

Temporal Traffic Dynamics Improve the Connectivity of Ad Hoc Cognitive Radio Networks

Wei Ren, Qing Zhao, *Fellow, IEEE*, and Ananthram Swami

Abstract—In an ad hoc cognitive radio network, secondary users access channels temporarily unused by primary users, and the existence of a communication link between two secondary users depends on the transmitting and receiving activities of nearby primary users. Using theories and techniques from continuum percolation and ergodicity, we analytically characterize the connectivity of the secondary network defined in terms of the almost sure finiteness of the multihop delay, and show the occurrence of a phase transition phenomenon while studying the impact of the temporal dynamics of the primary traffic on the connectivity of the secondary network. Specifically, as long as the primary traffic has some temporal dynamics caused by either mobility and/or changes in traffic load and pattern, the connectivity of the secondary network depends solely on its own density and is independent of the primary traffic; otherwise, the connectivity of the secondary network requires putting a density-dependent cap on the primary traffic load. We show that the scaling behavior of the multihop delay depends critically on whether or not the secondary network is instantaneously connected. In particular, we establish the scaling law of the minimum multihop delay with respect to the source–destination distance when the propagation delay is negligible.

Index Terms—Ad hoc cognitive radio (CR) network, connectivity, continuum percolation, ergodicity, multihop delay, traffic dynamics.

I. INTRODUCTION

IN SPECTRUM overlay networks, primary and secondary users share a common spectrum in a hierarchical manner to achieve spectrum efficiency and interoperability [1]. By sensing and learning the communication environment via their cognitive radios [2], secondary users identify and exploit instantaneous and local spectrum opportunities while avoiding unacceptable interference to primary users [1].

We analytically characterize the connectivity and multihop delay of the secondary network. The existence of a communication link between two secondary users depends on not only their separation, but also the occurrence of the spectrum oppor-

tunity determined by the transmitting and receiving activities of nearby primary users. It is this interaction with the primary network that makes the problem fundamentally different from, and the analysis considerably more complex than, their counterparts in homogeneous networks. A qualitative and quantitative characterization of the impact of primary traffic on the secondary network is thus critical for understanding the performance limit of ad hoc cognitive radio (CR) networks and is the main topic of this paper.

A. Main Results

We consider a Poisson-distributed secondary network overlaid with a Poisson-distributed primary network in an infinite two-dimensional Euclidean space.¹ We define connectivity via the finiteness of the minimum multihop delay (MMD) between two randomly chosen secondary users, referred to as finite-delay connectivity (fd-connectivity). Specifically, the network is fd-disconnected if the MMD between two randomly chosen secondary users is infinite almost surely (a.s.), and is fd-connected if the MMD is finite with a positive probability (wpp.). Notice that the MMD considered here is not the multihop delay for a specific routing protocol. Instead, it is the minimum multihop delay that can be achieved by *any* routing protocol. The MMD thus specifies a fundamental performance limit and provides a benchmark for comparison.

We consider temporal dynamics in the primary traffic that could be caused by mobility and/or changes in the traffic load and pattern. We assume that the secondary network is static. Under the Poisson model, the two key parameters that characterize the topological structure of the secondary network and the primary traffic load are the density λ_S of the secondary users and the sequence $\{\lambda_{PT}(t) : t \geq 0\}$ of the densities of the primary transmitters. The fd-connectivity of the secondary network can thus be characterized by a partition of the infinite-dimensional space $(\lambda_S, \{\lambda_{PT}(t) : t \geq 0\})$.

Although the above partition appears to be intractable, we show that as long as the primary traffic has some temporal dynamics (no matter how small the range of the dynamics is), the fd-connectivity of the secondary network depends solely on its own density λ_S and is independent of the densities $\{\lambda_{PT}(t)\}$ of the primary transmitters, as illustrated in Fig. 1(a). In other words, no matter how heavy the primary traffic is, the secondary network is fd-connected, as long as its density λ_S exceeds the critical density λ_c of a homogeneous network (i.e., in the absence of the primary network). Note that when $\lambda_S > \lambda_c$, there

Manuscript received January 24, 2011; revised April 29, 2012; accepted January 18, 2013; approved by IEEE/ACM TRANSACTIONS ON NETWORKING Editor V. Misra. Date of publication February 20, 2013; date of current version February 12, 2014. This work was supported in part by the Army Research Office under Grant W911NF-08-1-0467 and the National Science Foundation under Grant CCF-0830685.

W. Ren and Q. Zhao are with the Department of Electrical and Computer Engineering, University of California, Davis, CA 95616 USA (e-mail: qzhao@ece.ucdavis.edu).

A. Swami is with the Army Research Laboratory, Adelphi, MD 20783 USA. Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/TNET.2013.2244612

¹This infinite network model is equivalent in distribution to the limit of a sequence of finite networks with a fixed density as the area of the network increases to infinity, i.e., the so-called *extended network*.

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE FEB 2014		2. REPORT TYPE		3. DATES COVERED 00-00-2014 to 00-00-2014	
4. TITLE AND SUBTITLE Temporal Traffic Dynamics Improve the Connectivity of Ad Hoc Cognitive Radio Networks				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) University of California at Davis, Department of Electrical and Computer Engineering, Davis, CA, 95616				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT In an ad hoc cognitive radio network, secondary users access channels temporarily unused by primary users, and the existence of a communication link between two secondary users depends on the transmitting and receiving activities of nearby primary users. Using theories and techniques from continuum percolation and ergodicity, we analytically characterize the connectivity of the secondary network defined in terms of the almost sure finiteness of the multihop delay, and show the occurrence of a phase transition phenomenon while studying the impact of the temporal dynamics of the primary traffic on the connectivity of the secondary network. Specifically, as long as the primary traffic has some temporal dynamics caused by either mobility and/or changes in traffic load and pattern, the connectivity of the secondary network depends solely on its own density and is independent of the primary traffic; otherwise, the connectivity of the secondary network requires putting a density-dependent cap on the primary traffic load. We show that the scaling behavior of the multihop delay depends critically on whether or not the secondary network is instantaneously connected. In particular we establish the scaling law of the minimum multihop delay with respect to the source???destination distance when the propagation delay is negligible.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 13	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

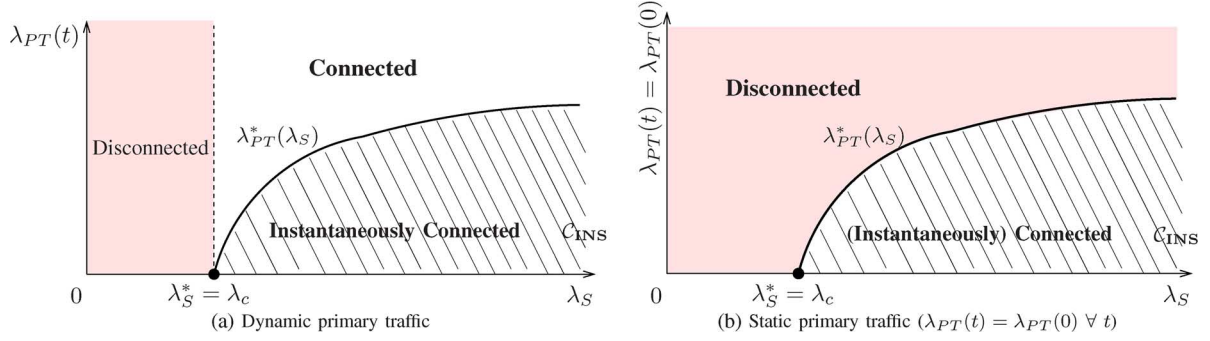


Fig. 1. FD-connectivity of ad hoc CR networks when the primary traffic has (a) temporal dynamics and (b) no temporal dynamics. The critical density λ_S^* of the secondary users is defined as the infimum density of the secondary users that ensures instantaneous connectivity under a *positive* density of the primary transmitters, and is equal to the critical density λ_c of a homogeneous network; the upper boundary $\lambda_{PT}^*(\lambda_S)$ is defined as the supremum density of the primary transmitters that ensures instantaneous connectivity with a *fixed* density λ_S of the secondary users.

is a.s. a unique infinite connected component (ICC) [3, Ch. 3] in the secondary network formed by topological links (a topological link exists between two users that are within communication range). We show that for any two secondary users in this ICC, the MMD is finite a.s. The intuition is that messages can traverse a topological path connecting the two secondary users by making stops in between to wait for spectrum opportunities, and more importantly, the waiting time is finite a.s. due to the temporal dynamics of the primary traffic. Since the percentage of the secondary users in this ICC is strictly positive, it follows that the MMD between two randomly chosen secondary users is finite wpp., i.e., the secondary network is fd-connected.

On the other hand, when the primary network is static, we show that the secondary network is fd-connected if and only if (iff.) it is *instantaneously* connected, as shown in Fig. 1(b). The secondary network is instantaneously connected if it has a unique ICC formed by communication links a.s. The existence of a communication link requires the existence of a topological link *and* the presence of a spectrum opportunity determined by the transmitting and receiving activities of nearby primary users. Due to this requirement, the instantaneous connectivity puts a cap on the tolerable primary traffic, which is an increasing function of the density λ_S of the secondary users [see Fig. 1(b)]. Moreover, given a static primary network, the set of communication links in the secondary network is fixed over time. It implies that if a topological link does not see an opportunity at the beginning, then it will never see it. Thus, messages from one secondary user can only reach another secondary user within the same connected component formed by communication links. If the secondary network is *instantaneously* connected, then wpp. two randomly chosen secondary users belong to this ICC formed by communication links,² and the MMD between them is finite; otherwise, they belong to two different finite connected components a.s., and they are inaccessible from each other, i.e., the MMD is infinite.

Although the primary traffic does not affect the fd-connectivity of the secondary network when it has temporal dynamics, it does affect the behavior of the MMD. Indeed, we show that the scaling behavior of the MDD with respect to the source–destination distance is starkly different depending on whether the

secondary network is instantaneously connected wpp. or not. Notice that the multihop delay in the secondary network consists of two components: the propagation delay and the waiting time at each hop for the occurrence of a spectrum opportunity. When the propagation delay is negligible, we show that if the secondary network is instantaneously connected wpp., the MMD is asymptotically independent of the source–destination distance; otherwise, the MMD scales at least linearly with the source–destination distance. We also study the case of nonnegligible propagation delay. Simulations show that the MMD-to-distance ratio for a secondary network that is instantaneously connected wpp. can be orders of magnitude smaller than that for a secondary network that is not instantaneously connected a.s.

These analytical results also provide important insights and design guidelines for practical systems. Since almost all primary networks have temporally dynamic traffic, it follows from the result on fd-connectivity that the accessibility between two secondary users is independent of the presence of the primary network, although it may incur a larger multihop delay. From the result on multihop delay, we can see that if the primary network has heavy traffic, then the secondary network can only be used for delay-tolerant applications; conversely, if a secondary network is deployed for delay-sensitive applications, then it should be operated within the instantaneous connectivity region C_{INS} for a positive portion of time,³ which imposes restrictions on the traffic load of the primary network or on the density of the secondary network.

B. Related Work

There have been only a few results on the connectivity of ad hoc CR networks. The Laplacian matrix is used to approximately characterize the graph connectivity in [5]. However, this does not characterize the multihop delay, and it does not take into account the impact of the receiving activities of the primary network on the secondary network.

Different types of connectivity of homogeneous networks (i.e., secondary network only) have been well studied in [6]–[14] and references therein. The theory of continuum

²It is shown in [4] that there exists either zero or one ICC formed by communication links in the secondary network a.s.

³Since $\{\lambda_{PT}(t)\}$ is ergodic, it follows from Fact A1 (in Appendix A) that the instantaneous connectivity of the secondary network wpp. is equivalent to that of the secondary network for a positive portion of time.

percolation has been used by Dousse *et al.* in analyzing the connectivity under the worst-case mutual interference [10], [11]. A dynamic connectivity graph for ALOHA networks is introduced in [14] to establish the scaling law of the delay with respect to the source–destination distance.

In [12] and [13], Kong and Yeh studied the connectivity and multihop delay of a homogeneous network with dynamic ON–OFF links based on continuum percolation theory. They showed that the scaling of the MMD behaves distinctly in two regimes, depending on whether the network is percolated. A major difference among [12], [13], and this paper is that in [12] and [13], the availability of a communication link is assumed to be independent of that of other communication links. This assumption does not apply to ad hoc cognitive radio networks considered in this paper. In a cognitive radio network, the communication links in the secondary network are spatially correlated since multiple secondary users may be affected by the same set of primary users. This leads to a dependent percolation model in the analysis, which is, in general, more difficult to deal with than the independent percolation model used in [12] and [13]. The proof of T2.1 in this paper is inspired by the proof of Lemma 9 in [12] and [13], but it is significantly simplified with the help of ergodicity theory.

In [15], we have studied the connectivity and multihop delay of ad hoc CR networks under the assumption that the realizations of the primary network are i.i.d. across slots, but we had not obtained the necessary and sufficient condition for the independence of the fd-connectivity of the secondary network from the primary traffic. In this paper, not only is this i.i.d. assumption replaced by a more realistic assumption under which the realizations of the primary network can be temporally correlated, but also the necessary and sufficient condition is provided. As detailed in Sections III–D and IV, the relaxation of this i.i.d. assumption significantly complicates the analysis, especially the one for fd-connectivity.

II. NETWORK MODEL

We consider a Poisson-distributed secondary network overlaid with a Poisson-distributed primary network in an infinite two-dimensional Euclidean space. The primary network adopts a synchronized slotted structure with slot length T_S . Thus, T_S can be considered to be the time constant of the spectrum opportunities that are determined by the transmitting and receiving activities of the primary users. Without loss of generality, we set $T_S = 1$.

Let $\Pi_{PT}(t)$ ($t \geq 0$) denote the point process of the primary transmitters at t . At $t = 0$, primary transmitters are distributed according to a two-dimensional Poisson point process $\Pi_{PT}(0)$ with density $\lambda_{PT}(0)$. The density $\lambda_{PT}(t)$ of the primary transmitters may vary due to various reasons. One can consider this as a birth–death process; nodes die when they have no more packets to send, and are (re)born when they do. We could also consider this from a duty-cycling perspective: Nodes sleep and wake up. Specifically, at each $t \geq 0$, each primary transmitter has a probability $q(t)$ of not transmitting at $t + 1$ due to lack of packets or perceived channel conditions; on the other hand, at $t + 1$, some primary nodes, which were silent in slot t , may

start transmitting. They are distributed, independent of $\Pi_{PT}(t)$, according to a two-dimensional Poisson point process with density $p(t)\lambda_{PT}(t)$, where the multiplicative factor $p(t) \geq 0$ is introduced for convenience and may exceed unity. Notice that since both $q(t)$ and $p(t)$ can only take nonnegative values, $\mathbb{E}[q(t)] > 0$ implies that $\mathbb{E}[p(t)] > 0$, and vice versa. Otherwise, each sample path of $\lambda_{PT}(t)$ will be either monotonically increasing or monotonically decreasing if either $\mathbb{E}(q(t)) = 0$, $\mathbb{E}(p(t)) > 0$ or $\mathbb{E}(q(t)) > 0$, $\mathbb{E}(p(t)) = 0$, which contradicts the stationarity and ergodicity of $\lambda_{PT}(t)$.

Not only the density of the primary transmitters, but also the positions of the remaining primary transmitters may change. The primary transmitters that continue transmitting will move to a new position at $t + 1$ according to a random displacement vector $\vec{m}(t)$ with a finite variance in each direction. The movements are modeled as i.i.d. for different primary transmitters, and for each primary transmitter, $\vec{m}(t)$ is i.i.d. across slots.⁴ Based on the thinning theorem and a stationarity theorem for dynamic Boolean models [17, Proposition 1.3], the primary transmitters that continue transmitting will form a two-dimensional Poisson point process with density $\lambda_{PT}(t)[1 - q(t)]$ at $t + 1$. Thus, by induction, the point process of transmitters $\Pi_{PT}(t + 1)$ containing the old primary transmitters at t and the new primary transmitters at $t + 1$ is Poisson with density $\lambda_{PT}(t + 1) = \lambda_{PT}(t)[1 - q(t) + p(t)]$. The random process $\{\lambda_{PT}(t)\}$ is assumed to be stationary and ergodic. The two related random processes $\{q(t) \in [0, 1]\}$ and $\{p(t) \geq 0\}$ are assumed to be stationary and ergodic; they may be correlated with $\{\lambda_{PT}(t)\}$.

The primary receivers are randomly (may not be uniformly) located within the transmission range⁵ R_p of their corresponding transmitters at each t , and their relative positions with respect to their corresponding transmitters can be either fixed or a stationary and ergodic random process over time. Based on the displacement theorem [18, Ch. 5], it can be shown that at each t , the primary receivers form another two-dimensional Poisson point process $\Pi_{PR}(t)$ with density $\lambda_{PT}(t)$, which is correlated with $\Pi_{PT}(t)$.

Secondary users are distributed according to a two-dimensional Poisson point process Π_S with density λ_S , which is independent of $\{\Pi_{PT}(t)\}$ and $\{\Pi_{PR}(t)\}$. The locations of the secondary users are static over time, and they have a uniform transmission range r_p .

III. CONNECTIVITY

In this section, we analytically characterize the connectivity of the secondary network. In particular, we show the occurrence of a phase transition phenomenon in terms of the impact of the temporal dynamics of the primary traffic on the fd-connectivity of the secondary network.

⁴This assumption of “i.i.d. across slots” is made so that we can use the classical central limit theorem (see Appendix B). Since the central limit theorem can be extended to the two cases of identical but weakly dependent distributions and martingales [16, Sec. 7.7], this assumption can be relaxed.

⁵Here, we assume that all the primary transmitters use the same power and the transmitted signals undergo isotropic path loss.

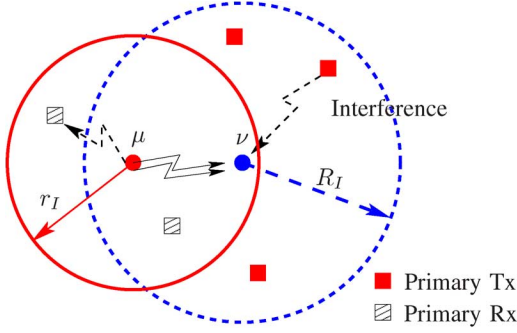


Fig. 2. Definition of spectrum opportunity.

A. Topological Link Versus Communication Link

A *topological link* exists between any two secondary users that are within each other's transmission range. Thus, topological links in the secondary network are independent of the primary network. In contrast, as discussed in the next paragraph, the existence of a *communication link* between two secondary users depends not only on the distance between them, but also on the availability of the communication channel, i.e., the presence of a spectrum opportunity. As a result, even in a static secondary network, communication links are time-varying due to the temporal dynamics of spectrum opportunities.

We consider the disk signal propagation and interference model as illustrated in Fig. 2. There exists an opportunity from μ , the secondary transmitter, to ν , the secondary receiver, if the transmission from μ does not interfere with *primary receivers* in the solid circle, and the reception at ν is not affected by *primary transmitters* in the dashed circle [19]. Referred to as the interference range of secondary users, the radius r_I of the solid circle centered at μ depends on the transmission power of μ and the interference tolerance of primary receivers, whereas the radius R_I of the dashed circle (the interference range of primary users) depends on the transmission power of primary users and the interference tolerance of the secondary user ν .

It follows from the above discussion that spectrum opportunities are *asymmetric*. Specifically, a channel that is an opportunity when μ is the transmitter and ν the receiver may not be an opportunity when ν is the transmitter and μ the receiver. Since unidirectional links are difficult to utilize, especially for applications with guaranteed delivery that require acknowledgments, we only consider bidirectional links in the secondary network when we define connectivity.

B. Instantaneous Connectivity Versus Topological Connectivity

In each slot t , we can obtain an undirected random graph $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$ consisting of all the secondary users and their communication links, which represents the instantaneous connectivity of the secondary network in this slot. As illustrated in Fig. 3, this graph $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$ is determined by the three Poisson point processes in slot t : Π_S , $\Pi_{PT}(t)$, and $\Pi_{PR}(t)$, where $\Pi_{PT}(t)$ and $\Pi_{PR}(t)$ are correlated.

We define the instantaneous connectivity of the secondary network in slot t as the a.s. existence of a unique ICC in $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$. Given the transmission power and the

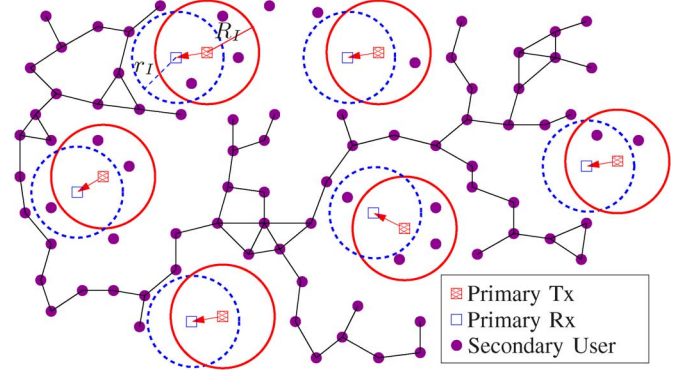


Fig. 3. Realization of the random graph $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$ that consists of all the secondary users and their communication links in slot t (denoted by solid lines). The solid circles denote the interference regions of the primary transmitters within which secondary users cannot successfully receive, and the dashed circles denote the required protection regions for the primary receivers within which secondary users should refrain from transmitting.

interference tolerance of both the primary and the secondary users (i.e., R_p , R_I , r_p , and r_I are fixed), the instantaneous connectivity region $\mathcal{C}_{INS}$ for slot t is defined as

$$\mathcal{C}_{INS} \triangleq \{(\lambda_S, \lambda_{PT}(t)) : \mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t) \text{ is connected}\}. \quad (1)$$

The upper boundary $\lambda_{PT}^*(\lambda_S)$ of $\mathcal{C}_{INS}$ is defined as

$$\lambda_{PT}^*(\lambda_S) \triangleq \sup \{\lambda_{PT}(t) : \mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t) \text{ is connected}\}. \quad (2)$$

The critical density of the secondary users, λ_S^* , is defined as

$$\lambda_S^* \triangleq \inf \{\lambda_S : \exists \lambda_{PT}(t) > 0 \text{ s.t. } \mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t) \text{ is connected}\}.$$

It is shown in [4] that λ_S^* equals the critical density λ_c of a *homogeneous* network. A detailed analytical characterization of $\mathcal{C}_{INS}$ is given in [4]. Let $\theta(\lambda_S, \lambda_{PT}(t))$ denote the probability that an arbitrary secondary user belongs to the ICC in $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$, if one exists, then we have that

$$\theta(\lambda_S, \lambda_{PT}(t)) \begin{cases} > 0, & \text{if } (\lambda_S, \lambda_{PT}(t)) \in \mathcal{C}_{INS} \\ = 0, & \text{otherwise.} \end{cases} \quad (3)$$

The fd-connectivity of the secondary network is defined by the finiteness of the MMD between two randomly chosen secondary users. To ensure finiteness of the multihop delay between two secondary users, it is necessary to have a path formed by topological links between them, otherwise they are not accessible from each other. Consider an undirected random graph $\mathcal{G}_S(\lambda_S)$ consisting of all the secondary users and their topological links. Notice that $\mathcal{G}_S(\lambda_S)$ depends only on the Poisson point process Π_S of the secondary network. Define the topological connectivity of the secondary network as the a.s. existence of a unique ICC in $\mathcal{G}_S(\lambda_S)$. It follows that fd-connectivity implies topological connectivity, i.e., topological connectivity is usually weaker than fd-connectivity. On the other hand, it is easy to show that instantaneous connectivity is usually stronger than fd-connectivity.

C. Connectivity under Static Primary Network

Consider a static primary network, i.e., the sets of the primary transmitters and receivers do not change over time, and their positions are also fixed. Then, as shown below, the necessary and sufficient condition for the connectivity of the secondary network is its instantaneous connectivity [see Fig. 1(b) for an illustration], i.e., the connectivity of $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$.

Proposition 1: Given a static primary network, i.e., $p(t) = q(t) = 0$ and $\vec{m}(t) = \vec{0} \forall t$, a necessary and sufficient condition for the connectivity of the secondary network is given by $(\lambda_S, \lambda_{PT}) \in \mathcal{C}_{INS}$ [defined by (1)], i.e., the MMD is finite wpp. iff. the network is instantaneously connected.

Proof: If $(\lambda_S, \lambda_{PT}) \in \mathcal{C}_{INS}$, then there exists a unique ICC formed by communication links a.s. in the secondary network. It implies that two randomly chosen secondary users belong to this ICC wpp. Since the set of communication links does not change, it follows that the MMD between them is finite wpp., i.e., the secondary network is connected.

If $(\lambda_S, \lambda_{PT}) \notin \mathcal{C}_{INS}$, then only finite connected components formed by communication links exist a.s. Thus, two randomly chosen secondary users belong to two different connected components a.s., and the MMD between them is infinite a.s. ■

D. Connectivity Under Dynamic Primary Network

Let $\vec{r}(t)$ denote the magnitude of the displacement vector $\vec{m}(t)$. Consider a dynamic primary network, where the dynamics can be caused by mobility ($\mathbb{E}[r^2(t)] > 0$) and/or changes in traffic load and pattern ($\mathbb{E}[q(t)] > 0$ implying that $\mathbb{E}[p(t)] > 0$). As illustrated in Fig. 1(a), we show in the following proposition that a necessary and sufficient condition for the connectivity of the secondary network is its topological connectivity, i.e., the connectivity of $\mathcal{G}_S(\lambda_S)$.

Proposition 2: Consider a dynamic primary network, i.e., $\mathbb{E}[q(t)] > 0$ or $\mathbb{E}[r^2(t)] > 0$. A necessary and sufficient condition for the connectivity of the secondary network is $\lambda_S > \lambda_c$, where λ_c is the critical density of homogeneous networks. ■

Remark: Since the multihop delay is a finite sum of single-hop delays, the a.s. finiteness of the MMD is implied by that of the single-hop delay. Under the primary network model where its realizations are correlated across slots, the single-hop delay is, however, difficult to analyze. In the proof, we use theories and techniques from ergodic theory to overcome this difficulty. Specifically, we establish the ergodicity of a measure-preserving (m.p.) dynamical system that consists of the probability space associated with the primary transmitters and an m.p. shift transformation in the time domain. A brief introduction to ergodic theory can be found in Appendix A.

Proof: If $\lambda_S \leq \lambda_c$, then there does not exist an infinite topologically connected component in $\mathcal{G}_S(\lambda_S)$ a.s. It follows that two randomly chosen secondary users belong to two different topologically connected components a.s., and the MMD between them is infinite a.s.

If $\lambda_S > \lambda_c$, then there exists a unique infinite topologically connected component C_T in $\mathcal{G}_S(\lambda_S)$ a.s. It follows that two randomly chosen secondary users μ and ν belong to C_T wpp. In other words, we can find a topological path L with finite hops from μ to ν wpp. Since the MMD $t(\mu, \nu)$ is bounded above by

the multihop delay $t^L(\mu, \nu)$ along the path L , it suffices to show the a.s. finiteness of $t^L(\mu, \nu)$, which is a direct consequence of the following lemma.

Lemma 1: Let $t_s(w_1, w_2)$ denote the single-hop delay from w_1 to w_2 , where w_1 and w_2 are connected via a topological link. If $\mathbb{E}[q(t)] > 0$ or $\mathbb{E}[r^2(t)] > 0$, then $t_s(w_1, w_2) < \infty$ a.s.

Proof of Lemma 1: In the following proof, by establishing the ergodicity, we transform the problem into finding an event that has nonzero instantaneous probability and that implies the availability of a given communication link in the secondary network. Since every primary receiver is within a certain range of its primary transmitter, a simple way to establish this is to select an event that requires a large region around the link to be free of primary transmitters.

Assume that the propagation delay $\tau \leq T_S = 1$ so that the spectrum opportunity lasts long enough to ensure the success of the transmission. Also assume that w_1 intends to transmit the message at $t = 0$. Thus, $t_s(w_1, w_2)$ is the waiting time $t_{sw}(w_1, w_2)$ for the presence of the first bidirectional opportunity plus the propagation delay τ , i.e.,

$$t_s(w_1, w_2) = t_{sw} + \tau = \arg \min_{t \in \{0, 1, 2, \dots\}} \{\mathbb{I}_E(t) = 1\} + \tau,$$

where $\mathbb{I}_E(t)$ is the indicator of the event that a bidirectional opportunity exists in the t th slot.

Next, we show the a.s. finiteness of t_{sw} . Let $\mathbb{I}(w, d, rx/tx)$ denote the event that there exist primary receivers/transmitters within distance d of the secondary user w , and $\overline{\mathbb{I}(w, d, rx/tx)}$ the complement of $\mathbb{I}(w, d, rx/tx)$. As illustrated in Fig. 4, the occurrence of the bidirectional opportunity E is given by

$$E \triangleq \overline{\mathbb{I}(w_1, r_1, rx)} \cap \overline{\mathbb{I}(w_1, R_1, tx)} \cap \overline{\mathbb{I}(w_2, r_1, rx)} \cap \overline{\mathbb{I}(w_2, R_1, tx)}.$$

Let O be the midpoint of the segment connecting w_1 and w_2 . Define the event F as

$$F \triangleq \overline{\mathbb{I}(O, R_M, tx)}$$

where $R_M = \max\{r_1 + R_p + (r_p/2), R_1 + (r_p/2)\}$. Let $\bar{t}_{sw}$ be the waiting time for the first occurrence of the event F , i.e.,

$$\bar{t}_{sw} = \arg \min_{t \in \{0, 1, 2, \dots\}} \{\mathbb{I}_F(t) = 1\}$$

where $\mathbb{I}_F(t)$ is the indicator of the event F during the t th slot. Since $F \subseteq E$, we have $t_{sw} \leq \bar{t}_{sw}$. Thus, we can show the a.s. finiteness of t_{sw} by proving the a.s. finiteness of $\bar{t}_{sw}$.

Consider the stationary random process $\{\Pi_{PT}(t) : t \geq 0\}$ where $\Pi_{PT}(t)$ is the Poisson point process formed by the primary transmitters in slot t . Based on a trivial generalization of the Kolmogorov extension theorem, we can construct a double-sided stationary random process $\{\Pi_{PT}(t) : t \in \mathbb{Z}\}$ that has the same finite dimensional distributions as $\{\Pi_{PT}(t) : t \geq 0\}$. Let $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT})$ be the probability space of $\{\Pi_{PT}(t) : t \in \mathbb{Z}\}$. Let $\pi_t(t \in \mathbb{Z})$ denote the realization of $\Pi_{PT}(t)$. $\forall \omega = \{\dots, \pi_{-1}, \pi_0, \pi_1, \dots\} \in \Omega_{PT}$, i.e., ω is any sample path of the double-sided stationary random process induced by the primary transmitters, define a shift transformation T as

$$(T\omega)_t = \pi_{t+1} \quad \forall t \in \mathbb{Z} \quad (4)$$

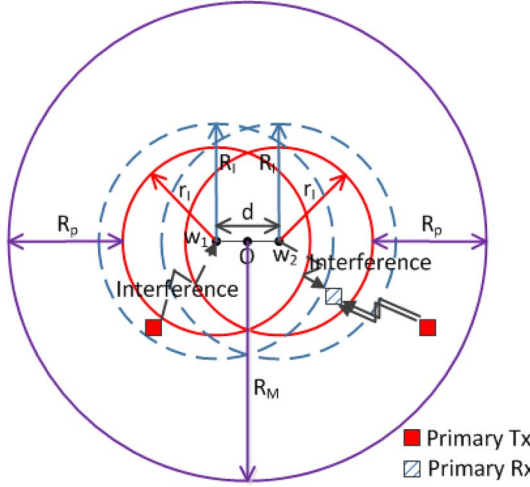


Fig. 4. Illustration of the bidirectional opportunity E and the constructed event F when $r_1 + R_p \geq R_1$. Here, d is the distance between the two secondary users w_1 and w_2 , which is bounded above by the transmission range r_p of the secondary users. The solid circles with radii r_1 (the interference range of the secondary users) denote the required protection regions for the primary receivers within which secondary users should refrain from transmitting, and the dashed circles with radii R_1 (the interference range of the primary users) denote the interference regions of the primary transmitters within which secondary users cannot successfully receive. The occurrence of F implies the occurrence of E because every primary receiver is within the transmission range R_p of its corresponding primary transmitter.

where $(T\omega)_t$ denotes the t th realization of $T\omega$. Since $\{\Pi_{PT}(t) : t \in \mathbb{Z}\}$ is time-stationary,⁶ it follows that $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT}, T)$ constitute an m.p. dynamical system. If $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT}, T)$ is ergodic, which will be shown in Lemma 2, it follows from Fact A1 that a.s.

$$\begin{aligned} \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=0}^{n-1} \mathbb{I}_F(k) &= \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=0}^{n-1} T^k \mathbb{I}_F(0) = \mathbb{E}[\mathbb{I}_F(0)] \\ &= \int_0^\infty \exp(-\lambda_{PT} \pi R_M^2) dF(\lambda_{PT}) > 0 \end{aligned}$$

where $F(\lambda_{PT})$ is the cumulative distribution function (CDF) of $\lambda_{PT}(0)$. Thus, $\bar{t}_{sw} < \infty$ a.s. ■

Now we only need to prove the following lemma to complete the proof of Lemma 1.

Lemma 2: Let $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT})$ be the probability space of $\{\Pi_{PT}(t) : t \in \mathbb{Z}\}$, and T the shift transformation defined by (4). If $\mathbb{E}[q(t)] > 0$ or $\mathbb{E}[r^2(t)] > 0$, then $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT}, T)$ is ergodic.

Sketch Proof of Lemma 2: We show that $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT}, T)$ is mixing, which implies its ergodicity [20, Proposition 2.5.1]. This is done by proving the asymptotic independence of one event from another transformed event. For details, please see Appendix B. ■

Remark: If $\mathbb{E}[q(t)] = \mathbb{E}[p(t)] = 0$ and $\mathbb{E}[r^2(t)] = 0$, i.e., the primary network is static, it can be easily shown that the m.p. dynamical system $(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT}, T)$ is not ergodic. Consider the following counterexample: Let $f \in L^1(\Omega_{PT}, \mathcal{F}_{PT}, P_{PT})$ be

the indicator function of the event that there does not exist any primary transmitter within the unit square B_1 centered at the origin at $t = 0$, then for any $\omega \in \Omega_{PT}$, the time average $\bar{f}$ of f is given by

$$\begin{aligned} \bar{f} &= \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=0}^{n-1} f(T^k \omega) \\ &= \begin{cases} 1, & \text{if no primary tx in } B_1 \text{ at } t = 0 \\ 0, & \text{otherwise.} \end{cases} \end{aligned}$$

However, the ensemble average $\mathbb{E}[f]$ of f is given by

$$\mathbb{E}[f] = \int_{\Omega_{PT}} f(\omega) dP_{PT} = \exp(-\lambda_{PT}).$$

Based on Fact A1, we have that the m.p. dynamical system in the static case is not ergodic. ■

Combining Propositions 1 and 2, we obtain a necessary and sufficient condition for the independence of the fd-connectivity of the secondary network with the primary network.

Theorem 1: Let λ_c be the critical density of homogeneous networks. Then, $\lambda_S > \lambda_c$ is a necessary and sufficient condition for the fd-connectivity of the secondary network iff. $\mathbb{E}[q(t)] > 0$ or $\mathbb{E}[r^2(t)] > 0$.

This theorem implies the occurrence of a phase transition phenomenon in the necessary and sufficient condition for the fd-connectivity. Specifically, if the primary network is static, the connectivity of the secondary network is equivalent to its instantaneous connectivity which depends on both its topology and the primary traffic; if the primary network is dynamic, the connectivity of the secondary network is equivalent to its topological connectivity, which depends solely on its topology and is independent of the primary traffic.

The key idea in the above proof is the ergodicity of the m.p. dynamical system associated with the primary transmitters. By taking one step further along the same line of the proof of Lemma 2, we can establish the ergodicity of the primary network consisting of both transmitters and receivers, leading to the ergodicity of the availability of a communication link in the secondary network. This ergodicity of the primary transmitters is thus essentially the ergodicity of the availability of a communication link in the secondary network. Thus, Theorem 1 is applicable to any ad hoc network where the availability of a communication link follows an arbitrary ergodic process.

IV. MULTIHOP DELAY

In this section, we analytically characterize the scaling behavior of the MMD with respect to the source-destination distance when the primary traffic is dynamic. Let $C(\mathcal{G}_S(\lambda_S))$ be the ICC in $\mathcal{G}_S(\lambda_S)$ when $\lambda_S > \lambda_c$ i.e., the secondary network is fd-connected. We seek to establish the scaling law of the MMD between two arbitrary users in $C(\mathcal{G}_S(\lambda_S))$ with respect to the distance between them. As shown below, the scaling behavior of the MMD is determined by whether the secondary network is instantaneously connected wpp. or not.

⁶The temporal stationarity of $\{\Pi_{PT}(t)\}$ is shown by [17, Proposition 1.3].

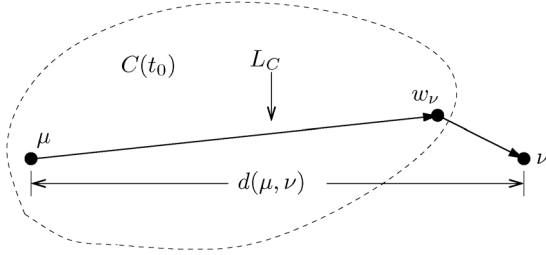


Fig. 5. Illustration of the constructed path L_C from μ to ν when $\Pr\{(\lambda_S, \lambda_{PT}(t)) \in C_{INS}\} > 0$. $C(t_0)$ is the ICC of $\mathcal{G}(\lambda_S, \lambda_{PT}(t_0))$ that first contains μ , and w_ν is the user in $C(t_0)$ that is closest to ν .

A. Negligible Propagation Delay

When the propagation delay $\tau = 0$, once a user has received the message, it can spread the message instantaneously throughout the connected component, formed by communication links, which contains it. Thus, if the secondary network is instantaneously connected during some time-slot, the source can route its message via the ICC such that the message can move a large number of hops toward the destination within this slot, leading to the multihop delay being asymptotically independent of the source–destination distance. On the other hand, if the secondary network is always not instantaneously connected, the message can move forward only a limited number of hops within each slot, which results in the linear scaling of the MMD. We state this formally next.

Theorem 2: Assume that $\tau = 0$, and $\mathbb{E}[q(t)] > 0$ or $\mathbb{E}[r^2(t)] > 0$. For any two secondary users $\mu, \nu \in C(\mathcal{G}_S(\lambda_S))$, the ICC of $\mathcal{G}_S(\lambda_S)$, let $t(\mu, \nu)$ denote the MMD from μ to ν and $d(\mu, \nu)$ the distance between μ and ν ; then, we have the following.

T2.1) If $\Pr\{\lambda_{PT}(t) < \lambda_{PT}^*(\lambda_S)\} > 0$ where $\lambda_{PT}^*(\lambda_S)$ is defined in (2)

$$\lim_{d(\mu, \nu) \rightarrow \infty} \frac{t(\mu, \nu)}{g(d(\mu, \nu))} = 0! \text{a.s.}$$

where $g(d)$ is any monotonically increasing function of d with $\lim_{d \rightarrow \infty} g(d) = \infty$.

T2.2) If $\Pr\{\lambda_{PT}(t) < \lambda'_{PT}\} = 0$ for some $\lambda'_{PT} > \lambda_{PT}^*(\lambda_S)$

$$\liminf_{d(\mu, \nu) \rightarrow \infty} \frac{\mathbb{E}[t(\mu, \nu)]}{d(\mu, \nu)} > 0.$$

$\Pr\{\lambda_{PT}(t) < \lambda_{PT}^*(\lambda_S)\} > 0$ implies $\Pr\{(\lambda_S, \lambda_{PT}(t)) \in C_{INS}\} > 0$, and $\Pr\{\lambda_{PT}(t) < \lambda'_{PT}\} = 0$ for some $\lambda'_{PT} > \lambda_{PT}^*(\lambda_S)$ implies $\Pr\{(\lambda_S, \lambda_{PT}(t)) \in C_{INS}\} = 0$, but not vice versa. We state the two conditions in the above way because we have not been able to establish whether the boundary point $(\lambda_S, \lambda_{PT}^*(\lambda_S)) \in C_{INS}$.

Proof Sketch: For T2.1, we use the ICC in $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t_0), t_0)$ during some slot t_0 to construct a path from μ to ν such that the multihop delay along this path is independent of the distance $d(\mu, \nu)$ (see Fig. 5 for an illustration). Let t_0 be the first slot such that μ belongs to the ICC $C(t_0)$ of $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t_0), t_0)$, and w_ν the user in $C(t_0)$ that is closest to ν . Since the propagation delay $\tau = 0$, the multihop delay from μ to w_ν is zero. It follows that

$$t(\mu, \nu) = t_0 + t(w_\nu, \nu).$$

Then, it suffices to show that t_0 and $t(w_\nu, \nu)$ are independent of $d(\mu, \nu)$, which we prove by using continuum percolation theory and ergodic theory. For details, please see Appendix C.

T2.2 is proven by using a coupling argument [3, Ch. 2] and then deriving a lower bound on $t(\mu, \nu)/d(\mu, \nu)$ by considering the fact that the message from μ can traverse only a finite distance toward ν during each slot. For details, please see Appendix D. ■

From the above, we see that the existence of the giant connected component can significantly reduce the multihop delay, especially when the destination is far away from the source.

B. Nonnegligible Propagation Delay

When the propagation delay $\tau > 0$, it takes at least time τ for the message to traverse a distance r_p , which imposes a lower bound τ/r_p on the ratio of the MMD to the source–destination distance. This implies that the MMD scales at least linearly with the source–destination distance.

The positive propagation delay τ also imposes an upper bound T_S/τ on the maximum number of hops that the message can traverse in a slot T_S . If the secondary network is instantaneously connected in this slot, this upper bound can be actually attained in the ICC consisting of communication links. Otherwise, this upper bound may not be attained due to the limited diameter of the finite connected components formed by communication links, especially when the propagation delay τ is small. In other words, there may not exist a connected component that has a path with T_S/τ hops. Thus, it can be expected that the MMD-to-distance ratio for a network that is instantaneously connected wpp. is much smaller than that for one that is not instantaneously connected a.s. (see Figs. 6(c), 6(d), 7(c), and 7(d) for an illustration).

V. SIMULATION RESULTS

We present two sets of simulation results. One set is to show the impact of connectivity on the scaling law of MMD with respect to the source–destination distance (see Figs. 6 and 7), and the other set is to show the impact of the temporal dynamics of the primary traffic on the MMD (see Fig. 8). The density λ_S of the simulated secondary network is larger than the critical density λ_c . Thus, the secondary network is either instantaneously connected or only connected but not instantaneously connected, depending on the density λ_{PT} of the primary transmitters. The area of the network is chosen to be large enough such that the asymptotic behavior can be observed. Without loss of generality, we assume that the source is located at the origin. Each node in the network is a potential destination. This allows us to simulate different realizations of the source–destination pair using one Monte Carlo run.

We consider two mobility models for the primary transmitters: the random walk model and the random waypoint model [21], where the former model has i.i.d. increments and the latter one does not. In the mobility model of the primary transmitters (see Section II), we assume that the movement increments are i.i.d. across slots. As we will see, the simulation results for the random waypoint model are very similar to those for the random walk model, which implies that the above assumption could be relaxed. For the random walk model, each primary

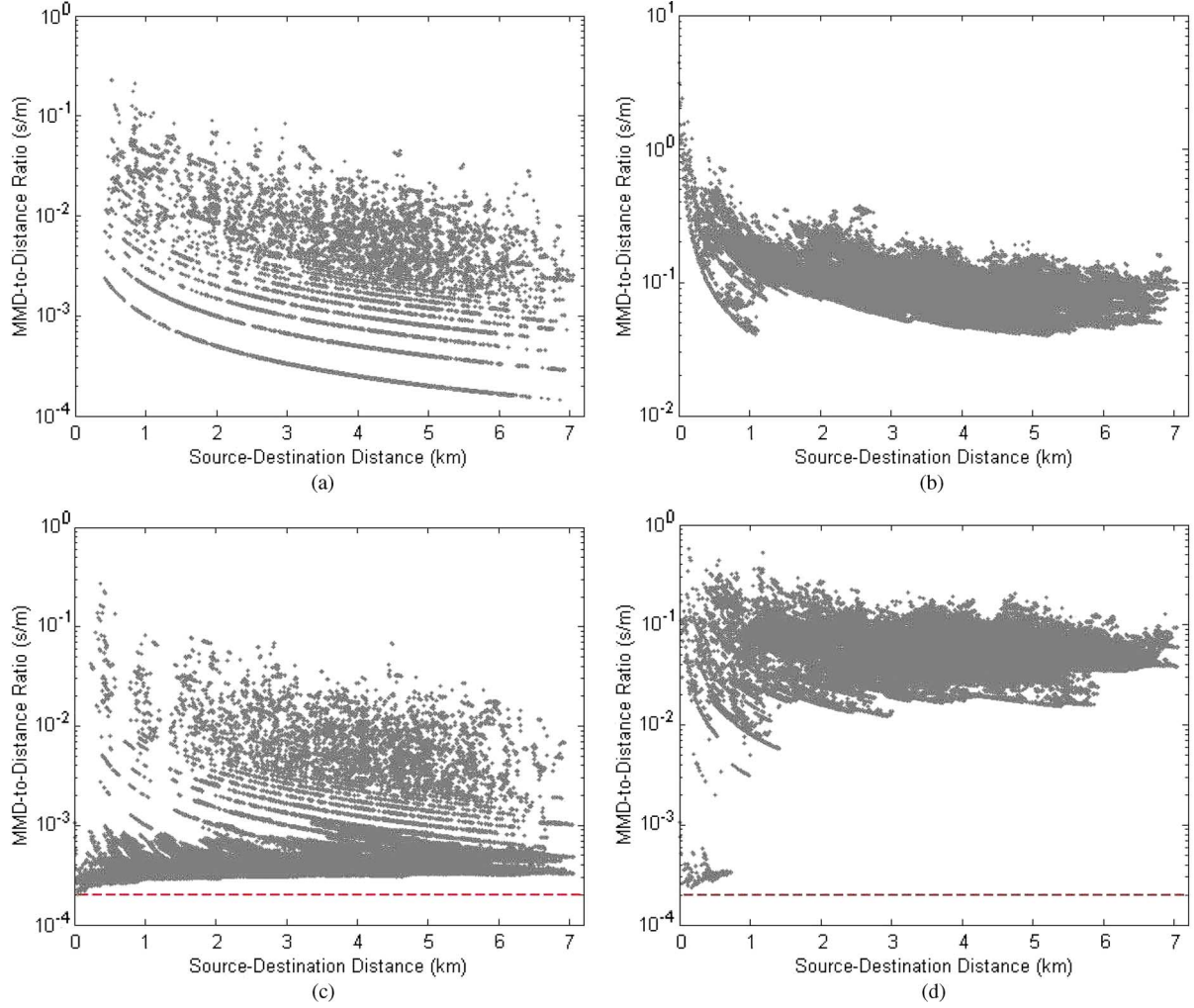


Fig. 6. MMD-to-distance ratio (in logarithmic scale) versus source-destination distance for random walk model. Notice that the MMD-to-distance ratio is obtained in one Monte Carlo run. The secondary users are distributed within a square $[-5 \text{ km}, 5 \text{ km}] \times [-5 \text{ km}, 5 \text{ km}]$ with density $\lambda_S = 700 \text{ km}^{-2}$. Given the transmission range $r_p = 50 \text{ m}$ of the secondary users, λ_S is larger than the critical density $\lambda_c(50) = 576 \text{ km}^{-2}$. Other simulation parameters are: $r_I = 80 \text{ m}$, $R_p = 50 \text{ m}$, $R_I = 80 \text{ m}$, $T_S = 1 \text{ s}$, $p_0^{rw} = 0.05$, $r_m = 5 \text{ m}$, $r_M = 30 \text{ m}$. (a) Instantaneously connected ($\lambda_{PT} = 5 \text{ km}^{-2}$, $\tau = 0$). (b) Not instantaneously connected ($\lambda_{PT} = 30 \text{ km}^{-2}$, $\tau = 0$). (c) Instantaneously connected ($\lambda_{PT} = 5 \text{ km}^{-2}$, $\tau = 0.01 \text{ s}$). (d) Not instantaneously connected ($\lambda_{PT} = 30 \text{ km}^{-2}$, $\tau = 0.01 \text{ s}$).

transmitter has a probability p_0^{rw} of staying at the current position in the next slot; otherwise, it will move to a new position according to a displacement vector whose magnitude is uniformly distributed within an interval $[r_m, r_M]$ and whose angle is uniformly distributed within $[0, 2\pi)$. When it reaches the simulation boundary, it bounces off the simulation border with an angle determined by the incoming direction.

For the random waypoint model, each primary transmitter chooses a random destination (not the destination for its transmission) uniformly distributed in the simulation area, which determines its displacement direction; then, it chooses a random speed uniformly distributed within an interval $[v_m, v_M]$ to move toward the destination. Upon reaching the destination, it may stay for a random number of slots $rgar$ is geometrically distributed with parameter $1 - p_0^{rw}$. The primary receivers are uniformly distributed within transmission range R_p of their corresponding transmitters in each slot.

Since it is difficult to identify the path with the MMD that depends on the topology of the secondary network and the

transmitting and receiving activities in the primary network in an intricate way, we obtain the MMD by considering a flooding scheme that tries every possible path from source to destination. During flooding, every user that has received the message (including the source) will transmit the message to its neighbors (i.e., within its transmission range) when it experiences a bidirectional spectrum opportunity with any of its neighbors. The transmission attempts will not stop until all its neighbors receive the message. The time that a user first receives the message during the flooding is the MMD from the source to this user. To highlight the impact of the waiting time for spectrum opportunities that is unique to CR networks, we do not consider the delay caused by scheduling, contention, or queuing. It can be shown that this flooding scheme gives us the MMD. We stress that flooding is used solely to determine the MMD and verify our scaling laws; flooding is not suggested as a routing protocol in the secondary network.

Figs. 6(a), 6(b), 7(a), and 7(b) show the MMD-to-distance ratio as a function of the source-destination distance when the

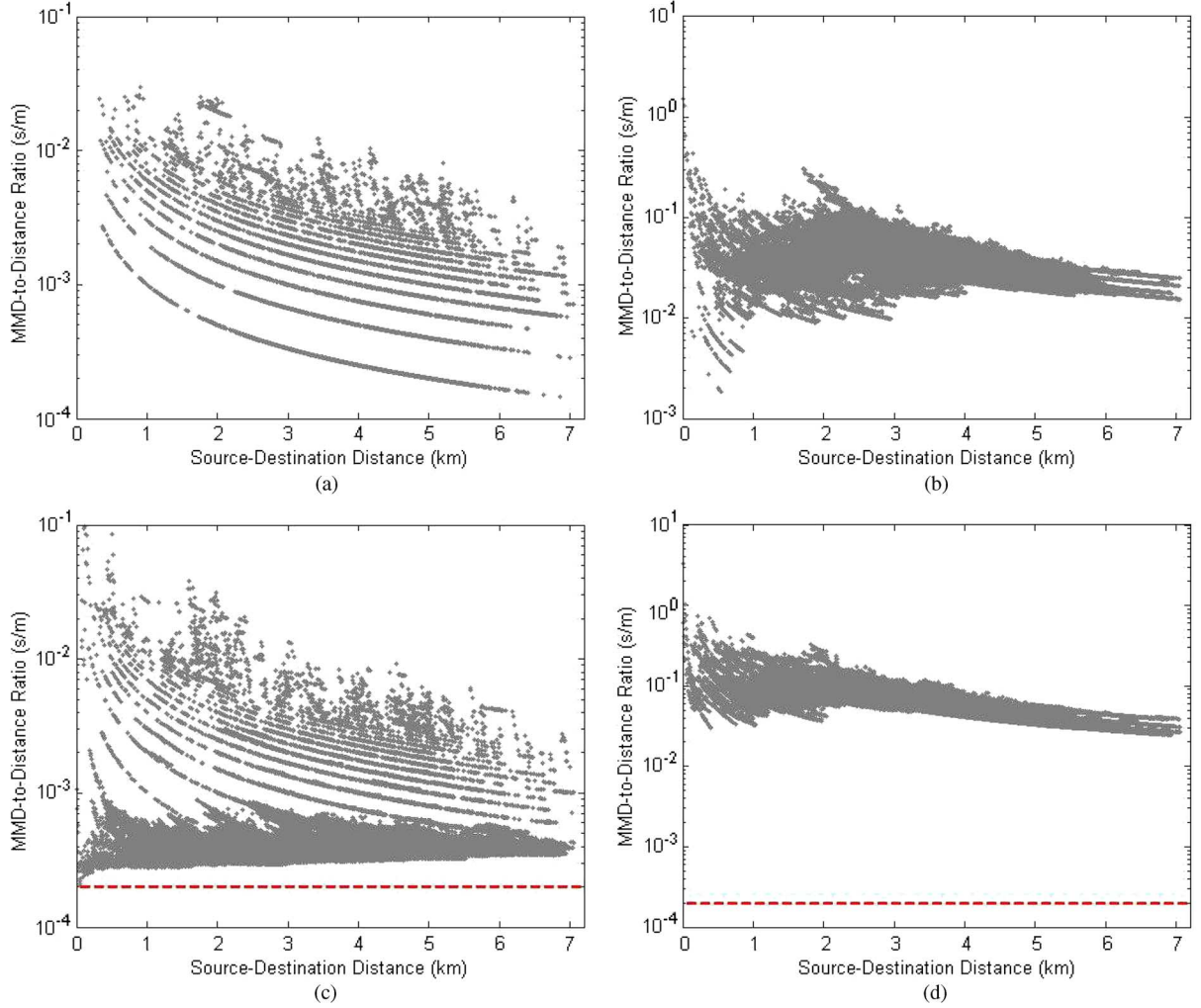


Fig. 7. MMD-to-distance ratio (in logarithmic scale) versus source-destination distance for random waypoint model. Notice that the MMD-to-distance ratio is obtained in one Monte Carlo run. The secondary users are distributed within a square $[-5 \text{ km}, 5 \text{ km}] \times [-5 \text{ km}, 5 \text{ km}]$ with density $\lambda_S = 700 \text{ km}^{-2}$. Given the transmission range $r_p = 50 \text{ m}$ of the secondary users, λ_S is larger than the critical density $\lambda_c(50) = 576 \text{ km}^{-2}$. Other simulation parameters are: $r_t = 80 \text{ m}$, $R_p = 50 \text{ m}$, $R_t = 80 \text{ m}$, $T_S = 1 \text{ s}$, $p_0^{\text{wsp}} = 0.05$, $v_m = 5 \text{ m/s}$, $V_M = 30 \text{ m/s}$. (a) Instantaneously connected ($\lambda_{PT} = 5 \text{ km}^{-2}$, $\tau = 0$). (b) Not instantaneously connected ($\lambda_{PT} = 30 \text{ km}^{-2}$, $\tau = 0$). (c) Instantaneously connected ($\lambda_{PT} = 5 \text{ km}^{-2}$, $\tau = 0.01 \text{ s}$). (d) Not instantaneously connected ($\lambda_{PT} = 30 \text{ km}^{-2}$, $\tau = 0.01 \text{ s}$).

propagation delay $\tau = 0$, where each dot represents a realization of the destination. We see that if the secondary network is instantaneously connected [Figs. 6(a) and 7(a)], the ratio decreases rapidly with distance and can be expected to go to zero. On the other hand, if the secondary network is not instantaneously connected [Figs. 6(b) and 7(b)], the ratio levels off as the distance increases and will approach a positive constant. Note that in Figs. 6(a) and 7(a), the MMD-to-distance ratios of different realizations of the destination are grouped into several continuous curves, each associated with a fixed MMD. Since the message is mainly delivered via the ICC consisting of communication links when the secondary network is instantaneously connected, the secondary users are actually grouped according to the first time that they are in an ICC. Figs. 6(a) and 7(a) tell us that due to the temporal dynamics of spectrum opportunities caused by the mobility of the primary network, every node will be part of an ICC within a few slots.

In Figs. 6(c), 6(d), 7(c), and 7(d), we compare the MMD-to-distance ratio in a network that is instantaneously connected and

in a network that is not when τ is nonzero but small. The four dashed lines in Figs. 6(c), 6(d), 7(c), and 7(d) denote the lower bound τ/r_p imposed by the propagation delay. Although the ratio for the network that is instantaneously connected does not go to zero due to the nonnegligible propagation delay, it is 100 times smaller than the ratio for the network that is not. Notice that a small group of dots is located at the bottom left corner of Fig. 6(d). This is because they are close to the source, and their corresponding secondary users happen to fall into the small connected component formed by communication links containing the source in the first few slots.

We also compare the expected MMD (denoted by $\mathbb{E}[\text{MMD}_1]$) under a mobile primary network that has fixed traffic load to the one (denoted by $\mathbb{E}[\text{MMD}_2]$) under a mobile primary network that has time-varying traffic load.⁷ The fractional difference $(\mathbb{E}[\text{MMD}_1] - \mathbb{E}[\text{MMD}_2])/\mathbb{E}[\text{MMD}_1]$ is expressed as a percentage in Fig. 8, where the secondary network is always

⁷For mobility, here we only consider the random walk model.

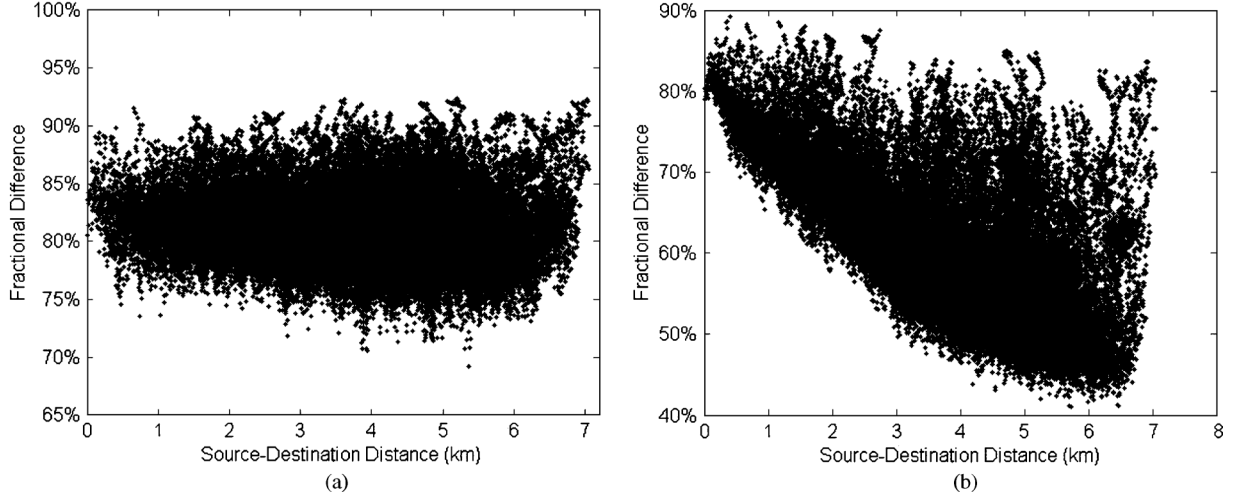


Fig. 8. Fractional difference $(\mathbb{E}[\text{MMD}_1] - \mathbb{E}[\text{MMD}_2])/\mathbb{E}[\text{MMD}_1]$ expressed as a percentage versus source–destination distance. The average value is taken over 1000 Monte Carlo runs. For $\mathbb{E}[\text{MMD}_1]$, the density $\lambda_{PT}(t)$ of the primary transmitters is fixed to be 5 km^{-2} . For $\mathbb{E}[\text{MMD}_2]$, $\lambda_{PT}(t)$ is uniformly distributed within $[0 \text{ km}^{-2}, 10 \text{ km}^{-2}]$ in each slot, and it is i.i.d. across slots. Other simulation parameters are the same as Fig. 6. (a) $\tau = 0$. (b) $\tau = 0.01 \text{ s}$.

instantaneously connected. The fact that the fractional difference is always positive implies that the introduction of another type of temporal dynamics reduces the expected MMD. Moreover, when the propagation delay $\tau = 0$ [see Fig. 8(a)], the fractional difference is more or less constant; when $\tau > 0$ [see Fig. 8(b)], it drops as the source–destination distance increases. Since the percentage of the secondary users in the ICC for case 2 is larger than the one for case 1 during many slots, the waiting time of each secondary user to become part of the ICC, which equals the MMD when $\tau = 0$, is uniformly decreased (irrespective of the distance from the source). However, when $\tau > 0$, the reduction of the expected MMD in case 2 is limited by the positive propagation delay τ .

VI. CONCLUSION AND FUTURE WORK

We say that the secondary or cognitive radio network is connected if the minimum multihop delay between an arbitrary source–destination pair is finite wpp. We have analytically characterized this connectivity. The impact of the primary traffic on the connectivity of the secondary network has been examined by establishing a necessary and sufficient condition for connectivity. Specifically, depending on whether the primary traffic has temporal dynamics or not, the connectivity of the secondary network is equivalent to its topological connectivity that is independent of the primary traffic, or its instantaneous connectivity that depends on the primary traffic. The temporal dynamics of the primary traffic can be caused by either mobility or changes in traffic load and pattern, and it is shown to significantly improve the connectivity of the secondary network in the sense that no matter how heavy the primary traffic is, the secondary network is connected as long as its density exceeds the critical density of homogeneous networks.

We have also studied the impact of connectivity on the multihop delay. When the propagation delay is negligible, depending on whether the secondary network is instantaneously connected with a positive probability or not, the scaling of the minimum multihop delay behaves differently in terms of the

scaling order. This scaling result is independent of the random positions of the source and the destination, and it only depends on the network parameters (e.g., the density of the secondary users and the traffic load of the primary network).

In the above analysis, we have assumed a disk signal propagation model that only incorporates path loss. If we take into account fading and shadowing, then a fixed transmission range does not hold, leading to a random connection model (RCM) [3, Ch. 1]. Since the RCM shares several basic properties (e.g., the ergodicity and the existence of the critical density) with the Boolean model used in this paper, we expect that the results established here can be extended to the RCM, although the derivations may become more complicated.

To highlight the impact of the waiting time for spectrum opportunities on the multihop delay of the secondary network, we have not considered delay caused by scheduling or contention among secondary users that shares similar flavor to that in conventional ad hoc networks. The results thus characterize the minimum delay and the fundamental performance limit. It is our hope that this paper will serve as a starting point to a more complete characterization of multihop delay that includes all these different factors.

APPENDIX A BASICS OF ERGODIC THEORY

The study object of ergodic theory is the so-called m.p. dynamical system (d.s.) $(\Omega, \mathcal{F}, \mu, T)$, which consists of a set Ω , a σ -algebra $\mathcal{F}$ of measurable subsets of Ω , a nonnegative measure μ on $(\Omega, \mathcal{F})$, and an invertible m.p. transformation $T : \Omega \rightarrow \Omega$ such that $\mu(T^{-1}F) = \mu(F) \forall F \in \mathcal{F}$. A set $F \in \mathcal{F}$ is said to be T -invariant if $T^{-1}F = F$.

An m.p.d.s. $(\Omega, \mathcal{F}, \mu, T)$ is said to be ergodic if for any invariant set, either itself or its complement has measure zero. We use the following fact frequently in the paper.

Fact A1 [20, Theorem 2.4.4]: An m.p.d.s. $(\Omega, \mathcal{F}, \mu, T)$ is ergodic, where $(\Omega, \mathcal{F}, \mu)$ is a probability space, iff.

$\forall f \in L^1(\Omega, \mathcal{F}, \mu)$ (i.e., f is a random variable with finite mean), and $\omega \in \Omega$ we have

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=0}^{n-1} f(T^k \omega) = \int_{\Omega} f d\mu \text{ a.s.}$$

If T is a shift transformation in the time domain, the above equation can be interpreted as the a.s. equality between the time average and the ensemble average.

An m.p.d.s. $(\Omega, \mathcal{F}, \mu, T)$ is said to be *mixing* if $\forall E, F \in \mathcal{F}$, $\mu(T^n E \cap F) - \mu(E)\mu(F) \rightarrow 0$ as $n \rightarrow \infty$. A mixing m.p.d.s. is ergodic [20, Proposition 2.5.1]. Typically, it is easier to establish ergodicity by showing that the m.p.d.s. is mixing.

APPENDIX B PROOF OF LEMMA 2

We consider two cases.

Case 1: $\mathbb{E}[q(t)] = \mathbb{E}[p(t)] = 0$, but $\mathbb{E}[r^2(t)] > 0$. Without loss of generality, we assume that $\mathbb{E}[m_x^2(t)] > 0$ where $m_x(t)$ is the x -component of the displacement $\vec{m}(t)$ in slot t . Consider two events F_1 and F_2 that depend only on the points of $\Pi_{PT}(0)$ within a box B_m centered at the origin with side length m . Let G_n denote the event that all the points of $\Pi_{PT}(0)$ within B_m are outside B_m in slot n , and H_K the event that there are at most K points of $\Pi_{PT}(0)$ within B_m . Let $X(n)$ be the x -component of the cumulative displacement associated with a point from $t = 0$ to $t = n$, i.e., $X(n) = \sum_{t=0}^n m_x(t)$. Then

$$\begin{aligned} \Pr\{G_n|H_K\} &\geq \sum_{k=0}^K (\Pr\{|X(n-1)| > m\})^k \Pr\{\#\text{points} \in B_m = k|H_K\} \\ &\geq (\Pr\{|X(n-1)| > m\})^K. \end{aligned}$$

Here, we have used the i.i.d. property of $X(t)$ across points. It follows from [22, Appendix B, Lemma B1] that

$$\lim_{n \rightarrow \infty} \Pr\{|X(n-1)| > m\} = 1.$$

Thus

$$\lim_{n \rightarrow \infty} \Pr\{G_n|H_K\} = 1. \quad (\text{B1})$$

If G_n occurs for some n , then obviously F_1 is independent of $T^n F_2$, i.e., $\Pr\{F_1 \cap T^n F_2|G_n \cap H_K\} = \Pr\{F_1|G_n \cap H_K\}\Pr\{T^n F_2|G_n \cap H_K\}$. Now since

$$\begin{aligned} \Pr\{F_1 \cap T^n F_2|H_K\} &= \Pr\{F_1 \cap T^n F_2|G_n \cap H_K\}\Pr\{G_n|H_K\} \\ &\quad + \Pr\{F_1 \cap T^n F_2|\overline{G_n} \cap H_K\}\Pr\{\overline{G_n}|H_K\} \end{aligned}$$

we have that

$$\begin{aligned} &\Pr\{F_1|G_n \cap H_K\}\Pr\{T^n F_2|G_n \cap H_K\}\Pr\{G_n|H_K\} \\ &\leq \Pr\{F_1 \cap T^n F_2|H_K\} \\ &\leq \Pr\{F_1|G_n \cap H_K\}\Pr\{T^n F_2|G_n \cap H_K\}\Pr\{G_n|H_K\} \\ &\quad + \Pr\{\overline{G_n}|H_K\}. \end{aligned}$$

Thus

$$\begin{aligned} &\lim_{n \rightarrow \infty} \Pr\{F_1 \cap T^n F_2|H_K\} \\ &= \lim_{n \rightarrow \infty} \Pr\{F_1|G_n \cap H_K\}\Pr\{T^n F_2|G_n \cap H_K\}\Pr\{G_n|H_K\} \\ &= \lim_{n \rightarrow \infty} \Pr\{F_1|G_n \cap H_K\} \lim_{n \rightarrow \infty} \Pr\{T^n F_2|G_n \cap H_K\}. \end{aligned}$$

Equation (B1) and the temporal stationarity of $\{\Pi_{PT}(t)\}$ yield

$$\begin{aligned} \lim_{n \rightarrow \infty} \Pr\{F_1|G_n \cap H_K\} &= \Pr\{F_1|H_K\} \\ \lim_{n \rightarrow \infty} \Pr\{T^n F_2|G_n \cap H_K\} &= \lim_{n \rightarrow \infty} \Pr\{T^n F_2|H_K\} \\ &= \Pr\{F_2|H_K\}. \end{aligned}$$

We thus have that

$$\lim_{n \rightarrow \infty} \Pr\{F_1 \cap T^n F_2|H_K\} = \Pr\{F_1|H_K\}\Pr\{F_2|H_K\}.$$

Since $\lim_{K \rightarrow \infty} \Pr\{H_K\} = 1$, as $K \rightarrow \infty$

$$\lim_{n \rightarrow \infty} \Pr\{F_1 \cap T^n F_2\} = \Pr\{F_1\}\Pr\{F_2\}.$$

Since any two arbitrary events F_1 and F_2 can be approximated by two sequences of events $\{F_1^m\}$ and $\{F_2^m\}$ that depend only on the realization of $\Pi_{PT}(0)$ inside B_m , the conclusion follows from [22, Appendix C, Lemma C1].

Case 2: $\mathbb{E}[q(t)] > 0$. Consider two events F_1 and F_2 that depend only on the points⁸ of $\{\Pi_{PT}(t) : -T \leq t \leq T\}$ within a box B_m centered at the origin with side length m . Let G_n denote the event that all the points that have visited B_m during $-T \leq t \leq T$ do not transmit in slot $n-T$, and H_K the event that there are at most K such points. Fixing a realization of $\{q(t)\}$, we have that

$$\begin{aligned} \Pr\{G_n|H_K\} &\geq \sum_{k=0}^K \left[1 - \prod_{i=T}^{n-T-1} (1 - q(i)) \right]^k \Pr\{\#\text{points} \in B_m = k|H_K\} \\ &\geq \left[1 - \prod_{i=T}^{n-T-1} (1 - q(i)) \right]^K. \end{aligned}$$

Since $\{q(t)\}$ is ergodic, based on Fact A1, we have that

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{t=0}^{n-1} q(t) = \mathbb{E}[q(t)] \text{ a.s.}$$

implying that $\lim_{n \rightarrow \infty} \sum_{t=0}^{n-1} q(t) = \infty$ a.s. It follows that

$$\lim_{n \rightarrow \infty} \prod_{i=T}^{n-T-1} (1 - q(i)) = 0 \text{ a.s.}$$

Thus

$$\lim_{n \rightarrow \infty} \Pr\{G_n|H_K\} = 1$$

and the rest of the proof follows along the same line of the one of Case 1.

⁸Since the set of the primary transmitters may change in every slot, it is not enough to only consider the points of $\Pi_{PT}(0)$.

APPENDIX C PROOF OF T2.1

We use the ICC consisting of communication links in $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t_0), t_0)$ during some slot t_0 to construct a path L_C from μ to ν such that the multihop delay along this path is independent of the distance $d(\mu, \nu)$ (see Fig. 5). Then, we analyze the multihop delay $t^C(\mu, \nu)$ along L_C .

Assume that μ starts trying to send the message at time $t = 0$. Let $C(t)$ be the ICC in $\mathcal{G}_H(\lambda_S, \lambda_{PT}(t), t)$ if it exists,⁹ and t_0 the first slot such that $\mu \in C(t_0)$. Let $(\Omega_P, \mathcal{F}_P, P_P)$ be the probability space of $\{\Pi_{PT}(t), \Pi_{PR}(t) : t \in \mathbb{Z}\}$, and define a shift transformation T similarly to (4). Given that $\mathbb{E}[q(t)] > 0$ or $\mathbb{E}[r^2(t)] > 0$, the ergodicity of the m.p. dynamical system $\{\Omega_P, \mathcal{F}_P, P_P, T\}$ follows along the same line of the proof of Lemma 2. Let F_t denote the event that $\mu \in C(t)$. Since $\Pr\{\lambda_{PT}(t) < \lambda_{PT}^*(\lambda_S)\} > 0$ implies that $\exists \lambda'_{PT} < \lambda_{PT}^*(\lambda_S)$ such that $\Pr\{\lambda_{PT}(t) \leq \lambda'_{PT}\} > 0$, we have that

$$\Pr\{F_0\} \geq \Pr\{\lambda_{PT}(t) \leq \lambda'_{PT}\} \theta(\lambda_S, \lambda'_{PT}) > 0$$

where $\theta(\lambda_S, \lambda'_{PT})$ is the probability that an arbitrary secondary users belongs to an ICC in $\mathcal{G}_H(\lambda_S, \lambda'_{PT}, t)$ given by (3). It follows from the arguments similar to those in showing that $\bar{t}_{sw} < \infty$ a.s. in the proof of Lemma 1, that $t_0 < \infty$ a.s.

Given $C(t_0)$, we define user w_ν as the user in $C(t_0)$ that is closest to ν , i.e.,

$$w_\nu \triangleq \arg \min_{w_i \in C(t_0)} d(w_i, \nu).$$

Notice that if $\nu \in C(t_0)$, then $w_\nu = \nu$.

As illustrated in Fig. 5, the constructed path L_C passes through w_ν , and the multihop delay $t^C(\mu, \nu)$ along the path L_C can be expressed as

$$t^C(\mu, \nu) = t_0 + t(\mu, w_\nu) + t(w_\nu, \nu) = t_0 + t(w_\nu, \nu),$$

where $t(w_\nu, \nu)$ is the MMD from w_ν to ν . In the last step, we have used $t(\mu, w_\nu) = 0$, which is due to the fact that $\mu, w_\nu \in C(t_0)$ and $\tau = 0$. Next, we prove the following lemma.

Lemma C1: $t(w_\nu, \nu)$ is finite a.s.

Proof Sketch: We first show that $d(w_\nu, \nu) < \infty$ a.s. by using the ergodicity of the network model, and then obtain an upper bound on the multihop delay $t^L(w_\nu, \nu)$ along the shortest path¹⁰ $L(w_\nu, \nu)$ from w_ν to ν . Since $t(w_\nu, \nu) \leq t^L(w_\nu, \nu)$, the a.s. finiteness of $t(w_\nu, \nu)$ follows from that of the upper bound on $t^L(w_\nu, \nu)$. The proof here is inspired by the proof of [13, Lemma 9], but with a much simpler proof of $d(w_\nu, \nu) < \infty$. For details, see [15, Appendix B]. ■

APPENDIX D PROOF OF T2.2

Let $t'(\mu, \nu)$ be the MMD from μ to ν when $\lambda_{PT}(t) = \lambda'_{PT} \forall t$. Since $\Pr\{\lambda_{PT}(t) < \lambda'_{PT}\} = 0$, i.e., $\Pr\{\lambda_{PT}(t) \geq \lambda'_{PT}\} = 1$,

⁹Since $\lambda_{PT}(t)$ is time-varying, it is possible that $C(t)$ does not exist in some slots.

¹⁰The shortest path is the path from the source to the destination with the minimum number of hops. Notice that the shortest path is not necessarily the minimum path since the probability of having an opportunity is a function of the hop length and a longer hop usually results in more waiting time.

every realization of the primary network with density $\lambda_{PT}(t)$ can be obtained from adding more primary transmitter–receiver pairs to a realization of the primary network with density λ'_{PT} [3, Ch. 2]. It follows that the expected single-hop delay under $\lambda_{PT}(t)$ is bounded below by the one under λ'_{PT} . Then, based on the above coupling argument, we have $\mathbb{E}[t(\mu, \nu)] \geq \mathbb{E}[t'(\mu, \nu)]$. It suffices to show that

$$\liminf_{d(\mu, \nu) \rightarrow \infty} \frac{\mathbb{E}[t'(\mu, \nu)]}{d(\mu, \nu)} > 0$$

which is shown in [15, Lemma 4].

REFERENCES

- [1] Q. Zhao and B. M. Sadler, “A survey of dynamic spectrum access,” *IEEE Signal Process. Mag.*, vol. 24, no. 3, pp. 79–89, May 2007.
- [2] I. J. Mitola and J. G. Maguire, “Cognitive radio: Making software radios more personal,” *IEEE Pers. Commun.*, vol. 6, no. 4, pp. 13–18, Aug. 1999.
- [3] R. Meester and R. Roy, *Continuum Percolation*. New York, NY, USA: Cambridge Univ. Press, 1996.
- [4] W. Ren, Q. Zhao, and A. Swami, “Connectivity of cognitive radio networks: Proximity vs. opportunity,” in *Proc. ACM MobiCom Workshop Cogn. Radio Netw.*, Sep. 2009, pp. 37–42.
- [5] A. Abbagnale, F. Cuomo, and E. Cipollone, “Measuring the connectivity of a cognitive radio ad-hoc network,” *IEEE Commun. Lett.*, vol. 14, no. 5, pp. 417–419, May 2010.
- [6] P. Gupta and P. R. Kumar, “Critical power for asymptotic connectivity in wireless networks,” in *Stochastic Analysis, Control, Optimization and Applications: A Volume in Honor*, W.H. Fleming, Ed., W. M. McEneaney, Ed., G. Yin and Q. Zhang, Eds. Boston, MA, USA: Birkhauser, 1998, pp. 547–566.
- [7] T. K. Phillips, S. S. Panwar, and A. N. Tantawi, “Connectivity properties of a packet radio network model,” *IEEE Trans. Inf. Theory*, vol. 35, no. 5, pp. 1044–1047, Sep. 1989.
- [8] J. Ni and S. A. G. Chandler, “Connectivity properties of a random network,” *Inst. Elect. Eng. Proc. Commun.*, vol. 141, no. 4, pp. 289–296, Aug. 1994.
- [9] C. Bettstetter, “On the minimum node degree and connectivity of a wireless multihop network,” in *Proc. ACM MobiHoc*, Jun. 2002, pp. 80–91.
- [10] O. Dousse, F. Baccelli, and P. Thiran, “Impact of interference on connectivity in ad hoc networks,” *IEEE/ACM Trans. Netw.*, vol. 13, no. 2, pp. 425–436, Apr. 2005.
- [11] O. Dousse, M. Franceschetti, N. Macris, R. Meester, and P. Thiran, “Percolation in the signal to interference ratio graph,” *J. Appl. Probab.*, vol. 43, no. 2, pp. 552–562, 2006.
- [12] Z. N. Kong and E. M. Yeh, “Connectivity and latency in large-scale wireless networks with unreliable links,” in *Proc. IEEE INFOCOM*, Apr. 2008, pp. 11–15.
- [13] Z. N. Kong and E. M. Yeh, “Connectivity, percolation, and information dissemination in large-scale wireless networks with dynamic links,” *IEEE Trans. Inf. Theory* Feb. 2009 [Online]. Available: <http://arxiv.org/abs/0902.4449>, submitted for publication
- [14] R. K. Ganti and M. Haenggi, “Dynamic connectivity in aloha ad hoc networks,” *IEEE Trans. Inf. Theory* Mar. 2010 [Online]. Available: <http://arxiv.org/abs/0808.4146>, submitted for publication
- [15] W. Ren, Q. Zhao, and A. Swami, “On the connectivity and multihop delay of ad hoc cognitive radio networks,” *IEEE J. Sel. Areas Commun.*, vol. 29, no. 4, pp. 805–818, Apr. 2011.
- [16] R. Durrett, *Probability: Theory and Examples*, 2nd ed. Pacific Grove, CA, USA: Duxbury Press, 1996.
- [17] J. V. D. Berg, R. Meester, and D. G. White, “Dynamic boolean models,” *Stochastic Process. Appl.*, vol. 69, no. 2, pp. 247–257, Sep. 1997.
- [18] J. F. C. Kingman, *Poisson Processes*. New York, NY, USA: Oxford Univ. Press, 1993.
- [19] Q. Zhao, “Spectrum opportunity and interference constraint in opportunistic spectrum access,” in *Proc. IEEE ICASSP*, Apr. 2007, vol. 3, pp. 605–608.
- [20] K. Petersen, *Ergodic Theory*. New York, NY, USA: Cambridge Univ. Press, 1989.

- [21] T. Camp, J. Boleng, and V. Davies, "A survey of mobility models for ad hoc network research," *Wireless Commun. Mobile Comput.*, vol. 2, no. 5, pp. 483–502, Sep. 2002.
- [22] W. Ren, Q. Zhao, and A. Swami, "Temporal traffic dynamics improve the connectivity of ad hoc cognitive radio networks," Tech. Rep. TR-10-02, Jul. 2010 [Online]. Available: <http://www.ece.ucdavis.edu/~qzhao/TR-10-02.pdf>



Wei Ren received the B.Sc. and M.Sc. degrees from Peking University, Beijing, China, in 2003 and 2006, respectively, and the Ph.D. degree from University of California, Davis, CA, USA, in 2011.

He is currently a Software Development Engineer with Microsoft Corporation. During his Ph.D. study, his research interests were in cognitive radio systems and wireless networks. He was also interested in algorithmic theory and optimization techniques for communications.



Qing Zhao (S'97–M'02–SM'08–F'13) received the Ph.D. degree in electrical engineering from Cornell University, Ithaca, NY, USA, in 2001.

In 2004, she joined the Department of Electrical and Computer Engineering, University of California (UC), Davis, CA, USA, where she is currently a Professor. Her research interests are in the general area of stochastic optimization, decision theory, and algorithmic theory in dynamic systems and communication and social networks.

Prof. Zhao received the 2010 *IEEE Signal Processing Magazine* Best Paper Award and the 2000 Young Author Best Paper Award from the IEEE Signal Processing Society. She holds the title of UC Davis Chancellor's Fellow and received the 2008 Outstanding Junior Faculty Award from the UC Davis College of Engineering. She was a plenary speaker at the 11th IEEE Workshop on Signal Processing Advances in Wireless Communications (SPAWC) 2010. She is also a coauthor of two papers that received student paper awards at ICASSP 2006 and the IEEE Asilomar Conference 2006.



Ananthram Swami received the B.Tech. degree from the Indian Institute of Technology (IIT), Bombay, India, the M.S. degree from Rice University, Houston, TX, USA, and the Ph.D. degree from the University of Southern California (USC), Los Angeles, CA, USA, all in electrical engineering.

He is with the U.S. Army Research Laboratory (ARL), Adelphi, MD, USA, as the Army's ST (Senior Research Scientist) for Network Science. He has held positions with Unocal Corporation, USC, CS-3, and Malgudi Systems. He was a Statistical Consultant to the California Lottery, developed a MATLAB-based toolbox for non-Gaussian signal processing, and has held visiting faculty positions with INP, Toulouse. His research interests are in the broad area of network science: the study of interactions and co-evolution, prediction and control of interdependent networks, with applications in composite tactical networks.

Dr. Swami is an ARL Fellow.