

UC Santa Cruz

UC Santa Cruz Previously Published Works

Title

Link Dynamics in MANETs with Restricted Node Mobility: Modeling and Applications

Permalink

<https://escholarship.org/uc/item/0wn8z9hd>

Author

Garcia-Luna-Aceves, J.J.

Publication Date

2009-09-01

Peer reviewed

Link Dynamics in MANETs with Restricted Node Mobility: Modeling and Applications

Xianren Wu, *Member, IEEE*, Hamid R. Sadjadpour, *Senior Member, IEEE*,
and J. J. Garcia-Luna-Aceves, *Fellow, IEEE*

Abstract—We present statistical models to accurately evaluate the distribution of the lifetime of wireless links in a mobile ad hoc network (MANET) in which nodes move randomly within constrained areas. We show that link lifetime can be computed through a two-state Markov model and further apply the computed statistics to the optimization of segmentation schemes of an information stream. Summarizing all these results, we further provide a comprehensive analysis on throughput, delay, and storage requirements for MANETs with restricted node mobility.

Index Terms—Link dynamics, restricted mobility, mobility modeling, Markov model, MANETs.

I. INTRODUCTION

MOBILITY brings fundamental changes to the performance and design of all aspects of protocol stacks in mobile ad hoc networks (MANET). Understanding the statistics of link lifetime is crucial for the accurate analysis of MANET parameters and protocols. For example, the performance of routing protocols in a MANET exhibits a direct relationship to the mean value of link lifetime [1]. Interestingly, as critical as link lifetime is for the performance of the protocol stack in MANET, not much effort has been devoted to the analytical modeling of link lifetime as a function of node mobility, which is a defining attribute of MANETs! As a result, link lifetime in MANET has been analyzed mostly through simulations, and the analytical models of channel access and routing protocols for MANETs have not represented the temporal nature of MANET links accurately. Similarly, most studies of routing-protocol performance have relied exclusively on simulations, or had to use limited models of link availability, to address the dynamics of paths impacting routing protocols.

Manuscript received March 20, 2007; revised May 17, 2008 and April 8, 2009; accepted April 8, 2009. The associate editor coordinating the review of this paper and approving it for publication was R. Berry.

X. Wu and H. R. Sadjadpour are with the University of California at Santa Cruz, Santa Cruz, CA 95064 (e-mail: {wuxr, hamid}@soe.ucsc.edu).

J. J. Garcia-Luna-Aceves is with the University of California at Santa Cruz, Santa Cruz, CA 95064, and is also affiliated with Palo Alto Research Center, 3333 Coyote Hill Road, Palo Alto, CA 94304 (e-mail: jj@soe.ucsc.edu).

This work was partially sponsored by the U.S. Army Research Office under grants W911NF-04-1-0224 and W911NF-05-1-0246, by the National Science Foundation under grant CCF-0729230, by the Defense Advanced Research Projects Agency through Air Force Research Laboratory Contract FA8750-07-C-0169, and by the Baskin Chair of Computer Engineering. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of the U.S. Government.

Digital Object Identifier 10.1109/TWC.2009.070309

This paper provides the most accurate analytical model of link lifetime in MANETs to date, and characterizes link lifetime as a function of node mobility. The importance of this model is twofold. First, it enables answering many questions regarding fundamental tradeoffs in throughput, delay, and storage requirements in MANETs, as well as the relationship between many cross-layer design choices (e.g., information block segmentation) and network dynamics (e.g., how long links last in a MANET). Second, it enables the development of analytical models for channel access and routing schemes by allowing such protocols to use link lifetime expressions that are accurate with respect to simulations based on widely-used mobility models.

Recently, Samar and Wicker [2], [3] pioneered the analytical evaluation of link dynamics. They further provided important insight on the importance of an analytical formulation of link dynamics on the optimization of protocols being designed. However, Samar and Wicker assume that communicating nodes maintain constant speed and direction in order to evaluate the distribution of link lifetime. This simplification overlooks the case in which either of the communicating nodes change their speeds or directions while they are in transmission range of one another. As a result, the results predicted by Samar and Wicker's model could deviate from reality, being more conservative and underestimating the distribution of link lifetime [2], [3], especially when the ratio R/v between the radius of the communication range R to the node speed v becomes large, such that nodes are likely to change their velocity and direction during an exchange.

The first contribution of this paper consists of introducing a two-state Markov model that better describes the mobility behaviors for communicating nodes. The proposed model shows a marked improvement in characterizing the statistics of link lifetime, while subsumes the model of Samar and Wicker [2], [3] as a special case. By characterizing link lifetime, we further study the crosslayer optimization problem on segmentation of the information stream and provide solutions to maximize the end-to-end throughput.

Another contribution of the paper is to provide a comprehensive coverage of MANETs with restricted mobility, where each node moves within a constrained area. These networks play an important role in the real world, where nodes usually travel only a portion of the entire network. As published in the *information assurance framework* [4] from the National Security Agency, such networks represent the more realistic scenarios for tactical users, especially for the users deployed in

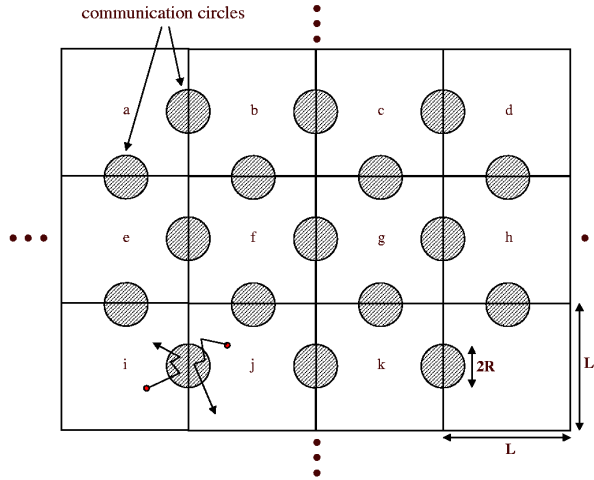


Fig. 1. Model of network structure.

the division and rear area. The only prior work addressing this issue of which we are aware is that of Groenevelt et al. [5]. It covers delay aspects of such networks, but only for the case of one-dimensional restricted mobility. For this reason, we strive to provide the first thorough analysis (to the best of authors' knowledge) of two-dimensional restricted mobility networks on link dynamics, optimal segmentation of information stream, throughput, delay, and storage tradeoffs.

The paper is organized as follows. Section II describes system models including network and mobility models in the paper. Section III presents the proposed two-state Markov model and results on link lifetime, along with simulation results for model validation. Section IV uses the derived statistics of link lifetime in section III for the problem of optimal segmentation of information stream. Section V provides a thorough analysis of throughput, delay and storage capacity of a MANET with restricted mobility, followed by concluding remarks in section VI.

II. SYSTEM MODEL

In many tactical applications [4], nodes of a MANET traverse only a small portion of the entire area covered by the network. We consider a square or rectangular area partitioned into squarelets similar to prior analytical models of MANETs and as depicted in Fig. 1. The entire network is divided into multiple squarelets, which we call *cells*, and each cell is of size $L \times L$.

Communication between nodes in neighboring cells is allowed around their cell boundaries and all nodes transmit with uniform power. According to the protocol model [6], the allowable communication circle should be deliberately designed to avoid excessive interference to nearby cells and to satisfy protocol model. Referring to the design in [7], a feasible solution is to choose circular regions centered at cell boundaries, as depicted in Fig. 1. It should be noted that nodes communicate only in the communication circle (even if they are able to when they are not in it).

A typical communication session between two nodes involves several control and data packet transmissions. Depending on the protocol, nodes may be required to transmit beacons

to their neighbors to synchronize their clocks for a variety of reasons (e.g., power management, frequency hopping). Nodes can find out about each other's presence by means of such beacons, or by the reception of other types of signaling packets (e.g., HELLO messages). Once a transmitter knows about the existence of a receiver, it can send data packets, which are typically acknowledged one by one, and the MAC protocol attempts to reduce or avoid those cases in which more than one transmitter sends data packets around a given receiver, which typically causes the loss of all such packets at the receiver. To simplify our modeling of link lifetimes, we assume that the proper mechanisms are in place for neighboring nodes to find each other, and that all transmissions of data packets are successful as long as they do not last beyond the lifetime of the wireless link between transmitter and receiver. Relaxing this simplifying assumption is the subject of future work, as it involves the modeling of explicit medium access control schemes (e.g., [8]).

Nodes are mobile, initially randomly and uniformly distributed over these cells. Nodes move according to the widely-used random direction mobility model (RDMM) [9]–[13], which improves on the random waypoint mobility model to have a uniform stationary spatial node distribution. Nodes movements are independent and identically distributed (iid) and can be described by a continuous-time stochastic process. The continuous movement of nodes is divided into mobility epochs during which a node moves at constant velocity, i.e., fixed speed and direction. But the speed and direction varies from epoch to epoch. The time duration of epochs is denoted by a random variable τ , assumed to be exponentially distributed with parameter λ_m . Its complementary cumulative distribution function (CCDF) $F_\tau(x)$ can be written as [11].

$$F_\tau(x) = \exp(-\lambda_m x) \quad (1)$$

The direction during each epoch is assumed to be uniformly distributed over $[0, 2\pi)$ and the speed of each epoch is uniformly distributed over $[v_{min}, v_{max}]$, where v_{min}, v_{max} specify the minimum and maximum speed of nodes respectively. Speed, direction and epoch time are mutually uncorrelated and independent over epochs. The motion of nodes will be reflected when hitting cell boundaries but in this case, it is still considered within the same epoch.

The movement of each node is restricted into the cell where it is initially located. Each source node randomly chooses its destination and in most cases, the source and destination nodes are not within the same cell. As a result, most data traffic need to travel across cells and links over neighboring cells are focal points for such networks. The analysis of this paper is focused on inter-cell links but our analysis can be also extended to intra-cell links or nodes with unrestricted mobility.

III. LINK LIFETIME

A bidirectional link exists between two nodes if they are within communication circle of each other. In this paper, we do not consider unidirectional links, given that the vast majority of channel access and routing protocols use only bidirectional links for their operation. Hence, we will refer to bidirectional links simply as links for the rest of this paper.

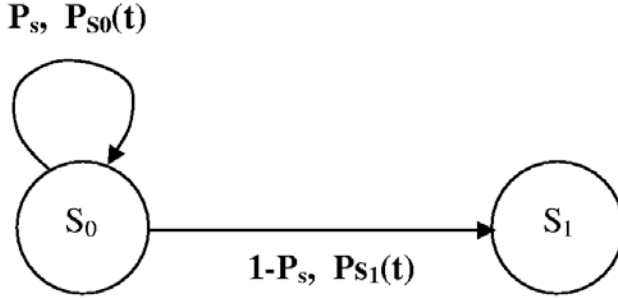


Fig. 2. Two-state Markov model for S-LLT evaluation.

When a data packet starts at time t_0 , positions of nodes (e.g. nodes m_a and m_b) during communication session could be anywhere inside communication circle. The end points of a link locate inside the communication circle and the distribution of them should be a distribution conditioned on such fact. Since we assume the uniform stationary distribution from RDMM model, the assumption lead to the result that the end points are also uniformly distributed in the communication circle as if in the stationary state. Let B (bits/s) be the transmission rate of a data packet, L_p be the length of the data packet, and $t_0 + T_a$ (or $t_0 + T_b$) denotes the moment a node m_a (or m_b) is moving out of communication range. A packet can be successfully transferred only if nodes m_a and m_b stay within communication circle during the entire communication session, that is,

$$L_p/B \leq T_L \quad (2)$$

$$T_L = \min(T_a, T_b). \quad (3)$$

T_L is the link lifetime (LLT) which dictates the maximum possible data transfer duration. Statistically, T_a and T_b specify the distribution of residence time that measures the duration of the time, for either nodes m_a or m_b , starting from a random location inside the communication circle with equal probability and continuously stay inside the communication circle before finally moving out of it.

Given that nodes m_a and m_b start from a random location inside the communication circle, then T_a and T_b , respectively, will be random variables representing the residual time that each node remains in the communication circle before first moving outside of it. Since the motions of nodes are iid, the distribution of T_a and T_b is the same. We then define such distribution as the *single-node link lifetime* (S-LLT) distribution. Furthermore, its complementary cumulative distribution function (CCDF) is denoted by $F_S(t)$, i.e., $F_S(t) = P(T_a \geq t) = P(T_b \geq t)$. Clearly, we can compute the link CCDF $F_L(t)$ as

$$F_L(t) = F_S^2(t). \quad (4)$$

The link outage probability P_{L_p} associated with a particular packet length L_p can be evaluated as

$$P_{L_p} = P(T_L \leq \frac{L_p}{B}) = 1 - F_L(\frac{L_p}{B}). \quad (5)$$

A. Single-Node Link Lifetime (S-LLT)

From the above, it follows that the essence of modeling link dynamics in MANETs consists of evaluating the distribution of S-LLT, because it reflects the link dynamics resulting from the motions of nodes. S-LLT measures the duration of time for a node to continuously stay inside the communication circle of another node. In our model, this range is a circle.

We also know that the movement of nodes consists of a sequence of mobility epochs. Let A_s be the starting point of current mobility epoch and its position will be uniformly distributed over communication circle [11]. The end point of the current epoch is denoted by A_d , and A_d may be anywhere in the cell, i.e., inside or out of the communication circle. In the case that A_d is located inside the communication circle, it serves as the starting point (i.e., new A_s) for the next epoch and the whole process is repeated. In the evaluation of S-LLT, the repeating procedure ends when the final A_d is out of the communication circle.

As illustrated in Fig. 2, the procedure for evaluating the S-LLT can be modeled as a two-state Markov process. The sojourn state S_0 represents the scenario where the end point A_d of current epoch is located inside the communication circle, while the departing state S_1 refers to the complementary scenario where A_d will be out of communication circle. Compared to the model by Samar and Wicker [2], [3], in which only the last scenario (i.e., state S_1) is considered, the two-state Markov model more accurately reflects the motion of nodes, which naturally expects better results in evaluating link dynamics.

Let P_s be the *residence probability*, which denotes the probability that A_d is located inside the communication circle. The probability distribution function (PDF) $p_{S_0}(t)$ specifies the distribution of sojourn time of mobility epochs when a node stays in state S_0 . Correspondingly, the PDF $p_{S_1}(t)$ is used to measure the distribution of departing times when nodes move out of communication circle and switch to state S_1 .

Before eventually moving out of the communication circle, i.e., being switched to the departing state S_1 , nodes may stay at the residence state S_0 multiple times. Let N_i be the integer variable counting the number of times for a node to remain in state S_0 , and $\{S_{0,0}, \dots, S_{0,N_i-1}\}$ be the associated random variables that specify the duration of time of mobility epochs for each return.

Clearly, $\{S_{0,0}, \dots, S_{0,N_i-1}\}$ are random variables of the same distribution but correlated. However, to make our problem more tractable, we assume that $\{S_{0,0}, \dots, S_{0,N_i-1}\}$ are statistically i.i.d random variables of distribution $p_{S_0}(t)$. Our simplifying assumption deviates the final result slightly from the real situation when the residence probability becomes larger. However, as we will see later, our model still provides good approximations, even with a large residence probability.

We define S_1 as the random variable measuring the departing time with distribution $p_{S_1}(t)$. Simply, one can evaluate conditional single-node link lifetime $T_S(N_i)$ as

$$T_S(N_i) = \sum_{n=0}^{N_i-1} S_{0,n} + S_1 \quad (6)$$

$$P(N_i = K) = P_s^K \quad (7)$$

The characteristic function $U_{T_S}(\theta)$ for the S-LLT T_S can now be evaluated as

$$\begin{aligned} U_{T_S}(\theta) &= E(e^{j\theta T_S}) \\ &= \sum_{k=0}^{\infty} E(e^{j\theta(\sum_{n=0}^{k-1} S_{0,n} + S_1)}) P(N_i = k) \\ &= \frac{U_1(\theta)}{1 - U_0(\theta)P_s} \end{aligned} \quad (8)$$

where $U_0(\theta)$ and $U_1(\theta)$ are the characteristic functions of $p_{S_0}(t)$ and $p_{S_1}(t)$ respectively.

When the ratio of radio range to node's speed is relatively small, A_d will be mostly located outside of the communication circle. Consequently, one will have $P_s \ll 1$. Given that $U_0(\theta)$ is the characteristic function of $p_{S_0}(t)$, one has $|U_0(\theta)| \leq 1$. Finally, it is clear that $U_0(\theta)P_s \ll 1$. Therefore, Eq. (8) can be approximated as

$$U_{T_S}(\theta) \approx U_1(\theta). \quad (9)$$

For clarification purposes, we call Eq. (8) as the Enhanced S-LLT (ES-LLT), which is based on the two-state Markov model. The approximation in Eq. (9) is called Approximated S-LLT (AS-LLT), and it reflects the scenario considered by Samar and Wicker [2], [3]. As we will see later, the analytical expression of AS-LLT is the same as the expression in [2], [3], except for a normalization factor.

To evaluate the S-LLT T_S , we need to evaluate P_s , $p_{S_0}(t)$, and $p_{S_1}(t)$, which we do next. Let z_d denote the least distance to be traveled by node to move out of the communication circle, starting from the position A_s with the direction and speed v being kept unchanged. The probability P_s can now be evaluated through z_d as

$$\begin{aligned} P_s &= E_v(P_s(v)), \\ P_s(v) &= \int_{z_d} P(\tau \leq \frac{x}{v}) p_{z_d}(x) dx = \int_{z_d} (1 - F_\tau(\frac{x}{v})) p_{z_d}(x) dx \\ &= \int_{z_d} (1 - \exp(-\lambda_m x/v)) p_{z_d}(x) dx \end{aligned} \quad (10)$$

where $P_s(v)$ is the conditional probability of P_s on v . $p_{z_d}(x)$ is PDF of z_d and from extending the work by Hong and Rappaport [14], we know that it can be calculated as

$$p_{z_d}(x) = \begin{cases} \frac{2}{\pi R^2} \sqrt{R^2 - (\frac{x}{2})^2}, & \text{for } 0 \leq x \leq 2R \\ 0, & \text{elsewhere} \end{cases} \quad (11)$$

where R specifies the radius of the communication circle.

$p_{S_0}(t)$ is the PDF of the time duration for nodes to return to the state S_0 . Conditioning on speed v and assuming that the starting time is at time 0, $p_S(t)$ is the probability of the node changing its velocity at time t conditioned on that A_d is located inside the communication circle. Hence,

$$p_{S_0}(t) = E_v(p_{S_0}(t|v)) \quad (12)$$

$$\begin{aligned} p_{S_0}(t|v) &= \frac{1}{P_s} P(t = \tau, z_d \geq v\tau|v) \\ &= \frac{1}{P_s} \lambda_m e^{-\lambda_m t} \int_{vt}^{2R} p_{z_d}(x) dx \\ &= \begin{cases} \frac{4\lambda_m e^{-\lambda_m t}}{\pi P_s} \left[\frac{\pi}{4} - \frac{vt}{4R} \sqrt{1 - (\frac{vt}{2R})^2} + \sin^{-1}(\frac{vt}{2R}) \right], & 0 \leq t \leq \frac{2R}{v} \\ 0, & \text{elsewhere} \end{cases} \end{aligned} \quad (13)$$

TABLE I
RESIDENCE PROBABILITY P_s .

R (m)	v (m/s)		
	1	10	20
100	0.09	0.01	0.005
200	0.17	0.02	0.01

where $p_{S_0}(t|v)$ is the conditional PDF on v .

$p_{S_1}(t)$ can be evaluated in much the same way as we have done for $p_{S_0}(t)$. Conditioning on speed v and assuming that the starting time is at time 0, $p_{S_1}(t)$ is simply the probability of the node moving out of the communication circle at time t with velocity being kept constant. Hence,

$$\begin{aligned} p_{S_1}(t) &= E_v(p_{S_1}(t|v)) \\ p_{S_1}(t|v) &= \frac{1}{1 - P_s} P(t = \frac{z_d}{v}, z_d \leq v\tau|v) \\ &= \frac{1}{1 - P_s} P(\tau \geq t) P(z_d = vt)(vt)' \\ &= \begin{cases} \frac{4e^{-\lambda_m t}}{\pi(1 - P_s)} \frac{v}{2R} \sqrt{1 - (\frac{vt}{2R})^2}, & 0 \leq t \leq \frac{2R}{v} \\ 0, & \text{elsewhere} \end{cases} \end{aligned} \quad (14)$$

where $p_{S_1}(t|v)$ is the conditional PDF on v . A detailed examination of Eq. (14) reveals that it shares the same core analytical expression of link lifetime distribution of Eq. (15) in [3], with the only exception that a normalization factor $e^{-\lambda_m t}/(1 - P_s)$ accounts for the probability of nodes leaving for state S_1 . It implies that AS-LLT formula, solely relying on $p_{S_1}(t)$, gives the same link lifetime distribution as in [3].

B. Model Validations

In the simulation, there are a total of 100 nodes randomly placed for each $1000m \times 1000m$ square cell. Each node has the same transmit power and two profiles of the radio transmission range are chosen for simulation. Both are within the coverage of IEEE 802.11 PHY layer and they are $\{200m, 100m\}$. After initial placement, nodes keep moving continuously according to the RDMM model. The mobility parameter λ_m is chosen to be $\lambda_m = 4$ and three different speeds are simulated $v \in \{1, 10, 20\}$ (m/s), from pedestrian speed to normal vehicle speed. Combining the power profile and velocity profile, six different scenarios are simulated $\{I : (200m, 1m/s); II : (100m, 1m/s); III : (200m, 10m/s); IV : (100m, 10m/s); V : (200m, 20m/s); VI : (100m, 20m/s)\}$.

Nodes are randomly activated to randomly choose destination node for data transmission. The traffic of activated nodes are supplied from a CBR source with a packet rate 0.5p/s. Given that the choice of specific MAC layer and routing protocol may affect the results, we assume perfect MAC and routing, rendering zero delays or losses due to such functionality, enabling the simulation to capture statistics solely due to mobility.

Table I describes the residence probability P_s for all six scenarios. As shown in Table I, the residence probability increases with the relative radius $\text{ReR} \frac{R}{v}$, indicating that it is more likely for nodes with larger ReR to stay inside the communication circle.

It is worthy of noting that the two-phase Markov model is a general model able to evaluate other networks with the

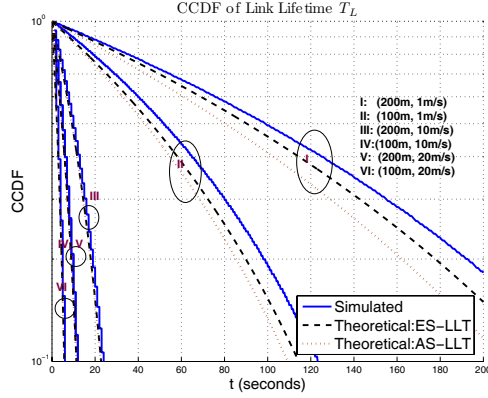


Fig. 3. Link lifetime T_L (inter-cell): simulated, ES-LLT(Markov), AS-LLT.

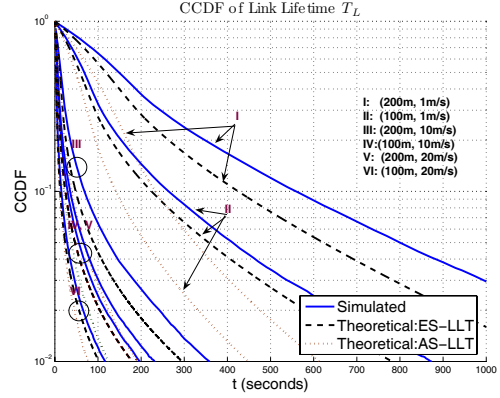


Fig. 4. Link Lifetime T_L (intra-cell): simulated, ES-LLT(Markov), AS-LLT.

two building blocks $p_{S_0}(t)$ and $p_{S_1}(t)$ adapted for the specific network and mobility models. We have applied the two-phase model to model random waypoint mobility model and obtained similar results. Figs. (3) and (4) present the results of link lifetime with ES-LLT and AS-LLT formula for both intra-cell¹ and inter-cell links. The results clearly confirm that the two-state Markov model is a powerful tool to accurately model link dynamics of link lifetime distribution as a function of node mobility. It can be observed that the ES-LLT formula, obtained from the Markov model, are closer to the simulations in all scenarios. On the other hand, the AS-LLT formula with the simplified assumptions corresponding to the model by Samar and Wicker [2], [3] gives good approximation to the simulations only for small values of $\text{ReR } \frac{R}{v}$ and greatly deviates from the simulations when $\text{ReR } \frac{R}{v}$ becomes large, i.e., larger residence probability P_s and larger possibility for nodes to stay inside communication circle. Furthermore, it is also clear that lifetime of inter-cell links is much shorter than intra-cell links, indicating the bottleneck effect from inter-cell links on network throughput.

IV. SEGMENTATION SCHEMES AND THEIR OPTIMIZATION

In wireless communication, source information stream usually needs to be segmented into a sequence of fixed-length information blocks for transmission. These information blocks will be further processed (e.g. channel encoding) to fit into various transmission schemes. Given that nodes move in a MANET, the data transfer can be temporarily broken if any link on the path to the destination is broken. An alternative path may not be available immediately and significant delay could be expected, particularly for sparse networks. However, it may have already resulted in timeout of the communication session and the upper layer needs to initiate another session to retransfer the whole information block. For example, for P2P application, the information is divided into multiple smaller blocks for transmission. If data transfer is failed in the middle of transmission of one block, the entire block usually needs to be retransmitted rather than continue transmission from the broken point.

¹Please refer to reference [15] for derivations of the results on intra-cell links.

Within the context of highly dynamic environment in MANETs, it is quite important to design segmentation schemes or use information data packet (information block) lengths that maximize the end-to-end throughput. If a data-packet length is too long, frequent link breaks could lead to significant packet dropout during the transfer. On the other hand, if data packet length is too short, the overhead from sub layer processing could significantly reduce the effective throughput. Judicious design of segmentation scheme as a function of link dynamics can be of great importance in maximizing throughput of MANETs, being a cross-layer optimization problem. However, this problem remain almost undeveloped since its solution necessitates knowledge of statistics of link lifetime. With the computed CCDF in section III, we are now able to provide segmentation schemes optimized on various systematic constraints.

We consider store-and-forward scheme similar to the one in MANETs with unlimited mobility [16], [17]. Source node splits information stream to relay nodes in its neighbor cells. Each relay stores information in the queue and delivers information from the queue only when it meets another relay nodes or the destination node in another cell. By doing it, usage of relays is minimized to improve the overall throughput. In the scheme, we are concerned with the lifetime of point-to-point links.

When the length of data packets (information streams) is constant, it is fairly natural to ask what will be the optimal packet length. For every packet length L_p , we know that there is an associated link outage probability P_{L_p} specifying the probability of link breakage during packet transfer. Every dropped information packet during link outage needs to be retransmitted and therefore reduces the effective throughput.

One segmentation scheme is to simply choose the maximum possible packet length L_0 that satisfies a pre-defined link outage probability requirement. We term this scheme as link outage priority design (LOPD) and it can be described as

$$L_0 = \max_{L_p} P_{L_p} \leq \omega_p \quad (16)$$

where ω_p is a constant to specify the link dropout probability requirement.

Alternatively, we can use a cost function $C(L_p, P_{L_p})$ that incorporates the negative effect from the packet retransmission

into evaluating the effective throughput $T(L_p)$ for a specific packet length L_p . It is worthy of noting that the cost function $C(L_p, P_{L_p})$ could be a systematic constraint from upper layer to consider the negative effects from delay and packet re-transmissions, etc. Further, optimizing the effective throughput $T(L_p)$ gives the optimal packet length L_0 . Consequently, this segmentation scheme is termed as link throughput priority design (LTPD).

In the LTPD design, when the packet length is L_p , we can describe the effective throughput $T(L_p)$ function as

$$T(L_p) = (1 - P_{L_p}) \cdot L_p - C(L_p, P_{L_p}) \cdot P_{L_p} \cdot L_p \quad (17)$$

The optimal packet length L_0 will be the one that maximizes the effective throughput

$$L_0 = \max_{L_p} T(L_p) \quad (18)$$

Normally, P_{L_p} is a monotonically decreasing function. When the cost function is chosen to be a constant penalty value, i.e., $C(L_p, P_{L_p}) = C$, by taking the derivative with respect to L_p , the optimal packet length L_0 is the value satisfying

$$1 - (1 + C)P_{L_0} = (1 + C)L_0 \left. \frac{dP_{L_p}}{dL_p} \right|_{L_p=L_0}. \quad (19)$$

In Fig. 5, we exploit the application of link lifetime distribution to the optimization of segmentation scheme using the same examples of the previous section. For illustration purpose, the cost function for our example of LTPD design is chosen as a constant penalty value 2, (i.e., $C(L_p, P_{L_p}) = 2$). However, it should be noted that the practical cost function can be much more complicated and determined by upper layer for a cross-layer optimization solution. Computing the optimum choice for $C(L_p, P_{L_p})$ is beyond the scope of this paper. The effective throughput $T(L_p)$ is computed for every L_p and drawn for all three methods: Simulation, ES-LLT (Markov model), and AS-LLT. As expected, ES-LLT approximates simulation very well, while AS-LLT tends to conservatively underestimate the effective throughput for larger ReR. In addition, all curves of the effective throughput (either Simulation, ES-LLT, or AS-LLT formula) are well behaved as convex functions with numerical solution readily available.

The optimized solutions $\frac{L_0}{B}$ of both LOPD and LTPD protocols on information segmentation are illustrated in Fig. 6. In the simulation, the link outage tolerance of LOPD design is set to be $\omega_p = 0.1$, i.e., the maximum link outage probability should be less than 10%. Two key observations are made: (1) For both LTPD and LOPD designs, the ES-LLT (Markov model) approaches the simulated optimal solution better and signifies substantial improvement of throughput over the AS-LLT model ([2], [3]); and (2) LTPD design suggests a balanced design between longer packet and larger retransmission rate to offer higher throughput over LOPD design. LOPD design, on the other hand, tends to be more conservative on the throughput but resulting in less packet retransmission.

Another important observation from Fig. 6 is that the optimal packet (information block) length design, obtained

from either the simulation or Markov ES-LLT formula, exhibit linear proportion to the ReR value $\frac{R}{v}$. It suggests that mathematically, the optimal information segmentation should follow the rule²

$$\frac{L_0}{B} = \Theta\left(\frac{R}{v}\right). \quad (20)$$

V. ANALYSIS OF THROUGHPUT, AVERAGE DELAY AND STORAGE

A. Throughput

We consider again the store-and-forward scheme in Section IV, where source node splits information stream to relay nodes in its neighbor cells, each relay stores information in the queue and delivers information from the queue only when it meets another relay nodes or the destination node in another cell. We also assume that every relay node maintains a separate queue for each source-destination pair and the queue is served in a First-Come-First-Serve (FCFS) manner. Because all cells resemble each other and nodes have iid movements, it is clear that all such queues are similar. Furthermore, we adopt a conservative scenario in which only one node per cell can act as the relay node of a specific route for later delay analysis. In reality, every node can act as a relay, which leads to less delay but a much more complex network of queues.

To facilitate our analysis, distribution of link interarrival time (LIT) for inter-cell links is summarized in the following Theorem and the proof is provided in Appendix.

Theorem 1: ³Let nodes A and B be moving independently of each other in two adjacent square cells of size $L \times L$. Their movement follow the RDMM model and are of average speed $E(v)$. Then LIT of such inter-cell links between nodes is approximated as an exponential distribution with parameter λ_I , where λ_I and the mean time of I are given by

$$\lambda_I \approx \frac{\pi^2 \cdot E(v) \cdot R^3}{2L^4} \quad (21)$$

$$E(I) \approx \frac{2L^4}{\pi^2 \cdot E(v) \cdot R^3} \quad (22)$$

For every cell, there should be at least one node inside the cell in order to maintain the connectivity of the network. Let $a(n) = \frac{L^2}{A_N}$ be the fractional cell size, where A_N is the overall size of the network and n is the number of nodes in the network. The connectivity requirement necessitates [18] that only when $a(n) \geq \frac{2 \log(n)}{n}$, each cell has at least one node with high probability (whp), i.e., with probability $\geq 1 - \frac{1}{n}$. In this case, each cell will have $\Theta(na(n))$ nodes inside whp [18].

²We recall the following notation: (i) $f(n) = O(g(n))$ means that there exists a positive constant c and integer N such that $f(n) \leq cg(n)$ for $n > N$. (ii) $f(n) = \Theta(g(n))$ means that there exist positive constants c_1, c_2 and M , such that $0 \leq c_1 g(n) \leq f(n) \leq c_2 g(n) \forall n > M$.

³When evaluating LIT, the starting positions of nodes should be inside the communication circle rather than from the uniform stationary state. It may take a transient phase of a duration of $\Theta(L/v)$ for nodes' position to converge to stationary state. Since the transient phase is usually short and the effect of the transient phase on the LIT is negligible, the proof of LIT will consider nodes' positions starting from the uniform stationary state. It should be noted that the evaluation of LIT becomes exact if both nodes start from their steady-state.

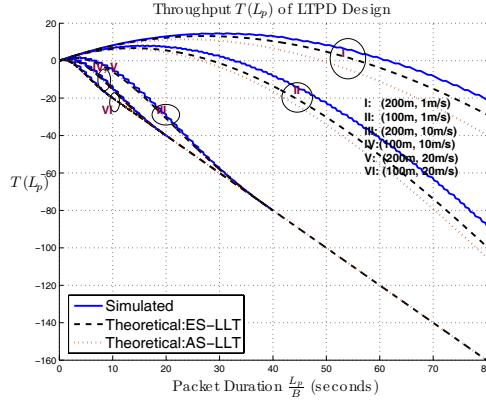


Fig. 5. LTPD design.

Recall that for inter-cell links, the size L_0 of a data packet should be chosen as $L_0 = \Theta(\frac{R \cdot B}{E(v)})$. With reference to Theorem 1, on average, every time duration of $E(I) = \Theta(\frac{L^4}{E(v) \cdot R^3})$ could have one data packet transferred. Accordingly, link throughput T_0 for one such pair of nodes can be computed as

$$T_0 = \frac{L_0}{I} = \Theta(\frac{R^4 B}{L^4}). \quad (23)$$

Normally, R is chosen on the same order of L , i.e., $\frac{R}{L} = \Theta(1)$. The above equation will be reduced to $T_0 = \Theta(B) = c_0$, where c_0 is a constant. Furthermore, from the connectivity constraint, there is at least one such link available for each node.

Due to limited mobility and transmission range, each packet needs to travel via multiple relays from source to destination following the path close to the straight line linking source and destination. Let the straight line connecting source with destination in the snapshot of initial network deployment be denoted as S-D line. Clearly, a source transmits data to its destination by multiple relays along the adjacent cells lying on its S-D line.

Let K be the average number of source-destination (S-D) lines passing through every cell and each source generates traffic $\Lambda(n)$ bits/s. To ensure that all required traffic is carried and recall that on average there are $\Theta(na(n))$ nodes in every cell, we need that

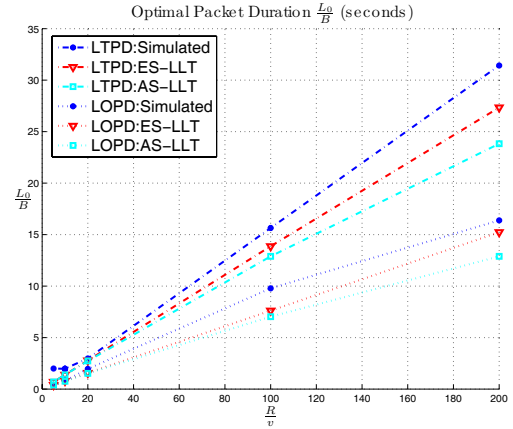
$$K \cdot \Lambda(n) \leq T_0 \cdot \Theta(na(n)) \Rightarrow \Lambda(n) = O(\frac{na(n)}{K}). \quad (24)$$

For every cell, the following lemma gives the number K of S-D lines passing through it.

Lemma 1: The number K of S-D lines passing through any cell is $\Theta(n\sqrt{a(n)})$, whp.

The proof of this lemma follows the proof of Lemma 3 in [18], because the S-D lines are determined from the initial network deployment, which is a snapshot of MANET and can also be treated as one configuration of a static wireless network.

The above analysis leads to the following conclusion on the throughput $\Lambda(n)$.

Fig. 6. Optimal packet duration $\frac{L_0}{B}$.

Theorem 2: For cell partitioned network with restricted mobility, we have $\Lambda(n) = O(\sqrt{a(n)})$ for generic mobility models. In particular, for a connected network whp, $\Lambda(n) = O(\sqrt{\frac{\log(n)}{n}})$.

B. Delay & Storage

Most packets need to travel across several cells before reaching their destinations and, therefore, must be stored in the queue of relay nodes. Consider an S-D queue at relay node m_r , a packet arrives when node m_r and the previous relay node (or the source node) simultaneously come into the communication circle; a packet departs when m_r meets another relay node (or the destination node) in the communication circle. Both the inter-arrival time and the inter-departure time are of the same order as link inter-arrival time (LIT). Since LIT can be characterized as exponentially distributed, each queue is characterized by a Poisson arrival process with exponential service time, thus being a M/M/1-FCFS queue.

For each S-D pair, queues at relay nodes construct a M/M/1-FCFS feedforward tandem network⁴. An important property of such a M/M/1-FCFS feedforward tandem network is the *Jackson's theorem* (see [19], page 150), i.e., if the tandem network with exponential service time is driven by a Poisson arrival process, every queue in the tandem network behaves as if it were an independent M/M/1-FCFS queue and thus can be analyzed individually. Recall the following properties for a M/M/1-FCFS queue (see [19], chapter 3) in the following lemma.

Lemma 2: Consider a discrete M/M/1-FCFS queue. Let $1 - \epsilon$ be the traffic intensity and λ be the exponential service rate of the queue, the average delay is given by

$$E(D) = \frac{1}{\lambda\epsilon} = \Theta(\frac{1}{\lambda}) \quad (25)$$

Furthermore, the mean and variance of the occupancy of the

⁴For delay to be finite, the arrival rate must be strictly less than the service rate but in this case, symmetric movements lead to a fully loaded tandem queue. To avoid this, we assume that if the available throughput is $\Lambda(n)$, each source generates traffic at a rate $(1 - \epsilon)\Lambda(n)$, for some $\epsilon > 0$. Please note that similar assumptions of a stable queue can be found in [20].

queue N_q is

$$E(N_q) = \frac{1 - \epsilon}{\epsilon} = \Theta(1), \quad (26)$$

$$\text{Var}(N_q) = \frac{1 - \epsilon}{\epsilon^2} = \Theta(1). \quad (27)$$

Without loss of generality, we can assume that the overall size of network is of unit area to analyze the network. In this case, we will have $A_N = 1$ and $L = \sqrt{a(n)}$. The average distance between S-D pairs is given by $\Theta(1)$ and the average number of hops for each packet is $\Theta(1/\sqrt{a(n)})$. Recall that every relay node carries information for $\Theta(n\sqrt{a(n)})$ S-D pairs and the service rate of each queue from LIT is $\lambda = \Theta(\frac{E(v)}{\sqrt{a(n)}})^5$. *Jackson's theorem* indicates that the delay for each S-D pair is the summation of delays occurred at relay nodes.

We can summarize the network performance in terms of average delay and storage in the following theorem.

Theorem 3: The average packet delay in a cell-partitioned network with restricted mobility and RDMM mobility models is given by

$$D(n) = \underbrace{\Theta\left(\frac{1}{\sqrt{a(n)}}\right)}_{\# \text{ of hops}} \cdot \underbrace{\Theta\left(\frac{\sqrt{a(n)}}{E(v)}\right)}_{\text{delay at each hop}} = \Theta\left(\frac{1}{E(v)}\right), \quad (28)$$

and the average information bit delay $D_b(n)$ is

$$D_b(n) = \frac{D(n)}{\Theta\left(\frac{RB}{E(v)}\right)} = \Theta\left(\frac{1}{RB}\right). \quad (29)$$

Furthermore, the mean and variance of the packet occupancy (i.e., storage requirement) is given by

$$E(N_p) = \text{Var}(N_p) = \Theta(n\sqrt{a(n)}), \quad (30)$$

and the corresponding bit storage requirement N_b is

$$E(N_b) = \text{Var}(N_b) = \Theta(n\sqrt{a(n)}) \cdot \Theta\left(\frac{RB}{E(v)}\right). \quad (31)$$

Summarizing the analysis, several important observations can be drawn here.

- By optimally segmenting the information, throughput of the network scales as $\Lambda(n) = O(\sqrt{a(n)})$ and packet-wise storage scales as $\Theta(n\sqrt{a(n)})$. Choices made to improve throughput will come with the price of increase in storage.
- Mobility can help alleviate the packet delay but won't be helpful to the bit-wise delay. It might be counter intuitive at the first glance. However, a detailed examination reveals that faster mobility brings more opportunities for nodes to deliver information packets but at the cost of reduced time for each communication. When information packets are optimally chosen, the negative effect from reduced communication time balances off the benefit from faster mobility. Eventually, to reduce the bit-wise delay, the only way is to increase the bandwidth and data rate for transmission.

⁵It can be obtained by substituting $\frac{R}{L} = \Theta(1)$ and $L = \sqrt{a(n)}$ into Eq. (21).

VI. CONCLUSIONS

We have presented an analytical framework for the characterization of link lifetime in MANETs with restricted mobility. Given the existence of prior attempts to incorporate link lifetime in the modeling of routing and clustering schemes [21]–[23], we believe that this new framework will find widespread use by researchers interested in the analytical modeling and optimization of channel access and routing protocols in MANETs.

We also apply the computed statistics from our framework to address the optimization of segmentation schemes as a function of link dynamics in a MANET. The optimized solutions obtained from the proposed analytical framework show a substantial improvement on network throughput. We summarize all these results to provide the first comprehensive analysis on throughput, average delay, and storage requirements for MANETs with restricted mobility.

ACKNOWLEDGEMENT

The authors would like to thank Dr. Robin Groenevelt and Prof. Philippe Nain of INRIA Institute for providing the simulation environment and the editor Professor Randall Berry and anonymous reviewers for their constructive comments.

APPENDIX

Proof of Theorem 1

The proof proceeds by modeling the meeting of two nodes in the communication circle as a geometric variable with some probability p of success and then taking the limit to derive the exponential distribution. The probability p will depend on the speeds and the positions of the two nodes. The probability p is obtained through summarizing the three exclusive scenarios analyzed below.

We first consider the case where node B is inside the communication region within the time duration $[t, t + \Delta_t]$, while node A moves into the communication circle with some probability p_1 . Because Δ_t is fairly small, we can assume that there is no change of directions within the duration Δ_t . The probability p_B that node B is located inside the communication circle at time t can be obtained from the stationary distribution,

$$p_B = \int \int_{S_B} \zeta(x, y) dx dy \quad (32)$$

where $\zeta(x, y)$ stands for the stationary spatial nodes' distribution and S_B (or S_A) denotes the semicircle of the communication circle in the cell B (or cell A). Meanwhile, we can also have similar definition of p_A . Because nodes are moving independently, the probability p_1 will be the product of p_B and p_{SA} . p_{SA} represents the probability of events that node A moves into the communication circle within time frame $[t, t + \Delta_t]$. It can be noted that we have neglected the probability of node B moving out of the communication circle within the time frame $[t, t + \Delta_t]$. In fact, the probability is on the same order of the third scenario and can be expressed as $o(\Delta_t)$.

Clearly, the probability p_{SA} varies with the initial location, speed v_A and direction ϕ_A of node A at time t . Without loss

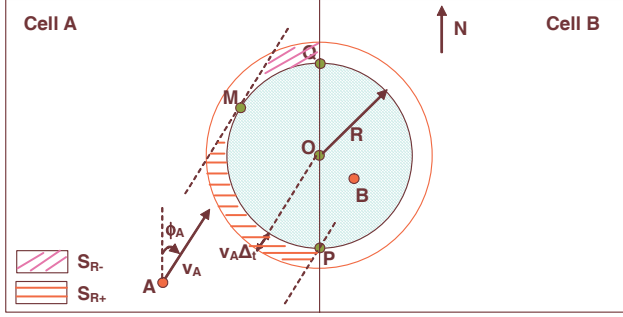


Fig. 7. Illustration of the first scenario.

of generality, we can assume $\phi_A \in [0, \pi]$ in our analysis. Conditioning on v_A and ϕ_A , within time duration $[t, t + \Delta_t)$, node A can at most travel towards the center point O for a distance of $v_A \Delta_t$. It implies that node A should be located inside the ring area in cell A in Fig. 7 for it to move into the communication circle within time duration $[t, t + \Delta_t)$.

To construct the ring area, we first draw two lines parallel to the direction ϕ_A . One line passes point P, while another line is a tangential line with respect to the circular communication circle at point M. For every point on *arcA*, we can draw a line passing through the point (termed as *cross* point) and in the meanwhile being parallel to the direction ϕ_A . One outmost point (called *verge* point) on the verge of the ring area can then be determined by looking for the point lying on the line with a distance of $v_A \Delta_t$ from the *cross* point. The *verge* point should be inside cell A while outside the communication region. To ensure that node A can move into the contact region S_A within time duration $[t, t + \Delta_t)$ with velocity v_A and direction ϕ_A , the location of node A at time t should be within the shaded area S_{R+} , i.e., the intersection area formed by the ring and the two parallel lines along direction ϕ_A in Fig. 7.

Let *arcPM* be the arc from point P to point M on the circumference. Conditioning on v_A and ϕ_A , the probability $p_{S_{R+}}$ for node A moving into the communication can now be computed as

$$p_{S_{R+}|\{v_A, \phi_A\}} = \int \int_{S_{R+}} \zeta(x, y) dx dy \approx v_A \cdot \Delta_t \cdot p_{arcPM} \quad (33)$$

where $p_{arcPM} = \int \int_{arcPM} \zeta(x, y) dx dy$. Consider the supplementary scenario where node A is of the same location and speed at time t but moving at direction $\phi_A - \pi$. Obviously, node A should now be within the supplementary area S_{R-} in Fig. 7. Let *arcQM* be the arc from point Q to M on the circumference. The complementary probability $p_{S_{R-}}$ can now be obtained as

$$p_{S_{R-}|\{v_A, \phi_A\}} = \int \int_{S_{R-}} \zeta(x, y) dx dy \approx v_A \cdot \Delta_t \cdot p_{arcQM} \quad (34)$$

where $p_{arcQM} = \int \int_{arcQM} \zeta(x, y) dx dy$.

Noting that *arcA* = *arcPM* + *arcQM*, where *arcA* (or *arcB*) is the circumference of the communication circle inside cell A (or cell B). We will have $p_{arcA} = p_{arcPM} + p_{arcQM}$, and averaging over all possible v_A and ϕ_A 's, the probability

p_{S_A} is given by

$$\begin{aligned} p_{S_A} &= E_{v_A} \left\{ \frac{1}{2\pi} \int_0^\pi (p_{S_{R+}|\{v_A, \phi_A\}} + p_{S_{R-}|\{v_A, \phi_A\}}) d\phi_A \right\} \\ &= E_{v_A} \left\{ \frac{1}{2\pi} v_A \cdot \Delta_t \cdot \int_0^\pi (p_{arcPM} + p_{arcQM}) d\phi_A \right\} \\ &= E_{v_A} \{ v_A \cdot \Delta_t \cdot p_{arcA} \cdot \frac{1}{2\pi} \int_0^\pi 1 d\phi_A \} \\ &= \frac{E(v_A)}{2} \cdot \Delta_t \cdot p_{arcA} \end{aligned} \quad (35)$$

The above leads to

$$p_1 = p_{S_A} \cdot p_B = \frac{E(v_A)}{2} \cdot \Delta_t \cdot p_{arcA} \cdot p_B \quad (36)$$

The next scenario for our proof consists of symmetric scenario where node A stays inside the communication circle within the time duration $[t, t + \Delta_t)$, while node B is going to move into the communication circle by some probability p_2 . Following similar derivation and analysis, p_2 can be calculated as

$$p_2 = \frac{E(v_B)}{2} \cdot \Delta_t \cdot p_{arcB} \cdot p_A \quad (37)$$

The last scenario we need to consider for our proof is the case where both node A and node B are located outside the communication region at time t but are going to move into the communication circle within time duration $[t, t + \Delta_t)$. In contrast to the two prior scenarios in which one node is within the communication circle while another one is located within the ring area at time t , in this case both nodes should be located within their respective ring area at time t .

It should be noted that the analytical procedure through geometric-variable analysis in the above scenarios can also be applied to analyze this scenario with minor modifications expected. For the purpose of succinctness, we will not elaborate on the derivations and our analysis shows that the probability p_3 for this case can be summarized as

$$\begin{aligned} p_3 &= E\{v_A \cdot v_B \cdot \Delta_t^2 \cdot p_{arcA} \cdot p_{arcB} \\ &\quad \cdot (\frac{1}{2\pi})^2 \int_0^\pi \int_0^\pi 1 d\phi_A d\phi_B\} \\ &= \frac{E(v_A)E(v_B)}{4} \cdot \Delta_t^2 \cdot p_{arcA} \cdot p_{arcB} = o(\Delta_t) \end{aligned} \quad (38)$$

Summarizing all three scenarios, we obtain that the probability p is given by

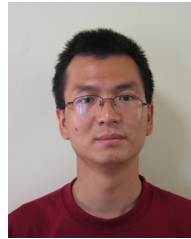
$$\begin{aligned} p &= p_1 + p_2 + p_3 \\ &= \frac{1}{2} \cdot \Delta_t \cdot (E(v_A) \cdot p_{arcA} \cdot p_B \\ &\quad + E(v_B) \cdot p_{arcB} \cdot p_A) + o(\Delta_t) \end{aligned} \quad (39)$$

Taking the limit $\Delta_t \rightarrow 0$ gives an exponential distribution with parameter $\lambda_F \approx \frac{E(v_A) \cdot p_{arcA} \cdot p_B + E(v_B) \cdot p_{arcB} \cdot p_A}{2}$.

Till now, we have arrived at a proof of Theorem 1 on general mobility models. For RDMM model, it should be noted that the stationary spatial nodes' distribution is uniform, i.e., $\zeta(x, y) = 1/L^2$ [12], [20]. It in turn gives $p_{arcA} = p_{arcB} = \int \int_{arcA} \zeta(x, y) dx dy = \frac{\pi \cdot R}{L^2}$ and $p_A = p_B = \int \int_{S_A} \zeta(x, y) dx dy = \frac{\pi \cdot R^2}{2L^2}$. By substituting these equations into the above proof, Theorem 1 follows.

REFERENCES

- [1] F. Bai, N. Sadagopan, and A. Helmy, "Important: a framework to systematically analyze the impact of mobility on performance of routing protocols for adhoc networks," in *Proc. IEEE Infocom*, San Francisco, CA, 2003.
- [2] P. Samar and S. B. Wicker, "On the behavior of communication links of a node in a multi-hop mobile environment," in *ACM MobiHoc 04*, pp. 148-156, 2004.
- [3] P. Samar and S. B. Wicker, "Link dynamics and protocol design in a multi-hop mobile environment," *IEEE Trans. Mobile Computing*, vol. 5, no. 9, pp. 1156-1172, 2006.
- [4] National Security Agency, *Assurance Technical Framework*, ch. 9, 3.1 edition, 2002.
- [5] R. B. Groenevelt, E. Altman, and P. Nain, "Relaying in mobile ad hoc networks: the Brownian motion mobility model," *Springer J. Wireless Networks (WINET)*, vol. 12, pp. 561-571, Oct. 2006.
- [6] P. Gupta and P. R. Kumar, "The capacity of wireless networks," *IEEE Trans. Inform. Theory*, vol. 46, no. 2, pp. 388-404, Mar. 2000.
- [7] R. M. Moraes, H. R. Sadjadpour, and J. J. Garcia-Luna-Aceves, "Mobility-capacity-delay trade-off in wireless ad-hoc networks." Elsevier, July 2005.
- [8] M. Carvalho and J. J. Garcia-Luna-Aceves, "A scalable model for channel access protocols in multihop ad hoc networks," in *Proc. ACM Mobicom 2004*, Philadelphia, PA, Sept. 2004.
- [9] S. Jiang, D. He, and J. Rao, "A prediction-based link availability estimation for mobile ad hoc networks," in *Proc. IEEE Infocom*, pp. 1745-1752, Apr. 2001.
- [10] S. Jiang, D. He, and J. Rao, "A prediction-based link availability estimation for routing metrics in manets," *IEEE/ACM Trans. Networking*, vol. 13, no. 6, pp. 1302-1312, Dec. 2005.
- [11] A. B. McDonald and T. F. Znati, "A mobility-based framework for adaptive clustering in wireless ad hoc networks," *IEEE J. Select. Areas Commun.*, vol. 17, no. 8, pp. 1466-1487, Aug. 1999.
- [12] C. Bettstetter, "Mobility modeling in wireless network: categorization, smooth movement and border effects," *ACM Mobile Computing Commun. Review*, vol. 5, pp. 55-67, Jan. 2001.
- [13] R. Guerin, "Channel occupancy time distribution in a cellular radio system," *IEEE Trans. Veh. Technol.*, vol. 35, no. 3, pp. 89-99, Aug. 1987.
- [14] D. Hong and S. S. Rappaport, "Traffic model and performance analysis for cellular mobile radio telephone systems with prioritized and non prioritized handoff procedures," *IEEE Trans. Veh. Technol.*, vol. 35, no. 3, pp. 77-92, Aug. 1986.
- [15] X. Wu, H. R. Sadjadpour, and J. J. Garcia-Luna-Aceves, "Analytical modeling of link and path dynamics and their implications on packet length in manets," in *Proc. SPECTS 2007*.
- [16] M. Grossglauser and D. Tse, "Mobility increases the capacity of ad-hoc wireless network," in *Proc. Twentieth Annual Joint Conf. IEEE Computer Commun. Societies*, vol. 3, pp. 1360-1369, Apr. 2001.
- [17] M. Grossglauser and D. Tse, "Mobility increases the capacity of adhoc wireless networks," *IEEE/ACM Trans. Networking*, vol. 10, no. 4, pp. 477-486, Aug. 2002.
- [18] A. E. Gammal, J. Mammen, B. Prabhaker, and D. Shah, "Throughput-delay trade-off in wireless network," in *Proc. IEEE Infocom*, vol. 1, pp. 464-475, 2004.
- [19] L. Kleinrock, *Queueing Systems, Volume 1: Theory*. John Wiley & Sons, Inc., 1975.
- [20] N. Bansal and Z. Liu, "Capacity, delay and mobility in wireless ad-hoc networks," in *Proc. IEEE Infocom*, 2003.
- [21] N. Sadagopan *et al.*, "Paths: analysis of path duration statistics and their impact on reactive manet routing protocols," in *Proc. ACM MobiHoc 03*, Annapolis, MD, June 2003.
- [22] A. Tsirigos and Z. J. Haas, "Analysis of multipath routing—part I: the effect on the packet delivery ratio," *IEEE Trans. Wireless Commun.*, vol. 3, no. 1, pp. 138-146, Jan. 2004.
- [23] D. Turgut *et al.*, "Longevity of routes in mobile ad hoc networks," in *Proc. VTC '01*, Greece, May 2001.



and coding theory.

Xianren Wu (S'04-M'08) received the B.S. degree in communication engineering from Nanjing University of Posts and Telecommunications, Nanjing, China, in 1998, the M.S. degree in information engineering from Beijing University of Posts and Telecommunications, Beijing, China, in 2001, and the Ph.D. degree in Electrical Engineering from University of California, Santa Cruz, in 2008. He received best paper award in SPECTS 2007 conference. His general research interest spans over mobile ad hoc networks, wireless communications



Hamid R. Sadjadpour (S'94-M'95-SM'00) received his B.S. and M.S. degrees from Sharif University of Technology with high honor and Ph.D. degree from University of Southern California in 1986, 1988 and 1996, respectively. After graduation, he joined AT&T as a member of technical staff, later senior technical staff member, and finally Principal member of technical staff at AT&T Lab. in Florham Park, NJ until 2001. In fall 2001, he joined University of California, Santa Cruz (UCSC) where he is now an Associate Professor.

Dr. Sadjadpour has served as technical program committee member in numerous conferences and as chair of communication theory symposium at WirelessCom 2005, and chair of communication and information theory symposium at IWCMC 2006, 2007 and 2008 conferences. He has been also Guest editor of EURASIP on special issue on Multicarrier Communications and Signal Processing in 2003 and special issue on Mobile Ad Hoc Networks in 2006, and is currently Associate editor for the JOURNAL OF COMMUNICATIONS AND NETWORKS (JCN). He has published more than 110 publications. His research interests include space-time signal processing, scaling laws for wireless ad hoc networks, performance analysis of ad hoc and sensor networks, and MAC layer protocols for MANETs. He is the co-recipient of International Symposium on Performance Evaluation of Computer and Telecommunication Systems (SPECTS) 2007 best paper award and the IEEE Fred W. Ellersick Award for Best Unclassified Paper at the 2008 Military Communications (MILCOM) conference. He holds more than 13 patents, one of them accepted in spectrum management of T1.E1.4 standard.



J.J. Garcia-Luna-Aceves (S'75-M'77-SM'02-F'06) received the B.S. degree in Electrical Engineering from the Universidad Iberoamericana, Mexico City, Mexico in 1977; and the M.S. and Ph.D. degrees in Electrical Engineering from the University of Hawaii at Manoa, Honolulu, HI in 1980 and 1983, respectively. He holds the Jack Baskin Endowed Chair of Computer Engineering at the University of California, Santa Cruz (UCSC), and is a Principal Scientist at the Palo Alto Research Center (PARC). Prior to joining UCSC in

1993, he was a Center Director at SRI International (SRI) in Menlo Park, California. He has been a Visiting Professor at Sun Laboratories and a Principal of Protocol Design at Nokia.

Dr. Garcia-Luna-Aceves holds 26 U.S. patents, and has published a book and more than 375 papers. He has directed 26 Ph.D. theses and 22 M.S. theses since he joined UCSC in 1993. He has been the General Chair of the ACM MobiCom 2008 Conference; the General Chair of the IEEE SECON 2005 Conference; Program Co-Chair of ACM MobiHoc 2002 and ACM MobiCom 2000; Chair of the ACM SIG Multimedia; General Chair of ACM Multimedia '93 and ACM SIGCOMM '88; and Program Chair of IEEE MULTIMEDIA '92, ACM SIGCOMM '87, and ACM SIGCOMM '86. He has served in the IEEE Internet Technology Award Committee, the IEEE Richard W. Hamming Medal Committee, and the National Research Council Panel on Digitization and Communications Science of the Army Research Laboratory Technical Assessment Board. He has been on the editorial boards of the IEEE/ACM TRANSACTIONS ON NETWORKING, the MULTIMEDIA SYSTEMS JOURNAL, and the JOURNAL OF HIGH SPEED NETWORKS.

He is an IEEE Fellow and an ACM Fellow, and is listed in Marquis Who's Who in America and Who's Who in The World. He is the co-recipient of the IEEE Fred W. Ellersick 2008 MILCOM Award for best unclassified paper. He is also co-recipient of Best Paper Awards at the IEEE MASS 2008, SPECTS 2007, IFIP Networking 2007, and IEEE MASS 2005 conferences, and of the Best Student Paper Award of the 1998 IEEE International Conference on Systems, Man, and Cybernetics. He received the SRI International Exceptional-Achievement Award in 1985 and 1989.