

Randomized Scheduling of Real-Time Traffic in Wireless Networks Over Fading Channels

Christos Tsanikidis^{ID}, *Graduate Student Member, IEEE*, and Javad Ghaderi^{ID}, *Senior Member, IEEE*

Abstract—Despite the rich literature on scheduling algorithms for wireless networks, algorithms that can provide deadline guarantees on packet delivery for general traffic and interference models are very limited. In this paper, we study the problem of scheduling real-time traffic under a *conflict-graph interference model* with *unreliable links* due to channel fading. Packets that are not successfully delivered within their deadlines are of no value. We consider traffic (packet arrival and deadline) and fading (link reliability) processes that evolve as an *unknown* finite-state Markov chain. The performance metric is efficiency ratio which is the fraction of packets of each link which are delivered within their deadlines compared to that under the optimal (unknown) policy. We first show a conversion result that shows classical non-real-time scheduling algorithms can be ported to the real-time setting and yield a constant efficiency ratio. In particular, Max-Weight Scheduling (MWS) yields an efficiency ratio of $1/2$. We then propose randomized algorithms that achieve efficiency ratios strictly higher than $1/2$, by carefully randomizing over the maximal schedules. Further, we propose low-complexity and myopic distributed randomized algorithms, and characterize their efficiency ratio. Simulation results are presented that verify that the randomized algorithms outperform classical ones such as MWS and GMS for scheduling real-time traffic over fading channels.

Index Terms—Scheduling, real-time traffic, Markov processes, stability, wireless networks.

I. INTRODUCTION

THERE has been vast research on scheduling algorithms in wireless networks which mostly focus on maximizing long-term throughput when packets have no strict delay constraints. *Max-Weight Scheduling* (MWS) policy is known to be throughput optimal in such settings, attaining any desired throughput vector in the feasible throughput region [2]. Further, greedy scheduling policies such as LQF [3], [4], or distributed policies such as CSMA [5], [6], [7] have been proposed, to alleviate the computational complexity of MWS and achieve a certain fraction of the throughput region. However, in many emerging applications, such as Internet of Things

(IoT), vehicular networks, and edge computing, delays and deadline guarantees on packet delivery also play an important role [8], [9], [10], as packets that are not received within specific deadlines are of little or no value, and are typically discarded by the application. This *discontinuity* in the packet value as a function of latency makes the problem significantly more challenging than traditional scheduling where packets do not have strict deadlines.

There is an increasing body of work attempting to address the above challenge, however they either assume a frame-based traffic model [11], [12], [13], [14], [15], relax the wireless interference graph constraints [16], or use greedy scheduling approaches like LDF [17], [18]. In the frame-based traffic model, time is divided into frames, and packet arrivals and their deadlines during a frame are assumed to be *known* at the beginning of the frame, and deadlines are constrained by the frame's length [11], [12], [13], [14], [15]. Under such assumptions, the optimal solution in each frame is a Max-Weight schedule, where the weight of each link is its *deficit* counter (i.e., a measure of how many more packets need to be transmitted from a link to meet its delivery ratio requirement). Note that unless the traffic is restricted to be periodic and synchronized across the users, such solutions are *non-causal*. Partial generalizations of the frame-based traffic are considered in [19] and [20] without performance guarantees. The optimal scheduling policy (and the real-time throughput region) for general traffic patterns and interference graphs is unknown and very difficult to characterize. *Largest-Deficit-First* (LDF) is a causal policy which extends the well-studied *Longest-Queue-First* (LQF) [3], [4] from traditional scheduling to real-time scheduling. The performance of LDF has been studied in terms of *efficiency ratio*, which is the fraction of the real-time throughput region guaranteed by LDF. Under *i.i.d.* packet arrivals and deadlines, *with no fading*, LDF was shown to achieve an efficiency ratio of at least $\frac{1}{\beta+1}$ [17], where β is the *interference degree* of the network (which is the maximum number of links that can be scheduled simultaneously out of a link and its neighboring links).

Recently, the work [21], [22] has shown that, through randomization, it is possible to design algorithms that can significantly improve over prior algorithms (such as LDF), in terms of both efficiency ratio and traffic assumptions. Specifically, [21] proposed two randomized scheduling algorithms, namely, AMIX-ND for collocated networks with an efficiency ratio of at least $\frac{e-1}{e} \approx 0.63$, and AMIX-MS for general interference graphs with an efficiency ratio of at least $\frac{|I|}{2|I|-1} > \frac{1}{2}$ ($|I|$ is the number of maximal independent sets).

Manuscript received 30 September 2021; revised 31 July 2022; accepted 10 November 2022; approved by IEEE/ACM TRANSACTIONS ON NETWORKING Editor S. Moharir. Date of publication 24 November 2022; date of current version 18 August 2023. This work was supported in part by NSF under Grant CNS-1652115 and in part by ARO under Grant W911NF1910379. An earlier version of this paper appeared in [DOI: 10.1109/INFOCOM42981.2021.9488917]. (Corresponding author: Javad Ghaderi.)

The authors are with the Department of Electrical Engineering, Columbia University, New York, NY 10027 USA (e-mail: c.tsanikidis@columbia.edu; jghaderi@columbia.edu).

This article has supplementary downloadable material available at <https://doi.org/10.1109/TNET.2022.3223315>, provided by the authors.

Digital Object Identifier 10.1109/TNET.2022.3223315

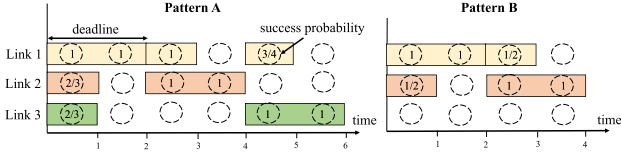


Fig. 1. An example of a Markovian traffic-and-fading process that alternates between two patterns. Each rectangle indicates a packet for a link. The left side of the rectangle corresponds to its arrival time, and its length corresponds to its deadline. The numbers in circles indicate the channel success probability of each link in each time slot.

However, the complexity of AMIX-MS can be prohibitive for implementation in large networks. Moreover, intrinsic wireless channel fading has not been considered in [21] and [22], and packet transmission over a link is assumed to be always reliable.

In this work, we consider an interference graph model of wireless network subject to fading, where packet transmissions over links are unreliable. We consider a joint traffic (packet arrival and deadline) and fading (link's success probability) process that evolves as an *unknown* Markov chain over a finite state space. Note that the addition of fading complicates the deadline-constrained scheduling problem significantly. In this case, a transmission over a link, even if all its interfering links are silent, might fail, which not only wastes the transmission, but also reduces the remaining deadline of the packet by one.

A natural question is whether the existing traditional scheduling algorithms (which focus on long-term throughput with no deadline constraints) can be used to provide guarantees for scheduling in this setting, *without assuming the knowledge of future traffic and fading process and without making frame-based assumptions*? We show that interestingly the answer is “yes”, but fading and deadlines might significantly degrade their performance guarantees.

For example, consider the initial two time slots in Pattern B in traffic-fading process of Figure 1 with two interfering links: link 1 has a packet with deadline 2 and link 2 has a packet with deadline 1. Link 1's success probability is 1 in the first two slots, and link 2's success probability is 0.5 in the first time slot. Not knowing the future traffic and fading, an opportunistic scheduler would prioritize link 1 in the first time slot and subsequently in the second time slot, but the optimal policy would always schedule link 2 in the first time slot and packet of link 1 in the second time slot.

A key insight of our work is that *careful randomization* in decision is crucial to hedge against the risk of poor decision due to lack of the knowledge of future traffic and fading. We design randomized algorithms that *strictly* outperform existing scheduling algorithms.

A. Contributions

Our main contributions can be summarized as follows:

- **Application of Traditional Scheduling Algorithms to Real-Time Scheduling.** Each non-empty link (i.e., with unexpired packets for transmission) is associated with a weight which is the product of its deficit counter and channel fading probability at that time, otherwise if the

link is empty, the weight is zero. We show that any algorithm that provides a ψ -approximation to Max-Weight Schedule (MWS) under such weights, achieves an efficiency ratio of at least $\frac{\psi}{\psi+1}$ for real-time scheduling under any Markov traffic-fading process. As a consequence, MWS policy achieves an efficiency ratio of $\frac{1}{2}$, and GMS (Greedy Maximal Scheduling, which extends LDF) provides an efficiency ratio of at least $\frac{1}{\beta+1}$.

- **Randomized Scheduling of Real-Time Traffic Over Fading Channels.** We extend [21], [22] to show the power of randomization for scheduling real-traffic traffic over fading channels. By carefully randomizing over the maximal schedules, the algorithms can achieve an efficiency ratio of at least $\frac{|X|}{2|X|-1} > \frac{1}{2}$ in any general graph under any *unknown* Markov traffic-fading process. In the special case of a collocated network with i.i.d. channel success probability q , the algorithm can achieve an efficiency ratio of at least $(\frac{q}{1-e^{-q}} + 1 - q)^{-1}$, which ranges from 0.5 to 0.632, as q varies from 0 to 1.
- **Low-Complexity and Distributed Randomized Algorithms.** To address the high complexity of the randomized algorithm in general graphs, we propose a low-complexity covering-based randomization and a myopic distributed randomization. Given a coloring of the graph using χ colors, we can achieve an efficiency ratio of at least $\frac{1}{2\chi-1}$, by randomizing over χ schedules. Moreover, we show that a myopic distributed randomization, which is simple and easily implementable, can achieve an efficiency ratio of approximately $\frac{1}{\Delta+1}$ in any graph with maximum degree Δ .

B. Notations

Some of the notations used in this paper are as follows. Given $\mathbf{y} \in \mathbb{R}^n$, $\|\mathbf{y}\| = \sum_{i=1}^n |y_i|$. We use $\text{int}(A)$ to denote the interior of set A . $[x]^+ = \max\{x, 0\}$. $\mathbb{1}(E)$ is the indicator function of event E . $|A|$ is the cardinality of set A . $\mathbb{E}^Y[\cdot]$ is used to indicate that expectation is taken with respect to random variable Y .

C. Organization

The rest of the paper is organized as follows. We start with the model and definitions in Section II. We state the main results and present the scheduling algorithms in Section III. Section IV is devoted to the proofs of main results. The simulation results are presented in Section V. Finally we end the paper with conclusions and future research directions in Section VI.

II. MODEL AND DEFINITIONS

A. Wireless Network and Interference Model

Consider a set of K links, denoted by the set $\mathcal{K}$. We assume time is divided into slots, and in every slot t , each link $l \in \mathcal{K}$ can attempt to transmit at most one packet. To model interference between links, we use the standard *interference graph* $G_I = (\mathcal{K}, E_I)$: Each vertex of G_I is a link, and there is an edge $(l_1, l_2) \in E_I$ if links l_1, l_2 interfere with each other.

Hence, two links that share an edge in G_I , cannot transmit packets at the same time. Let $\mathcal{I}$ be the set of all maximal independent sets of graph G_I . Also let $\mathcal{B}(t) \subseteq \mathcal{K}$ denote the set of nonempty links (i.e., links that have packets available to transmit) at time t . We use $M(t)$ to denote the set of links scheduled at time t . By definition, $M(t)$ is a valid *schedule* if links in $M(t)$ are nonempty and form an independent set of G_I , i.e.,

$$M(t) \subseteq (\mathcal{B}(t) \cap D), \text{ for some } D \in \mathcal{I}. \quad (1)$$

A schedule is said to be maximal if no nonempty link can be added to the schedule without violating the interference constraints.

B. Fading Model

Transmission over a link is unreliable due to wireless channel fading. To capture the channel fading over link l , we use an ON-OFF model where link l at time t is ON ($C_l(t) = 1$) with probability $q_l(t)$, otherwise it is OFF ($C_l(t) = 0$). If scheduled, transmission over link l at time t is successful only if link l is ON. Let $\mathbf{q}(t) = (q_l(t), l \in \mathcal{K})$ be the vector of success probabilities of the links. We assume that at any time slot t , the link's success probability is known to the scheduler before making a decision. At the end of time slot t , link's transmitter receives a feedback from its receiver indicating whether transmission was successful or not. A special case of this model is when $q_l(t) \in \{0, 1\}$, $\forall l \in \mathcal{K}$, in which case the channel state $C_l(t)$ is deterministically known to the scheduler, which requires periodic channel state estimation. Another special case is when $C_l(t)$ is i.i.d. with some probability q_l which eliminates the need for periodic estimation. Similar models have been used in, e.g., [16] and [23].

Overall, we define $I(t) = (I_l(t), l \in \mathcal{K})$ to denote the successful packet transmissions over the links at time t . Note that by definition, $I_l(t) = \mathbb{1}(l \in M(t))C_l(t)$, where $M(t)$ is the valid schedule selected at time t , and $C_l(t)$ is the channel state of link l .

C. Traffic Model

We assume a single-hop real-time traffic. We use $a_l(t) \leq a_{\max}$ to denote the number of packets arriving to link l at time slot t , where $a_{\max} < \infty$ is a constant. Each arriving packet has a deadline¹ which indicates the maximum delay that the packet can tolerate before successful transmission. A packet arriving at time t with deadline d has to be successfully transmitted before the end of time slot $t + d - 1$, otherwise it will be discarded. We define the traffic process $\tau(t) = (\tau_{l,d}(t), d = 1, \dots, d_{\max}, l \in \mathcal{K})$, where $\tau_{l,d}(t)$ is the number of packets with deadline d arriving to link l at time t , and $d_{\max} < \infty$.

D. Traffic and Fading Process

In general, we assume that the joint traffic and fading process $\mathbf{z}(t) = (\mathbf{q}(t), \tau(t))$ evolves as an “unknown” irreducible Markov chain over a finite state space $\mathcal{Z}$. See Figure 1 for an example of a Markovian traffic-fading process.

¹The term “relative deadline” or “slack” is often used instead.

Without loss of generality, we make the following assumption to make this Markov chain *non-trivial*: For every link l , there are two states $\mathbf{z}^l, \mathbf{z}'^l \in \mathcal{Z}$ such that $\mathbf{z}^l$ has a packet arrival with deadline d , and $\mathbf{z}'^l$ has $q_l > 0$, and there is a positive probability that $\mathbf{z}(t)$ can go from $\mathbf{z}^l$ to $\mathbf{z}'^l$ in at most d time slots. This assumption simply states that it is possible to successfully transmit some packets of every link l within their deadlines. If a link does not satisfy this condition, we can simply remove it from the system. Note that, without loss of generality, $\mathbf{z}(t)$ is assumed to be an irreducible finite-state Markov chain,² hence it is positive recurrent [24], and the time-average of any bounded function of $\mathbf{z}(t)$ is well defined, in particular the packet arrival rate for link l :

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{s=1}^t a_l(s) =: \bar{a}_l. \quad (2)$$

E. Buffer Dynamics

The buffer of link l at time t , denoted by $\Psi_l(t)$, contains the existing packets at link l which have not expired yet and also the newly arrived packets at time t . The remaining deadline of each packet in $\Psi_l(t)$ decreases by one at every time slot, until the packet is *successfully* transmitted or reaches the deadline 0, which in either case the packet is removed from $\Psi_l(t)$. We also define $\Psi(t) = (\Psi_l(t); l \in \mathcal{K})$.

F. Delivery Requirement and Deficit

As in [11], [12], [13], [14], and [17], we assume that there is a minimum delivery ratio requirement p_l (QoS requirement) for each link $l \in \mathcal{K}$. This means we must successfully deliver at least p_l fraction of the incoming packets on each link l *within their deadlines*. Formally,

$$\liminf_{t \rightarrow \infty} \frac{\sum_{s=1}^t I_l(s)}{\sum_{s=1}^t a_l(s)} \geq p_l. \quad (3)$$

We define a deficit $w_l(t)$ which measures the number of successful packet transmissions owed to link l up to time t to fulfill its minimum delivery ratio requirement. As in [14], [17], and [21], the deficit evolves as

$$w_l(t+1) = \left[w_l(t) + \tilde{a}_l(t) - I_l(t) \right]^+, \quad (4)$$

where $\tilde{a}_l(t)$ indicates the amount of deficit increase due to packet arrivals. For each packet arrival, we should increase the deficit by p_l on average. For example, we can increase the deficit by exactly p_l for each packet arrival to link l , or use a coin tossing process as in [14] and [17], i.e., each packet arrival at link l increases the deficit by one with probability p_l , and zero otherwise. We refer to $\tilde{a}_l(t)$ as the *deficit arrival* process for link l . Note that it holds that

$$\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{s=1}^t \tilde{a}_l(s) = \bar{a}_l p_l := \lambda_l, \quad l \in \mathcal{K}. \quad (5)$$

We refer to λ_l as the deficit arrival rate of link l . Note that an arriving packet is always added to the link's buffer, regardless of whether and how much deficit is added for that packet. Also

²If the chain is not irreducible, since it is finite state, there is at least one recurrent class. Starting from any state, the chain eventually enters a recurrent class and remains there for the remaining time, and hence we can truncate the chain to the recurrent class.

note that in (4) each time a packet is transmitted successfully from link l , i.e., $I_l(t) = 1$, the deficit is reduced by one. The dynamics in (4) define a deficit queueing system, with bounded increments/decrements, whose stability, e.g., in the sense

$$\limsup_{t \rightarrow \infty} \frac{1}{t} \sum_{s=1}^t \mathbb{E}[w_l(s)] < \infty, \quad (6)$$

implies (3) holds [25]. Define the vector of deficits as $\mathbf{w}(t) = (w_l(t), l \in \mathcal{K})$. The system state at time t is then defined as

$$\mathcal{S}(t) = (\mathbf{q}(t), \tau(t), \Psi(t), \mathbf{w}(t)). \quad (7)$$

Objective: Define $\mathcal{P}_C$ to be the set of all causal policies, i.e., policies that do not know the information of future arrivals, deadlines, and channel success probabilities in order to make scheduling decisions. For a given traffic-fading process $\mathbf{z}(t)$, with fixed $\bar{a}_l$, defined in (2), we are interested in causal policies that can stabilize the deficit queues for the largest set of delivery rate vectors $\mathbf{p} = (p_l, l \in \mathcal{K})$, or equivalently largest set of $\boldsymbol{\lambda} = (\lambda_l := \bar{a}_l p_l, l \in \mathcal{K})$ possible. For a given traffic process, we say the rate vector $\boldsymbol{\lambda} = (\lambda_l, l \in \mathcal{K})$ is supportable under some policy $\mu \in \mathcal{P}_C$ if all the deficit queues remain stable for that policy. Then one can define the supportable real-time rate region of the policy μ as

$$\Lambda_\mu = \{\boldsymbol{\lambda} \geq 0 : \boldsymbol{\lambda} \text{ is supportable by } \mu\}. \quad (8)$$

The supportable *real-time rate region* under all the causal policies is defined as $\Lambda = \bigcup_{\mu \in \mathcal{P}_C} \Lambda_\mu$. The overall performance of a policy μ is evaluated by the *efficiency ratio* defined as

$$\gamma_\mu^* = \sup\{\gamma : \gamma\Lambda \subseteq \Lambda_\mu\}. \quad (9)$$

For a casual policy μ , we aim to provide a *universal lower bound* on the efficiency ratio that holds for “all” Markovian traffic-fading processes, *without* knowing the Markov chain.

III. SCHEDULING ALGORITHMS AND MAIN RESULTS

Recall that $\mathcal{I}$ is the set of all maximal independent sets of interference graph G_I , and $\mathcal{B}(t)$ is the set of links that have packets to transmit at time t . At any time t , we define $\mathcal{M}(t)$ to be the set of maximal schedules. Formally, $\mathcal{M}(t) = \mathcal{I}|_{\mathcal{B}(t)}$ where $\mathcal{I}|_{\mathcal{B}(t)}$ is the set of maximal independent sets of the induced subgraph of G_I on the vertex set $\mathcal{B}(t)$.

Definition 1 (Gain of a Schedule): The *gain* of maximal schedule $M \in \mathcal{M}(t)$ is defined as the total deficit of links with a successful packet transmission at time t , i.e.,

$$\mathcal{G}_M(t) := \sum_{l \in M} C_l(t) w_l(t).$$

Note that since link l 's channel state $C_l(t)$ is a random variable, the gain of the schedule is also random.

Definition 2 (Weight of a Schedule): We define the *weight* of maximal schedule M to be its expected gain, conditional on state $\mathcal{S}(t)$, i.e.,

$$W_M(t) := \mathbb{E}[\mathcal{G}_M(t) | \mathcal{S}(t)] = \sum_{l \in M} w_l(t) q_l(t). \quad (10)$$

Note that the above expectation is with respect to the *randomness of the fading channel*.

Definition 3 (MWS): *Max-Weight Schedule* (MWS) at time t is defined as

$$M^*(t) := \arg \max_{M \in \mathcal{M}(t)} W_M(t).$$

We use $W^*(t) := W_{M^*(t)}(t)$ to denote its weight. The MWS policy is the policy that selects a MWS at any time.

In general, given a Markov policy ALG,³ let $\pi_M(t)$ be the probability that ALG selects maximal schedule M at time t . Hence, the probability that link l is scheduled at time t is

$$\phi_l(t) := \sum_{M \in \mathcal{M}(t)} \pi_M(t) \mathbb{1}(l \in M). \quad (11)$$

With a minor abuse of notation, we will use $\mathcal{G}_{ALG}(t)$ to denote the gain of the schedule selected by policy ALG at time t . Then the expected gain of ALG is given by

$$\mathbb{E}[\mathcal{G}_{ALG}(t) | \mathcal{S}(t)] = \sum_{M \in \mathcal{M}(t)} \pi_M(t) W_M(t) \quad (12)$$

$$= \sum_{l \in \mathcal{K}} \phi_l(t) q_l(t) w_l(t). \quad (13)$$

In the above, the expectation is with respect to the randomness of the channel and the decisions of policy ALG. Without loss of generality, we consider natural policies that transmit the earliest-deadline packet from every selected link in the schedule. This is because, similar to [21] and [22], if a policy transmits a packet that is not the earliest-deadline, that packet can be replaced with the earliest-deadline packet of that link, and this only improves the state. Similarly, the optimal policy always selects a maximal schedule at any time.

A. Converting Classical Non-Real-Time Algorithms for Real-Time Scheduling

We first state a theorem that, surprisingly, allows us to convert non-real-time scheduling policies to real-time scheduling policies, hence enabling the use of numerous policies from the literature of traditional non-real-time scheduling (i.e., MWS and its approximations). Specifically, policies whose expected gain is ψ -fraction of the expected gain of MWS (i.e., $W^*(t)$), yield an efficiency ratio of at least $\psi/(\psi + 1)$. The result is stated formally in Theorem 1 below.

Theorem 1: Consider any policy ALG such that, at every time t , $\mathbb{E}[\mathcal{G}_{ALG}(t) | \mathcal{S}(t)] \geq \psi W_{M^*(t)}$, whenever $\|\mathbf{w}(t)\| \geq W'$, for some finite W' . Then

$$\gamma_\mu^* \geq \frac{\psi}{\psi + 1}.$$

The proof of Theorem 1 is provided in Section IV-D.

Note that this conversion results in deadline-oblivious policies which can be preferred in cases where information about the deadlines of packets is either not accurate or not available. We remark that for certain policies, the result of Theorem 1 is tight as seen in Corollary 1.1 below, whereas for other policies the bound can be loose, as we will see later in Remark 7. We now state some of the implications of Theorem 1.

³A Markov policy is a policy that chooses the action at time t as a function of the current state $\mathcal{S}(t)$.

Corollary 1.1: MWS policy provides efficiency ratio $\gamma_{\text{MWS}}^* = \frac{1}{2}$ for real-time scheduling under Markov traffic-fading processes.

Proof: Using Theorem 1 for MWS with $\psi = 1$, we directly obtain $\gamma_{\text{MWS}}^* \geq 1/2$. We can get the opposite inequality through an adversarial example. If we consider a simple network with two interfering links without fading, then MWS reduces to LDF, which has been shown to have $\gamma^* \leq 1/2$ for Markovian traffic with deterministic deficit admission [17], [21] (recall that efficiency ratio γ^* is defined as a universal bound for all traffic-fading processes for any given graph). $\square$

Definition 4 (GMS Policy): The Greedy Maximal Scheduling (GMS) policy is defined as follows. Order the nonempty links $l \in \mathcal{B}(t)$ in the decreasing order of the product $w_l(t)q_l(t)$. Then construct a schedule recursively by including the nonempty link with the largest $w_l(t)q_l(t)$, removing its interfering links, and repeating the same procedure over the remaining links.

The following corollary extends the result of [17], which was for i.i.d. traffic without fading (i.e., for LDF), to general Markov traffic and fading processes.

Corollary 1.2: GMS provides an efficiency ratio $\gamma_{\text{GMS}}^* \geq \frac{1}{\beta+1}$ for real-time scheduling under any Markov traffic-fading process, where β is the interference degree of the graph.

Proof: The schedule M constructed by GMS satisfies $\sum_{l \in M} w_l(t)q_l(t) \geq \frac{1}{\beta} W_{M^*}(t)$ and thus we can apply Theorem 1 with $\psi = 1/\beta$ to obtain the result. The details can be found in Appendix A in the Supplementary Materials. $\square$

Note that many other results from traditional scheduling literature can be converted to real-time scheduling using Theorem 1. For example, the distributed greedy policy DIST-GREEDY in [26] also obtains $\gamma_{\text{DISTGREEDY}}^* \geq \frac{1}{\beta+1}$ but has a lower complexity than GMS.

B. Randomized Real-Time Scheduling Algorithms

In this section, we extend AMIX policies for collocated networks and general networks introduced in [21] and [22] to incorporate fading. We refer to the generalized policies as FAMIX (*Fading-based Adaptive MIX*). We describe these policies below.

1) *FAMIX-MS: Randomized Scheduling in General Graphs:* Let $R := |\{M \in \mathcal{M}(t), W_M(t) > 0\}|$ be the number of maximal schedules at time t with positive weight (based on Definition (10)). We discuss the nontrivial case of $R > 0$. We index and order the maximal schedules such that $M^{(i)} \in \mathcal{M}(t)$ has the i -th largest weight, i.e.,

$$W_{M^{(1)}}(t) \geq W_{M^{(2)}}(t) \cdots \geq W_{M^{(R)}}(t). \quad (14)$$

Define the *subharmonic* average of weight of the first n maximal schedules, $n \leq R$, to be

$$C_n(t) = \frac{n-1}{\sum_{i=1}^n (W_{M^{(i)}}(t))^{-1}}. \quad (15)$$

Then select schedule $M^{(i)}$ with probability

$$\pi_{M^{(i)}}^{\bar{n}}(t) \equiv \pi_i^{\bar{n}}(t) = \begin{cases} 1 - \frac{C_{\bar{n}}(t)}{W_{M^{(i)}}(t)}, & 1 \leq i \leq \bar{n} \\ 0, & \bar{n} < i \leq R \end{cases} \quad (16)$$

where $\bar{n}$ is the largest n such that $\{\pi_i^n(t), 1 \leq i \leq R\}$ defines a valid probability distribution, i.e., $\pi_i^n(t) \geq 0$, and $\sum_{i=1}^R \pi_i^n(t) = 1$. It is not hard to verify that this is equivalent to finding the largest n such that $\pi_n^n(t) \geq 0$, since $\pi_i^n(t)$ is decreasing with respect to i for fixed n .

Remark 1: It has been shown in [21] that $\bar{n}$ in such a distribution can be found through a binary search by looking for the first value of n with $\pi_n^n(t) \geq 0$. The algorithm FAMIX-MS above and AMIX-MS in [21] and [22] are similar, but they crucially differ in the definition of weight of a schedule (10).

The theorem below states a lower bound on the efficiency ratio of FAMIX-MS.

Theorem 2: In a general interference graph G_I with maximal independent sets $\mathcal{I}$, the efficiency ratio of FAMIX-MS is

$$\gamma_{\text{FAMIX-MS}}^* \geq \frac{|\mathcal{I}|}{2|\mathcal{I}| - 1} > \frac{1}{2}.$$

The proof of Theorem 2 is presented in Section IV-E.

Remark 2: Theorem 2 shows that with randomization we can do *strictly* better than MWS (Corollary 1.1). In particular, in the case of a complete bipartite interference graph (which has two maximal independent sets), FAMIX-MS yields an efficiency ratio of at least $2/3$ while MWS yields $1/2$.

The main disadvantage of FAMIX-MS is that it is computationally demanding as it could require randomization over all the maximal schedules. Therefore, FAMIX-MS would be better suited for networks with a small number of maximal schedules. Note that MWS also suffers from analogous complexity issues. In Section III-C, we will discuss how to design low-complexity randomized algorithms that can obtain nontrivial efficiency ratios in larger interference graphs.

2) *FAMIX-ND: Randomized Scheduling in Collocated Graphs:* Extending AMIX-ND [21] to fading channels is more challenging. In particular, the derivation in [21] relied on two main ideas: (1) it is sufficient to consider only a restricted set of “non-dominated” links for transmission, and (2) there is an ordering among the non-dominated links such that given two non-dominated links, having a packet in the buffer of one of the links is always preferred to that of the other link. Finding such domination relationship for a general Markov fading process is difficult. Here, we describe an extension under a simplified fading process, where $q_l(t) = q_l = q$, i.e., the links’ success probabilities are fixed and equal to q . We allow the channel state across links to be either independent or positively correlated, i.e., $\Pr[C_{l_2}(t) = 1 | C_{l_1}(t) = 1] \geq q_{l_2}$. The above setting could be a reasonable approximation in collocated networks where links have similar reliabilities, and an active channel for one link implies a better condition for the overall shared wireless medium.

Let $e_l(t)$ denote the deadline of the earliest-deadline packet of link l at time t . We say that link l_1 dominates link l_2 at time t if $e_{l_1}(t) \leq e_{l_2}(t)$, $w_{l_1}(t) \geq w_{l_2}(t)$, i.e., link l_1 is more urgent and has a higher deficit. Based on this definition, the set of non-dominated links at any time can be found though a simple recursive procedure as in [21], i.e., add the largest-deficit nonempty link, remove all the links dominated by it from consideration, and repeat the same procedure

over the remaining links. The following theorem describes FAMIX-ND and its efficiency ratio for a collocated network with a channel success probability q .

Theorem 3: Consider a collocated network, where $q_l = q$, $\forall l \in \mathcal{K}$, and channels of links are independent or positively correlated. Order and re-index the non-dominated links such that

$$w_1(t) \geq w_2(t) \geq \dots \geq w_n(t).$$

Starting from $i = 1$, assign probability $\pi_i(t)$ to the i -th non-dominated link,

$$\pi_i(t) = \min \left\{ \frac{1}{q} \left(1 - \frac{w_{i+1}(t)}{w_i(t)} \right), 1 - \sum_{j < i} \pi_j(t) \right\}. \quad (17)$$

FAMIX-ND selects the i -th non-dominated link with probability $\pi_i(t)$ and transmits its earliest-deadline packet. Then

$$\gamma_{\text{FAMIX-ND}}^* \geq \left(\frac{q}{1 - e^{-q}} + (1 - q) \right)^{-1} := h(q). \quad (18)$$

The proof of Theorem 3 is presented in Section IV-F.

Note that FAMIX-ND does not assign any probability to dominated links. In (17), we assign probability to non-dominated links, starting from the largest-deficit link $i = 1$, until we exhaust all the available probability, i.e., the i -th non-dominated link will receive a positive probability if $\sum_{j < i} \pi_j(t) < 1$.

Remark 3: Note that due to channel uncertainty, FAMIX-ND boosts the transmission probability of larger-deficit links. Intuitively, as q becomes small, deadlines of packets are “effectively” reduced, as each packet will need to be transmitted several times before success. For example, in the case that $q \leq 1 - \frac{w_2(t)}{w_1(t)}$, FAMIX-ND will transmit the packet of the largest-deficit link with probability 1.

Remark 4: The lower bound $h(q)$ on efficiency ratio in (18) is a monotone function of q , which increases from 0.5 to $\frac{e-1}{e}$, as q goes from 0 to 1. For $q = 1$, this recovers the result of [21] for non-fading channels, i.e., $\gamma_{\text{FAMIX-ND}}^* \geq \frac{e-1}{e}$.

In the case of unequal $q_i \in (q_{\min}, q_{\max})$, using (17) by replacing q with $q_{\min}$, will give an efficiency ratio of at least $\left(\frac{q_{\max}}{1 - e^{-q_{\min}}} + 1 - q_{\max} \right)^{-1} := h(q_{\min}, q_{\max})$. Depending on $q_{\max}$, $q_{\min}$, K , we can choose either FAMIX-ND or FAMIX-MS and achieve $\gamma^* \geq \max\{h(q_{\min}, q_{\max}), \frac{K}{2K-1}\}$.

C. Low-Complexity and Distributed Randomized Variants

The general algorithm FAMIX-MS in Section III-B potentially randomizes over all the maximal schedules. This can be computationally expensive in large networks that may have many maximal schedules. In this section, we present variants of FAMIX-MS that only need to consider a subset of the maximal schedules or are distributed.

1) Covering-Based Randomized Algorithms: We propose two variants that only need to consider a subset of the maximal schedules.

Let $\text{FAMIX-MS}|_{\mathcal{M}_0(t)}$ denote a policy that, at any time t , selects a schedule from a subset $\mathcal{M}_0(t) \subseteq \mathcal{M}(t)$, according to probabilities of FAMIX-MS computed for $\mathcal{M}_0(t)$. Proposition 1 below states a sufficient condition under which

randomization over a subset of the maximal schedules can provide a related approximation on the efficiency ratio.

Proposition 1: Consider a subset of maximal schedules $\mathcal{M}_0(t) \subseteq \mathcal{M}(t)$, and policy $\text{ALG} = \text{FAMIX-MS}|_{\mathcal{M}_0(t)}$. Suppose for every $M \in \mathcal{M}(t) \setminus \mathcal{M}_0(t)$, we have

$$\begin{aligned} \psi(t) \sum_{l \in M} w_l(t) q_l(t) \bar{\phi}_l(t) &\leq \max_{M' \in \mathcal{M}_0(t)} \sum_{l \in M'} q_l(t) w_l(t) \bar{\phi}_l(t) \\ &\quad + (1 - \psi(t)) \mathbb{E}[\mathcal{G}_{\text{ALG}}(t) | \mathcal{S}(t)], \end{aligned} \quad (19)$$

for some $\psi(t) \in (0, 1]$, where $\phi_l(t)$ was defined in (11) and $\bar{\phi}_l(t) = 1 - \phi_l(t)$. Then

$$\gamma_{\text{ALG}}^* \geq \min_{t \geq 0} \left\{ \psi(t) \frac{|\mathcal{M}_0(t)|}{2|\mathcal{M}_0(t)| - 1} \right\}. \quad (20)$$

The proof of Proposition 1 is presented in Section IV-G.

Proposition 1 requires selecting a suitable subset of maximal schedules $\mathcal{M}_0(t)$ at any time t , such that Condition (19) holds for some constant $\psi(t)$ as large as possible. Next, we present a special case in which Condition (19) holds, by considering a small set of schedules such that any other maximal schedule can be covered by them.

Lemma 4: Suppose $\mathcal{M}_0(t)$ is selected such that every $M \in \mathcal{M}(t) \setminus \mathcal{M}_0(t)$ is covered by at most ζ maximal schedules from $\mathcal{M}_0(t)$, i.e.,

$$M \subseteq \cup_{M' \in S_M} M', \text{ for some } S_M \subseteq \mathcal{M}_0(t) : |S_M| \leq \zeta. \quad (21)$$

Then Condition (19) holds with fixed $\psi(t) = \frac{1}{\zeta}$.

Proof: Using the covering definition (21), we have

$$\begin{aligned} \sum_{l \in M} q_l(t) w_l(t) \bar{\phi}_l(t) &\leq \sum_{M' \in S_M} \sum_{l \in M'} q_l(t) w_l(t) \bar{\phi}_l(t) \\ &\leq \zeta \max_{M' \in \mathcal{M}_0(t)} \sum_{l \in M'} q_l(t) w_l(t) \bar{\phi}_l(t), \end{aligned} \quad (22)$$

and hence condition (19) trivially holds for $\psi(t) = \frac{1}{\zeta}$. $\square$

In general, $\mathcal{M}_0(t)$ can be adaptive and constructed based on the link deficits and fading probabilities. Here, we apply Proposition 1 and Lemma 4 for a constant $\psi(t) = \psi$ and a family $\mathcal{M}_0(t)$ induced by fixed subset of the independent sets $\mathcal{I}_0 \subseteq \mathcal{I}$, i.e., $\mathcal{M}_0(t) = \mathcal{I}_0|_{\mathcal{B}(t)}$. With a minor abuse of notation, we refer to such algorithms as $\text{FAMIX-MS}|_{\mathcal{I}_0}$. Below, we present two such covering-based algorithms.

Corollary 4.1 (Coloring-Based Randomization): Consider a coloring of graph G_I with χ colors, which partitions the vertices of G_I into χ independent sets $\{D'_1, \dots, D'_\chi\}$. Extend these independent sets arbitrarily so that they are maximal $\{D_1, \dots, D_\chi\} := \mathcal{I}_0$. Then $\gamma_{\text{FAMIX-MS}|_{\mathcal{I}_0}}^* \geq \frac{1}{2\chi-1}$.

Proof: Any maximal schedule in $\mathcal{M} \setminus \mathcal{M}_0$ can be covered by at most χ maximal schedules in $\mathcal{M}_0$, thus $\psi = \frac{1}{\chi}$ by Lemma 4. Further, $\frac{|\mathcal{M}_0(t)|}{2|\mathcal{M}_0(t)|-1} \geq \frac{\chi}{2\chi-1}$. Thus applying Proposition 1, we get the stated efficiency ratio. $\square$

Refer to Figure 2 for an example of a valid coloring in a graph using $\chi = 4$ and an extension to 4 maximal independent sets $\mathcal{I}_0$. In this example, $\gamma_{\text{FAMIX-MS}|_{\mathcal{I}_0}}^* \geq \frac{1}{7}$.

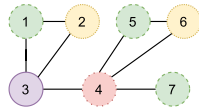


Fig. 2. An example of coloring of an interference graph using $\chi = 4$ colors. The corresponding independent sets are $D'_1 = \{4\}$, $D'_2 = \{3\}$, $D'_3 = \{1, 5, 7\}$, $D'_4 = \{2, 6\}$. These can be extended arbitrarily so that they become maximal, e.g., $D_1 = \{4, 2\}$, $D_2 = \{3, 7, 6\}$, $D_3 = \{1, 5, 7\}$, $D_4 = \{2, 6, 7\}$.

Remark 5: In general, finding an efficient coloring might be computationally demanding, but it needs to be done only once for a given G_I . There are many interesting families of graphs for which coloring can be solved efficiently. For example, for a (not necessarily complete) bipartite graph (e.g. a tree) where $\chi = 2$, we obtain $\gamma^* \geq \frac{1}{3}$. This performs much better than LDF whose efficiency ratio in a tree with maximum degree β is $\gamma^* \geq \frac{1}{2\sqrt{\beta-1}+1}$ in the case without fading (which is a special case of our setting). Another family that admits an efficient coloring are planar graphs where, by the four-color theorem, always have a 4-coloring which can be found in polynomial time [27]. Further, we remark that the independent sets D'_i could be extended adaptively at every time t , e.g., using GMS.

In the following corollary, we describe a partition-based randomization variant.

Corollary 4.2 (Partition-based Randomization): Partition the vertices of graph G_I into two sets $\mathcal{K}_1$ and $\mathcal{K}_2$, each with at most $\lceil K/2 \rceil$ vertices. Consider the set of maximal independent sets of each induced subgraph, denoted by $\mathcal{I}'_1, \mathcal{I}'_2$. Then extend these independent sets arbitrarily so that they form maximal independent sets for the whole graph G_I . Denote the new sets of maximal independent sets by $\mathcal{I}_1, \mathcal{I}_2$. Then $\gamma_{\text{FAMIS-MS}}^*|_{\mathcal{I}_1 \cup \mathcal{I}_2} \geq \frac{1}{2} \frac{|\mathcal{I}_1| + |\mathcal{I}_2|}{2(|\mathcal{I}_1| + |\mathcal{I}_2|) - 1} > \frac{1}{4}$.

Proof: By construction, every maximal independent set in G_I can be covered by at most 2 maximal independent sets in $\mathcal{I}_1 \cup \mathcal{I}_2$, and hence we have $\psi = 1/2$ by Lemma 4. Following similar arguments as in Corollary 4.1, we obtain the result. $\square$

Remark 6: By Corollary 4.2, we only need to randomize over at most $|\mathcal{I}_1| + |\mathcal{I}_2|$ maximal schedules at any time which can be much smaller than $|\mathcal{I}|$ in large graphs. The corollary can be generalized to more than 2 partitions to allow a trade-off between the guaranteed efficiency ratio and complexity.

2) *Myopic Distributed Randomized Algorithm:* We present a myopic distributed randomized algorithm that has a constant complexity.

Assume each slot is divided in two parts, a control part of duration T_C , and a packet transmission part with duration normalized to 1. Assume the deadlines of packets are expressed in terms of such time slots, each time slot having a duration $(1 + T_C)$. At the beginning of the control phase, every non-empty link $l \in \mathcal{B}(t)$ starts a timer $T_l \sim \text{Exp}(\nu_l)$, where $\text{Exp}(\nu)$ denotes an exponential distribution with rate ν . Once the timer of a link l runs down to zero, it broadcasts an announcement informing its neighbors that it will participate in data transmission, *unless* it has heard an earlier announcement from its neighboring links, or the control phase ends.

The next corollary states the efficiency ratio for the uniform timer rates.

Corollary 4.3: Consider the myopic randomized algorithm with control duration $T_C > 0$, and the same timer rate $\nu_l = \nu$ for every link l . If the maximum degree of G_I is Δ , then

$$\gamma_{\text{MYOPIC}}^* \geq \frac{1}{\frac{\Delta - \delta}{1 - \delta} + 1}, \quad (23)$$

where $\delta = e^{-T_C \nu}$.

The proof of Corollary 4.3 is provided in Section IV-H.

Note that theoretically, we can scale up the timer rate ν , so that $\delta \rightarrow 0$ for any $T_C > 0$, hence $\gamma_{\text{MYOPIC}}^* \gtrsim \frac{1}{\Delta + 1}$.

Remark 7: We note that a direct application of Theorem 1 to the myopic algorithm would yield an efficiency ratio $\gtrsim \frac{1}{\Delta + 2}$ as the myopic algorithm obtains a $\psi \approx \frac{1}{\Delta + 1}$ approximation to MWS (to see this, note that every link has roughly a probability of $\frac{1}{\Delta + 1}$ for getting service. Thus all the links of the MWS are included with this probability). Hence, Corollary 4.3 serves as an example that shows a more careful analysis can potentially improve the bound on the efficiency ratio obtained from Theorem 1.

IV. ANALYSIS TECHNIQUES AND PROOFS

We provide an overview of the techniques in our proofs.

A. Frame Construction

A key step in the analysis of our scheduling algorithms is a frame construction similar to the one in [21], but based on the *joint* traffic-fading process. The definition of frame is as follows

Definition 5 (Frames and Cycles): Starting from an initial traffic and fading state tuple $(\tau(0), \mathbf{q}(0)) = \mathbf{z} \in \mathcal{Z}$, let t_i denote the i -th return time of traffic-fading Markov chain $\mathbf{z}(t)$ to $\mathbf{z}$, $i = 1, \dots$. By convention, define $t_0 = 0$. The i -th cycle $\mathcal{C}_i$ is defined from the beginning of time slot $t_{i-1} + 1$ until the end of time slot t_i , with cycle length $C_i = t_i - t_{i-1}$. Given a fixed $k \in \mathbb{N}$, we define the i -th frame $\mathcal{F}_i^{(k)}$ as k consecutive cycles $\mathcal{C}_{(i-1)k+1}, \dots, \mathcal{C}_{ik}$, i.e., from the beginning of slot $t_{(i-1)k} + 1$ until the end of slot t_{ik} . The length of the i -th frame is denoted by $F_i^{(k)} = \sum_{j=(i-1)k+1}^{ik} C_j$. Define $\mathcal{J}(\mathcal{F}^{(k)})$ to be the space of all possible $(\tau(t), \mathbf{q}(t))$ patterns during a frame $\mathcal{F}^{(k)}$. Note that these patterns start after $\mathbf{z}$ and end with $\mathbf{z}$.

By the strong Markov property and the positive recurrence of traffic-fading Markov chain $\mathbf{z}(t)$, frame lengths $F_i^{(k)}$ are i.i.d with mean $\mathbb{E}[F^{(k)}] = k\mathbb{E}[C]$, where $\mathbb{E}[C]$ is the mean cycle length which is a bounded constant [24]. In fact, since state space $\mathcal{Z}$ is finite, all the moments of C (and $F^{(k)}$) are finite. We choose a fixed k , and, when the context is clear, drop the dependence on k in the notation.

Define the class of *non-causal $\mathcal{F}$ -framed* policies $\mathcal{P}_{NC}(\mathcal{F})$ to be the policies that, at the beginning of each frame $\mathcal{F}_i$, have complete information about the traffic-fading pattern in that frame, but have a restriction that they drop the packets that are still in the buffer at the end of the frame. Note that the number of such packets is at most $d_{\max} a_{\max} K$, which is negligible compared to the average number of packets in the frame, $\bar{a}_l \mathbb{E}[F] = \bar{a}_l k \mathbb{E}[C]$, as $k \rightarrow \infty$. Define the rate region

$$\Lambda_{NC}(\mathcal{F}) = \bigcup_{\mu \in \mathcal{P}_{NC}(\mathcal{F})} \Lambda_{\mu}. \quad (24)$$

Given a policy $\mu \in \mathcal{P}_{NC}(\mathcal{F})$, the time-average real-time service rate $\bar{I}_l$ of link l is well defined. By the renewal reward theorem (e.g. [28], Theorem 5.10), and boundedness of $\mathbb{E}[F]$,

$$\lim_{t \rightarrow \infty} \frac{\sum_{s=1}^t I_l(s)}{t} = \frac{\mathbb{E}[\sum_{t \in \mathcal{F}} I_l(t)]}{\mathbb{E}[F]} = \bar{I}_l. \quad (25)$$

Similarly for the deficit arrival rate λ_l , defined in (5),

$$\frac{\mathbb{E}[\sum_{t \in \mathcal{F}} \tilde{a}_l(t)]}{\mathbb{E}[F]} = \lambda_l, \quad l \in \mathcal{K}. \quad (26)$$

In Definition 5, each frame consists of k cycles. Using similar arguments as in [17] and [21], it is easy to see that

$$\liminf_{k \rightarrow \infty} \Lambda_{NC}(\mathcal{F}^{(k)}) \supseteq \text{int}(\Lambda),$$

Hence, if we prove that for a causal policy ALG, there exists a constant ρ , and a large k_0 , such that for all $k \geq k_0$,

$$\rho \text{int}(\Lambda_{NC}(\mathcal{F}^{(k)})) \subseteq \Lambda_{ALG}, \quad (27)$$

then it follows that $\Lambda_{ALG} \supseteq \rho \text{int}(\Lambda)$. For our algorithms, we find a ρ such that (27) holds for any traffic-fading process under our model. Then it follows that $\gamma_{ALG}^* \geq \rho$.

B. Lyapunov Argument

To prove (27), we rely on comparing the expected gain of ALG with that of the non-causal policy that maximizes the expected gain over the frame (which we refer to as *max-gain policy*). The following proposition, which is similar to that in [21], will be used to prove the main results. Its proof is similar to the proof in [21] with minor modifications to account for channel uncertainty. It is provided in Appendix B in Supplementary Materials for completeness.

Proposition 2: Consider a frame $\mathcal{F} \equiv \mathcal{F}^{(k)}$, for a fixed k based on the returns of the traffic-fading process $\mathbf{z}(t)$ to a state $\mathbf{z}$. Define the norm of initial deficits at the beginning of a frame $\|\mathbf{w}(t_0)\| = \sum_{l \in \mathcal{K}} w_l(t_0)$. Suppose for a causal policy ALG, given any $\epsilon > 0$, there is a W' such that when $\|\mathbf{w}(t_0)\| > W'$,

$$\frac{\mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | \mathcal{S}(t_0)]}{\mathbb{E}[\sum_{t \in \mathcal{F}} \hat{\mathcal{G}}_{\mu^*}(t) | \mathcal{S}(t_0)]} \geq \rho - \epsilon, \quad (28)$$

where $\mathcal{S}(t_0) = (\mathbf{q}(t_0), \boldsymbol{\tau}(t_0), \Psi(t_0), \mathbf{w}(t_0))$, and μ^* is the non-causal policy that maximizes the gain over the frame. Then for any $\lambda \in \rho \text{int}(\Lambda_{NC}(\mathcal{F}))$, the deficit queues are bounded in the sense of (6).

C. Amortized Gain Analysis

To use Proposition 2, we need to analyze the achievable gain of ALG and the non-causal policy μ^* over a frame. Since comparing the gains of the two policies directly is difficult, we adapt an amortized analysis technique from [21], initially extended from [29], [30], [31], and [32]. The general idea is as follows. Let $(\mathbf{q}(t), \boldsymbol{\tau}(t), \Psi(t), \mathbf{w}(t))$ be the state under our algorithm at time $t \in \mathcal{F}$, and $(\Psi^{\mu^*}(t), \mathbf{q}(t), \boldsymbol{\tau}(t), \Psi(t), \mathbf{w}^{\mu^*}(t))$ be the state under the optimal policy μ^* . The traffic-fading process $\mathbf{z}(t) = (\mathbf{q}(t), \boldsymbol{\tau}(t))$ is identical for both algorithms as it is independent of the actions of the scheduling policy. We change the state of μ^*

(by modifying its buffers and deficits) to make it identical to $(\mathbf{q}(t), \boldsymbol{\tau}(t), \Psi(t), \mathbf{w}(t))$, but also give μ^* an additional gain that ensures the change is advantageous for μ^* considering the rest of the frame. Let $\hat{\mathcal{G}}_{\mu^*}(t)$ denote the amortized gain of μ^* at time t with any compensated gain, which has the property that

$$\mathbb{E}\left[\sum_{t \in \mathcal{F}} \hat{\mathcal{G}}_{\mu^*}(t) | J, \mathcal{S}(t_0)\right] \geq \mathbb{E}\left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | J, \mathcal{S}(t_0)\right], \quad (29)$$

given any traffic-fading pattern $J \in \mathcal{J}(\mathcal{F})$ and initial frame state $\mathcal{S}(t_0)$. Then, the following proposition will be useful in bounding the gain and thus the efficiency ratio of our policies.

Proposition 3: Consider a Markov policy ALG that for any traffic-fading pattern $J \in \mathcal{J}(\mathcal{F})$, at any time $t \in \mathcal{F}$, satisfies

$$\rho \mathbb{E}[\hat{\mathcal{G}}_{\mu^*}(t) | J, \mathcal{S}(t)] \leq \mathbb{E}[\mathcal{G}_{ALG}(t) | J, \mathcal{S}(t)] + \mathcal{E}_F \quad (30)$$

for some $\mathcal{E}_F$ which is a measurable function of the frame length F , with $\mathbb{E}[F\mathcal{E}_F] < \infty$. Then $\gamma_{ALG}^* \geq \rho$.

Proof: First note that

$$\mathbb{E}[\mathcal{G}_{ALG}(t) | J, \mathcal{S}(t)] = \mathbb{E}[\mathcal{G}_{ALG}(t) | J, \mathcal{S}(t), \mathcal{S}(t_0)],$$

by the Markov property of ALG. Further, $\mathbb{E}[\hat{\mathcal{G}}_{\mu^*}(t) | J, \mathcal{S}(t)] = \mathbb{E}[\hat{\mathcal{G}}_{\mu^*}(t) | J, \mathcal{S}(t), \mathcal{S}(t_0)]$, since the amortized gain of the max-gain policy does not depend on the past state given the current state and future traffic-fading pattern (note that this amortized gain might depend on policy ALG, which itself is Markov). Hence, taking expectation of both sides of (30), conditional on $J, \mathcal{S}(t_0)$, and using the law of iterated expectations, we get

$$\rho \mathbb{E}[\hat{\mathcal{G}}_{\mu^*}(t) | J, \mathcal{S}(t_0)] \leq \mathbb{E}[\mathcal{G}_{ALG}(t) | J, \mathcal{S}(t_0)] + \mathbb{E}[\mathcal{E}_F | J, \mathcal{S}(t_0)]$$

Summing over the frame, and using the fact that $\mathbb{E}[\mathcal{E}_F | J, \mathcal{S}(t_0)] = \mathcal{E}_F$ since F is determined by J , we have

$$\rho \mathbb{E}\left[\sum_{t \in \mathcal{F}} \hat{\mathcal{G}}_{\mu^*}(t) | J, \mathcal{S}(t_0)\right] \leq \mathbb{E}\left[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | J, \mathcal{S}(t_0)\right] + F\mathcal{E}_F.$$

Using the definition (29), and taking expectations over the randomness of the traffic-fading pattern, we obtain

$$\rho \mathbb{E}\left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0)\right] \leq \mathbb{E}\left[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | \mathcal{S}(t_0)\right] + \mathbb{E}[F\mathcal{E}_F].$$

Note that by assumption, $\mathbb{E}[F\mathcal{E}_F] < \infty$. If we show that $\lim_{\|\mathbf{w}(t_0)\| \rightarrow \infty} \mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0)] = \infty$, then

$$\rho \leq \lim_{\|\mathbf{w}(t_0)\| \rightarrow \infty} \frac{\mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | \mathcal{S}(t_0)]}{\mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0)]},$$

and using Proposition 2, we obtain that $\gamma_{\mu} \geq \rho$.

To that end, consider the link $l_1 := \arg \max_l w_l(t_0)$. Since μ^* maximizes the expected gain over the frame, its expected gain cannot be lower than a policy that simply schedules link l_1 whenever it has available packets. Choose $\mathcal{F} \equiv \mathcal{F}^k$ with any $k \geq 2$. Then by the non-trivial traffic-fading Markov chain assumption (Section II), there is a traffic-fading pattern J' of length $F_{J'}$ with some nonzero probability of occurring $\Pr(J') > 0$, in which a packet arrives for link l_1 at some time $t_1 \in \mathcal{F}$, and at a later time $t_2 \in \mathcal{F}$ before the packet

expires, we have $q_l(t_2) > 0$, i.e., $l_1 \in \mathcal{B}(t_2)$ and $\mathbb{E}[C_{l_1}(t_2)] = q_{l_1}(t_2) > 0$. Then trivially

$$\begin{aligned} \mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0) \right] &\geq \mathbb{E} \left[w_{l_1}(t_2) C_{l_1}(t_2) | J', \mathcal{S}(t_0) \right] \Pr(J') \\ &\geq (w_{l_1}(t_0) - F_{J'}) q_{l_1}(t_2) \Pr(J') \\ &\geq \left(\frac{\|\mathbf{w}(t_0)\|}{K} - F_{J'} \right) q_{l_1}(t_2) \Pr(J'), \end{aligned}$$

which shows $\lim_{\|\mathbf{w}(t_0)\| \rightarrow \infty} \mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0) \right] = \infty$. $\square$

The following lemma describes a generic amortized gain computation for general networks that allows us to modify the state of the max-gain policy during a frame to match the state of the considered policy ALG. Recall from (11) that $\phi_l(t)$ is the probability that link l is scheduled under ALG, and $\pi_M(t)$ is the probability that schedule M is selected. Below we drop their dependence on time to simplify the notation.

Lemma 5: For any pattern $J \in \mathcal{J}(\mathcal{F})$ in a frame $\mathcal{F}$, given a Markov policy ALG, the amortized gain of the max-gain policy μ^ (w.r.t. ALG) if it selects maximal schedule $M \in \mathcal{M}$ at time t in the frame, is given by*

$$\begin{aligned} \mathbb{E} \left[\hat{\mathcal{G}}_{\mu^*}^{(M)}(t) | J, \mathcal{S}(t) \right] &= W_M(t) + \sum_{l \notin M} w_l(t) \phi_l q_l(t) + \mathcal{E}_m \\ &\leq W_M(t)(1 - \pi_M) + \mathbb{E}[\mathcal{G}_{ALG}(t) | \mathcal{S}(t)] \\ &\quad + \mathcal{E}_m, \end{aligned}$$

where $\mathcal{E}_m := K(F + a_{\max} d_{\max})$.

Proof: The main idea is similar to the one in [21], but we have to account for fading. Suppose μ^* attempts to transmit the earliest-deadline packets of links from schedule M whereas ALG attempts the earliest-deadline packets from schedule M' . We need to modify the state of μ^* , i.e., the buffers and deficits, so it is identical with the state of ALG. To achieve this, we allow μ^* to additionally transmit the packets of links successfully transmitted by ALG but not μ^* , i.e., $S_1 = (M'(t) \setminus M(t)) \cap \{l \in \mathcal{K} : C_l(t) = 1\}$. Transmitting such packets for μ^* might not be advantageous for its total gain as the weight of these packets can increase by $a_{\max} d_{\max}$ before they could be transmitted. Thus giving an additional reward $K a_{\max} d_{\max}$ to μ^* guarantees that the modification is advantageous. Further, we insert the packets transmitted successfully by μ^* but not ALG, i.e., $S_2 = (M(t) \setminus M'(t)) \cap \{l \in \mathcal{K} : C_l(t) = 1\}$, back to its buffers (which is advantageous for μ^*). Further to make the deficits identical, we increase the deficit counters of μ^* for links in S_2 , which is advantageous for the total-gain within the frame. Additionally, we decrease the deficit counters of μ^* for the links in S_1 . This change might not be advantageous for the total gain of μ^* , thus we give it extra reward for every possible subsequent transmission over links, which is at most KF . Hence the total additional compensation is $\mathcal{E}_m = K(F + a_{\max} d_{\max})$.

Following the above argument, the expected amortized gain of μ^* , when it selects M , is bounded as

$$\begin{aligned} \mathbb{E} \left[\hat{\mathcal{G}}_{\mu^*}^{(M)}(t) | J, \mathcal{S}(t) \right] \\ = \sum_{l \in M} q_l(t) w_l(t) \end{aligned}$$

$$\begin{aligned} &+ \sum_{M' \in \mathcal{M} \setminus \{M\}} \pi_{M'}(t) \sum_{l \in M' \setminus M} q_l(t) w_l(t) + \mathcal{E}_m \\ &\stackrel{(a)}{=} W_M(t) + \sum_{l \notin M} w_l(t) \phi_l q_l(t) + \mathcal{E}_m \end{aligned} \quad (31)$$

where (a) follows from definitions of ϕ_l and $W_M(t)$. The inequality in the lemma's statement follows by noting that

$$(31) \leq W_M(t) + \sum_{M' \in \mathcal{M} \setminus \{M\}} W_{M'}(t) \pi_{M'} + \mathcal{E}_m,$$

since $\sum_{l \in M' \setminus M} q_l(t) w_l(t) \leq W_{M'}(t)$, and by using (12). $\square$

D. Proof of Theorem 1: Conversion Result

In the rest of proofs, for notional compactness, we define $\mathbb{E}_{t,J}[\cdot] := \mathbb{E}[\cdot | \mathcal{S}(t), J]$ and $\mathbb{E}_t[\cdot] := \mathbb{E}[\cdot | \mathcal{S}(t)]$. Now consider a pattern $J \in \mathcal{J}(\mathcal{F})$. When $\|\mathbf{w}(t)\| \geq W'$, the amortized gain of μ^* , if it selects schedule M , is bounded as

$$\begin{aligned} \mathbb{E}_{t,J} \left[\hat{\mathcal{G}}_{\mu^*}^{(M)}(t) \right] &\stackrel{(a)}{\leq} W_M(t)(1 - \pi_M) + \mathbb{E}_t[\mathcal{G}_{ALG}(t)] + \mathcal{E}_m \\ &\leq W_{M^*}(t) + \mathbb{E}_t[\mathcal{G}_{ALG}(t)] + \mathcal{E}_m \\ &= \mathbb{E}_t[\mathcal{G}_{ALG}(t)] \left(\frac{W_{M^*}(t)}{\mathbb{E}_t[\mathcal{G}_{ALG}(t)]} + 1 \right) + \mathcal{E}_m \\ &\stackrel{(b)}{\leq} \mathbb{E}_t[\mathcal{G}_{ALG}(t)](1/\psi + 1) + \mathcal{E}_m, \end{aligned}$$

where in (a) we used Lemma 5, and in (b) we used the main assumption that ALG obtains ψ fraction of the maximum weight schedule. This inequality does not depend on the particular choice M . Similarly in the case that $\|\mathbf{w}(t)\| \leq W'$, it can be seen from (13) and Lemma 5 that $\mathbb{E}_{t,J} \left[\hat{\mathcal{G}}_{\mu^*}^{(M)}(t) \right] \leq 2W' + \mathcal{E}_m$. Consequently in either case, $\mathbb{E}_{t,J} \left[\hat{\mathcal{G}}_{\mu^*}^{(M)}(t) \right] \leq 2W' + \mathcal{E}_m + \mathbb{E}_t[\mathcal{G}_{ALG}(t)](1/\psi + 1)$. Applying Proposition 3 with $\rho = (1/\psi + 1)^{-1}$, we obtain the result.

E. Analysis of FAMIX-MS: Proof of Theorem 2

Using Lemma 5, and probabilities of FAMIX-MS indicated in (16), in both cases of $i \leq \bar{n}$ and $i > \bar{n}$, the amortized gain of μ^* can be bounded as

$$\begin{aligned} \mathbb{E}_{t,J} \left[\hat{\mathcal{G}}_{\mu^*}^{(M_i)}(t) \right] &\leq W_{M^{(i)}}(t)(1 - \pi_i) + \mathbb{E}_t[\mathcal{G}_{ALG}(t)] + \mathcal{E}_m \\ &\leq C_{\bar{n}}(t) + \mathbb{E}_t[\mathcal{G}_{ALG}(t)] + \mathcal{E}_m. \end{aligned}$$

Thus regardless of the schedule selected by μ^* , we have

$$\mathbb{E}_{t,J} \left[\hat{\mathcal{G}}_{\mu^*}(t) \right] \leq \mathbb{E}_t[\mathcal{G}_{ALG}(t)] \left(\frac{C_{\bar{n}}(t)}{\mathbb{E}_t[\mathcal{G}_{ALG}(t)]} + 1 \right) + \mathcal{E}_m$$

Note that $|\mathcal{I}| \geq |\mathcal{M}| \geq \bar{n}$, hence it suffices to show that

$$\frac{C_{\bar{n}}(t)}{\mathbb{E}_t[\mathcal{G}_{ALG}(t)]} \leq \frac{\bar{n} - 1}{\bar{n}}, \quad (32)$$

as then by Proposition 3 it will follow that $\gamma_{\text{FAMIX-MS}}^* \geq \left(\frac{|\mathcal{I}| - 1}{|\mathcal{I}|} + 1 \right)^{-1} = \frac{|\mathcal{I}|}{2|\mathcal{I}| - 1}$. To show (32), note that

$$\frac{C_{\bar{n}}(t)}{\mathbb{E}_t[\mathcal{G}_{ALG}(t)]} = \frac{\bar{n} C_{\bar{n}}(t) - (\bar{n} - 1) \mathbb{E}_t[\mathcal{G}_{ALG}(t)]}{\bar{n} \mathbb{E}_t[\mathcal{G}_{ALG}(t)]} + \frac{\bar{n} - 1}{\bar{n}}$$

hence, we need to show

$$\bar{n}C_{\bar{n}}(t) - (\bar{n} - 1)\mathbb{E}_t[\mathcal{G}_{ALG}(t)] \leq 0. \quad (33)$$

The inequality (33) holds because by (12) and (16), it follows that $\mathbb{E}_t[\mathcal{G}_{ALG}(t)] = \sum_{j=1}^{\bar{n}} W_{M(j)}(t) - \bar{n}C_{\bar{n}}(t)$ and hence (33) becomes

$$C_{\bar{n}}(t)\bar{n}^2 \leq (\bar{n} - 1) \sum_{j=1}^{\bar{n}} W_{M(j)}(t),$$

which holds by (15) and the arithmetic mean-harmonic mean inequality applied to $\{W_{M(i)}(t), i = 1, \dots, \bar{n}\}$.

F. Analysis of FAMIX-ND: Proof of Theorem 3

We first state the following Lemma that allows us to focus on policies that transmit from non-dominated links. Due to the uncertainty of the channels, our analysis has to deviate from the one in [21]. Here, we rely on a coupling argument to prove the following lemma.

Lemma 6: Given $J \in \mathcal{J}(\mathcal{F})$, let $\tilde{\mu}^$ be the maximum-gain policy that transmits only from non-dominated links at any time (we refer to $\tilde{\mu}^*$ as max-gain ND-policy) and the maximum-gain policy μ^* that can transmit any packet, then*

$$\mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{\tilde{\mu}^*}(t) | \mathcal{S}(t_0), J] \geq \mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0), J] - (a_{\max} + 1)F^2.$$

Proof: First we consider the case where deficits do not vary over time regardless of arrivals or transmissions. We argue that in this case transmitting from a non-dominated link has always higher total expected gain from any state considering the remaining frame for any pattern J . Let $V_l(t, \Psi, J)$ denote the maximum expected gain over the remaining frame for the given pattern J , if at the current time slot $t \in \mathcal{F}$ with buffers Ψ , link l is scheduled. Further, define $V(t, \Psi, J) = \max_l \{V_l(t, \Psi, J)\}$. Finally define $V_{ND}(t, S, J)$ to be the maximum expected gain for the remaining frame over policies that only schedule non-dominated links. Consider the earliest-deadline packets P_u, P_v of links u, v , with u dominating v . We argue that transmitting P_u yields a higher expected gain than that of P_v . If P_u is scheduled we have:

$$V_u(t, \Psi, J) = q(w_u + V(t+1, \Psi'_{(P_u)}, J)) + (1-q)V(t+1, \Psi', J),$$

where Ψ' are the updated buffers due to regular buffer dynamics but ignoring the scheduled packet by our policy, and $\Psi'_{(P)}$ indicates buffer Ψ' without packet P . Similar expression can be written if P_v is scheduled. Hence, to show $V_u(t, \Psi, J) \geq V_v(t, \Psi, J)$, equivalently we can show

$$w_u + V(t+1, \Psi'_{(P_u)}, J) \geq w_v + V(t+1, \Psi'_{(P_v)}, J). \quad (34)$$

Now note that by the domination definition, the deadline of packet v is at least as long as that of u , hence the optimal policy μ for $V(t+1, \Psi'_{(P_v)}, J)$ that can attempt packet u can be used to construct a policy μ' for $V(t+1, \Psi'_{(P_u)}, J)$ that attempts v instead of u whenever μ would have scheduled P_u . This is possible since P_v does not expire before P_u . Now consider a coupling for the channel outcomes in μ and μ' such that channel $C_u^\mu(t)$ under μ is identical to $C_v^{\mu'}(t)$ under μ' . This

does not affect the marginal success probability of links as both links u, v have the same probability of success q . Further each success when transmitting v under μ' results in reward w_v , instead of w_u under μ , therefore the gain of policy μ' is stochastically larger than that of policy μ , and hence in the expectation,

$$V(t+1, \Psi'_{(P_u)}, J) \geq (w_v - w_u) + V(t+1, \Psi'_{(P_v)}, J)$$

Thus in the fixed-deficit case we have shown that $V_{ND}(t, \Psi, J) \geq V(t, \Psi, J)$.

Now consider the case with time-varying deficits, and let $\tilde{V}(t, \Psi, J)$ denote the corresponding quantities. In this notation $\tilde{V}(t_0, \Psi_{t_0}, J) = \mathbb{E}[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | J, \mathcal{S}(t_0)]$. If pattern J has length F , then (a): $V(t, \Psi, J) \geq \tilde{V}(t, \Psi, J) - a_{\max}F^2$ since any transmission under any policy in the varying-deficit case will yield a gain of at most $a_{\max}F$ higher than the fixed-deficit case, since the deficit cannot increase more than $a_{\max}F$ during a frame, and we can have at most F such transmissions. Similarly we obtain (b): $\tilde{V}_{ND}(t, \Psi, J) + F^2 \geq V_{ND}(t, \Psi, J)$, as in the time-varying deficit case every transmission can result in at most F gain less than the fixed-deficit case, and we can have at most F such transmissions. By (a) and (b), we obtain: $\tilde{V}_{ND}(t, \Psi, J) \geq \tilde{V}(t, \Psi, J) - (a_{\max} + 1)F^2$. $\square$

Proof of Theorem 3. In view of Lemma 6, we perform an amortized gain analysis for FAMIX-ND in comparison with ND-policies. Due to the presence of fading, the gain analysis performed in [21] is not applicable in our case and the state of the channel needs to be considered. Suppose FAMIX-ND decides to schedule the earliest-deadline packet $P_f = (w_f, e_f)$ of link f and ND-policy $\tilde{\mu}^*$ transmits a packet $P_z = (w_z, e_z)$ from a different non-dominated link z ($z \neq f$). The state of $\tilde{\mu}^*$ and FAMIX-ND will be different in the following four cases,

- 1) $C_f(t) = 1, C_z(t) = 0$: In this case, we make the buffers of the two identical by allowing $\tilde{\mu}^*$ to also transmit P_f , and give extra reward $a_{\max}d_{\max}$. Further we reduce the deficit in $\tilde{\mu}^*$ for link f by 1 and we give reward F to $\tilde{\mu}^*$ similarly as in the proof of Lemma 5.
- 2) $C_f(t) = 0, C_z(t) = 1$: We replace P_z in the buffer of $\tilde{\mu}^*$ and increase the deficit of link z in $\tilde{\mu}^*$ which are advantageous for $\tilde{\mu}^*$.
- 3) $e_f \leq e_z, w_f \leq w_z$, and $C_f(t) = 1, C_z(t) = 1$: We replace P_f from link f with P_z in link z in buffers of $\tilde{\mu}^*$. Packet P_z has higher deadline and higher weight at time t . Since both packets will expire in at most $d_{\max}$ slots, the deficit of f can only increase by at most $d_{\max}a_{\max}$ before P_f expires, whereas the deficit of z can decrease by at most $d_{\max}$. Therefore giving $\tilde{\mu}^*$ additional compensation of $(1 + a_{\max})d_{\max}$ will guarantee that the modification is advantageous. Further, we decrease the deficit of link f by one ($w_f - 1$ in $\tilde{\mu}^*$) and we increase the deficit of link z by one ($w_z + 1$ in $\tilde{\mu}^*$). To compensate for the decrease in deficit of link f we give $\tilde{\mu}^*$ extra gain of F . Hence, the total compensation is bounded by $F + (a_{\max} + 1)d_{\max}$.
- 4) $e_z \leq e_f, w_z \leq w_f$ and $C_f(t) = 1, C_z(t) = 1$: In this case, we allow $\tilde{\mu}^*$ to additionally transmit packet P_f at time t and give extra reward $a_{\max}d_{\max}$, and inject a

copy of packet p_z to the buffer of link z . This makes the buffers identical, but results in the decrease of deficit of link f by one. To guarantee the change is advantageous for $\tilde{\mu}^*$ considering the rest of the frame, we give it extra reward F .

In all cases, the additional compensation is bounded by $\mathcal{E}_0 = F + (a_{max} + 1)d_{max}$. Let $q_{ij} = \Pr(C_i(t) = 1, C_j(t) = 1) \geq q^2$, or equivalently $q_{ij} = q^2 + \epsilon_{ij}$ for some $\epsilon_{ij} \geq 0$, by the positive correlation assumption, and hence $\Pr(C_i(t) = 1, C_j(t) = 0) = q - q_{ij} = q(1 - q) - \epsilon_{ij}$. Then,

$$\begin{aligned} \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(i)}(t)] &\leq qw_i + \sum_j \pi_j (q - q_{ij}) w_j + \sum_{j < i} w_j \pi_j q_{ij} + \\ \mathcal{E}_0 &= qw_i + (1 - q) \sum_j \pi_j q w_j + q \sum_{j < i} w_j \pi_j q \\ &\quad - \sum_j \pi_j \epsilon_{ij} w_j + \sum_{j < i} \pi_j \epsilon_{ij} w_j + \mathcal{E}_0 \\ &\leq q[w_i + (1 - q) \sum_j \pi_j w_j + q \sum_{j < i} w_j \pi_j] + \mathcal{E}_0 \end{aligned}$$

Now notice that for $\pi_i \leq \frac{1}{q}(1 - \frac{w_{i+1}}{w_i})$ the right-hand-side of the above inequality is maximized for $i = 1$. This can be verified by computing the difference of the gains for choices $i, i + 1$. Hence, we have

$$\begin{aligned} \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(i)}(t)] &\leq \max_i \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(i)}(t)] = \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(1)}(t)] \\ &= qw_1 + (1 - q) \mathbb{E}_t [\mathcal{G}_{ALG}(t)] + \mathcal{E}_0. \end{aligned} \quad (35)$$

Let $\bar{n}$ be the number of links with positive probability in (17). For the gain of FAMIX-ND, we have

$$\begin{aligned} \mathbb{E}_t [\mathcal{G}_{ALG}(t)] &= \sum_{i=1}^{\bar{n}} q \pi_i w_i \\ &= \sum_{i=1}^{\bar{n}-1} (w_i - w_{i+1}) + q(1 - \sum_{i=1}^{\bar{n}-1} \pi_i) w_{\bar{n}} \\ &= w_1 + w_{\bar{n}}(-1 + q - q \sum_{i=1}^{\bar{n}-1} \pi_i) \\ &\stackrel{(a)}{=} w_1 \left(1 - (1 - q + q \sum_{i=1}^{\bar{n}-1} \pi_i) \prod_{i=1}^{\bar{n}-1} (1 - \pi_i q) \right) \\ &\stackrel{(b)}{\geq} w_1 (1 - \left(\frac{\bar{n} - q}{\bar{n}} \right)^{\bar{n}}) \geq w_1 (1 - e^{-q}), \end{aligned}$$

where in (a), by the definition of π_i in (17), we used: $w_{\bar{n}} = w_1 \prod_{i=1}^{\bar{n}-1} (1 - \pi_i q)$. In (b) we used the geometric-arithmetic inequality. Using the above relation and (35) we get

$$\mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(i)}(t)] \leq \mathbb{E}_t [\mathcal{G}_{ALG}(t)] \left(\frac{q}{1 - e^{-q}} + (1 - q) \right) + F \mathcal{E}_0.$$

Using similar arguments as in the proof of Proposition 3, we can sum over the entire frame to obtain

$$\begin{aligned} \mathbb{E} \left[\sum_{t \in \mathcal{F}} \hat{G}_{\mu^*}^{(i)}(t) | J, \mathcal{S}(t_0) \right] \\ \leq \mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | J, \mathcal{S}(t_0) \right] \left(\frac{q}{1 - e^{-q}} + (1 - q) \right) + F \mathcal{E}_0. \end{aligned}$$

Due to the amortized analysis, (29) holds. Then by using Lemma 6 we obtain

$$\mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | J, \mathcal{S}(t_0) \right] - F^2(a_{max} + 1)$$

$$\leq \mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | J, \mathcal{S}(t_0) \right] \left(\frac{q}{1 - e^{-q}} + (1 - q) \right) + F \mathcal{E}_0.$$

Let $\rho = (\frac{q}{1 - e^{-q}} + 1 - q)^{-1}$. After taking expectation with respect to the randomness of the frame, rearranging terms and dividing appropriately, we obtain

$$\rho \leq \frac{\mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{ALG}(t) | \mathcal{S}(t_0) \right]}{\mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0) \right]} + \frac{\mathbb{E} [F \mathcal{E}_0 + F^2(a_{max} + 1)]}{\mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0) \right]} \rho$$

The proof is concluded as in the proof of Proposition 3, by $\lim_{\|\mathbf{w}(t_0)\| \rightarrow \infty} \mathbb{E} \left[\sum_{t \in \mathcal{F}} \mathcal{G}_{\mu^*}(t) | \mathcal{S}(t_0) \right] = \infty$, and applying Proposition 2.

G. Proof of Proposition 1: Low-Complexity Variants

Suppose that under ALG we have

$$\psi(t) \max_{M \in \mathcal{M}} \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(M)}(t)] \leq \max_{M \in \mathcal{M}_0(t)} \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(M)}(t)]. \quad (36)$$

Since ALG randomizes according to the probabilities of FAMIX-MS over the maximal schedules in $\mathcal{M}_0(t)$, we have

$$\xi(t) \max_{M \in \mathcal{M}_0(t)} \{ \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(M)}(t)] \} \leq \mathbb{E}_t [\mathcal{G}_{ALG}(t)] + \mathcal{E}, \quad (37)$$

where $\xi(t) := \frac{|\mathcal{M}_0(t)|}{2|\mathcal{M}_0(t)| - 1}$, by using the same arguments as in the proof of Theorem 2 applied to $\mathcal{M}_0(t)$. By (36) and (37),

$$\min_{t \geq 0} \{ \xi(t) \psi(t) \} \mathbb{E}_{t,J} [\hat{G}_{\mu^*}(t)] \leq \mathbb{E}_t [\mathcal{G}_{ALG}(t)] + \mathcal{E}.$$

Then by Proposition 3, it follows that $\gamma \geq \min_{t \geq 0} \{ \psi(t) \xi(t) \}$. Hence to conclude the proof, we simply need to argue that (19) implies (36) Using Lemma 5, and (13), it can be shown that

$$\mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(M)}(t)] = \sum_{l \in M} w_l(t) q_l(t) \bar{\phi}_l + \mathbb{E}_t [\mathcal{G}_{ALG}(t)] + \mathcal{E}_m$$

Plugging the above expression in (36), it can be seen that (19) is a sufficient condition for (36) to hold.

H. Proof of Corollary 4.3: Myopic Distributed Algorithm

The amortized gain can be written as

$$\begin{aligned} \mathbb{E}_{t,J} [\hat{G}_{\mu^*}^{(M)}(t)] &= \sum_{l \in M(t)} q_l(t) w_l(t) \bar{\phi}_l + \mathbb{E} [\mathcal{G}_{ALG}(t)] + \mathcal{E}_m \\ &\stackrel{(a)}{=} \mathbb{E}_{t,J} [\mathcal{G}_{ALG}(t)] \left(\frac{\sum_{l \in M(t)} q_l(t) w_l(t) \bar{\phi}_l}{\sum_l \phi_l q_l(t) w_l(t)} + 1 \right) + \mathcal{E}_m \end{aligned} \quad (38)$$

where in (a), we used (13).

Let $A_l = \bigcap_{z \in N(l)} \{T_l < T_z\}$ be the event that the timer of l expires earlier than that of its neighbors. We can obtain a bound on the probability of link getting scheduled as follows:

$$\begin{aligned} \phi_l &\geq \Pr [A_l \cap \{T_l < T_C\}] = \Pr [A_l] \Pr [T_l < T_C | A_l] \\ &\stackrel{(a)}{\geq} \frac{\nu_l}{\nu_l + \sum_{z \in N(z)} \nu_z} (1 - e^{-(\nu_l + \sum_{z \in N(z)} \nu_z) T_C}) \\ &\stackrel{(b)}{\geq} \frac{\nu_l}{\nu_l + \sum_{z \in N(z)} \nu_z} (1 - \delta), \end{aligned} \quad (39)$$

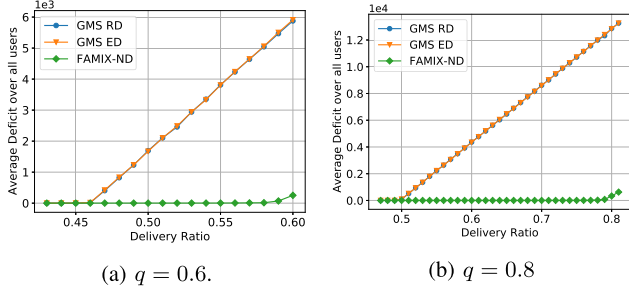


Fig. 3. Performance comparison of various algorithms in a simple two-link collocated network under two different channel success probabilities q .

where (a) follows from the fact that $T_l|A_l \sim \text{Exp}(\nu_l + \sum_{z \in N(z)} \nu_z)$ (easy to verify), and (b) from the choice of T_C . Hence, using (39), in the case that $\nu_l = \nu$, we get $\phi_l \geq \frac{1-\delta}{1+\Delta}$. Hence,

$$\begin{aligned} & \frac{\sum_{l \in M(t)} q_l(t) w_l(t) (1 - \phi_l)}{\sum_l \phi_l q_l(t) w_l(t)} \\ & \leq \frac{\sum_{l \in M(t)} q_l(t) w_l(t) (1 - \frac{1-\delta}{\Delta+1})}{\sum_l q_l(t) w_l(t) \frac{1-\delta}{\Delta+1}} \\ & \leq \frac{\sum_{l \in M(t)} q_l(t) w_l(t) (1 - \frac{1-\delta}{\Delta+1})}{\sum_{l \in M(t)} q_l(t) w_l(t) \frac{1-\delta}{\Delta+1}} \\ & \leq \frac{\Delta - \delta}{1 - \delta}, \end{aligned}$$

where the first inequality is due to the fact that the ratio is a decreasing function of ϕ_l , thus using (38),

$$\mathbb{E}_{t,J} [\hat{\mathcal{G}}_{\mu^*}^{(M)}(t)] \leq \mathbb{E}_{t,J} [\mathcal{G}_{ALG}(t)] \left(\frac{\Delta - \delta}{1 - \delta} + 1 \right) + \mathcal{E}_m.$$

The using Proposition 3, we obtain the final result.

V. SIMULATION RESULTS

We performed simulations under several networks and traffic-fading scenarios. We compare the performance of our algorithms FAMIX-ND and FAMIX-MS with GMS and MWS (see Section III). Since GMS does not specify the tie breaking rule when two or more links have the same weight, we consider two variants, GMS-ED and GMS-RD. The first variant, GMS-ED, break ties by transmitting the packet with the earliest deadline (ED) among all the links with equal weight. GMS-RD on the other hand break ties randomly (RD).

A. Simulation Results for Collocated Networks

In collocated network, MWS coincides with GMS, hence we compare FAMIX-ND with GMS (-ED, -RD).

First, we consider a simple network with two interfering links. Traffic is generated according to Pattern B of Figure 1, but with i.i.d. channel success probability q . We set an equal target delivery ratio p for both links. Figure 3a shows the average deficit under different algorithms for different values of p when $q = 0.6$. While GMS-RD and GMS-ED cannot achieve p greater than 0.47, FAMIX-MS can achieve p up to 0.58. Note that the optimal policy cannot achieve $p > 0.6$ in

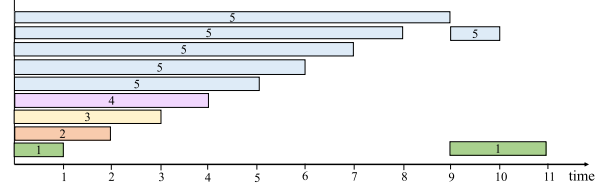


Fig. 4. Traffic pattern C used in the collocated network in Figure 5.

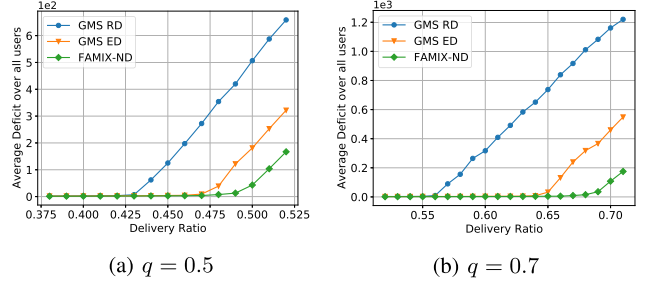


Fig. 5. Performance comparison in a 5-link collocated network using the traffic pattern of Figure 4, under two different channel success probabilities q .

this scenario due to the fading probability $q = 0.6$, and hence FAMIX-ND is near optimal. For comparison, Figure 3b shows the evaluation result for $q = 0.8$. We observe that the gap between the maximum achievable delivery ratio of GMS and FAMIX-ND widens for larger values of q . This behavior is consistent with the lower bounds on the efficiency ratios presented in this paper.

We observe a similar behavior for larger collocated networks. For example, consider a collocated network with 5 links, and a channel success probability of $q = 0.5$ for each link. The traffic pattern is indicated in Figure 4 and the evaluation result is shown in Figure 5a. This is the same traffic pattern used in [22]. Here, we observe that even with fading, the randomized policies introduced in this work outperform other algorithms, i.e., FAMIX-ND can support the highest delivery ratio requirement, followed by GMS-ED, and then by GMS-RD. Further, increasing the channel success probability to $q = 0.7$ widens the performance gap between the policies, as seen in Figure 5b. Note that the value of q sets a trivial upper bound of the maximum achievable delivery ratio of all policies, therefore FAMIX-ND is near optimal in all the simulations.

Finally we simulate a network with $q = 0.8$, 5 links and Bernoulli arrivals and random deadlines, uniformly chosen between 1 to 6. The results are shown in Figure 6. Despite the presence of randomness in the traffic, FAMIX-ND still demonstrates good performance compared to the maximum achievable delivery ratio (achieving an efficiency ratio of at least $0.73/0.80 \approx 0.91$) and to GMS.

B. Simulation Results for General Networks

1) *Efficiency Ratio Comparison in Random Geometric Graphs:* In Figure 7, we generate 100 (generalized) random geometric graphs to represent interference graphs of networks. Each network has 50 links. Links are assigned uniformly random locations in a square area, and are selected to interfere

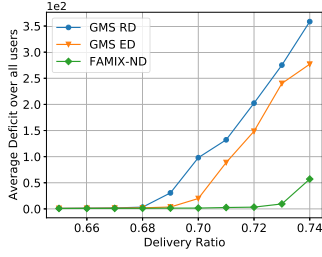


Fig. 6. Performance comparison in a 5-link collocated network, with $q = 0.8$, and packets with random deadlines.

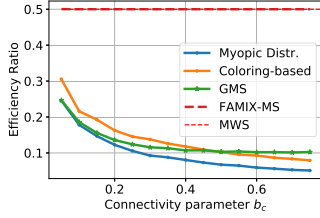


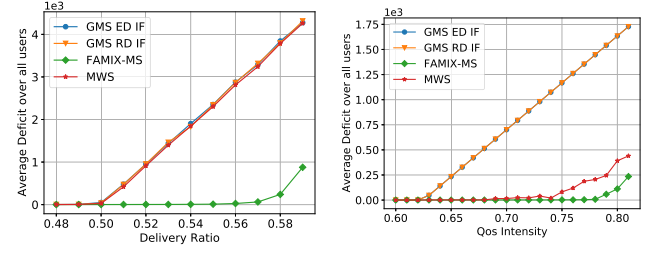
Fig. 7. Comparison of efficiency ratios of different algorithms in random graphs with varying connectivity parameter b_c .

with each other based on their distance (a.k.a. Waxman model). More specifically, two randomly located nodes at x_a and x_b have a probability of being connected given by $b_c e^{-||x_a - x_b||/L}$ where L and b_c are two constants. In our experiments, we fix $L = 0.3 \max_{u,v} ||x_u - x_v||$ and vary b_c to affect the sparsity of the graph. We refer to b_c as the connectivity parameter. We calculate the average efficiency ratio bound discussed in this work for Coloring Variant, Myopic Distributed Variant, GMS, FAMIX-MS and MWS. We observe that in the sparser graphs the coloring variant has a better efficiency ratio than GMS. FAMIX-MS is slightly above MWS, but as these graphs have a larger number of independent sets, the difference is not significant.

2) *Results in Other Networks:* In general networks, MWS and GMS can choose different schedules. In this case, for MWS, we transmit the earliest deadline packet of every link scheduled by MWS.

First, consider a sparse network with 5 links with interference graph G_I with edges $\{(l_1, l_2), (l_2, l_3), (l_2, l_4), (l_4, l_5)\}$. The traffic-fading process for $\{l_1, l_3, l_4\}$ is as in link 2 of Pattern B in Figure 1, and for links $\{l_2, l_5\}$ is as in link 1. We set an equal target delivery ratio p for all the links. Figure 8a shows the results. We can see that FAMIX-MS can support a significantly higher delivery ratio than the other algorithms. In this case, the optimal policy cannot achieve $p > 0.75$. This is because, for each link, for every two arrivals, the optimal policy cannot successfully transmit more than 1.5 packets on average.

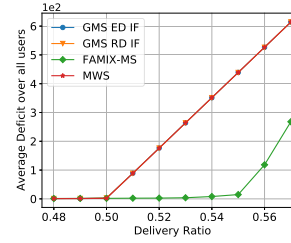
Next, consider a simple path network with links $\mathcal{K} = \{l_1, l_2, l_3\}$ with the interference graph G_I with edges $\{(l_1, l_2), (l_1, l_3)\}$. The traffic-fading process is indicated in Pattern A from Figure 1. Then GMS becomes unstable for $\mathbf{p} = (\frac{1}{3} + \epsilon, \frac{1}{2} + \epsilon, \frac{1}{2} + \epsilon)$, for any $\epsilon > 0$, whereas the optimal policy can achieve $\mathbf{p} = (\frac{11}{12}, \frac{5}{6}, \frac{5}{6})$. This is because



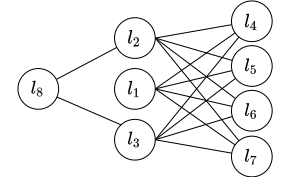
(a) A 5-link sparse network.

(b) 3-link line network.

Fig. 8. Performance comparison in two non-collocated networks, under Markovian traffic-fading.



(a) Evaluation results



(b) Interference graph

Fig. 9. (a) Performance comparison in an 8-link network with the interference graph shown in (b), with the mix of i.i.d. and Markovian traffic-fading processes.

in this traffic scenario, there is a way to attempt to transmit every packet exactly once, and the only packet losses will be due to fading. Then l_1 will have 2.75 packets successfully transmitted in every 3 arrivals, hence a delivery ratio of $\frac{2.75}{3} = \frac{11}{12}$. Similarly l_2 and l_3 will have a delivery ratio of $\frac{1+2/3}{2} = \frac{5}{6}$. This example illustrates the existence of simple examples with non-equal delivery ratio requirements where GMS performs poorly. Similar examples can be constructed for more links. To further study cases with non-equal delivery ratio, we consider the same network and set a vector of delivery ratio requirements $\mathbf{p} = (0.9, 0.8, 0.8)p$ for varying values of p which we refer to as *QoS intensity*. The result is illustrated in Figure 8. As we can see, MWS significantly outperforms GMS in this network, in contrast to the earlier 5-link sparse network. Further, as seen in Figure 8, FAMIX-MS supports larger values of QoS intensity p compared to MWS and GMS.

Finally, we consider a network with 8 links $\{l_1, \dots, l_8\}$, with interference graph shown in Figure 9. Links l_1, l_2 have the traffic-fading of link 1 in Pattern B of Figure 1b, and links l_4, l_5, l_8 have the traffic-fading of link 2 from the same pattern. In addition, we add two different i.i.d. traffic sources. The first traffic source has one packet arrival with probability 0.3 with deadline 1, at links l_3 and l_6 . The other source has 7 packet arrivals with probability 0.1, each with deadline 10, for link l_7 . Further, l_3 has the same channel success probabilities as l_1 , and l_6, l_7 have the same channel success probabilities as l_4 (with the success probabilities in the slots that are not specified assumed to be 1). The results are shown in Figure 9a FAMIX-MS retains a considerable advantage even with the presence of the i.i.d. sources. We observed

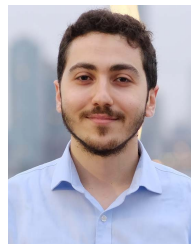
similar behavior when varying the packet arrival probability for the i.i.d. sources. This indicates that the results are robust to traffic perturbations and under different interference graphs.

VI. CONCLUSION

We considered scheduling of real-time traffic over fading channels, where traffic (arrival and deadline) and links' reliability evolves as an *unknown* finite-state Markov chain. We provided a conversion result that shows classical non-real-time scheduling algorithms like MWS and GMS can be ported to this setting and characterized their efficiency ratio. We then extended the randomized algorithms from [21] to fading channels. We further proposed low-complexity and myopic distributed randomized algorithms and characterized their efficiency ratio. Improving the bounds in this paper and investigating more efficient low-complexity and distributed randomized algorithms could be an interesting future research.

REFERENCES

- [1] C. Tsanikidis and J. Ghaderi, "Randomized scheduling of real-time traffic in wireless networks over fading channels," in *Proc. IEEE Conf. Comput. Commun.*, May 2021, pp. 1–10.
- [2] L. Tassioulas and A. Ephremides, "Stability properties of constrained queueing systems and scheduling policies for maximum throughput in multihop radio networks," in *Proc. 29th IEEE Conf. Decis. Control*, Dec. 1990, pp. 2130–2132.
- [3] C. Joo, X. Lin, and N. B. Shroff, "Understanding the capacity region of the greedy maximal scheduling algorithm in multihop wireless networks," *IEEE/ACM Trans. Netw.*, vol. 17, no. 4, pp. 1132–1145, Aug. 2009.
- [4] A. Dimakis and J. Walrand, "Sufficient conditions for stability of longest-queue-first scheduling: Second-order properties using fluid limits," *Adv. Appl. Probab.*, vol. 38, no. 2, pp. 505–521, 2006.
- [5] J. Ghaderi and R. Srikant, "On the design of efficient CSMA algorithms for wireless networks," in *Proc. 49th IEEE Conf. Decis. Control (CDC)*, Dec. 2010, pp. 954–959.
- [6] J. Ni, B. Tan, and R. Srikant, "Q-CSMA: Queue-length-based CSMA/CA algorithms for achieving maximum throughput and low delay in wireless networks," *IEEE/ACM Trans. Netw.*, vol. 20, no. 3, pp. 825–836, Jun. 2012.
- [7] D. Shah and J. Shin, "Delay optimal queue-based CSMA," *ACM SIGMETRICS Perform. Eval. Rev.*, vol. 38, no. 1, pp. 373–374, Jun. 2010.
- [8] C. Lu et al., "Real-time wireless sensor-actuator networks for industrial cyber-physical systems," *Proc. IEEE*, vol. 104, no. 5, pp. 1013–1024, May 2015.
- [9] J. Song et al., "WirelessHART: Applying wireless technology in real-time industrial process control," in *Proc. IEEE Real-Time Embedded Technol. Appl. Symp.*, Apr. 2008, pp. 377–386.
- [10] J. Gubbi, R. Buyya, S. Marusic, and M. Palaniswami, "Internet of Things (IoT): A vision, architectural elements, and future directions," *Future Generat. Comput. Syst.*, vol. 29, no. 7, pp. 1645–1660, 2013.
- [11] I.-H. Hou, V. Borkar, and P. R. Kumar, "A theory of QoS for wireless," in *Proc. 28th Conf. Comput. Commun.*, Rio de Janeiro, Brazil, Apr. 2009, pp. 486–494.
- [12] I.-H. Hou and P. R. Kumar, "Admission control and scheduling for QoS guarantees for variable-bit-rate applications on wireless channels," in *Proc. 10th ACM Int. Symp. Mobile Ad Hoc Netw. Comput.*, New Orleans, LO, USA, 2009, pp. 175–184.
- [13] I.-H. Hou and P. R. Kumar, "Scheduling heterogeneous real-time traffic over fading wireless channels," in *Proc. IEEE INFOCOM*, San Diego, CA, USA, Mar. 2010, pp. 2606–2614.
- [14] J. J. Jaramillo and R. Srikant, "Optimal scheduling for fair resource allocation in ad hoc networks with elastic and inelastic traffic," in *Proc. IEEE INFOCOM*, San Diego, CA, USA, Mar. 2010, pp. 2231–2239.
- [15] B. Li and A. Eryilmaz, "Optimal distributed scheduling under time-varying conditions: A fast-CSMA algorithm with applications," *IEEE Trans. Wireless Commun.*, vol. 12, no. 7, pp. 3278–3288, Jul. 2013.
- [16] R. Singh and P. R. Kumar, "Throughput optimal decentralized scheduling of multihop networks with end-to-end deadline constraints: Unreliable links," *IEEE Trans. Autom. Control*, vol. 64, no. 1, pp. 127–142, Jan. 2018.
- [17] X. Kang, W. Wang, J. J. Jaramillo, and L. Ying, "On the performance of largest-deficit-first for scheduling real-time traffic in wireless networks," *IEEE/ACM Trans. Netw.*, vol. 24, no. 1, pp. 72–84, Feb. 2016.
- [18] X. Kang, I.-H. Hou, and L. Ying, "On the capacity requirement of largest-deficit-first for scheduling real-time traffic in wireless networks," in *Proc. 16th ACM Int. Symp. Mobile Ad Hoc Netw. Comput.*, Jun. 2015, pp. 217–226.
- [19] L. Deng, C.-C. Wang, M. Chen, and S. Zhao, "Timely wireless flows with general traffic patterns: Capacity region and scheduling algorithms," *IEEE/ACM Trans. Netw.*, vol. 25, no. 6, pp. 3473–3486, Dec. 2017.
- [20] R. Li and A. Eryilmaz, "Scheduling for end-to-end deadline-constrained traffic with reliability requirements in multihop networks," *IEEE/ACM Trans. Netw.*, vol. 20, no. 5, pp. 1649–1662, Oct. 2012.
- [21] C. Tsanikidis and J. Ghaderi, "On the power of randomization for scheduling real-time traffic in wireless networks," in *Proc. IEEE Conf. Comput. Commun.*, Jul. 2020, pp. 59–68.
- [22] C. Tsanikidis and J. Ghaderi, "On the power of randomization for scheduling real-time traffic in wireless networks," *IEEE/ACM Trans. Netw.*, vol. 29, no. 4, pp. 1703–1716, Aug. 2021.
- [23] I.-H. Hou, "Scheduling heterogeneous real-time traffic over fading wireless channels," *IEEE/ACM Trans. Netw.*, vol. 22, no. 5, pp. 1631–1644, Oct. 2014.
- [24] E. B. Dynkin, *Theory of Markov Processes*. Chelmsford, MA, USA: Courier Corporation, 2012.
- [25] M. J. Neely, "Stability and probability 1 convergence for queueing networks via Lyapunov optimization," *J. Appl. Math.*, vol. 2012, pp. 1–35, Dec. 2012.
- [26] C. Joo, X. Lin, J. Ryu, and N. B. Shroff, "Distributed greedy approximation to maximum weighted independent set for scheduling with fading channels," *IEEE/ACM Trans. Netw.*, vol. 24, no. 3, pp. 1476–1488, Jun. 2016.
- [27] N. Robertson, D. Sanders, P. Seymour, and R. Thomas, "A new proof of the four-colour theorem," *Electron. Res. Announcements Amer. Math. Soc.*, vol. 2, no. 1, pp. 17–25, 1996.
- [28] S. M. Ross, *Applied Probability Models With Optimization Applications*. Chelmsford, MA, USA: Courier Corporation, 2013.
- [29] F. Y. L. Chin, M. Chrobak, S. P. Y. Fung, W. Jawor, J. Sgall, and T. Tichý, "Online competitive algorithms for maximizing weighted throughput of unit jobs," *J. Discrete Algorithms*, vol. 4, no. 2, pp. 255–276, 2006.
- [30] L. Z. Je, "One to rule them all: A general randomized algorithm for buffer management with bounded delay," in *Proc. Eur. Symp. Algorithms*. Cham, Switzerland: Springer, 2011, pp. 239–250.
- [31] M. Bienkowski, M. Chrobak, and Ł. Jeż, "Randomized competitive algorithms for online buffer management in the adaptive adversary model," *Theor. Comput. Sci.*, vol. 412, no. 39, pp. 5121–5131, Sep. 2011.
- [32] Ł. Jeż, F. Li, J. Sethuraman, and C. Stein, "Online scheduling of packets with agreeable deadlines," *ACM Trans. Algorithms*, vol. 9, no. 1, pp. 1–11, Dec. 2012.



Christos Tsanikidis (Graduate Student Member, IEEE) received the B.Sc. degree from the National Technical University of Athens in 2018 and the master's degree from Columbia University in 2020, where he is currently pursuing the Ph.D. degree with the Department of Electrical Engineering. He is researching scheduling and algorithmic problems in the area of computer networks. He was a recipient of a Best Paper Award and a Student Travel Grant from IEEE INFOCOM 2020.



Javad Ghaderi (Senior Member, IEEE) is currently an Associate Professor in electrical engineering at Columbia University. His research interests include network algorithms, control, and optimization. He was a recipient of several awards, including Best Student Paper Finalist Award at the 2013 American Control Conference, Best Paper Award at ACM CoNEXT 2016, NSF CAREER Award in 2017, Best Paper Award at IEEE INFOCOM 2020, and Best Student Paper Award at IFIP Performance 2020.