes the expected energy consumed subject to a constraint on the expected delivery delay.

K-transmissions MTR. The K-transmissions MTR limits the energy consumption of the network by limiting the number of copies that the source can generate to K . Thus, if the source meets a relay node before meeting the destination and the number of generated copies is less than K , then it transmits a new copy of the packet to this relay node. Alike the K-copies MTR, the limitation of the energy consumption of the K-transmissions MTR comes at the expense of an increase in the packet delivery delay. A similar scheme was introduced in [SH05] to limit the energy consumption of epidemic routing.

Two-hop relay protocol with erasure coding. In the two-hop relay protocol with erasure coding instead of generating packet copies the source introduces some redundancy in its transmissions and sends more data than the actual information. Upon receiving a piece of data (packet), the source produces n blocks of data. The transmission of the packet is completed when the destination receives the k th block, regardless of the identity of the $k \leq n$ blocks it has received [WJMF05]. More details are provided in Section 6.5 where we will show that erasure coding reduces the variance of the delivery delay.

The rest of this chapter is organized as follows: Section 6.2 summarizes the system model and defines the performance metrics of interest (delivery delay, overhead in terms of the number of packet copies and energy). In Sections 6.3 and 6.4 we develop Markovian analysis that yield closed-form expressions of these three performance metrics for the K-copies MTR and the K-transmissions MTR. Section 6.5 evaluates the performance of the two-hop relay protocol with erasure coding and compares it with the multicopy scheme. Section 6.6 concludes this chapter.

6.2 The system model

We consider the system model introduced in Chapter 5 Section 5.2. The following is a summary of this model.

There are $N + 1$ mobile nodes consisting of one source node, one destination node, and $N - 1$ relay nodes. All nodes move independently of each other according to the same random mobility model. Two nodes may only communicate at certain points in time, called *meeting times*. The time that elapses between two consecutive meeting times of a given pair of nodes is called the *inter-meeting time*. All inter-meeting times are independent and identically distributed (iid) random variables (rvs) with a common exponential probability distribution with rate $\lambda > 0$.

We consider the scenario where the source has a single packet to transmit to the destination.

Each packet copy has a time-to-live (TTL) associated with it: when a TTL expires, the relay node that holds the copy drops it. This relay node then becomes eligible to receive another copy. We assume that the source cannot transmit a copy to a relay node that already holds a copy, and that TTLs are iid with a common exponential distribution with rate $\mu > 0$. Note that there is no TTL associated with the packet of the source.

We define the (packet) delivery delay T_d as the first time when the destination receives the original packet sends by the source or a copy sends by a relay node, whichever arrives first to the destination. In addition to T_d , let C_d and G_d denote the number of copies of the packet before the delivery to the destination and G_d the total number of transmissions before the delivery to the destination. Note that the latter metric is related to the energy needed to deliver a packet to the destination.

For the energy consumption model, we will only consider the energy consumption due to packet transmission and decoding. Let p_t be the energy needed at the sender to transmit a packet to another node and let p_r be the energy needed at the receiver to decode a packet. The energy consumed by the source before the packet is delivered to the destination $P_s = p_t \overline{G_d}$ where $\overline{G_d}$ denotes the expectation of G_d , since the source needs to generate on the average $\overline{G_d}$ copies of the packet before one copy reaches the destination. The energy consumed by all nodes before the delivery time is given by $P_d = (p_t + p_r) \overline{G_d}$.

6.3 K-copies MTR protocol

In this section we consider the K-copies MTR protocol with exponentially distributed node inter-meeting times of parameter $\lambda > 0$ and exponentially distributed TTLs of parameter $\mu > 0$, and where the number of copies of a packet in the network may not exceed K (including the packet at the source), where K is an arbitrary integer less than or equal to N . We recall that we only focus on the transmission of a single packet between a given source and a given destination, and that the packet at the source has no TTL (only copies have a TTL).

In this section we derive closed-form expressions for the n^{th} order-moment of T_d for

all $n \geq 1$, the probability distribution of C_d , and the expectation of G_d . We conclude this section by showing how these results can be used to find the value of K that minimizes $E[G_d]$, subject to a constraint on $E[T_d]$.

Under the above assumptions it is easily seen that the system can be modeled as a finite-state absorbing Markov chain $\mathbf{M}_{Kc} = \{I(t), t \geq 0\}$, where $\mathbf{M}_{Kc} \in \{1, 2, \dots, K\}$ gives the number of copies of the packet at time t if $t < T_d$, and $I(t) = a$ if $t \geq T_d$. Thus, the state a is the absorbing state of $\mathbf{M}_{Kc}$. Let $\mathbf{P}_{Kc} = [p_{Kc}(i, j)]$ be the infinitesimal generator of $\mathbf{M}_{Kc}$. From the transition rate diagram of $\mathbf{M}_{Kc}$ in Figure 6.1 we readily find

$$\begin{aligned} p_{Kc}(i, i+1) &= (N-i)\lambda, \quad i = 1, \dots, K-1, \\ p_{Kc}(i, i-1) &= (i-1)\mu, \quad i = 2, \dots, K, \\ p_{Kc}(i, i) &= -[N\lambda + (i-1)\mu], \quad i = 1, \dots, K-1, \\ p_{Kc}(K, K) &= -[K\lambda + (K-1)\mu], \\ p_{Kc}(i, a) &= i\lambda, \quad i = 1, \dots, K, \\ p_{Kc}(i, j) &= 0, \text{ otherwise.} \end{aligned}$$

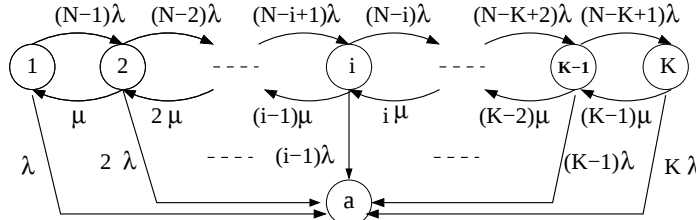


Figure 6.1: Transition diagram of the absorbing Markov chain $\mathbf{M}_{Kc}$.

The infinitesimal generator of $\mathbf{M}_{Kc}$ can be written as

$$\mathbf{P}_{Kc} = \left(\begin{array}{c|c} \mathbf{Q}_{Kc} & \mathbf{R}_{Kc} \\ \hline \mathbf{0} & 0 \end{array} \right), \quad (6.1)$$

where $\mathbf{Q}_{Kc} = [p_{Kc}(i, j)]_{1 \leq i, j \leq K}$, $\mathbf{R}_{Kc} = (p_{Kc}(1, a), \dots, p_{Kc}(K, a))^T$, and $\mathbf{0}$ is the row vector of dimension $K+1$ whose all components are equal to 0.

We will show below that for any initial state $I(0)$, $E[T_d^n]$, $P(C_d = j)$ and $E[G_d]$ can be derived in closed-form if one has a closed-form expression for $\mathbf{Q}_{Kc}^{-1}$, the inverse of the matrix $\mathbf{Q}_{Kc}$ that will be determined in closed-form. Let $q_{Kc}(i, j)$ be the (i, j) -entry of $\mathbf{Q}_{Kc}^{-1}$.

We now derive $\mathbf{Q}_{Kc}^{-1}$ in closed-form. We note that $\mathbf{Q}_{Kc}$ can be decomposed as

$$\mathbf{Q}_{Kc} = \hat{\mathbf{Q}}_K + buu^T, \quad (6.2)$$

where $\hat{\mathbf{Q}}_K$ is the K -by- K sub-matrix composed of the first K rows and columns of the generator $\mathbf{Q}$ when $K = N$, $u = (0, \dots, 0, 1)^T$ and $b = \lambda(N - K)$. By applying the Sherman-

Morrison formula [PFTV92] we find that

$$\begin{aligned}\mathbf{Q}_{Kc}^{-1} &= \hat{\mathbf{Q}}_K^{-1} - \frac{b}{1 + bu^T \hat{\mathbf{Q}}_K^{-1} u} \hat{\mathbf{Q}}_K^{-1} u u^T \hat{\mathbf{Q}}_K^{-1} \\ &= \hat{\mathbf{Q}}_K^{-1} - \frac{b}{1 + b\hat{q}_K(K, K)} \mathbf{A}_K,\end{aligned}$$

where $\hat{q}_K(i, j)$ is the (i, j) -entry of the matrix $\hat{\mathbf{Q}}_K^{-1}$, and $\mathbf{A}_K$ is a K -by- K matrix with (i, j) -entry equal to $\hat{q}_K(i, K)\hat{q}_K(K, j)$. Equivalently,

$$q_{Kc}(i, j) = \hat{q}_K(i, j) - \frac{\lambda(N - K)\hat{q}_K(i, K)\hat{q}_K(K, j)}{1 + \lambda(N - K)\hat{q}_K(K, K)}. \quad (6.3)$$

It remains to find the entries $\hat{q}_K(i, j)$ of $\hat{\mathbf{Q}}_K^{-1}$ in closed-form.

The matrix $-\hat{\mathbf{Q}}_N^{-1}$ was obtained in closed-form in Lemma 2 of Chapter 5 and it was denoted by $\mathbf{Q}^*$ with (i, j) -entry $q^*(i, j)$. On the other hand, simple algebra shows that the (i, j) -entry of $\hat{\mathbf{Q}}_K^{-1}$ is related to the (i, j) -entry of $\hat{\mathbf{Q}}_N^{-1}$, equal to $-q^*(i, j)$, as follows:

$$\hat{q}_K(i, j) = -q^*(i, j) + \frac{\lambda(N - K)q^*(i, K)q^*(K + 1, j)}{1 + \lambda(N - K)q^*(K + 1, K)} \quad (6.4)$$

for $1 \leq i \leq K$ and $1 \leq j \leq K$, where $q^*(i, j)$ is defined by (see Lemma 2 Chapter 5)

$$q^*(i, j) = -\frac{1}{\mu} \frac{1}{\binom{N-1}{i-1} \rho^{i-1}} \sum_{k=1}^N \frac{\Psi_i^k \Psi_j^k}{z_k \Psi^k \tau^2(\Psi^k)^T}, \quad (6.5)$$

where $\rho := \lambda/\mu$,

$$z_k = \frac{-N(2\rho + 1) + 1 - (N + 1 - 2k)\sqrt{4\rho + 1}}{2},$$

for $k = 1, 2, \dots, N$, $\Psi^k = (\Psi_1^k, \dots, \Psi_N^k)$ with

$$\begin{aligned}\Psi_i^k &= \sum_{l=l_0}^{l_1} \binom{k-1}{l} \binom{N-k}{i-1-l} (-1)^{i-1} \left(\frac{x_1}{x_2}\right)^{k-1-l} x_2^{N-i}, \\ x_1 &= \frac{-1 - \sqrt{1+4\rho}}{2\rho}, \quad x_2 = \frac{-1 + \sqrt{1+4\rho}}{2\rho},\end{aligned}$$

for $l_0 = \max(0, i - 1 - N + k)$ and $l_1 = \min(i - 1, k - 1)$, and where $\tau = (\tau_1, \dots, \tau_N)$ with $\tau_i = \left(\binom{N-1}{i-1} \rho^{i-1}\right)^{-1/2}$.

Equations (6.5), (6.4), and (6.3) together give the closed-form expression of $q_{Kc}(i, j)$. Therefore, the matrix $\mathbf{Q}_{Kc}^{-1}$ exists and that state a is the only absorbing state of $\mathbf{M}_{Kc}$, see Lemma 1.a Chapter 5.

6.3.1 Delivery delay

Given that there are i copies of the packet at time $t = 0$ (i.e. $I(0) = i$) the expected delivery delay of the packet, which is by definition equal to the expected time of $\mathbf{M}_{Kc}$ before absorption in state a , is given by (see Lemma 1.b, Chapter 5)

$$E_i[T_d] = -\alpha_i \mathbf{Q}_{Kc}^{-1} \mathbf{e} = -\sum_{j=1}^K q_{Kc}(i, j), \quad (6.6)$$

where $\mathbf{e}$ is K -dimensional column vector with all entries equal to 1, and α_i is a K -dimensional row vector with all entries equal to 0 except the i th one that is equal to 1. More generally, (see Lemma 1.b, Chapter 5)

$$E_i[T_d^n] = (-1)^n n! (\alpha_i \mathbf{Q}_{Kc}^{-n} \mathbf{e}) = (-1)^n n! \sum_{j=1}^K q_{Kc}^{(n)}(i, j) \quad (6.7)$$

for $i = 1, \dots, K$, where $q_{Kc}^{(n)}(i, j)$, the (i, j) -entry of $\mathbf{Q}_{Kc}^{-n}$, can be expressed in closed-form in terms of the entries of $\mathbf{Q}_{Kc}^{-1}$ which are given in (6.3)-(6.5).

6.3.2 Number of copies at delivery instant

Given that there are i copies of the packet at time $t=0$, the probability distribution of the number copies just at delivery time is (Hint: split the absorbing state a into K absorbing states $a_1, \dots, a_K$ and compute the probability of absorption in state a_i – see Section 5.3.2 for details)

$$P_i[C_d = j] = -j \lambda q_{Kc}(i, j) \quad (6.8)$$

for $1 \leq i, j \leq K$. In particular, the expected number of copies at delivery time given that $I(0) = 1$ is

$$E[C_d] = \sum_{j=1}^K j P_1[C_d = j] = -\lambda \sum_{j=1}^K j^2 q_{Kc}(i, j). \quad (6.9)$$

6.3.3 Number of transmissions before absorption

Given that $I(0) = i$, the expected total number of transmissions before absorption is given by (use (5.15) with $m(i, k) := (\lambda N + (k-1)\mu)q_{Kc}(i, k)$ for $1 \leq k \leq K-1$ and $m(i, K) := (\lambda K + (K-1)\mu)q_{Kc}(i, K)$)

$$\begin{aligned} E_i[G_d] &= \frac{1}{2} \left[\sum_{k=1}^{K-1} (\lambda N + (k-1)\mu) q_{Kc}(i, k) + (K\lambda + (K-1)\mu) q_{Kc}(i, K) + E_i[C_d] - i - 1 \right] \\ &= \frac{1}{2} \left[(\lambda N - \mu) E_i[T_d] + E_i[C_d] + \lambda(N - K) q_{Kc}(i, K) + \frac{1}{\rho} - i - 1 \right], \end{aligned} \quad (6.10)$$

where the second equality follows from (6.6)-(6.8).

6.3.4 Minimizing the consumed energy

The total number of copies that are transmitted is directly related to the energy consumed to transmit one packet. Our objective now is to use the results obtained previously to find the value of K (recall that K is the maximum number of copies in the network at any time) which minimizes the total number of copies that are transmitted before the delivery of the packet to the destination, subject to a constraint on the expected delivery delay, that is,

$$\min_{\{K: E^{(K)}[T_d] \leq C\}} E^{(K)}[G_d],$$

where $E^{(K)}[T_d] := E_1[T_d]$ and $E^{(K)}[G_d] := E_1[G_d]$. The superscript (K) emphasizes the dependency in the variable K . Since the integer mappings $K \rightarrow E^{(K)}[T_d]$ and $K \rightarrow E^{(K)}[G_d]$ are strictly decreasing and strictly increasing, respectively, the solution to this constrained optimization problem is obtained for the smallest integer K in $[1, N]$ such that $E^{(K)}[T_d] \leq C$ if $E^{(N)}[T_d] \leq C$, and has no solution if $E^{(N)}[T_d] > C$.

Figure 6.2 reports the optimal value of K (called K_{opt}) for $50 \leq N \leq 200$ with $\rho = 1$ and for three values of λ ; Figure 6.3 reports K_{opt} when the constraint C lies in $(0, 300)$, for $\rho = 1$ and for three values of λ . Figure 6.3 shows that K_{opt} is inversely proportional to C and is proportional to λ .

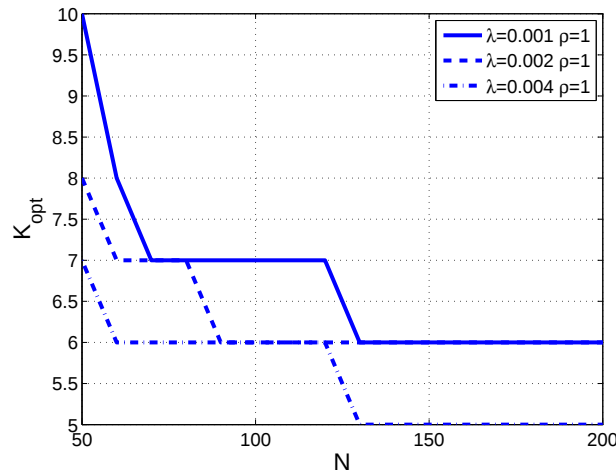


Figure 6.2: The optimal maximum number of copies in the network (K_{opt}) as a function of the number of nodes (N) for $C = E^{(50)}[T_d] + 10$ (sec.).

6.4 K-transmissions MTR protocol

In this section we evaluate the K-transmissions MTR that limits the energy consumption by limiting the number of copies that the source can generate before the packet reaches the destination. We now assume that the source can generate at most K copies of the packet,

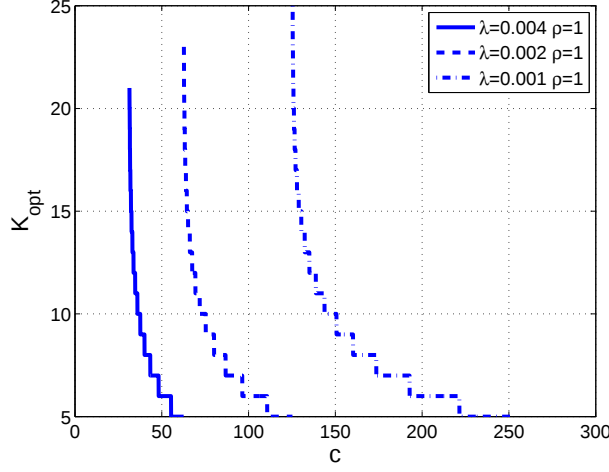
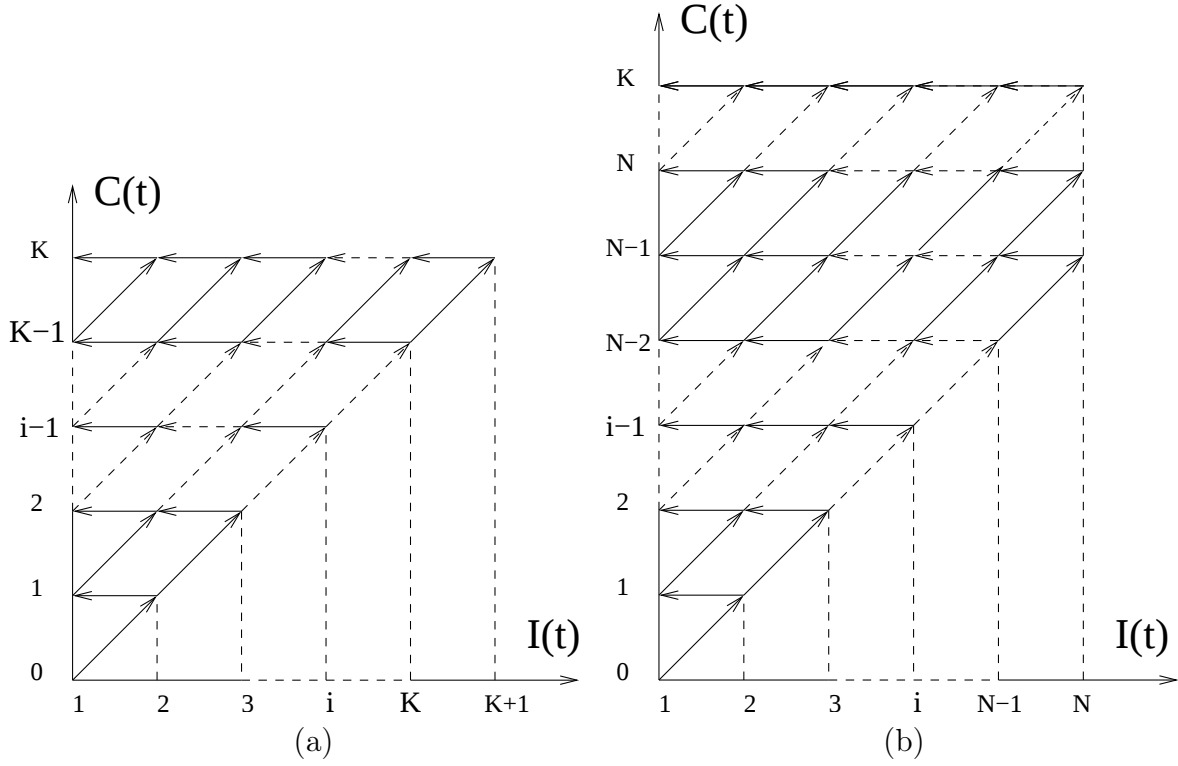


Figure 6.3: The optimal maximum number of copies in the network (K_{opt}) as a function of constraint (C) on the expected delivery delay ($E[T_d]$) for $N = 100$ and $C \in [E^{(100)}[T_d], 2E^{(100)}[T_d]]$.

where K is an arbitrary integer. The node inter-meeting times are exponentially distributed with parameter $\lambda > 0$ and the TTLs are exponentially distributed with parameter $\mu > 0$. At time $t = 0$, the source is ready to transmit the packet and it is the only node that carries the packet. Alike the K -copies MTR protocol in Section 6.3, we will compute the n^{th} order-moment of the delivery delay, $E[T_d^n]$, and the distribution of the total number of copies generated, $P(G_d = c)$, at the delivery time.

Under the above assumptions the behavior of the K -transmissions MTR can be modeled as a two-dimensional, finite-state, absorbing Markov chain $\mathbf{M}_{Ktx} = \{(I(t), C(t)), t \geq 0\}$ with state (i, c) , where $i \in \{1, 2, \dots, N\}$ gives the number of copies in the network, and $c \in \{0, 1, \dots, K\}$ records the total number of copies generated by the source for $t < T_d$, and $(I(t), C(t)) = a$ for $t \geq T_d$. The initial state of $\mathbf{M}_{Ktx}$ is $(I(0), C(0)) = (1, 0)$. In this case for $i \in \{1, \dots, N\}$ c should lie in $\{i - 1, \dots, K\}$. Figure 6.4 displays the transition rate diagram of $\mathbf{M}_{Ktx}$ in the two cases: (a) $K \leq N - 1$ and (b) $K \geq N$. In Figure 6.4 the x-axis represents $I(t)$ and the y-axis represents $C(t)$ for $t < T_d$, i.e. before absorption. Note that the absorption is possible to occur with transition rate $i\lambda$ in any state (i, c) of Figure 6.4. The infinitesimal generator matrix, $\mathbf{P}_{Ktx} = [p_{Ktx}((i, c), \cdot)]$, of $\mathbf{M}_{Ktx}$ is given by

$$\begin{aligned}
 p_{Ktx}((i, c), (i + 1, c + 1)) &= (N - i)\lambda, \quad 1 \leq i \leq K_m \text{ and } i - 1 \leq c \leq K - 1, \\
 p_{Ktx}((i, c), (i - 1, c)) &= (i - 1)\mu, \quad 2 \leq i \leq K_m \text{ and } i - 1 \leq c \leq K - 1, \\
 p_{Ktx}((i, c), a) &= i\lambda, \quad 1 \leq i \leq K_m \text{ and } i - 1 \leq c \leq K - 1, \\
 p_{Ktx}((i, c), (i, c)) &= -(N\lambda + (i - 1)\mu), \quad 1 \leq i \leq K_m \text{ and } i - 1 \leq c \leq K - 1, \\
 p_{Ktx}((N, c), (N - 1, c)) &= (N - 1)\mu \mathbf{1}_{\{K \geq N\}}, \quad N - 1 \leq c \leq K - 1,
 \end{aligned}$$


 Figure 6.4: Transition rate diagram of $\mathbf{M}_{Ktx}$: (a) $K \leq N - 1$, and (b) $K \geq N$.

$$\begin{aligned}
 p_{Ktx}((N, c), a) &= N\lambda \mathbf{1}_{\{K \geq N\}}, \quad N - 1 \leq c \leq K - 1, \\
 p_{Ktx}((N, c), (N, c)) &= -(N\lambda + (N - 1)\mu) \mathbf{1}_{\{K \geq N\}}, \quad N - 1 \leq c \leq K - 1, \\
 p_{Ktx}((i, K), (i - 1, K)) &= (i - 1)\mu, \quad 2 \leq i \leq K_m + 1, \\
 p_{Ktx}((i, K), a) &= i\lambda, \quad 1 \leq i \leq K_m + 1, \\
 p_{Ktx}((i, K), a) &= -(i\lambda + (i - 1)\mu), \quad 1 \leq i \leq K_m + 1, \\
 p_{Ktx}((i, c), (j, d)) &= 0, \quad \text{otherwise,}
 \end{aligned}$$

with $K_m := \min(K, N - 1)$, and where a is the absorbing state. Let L denote the cardinality of the state space of $\mathbf{M}_{Ktx}$ excluding the absorbing state a , then

$$L = \frac{1}{2} \left[(K + 1)(K + 2) \mathbf{1}_{\{K \leq N - 1\}} + N(2K - N + 3) \mathbf{1}_{\{K \geq N\}} \right]. \quad (6.11)$$

If we label the transient states $(1, 0)$ as 1, $(2, 1)$ as 2, $\dots$, (i, c) as $\frac{(c+1)(c+2)}{2} - i + 1$ for $c \leq K_m$ and $i \leq c + 1$, $\dots$, (i, c) as $\frac{N(2c - N + 1)}{2} + N - i + 1$ for $c > K_m$ and $i \leq N$, then we can write the matrix $\mathbf{P}_{Ktx}$ as

$$\mathbf{P}_{Ktx} = \left(\begin{array}{c|c} \mathbf{Q}_{Ktx} & \mathbf{R}_{Ktx} \\ \hline \mathbf{0} & 0 \end{array} \right),$$

where $\mathbf{Q}_{Ktx} = [p_{Ktx}((i, c), \cdot)]_{1 \leq i, c \leq L}$, $\mathbf{R}_{Ktx} = (p_{Ktx}((i, c), a))_{1 \leq i, c \leq L}$, and $\mathbf{0}$ is the 1-by- L zero matrix.

We now derive $(\mathbf{Q}_{Ktx})^{-1}$ in closed-form. $\mathbf{Q}_{Ktx}$ can be seen as an upper-bidiagonal block matrix of canonical form given by

$$\mathbf{Q}_{Ktx} = \begin{pmatrix} \mathbf{A}_0 & \mathbf{B}_0 & & & \\ & \ddots & \ddots & & \\ & & \mathbf{A}_c & \mathbf{B}_c & \\ & & & \ddots & \ddots \\ & & & & \mathbf{A}_{K-1} & \mathbf{B}_{K-1} \\ & & & & & \mathbf{A}_K \end{pmatrix}, \quad (6.12)$$

where $\mathbf{A}_c$, $0 \leq c \leq K$, is a N_c -by- N_c upper-bidiagonal matrix with $N_c := \min(c+1, N)$, and where $\mathbf{B}_c$, $0 \leq c \leq K-1$, is a N_c -by- N_{c+1} diagonal matrix. The matrix $\mathbf{A}_c$ represents the rate of transition between the states of the subspace $\{(i, c) : 1 \leq i \leq N_c\}$, and $\mathbf{B}_c$ denotes the rate of transition of the states of $\{(i, c) : 1 \leq i \leq N_c\}$ to the states of $\{(i, c+1) : 1 \leq i \leq N_c\}$. Observe that $\mathbf{A}_c$, for $0 \leq c \leq K$, is a non-singular matrix since its is an upper-bidiagonal matrix with diagonal entries that are all strictly negative. Based on this observation we will show in the following that $\mathbf{Q}_{Ktx}$ is non-singular.

From (6.12) it is easy to show that $\mathbf{Q}_{Ktx}^{-1}$ (if exist) is an upper triangular matrix of form

$$\mathbf{Q}_{Ktx}^{-1} = \begin{pmatrix} \mathbf{A}_0^{-1} & \mathbf{U}_{0,1} & \cdots & & \mathbf{U}_{0,K} \\ & \ddots & \ddots & & \\ & & \mathbf{A}_c^{-1} & \mathbf{U}_{c,c+1} & \cdots & \mathbf{U}_{c,K} \\ & & & \ddots & \ddots & \\ & & & & \mathbf{A}_{K-1}^{-1} & \mathbf{U}_{K-1,K} \\ & & & & & \mathbf{A}_K^{-1} \end{pmatrix}, \quad (6.13)$$

where

$$\begin{aligned} \mathbf{U}_{c,l} &= -\mathbf{U}_{c,l-1} \mathbf{B}_{l-1} \mathbf{A}_l^{-1} \\ &= (-1)^{l-c} \left(\prod_{m=c}^{l-1} \mathbf{A}_m^{-1} \mathbf{B}_m \right) \mathbf{A}_l^{-1}, \end{aligned} \quad (6.14)$$

for $0 \leq c \leq K-1$ and $c+1 \leq l \leq K$. Since $\mathbf{A}_c$ is a non-singular matrix for $0 \leq c \leq K$, then $\mathbf{U}_{c,l}$ exists and the matrix $\mathbf{Q}_{Ktx}$ is non-singular. Thus, from Lemma 1.a of Chapter 5, we have that the states of $\mathbf{M}_{Ktx}$ are all transient except the state a that is absorbing. Once the matrix $\mathbf{Q}_{Ktx}^{-1}$ has been computed the main performance metrics can easily be deduced, as shown below. In the following, $q_{Ktx}(i, j)$ will denote the (i, j) -entry of $\mathbf{Q}_{Ktx}^{-1}$.

6.4.1 Expected delivery delay

The expected delivery delay is by definition the expected time to absorption which is given by (Chapter 5, Lemma 1.c)

$$\overline{T_d} = - \sum_{i=1}^L q_{Ktx}(1, i). \quad (6.15)$$

More generally, the n^{th} order-moment of T_d reads (Chapter 5, Lemma 1.b)

$$E[T_d^n] = (-1)^n n! (\alpha_1 \mathbf{Q}_{Ktx}^{-n} \mathbf{e}) = (-1)^n n! \sum_{j=1}^L q_{Ktx}^{(n)}(1, j) \quad (6.16)$$

where $\mathbf{e}$ is L -dimensional column vector with all entries equal to 1, α_1 is a L -dimensional row vector with all entries equal to 0 except the first one that is equal to 1, and $q_{Ktx}^{(n)}(i, j)$, the (i, j) -entry of $\mathbf{Q}_{Ktx}^{-n}$, can be expressed in closed-form in terms of the entries of $\mathbf{Q}_{Ktx}^{-1}$.

6.4.2 Distribution of the number of copies generated

Let G_d denote the rv that represents the number of copies generated by the source before the delivery time. The probability density of G_d , $P(G_d = c)$, is the probability that the last state visited before absorption is in the subspace $\{(i, c) : 1 \leq i \leq N_c\}$ with $N_c = \min(c + 1, N)$. In the following, to find $P(G_d = c)$, we will first compute the probability that the last state visited before absorption is (i, c) , and then we will sum this probability over all the possible values of i .

If we split the absorbing state a into L absorbing states $a_1, \dots, a_L$ where L is the cardinality of the state space of $\mathbf{M}_{Ktx}$ (excluding the absorption state). Then, the probability that the absorption occurs in state a_l , where $l = l_1 := 1 - i + (c + 1)(c + 2)/2$ for $K \leq N - 1$ and $l = l_2 := N - i + 1 + N(2c - N + 1)/2$ for $K \geq N$, is equal to the probability that the last state visited before absorption is (i, c) . On the other hand the probability that the last state visited before absorption is (i, c) can be written as (see Section 5.3.2 for details)

$$P(\text{absorption in } (i, c)) = -\lambda i q_{Ktx}(1, l), \quad (6.17)$$

where $l = l_1 \mathbf{1}_{\{K \leq N-1\}} + l_2 \mathbf{1}_{\{K \geq N\}}$. The sum of the above probability over all the possible values of i gives the probability density of G_d that reads

$$P(G_d = c) = -\lambda \sum_{i=1}^{N_c} i (q_{Ktx}(1, l_1) \mathbf{1}_{\{K \leq N-1\}} + q_{Ktx}(1, l_2) \mathbf{1}_{\{K \geq N\}}), \quad (6.18)$$

where $N_c = \min(c + 1, N)$.

The expected value of G_d is equal to

$$\overline{G_d} := -\lambda \sum_{c=0}^K c \sum_{i=1}^{N_c} i(q_{Ktx}(1, l_1) \mathbf{1}_{K \leq N-1} + q_{Ktx}(1, l_2) \mathbf{1}_{K \geq N}). \quad (6.19)$$

Based on the energy consumption model introduced in Section 6.2, the energy consumed by the source before the packet is transmitted to the destination is $p_t \overline{G_d}$ while the energy consumed by all nodes during this period is $(p_t + p_d) \overline{G_d}$.

6.4.3 Numerical evaluation

For different values of λ , the inter-meeting time rate, Figure 6.5 plots the expected delivery time, $\overline{T_d}$, under the K -limited MTR protocol for different values of K , the maximum number of copies that the source may generate. For each λ , we observe there exists a threshold K_0 such that $\overline{T_d}$ is almost constant when $K \geq K_0$ ($K_0 \sim 20$ for $\lambda = 0.001$). Also, this constant value of $\overline{T_d}$ when $K \geq K_0$ is nothing than the mean delivery delay obtained in (5.10) of the original MTR protocol, i.e. for $K = \infty$. Figure 6.6 plots the expected delivery time, $\overline{G_d}$, under the K -limited two-hop relay protocol for different values of K . Similarly to $\overline{T_d}$, it exists a threshold K_0 such that $\overline{G_d}$ is almost constant when $K \geq K_0$. Further, Figure 6.6 shows that for different value of λ the asymptotic of $\overline{G_d}$ as K increases is almost independent of λ . This observation supports the result of $\overline{G_d}$ found in Section 5.5.3 that is approximatively equal to $\sqrt{\frac{\pi N}{2}}$ for N large.

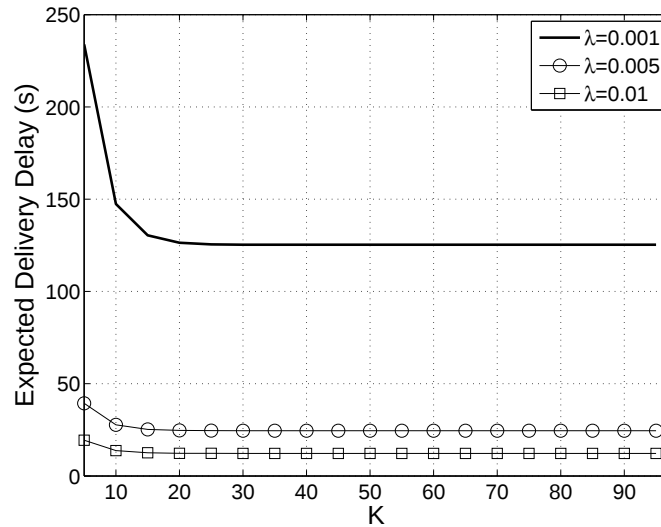


Figure 6.5: Expected delivery delay under K -transmissions MTR protocol for different values of K ($N=100$, $\mu=0.001$).

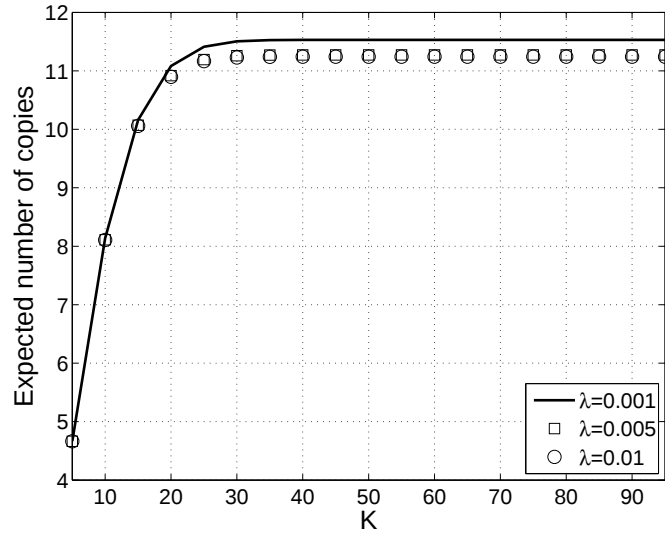


Figure 6.6: Expected total number of copies generated before absorption under K -transmissions MTR protocol for different values of K ($N=100$, $\mu=0.001$).

6.4.4 Asymptotic analysis

We conclude this section by studying the behavior of $\mathbf{M}_{Ktx}$ for large value of N and finite value of $K < N$. Let consider the path, MP, that joins the states $(1, 0)$ and $(K + 1, K)$, see Figure 6.4.a. Based on this Figure, MP is the unique path that joins these two states. The probability of MP reads

$$P_{MP} = \prod_{i=1}^K \frac{(N-i)\lambda}{N\lambda + (i-1)\mu}, \quad (6.20)$$

so that

$$\lim_{N \rightarrow \infty} P_{MP} = \prod_{i=1}^K \lim_{N \rightarrow \infty} \frac{(N-i)\lambda}{N\lambda + (i-1)\mu} = 1. \quad (6.21)$$

This shows that as N is large the system has a deterministic path MP *w.p.1*. In this case the state space of $\mathbf{M}_{Ktx}$ is reduced to $\{(i, i-1) : 1 \leq i \leq K\} \cup \{(i, K) : 1 \leq i \leq K+1\} \cup \{a\}$, where $\{a\}$ is the absorption state. Figure 6.7 displays the transition rate diagram of $\mathbf{M}_{Ktx}$ for N large, referred to as $\mathbf{M}_{Ktx}^{(N)}$. Conditioned on the last state before absorption, it is easy to show that the Laplace Stieltjes of T_d of $\mathbf{M}_{Ktx}^{(N)}$ reads

$$\begin{aligned} E[e^{-sT_d}] &= (N-1)! \sum_{i=1}^K \frac{i}{(N-i)!} \left(\frac{\lambda}{N\lambda + s} \right)^i \\ &+ \frac{(N-1)!K!}{(N-K-1)!} \frac{(\lambda\mu)^{K+1}}{(N\lambda + s)^K} \sum_{i=1}^{K+1} \frac{i^2}{i!\mu^i} \prod_{j=i}^{K+1} \frac{1}{j\lambda + (j-1)\mu + s}, \end{aligned} \quad (6.22)$$

and the probability density of G_d is given by

$$\begin{aligned}
 P(G_d = c) &= \frac{(N-1)!(c+1)}{(N-c-1)!N^{c+1}}, \quad 0 \leq c \leq K-1. \\
 P(G_d = K) &= \frac{(N-1)!}{(N-K-1)!N^K}
 \end{aligned} \tag{6.23}$$

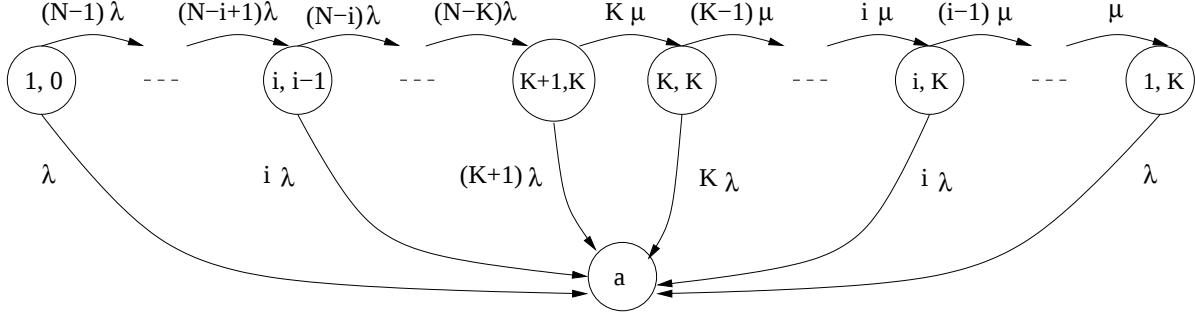


Figure 6.7: Transition rate diagram of $M_{Ktx}^{(N)}$.

6.5 Two-hop relay with Erasure Coding

We now consider a system where the source introduces some redundancy in its transmissions and sends more data than the actual information. The advantage of this mechanism is that it can considerably reduce the variance of the packet delivery delay at the expense of an increase of the expected delivery delay.

One of these techniques is known as erasure coding [Mit04]. Erasure coding with replication factor r works as follow. Upon receiving a packet of size M , the source produces $n = r \cdot M/b$ equal sized code blocks of size b , such that any of the $k = (1 + \epsilon) \cdot M/b$ code blocks can be used to reconstruct the packet. Here ϵ is a small constant, close to zero, that depends on the coding/decoding algorithm used [Mit04]. Thus, the destination is able to retrieve that packet if it receives $k < n$ blocks. On the other hand, when $k = 1$, the size of a block becomes almost equal to M , the packet size, and in this case the destination needs to receive one block in order to decode the packet. Thus for $k = 1$, the erasure coding scheme is the same as a simple multicopy scheme in which the source sends exactly one copy of a packet to n different relay nodes [WJMF05]. We will exploit this observation to compare erasure coding with the multicopy mechanism in Section 6.5.2.

Throughout this section the stochastic model is the following. There are N relay nodes, one source node, and one destination node. We assume that the source cannot send directly a packet to the destination. Inter-meeting times between any pair of nodes are exponentially distributed with rate $\lambda > 0$, except for the pair source-destination. Under this setting, the only way to forward the data from the source to the destination is through the relay nodes. The source has only one packet to send to the destination, and the source

implements the erasure coding algorithm with replication factor r and parameter k . Hence, the destination needs to receive $k < n$ of the blocks in order to decode the original packet. The forwarding mechanism used to deliver the blocks to the destination is the standard two-hop relay protocol. We assume that there is only one copy of a block in the network, which is either carried by the source or by a relay node that receives it after meeting the destination. A relay node can only relay one block at a time, and it is possible for a relay node that already delivered a block to the destination to receive a new block when it again encounters the source. There is no TTL associated with the blocks.

Let T_d and G_d be the delivery delay and the total number of source-relay transmissions at the time when k th block reaches the destination, respectively. It is worth pointing out that since we have assumed that transmissions are instantaneous, T_d gives a lower bound on the delivery delay obtained in a realistic network. Introduce the joint transform

$$H(s, z) := E[e^{-sT_d} z^{G_d}], \quad s \geq 0, |z| \leq 1.$$

We now evaluate $H(s, z)$.

Let $A(t) = (B(t), R(t))$ denote the two-dimensional process such that $A(t) = (m, l)$, $1 \leq m \leq n$, $1 \leq l \leq k-1$, $m+l \leq n$, if there are m relay nodes that hold m blocks (one block for relay node) and the destination has received l blocks at time $t < T_d$, and $A(t) = a$ when $t \geq T_d$ (a is an absorbing state). Under the above assumptions, $\{A(t), t \geq 0\}$ is a finite-state absorbing Markov process. Figure 6.8 displays the transition diagram of this Markov chain, where the y-axis represents $R(t)$ and the x-axis represents the sum $B(t) + R(t)$. More precisely, a point (i, l) , $i \geq l$, in the transition diagram means that the destination has received l blocks and that there are $i - l$ relay nodes that hold $i - l$ blocks (one block per each).

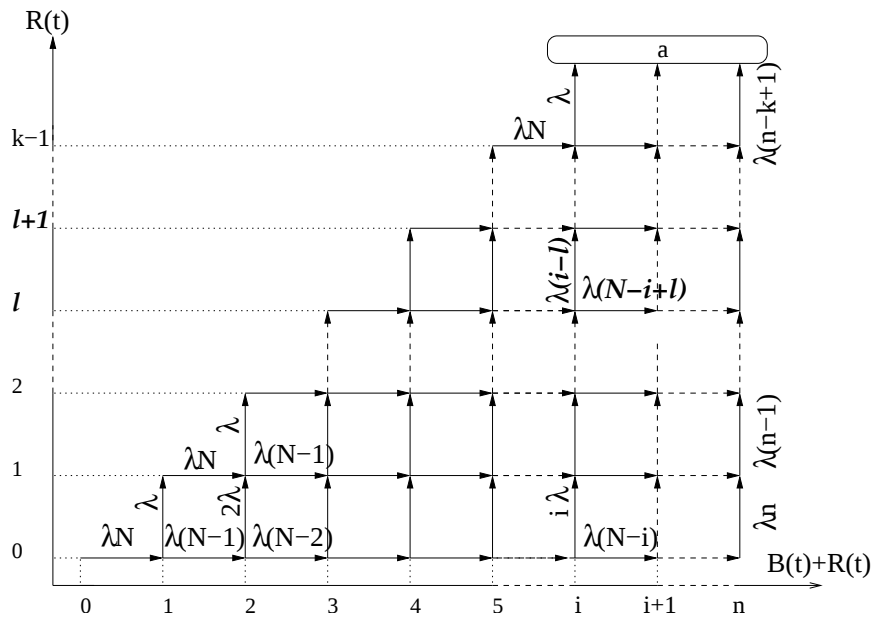


Figure 6.8: Transition diagram of the Markov chain $\{A(t), t \geq 0\}$.

6.5.1 Computation of $H(s, z)$

Let $j_i \geq 0$ denote the number of jumps (transitions) along the horizontal line of index $i \in \{0, \dots, k-1\}$. Let S_i denote the total number of steps along the lines of index less than or equal to i , namely,

$$S_i = \sum_{l=0}^i j_l,$$

where by convention $S_{-1} = 0$. It is clear from the transition diagram that $j_i \leq n - S_{i-1}$. Given S_{i-1} , the probability of making j_i jumps along the horizontal line i is

$$P(L_i = j_i) = \begin{cases} P_1(j_i) &= \frac{(N-S_{i-1}+i)!}{(N-S_i+i)!} \frac{S_i-i}{N^{j_i+1}} & , \quad S_i < n \\ P_2(j_i) &= \frac{(N-S_{i-1}+i)!}{(N-n+i)!} \frac{1}{N^{j_i}} & , \quad S_i = n, \end{cases} \quad (6.24)$$

for $0 \leq i \leq k-1$. Let m^* denote the index of the horizontal line such that $S_{m^*-1} < n$ and $S_{m^*} = n$. Note that when m^* exists, the probability of any path from the point $(0, 0)$ to the absorption state a gives

$$P_2(j_{m^*}) \prod_{i=0}^{m^*-1} P_1(j_i). \quad (6.25)$$

In the case where m^* does not exist, the probability of any path from the point $(0, 0)$ to the absorption state a reads

$$\prod_{i=0}^{k-1} P_1(j_i). \quad (6.26)$$

Conditioning on all possible paths, the joint transform of the total amount of time spent before the delivery of the k^{th} block of the packet and the number of source-relay transmissions in this process is then

$$\begin{aligned} H(s, z) &= \sum_{j_0=1}^{n-1} \cdots \sum_{j_{k-1}=0}^{n-1} \left(\frac{zN\lambda}{s + N\lambda} \right)^{S_{k-1}+k} \times \prod_{l=0}^{k-1} P_1(j_l) \\ &+ \sum_{j_0=1}^n \cdots \sum_{j_{m^*-1}=0}^{n-S_{m^*-2}} \left(\frac{zN\lambda}{s + N\lambda} \right)^{n+m^*} \prod_{l=m^*}^{k-1} \frac{\lambda(n-l)}{s + \lambda(n-l)} \times P_2(n - S_{m^*-1}) \prod_{l=0}^{m^*-1} P_1(j_l). \end{aligned} \quad (6.27)$$

6.5.2 Erasure coding vs simple multicopy scheme

We now evaluate the expectation and the variance of T_d for different values of k , n , and N . Let $\sigma_{T_d} := \sqrt{\text{var}(T_d)}/E[T_d]$ denote the normalized standard deviation of T_d . As noted earlier in this section, erasure coding when $k = 1$ is similar to a simple multicopy scheme with maximum number of transmissions equal to n . Table 6.1 shows that $E[T_d]$ increases

(N, n)	$(20, 10)$			$(40, 10)$		
k	1	2	5	1	2	5
$E[T_d]$ (sec.)	31.51	42.04	96.69	22	33.7	79.76
σ_{T_d}	8.7	4.9	1.17	2.01	1.38	0.63

Table 6.1: Erasure coding ($k > 1$) vs multicopy scheme ($k = 1$) for different value of n and N and $\lambda = 0.01$.

with k and that the normalized standard deviation decreases with k . For instance, for $N = 20$ and $n = 10$ when k increases from 1 to 5 the delivery delay increases by a factor of 3 while the normalized standard deviation decreases by a factor of 7.5, thereby showing that erasure coding has a much lower variability than a simple multicopy scheme. A similar result was found in [WJMF05] under the assumption that at time 0 the source instantaneously transmits all of its n blocks to n different relay nodes.

6.5.3 Asymptotic analysis

We conclude this section by investigating the behavior of $T_N^*(s) := E[e^{-sT_d}]$, the LST of T_d , as N is large. Figure 6.9 gives the evolution of the most probable path as N increases for given values of n and k . It is clear from Figure 6.9 that as N becomes large the most probable path is the one where all n blocks are first transmitted to n relay nodes, and then these relay nodes starts to deliver these blocks to the destination. The probability of the most probable path (called MP) for large N is

$$P_{MP} = \prod_{i=1}^{n-1} \frac{N-i}{N}, \quad (6.28)$$

so that

$$\lim_{N \rightarrow \infty} P_{MP} = \lim_{N \rightarrow \infty} \prod_{i=1}^{n-1} \frac{N-i}{N} = \prod_{i=1}^{n-1} \lim_{N \rightarrow \infty} \frac{N-i}{N} = 1.$$

This implies that as N is large the system has a deterministic path MP *w.p.1*. In this case we easily see that

$$T_N^*(s) \approx \left(\frac{N\lambda}{s + N\lambda} \right)^n \prod_{l=0}^{k-1} \frac{\lambda(n-l)}{s + \lambda(n-l)} \quad (6.29)$$

as N is large. Figure 6.10 shows the rate of convergence of P_{MP} to 1 as the number of nodes N increases and for different values of n , $k = 5$.

6.6 Concluding remarks

In this chapter, we have evaluated the performance of a class of two-hop relay protocols for mobile ad hoc networks. The interest is on the multicopy two-hop relay protocol, where

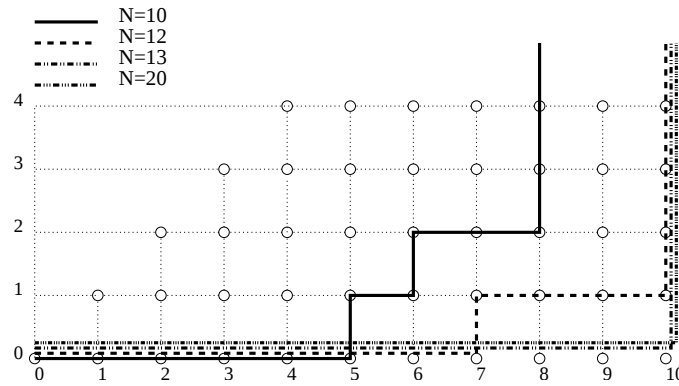


Figure 6.9: Most probable path of $\{A(t), t \geq 0\}$ when $n = 10$ and $k = 5$.

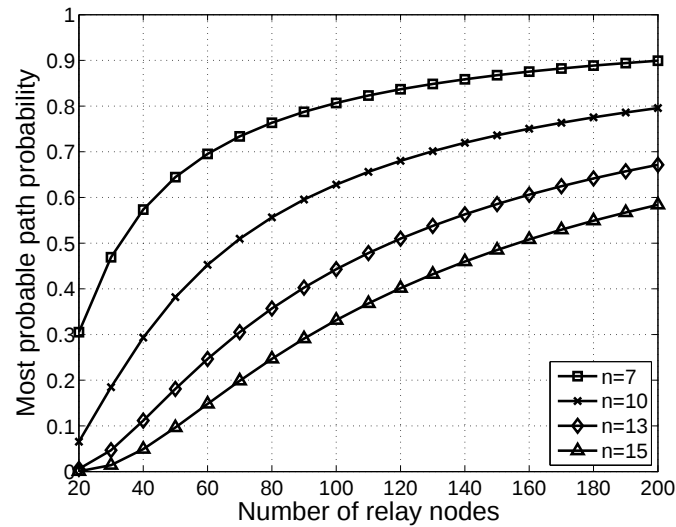


Figure 6.10: Evolution of the probability of the most probable path as function of N for different values of n and in the case where $k = 5$.

the source may generate multiple copies of a packet, and on the two-hop relay protocol with erasure coding, where a piece of information is fragmented into n blocks. Closed-form expressions of the n^{th} order-moment of the time to deliver a packet to its destination, the distribution of the number of copies of the packet at delivery instant, and of the expected total number of copies that the source generates; the latter number will be larger than the former when time-to-live (TTLs) are associated with the copies of a packet, a situation that we addressed. Finally, we characterized the delivery delay in the two-hop relay protocol with erasure coding and compare its performance with the multicopy two-hop relay protocol. We found that the delivery delay in the case of erasure coding has much lower dispersion than the delivery delay of the multicopy scheme.

In the following chapter, we will investigate the impact of constant TTLs and arbitrarily inter-meeting times on the performance of the multicopy two-hop relay protocol.

Chapter 7

Impact of Constant TTL and Arbitrary Inter-meeting on MTR

In this chapter, we evaluate the impact of constant TTLs as opposed to exponential TTLs, and we develop an approximation analysis in the case where the inter-meeting times are arbitrarily distributed. In particular, we show that exponential inter-meeting times yield stochastically smaller delivery delays than hyper-exponential inter-meeting times, and that exponential TTLs yield stochastically larger delivery delay than constant TTLs.

7.1 Introduction

This chapter evaluates the impact of constant TTLs and arbitrary inter-meeting times on the performance of the multicopy two-hop relay (MTR) protocol. The following assumptions will be enforced throughout: All inter-meeting times are independent and identically distributed (i.i.d.) random variables (rvs) with a common cumulative probability distribution (CDF) $G(\cdot)$. The transmission of the packets occurs at the nodes meeting times and it is assumed to be instantaneous. This assumption will be justified in Delay Tolerant Networks, where the incurred delay to send a packet may be very large with respect to the transmission times [dtn].

In Section 7.2 we consider the MTR protocol with the assumption that TTLs are all constant and all equal, and inter-meeting times are exponentially distributed, i.e. $G(t) = 1 - e^{-\lambda t}$. We develop a (non-Markovian) analysis which allows us to determine the CDF of the delivery delay. We observe that constant TTLs generate stochastically smaller delivery delays than exponential TTLs.

In Section 7.3 we consider the same setting as in Section 7.2 but we relax the assumption that the inter-meeting times are exponentially distributed. We assume that they are arbitrarily distributed, namely, the CDF $G(\cdot)$ is arbitrary. We develop an approximation formula for the CDF of the delivery delay, and check its accuracy in the case where $G(\cdot)$ is the hyper-exponential distribution. We observe that the delivery delays with exponential inter-meeting times are stochastically smaller than the delivery delays with hyper-exponential inter-meeting times.

The model is the following: there are $N + 1$ mobile nodes consisting of one source node, one destination node, and $N - 1$ relay nodes. We consider the scenario where the source has a single packet to transmit to the destination. There is no TTL associated with the packet of the source. The performance metric of interest is the (packet) delivery delay T_d defined as the first time when the destination receives the original packet sent by the source or a copy sent by a relay node, whichever arrives first to the destination. The source uses multicopy two-hop relay (MTR) protocol in order to forward its packet to the destination.

7.2 Impact of constant TTL

We assume that the TTLs are *constant* and all equal to $0 < T < \infty$. The nodes inter-meeting times are exponentially distributed with rate $\lambda > 0$. As a result, the stochastic process $\mathbf{M}$ of Section 5.3 of the number of copies in the network is no longer a Markov process and a different approach has to be used in order to evaluate the performance of the protocol.

For convenience we label the nodes so that node 0 is the source, node N is the destination and nodes $1, 2, \dots, N - 1$ are the relay nodes. We have

$$T_d \stackrel{st}{=} \min(X_{sd}, D_1, \dots, D_{N-1}), \quad (7.1)$$

where X_{sd} is an exponential rv with intensity λ representing the inter-meeting time between

the source and the destination, and D_i is the time needed for relay node $i = 1, 2, \dots, N-1$ to deliver a copy of the packet to the destination. Moreover, the rvs $X_{sd}, D_1, \dots, D_{N-1}$ are mutually independent and the rvs $D_1, \dots, D_{N-1}$ are identically distributed. Hence,

$$P(T_d \leq t) = 1 - e^{-\lambda t} P(D_i > t)^{N-1}. \quad (7.2)$$

We now compute $P(D_i > t)$.

Let R be a rv representing the number of times the relay node i has received a copy of the packet before it transmits it to the destination.

Let Y be the time during which the relay node i holds the copy that it transmits to the destination. If $R = m+1$, then Y is the time during which the relay node i holds the $(m+1)$ st copy before delivering it to the destination. Clearly, $P(Y > t) = P(X_{id} > t | X_{id} < T)$ where X_{id} is a rv representing the inter-meeting time between the relay node i and the destination. By assumption, X_{id} is an exponential rv with intensity λ . Therefore,

$$P(Y > t) = \frac{e^{-\lambda t} - e^{-\lambda T}}{1 - e^{-\lambda T}} \mathbf{1}_{\{t \leq T\}}, \quad (7.3)$$

and the density function $f_Y(t)$ of Y is given by

$$f_Y(t) = \frac{\lambda e^{-\lambda t}}{1 - e^{-\lambda T}} \mathbf{1}_{\{t \leq T\}}. \quad (7.4)$$

Conditioned on $R = m+1$, we see that X_i is the sum of $m+1$ inter-meeting times which are i.i.d. and exponential distributed rvs with parameter λ , m constant TTLs of length T , and Y . Therefore,

$$P(D_i > t) = \sum_{m=0}^{\infty} P(E_{m+1} + mT + Y > t) P(R = m+1), \quad (7.5)$$

where E_k is a k -stage Erlang rv with parameter λ (i.e. each stage is an exponential rv with parameter λ).

Taking advantage of the memoryless property of the exponential distribution, it easy to see that R is a geometrically distributed rv with parameter p , where p is the probability that a relay node drops its copy before it meets the destination. Clearly $p = P(T < X_{id}) = \exp(-\lambda T)$, and

$$P(R = m) = (1 - p)p^{m-1}, \quad m = 1, 2, \dots \quad (7.6)$$

On the other hand, for all $x > 0$,

$$\begin{aligned} P(E_{m+1} + Y > x) &= \int_0^{\infty} P(E_{m+1} + y > x) f_Y(y) dy \\ &= P(Y > x) + \int_{y=0}^x P(E_{m+1} > x - y) f_Y(y) dy \\ &= P(Y > x) + \int_{y=0}^x \left(\sum_{k=0}^m e^{-\lambda(x-y)} \frac{(\lambda(x-y))^k}{k!} \right) \frac{\lambda e^{-\lambda y}}{1 - e^{-\lambda T}} \mathbf{1}_{\{y \leq T\}} dy \end{aligned}$$

$$= P(Y > x) + \frac{e^{-\lambda x}}{1-p} \sum_{k=0}^m \frac{(\lambda x)^{k+1} - (\lambda[x-T]^+)^{k+1}}{(k+1)!}, \quad (7.7)$$

where $[x]^+ := \max(0, x)$. The latter result together with (7.5) and (7.23) yields

$$P(D_i > t) = e^{-\lambda t} \Psi(\lambda, T, t) \quad (7.8)$$

where

$$\Psi(\lambda, T, t) := 1 + \sum_{m=0}^{\lfloor \frac{t}{T} \rfloor} \sum_{k=0}^m \frac{(A_m)^{k+1} - (B_m)^{k+1}}{(k+1)!} \quad (7.9)$$

and $A_m := \lambda(t - mT)$ and $B_m := \lambda[t - (m+1)T]^+$. Finally (use (C.4) and (7.8)),

$$P(T_d < t) = 1 - e^{-\lambda N t} \Psi(\lambda, T, t)^{N-1}, \quad t > 0. \quad (7.10)$$

In particular, the expected delivery delay for constant TTL is

$$E[T_d] = \int_0^\infty e^{-\lambda N t} \Psi(\lambda, T, t)^{N-1} dt. \quad (7.11)$$

For $N = 50$ and for two different values of the inter-meeting time parameter λ , Figure 7.1 displays the expected delivery delay with constant TTL (equal to T) and the expected delivery delay when the TTL is exponentially distributed with parameter $\mu = 1/T$, both as a function of $\rho = \lambda T$. We observe that the expected delivery delay is always higher with an exponential TTL than with a constant TTL. An intuitive explanation is that in the case of an exponential TTL there is high probability (equal to $1 - e^{-1} \approx 0.63$) that the sampled rv is smaller than its expected value, which in turn increases the delivery delay of the packet.

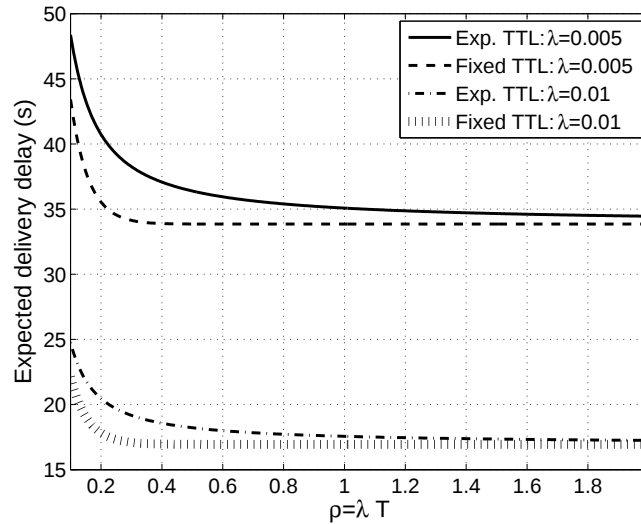


Figure 7.1: Expected delivery delay under constant and exponential TTL ($N = 50$).

Observe that the expected delivery delay under both constant and exponential TTL converges to an asymptotic value as N is large. A numerical analysis shows that this asymptotic value is approximately equal to $\frac{1}{\lambda} \sqrt{\frac{\pi}{2N}}$, and is independent of T . The same value was obtained in Section 5.5 using a mean field approach.

Figure 7.2 shows that the T_d with constant $TTL = T$ is stochastically smaller than the T_d with exponentially distributed TTL with parameter $\mu = 1/T$.

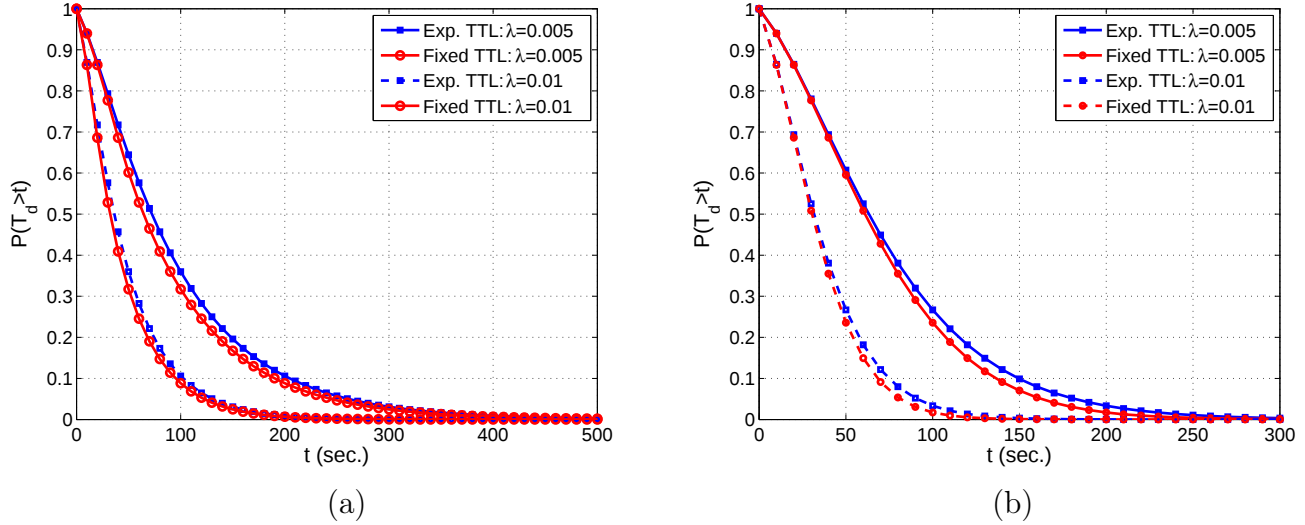


Figure 7.2: complementary CDF of T_d for the TTLs exponentially distributed and for the TTLs fixed in the case where $N = 10$: (a) $\lambda T = 0.2$, (b) $\lambda T = 1$.

Remark 5 (Model extensions) Assume that the inter-meeting times between relay node i and the source (resp. destination) are exponentially distributed with parameter λ_i and that the TTL at relay node i is constant and equal to T_i , for $i = 1, \dots, N - 1$. A similar analysis to the one above gives, for $t > 0$,

$$P(T_d < t) = 1 - e^{-\lambda t - \sum_{i=1}^{N-1} \lambda_i t} \prod_{i=1}^{N-1} \Psi(\lambda_i, T_i, t). \quad (7.12)$$

7.3 Impact of arbitrary inter-meeting time distribution

Chaintreau et al [CHC⁺06] have recently observed that the inter-meeting time distributions in a mobile ad hoc networks, where nodes are humans moving in a conference space, show a heavier-than-exponential tail. More precisely, the tail of the inter-meeting time distribution is hyperbolic in some finite range, after this finite range, the tail of the distribution exhibits a sharp decay.

This finding has motivated us to investigate the impact of arbitrary inter-meeting time distribution on the performance of MTR protocol. Throughout this section we assume that

for any pair of nodes their inter-meeting times are i.i.d. with distribution $G(t)$, and all inter-meetings are mutually independent. Let X be a generic rv with distribution G . Also define $G^*(s) = E[e^{-sX}]$ the LST of X .

We assume that TTLs are constant and all equal to T , and that there is no restriction on the number of copies of the packet in the network. We assume that node 0 is the source, node N is the destination and nodes $1, \dots, N-1$ are relay nodes.

We assume that the source, destination and relay node i are in steady-state at time $t = 0$, and that the relay node i does not hold a copy of the packet at $t = 0$ (only the source holds the original packet at $t = 0$). Similarly to Section 7.2 the delivery delay T_d is given by

$$T_d \stackrel{st}{=} \min(X_{sd}, D_1, \dots, D_{N-1}). \quad (7.13)$$

The rvs $D_1, \dots, D_{N-1}$, which have been defined in Section 7.2 are i.i.d. and independent of X_{sd} . Hence,

$$P(T_d > t) = (1 - G_e(t)) P(D_i > t)^{N-1}, \quad (7.14)$$

where $G_e(t)$ the CDF of X_{sd} that is the excess of the probability distribution $G(t)$, that is,

$$G_e(t) = \frac{1}{E[X]} \int_0^t (1 - G(u)) du. \quad (7.15)$$

We need to determine $P(D_i > t)$. We shall actually find an approximation formula for $P(D_i > t)$ since finding an exact expression is a very difficult task, unless $G(t)$ is the exponential distribution. From now on i is fixed in $\{1, \dots, N-1\}$.

Similarly to Section 7.2 we define the rv R as the number of copies that a relay node has received before it transmits a copy to the destination. On the event $R = m+1$, let $a_k > 0$ be the arrival time of the k th copy to the relay node i for $k = 1, \dots, m+1$, let $d_k > a_k$ be the time where the k th copy is dropped by relay node i for $k = 1, \dots, m$, and let e_{m+1} be the time where copy $m+1$ reaches the destination. Define $\hat{X} = a_1$, $Z_k = a_{k+1} - d_k$ for $k = 1, \dots, m$, and $\hat{Z} = e_{m+1} - a_{m+1}$. Clearly,

$$D_i = \hat{X} + Z_1 + \dots + Z_m + mT + \hat{Z} \quad (7.16)$$

on the event $R = m+1$. Given that $R = m+1$, the rvs $\hat{X}, Z_1, \dots, Z_m, \hat{Z}$ are mutually independent; moreover the rvs $Z_1, \dots, Z_m$ are i.i.d.

Let $D^*(s) := E[e^{-sD_i}]$ be the LST of D_i . We have

$$D^*(s) = E[e^{-s\hat{X}}] E[e^{-s\hat{Z}}] \sum_{m \geq 0} (e^{-sT} E[e^{-sZ_k}])^m P(R = m+1). \quad (7.17)$$

1. *Evaluation of $Z_T^*(s) := E[e^{-sZ_k}]$.* Recall that X denotes a generic inter-meeting time and that its density probability is $g(\cdot)$.

Let $h_T(t) := dP(Z_k < t)/dt$ be the probability density of Z_k . The reason why we indicate the dependency on the parameter T in $h_T(t)$ will soon become apparent. If the source does not meet the relay node i in $(a_k, a_k + T)$ then $Z_k = a_{k+1} - a_k - T \stackrel{st}{=} X - T$, otherwise $Z_k = a_{k+1} - a_k - (T - (a'_k - a_k))$ where a'_k is first time the source meets the relay node i in $(a_k, a_k + T)$. The latter rewrites $Z_k \stackrel{st}{=} X_1 + X_2 - T$ with $X_j \stackrel{st}{=} X$ for $j = 1, 2$. From this we deduce that $h_T(t)$ satisfies the following renewal equation

$$h_T(t) = g(T + t) + \int_0^T g(u) h_{T-u}(t) du. \quad (7.18)$$

Multiplying both sides of this equation by e^{-st} and integrating over t in $[0, \infty)$ gives

$$Z_T^*(s) = \int_0^\infty e^{-st} g(T + t) dt + \int_0^T g(u) Z_{T-u}^*(s) du. \quad (7.19)$$

We have shown that $Z_T^*(s)$ satisfies an integral equation (of Fredholm type) from which $Z_T^*(s)$ can be obtained numerically using standard techniques [Mik64].

2. Probability distribution of R . Finding the probability distribution of R is difficult. We will first assume that R is a geometric rv with parameter $\pi = 1 - P(R = 1)$. Let $\hat{Y}$ be the first time after $t = 0$ that the relay node meets the destination. Since at $t = 0$ the source and the relay node are in steady-state the distribution of $\hat{Y}$, $G_e(t)$, is the excess probability distribution of $G(t)$. Let $g_e(t)$ denotes the density of $G_e(t)$. $P(R = 1)$ is the probability that the relay node meets the source at $a_1 = \hat{X}$ and the last meeting between the relay node and the destination is at $e_1 \in [a_1, a_1 + T]$. Conditioned on $\hat{X}$, the probability of $\{R = 1\}$ verifies the following renewal equation

$$P(R = 1 | \hat{X} = x) = G_e(x + T) - G_e(x) + \int_0^x q(x - u) g_e(u) du, \quad (7.20)$$

where $q(x - u)$ is the solution of the following renewal equation

$$q(x - u) = G(x - u + T) - G(x - u) + \int_0^{x-u} q(x - u - v) g(v) dv. \quad (7.21)$$

Hence,

$$\pi = 1 - P(R = 1) = 1 - \int_0^\infty P(R = 1 | \hat{X} = x) g_e(x) dx. \quad (7.22)$$

It is possible to compute the integral equation of π numerically. However, for sake of simplicity, we will assume that the destination node is at equilibrium at time a_1 , so that $\pi = 1 - G_e(T)$. In summary, we shall approximate the probability distribution of R by

$$P(R = m + 1) \approx (1 - \pi) \pi^m, \quad m \geq 0. \quad (7.23)$$

Based on the above results we may approximate $D^*(s)$ in (7.17) by

$$\begin{aligned} D^*(s) &\approx E[e^{-s\hat{X}}] E[e^{-s\hat{Z}}] \sum_{m \geq 0} (1 - \pi) (\pi e^{-sT} Z_T^*(s))^m \\ &= \frac{1 - \pi}{sE[X]} \frac{(1 - G^*(s)) E[e^{-s\hat{Z}}]}{1 - \pi e^{-sT} Z_T^*(s)}, \end{aligned} \quad (7.24)$$

where $Z_T^*(s)$ is the solution of the integral equation (7.19), $\pi = 1 - G_e(T)$, and

$$E[e^{-s\hat{X}}] = \frac{1 - G^*(s)}{sE[X]}. \quad (7.25)$$

It remains to evaluate $E[e^{-s\hat{Z}}]$. Again, this is not an easy task. Clearly, $e^{-sT} \leq E[e^{-s\hat{Z}}] \leq 1$. For sake of simplicity, we will replace $E[e^{-s\hat{Z}}]$ by 1. This gives the final approximation

$$D^*(s) \approx \frac{1 - \pi}{sE[X]} \frac{(1 - G^*(s))}{1 - \pi e^{-sT} Z_T^*(s)}. \quad (7.26)$$

Finally, $P(D_i > t)$ is obtained by inverting $(1 - D^*(s))/s$, since

$$\frac{1 - D^*(s)}{s} = \int_0^\infty e^{-st} P(D_i > t) dt, \quad (7.27)$$

and with the help of the complex inversion formula [Spi, Chap. 7], which yields

$$P(D_i > t) = \frac{1}{2\pi i} \int_{\gamma - i\infty}^{\gamma + i\infty} e^{ts} \frac{1 - D^*(s)}{s} ds, \quad t > 0, \quad (7.28)$$

where the integration has to be performed along a line $s = \gamma$ in the complex plane (in (7.28) i denotes the imaginary complex number). The real number γ must be chosen so that $s = \gamma$ lies to the right of all singularities.

The approximation for $F_{T_d}(t) := P(T_d > t)$ has been computed when $G(t)$ is an hyper-exponential distribution, namely,

$$G(t) = 1 - \sum_{l=1}^H p_l e^{-\nu_l t}, \quad (7.29)$$

and compared to simulation results. The evaluation of the integral in (7.28) has been performed by using the procedure described in [DA68].

Numerical results are reported in Figure 7.3 (for $N = 10$) and in Figure 7.4 (for $N = 50$). Two hyper-exponential distributions, represented by the t -tuple $(H, \nu_1, \dots, \nu_H, p_1, \dots, p_N)$, have been considered: (H1) $(3, 0.09, 0.08, 0.07, 0.6, 0.3, 0.1)$ with mean $M := E[X] = 22.83 \text{sec.}$,

and (H2) $(3, 0.05, 0.04, 0.03, 0.6, 0.3, 0.1)$ with mean $M := E[X] = 11.84\text{sec.}$. The simulations were done using a C program for a network composed of one source, one destination and $N - 1$ relay nodes.

Let $T_d(\text{app})$ and $T_d(\text{sim})$ be the approximate and simulated delivery delays, respectively. Let $F_{T_d}(\text{app})(\cdot)$ (resp. $F_{T_d}(\text{sim})(\cdot)$) denotes the complementary cumulative distribution function (CCDF) of $T_d(\text{app})$ (resp. $T_d(\text{sim})$).

Figure 7.3 displays the mappings $t \rightarrow F_{T_d(\text{app})}(t)$ and $t \rightarrow F_{T_d(\text{sim})}(t)$ for different values of the expected inter-meeting time for the hyper-exponential distributions H1 and H2. We observe that the approximation is accurate for moderate value of N .

Figure 7.4 compares $F_{T_d(\text{sim})}(t)$ with the CCDF of T_d in the case where the inter-meeting times are exponentially distributed, where the latter distribution has been obtained by using (C.2). We conclude from these results that T_d under exponential inter-meeting times is stochastically smaller than T_d under hyper-exponential inter-meeting times. This is related to the fact that the hyper-exponential distribution has a fatter tail than the exponential distribution.

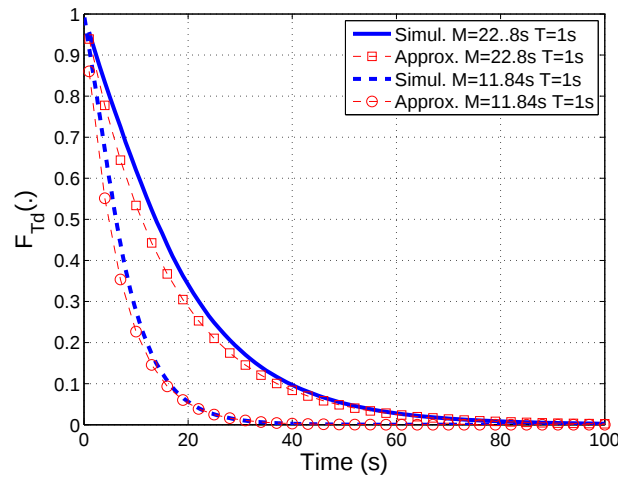


Figure 7.3: Mappings $t \rightarrow F_{T_d(\text{app})}(t)$ and $t \rightarrow F_{T_d(\text{sim})}(t)$ for two different hyper-exponential distributions ($N = 10$).

7.4 Concluding remarks

This chapter evaluated the impact of constant TTLs as opposed to exponential TTLs, and we develop an approximation analysis in the case where the inter-meeting times are arbitrarily distributed. In particular, we showed that exponential inter-meeting times yield stochastically smaller delivery delays than hyper-exponential inter-meeting times; we also showed that exponential TTLs yield stochastically larger delivery than constant TTLs.

In the following chapter we will conclude the thesis.

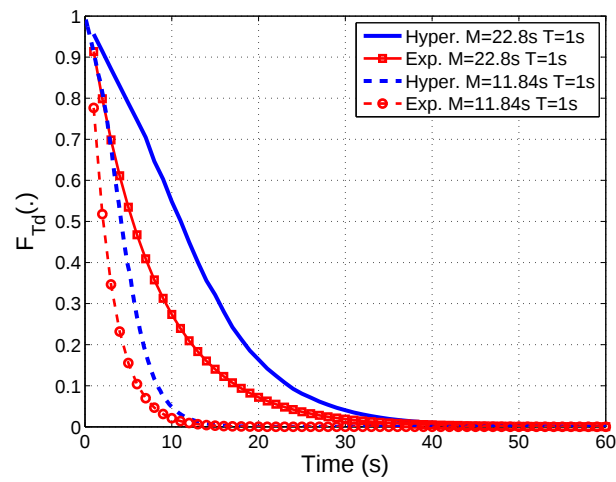


Figure 7.4: Comparison of CCDF of T_d in the case of hyper-exponential (simulated CCDF) and exponential (CCDF of (C.2)) inter-meeting time distributions ($N = 50$).

Chapter 8

Conclusion

8.1 Summary of the results

This thesis focuses on the impact of the mobility in mobile ad hoc networks. We showed through extensive survey of the literature that mobility is the main cause of TCP performance degradation over mobile ad hoc network as it induces frequent route failures. To overcome this problem most of the TCP proposals over ad hoc networks suggest to freeze TCP in the period of failures and to reactivate it on route re-establishments. This motivates us to evaluate the performance gain achieved when the mobility is used to relay packets between nodes. This is done by allowing the nodes to relay the packets of the others through the two-hop relay protocol.

Much of the rest of the thesis was devoted to study the two-hop relay protocol. The performance metric of interest were the throughput, the relay buffer size, the packet delivery delay, and the overhead in terms of packet and energy consumed by the network. In Chapter 3, we showed that relay throughput depends on the node mobility models via their stationary distribution of location. Furthermore, the random mobility models that results in a uniform stationary distribution of node location achieve the lowest relay throughput.

The improvement of the network throughput when the two-hop relay is used comes unfortunately at the expense of an increase in the packet delivery delay. To reduce the delivery delay of packets between source-destination pairs, Chapter 5 introduced the multicopy two-hop relay (MTR) protocol that generates multiple copies of a packet. This is done with the assumption that copies in the relay nodes have limited lifetime. Modeling the system as an absorbing Markov chain yields a closed form expression for the distribution of the delivery delay, distribution of number of copies at delivery instant, and the expected total number of times of generated copies at the delivery time. The latter metric is related to the necessary energy to deliver the packet to the destination.

In Chapter 6 we evaluated the performance of two extensions of MTR protocol called K-copies and K-transmissions MTR that limits the consumed energy by limiting to K either the maximum number of the copies in the network or the total number the copies generated by the source. These performance results were used to find the optimal value of K that minimizes the energy consumption subject to a constraint on the delivery delay. In addition, we evaluated erasure coding a scheme that we showed that gives a small variation of the delivery delay compared with the simple multicopy scheme like in MTR.

In Chapter 7 we studied the impact of considering constant TTL and arbitrary distribution of nodes inter-meeting times on MTR protocol. In this case a non-Markovian analysis was done. Especially, we showed that constant TTL yields stochastically smaller delivery delay than exponential TTLs, and heavy tailed inter-meeting times yield stochastically larger delivery delay than exponential inter-meeting times.

8.2 Future work

In this thesis we have evaluated the performance of relaying when the relay nodes are mobile. Another research direction that can be considered is when the relay nodes are stationary. In this situation the question that arises is “what is the best placement of the relay nodes in order to improve the network performance and to minimize the interference between relay nodes?”. I started to investigate this issue.

In the second part of the thesis we have considered a single packet in the system and only one source. This assumption neglects the queueing delay and the impact of other packets in the network. Relaxing this assumption in MANETs will require a detailed queueing analysis of the relay buffer of nodes. This will constitute my future research activities in the domain of queueing theory.

On the other hand, we have assumed that when two nodes meet they can exchange data successfully with probability one. Assuming that a node can apply a certain policy to decide to transmit or not will be important. Since, in this case one can try to find the optimal policy that the nodes may select to minimize a cost function subject to some constraints on energy or delay. I plan to work on this direction using the formalism of dynamic programming.

Finally, we believe that power control of the mobile ad hoc networks in the context of limited energy of batteries of the mobile nodes is a hot research topic. Our main objective in this topic will be to increase the lifetime of the batteries and at the same time to decrease the delivery delay of the packets. In order to realize such objective, one can use queueing theory and dynamic programming theory.

Appendix A

IEEE 802.11 WLANs standard

The IEEE 802.11 standard [IEE] defines a physical layer (PHY) and medium access control (MAC) protocols for WLANs.

The IEEE 802.11 standard encompasses three different PHY specifications: frequency hopping spread spectrum (FHSS), direct sequence spread spectrum (DSSS), and infrared (IR). The most commonly deployed PHY is the DSSS working at the 2.4 GHz ISM band. The basic data rate for DSSS is 1 Mbps (Mega bits per second) achieved by means of a differential binary phase shift keying (DBPSK) modulation. Similarly, a data rate of 2 Mbps is provided using differential quadrature phase shift keying (DQPSK) modulation. Higher rates of 5.5 and 11 Mbps are added in the 802.11b version using complementary code keying (CCK). Using orthogonal frequency division multiplexing (OFDM) in the 5 GHz band, 802.11a supports data rate up to 54 Mbps. But due to the incompatibility problem between 802.11b and 802.11a, the 802.11g was introduced to work in the 2.4 GHz ISM band and to support data rates up to 54 Mbps.

The IEEE 802.11 MAC layer specifications, common to all PHYs and data rates, coordinates the communications between stations and controls the access to the channel. The Distributed Coordination Function (DCF) describes the default MAC protocol operations. DCF is based on a carrier-sense, multiple access, collision avoidance (CSMA/CA) scheme. Both the PHY and the MAC layers cooperate to implement collision avoidance procedures. The PHY implements the clear channel assessment (CCA) algorithm to determine if the channel is idle (known as physical carrier sensing mechanism). In general, CCA includes one of the 3 following methods: The first method is energy above threshold, in which CCA reports a busy medium upon detection of any energy above an energy detection (ED) threshold. The second method is carrier sense only, in which CCA reports a busy medium only

upon detection of a signal. The signal power may be above or below the ED threshold. The third method is a combination of the two first methods. The MAC provides a virtual carrier sense mechanism. This mechanism is referred to as the network allocation vector (NAV). The NAV maintains a prediction of future traffic on the medium. This prediction is based on the information sent in the RTS and CTS control frames. The NAV setting is shown in Fig. A.1. Communication is established when one of the wireless nodes sends the RTS frame. The receiving station issues a CTS frame that echoes the sender's address. As generally the antenna used is omni-directional, all nodes within the transmission range of the sender are able to decode the RTS. Then they will update their NAV, as well as nodes within the transmission range of the receiver able to decode the CTS. The exchange of RTS/CTS prior to current data frames provides a channel reservation, in order to avoid interference caused by hidden nodes.

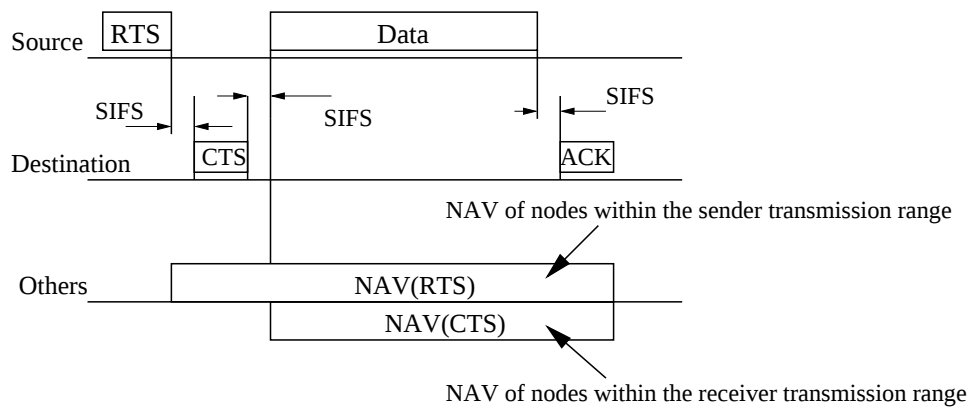


Figure A.1: RTS/CTS/DATA/ACK and NAV setting.

According to the 802.11 protocol specifications, the MAC must implement the basic access method (depicted in Fig. A.2) as follow. If a node want to access the channel for transmission, it has to wait until the channel becomes idle through the CSMA/CA algorithm. If the channel is sensed idle for a period greater than a DCF inter-frame space (DIFS), the node goes into a random *backoff* time. The backoff duration, expressed in slot, is uniformly chosen in the interval $\{0, 1, \dots, CW\}$, where contention window CW is between CW_{min} and CW_{max} . For example in 802.11b, $CW_{min} = 7$ and $CW_{max} = 1023$. If the medium is sensed busy at any time during the backoff, the backoff procedure is suspended. It is resumed after the medium has been idle for the duration of DIFS period. Upon the successful reception of a transmitted frame, the destination node returns an ACK frame after a Short inter-frame space (SIFS). If an ACK is not received within an ACK timeout interval, the sender node assumes that the transmission is failed, and it needs to retransmit data frame by repeating the basic access procedure. In case of repeated transmission failures, the node can retransmit a short frame up to 7 times¹, and a long frame it has 5 times². After that, the node will

¹e.g RTS and frame of length less than the threshold dot11RTSThreshold.

²Frame of length greater than threshold dot11RTSThreshold.

drop this data frame, and the MAC layer will notify Logical Link Control (LLC) layer about a link failure, which in turn will inform the routing layer about the link failure. In fact, a node that fails transmitting a frame assumes that the failure is due to collisions caused by other nodes transmissions, which are contending to access the channel. So to resolve the contention, before repeating the transmission procedure, a node duplicates its contention window. By duplicating the CW, the probability that two nodes select the same backoff time becomes smaller, then collisions are avoided.

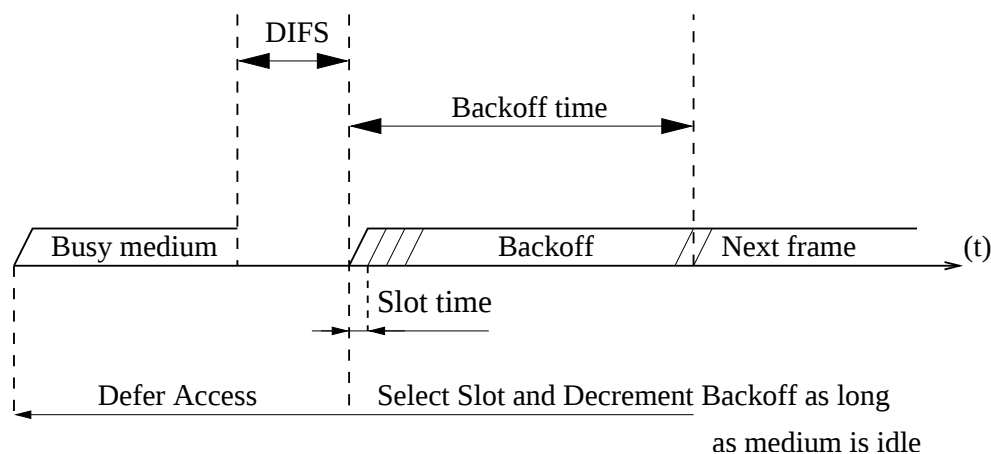


Figure A.2: Basic access method.

Errors are detected by checking the Frame Check Sequence (FCS) that is appended to the frame payload. In case of frame error detection, the packet is dropped and it is retransmitted later. A node that receives an erroneous frame has to determine idle medium for a duration of extended IFS (EIFS) before resuming the backoff procedure. In Tab. A.1, we list the features of the 802.11b “Wireless LAN World PC Card” of AGERE systems.

Frequency range	2400MHz to 2484MHz
Slot time	20 μs
SIFS	10 μs
DIFS	50 μs
Transmission power	15dBm
Data rate	1,2,5.5,11Mbps
Outdoor tx range	550,400,270,160m

Table A.1: Wireless LAN PC card of AGERE systems characteristics.

Appendix B

Routing Protocols in Ad Hoc Networks

First, we have to note that in Ad hoc networks each node acts as a router for other nodes. The traditional link-state and distance-vector algorithms do not scale well in large MANETs. This is because periodic or frequent route updates in large networks may consume a significant part of the available bandwidth, increase channel contention and require each node to frequently recharge its power supply.

To overcome the problems associated with the link-state and distance-vector algorithms a number of routing protocols has been proposed for MANETs. These protocols can be classified into three different groups: proactive, reactive and hybrid. In proactive routing protocols, the routes to all destination are determined at the start up, and maintained by means of periodic route update process. In reactive protocols, routes are determined when they are required by the source using a route discovery process. Hybrid routing protocols combine the basic properties of the first two classes of protocols into one. The IETF working group *MANET* is aiming to standardize IP routing protocol functionality suitable for wireless routing application within both static and dynamic topologies [IET]. For a review of routing protocols in MANETs, we refer the reader to references [Per01] and [AWD04].

B.1 Proactive routing protocols

In proactive routing protocols, each node maintains information on routes to every other node in the network. The routing information is usually kept in different tables. These tables are periodically updated if the network topology changes. The difference between these protocols lies in the way the routing information is updated and in the type of information kept at each routing table. Among these protocols we find: Destination-Sequenced Distance-Vector (DSDV), Optimized Link State Routing (OLSR), Topology Broadcast Reverse Path

forwarding (TBRFP), Wireless Routing Protocol (WRP), Global state routing (GSR), Fish-eye State Routing (FSR), Source-Tree Adaptive Routing (STAR), Distance Routing Effect Algorithm for Mobility (DREAM). For a review of reactive protocols see [AWD04]. In the next paragraphs, we will show how the distance-vector protocol DSDV and the link-state protocol OLSR work.

Destination-Sequenced Distance-Vector (DSDV) [PW94] is a point to point distance vector protocol. DSDV provides a single path to destination, which is selected using the distance vector shortest path routing algorithm based on the number of hops. To reduce the amount of overhead transmitted through the network, two types of *update* packets are used. These are referred to as a “full dump” and incremental packets. The full dump carries all the available routing information, while the incremental packet carries only the information changed since the last full dump. Each routes has a sequence number stamped by the destination node, and only routes that have the most recent sequence number may be forwarded through *update* packets. The requirement of the *periodic* update messages introduces a large amount of overhead, which makes DSDV not suitable for large networks.

Optimized Link State Routing (OLSR) [CJ03] is a point to point link state protocol. Each node maintains topology information about the network by periodically exchanging link-state messages. The novelty of OLSR is that it minimizes the size of each control message and the number of re-broadcasting during route update, by employing the multi-point relaying (MPR) strategy. To do this, during each topology update, each node in the network selects a set of neighboring nodes to retransmit its packets. This set of nodes is called the multipoint relays of that node. Any node which is not in the set can read and process each routing packets but do not retransmit them. To select the MPRs, each node periodically broadcasts a list of its one-hop neighbors using “hello” messages. The use of the MPR strategy makes OLSR scalable, since only MPR nodes rebroadcast the update packets.

B.2 Reactive routing protocols

Reactive or on-demand routing protocols have been designed to reduce the overhead in proactive protocols. The overhead reduction is accomplished by establishing routes on-demand, and by maintaining only active routes. Route discovery usually takes place by flooding a route request packet through the network. When a node with a route to the destination (or the destination itself) is reached, a route reply packet is sent back to the source. Representative reactive routing protocols are: Dynamic Source Routing (DSR), Ad hoc On Demand Distance Vector (AODV), Temporally Ordered Routing Algorithm (TORA), Associativity Based routing (ABR), Signal Stability Routing (SSR). For a review of reactive protocols see [AWD04]. In the rest of this section, we will show how the DSR and AODV protocols work. These protocols are adopted by the MANET IETF working group.

DSR is a loop free [JMH03], source based, on demand routing protocol, where each node maintains a route cache that contains the routes discovered by the node. The route dis-

covery process is only initiated when a source does not have a valid route to the destination. Entries in the route cache are continually updated as new routes are learned.

AODV [PBRD03] is a reactive improvement of the DSDV protocol. AODV minimizes the number of route broadcasts by creating routes on-demand, instead of maintaining a complete list of routes as in the DSDV protocol. Similar to DSR, route discovery is initiated on-demand. The route request is then forwarded by the source to the neighbors, and so on, until either the destination or an intermediate node with a fresh route to the destination, is located.

DSR has a potentially larger control overhead and memory requirements than AODV, since each DSR packet must carry full routing path information, whereas AODV packets contain only the destination address. On the other hand, DSR can use both asymmetric and symmetric links during routing, while AODV only works with symmetric links. In addition, nodes in DSR maintain in their cache multiple routes to a destination, a feature helpful during route failures. In general, both AODV and DSR work well in small to medium networks with moderate mobility.

Appendix C

Présentation des Travaux de Thèse

C.1 Introduction

Pour près d'un quart de siècle, la communication mobile a vécu une croissance explosive, particulièrement dans la dernière décennie. En particulier, les réseaux ad hoc, l'un des secteurs des réseaux de communication mobile, ont attiré une importante communauté scientifique.

Les réseaux ad hoc sont des systèmes répartis complexes qui se composent de plusieurs nœuds mobiles ou stationnaires. Ces nœuds communiquent entre eux à travers des ondes radio, et ils ont la capacité de collaborer et de s'organiser librement sans l'intermédiaire d'aucune unité centrale. De cette façon, ils forment des réseaux arbitraires et provisoires qui peuvent fournir des services dans des zones où aucune infrastructure n'est préexistante. Selon le cas où les nœuds sont stationnaires ou mobiles, un réseau ad hoc s'appelle un réseau ad hoc mobile (MANET en anglais) ou un réseau ad hoc fixe (SANET en anglais). Dans ce qui suit, le terme réseau ad hoc représente à la fois MANET et SANET.

La capacité de routage qu'ont les nœuds est le point fondamental des réseaux ad hoc. On appellera la capacité de routage multisauts: un message partant d'un nœud source pour atteindre un nœud destination pouvant traverser plusieurs nœuds intermédiaires. Le chemin traversé par le message dans ce cas s'appelle la route multisauts. Cependant, dans certains scénarios, les routes multisauts subissent fréquemment des coupures à cause de la mobilité des nœuds, de la faible densité des nœuds, ou du canal radio. Ceci génère plusieurs défis aux protocoles des réseaux ad hoc.

D'une autre part, le réseau local sans-fil est un réseau sans fil à un seul saut, généralement avec une petite portée de transmission. Mais certaines normes du réseau local sans-fil peuvent parfois posséder deux modes de fonctionnement: un mode centralisé et un mode

ad hoc. Dans le mode centralisé, les communications entre deux nœuds ou entre un nœud et l'Internet passent à travers un point d'accès. Au contraire, le mode ad hoc est un mode pair à pair où chaque nœud peut servir comme un routeur pour les autres. Ce dernier mode est employé dans les réseaux ad hoc pour permettre les communications à multisauts. C'est pour cette raison, certaines normes du réseau local sans-fil s'appellent la technologie prometteuse des réseaux ad hoc. Parmi celles ci, on trouve le fameux standard de communication sans-fil IEEE 802.11 [IEE], connu commercialement sous le nom de WIFI.

C.2 Contributions et organisation

Cette thèse s'intéresse à l'impact de la mobilité dans les réseaux ad hoc mobiles. Elle est divisée en deux parties. La première partie dresse un état de l'art de TCP dans les réseaux ad hoc, et elle constitue en même temps la motivation de la deuxième partie. Dans cette partie, on étudie les performances du protocole de relais à deux sauts qui s'appuie sur la mobilité des nœuds pour transmettre les paquets entre les nœuds [GT02].

C.2.1 Première partie

La première partie est constituée seulement du chapitre 2.

Chapitre 1: Etat-de-l'Art de TCP dans les Réseaux Ad Hoc. Dans ce chapitre, nous dressons l'état de l'art de TCP dans les réseaux ad hoc et en particulier dans les réseaux ad hoc mobiles.

TCP (*Transport Control Protocol* en anglais) a été conçu pour fournir un service de transport de bout en bout fiable. Au cours des ans, des versions successives de TCP ont vu le jour, dans le but d'améliorer les performances du protocole. Toutefois, ces améliorations ont été réalisées uniquement dans le contexte des réseaux filaires, comme l'Internet. L'avènement des réseaux sans fil, en particulier les réseaux ad hoc mobiles, a mis à jour des comportements indésirables de TCP, qui dégradent fortement les performances (débit, temps de réponse, délai, etc.). Plusieurs contributions visant à adapter TCP aux réseaux sans fil ont récemment été proposées. L'objectif de cette étude est d'identifier les problèmes spécifiques liés à l'utilisation de TCP dans les réseaux ad hoc et à faire le point sur les dernières solutions proposées pour y remédier.

Dans les réseaux ad hoc, le problème principal du TCP est qu'il effectue du contrôle de congestion en cas de pertes qui ne sont pas nécessairement liés à la congestion. Puisque le taux d'erreurs des bits est très bas dans les réseaux filaires, presque toutes les versions de TCP de nos jours supposent que les pertes de paquets sont nécessairement dues à la congestion. En conséquence, quand un paquet est détecté comme perdu, TCP diminue son débit de transmission en ajustant sa fenêtre de congestion. Malheureusement, les réseaux sans fil souffrent de plusieurs types de pertes qui ne sont pas liées à la congestion, faisant en sorte que TCP n'est pas adapté à cet environnement. Au cours des dernières années, plusieurs propositions ont apparu pour améliorer les performances de TCP dans les réseaux sans fil qui ne sont pas nécessairement de type ad hoc. Comme par exemple le cas des réseaux

sans fil local [AHM03, BPSK97, BSAK95, BB97], des réseaux cellulaires mobiles [BS97, BSK95], et des réseaux satellites [DMT96, HK99]. D'une part, les réseaux ad hoc ont plusieurs problèmes communs avec ces réseaux, en particulier le taux d'erreurs élevé et les canaux radio asymétriques. D'une autre part, il y a des nouveaux problèmes qui sont dus à la mobilité et à la communication multisauts des nœuds, tels que la partition du réseau, la coupure des routes, et les nœuds cachés.

En examinant les performances de TCP dans les réseaux ad hoc, nous avons d'abord identifié quatre principaux problèmes: (i) TCP n'est pas capable de différencier entre les pertes causées par les coupures de routes et celles dues à la congestion dans le réseau. (ii) TCP souffre des coupures fréquentes des routes. (iii) la saturation du canal radio. (iv) TCP n'est pas équitable. Notez que les deux premiers problèmes sont les principales causes de la dégradation de performance de TCP dans MANETs. Tandis que les deux derniers sont les principales causes de la dégradation de performance de TCP dans SANETs. En s'appuyant sur ces quatre problèmes, les solutions de TCP proposées dans la littérature dans les réseaux ad hoc ont été groupées en quatre ensembles. Dans chaque ensemble, les solutions ont été groupées en deux sous ensembles selon le cas où ces solutions nécessitent la modification d'une seule couche du modèle OSI (Open Systems Interconnection en anglais), ou la modification et la collaboration entre au moins deux couches du modèle.

Par exemple, des améliorations considérables ont été effectuées quand TCP différencie entre les pertes de paquets dues à la congestion et celles qui sont dues aux caractéristiques spécifiques de MANET. Afin de réaliser cette approche dans [CRVP98a, HV02, LS01, WZ02, KTC01] il est proposé que quand le protocole du routage détecte qu'une route est coupée, il informe TCP à ce propos. Lors de la réception de cette information, ce dernier entre dans un état d'attente, où il cesse de transmettre les paquets et gèle les valeurs de tous ses paramètres, comme la taille de la fenêtre de congestion et le temporisateur de retransmission. Après la reconstruction de la route avec la destination, TCP se réactive et entre dans son état normal de fonctionnement. Dans les solutions qui visent à modifier une seule couche du modèle OSI, ces modifications et les nouveaux mécanismes sont implémentés seulement dans cette seule couche. Par exemple dans [AJ03], l'utilisation des acquittements adaptatifs retardés proposés par Altman et Jimenez réduit la saturation des canaux radios dans SANETs. Dans [FZL⁺03], Fu et Al proposent deux nouvelles techniques de la couche liaison afin de résoudre le problème du partage équitable des ressources dans SANETs.

La principale conclusion est que la mobilité dégrade les performances de TCP dans MANETs, à cause de problèmes de routage et de partitionnement du réseau qu'elle occasionne. Partant de ce constat, nous proposons et analysons dans la deuxième partie de la thèse des schémas de transmission qui s'appuient sur la mobilité des nœuds. Plus précisément, chaque nœud peut servir de relai en l'absence de route directe entre la source et la destination selon le protocole de relais à deux sauts.

C.2.2 Deuxième partie

La deuxième partie est constituée des chapitres 3-7. Elle évalue les performances du protocole de relais à deux sauts à l'aide des modèles stochastiques. Ce protocole s'appuie sur la

mobilité des nœuds pour transmettre les paquets de leurs sources à leurs destinations.

Chapitre 2: Effet de la Mobilité sur le débit des Protocoles à base de Relais dans les Réseaux Ad Hoc. Dans ce chapitre, nous considérons un réseau ad hoc mobile constitué de trois types de nœuds: nœuds sources, nœuds destinations et nœuds relais. Le rôle du nœud relais est de transmettre les paquets entre la source et la destination en l'absence d'une route directe entre ces derniers. Nous supposons que les nœuds se déplacent indépendamment les uns des autres. Nous définissons le débit du nœud relais comme le taux maximal auquel celui-ci peut transmettre les paquets. Nous trouvons que ce débit ne dépend que de la distribution des positions des nœuds à l'état stationnaire. De plus, nous montrons que ce débit atteint une valeur minimale quand la distribution de position est uniforme. Pour deux modèles de mobilité (Random Waypoint, Random Direction) nous établissons des formules approchées du débit du nœud relais ainsi nous trouvons la condition de stabilité de la file d'attente d'un nœud relais.

Dans un réseau ad hoc mobile (MANET), puisque le réseau est dynamique et les nœuds sont mobiles, l'état des liens évolue avec le temps. Un lien entre deux nœuds est fonctionnel quand ces nœuds sont à l'intérieur de la portée de transmission de l'une de l'autre, et il n'est pas fonctionnel dans le cas contraire. L'établissement d'une route multisauts d'un nœud source à un nœud destination exige la disponibilité d'un certain nombre de liens qui sont tous fonctionnels. En effet, dans MANETs, les coupures de route sont souvent occasionnées particulièrement quand les nœuds se déplacent fréquemment et la densité des nœuds est faible. C'est un scénario typique des réseaux tolérants les délais (DTNs en anglais) [dtn].

Grossglauser et Tse [GT02] ont observé qu'on peut s'appuyer sur la mobilité pour améliorer la capacité du réseau. Leur idée était d'évaluer le gain du débit achevé en se servant des nœuds mobiles en tant que des nœuds relais. Leur mécanisme de relais est appelé le protocole de relais à deux sauts. Son fonctionnement est comme le suivant: s'il n'y a aucune route entre le nœud source et le nœud destination, le nœud source transmet ses paquets à tous ses nœuds voisins (appelés les nœuds *relais*) pour les transférer à la destination. Selon ce protocole, un nœud relais a seulement le droit de transmettre les paquets qu'il a à leurs destinations, ce qui justifie le nom de ce protocole relais à deux sauts. Après la proposition de ce protocole, plusieurs études de performance ont été faites. Le but de ces dernières était d'évaluer le passage à l'échelle du débit et du délai quand le nombre des nœuds est très grand, à ce titre voir [GK00, GMPS04]. Dans ce chapitre, on s'intéresse particulièrement à l'effet de la mobilité sur le débit maximal d'un nœud relais dans le cas où le nombre des nœuds est fini.

Ce chapitre étudie l'impact du modèle de la mobilité des nœuds sur le *débit de relais* des paquets en utilisant le protocole de relais à deux sauts. Particulièrement, on s'intéresse au débit de relais maximal d'un nœud, c-à-d, le taux maximal de relais de données que un nœud peut contribuer à la communication entre deux autres nœuds. Le relais des données pour d'autres nœuds exige que le nœud relais alloue ses propres ressources en terme d'énergie et de la taille de la file d'attente. La consommation d'énergie du protocole de relais à deux sauts sera établie dans le chapitre 5 dans un cas plus général où un paquet peut avoir plusieurs copies dans le réseau. En particulier, un nœud relais doit garder les données à

transmettre dans sa file d'attente de relais. Par conséquent, l'étude de la taille de la file d'attente d'un nœud relais forme un important axe de recherche, ce problème sera détaillé dans le Chapitre 4.

Les modèles de mobilité considérés dans ce chapitre sont aléatoires. Particulièrement, on s'intéresse au modèle random waypoint (RWP) et au modèle de la direction aléatoire (RD). Dans le modèle RWP [BMJ⁺98] chaque nœud est assigné à une position initiale dans un plan, et il voyage vers la destination choisie aléatoirement à une vitesse constante. La distribution de la vitesse d'un nœud est uniforme indépendamment de sa position et de celle de la destination. Après l'arrivée à la destination, le nœud peut faire un arrêt de mouvement pendant un temps aléatoire. Après l'écoulement de ce temps, une nouvelle destination et vitesse sont choisies, indépendamment des valeurs antérieures. Selon RWP, la distribution stationnaire des positions n'est pas uniforme, et elle atteint un point maximal au centre du plan où les nœuds bougent. Dans le modèle RD [PNL05] à chaque nœud est assigné une direction, une vitesse, et un temps de voyage à l'état initial. Le nœud voyage dans cette direction à la vitesse choisie pour une durée donnée. Quand le nœud termine son voyage, le nœud peut faire un arrêt de mouvement pendant un temps aléatoire. Après un nouveau temps de voyage, une nouvelle vitesse, et une nouvelle direction sont choisis indépendamment de ses valeurs antérieures. Quand un nœud atteint une frontière, il se reflète ou il fait une réapparition instantanée de l'autre côté du plan. Selon RD, la distribution stationnaire des positions des nœuds est uniforme.

Le point de départ est l'observation qui relie l'évolution de la taille de la file du nœud relais à l'évolution de la quantité du travail en charge dans une file d'attente G/G/1. Les temps de service et les temps entre les arrivées dans ce système sont déterminés en fonction du modèle de la mobilité des nœuds.

Les principaux résultats sont les suivants:

- Nous trouvons que le débit de relais ne dépend que de la distribution de position des nœuds à l'état stationnaire.
- Nous montrons que le débit de relais atteint une valeur minimale quand la distribution de position est uniforme.

Chapitre 3: Effet de la Mobilité sur la Taille des Files de Relais des Protocoles à base de Relais dans les Réseaux Ad Hoc. Dans ce chapitre, on considère un réseau ad hoc mobile constitué de trois types de nœuds: source, destination, et nœud relais. On étudie l'évolution de la taille de la file d'attente d'un nœud relais en fonction de la mobilité. Nous trouvons des formules approchées de la taille moyenne de la file d'attente de relais, et nous prouvons que cette taille moyenne dépend du temps moyen de contact entre les nœuds et aussi de sa variance. Les modèles de mobilité sont la marche aléatoire [Fel68] (random walk en anglais) et le modèle de direction aléatoire (RD) introduit dans le chapitre précédent.

L'objectif dans ce chapitre est d'étudier l'évolution de la taille de la file d'attente d'un nœud relais (RB). À cet effet, nous trouvons la taille moyenne de RB, dans le cas de forte

charge, incorporée à des instants de temps qu'on appelle les temps de cycle. Notez que la taille moyenne de RB dans le cas de forte charge sert également comme une limite supérieure de la taille moyenne de RB [Kle76]. En outre, nous prouvons numériquement que sous les modèles de mobilité considérés, la taille moyenne de RB, incorporée aux instants des cycles, converge vers la moyenne temporelle du RB dans le cas de forte charge. Cette étude a été réalisée pour les deux modèles de mobilité considérés.

Comme dans le chapitre précédent, le point de départ est l'observation qui relie l'évolution de la taille de la file d'attente de nœud relais à l'évolution de la quantité du travail en charge dans une file d'attente $G/G/1$. Le temps moyen d'attente dans la file $G/G/1$ sera approché par sa valeur correspondante dans une file $GI/G/1$, c.-à-d., quand les temps entre les arrivées sont indépendants des temps de services. En utilisant l'inégalité de Kingman dans le cas d'une file $GI/G/1$ [Kle76, P. 29], on trouve que la taille moyenne de RB, $E[\tilde{B}]$, vérifie l'inégalité suivante

$$E[\tilde{B}] \leq \frac{r_d^2 Var(\alpha_n) + r_s^2 Var(\sigma_n)}{2(r_d E[\alpha_n] - r_s E[\sigma_n])}, \quad (C.1)$$

où r_d et r_s sont les taux de transmission des nœuds source et relais respectivement, α_n est le temps total qu'un nœud source reste en contact avec le nœud relais dans un cycle, et σ_n est le temps total que le nœud relais reste en contact avec la destination dans un cycle. Les expressions de la variance et de la valeur moyenne de α_n et σ_n ont été trouvées dans le cas de la marche aléatoire et de la direction aléatoire. Notez que la valeur moyenne de $E[\tilde{B}]$ tend vers le second terme de (C.1) dans le cas de forte charge, c.-à-d., quand $r_d E[\alpha_n] \approx r_s E[\sigma_n]$ avec $r_d E[\alpha_n] < r_s E[\sigma_n]$.

Chapitre 4: Évaluation de Performances du Protocole de Relais à Deux Sauts avec la Durée de Vie des Paquets Limitée. Nous considérons un réseau ad hoc mobile. Le protocole de relais utilisé est le protocole de relais à deux sauts avec l'option qu'un paquet peut avoir plusieurs copies dans le réseau. De plus, nous supposons que la durée de vie des paquets transmis par les nœuds relais est limitée. Tous les nœuds se déplacent indépendamment les uns des autres. Nous établissons les formules de la distribution de probabilité des temps de délai pour l'envoi des paquets à la destination, le temps de délai pour recevoir l'accusé de réception d'un paquet par un nœud source, ainsi que la quantité d'énergie consommée par le réseau pour l'envoi des paquets. Ces formules ont été trouvées en utilisant un modèle markovien, ainsi qu'un modèle fluide. En comparant les résultats théoriques trouvés et les résultats des simulations, nous avons trouvé que notre modèle est précis surtout pour les petites valeurs de portée de transmission des nœuds et lorsque le nombre des nœuds est fini.

L'objectif de ce chapitre est d'évaluer les performances d'une variante du protocole de relais à deux sauts qui vise à réduire le temps de délai de l'envoi d'un paquet d'une source à une destination. Ceci a été réalisé en permettant au nœud source de générer plusieurs copies du paquet et de transmettre ces copies aux nœuds relais (une copie par nœud) [SH05, ZNKT06, GNK05]. Le reste de ce chapitre est consacré à étudier cette variante

que l'on appellera le protocole de multicopie et de relais à deux sauts (MTR). L'intérêt porte sur le temps de délai de l'envoi, le temps de l'aller retour des paquets (Round-Trip-Time en anglais), et les coûts induits par le protocole MTR. Ceci sera fait dans le cas où les paquets dans les nœuds relais ont une durée de vie limitée dans le réseau au contraire des autres études [GNK05, SH05].

Un autre protocole de relais lié au protocole MTR est le routage épidémique [SH03, ZNKT06]. Ce protocole est identique à MTR, sauf que dans celui-ci un nœud relais peut transmettre un paquet à n'importe quel nœud, y compris un autre nœud de relais. Le routage épidémique diminue le délai de l'envoi des paquets à leurs destinations au coût d'une augmentation de l'énergie consommée par le réseau. Les performances du protocole MTR et du protocole de routage épidémique sont également comparées dans ce chapitre.

Nous considérons le modèle présenté dans [GNK05]. Dans ce modèle, les caractéristiques de MANETs sont capturées par un simple paramètre, $1/\lambda$, qui représente le temps écoulé entre deux rencontres consécutives de n'importe quelle paire de nœuds, appelé le temps entre-rencontre. Plus précisément, le réseau est composé de $N + 1$ nœuds: un nœud source, un nœud destination, et $N - 1$ nœuds relais. Les nœuds peuvent seulement communiquer pendant certains instants de temps, appelés les périodes de rencontre. Dans [GNK05], les temps entre-rencontres sont mutuellement indépendants et identiquement distribués avec une distribution exponentielle d'intensité $\lambda > 0$.

Dans ce chapitre nous adressons le scénario où le nœud source veut envoyer un paquet au nœud destination. À cet effet, la source peut utiliser les nœuds relais selon le protocole MTR.

En plus du modèle dans [GNK05], on suppose dans ce chapitre que chaque copie du paquet a une durée de vie limitée (time-to-live TTL en anglais). Quand le TTL d'une copie, expire la copie est détruite. On a fait l'hypothèse que les TTLs des copies ont une distribution exponentielle d'intensité $\mu > 0$. Le paquet de la source n'a aucune TTL associée, de sorte que la source puisse toujours envoyer une copie à un autre nœud (si le paquet à la source a un TTL, il y a alors une probabilité non nulle que le nœud destination ne recevra jamais le paquet. Ce scénario n'est pas considéré dans ce chapitre).

Nous supposons que la source est prête à transmettre le paquet à la destination à l'instant $t = 0$. Le délai de l'envoi (de paquet), T_d , est le premier instant de temps, après $t = 0$, où la destination reçoit le paquet (ou une copie du paquet). En plus de T_d , nous évaluerons, C_d , le nombre de copies dans le système à l'instant de la livraison du paquet, et G_d , le nombre total de copies générées par la source avant la livraison du paquet (Section 5.3). Le dernier est lié à l'énergie nécessaire pour livrer le paquet à la destination. À noter que ces trois métriques sont indépendantes de mécanisme de diffusion de l'anti-paquet, appelé VACCIN, puisque ce dernier est activé après la livraison du paquet à la destination. En considérant VACCIN, le temps d'aller retour de paquet, T_r , défini comme le premier instant après $t = 0$ où la source reçoit l'anti-paquet, est également évalué.

Le formalisme employé dans ce chapitre pour trouver les quatre métriques est la théorie de la chaîne de Markov absorbante à temps continu et à nombre d'états fini [Neu81]. La chaîne transite à l'état absorbant quand la destination reçoit le paquet pour la première fois.

Ainsi, dans le cas où le nombre des nœuds est grand, des formules approchées des métriques considérées ont été établies en utilisant un modèle fluide.

Chapitre 5: Evaluation des Performances d'une Classe de Protocole de Relais à Deux Sauts.

Dans ce chapitre, nous évaluons les performances d'une classe de protocoles de relais à deux sauts dans MANETs avec des modèles simples. L'intérêt porte sur le protocole de multicopie et de relais à deux sauts (MTR), selon lequel la source peut générer plusieurs copies d'un même paquet, et sur le protocole de relais de deux sauts avec le codage des données, avec lequel les données sont distribuées dans des blocs. Les métriques de performance d'intérêt sont le délai de l'envoi d'un paquet à sa destination, le nombre de copies à l'instant de la livraison, et le nombre total de copies générées par la source; la dernière métrique sera plus grande quand les copies ont des durées de vie limitées, une situation que nous adressons. Nous étudions également les cas où le nombre de copies dans le réseau est limité, et où le nombre de copies générées est limité afin de réduire la consommation d'énergie.

Plus précisément ce chapitre évalue deux extensions du protocole de multicopie et de relais à deux sauts (MTR) qui était déjà présenté dans le chapitre 5. La première extension est appelée *K-copies MTR*. Dans cette extension, le nombre de copies dans le réseau a une borne supérieure égale à K . Cependant, la deuxième extension est appelée *K-transmissions MTR*, pour laquelle le nombre de copies générées dans le réseau est limité à K . En outre, la performance du protocole de relais à deux sauts avec le codage de perte (erasure coding en anglais) est évaluée et comparée avec le mécanisme de multicopie comme celui de MTR. Commençons maintenant la présentation de ces trois extensions.

K-copies MTR. Dans *K-copies MTR*, la source du paquet limite le nombre de copies dans le réseau à K . Ainsi, si la source rencontre un nœud relais avant de rencontrer la destination, et si le nombre de copies dans le réseau est inférieur à K , alors elle envoie une copie à ce nœud relais; le nœud relais enverra à son tour le paquet à la destination quand elle se met en contact avec lui. Notez que, dans le protocole de relais à deux sauts, la source est le seul nœud qui génère des copies, et ainsi, elle ajuste la valeur de TTL d'une copie. Donc, avant la livraison du paquet à la destination, la source sait exactement le nombre actuel de copies dans le réseau. En s'appuyant sur ceci, la source décide de transmettre une nouvelle copie à un nœud relais qu'elle rencontre seulement dans le cas où le nombre de copies est inférieur à K . Notez que, la limitation du nombre de copies dans le réseau des *K-copies MTR* est au coût d'une augmentation du délai de la livraison de paquet. Dans la section 6.3.4, nous avons trouvé la valeur optimale de K , le nombre maximum des copies de paquet dans le réseau, qui minimise l'énergie consommée en présence d'une contrainte sur le délai.

K-transmissions MTR. Le *K-transmissions MTR* limite l'énergie consommée dans le réseau en limitant le nombre de copies que la source peut produire à K . Ainsi, si la source rencontre un nœud relais avant de rencontrer la destination, et si le nombre de copies produites est inférieur à K , alors elle transmet une nouvelle copie à ce nœud. De la même façon que *K-copies MTR*, cette limitation de la consommation d'énergie de cette extension implique avec une augmentation du délai de la livraison. Cette extension est semblable à

celle de [SH05] pour limiter la consommation d'énergie du routage épidémique.

Protocole de Relais à deux sauts avec le codage de pertes. Dans cette extension, au lieu de copier le paquet plusieurs fois, la source produit une certaine redondance dans ses transmissions et envoie plus de données que l'information réelle. Ainsi, lors de la réception d'un morceau de données (paquet), la source génère n blocs. La transmission du paquet est accomplie quand la destination reçoit le k -ième bloc ($k \leq n$), indépendamment de l'identité des blocs qu'elle a reçus [WJMF05]. Plus de détails sont fournis dans la Section 6.5 où nous avons prouvé que le codage de perte réduit la dispersion du délai de la livraison.

Chapitre 6: L'impact de TTL Constant et de Temps Entre-Rencontre Arbitraire sur MTR. Dans ce chapitre, nous évaluons l'impact de TTLs constants au contraire des TTLs exponentiels qu'on a utilisé dans le chapitre précédent, et nous développons une analyse approximative dans le cas où les temps entre-rencontres sont arbitrairement distribués. La conclusion est que, l'entre-rencontre exponentiel produit des délais stochastiquement inférieurs à ceux de l'entre-rencontre hyper-exponentiel, et le TTL exponentiel produit des délais stochastiquement supérieurs à ceux d'un TTL constant.

Ce chapitre évalue l'impact des TTLs constants et les temps de l'entre-rencontres arbitraires sur les performances du protocole de multicopie et de relais à deux sauts (MTR). Les hypothèses suivantes seront considérées. Tous les temps d'entre-rencontres sont indépendants et identiquement distribués avec une distribution de probabilité commune cumulée (CDF) $G(\cdot)$. La transmission des paquets se produit dans les périodes de rencontre entre les nœuds et elle est instantanée. Cette hypothèse sera justifiée dans les réseaux tolérants les délais (DTN), où le temps nécessaire pour envoyer le paquet à sa destination peut être très grands par rapport aux temps de transmission [dtn].

Dans la section 7.2 nous considérons le protocole MTR dans le cas où les TTLs sont tous constants et tous égaux, et les temps entre-rencontres sont exponentiellement distribués, c.-à-d. $G(t) = 1 - e^{-\lambda t}$. Nous développons l'analyse (non-Markovienne) qui nous permet de déterminer le CDF du délai de la livraison. Nous observons que les TTLs constants produisent des délais de livraison stochastiquement inférieurs à ceux de TTLs exponentiels. Plus précisément, nous avons trouvé que le CDF des délais de livraison, T_d , dans le cas où il y a $N - 1$ nœuds relais est

$$P(T_d < t) = 1 - e^{-\lambda N t} \Psi(\lambda, T, t)^{N-1}, \quad t > 0, \quad (\text{C.2})$$

avec

$$\Psi(\lambda, T, t) := 1 + \sum_{m=0}^{\lfloor \frac{t}{T} \rfloor} \sum_{k=0}^m \frac{(A_m)^{k+1} - (B_m)^{k+1}}{(k+1)!}, \quad (\text{C.3})$$

$$A_m := \lambda(t - mT), \text{ et } B_m := \lambda[t - (m+1)T]^+.$$

Dans la section 7.3, nous considérons le même scénario que celui de la section 7.2, mais nous relâçons l'hypothèse que les temps d'entre-rencontres sont exponentiellement distribués. Nous supposons qu'ils sont arbitrairement distribués, c.-à-d., le CDF $G(\cdot)$ est

arbitraire. Nous développons une formule approchée pour le CDF du délai de la livraison, et nous vérifions son exactitude dans le cas où $G(\cdot)$ est la distribution hyper-exponentielle. Nous observons que les délais de la livraison avec l'entre-rencontre exponentiel sont stochastiquement plus petits que ceux avec des temps d'entre-rencontres hyper-exponentiels. Plus précisément, nous trouvons une formule approchée de la transformée de Laplace du D_i , le temps nécessaire pour qu'un nœud relais envoie un paquet à la destination. En inversant cette transformée, nous obtenons la CDF de D_i . Ce qui nous donne la CDF du délai de la livraison dans le cas où il y a $N - 1$ relais qui est

$$P(T_d \geq t) = (1 - G_e(t))P(D_i > t)^{N-1}, \quad (\text{C.4})$$

avec $G_e(t)$ est la CDF de temps d'excès de l'entre-rencontres.

C.3 Conclusion

Cette thèse s'intéresse à l'impact de la mobilité dans les réseaux ad hoc mobiles. Elle comporte deux parties. Dans la première partie nous dressons l'état d'art de TCP dans les réseaux ad hoc. La principale conclusion est que la mobilité dégrade les performances de TCP, à cause de problèmes de routage et de partitions du réseau qu'elle occasionne. Pour surmonter ce problème, la plupart des solutions de TCP suggèrent de mettre TCP en état de veille dans les périodes des échecs du routage et de le réactiver quand le réseau est re-connecté. En s'appuyant sur ce constat, nous avons évalué dans la deuxième partie le gain réalisé quand la mobilité des nœuds est exploitée. Ceci est réalisé en permettant les nœuds de transmettre les paquets des autres à travers le protocole de relais à deux sauts.

Une grande partie du reste de la thèse a été consacré pour étudier le protocole de relais de deux sauts. Nous nous sommes tout d'abord intéressés aux performances des nœuds relais (débit et taille moyenne des files) en utilisant le formalisme des files d'attente. Un des résultats principaux est que le débit des nœuds relais dépend des modèles de mouvements aléatoires à travers leurs distributions stationnaires de position, et que ce débit est minimisé quand les nœuds bougent selon des modèles de mouvements aléatoires qui ont une distribution stationnaire uniforme de position.

L'amélioration du débit de réseau quand le relais à deux sauts est utilisé vient malheureusement avec une augmentation du délai de la livraison de paquet. Pour réduire ce délai entre une source et une destination, le chapitre 5 présente le protocole de multicopies et de relais à deux sauts (MTR), avec lequel la source génère plusieurs copies du paquet et qui ont tous des durées de vie limitées (TTLs). En utilisant le formalisme de la chaîne de Markov absorbante, nous avons réussi à trouver la distribution du délai de la livraison d'un paquet, la distribution du nombre de copies à l'instant de la livraison, et le nombre total moyen des copies générées avant la livraison. La dernière métrique est liée à l'énergie nécessaire pour livrer le paquet à sa destination.

Pour limiter la consommation d'énergie du protocole MTR, le chapitre 6 évalue les performances de deux extensions de MTR. La première est K-copies MTR pour laquelle le nombre maximal des copies dans le réseau est limité à K . Dans la deuxième, appelée

K-transmissions MTR, le nombre total des copies générées par la source est limité à K . Ces évaluations ont permis de trouver la valeur optimale de K qui minimise la consommation d'énergie en présence d'une contrainte sur les délais de la livraison. D'un autre côté, nous avons aussi évalué le codage de pertes (erasure coding en anglais), une approche qui réduit la variation du délai en la comparant avec la multicoPIe comme celle de MTR.

Dans le Chapitre 7, nous avons étudié l'impact d'avoir des TTLs constants et une distribution arbitraire d'entre-rencontres de nœuds sur MTR. Dans ce cas une analyse non-Markovienne était faite. En particulier, nous avons prouvé que les TTLs constants génèrent des délais stochastiquement plus petits que ceux des TTLs exponentiels, au contraire des temps d'entre-rencontres hyper-exponentiels qui rapportent des délais stochastiquement plus grands que ceux des temps exponentiels d'entre-rencontres.

C.4 Perspectives

Dans la deuxième partie de la thèse, nous avons considéré un seul paquet dans le système et une seule source. Cette hypothèse néglige les délais dus aux autres paquets dans le réseau. La relaxation de cette hypothèse dans MANET nécessite une analyse plus détaillée de la file d'attente d'un nœud relais. Dans ce cas, ce délai dépendra des trafics des sources ainsi de la mobilité des nœuds qui ont participé à la transmission du paquet à la destination.

D'autre part, nous avons supposé que quand deux nœuds se rencontrent, ils peuvent échanger des données avec succès. En supposant qu'un nœud peut suivre une certaine politique pour décider de transmettre ou non le paquet sera utile. Parceque dans ce cas, on peut trouver la politique optimale que les nœuds peuvent choisir afin de réduire une fonction de coût de l'énergie et du délai.

Dans cette thèse, nous avons évalué les performances du protocole à deux sauts quand les nœuds relais sont mobiles. Une autre direction de recherche qui peut être aussi considérée est quand les nœuds relais ne bougent pas, alors que la source et la destination bougent. Dans cette situation, la question qui surgit: quel est le meilleur emplacement des nœuds relais afin d'améliorer le délai ou le débit en tenant compte des possibles interférences entre les nœuds relais quand ils sont proches les uns des autres.

En conclusion, nous croyons que le contrôle de puissance dans les réseaux ad hoc mobiles dans le contexte des batteries à énergie limitée des nœuds est un axe de recherche important. Notre objectif principal dans cet axe sera d'augmenter la durée de vie des batteries et en même temps de diminuer le délai. Afin de réaliser un tel objectif, on peut employer le formalisme de programmation dynamique et de la file d'attente.

List of Acronyms

ACK	Acknowledgment
AODV	Ad hoc On-demand Distance Vector
ARF	Automatic Rate Fall-back
ARQ	Automatic Repeat Request
ATCP	Ad hoc TCP Proposal
a.s.	almost surely
cf.	confer
CCDF	Complementary Cumulative Distribution Function
CDF	Cumulative Distribution Function
CTS	Clear To Send
CW	Congestion Window Size of TCP
DSDV	Dynamic destination Sequenced Distance Vector
DSR	Dynamic Source Routing
ECN	Explicit Congestion Notification
ELFN	Explicit Link Failure Notification
ERDN	Explicit Routing Disconnection Notification
ERSN	Explicit Routing Successful Notification
FEC	Forward Error Correction
FIFO	First In First Out
e.g.	for example
i.e	id est
iff	if and only if
i.i.d.	independent and identically distributed
LST	Laplace Stieltjes Transform
MAC	Medium Access Control
MANET	Mobile Ad Hoc network
MP	Most Probable
MTR	Multicopies Two-hop Relay
NS	Network Simulator
OLSR	Optimized Link State Routing
OOO	Out-Of-Order Event

RB	Relay Buffer
RD	Random Direction
RRC	Route Reconstruction
RTO	TCP's Retransmission Timer
RTT	Round trip time
TTL	Time To Live
r.h.s	right hand side
rv	random variable
RWP	Random Waypoint
SANET	Static Ad Hoc network
w.p.1	with probability one
w.r.t.	with respect to

Bibliography

- [ACG03] G. Anastasi, M. Conti, and E. Gregori. IEEE 802.11 ad hoc networks: performance measurements. In *Proc. of the Workshop on Mobile and Wireless Network (MWM 2003) in conjunction with ICDCS 2003*, May 2003.
- [AGS99] M. Allman, D. Glover, and L. Sanchez. Enhancing TCP over satellite channels using standard mechanisms. RFC 2488, Jan. 1999.
- [AHM03] L. Andrew, S. Hanly, and R. Mukhtar. CLAMP: Differentiated capacity allocation in access networks. In *Proc. of IEEE Int. Performance Computing and Communications Conf.*, pages 451–458, Phoenix, AZ, USA, Apr. 2003.
- [AJ03] E. Altman and T. Jiménez. Novel delayed ACK techniques for improving TCP performance in multihop wireless networks. In *Proc. of the Personal Wireless Communications*, pages 237–253, Venice, Italy, Sep. 2003.
- [AMS82] D. Anick, D. Mitra, and M. M. Sondhi. Stochastic theory of a data-handling system with multiple sources. *Bell System Technical Journal*, 61(8):1871–1896, Oct. 1982.
- [APSS04] V. Anantharaman, S.-J. Park, K. Sundaresan, and R. Sivakumar. TCP performance over mobile ad hoc networks: A quantitative study. *Journal of Wireless Communications and Mobile Computing*, 4(2):203–222, Mar. 2004.
- [AWD04] M. Abolhasan, T. Wysocki, and E. Dutkiewicz. A review of routing protocols for mobile ad hoc networks. *Journal of Ad Hoc Networks, Elsevier*, 2(1):1–22, Jan. 2004.
- [BAD00] C. Barakat, E. Altman, and W. Dabbous. On TCP performance in heterogeneous networks: a survey. *IEEE Communications Magazine*, 38(1):40–46, Jan. 2000.
- [BB87] F. Baccelli and P. Brémaud. *Palm Probabilities and Stationary Queues*. Springer-Verlag, 1987.

- [BB97] A. V. Bakre and B.R. Badrinath. Implementation and performance evaluation of indirect TCP. *ACM/IEEE Transactions on Networking*, 46(3):260–278, Mar. 1997.
- [BDSZ94] V. Bharghavan, A. Demers, S. Shneker, and L. Zhang. MACAW: a media access protocol for wireless LAN's. In *Proc. of ACM SIGCOMM*, pages 212–225, London, UK, 1994.
- [Bet01] C. Bettstetter. Mobility modeling in wireless networks: Categorization, smooth movement, border effects. *ACM Mobile Computing and Communications Review*, 5(3):55–67, July 2001.
- [BHPC04] C. Bettstetter, H. Hartenstein, and X. Pérez-Costa. Stochastic properties of the random waypoint mobility model. *ACM/Kluwer Wireless Networks, Special Issue on Modeling and Analysis of Mobile Networks*, 10(5):555–567, Sept. 2004.
- [BK01] R. Boppana and S. Konduru. An adaptive distance vector routing algorithm for mobile, ad hoc networks. In *Proc. of IEEE INFOCOM*, Anchorage, Alaska, USA, Apr. 2001.
- [blu] Bluetooth Special Interest Group. Web site: <http://www.bluetooth.com>.
- [BMJ⁺98] J. Broch, A. D. Maltz, B. D. Johnson, Y.-C.Hu, and Jetcheva. A performance comparison of multi-hop wireless ad hoc network routing protocols. In *Proc. of ACM MOBIKOM*, pages 85–97, Dallas, TX, Oct. 1998.
- [BPSK97] H. Balakrishnan, V. Padmanabhan, S. Seshan, and R. Katz. A comparison of mechanisms for improving TCP performance over wireless links. *ACM/IEEE Transactions on Networking*, 5(6):756–769, Dec. 1997.
- [BS97] K. Brown and S. Singh. M-TCP: TCP for mobile cellular networks. *ACM SIGCOMM Computer Communication Review*, 27(5):19–43, Oct. 1997.
- [BSAK95] H. Balakrishnan, S. Seshan, E. Amir, and R. Katz. Improving TCP/IP performance over wireless networks. In *Proc. of ACM MOBIHOC*, pages 2–11, Berkeley, CA, USA, 1995.
- [BSK95] H. Balakrishnan, S. Seshan, and R. Katz. Improving reliable transport and handoff performance in cellular wireless networks. *ACM Wireless Networks*, 1(4):469–481, Dec. 1995.
- [BV05] J.-Y. Le Boudec and M. Vojnovic. Perfect simulation and stationarity of a class of mobility models. In *Proc. of IEEE INFOCOM*, Miami, FL, Apr. 2005.
- [CDA03] C. Cordeiro, S. Das, and D. Agrawal. COPAS: Dynamic contention-balancing to enhance the performance of tcp over multi-hop wireless networks. In *Proc. of IC3N*, pages 382–387, Miami, USA, Oct. 2003.

- [CHC⁺06] A. Chaintreau, P. Hui, J. Crowcroft, C. Diot, R. Gass, and J. Scott. Impact of human mobility on the design of opportunistic forwarding algorithm. In *Proc. Proc. of INFOCOM 2006*, Barcelona, Spain, Apr. 2006.
- [CJ03] T. Clausen and P. Jacquet. Optimized Link State Routing Protocol (OLSR). RFC 3626, Category: Experimental, Oct. 2003.
- [CJWK02] K. Chin, J. Judge, A. Williams, and R. Kermode. Implementation experience with MANET routing protocols. *ACM SIGCOMM Computer Communication Review*, 32(5):49–59, Nov. 2002.
- [CRVP98a] K. Chandran, S. Raghunathan, S. Venkatesan, and R. Prakash. A Feedback based scheme for improving TCP performance in ad hoc wireless networks. In *Conference on Distributed Computing Systems*, pages 472–479, Amsterdam, Netherlands, May 1998.
- [CRVP98b] K. Chandran, S. Raghunathan, S. Venkatesan, and R. Prakash. A feedback based scheme for improving TCP performance in Ad-Hoc wireless networks. In *Proc. of the International Conference on Distributed Computing Systems (ICDCS'98)*, Amsterdam, Netherlands, May 1998.
- [CXN03] K. Chen, Y. Xue, and K. Nahrstedt. On setting TCP's congestion window limit in mobile Ad Hoc networks. In *Proc. of IEEE ICC*, Anchorage, Alaska, USA, May 2003.
- [DA68] H. Dubner and J. Abate. Numerical inversion of laplace transforms by relating them to the finite fourier cosine transform. *Journal of the ACM*, 15(1):115–123, Jan. 1968.
- [DB01] T. Dyer and R. Boppana. A comparison of TCP performance over three routing protocols for mobile ad hoc networks. In *Proc. of ACM MOBIHOC*, pages 56–66, Long Beach, CA, USA, 2001.
- [DMT96] R. Durst, G. Miller, and E. Travis. TCP extensions for space communications. In *Proc. of ACM MOBICOM*, pages 15–26, Rye, NY, USA, 1996.
- [dtn] Delay tolerant research group. Web site: <http://www.dtnrg.org>.
- [Fel68] W. Feller. *Probability Theory and its Applications*. New York: John Wiley and Sons, Vol. 1 third ed., 1968.
- [FF96] K. Fall and S. Floyd. Simulation based comparisons of Tahoe, Reno, and Sack TCP. In *Computer Communications review*, Jul. 1996.
- [FJ93] S. Floyd and V. Jacobson. Random early detection gateways for congestion avoidance. *ACM/IEEE Transactions on Networking*, 1(4):397–413, Aug. 1993.

- [FZL⁺03] Z. Fu, P. Zerfos, H. Luo, S. Lu, L. Zhang, and M. Gerla. The impact of multihop wireless channel on TCP throughput and loss'. In *Proc. of IEEE INFOCOM*, San Francisco, USA, Apr. 2003.
- [GAGPK01] T. Goff, N. Abu-Ghazaleh, D. Phatak, and R. Kahvecioglu. Preemptive routing in ad hoc networks. In *Proc. of ACM MOBICOM*, pages 43–52, Rome, Italy, 2001.
- [GK00] P. Gupta and P. R. Kumar. The capacity of wireless networks. *ACM/IEEE Transactions on Information Theory*, 46(2), Mar. 2000.
- [glo] Global Mobile Information Systems Simulation Library GloMoSim. Web site: <http://pcl.cs.ucla.edu/projects/glomosim/>.
- [GMPS04] A. El Gamal, J. Mammen, B. Prabhakar, and D. Shah. Throughput-delay trade-off in wireless networks. In *Proc. of IEEE INFOCOM*, Hong Kong, Apr. 2004.
- [GNK05] R. Groenevelt, P. Nain, and G. Koole. The message delay in mobile ad hoc networks. *Performance Evaluation*, 62(1-4):210–228, Oct. 2005.
- [Gro05] R. Groenevelt. Stochastic models in mobile ad hoc networks. Technical report, Ph.D. Thesis, University of Nice Sophia Antipolis, Apr. 2005.
- [GS97] C. Grinstead and J. Snell. *Introduction to Probability*. American Mathematical Society, 1997.
- [GT02] M. Grossglauser and D. Tse. Mobility increases the capacity of ad hoc wireless networks. *ACM/IEEE Transactions on Networking*, 10(4):477–486, Aug. 2002.
- [GTB99] M. Gerla, K. Tang, and R. Bagrodia. TCP performance in wireless multi-hop networks. In *Proc. of the IEEE WMCSA*, New Orleans, LA, USA, 1999.
- [HAN04a] A. Al Hanbali, E. Altman, and P. Nain. A survey of TCP over ad hoc networks. *IEEE Communications Surveys and Tutorials*, 7(3):22–36, August 2004.
- [HAN04b] A. Al Hanbali, E. Altman, and P. Nain. A survey of TCP over mobile ad hoc networks. Technical Report RR-5182, INRIA, Sophia-Antipolis, May 2004.
- [Heo96] J. Heo. Improving the start up behavior of a congestion control scheme for tcp. In *Proc. of ACM SIGCOMM*, CA, USA, Aug. 1996.
- [HK99] T. Henderson and R. Katz. Transport protocols for Internet-compatible satellite networks. *IEEE JSAC*, 17(2):345–359, Feb. 1999.
- [HKG⁺05] A. Al Hanbali, A. A. Kherani, R. Groenoveit, P. Nain, and E. Altman. Impact of mobility on the performance of message relaying in ad hoc networks. Technical Report RR-5480, INRIA, Sophia-Antipolis, Jan. 2005.

- [HKG⁺06] A. Al Hanbali, A. A. Kherani, R. Groenovelt, P. Nain, and E. Altman. Impact of mobility on the performance of relaying in ad hoc networks. In *Proc. of IEEE INFOCOM*, Barcelona, Spain, Mar. 2006.
- [HNA06a] A. Al Hanbali, P. Nain, and E. Altman. Performance evaluation of packet relaying in ad hoc network. Technical Report RR-5860, INRIA, sophia-Antipolis, Mar. 2006.
- [HNA06b] A. Al Hanbali, P. Nain, and E. Altman. Performance of two-hop relay routing with limited packet lifetime. In *Proc. IEEE IWQOS*, pages 295–296, Short paper, New Haven, CT, USA, June 2006.
- [HNA06c] A. Al Hanbali, P. Nain, and E. Altman. Performance of two-hop relay routing with limited packet lifetime. In *Proc. IEEE/ACM VALUETOOLS*, Pisa, Italy, Oct. 2006.
- [HV02] G. Holland and N. Vaidya. Analysis of TCP performance over mobile ad hoc networks. *ACM Wireless Networks*, 8(2):275–288, Mar. 2002.
- [hyp] Broadband radio access networks (BRAN): High performance Local Area Network (HiperLAN) type 2. Technical Report 101 683 V1.1.1, ETSI.
- [IEE] IEEE. 802.11 WLAN standard. <http://standards.ieee.org/getieee802>.
- [IET] IETF. working group MANET. Web site: <http://www.ietf.org/>.
- [Jac98] V. Jacobson. Congestion avoidance and control. In *Proc. of ACM SIGCOMM*, Vancouver, Canada, Aug. 1998.
- [JMH03] D. Johnson, D. Maltz, and Y. Hu. The Dynamic Source Routing Protocol for Mobile Ad Hoc Networks (DSR). Internet draft, Apr. 2003.
- [JSAC01] C. Jones, K. Sivalingam, P. Agarwal, and J. Chen. A survey of energy efficient network protocols for wireless and mobile networks. *ACM Wireless Networks*, 7(4):343–358, 2001.
- [KK05] V. Kawadia and P. Kumar. A cautionary perspective on cross layer design. *IEEE Wireless Communication Magazine*, 12(1):3–11, Feb. 2005.
- [KKFT02] S. Kopparty, S. Krishnamurthy, M. Faloutous, and S. Tripathi. Split TCP for mobile ad hoc networks. In *Proc. of IEEE GLOBECOM*, Taipei, Taiwan, Nov. 2002.
- [KKT03] F. Klemm, S. Krishnamurthy, and S. Tripathi. Alleviating effects of mobility on tcp performance in ad hoc networks using signal strength based link management. In *Proc. of the Personal Wireless Communications*, pages 611–624, Venice, Italy, Sep. 2003.

- [Kle76] L. Kleinrock. *Queueing Systems, Volume II : Computer Applications*. John Wiley, 1976.
- [KM97] A. Kamerman and L. Monteban. Wavelan ii: A high-performance wireless lan for the unlicensed band. *Bell Labs Technical Journal*, pages 118–133, Summer 1997.
- [KTC01] D. Kim, C. Toh, and Y. Choi. TCP-BuS: Improving TCP performance in wireless ad hoc networks. *Journal of Communications and Networks*, 3(2):175–186, Jun. 2001.
- [Kur70] T. G. Kurtz. Solutions of ordinary differential equations as limits of pure jump markov processes. *Applied Probability*, 7:49–58, 1970.
- [LG01] S. Lu and M. Gerla. Split multipath routing with maximally disjoint paths in ad hoc networks. In *Proc. of IEEE ICC*, Helsinki, Finland, Jun. 2001.
- [LNT02] H. Lundgren, E. Nordstro, and C. Tschudin. Coping with communication gray zones in IEEE 802.11b based ad hoc networks. In *Proc. of the ACM Workshop on Wireless Mobile Multimedia*, pages 49–55, Atlanta, GA, USA, Sep. 2002.
- [Loy62] R.M. Loynes. The stability of a queue with non-independent inter-arrival and service times. In *Proc. Cambridge Philosophical Soc.*, pages 497–520, 1962.
- [LS01] J. Liu and S. Singh. ATCP: TCP for mobile ad hoc networks. *IEEE JSAC*, 19(7):1300–1315, Jul. 2001.
- [LXG03] H. Lim, K. Xu, and M. Gerla. TCP performance over multipath routing in mobile ad hoc networks. In *Proc. of IEEE ICC*, Anchorage, Alaska, May 2003.
- [Mik64] S. G. Mikhlin. *Integral Equations*. Pergamon Press, Oxford, 1964.
- [Mit04] M. Mitzenmacher. Digital fountains: A survey and look forward. In *Proc. of IEEE Information Theory Workshop*, San Antonio, TX, USA, Oct. 2004.
- [MSB00] J. Monks, P. Sinha, and V. Bharghavan. Limitations of TCP-ELFN for ad hoc networks. In *Proc. of Mobile and Multimedia Communications*, Tokyo, Japan, Oct. 2000.
- [Neu81] M. Neuts. *Matrix-Geometric Solutions in Stochastic Models: An Algorithmic Approach*. Johns Hopkins University Press, 1981.
- [NM81] M. Neuts and K. Meier. On the use of phase type distribution in reliability modelling of systems of small number components. *OR Spectrum*, 2(4):227–234, Dec. 1981.
- [ns2] The Network Simulator NS-2. Web site: <http://www.isi.edu/nsnam/ns/>.

- [PA00] V. Paxson and M. Allman. Computing TCP's retransmission timer. RFC 2988, Category: Standard Track, Nov. 2000.
- [PBRD03] C. Perkins, E. Belding-Royer, and S. Das. Ad hoc On-Demand Distance Vector (AODV) Routing. RFC 3561, Category: Experimental, Jul. 2003.
- [Per01] C. Perkins. *AD HOC Networking*. Adison-Wesly, Upper Saddle River, NJ, USA, 2001.
- [PFTV92] W. H. Press, B. P. Flannery, S. A. Teukolsky, and W. T. Vetterling. *Numerical Recipes in C: The Art of Scientific Computing*. Cambridge University Press, 1992.
- [PNL05] B. Liu P. Nain, D. Towsley and Z. Liu. Properties of random direction models. In *Proc. of IEEE INFOCOM*, Miami, FL, Mar. 2005.
- [PW94] C. Perkins and T. Watson. Highly dynamic Destination-Sequenced Distance-Vector routing (DSDV) for mobile computers. In *Proc. of ACM SIGCOMM*, London, UK, 1994.
- [RFB01] K. Ramakrishnan, S. Floyd, and D. Black. The addition of explicit congestion notification (ECN) to IP. RFC 3168, Category: Standards Track, Sep. 2001.
- [rfc81] Transmission Control Protocol. RFC 793, Sep. 1981.
- [SH03] T. Small and Z. J. Haas. The shared wireless infostation model: A new ad hoc networking paradigm. In *Proc. of ACM MOBIHOC*, Anapolis, MD, USA, 2003.
- [SH05] T. Small and Z. J. Haas. Resource and performance tradeoffs in delay-tolerant wireless networks. In *Proc. ACM SIGCOM Workshop on Delay-Tolerant Networks*, Philadelphia, PA, USA, Aug. 2005.
- [Spi] M. R. Spiegel. *Shaum's Outline of Theory and Problems of Laplace Transforms*. McGraw-Hill, New York.
- [TG99] K. Tang and M. Gerla. Fair sharing of MAC under TCP in wireless Ad Hoc networks. In *Proc. of IEEE Multiclass Mobility and Teletraffic for Wireless Communications Workshop*, Venice, Italy, Oct. 1999.
- [TK75] F. Tobagi and L. Kleinrock. Packet switching in radio channels: Part ii - the hidden terminal problem in Carrier Sense Multiple-Access modes and the busy-tone solution. *ACM/IEEE Transactions on Networking*, 23(12):1417–1433, 1975.
- [Toh97] C. Toh. Associativity-based routing for ad hoc mobile networks. *Journal of Wireless Personal Communications*, 4(2):103–139, Mar. 1997.

- [WJMF05] Y. Wang, S. Jain, M. Martonosi, and K. Fall. Erasure-coding based routing for opportunistic networks. In *Proc. SIGCOMM Workshop on DTN*, Philadelphia, PA, USA, Aug. 2005.
- [WZ02] F. Wang and Y. Zhang. Improving TCP performance over mobile ad hoc networks with out-of-order detection and response. In *Proc. of ACM MOBIHOC*, pages 217–225, Lausanne, Switzerland, Jun. 2002.
- [XBLG02] K. Xu, S. Bae, S. Lee, and M. Gerla. TCP behavior across multihop wireless networks and the wired internet. In *Proc. of the ACM Workshop on Wireless Mobile Multimedia*, pages 41–48, Atlanta, GA, USA, Sep. 2002.
- [XGQS03] K. Xu, M. Gerla, L. Qi, and Y. Shu. Enhancing TCP fairness in ad hoc wireless networks using neighborhood red. In *Proc. of ACM MOBICOM*, pages 16–28, San Diego, CA, USA, Sep. 2003.
- [XMB03] K. Xu, M. Gerla, and S. Bae. Effectiveness of RTS/CTS handshake in IEEE 802.11 based ad hoc networks. *Ad Hoc Networks Journal, Elsevier*, 1(1):107–123, Jul. 2003.
- [XS01] S. Xu and T. Saadawi. Does the IEEE 802.11 MAC protocol work well in multihop wireless ad hoc networks? *IEEE Communications Magazine*, 39(6):130–137, Jun. 2001.
- [XS02] S. Xu and T. Saadawi. Performance evaluation of TCP algorithms in multihop wireless packet networks. *Journal of Wireless Communications and Mobile Computing*, 2(1):85–100, 2002.
- [YSY03] L. Yang, W. Seah, and Q. Yin. Improving fairness among TCP flows crossing wireless ad hoc and wired networks. In *Proc. of ACM MOBIHOC*, pages 57–63, Annapolis, Maryland, USA, Jun. 2003.
- [ZNKT06] E. Zhang, G. Neglia, J. Kurose, and D. Towsley. Performance modeling of epidemic routing. In *Proc. of Networking*, pages 827–839, Coimbra, Portugal, May 2006.

Résumé

Cette thèse s'intéresse à l'impact de la mobilité sur les performances des réseaux ad hoc mobiles (MANETs en anglais). Elle comporte deux parties. La première partie de la thèse dresse un état-de-l'art du protocole TCP dans MANETs. La principale conclusion est que la mobilité dégrade les performances de TCP, à cause de problèmes de routage et de partitions du réseau qu'elle occasionne. Partant de ce constat, nous proposons et analysons dans la deuxième partie de la thèse des schémas de transmission qui s'appuient sur la mobilité. Plus précisément, chaque noeud peut servir de relais en l'absence de route directe entre la source et la destination. Nous nous sommes tout d'abord intéressés aux performances des noeuds relais (débit et taille moyenne des files) en utilisant le formalisme des files d'attente. Un des résultats principaux est que le débit des noeuds relais est minimisé quand les noeuds bougent selon des modèles de mouvements aléatoires qui ont une distribution stationnaire uniforme de position. Pour optimiser les performances du protocole de relais à deux sauts, particulièrement le délai de transmission, nous avons ensuite étudié le cas où un paquet peut avoir plusieurs copies dans le réseau, sous l'hypothèse où ces copies ont des durées de vie limitée. Les performances (délai, énergie consommée) ont été obtenus en utilisant le formalisme des chaînes de Markov absorbantes, ainsi que des modèles fluides. Nous avons appliqué nos résultats pour optimiser la consommation d'énergie en présence de contraintes sur les délais.

Mots-clés: Évaluation des performances, réseaux ad hoc mobiles, TCP, protocole de relais à deux sauts, analyse Markovienne, théorie de la file d'attente, modèle fluide.

Abstract

This thesis deals with the mobility impact on the performance of mobile ad hoc network (MANET). It contains two parts. The first part surveys the TCP protocol over MANET. The main conclusion is that mobility degrades the TCP performance. Since it induces frequent route failures and extended network partitions. These implications were the motivation in the second part to introduce and evaluate new transmission schemes that rely on the mobility to improve the capacity of MANET. More precisely, in the absence of a direct route between two nodes the rest of the nodes in the network can serve as the relay nodes. In the beginning, the focus was on the performance of the relay nodes (throughput and relay buffer size) using a detailed queueing analysis. One of the main results was that random mobility models that have uniform stationary distribution of nodes location achieve the lowest throughput of relaying. Next, in order to optimize the performance of the two-hop relay protocol, especially the delivery delay of packets, we evaluated the multicopy extension under the assumption that the lifetime of the packets is limited. The performance results (delivery delay, round trip time, consumed energy) were derived using the theory of absorbing Markov chains and the fluid approximations. These results were exploited to optimize the total energy consumed subject to a constraint on the delivery delay.

Keywords: Performance evaluation, mobile ad hoc Networks, TCP, two-hop relay protocol, epidemic routing, Markovian analysis, queueing theory, fluid model.