

On Wireless Scheduling Algorithms for Minimizing the Queue-Overflow Probability

V. J. Venkataramanan and Xiaojun Lin, *Member, IEEE*

Abstract—In this paper, we are interested in wireless scheduling algorithms for the downlink of a single cell that can minimize the queue-overflow probability. Specifically, in a large-deviation setting, we are interested in algorithms that maximize the asymptotic decay rate of the queue-overflow probability, as the queue-overflow threshold approaches infinity. We first derive an upper bound on the decay rate of the queue-overflow probability over all scheduling policies. We then focus on a class of scheduling algorithms collectively referred to as the “ α -algorithms.” For a given $\alpha \geq 1$, the α -algorithm picks the user for service at each time that has the largest product of the transmission rate multiplied by the backlog raised to the power α . We show that when the overflow metric is appropriately modified, the minimum-cost-to-overflow under the α -algorithm can be achieved by a simple linear path, and it can be written as the solution of a vector-optimization problem. Using this structural property, we then show that when α approaches infinity, the α -algorithms asymptotically achieve the largest decay rate of the queue-overflow probability. Finally, this result enables us to design scheduling algorithms that are both close to optimal in terms of the asymptotic decay rate of the overflow probability and empirically shown to maintain small queue-overflow probabilities over queue-length ranges of practical interest.

Index Terms—Asymptotically optimal algorithms, cellular system, large deviations, queue-overflow probability, wireless scheduling.

I. INTRODUCTION

LINK scheduling is an important functionality in wireless networks due to both the shared nature of the wireless medium and the variations of the wireless channel over time. In the past, it has been demonstrated that by carefully choosing the scheduling decision based on the channel state and/or the demand of the users, the system performance can be substantially improved (see, e.g., the references in [2]). Most studies of scheduling algorithms have focused on optimizing the long-term average throughput of the users or, in other words, stability. Consider the downlink of a single cell in a cellular network. The base-station transmits to N users. There is a queue Q_i associated with each user $i = 1, 2, \dots, N$. Due to interference, at any given time, the base-station can only serve the queue of one user. Hence, this system can be modeled as a single server

serving N queues. Assume that data for user i arrives at the base-station at a constant rate λ_i . Furthermore, assume a slotted model, and in each time-slot the wireless channel can be in one of M states. In each state $m = 1, 2, \dots, M$, if the base-station picks user i to serve, the corresponding service rate is F_m^i . Hence, at each time-slot, Q_i increases by λ_i , and if it is served and the channel is at state m , Q_i decreases by F_m^i . We assume that perfect channel information is available at the base-station. In a stability problem [3]–[5], the goal is to find algorithms for scheduling the transmissions such that the queues are stabilized at given offered loads. An important result along this direction is the development of the so-called “throughput-optimal” algorithms [3]. A scheduling algorithm is called *throughput-optimal* if, at any offered load under which any other algorithm can stabilize the system, this algorithm can stabilize the system as well. It is well known that the following class of scheduling algorithms are throughput-optimal [3]–[5]: For a given $\alpha \geq 1$, the base-station picks the user for service at each time that has the largest product of the transmission rate multiplied by the backlog raised to the power α . In other words, if the channel is in state m , the base-station chooses the user i with the largest $(Q_i)^\alpha F_m^i$. To emphasize the dependency on α , in the sequel we will refer to this class of throughput-optimal algorithms as α -algorithms.

While stability is an important first-order metric of success, for many delay-sensitive applications, it is far from sufficient. In this paper, we are interested in the probability of queue overflow, which is equivalent to the delay-violation probability under certain conditions. The question that we attempt to answer is the following: Is there an optimal algorithm in the sense that, at any given offered load, the algorithm can achieve the smallest probability that any queue overflows, i.e., the smallest value of $\mathbf{P}[\max_{1 \leq i \leq N} Q_i(T) \geq B]$, where B is the overflow threshold? Note that if we impose a quality-of-service (QoS) constraint on each user in the form of an upper bound on the queue-overflow probability, then the above optimality condition will also imply that the algorithm can support the largest set of offered loads subject to the QoS constraint.

Unfortunately, calculating the exact queue distribution is often mathematically intractable. In this paper, we use large-deviation theory [11], [12] and reformulate the QoS constraint in terms of the asymptotic decay rate of the queue-overflow probability as B approaches infinity. In other words, we are interested in finding scheduling algorithms that can achieve the smallest possible value of

$$\limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} \left[\max_{1 \leq i \leq N} Q_i(T) \geq B \right]. \quad (1)$$

Our main results are the following. We show that there exists an optimal decay rate I_{opt} such that for any scheduling algorithm

$$\liminf_{B \rightarrow \infty} \frac{1}{B} \log \left(\mathbf{P} \left[\max_{1 \leq i \leq N} Q_i(0) \geq B \right] \right) \geq -I_{\text{opt}}.$$

Manuscript received September 22, 2008; revised April 29, 2009; approved by IEEE/ACM TRANSACTIONS ON NETWORKING Editor J. Walrand. First published January 19, 2010; current version published June 16, 2010. This work was supported in part by NSF grants CNS-0626703, CNS-0643145, CNS-0721484, and CCF-0635202 and a grant from Purdue Research Foundation. An earlier version of this paper has appeared in the 45th Annual Allerton Conference on Communication, Control, and Computing, September 2007 [1].

The authors are with Center for Wireless Systems and Applications (CWSA) and the School of Electrical and Computer Engineering, Purdue University, West Lafayette, IN 47906 USA (e-mail: vvenkat@ecn.purdue.edu; linx@ecn.purdue.edu).

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.2009.2037896

Furthermore, for α -algorithms

$$\lim_{\alpha \rightarrow \infty} \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P}_0^\alpha \left[\max_{1 \leq i \leq N} Q_i(T) \geq B \right] \leq -I_{\text{opt}}$$

where $\mathbf{P}_0^\alpha$ is the probability measure for the α -algorithm. Hence, when α approaches infinity, the α -algorithms asymptotically achieve the largest decay rate of the queue-overflow probability.

For the above problem, it is natural to use the large-deviation theory¹ because the overflow probability that we are interested in is typically very small [11], [12]. Large-deviation theory has been successfully applied to wireline networks (see, e.g., [13]–[19]) and to wireless scheduling algorithms that only use the channel state to make the scheduling decisions [20]–[22]. However, when applied to wireless scheduling algorithms that also use the queue-length to make scheduling decisions (e.g., the α -algorithms), this approach encounters a significant amount of technical difficulty. Specifically, in order to apply the large-deviation theory to queue-length-based scheduling algorithms, one has to use sample-path large deviation and formulate the problem as a multidimensional calculus-of-variations (CoV) problem for finding the “most likely path to overflow.” The decay rate of the queue-overflow probability then corresponds to the cost of this path, which is referred to as the “minimum cost to overflow.” Unfortunately, for many queue-length-based scheduling algorithms of interest, this multidimensional calculus-of-variations problem is very difficult to solve. In the literature, only some restricted cases have been solved: Either restricted problem structures are assumed (e.g., symmetric users and ON–OFF channels [23]), or the size of the system is very small (only two users) [24]. In this paper, to circumvent the difficulty of the multidimensional CoV problem, we apply a novel technique introduced in [26]. Specifically, we use a Lyapunov function to map the multidimensional CoV problem to a one-dimensional problem, which allows us to bound the minimum cost to overflow by solutions of simple vector optimization problems. This technique is of independent interest and may be useful for analyzing other queue-length-based scheduling algorithms as well.

In a recent work [25],² the author shows that the “exponential rule” can maximize the decay rate of the queue-overflow probability over all scheduling policies. The results in this paper are comparable but different. The advantage of working with the α -algorithms instead of the exponential rule is that the α -algorithms are scale-invariant (i.e., the outcome of the scheduling decision does not change if all queue lengths are multiplied by a common factor). Hence, we can use the standard sample-path large-deviation principle (LDP) instead of the refined LDP used in [25] that is more technically involved. In addition, our results highlight the role that the exponent α plays in determining the asymptotic decay rate. Finally, using the insight of our main result, we design a scheduling algorithm that is both close to optimal in terms of the asymptotic decay rate of the overflow probability and empirically shown to maintain small queue-overflow probabilities over queue-length ranges of practical interest.

¹Alternatively, one may use other asymptotic techniques such as heavy-traffic limits [6]–[8] or focus on order-optimal bounds on the expected queue length/packet delay [9], [10].

²Note that this work is published after our preliminary results reported in [1].

II. THE SYSTEM MODEL AND ASSUMPTIONS

We consider the downlink of a single cell in which a base-station serves N users. We assume a slotted system, and we assume that the state of the channel at each time slot is chosen i.i.d from one of M possible states. Let $C(t)$ denote the state of the channel at time $t = 1, 2, \dots$, and let $p_m = \mathbf{P}[C(t) = m]$, $m = 1, 2, \dots, M$. Let $\mathbf{p} = [p_1, \dots, p_M]$. We assume that the base-station can serve one user at a time. Let F_m^i denote the service rate for user i when it is picked for service and the channel state is m .

We assume that data for user i arrive as fluid at a constant rate λ_i . Let $\boldsymbol{\lambda} = [\lambda_1, \dots, \lambda_N]$. Let $Q_i(t)$ denote the backlog of user i at time t , and let $\mathbf{Q}(t) = [Q_1(t), \dots, Q_N(t)]$. In general, the decision of picking which user to serve is a function of the global backlog $\mathbf{Q}(t)$ and the channel state $C(t)$. Let $U(t)$ denote the index of the user picked for service at time t . The evolution of the backlog for each user i is then governed by

$$Q_i(t+1) = \left[Q_i(t) + \lambda_i - \sum_{m=1}^M \mathbf{1}_{\{C(t)=m, U(t)=i\}} F_m^i \right]^+ \quad (2)$$

where $[\cdot]^+$ denotes the projection to $[0, +\infty)$. Note that $\sum_{m=1}^M \sum_{i=1}^N \mathbf{1}_{\{C(t)=m, U(t)=i\}} = 1$ since only one user can be served at a time.

A particular class of scheduling algorithms that we will focus on are collectively referred to as the “ α -algorithms,” where α is a parameter that takes values from the set of natural numbers. Given α , the behavior of the algorithm is as follows. When the backlog of the users is $\mathbf{Q}(t)$ and the state of the channel is $C(t) = m$, the algorithm chooses to serve the user i for which the product $Q_i^\alpha(t) F_m^i$ is the largest. If there are several users that achieve the largest $Q_i^\alpha(t) F_m^i$ together, one of them is chosen arbitrarily. It is well known that this class of algorithms are throughput-optimal, i.e., they can stabilize the system at the largest set of offer loads $\boldsymbol{\lambda}$ [3]–[5]. Note that although these algorithms do not explicitly keep track of past history; they do so implicitly by their dependence on $\mathbf{Q}(t)$. Hence, they are able to stabilize the system without explicit knowledge of the operating conditions such as arrival rate and channel probabilities.

Consider the system when it is operated at a given offered load $\boldsymbol{\lambda}$ and is stable under a given scheduling algorithm. Specifically, we assume that there is a positive number $\epsilon > 0$ such that $\boldsymbol{\lambda}(1 + \epsilon)$ is in the capacity region of the system. This implies (refer to [3]) that there exists $[\hat{\gamma}_m^i] \geq 0$ such that $\sum_{i=1}^N \hat{\gamma}_m^i = 1$ for all $m = 1, \dots, M$ and

$$\lambda_i(1 + \epsilon) \leq \sum_{m=1}^M p_m \hat{\gamma}_m^i F_m^i \text{ for all } i = 1, \dots, N. \quad (3)$$

In this paper, we are interested in the probability that the largest backlog exceeds a certain threshold B , i.e.,

$$\mathbf{P} \left[\max_{1 \leq i \leq N} Q_i(0) \geq B \right]. \quad (4)$$

Note that the probability in (4) is equivalent to a delay-violation probability when the arrival rates λ_i are constant because the two types of events are related by (see [23] and [27]) $\mathbf{P}[\text{Delay at link } i \geq d_i] = \mathbf{P}[Q_i(0) \geq \lambda_i d_i]$. The focus of this paper is in scheduling algorithms that minimize (4).

The problem of calculating the exact probability $\mathbf{P}[\max_{1 \leq i \leq N} Q_i(0) \geq B]$ is often mathematically intractable. In this work, we are interested in using large-deviation theory to compute estimates of this probability. Specifically, we will use the following limits:

$$I_0(\lambda) \triangleq -\liminf_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} \left[\max_{1 \leq i \leq N} Q_i(0) \geq B \right] \quad (5)$$

$$J_0(\lambda) \triangleq -\limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} \left[\max_{1 \leq i \leq N} Q_i(0) \geq B \right]. \quad (6)$$

In essence, $I_0(\lambda)$ and $J_0(\lambda)$ are upper and lower bounds, respectively, of the decay rate of (4) as the overflow threshold B approaches infinity. In the following sections, we will show that no scheduling algorithm can have a decay rate larger than a certain value I_{opt} (defined in Section IV), i.e., $I_0(\lambda) \leq I_{\text{opt}}$. Then, we will show that the α -algorithms asymptotically achieve the decay rate I_{opt} . In other words, for the α -algorithms, $J_0(\lambda)$ approaches I_{opt} , as $\alpha \rightarrow \infty$.

III. PRELIMINARIES

Since the channel states are i.i.d. in time, the following sample-path LDP holds for the channel-state process. Specifically, we define the empirical measure process $\mathbf{S}(t) = [S_m(t), m = 1, \dots, M]$ as follows:

$$S_m(t) = \int_0^t \mathbf{1}_{\{C(\lfloor \tau \rfloor) = m\}} d\tau$$

where $\lfloor \tau \rfloor$ represents the largest integer no greater than τ . Then, for any nonnegative integer B , define the scaled channel-rate process

$$\mathbf{s}^B(t) = \frac{\mathbf{S}(Bt)}{B}. \quad (7)$$

It is easy to see that $\mathbf{s}^B(\cdot)$ is Lipschitz-continuous, and hence its derivative exists almost everywhere. For any given $T > 0$, let $\tilde{\Psi}_T$ denote the space of mappings from $[0, T]$ to $\mathbb{R}^M$, equipped with the essential supremum norm [12, pp. 176, 352]. Let $\mathcal{P}_M$ denote the set of probability vectors of dimension M , i.e., $\phi = [\phi_m, m = 1, \dots, M] \in \mathcal{P}_M$ implies that $\phi \geq 0$ and $\sum_{m=1}^M \phi_m = 1$. For any $\phi \in \mathcal{P}_M$, define³

$$H(\phi || \mathbf{p}) = \sum_{m=1}^M \phi_m \log \frac{\phi_m}{p_m}$$

with the convention that $0 \log 0 = 0$. Then, as $B \rightarrow \infty$, it is well known that the sequence of scaled channel-rate processes $\mathbf{s}^B(\cdot)$ on the interval $[0, T]$ satisfies a sample-path LDP with good rate function [12, Mogulskii's Theorem (Theorem 5.1.2), p. 176]

$$I_s^T(\mathbf{s}(\cdot)) = \begin{cases} \int_0^T H(\dot{\mathbf{s}}(t) || \mathbf{p}) dt, & \text{if } \mathbf{s}(\cdot) \in \mathcal{AC} \\ \infty, & \text{otherwise} \end{cases}$$

³This is commonly known as the relative entropy between ϕ and $\mathbf{p}$.

where $\mathcal{AC}$ denotes the set of absolute continuous functions in $\tilde{\Psi}_T$. This LDP means that for any set $\tilde{\Gamma}$ of trajectories in $\tilde{\Psi}_T$, the following inequality holds:

$$\begin{aligned} -\inf_{\mathbf{s}(\cdot) \in \tilde{\Gamma}^\circ} I_s^T(\mathbf{s}(\cdot)) &\leq \liminf_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \tilde{\Gamma}] \\ &\leq \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \tilde{\Gamma}] \\ &\leq -\inf_{\mathbf{s}(\cdot) \in \tilde{\Gamma}} I_s^T(\mathbf{s}(\cdot)) \end{aligned} \quad (8)$$

where $\tilde{\Gamma}^\circ$ and $\tilde{\Gamma}$ denote the interior and closure, respectively, of the set $\tilde{\Gamma}$. In addition, if $\tilde{\Gamma}$ is a *continuity set* [12, p. 5], the two bounds meet, and we then have

$$\lim_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \tilde{\Gamma}] = -\inf_{\mathbf{s}(\cdot) \in \tilde{\Gamma}} I_s^T(\mathbf{s}(\cdot)). \quad (9)$$

Hence, the large-deviation rate function $I_s^T(\mathbf{s}(\cdot))$ characterizes how “rarely” the trajectory $\mathbf{s}(\cdot)$ occurs.

Using a similar scaling as $\mathbf{s}^B(\cdot)$, define the scaled backlog process

$$\mathbf{q}^B(t) = \frac{\mathbf{Q}(Bt)}{B}, \text{ for } t = 0, \frac{1}{B}, \frac{2}{B}, \dots \quad (10)$$

and by linear interpolation otherwise. Hence, for each $\mathbf{s}^B(\cdot)$ and a given initial condition $\mathbf{q}^B(0)$, we can use (2) to determine the corresponding $\mathbf{q}^B(\cdot)$. As $B \rightarrow \infty$, we will have a sequence of $\mathbf{s}^B(\cdot)$ and $\mathbf{q}^B(\cdot)$. It is easy to see that both $\mathbf{s}^B(\cdot)$ and $\mathbf{q}^B(\cdot)$ are Lipschitz-continuous. Hence, there must exist a subsequence that converges uniformly over the interval $[0, T]$. We use $(\mathbf{s}(\cdot), \mathbf{q}(\cdot))$ to denote such a limit, and we refer to it as a *fluid sample path*.

In essence, the goal of the rest of the paper is to use the known sample-path LDP of $\mathbf{s}^B(\cdot)$ to characterize that of $\mathbf{q}^B(\cdot)$ and that of the queue-overflow probability. In [1], we assume that a sample-path LDP also holds for $\mathbf{q}^B(\cdot)$. Unfortunately, such an assumption appears to be difficult to verify. Instead, in this paper, we will use a different approach to establish the desirable results.

IV. AN UPPER BOUND ON THE DECAY RATE OF THE OVERFLOW PROBABILITY

In this section, we first present an upper bound I_{opt} on $I_0(\lambda)$ [defined in (5)] under a given offered load λ . This value I_{opt} bounds from above the decay rate for the overflow probability of the stationary backlog process $\mathbf{Q}(t)$ over all scheduling policies. For every probability vector $\phi \in \mathcal{P}_M$, define the following optimization problem:

$$\begin{aligned} w(\phi) &\triangleq \inf_{[\tilde{\gamma}_m^i]} \max_{1 \leq i \leq N} \left[\lambda_i - \sum_{m=1}^M \phi_m \tilde{\gamma}_m^i F_m^i \right]^+ \\ \text{subject to } &\sum_{i=1}^N \tilde{\gamma}_m^i = 1 \text{ for all } m = 1, \dots, M \\ &\tilde{\gamma}_m^i \geq 0 \text{ for all } i = 1, \dots, N \text{ and } m. \end{aligned}$$

Here, $\tilde{\gamma}_m^i$ can be interpreted as some long-term fraction of time that user i is served when the channel state is m . Hence, if the channel-rate process is given by $\mathbf{s}(t) = \phi t$, then $[\lambda_i - \sum_{m=1}^M \phi_m \tilde{\gamma}_m^i F_m^i]^+$ denotes the long-term growth rate of the backlog of user i . Furthermore, if all queues start empty, then $w(\phi)$ is the minimum rate of growth of the backlog of the largest queue.

Next, define I_{opt} as

$$I_{\text{opt}} \triangleq \inf_{\{\phi \in \mathcal{P}_M | w(\phi) > 0\}} \frac{H(\phi \| \mathbf{p})}{w(\phi)}. \quad (11)$$

Given a fixed offered load λ , assume that the backlog process $Q(\cdot)$ under a given scheduling policy is stationary and ergodic. We will show the following result.⁴

Proposition 1: Under any scheduling policy, the following holds:

$$\liminf_{B \rightarrow \infty} \frac{1}{B} \mathbf{P} \left(\max_{1 \leq i \leq N} Q_i(0) \geq B \right) \geq -I_{\text{opt}}. \quad (12)$$

In other words, I_{opt} is an upper bound for the decay rate of the overflow probability over all scheduling policies. This upper bound, although in a different form, is equal to the one derived in [25].

Toward this end, we first show that the function $w(\cdot)$ provides a lower bound on the backlog of the largest queue, as proved in the following lemma.

Lemma 2: For any $\epsilon > 0$, there exists $B_0 > 0$ such that for all $B \geq B_0$ and all scaled channel-rate process $\mathbf{s}^B(\cdot)$ (with $\mathbf{s}^B(0) = 0$), the following holds

$$\max_{1 \leq i \leq N} q_i^B(T) \geq T w \left(\frac{\mathbf{s}^B(T)}{T} \right) - \epsilon, \quad \text{for all } T > 0.$$

Proof: Note that the queue backlog process is related to the channel-state process by (2). Take the scaling in (7) and (10). Then, given $\mathbf{s}^B(\cdot)$, at any time t such that Bt is an integer, we must have

$$q_i^B(t) \geq \left[\lambda_i t - \int_0^t \sum_{m=1}^M \dot{s}_m^B(\tau) \mathbf{1}_{\{U(\lfloor B\tau \rfloor) = i\}} F_m^i d\tau \right]^+.$$

For any $T > 0$, there must exist a value of t such that Bt is an integer and $|t - T| \leq 1/B$. Hence, for any $\epsilon > 0$, there must exist $B_0 > 0$ such that for all $B \geq B_0$

$$q_i^B(T) \geq \left[\lambda_i T - \int_0^T \sum_{m=1}^M \dot{s}_m^B(\tau) \mathbf{1}_{\{U(\lfloor B\tau \rfloor) = i\}} F_m^i d\tau \right]^+ - \epsilon.$$

Let $\phi_m = s_m^B(T)/T$, $m = 1, \dots, M$. If $\phi_m > 0$, let

$$\tilde{\gamma}_m^i = \frac{1}{s_m^B(T)} \int_0^T \dot{s}_m^B(\tau) \mathbf{1}_{\{U(\lfloor B\tau \rfloor) = i\}} d\tau.$$

⁴Note that Proposition 1 also holds trivially if the system is unstable.

Otherwise, let $\tilde{\gamma}_m^1 = 1$ and $\tilde{\gamma}_m^i = 0$ for $i \geq 2$. We then have $q_i^B(T) \geq T \left[\lambda_i - \sum_{m=1}^M \phi_m \tilde{\gamma}_m^i F_m^i \right]^+ - \epsilon$. Taking the maximum over all $i = 1, \dots, N$, we have

$$\max_{1 \leq i \leq N} q_i^B(T) \geq T \max_{1 \leq i \leq N} \left[\lambda_i - \sum_{m=1}^M \phi_m \tilde{\gamma}_m^i F_m^i \right]^+ - \epsilon.$$

Finally, since $\sum_{i=1}^N \mathbf{1}_{\{U(\lfloor B\tau \rfloor) = i\}} = 1$, $[\tilde{\gamma}_m^i]$ is a feasible point for the optimization problem $w(\mathbf{s}^B(T)/T)$. Thus, we obtain the lower bound that $\max_{1 \leq i \leq N} q_i^B(T) \geq T w(\mathbf{s}^B(T)/T) - \epsilon$. ■

In addition, it is easy to show that the value of $w(\phi)$ is continuous with respect to ϕ as stated below in Lemma 3. Let $\|\phi\|$ denote the Euclidean norm of ϕ .

Lemma 3: Let ϕ^1 and ϕ^2 be vectors from $\mathcal{P}_M$. The optimization problem $w(\cdot)$ is continuous in the sense that for any $\epsilon > 0$ and $\|\phi^1 - \phi^2\| < \epsilon$, the following holds:

$$|w(\phi^1) - w(\phi^2)| \leq \epsilon \sum_{i=1}^N \sum_{m=1}^M F_m^i.$$

The intuition behind Lemma 3 comes from the fact that the function $\max_{1 \leq i \leq N} [\lambda_i - \sum_{m=1}^M \phi_m \tilde{\gamma}_m^i F_m^i]^+$ is continuous in ϕ for any $[\tilde{\gamma}_m^i]$. The detailed proof is provided in our technical report [28]. We can now prove Proposition 1.

Proof (of Proposition 1): For any $\delta > 0$, we can find ϕ_δ from $\{\phi \in \mathcal{P}_M | w(\phi) > 0\}$ such that $(H(\phi_\delta \| \mathbf{p})/w(\phi_\delta)) < I_{\text{opt}} + \delta$. Define $\mathbf{s}_\delta(t) \triangleq t\phi_\delta$ for $t \geq 0$. Let ϵ be some positive number, and let $T = (1 + \epsilon + \epsilon \sum_{m=1}^M \sum_{i=1}^N F_m^i)/w(\phi_\delta)$. Let $B_T(\mathbf{s}_\delta(\cdot), \epsilon)$ be the set of functions in the space Ψ^T such that $\sup_{t \in [0, T]} \|\mathbf{s}(t) - \mathbf{s}_\delta(t)\| < \epsilon$. Therefore, for any B , $\mathbf{s}^B(\cdot) \in B_T(\mathbf{s}_\delta(\cdot), \epsilon)$ implies $\|\mathbf{s}^B(T)/T - \phi_\delta\| < (\epsilon/T)$. By Lemma 3, this in turn implies that

$$T w \left(\frac{\mathbf{s}^B(T)}{T} \right) \geq T w(\phi_\delta) - \epsilon \sum_{m=1}^M \sum_{i=1}^N F_m^i. \quad (13)$$

Now, using Lemma 2, for all $B > B_0$ and $\mathbf{s}^B(\cdot)$, we have $\max_{1 \leq i \leq N} q_i^B(T) \geq T w(\mathbf{s}^B(T)/T) - \epsilon$. Hence, by (13) we conclude that, for all $B > B_0$ and $\mathbf{s}^B(\cdot) \in B_T(\mathbf{s}_\delta(\cdot), \epsilon)$, we have

$$\max_{1 \leq i \leq N} q_i^B(T) \geq T w(\phi_\delta) - \epsilon - \epsilon \sum_{m=1}^M \sum_{i=1}^N F_m^i = 1 \quad (14)$$

where equality holds by the definition of T . Therefore

$$\begin{aligned} \mathbf{P} \left(\max_{1 \leq i \leq N} Q_i(0) \geq B \right) &= \mathbf{P} \left(\max_{1 \leq i \leq N} Q_i(BT) \geq B \right) \\ &= \mathbf{P} \left(\max_{1 \leq i \leq N} q_i^B(T) \geq 1 \right) \\ &\geq \mathbf{P} \left(\mathbf{s}^B(\cdot) \in B_T(\mathbf{s}_\delta(\cdot), \epsilon) \right). \end{aligned}$$

By the LDP for $\mathbf{s}^B(\cdot)$ [see Inequality (8)], we then have

$$\begin{aligned} & \liminf_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} \left(\max_{1 \leq i \leq N} Q_i(0) \geq B \right) \\ & \geq \liminf_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in B_T(\mathbf{s}_\delta(\cdot), \epsilon)] \\ & \geq - \inf_{\mathbf{s}(\cdot) \in B_T(\mathbf{s}_\delta(\cdot), \epsilon)} \int_0^T H(\dot{\mathbf{s}}(t) \| \mathbf{p}) dt \\ & \geq - \int_0^T H(\dot{\mathbf{s}}_\delta(t) \| \mathbf{p}) dt = -TH(\phi_\delta \| \mathbf{p}) \\ & \geq - \left(1 + \epsilon + \epsilon \sum_{m=1}^M \sum_{i=1}^N F_m^i \right) (I_{\text{opt}} + \delta). \end{aligned}$$

Since δ and ϵ can be arbitrarily small, we conclude that

$$\liminf_{B \rightarrow \infty} \frac{1}{B} \log \left(\mathbf{P} \left(\max_{1 \leq i \leq N} Q_i(0) \geq B \right) \right) \geq -I_{\text{opt}}.$$

■

V. A LOWER BOUND ON THE DECAY RATE OF THE OVERFLOW PROBABILITY FOR α -ALGORITHMS

In this section, we will use the following modified queue-overflow event: $\{V_\alpha(\mathbf{q}^B(t)) \geq 1\}$, where $V_\alpha(\mathbf{q}) \triangleq (\sum_{i=1}^N (q_i)^{\alpha+1})^{1/(\alpha+1)}$. Note that this overflow event is different from the queue-overflow event $\{\max_{1 \leq i \leq N} q_i^B(t) \geq 1\}$ that is used in earlier sections. It turns out that computing the large-deviation decay rate for $\mathbf{P}\{\max_{1 \leq i \leq N} q_i^B(t) \geq 1\}$ requires solving a CoV problem that is very difficult. The reason to use the modified overflow metric $V_\alpha(\mathbf{q}^B(t))$ is that the corresponding decay rate is much easier to compute and $V_\alpha(\mathbf{q}^B(t))$ approximates the function $\max_{1 \leq i \leq N} q_i^B(t)$ when α is large. To see this, note that as $\alpha \rightarrow \infty$, the difference between $V_\alpha(\mathbf{q}^B(t))$ and $\max_{1 \leq i \leq N} q_i^B(t)$ decreases to 0. Furthermore, the function $V_\alpha(\mathbf{Q})$ is a Lyapunov function for the α -algorithm. Hence, the theory developed in [26] applies and enables us to provide analytical results for this modified overflow metric. On the other hand, even though $\max_{1 \leq i \leq N} Q_i$ may be viewed as a Lyapunov function for some throughput-optimal algorithm, e.g., the exponential rule [25], the algorithm is typically not scale-invariant. Hence, it appears to be difficult to apply the theory of [26] directly on $\max_{1 \leq i \leq N} Q_i$.

A. A General Lower Bound

We first provide a lower bound that relates the decay rate of the overflow probability to the “minimum cost to overflow” among all fluid sample paths. For ease of exposition, instead of considering the stationary system, we consider a system that starts at time 0 (although the results can also be extended to the stationary system, as we will comment later). Specifically, let $\mathbf{Q}(0) = 0$. Let $\mathbf{P}_0$ denote the probability measure conditioned on $\mathbf{Q}(0) = 0$. For any $T > 0$, let $\hat{\Gamma}_T$ denote the set of fluid sample paths $(\mathbf{s}(\cdot), \mathbf{q}(\cdot))$ on the interval $[0, T]$ such that $\mathbf{q}(0) = 0$ and $V_\alpha(\mathbf{q}(T)) \geq 1$. We then have the following lower bound,

which is comparable to [25, Theorem 7.1] although we do not need to use the refined LDP.

Proposition 4: Consider $\hat{\Gamma}_T$ as defined earlier. Then, the following holds:

$$\begin{aligned} & \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P}_0 [V_\alpha(\mathbf{q}^B(T)) \geq 1] \\ & \leq - \inf_{(\mathbf{s}(\cdot), \mathbf{q}(\cdot)) \in \hat{\Gamma}_T} \int_0^T H(\dot{\mathbf{s}}(t) \| \mathbf{p}) dt. \end{aligned} \quad (15)$$

Remark: The infimum on the right-hand side of (15) is often called the “minimum cost to overflow.” This result reflects the well-celebrated large-deviation philosophy that “rare events occur in the most likely way.” Specifically, Proposition 4 states that the probability of queue overflow is determined mostly by the smallest cost among all fluid sample paths that overflow. This fluid sample path is often referred to as the “most likely path to overflow.”

Proof: Fix $T > 0$. Recall that we have set $\mathbf{q}^B(0) = 0$ for all B . Let $\tilde{\Gamma}^B$ be the set of channel rate processes $\mathbf{s}^B(\cdot)$ such that the corresponding backlog process satisfies $V_\alpha(\mathbf{q}^B(T)) \geq 1$. For all $n \geq 1$, we have

$$\limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \tilde{\Gamma}^B] \quad (16)$$

$$\leq \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \cup_{B=n}^\infty \tilde{\Gamma}^B]. \quad (17)$$

By the LDP for $\mathbf{s}^B(\cdot)$ [see (8)], we have

$$\begin{aligned} & \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \cup_{B=n}^\infty \tilde{\Gamma}^B] \\ & \leq - \inf_{\mathbf{s}(\cdot) \in \cup_{B=n}^\infty \tilde{\Gamma}^B} \int_0^T H(\dot{\mathbf{s}}(t) \| \mathbf{p}) dt. \end{aligned}$$

Note that the sequence of sets $\cup_{B=n}^\infty \tilde{\Gamma}^B$ is decreasing in n . We therefore have

$$\begin{aligned} & \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [\mathbf{s}^B(\cdot) \in \tilde{\Gamma}^B] \\ & \leq - \lim_{n \rightarrow \infty} \inf_{\mathbf{s}(\cdot) \in \cup_{B=n}^\infty \tilde{\Gamma}^B} \int_0^T H(\dot{\mathbf{s}}(t) \| \mathbf{p}) dt. \end{aligned} \quad (18)$$

It remains to show that the right-hand-side of (18) is no greater than that of (15). For each n , we can find $\mathbf{y}_n(\cdot) \in \cup_{B=n}^\infty \tilde{\Gamma}^B$ such that

$$\int_0^T H(\dot{\mathbf{y}}_n(t) \| \mathbf{p}) dt < \inf_{\mathbf{s}(\cdot) \in \cup_{B=n}^\infty \tilde{\Gamma}^B} \int_0^T H(\dot{\mathbf{s}}(t) \| \mathbf{p}) dt + \frac{1}{n}. \quad (19)$$

Since $\mathbf{y}_n(\cdot)$ is equicontinuous, we can find a subsequence that converges uniformly on $[0, T]$. For ease of exposition, we slightly abuse notation and denote this subsequence by $\mathbf{y}_n(\cdot)$. Let $\mathbf{y}^*(\cdot)$ denote its limit, i.e., $\lim_{n \rightarrow \infty} \mathbf{y}_n(\cdot) = \mathbf{y}^*(\cdot)$. Since the cost function $\int_0^T H(\cdot \| \mathbf{p}) dt$ is lower semi-continuous, we have

$$\liminf_{n \rightarrow \infty} \int_0^T H(\dot{\mathbf{y}}_n(t) \| \mathbf{p}) dt \geq \int_0^T H(\dot{\mathbf{y}}^*(t) \| \mathbf{p}) dt. \quad (20)$$

For each $\mathbf{y}_n(\cdot)$, since it belongs to the closure of $\cup_{B=n}^{\infty} \tilde{\Gamma}^B$, we can find a sequence $\mathbf{y}_{n,m}(\cdot) \in \cup_{B=n}^{\infty} \tilde{\Gamma}^B$, $m = 1, 2, \dots$, such that $\mathbf{y}_n(\cdot) = \lim_{m \rightarrow \infty} \mathbf{y}_{n,m}(\cdot)$. Then, from all $\mathbf{y}_{n,m}(\cdot)$, $n = 1, 2, \dots$, $m = 1, 2, \dots$, we can find another sequence $\mathbf{y}_{n,m_n}(\cdot)$, $n = 1, 2, \dots$, such that $\lim_{n \rightarrow \infty} \mathbf{y}_{n,m_n}(\cdot) = \mathbf{y}^*(\cdot)$. (For example, we can let $m_1 = 1$. Then, given m_n , we can choose m_{n+1} such that $\sup_{t \in [0, T]} \|\mathbf{y}_{n+1, m_{n+1}}(t) - \mathbf{y}_{n+1}(t)\| < (\sup_{t \in [0, T]} \|\mathbf{y}_{n, m_n}(t) - \mathbf{y}_n(t)\|/2)$.) For notational convenience, let $\hat{\mathbf{y}}_n(\cdot)$ denote the sequence $\mathbf{y}_{n, m_n}(\cdot)$ from here on.

For each n , let $\hat{\mathbf{q}}_n(\cdot)$ be the backlog process corresponding to the channel rate process $\hat{\mathbf{y}}_n(\cdot)$. By construction, $\hat{\mathbf{q}}_n(0) = 0$ and $V_\alpha(\hat{\mathbf{q}}_n(T)) \geq 1$ for all n . Since the backlog processes are equicontinuous, we can find a subsequence of $(\hat{\mathbf{y}}_n, \hat{\mathbf{q}}_n)$ such that this subsequence converges to $(\mathbf{y}^*(\cdot), \mathbf{q}^*(\cdot))$ uniformly over the interval $[0, T]$, where $\mathbf{q}^*(\cdot)$ satisfies $\mathbf{q}^*(0) = 0$ and $V_\alpha(\mathbf{q}^*(T)) \geq 1$. Therefore, $(\mathbf{y}^*(\cdot), \mathbf{q}^*(\cdot))$ is in $\tilde{\Gamma}_T$, and thus

$$\int_0^T H(\mathbf{y}^*(t) \|\mathbf{p}\|) dt \geq \inf_{(\mathbf{s}(\cdot), \mathbf{q}(\cdot)) \in \tilde{\Gamma}_T} \int_0^T H(\mathbf{s}(t) \|\mathbf{p}\|) dt.$$

Combining with (19) and (20), we conclude that

$$\begin{aligned} \lim_{n \rightarrow \infty} \inf_{\mathbf{s}(\cdot) \in \cup_{B=n}^{\infty} \tilde{\Gamma}^B} \int_0^T H(\mathbf{s}(t) \|\mathbf{p}\|) dt \\ \geq \liminf_{n \rightarrow \infty} \int_0^T H(\hat{\mathbf{y}}_n(t) \|\mathbf{p}\|) dt \\ \geq \inf_{(\mathbf{s}(\cdot), \mathbf{q}(\cdot)) \in \tilde{\Gamma}_T} \int_0^T H(\mathbf{s}(t) \|\mathbf{p}\|) dt. \end{aligned}$$

This along with (18) proves the proposition. $\blacksquare$

B. Bounding the Minimum Cost to Overflow Through Lyapunov Functions

Finding the minimum cost to overflow in (15) is a multidimensional CoV problem, which is often very difficult [23], [24], [29]. In this section, we first use the idea of [26] to show another much simpler lower bound (Proposition 6). We will exploit the fact that V_α is a Lyapunov function of the system operated under the α -algorithm. We will then show that this lower bound is indeed equal to the minimum cost to overflow, and it can be attained by a simple linear trajectory.

We begin with a result that characterizes the relationship between $V_\alpha(\mathbf{q}(\cdot))$ and the channel-rate process $\mathbf{s}(\cdot)$.

Proposition 5: Let $(\mathbf{s}(\cdot), \mathbf{q}(\cdot))$ be any fluid sample path. Except for a set $\mathcal{T}_0$ of measure zero, at any time $t \notin \mathcal{T}_0$ and $\mathbf{q}(t) \neq 0$, the drift of the Lyapunov function $V_\alpha(\mathbf{q}(t))$ is

$$\begin{aligned} \dot{V}_\alpha(\mathbf{q}(t)) = \left(\sum_{i=1}^N (q_i(t))^{\alpha+1} \right)^{\frac{-\alpha}{\alpha+1}} \left[\sum_{i=1}^N (q_i(t))^\alpha \lambda_i \right. \\ \left. - \sum_{m=1}^M \dot{s}_m(t) \max_{1 \leq k \leq N} ((q_k(t))^\alpha F_m^k) \right]. \quad (21) \end{aligned}$$

The proof is provided in our technical report [28].

Remark: An intuitive way to understand Proposition 5 is as follows. From (2), if we take the scaling in (7) and (10) and let

$B \rightarrow \infty$, we would expect that the limiting fluid sample path will follow an ordinary differential equation as follows: There exists $\tilde{\gamma}_m^i(t)$, $i = 1, \dots, N$, $m = 1, \dots, M$ such that

$$\dot{q}_i(t) = \lambda_i - \sum_{m=1}^M \dot{s}_m(t) \tilde{\gamma}_m^i(t) F_m^i$$

if $q_i(t) > 0$ or $\lambda_i - \sum_{m=1}^M \dot{s}_m(t) \tilde{\gamma}_m^i(t) F_m^i \geq 0$; $\dot{q}_i(t) = 0$, otherwise; $[\tilde{\gamma}_m^i(t)]$ are nonnegative and satisfy

$$\begin{aligned} \sum_{i=1}^N \tilde{\gamma}_m^i(t) &= 1 \text{ for all } m = 1, \dots, M, \\ \tilde{\gamma}_m^i(t) &= 0 \text{ whenever } (q_i(t))^\alpha F_m^i < \max_{1 \leq k \leq N} (q_k(t))^\alpha F_m^k. \end{aligned} \quad (22)$$

The variables $\tilde{\gamma}_m^i(t)$ can be viewed as the fraction of time that user i is served when channel state is m in an infinitesimal interval immediately after t . Then, using the definition of $V_\alpha(\cdot)$, at any time t when $\mathbf{q}(t)$ is differentiable, we must have

$$\begin{aligned} \dot{V}_\alpha(\mathbf{q}(t)) &= \left(\sum_{i=1}^N (q_i(t))^{\alpha+1} \right)^{\frac{-\alpha}{\alpha+1}} \\ &\times \left[\sum_{i=1}^N (q_i(t))^\alpha \lambda_i - \sum_{m=1}^M \dot{s}_m(t) \sum_{i=1}^N (q_i(t))^\alpha \tilde{\gamma}_m^i(t) F_m^i \right]. \end{aligned}$$

Using (22), (21) then follows. In our technical report [28], we provide the proof of Proposition 5, which makes this argument more precise.

Next, for any $\phi \in \mathcal{P}_M$, let $\mathbf{x} = [x_i, i = 1, \dots, N]$, and let

$$\begin{aligned} a(\phi) &= \max_{\mathbf{x} \geq 0} \left[\sum_{i=1}^N x_i^\alpha \lambda_i - \sum_{m=1}^M \phi_m \max_{1 \leq k \leq N} (x_k^\alpha F_m^k) \right] \\ \text{subject to } &\sum_{i=1}^N x_i^{\alpha+1} \leq 1. \end{aligned} \quad (23)$$

We will show soon that the Lyapunov drift on the right-hand side of (21) must be no larger than $a(\dot{\mathbf{s}}(t))$. Furthermore, let

$$J_\alpha \triangleq \inf_{\{\phi \in \mathcal{P}_M \mid a(\phi) > 0\}} \frac{H(\phi \|\mathbf{p}\|)}{a(\phi)}. \quad (24)$$

Then, intuitively J_α can be interpreted as a lower bound on unit cost to raise $V_\alpha(\mathbf{q}(t))$. In order to overflow, we must raise $V_\alpha(\mathbf{q}(t))$ from 0 to 1. Hence, J_α should be a lower bound on the minimum cost to overflow, which is indeed the case as we show in the following proposition.

Proposition 6: For any $T > 0$, the following holds:

$$\limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P}_0 [V_\alpha(\mathbf{q}^B(T)) \geq 1] \leq -J_\alpha. \quad (25)$$

Remark: Note that the event $V_\alpha(\mathbf{q}^B(T)) \geq 1$ is equivalent to $V_\alpha(\mathbf{Q}(BT)) \geq B$. As $T \rightarrow \infty$, we would expect that the probability $\mathbf{P}_0[V_\alpha(\mathbf{q}^B(T)) \geq 1]$ approaches the stationary overflow probability $\mathbf{P}[V_\alpha(\mathbf{q}^B(0)) \geq 1]$. Since J_α is independent of T , we would then expect that J_α becomes a lower bound for the decay rate of the stationary overflow probability, i.e.,

$$\limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P} [V_\alpha(\mathbf{q}^B(0)) \geq 1] \leq -J_\alpha.$$

This convergence can indeed be shown using the so-called Freidlin–Wentzell theory [11], [25]. However, the details are quite technical. Due to space constraints, we do not provide the details here. Interested readers can refer to our technical report [28].

Proof (of Proposition 6): Fix $T > 0$. Recall the definition of $\hat{\Gamma}_T$ in Section V-A. For any fluid sample path $(\mathbf{s}(\cdot), \mathbf{q}(\cdot))$ in $\hat{\Gamma}_T$ (which overflows at time T), we will show that the cost of the path $\int_0^T H(\mathbf{s}(t)|\mathbf{p})dt$ is at least J_α . The result of the proposition then follows from Proposition 4. Toward this end, note that since the backlog process $\mathbf{q}(\cdot)$ is Lipschitz-continuous, it is differentiable almost everywhere. According to Proposition 5, for any t such that $t \notin \mathcal{T}_0$ and $\mathbf{q}(t) \neq 0$, we must have

$$\begin{aligned} \dot{V}_\alpha(\mathbf{q}(t)) &= \left(\sum_{i=1}^N (q_i(t))^{\alpha+1} \right)^{\frac{-\alpha}{\alpha+1}} \left[\sum_{i=1}^N (q_i(t))^\alpha \lambda_i \right. \\ &\quad \left. - \sum_{m=1}^M \dot{s}_m(t) \max_{1 \leq k \leq N} ((q_k(t))^\alpha F_m^k) \right] \\ &= \sum_{i=1}^N \tilde{q}_i^\alpha \lambda_i - \sum_{m=1}^M \dot{s}_m(t) \max_{1 \leq k \leq N} (\tilde{q}_k^\alpha F_m^k) \end{aligned}$$

where $\tilde{q}_i = q_i(t) [\sum_{i=1}^N (q_i(t))^{\alpha+1}]^{-(1/(\alpha+1))}$, $i = 1, \dots, N$. Since $\sum_{i=1}^N \tilde{q}_i^{\alpha+1} = 1$, $\tilde{\mathbf{q}} = [\tilde{q}_i]$ is a feasible point that satisfies the constraint in (23). We then have

$$\dot{V}_\alpha(\mathbf{q}(t)) \leq a(\dot{\mathbf{s}}(t)).$$

Hence, if $\dot{V}_\alpha(\mathbf{q}(t)) > 0$, we must have $a(\dot{\mathbf{s}}(t)) > 0$. Then, using the definition of J_α in (24), we have

$$H(\dot{\mathbf{s}}(t)|\mathbf{p}) \geq J_\alpha \dot{V}_\alpha(\mathbf{q}(t)).$$

On the other hand, if $\dot{V}_\alpha(\mathbf{q}(t)) \leq 0$, the above inequality also holds trivially. Hence, the cost of the path must satisfy

$$\int_0^T H(\dot{\mathbf{s}}(t)|\mathbf{p}) dt \geq J_\alpha \int_0^T \dot{V}_\alpha(\mathbf{q}(t)) dt.$$

Recall that any fluid sample path in $\hat{\Gamma}_T$ must satisfy $\mathbf{q}(0) = 0$ and $V_\alpha(\mathbf{q}(T)) \geq 1$. Hence

$$\int_0^T \dot{V}_\alpha(\mathbf{q}(t)) dt \geq 1.$$

The result of the proposition then follows. ■

Remark: We briefly comment on why it is critical to use a Lyapunov function in the above procedure. Although a result similar to Proposition 6 could also be derived if we replace $V_\alpha(\cdot)$ by any function of $\mathbf{q}(t)$, such a result is only useful when the lower bound J_α is positive (otherwise the bound is trivial). The fact that $V_\alpha(\cdot)$ is a Lyapunov function is the key to ensure this property. To see this, note that if $\phi = \mathbf{p}$, then the drift of the Lyapunov function will be negative for any $\mathbf{q}(t)$ (which is required for the stability of the system), implying that the value of $a(\mathbf{p}) = 0$. Hence, for the constraint in (24) to be satisfied, ϕ must be away from $\mathbf{p}$. As a result, the objective function of (24) must be positive. We will see soon that this then implies that the infimum in (24) is also positive.

C. The Path to Overflow That Attains the Lower Bound J_α

In this subsection, we further simplify J_α , and then show that J_α is equal to the minimum cost to overflow in (15). We define the following optimization problem. Let $\mathbf{y} = [y_1, \dots, y_N]$. For any $\phi \in \mathcal{P}_M$, define

$$\begin{aligned} w_\alpha(\phi) &= \min_{\mathbf{y} \geq 0, [\tilde{\gamma}_m^i] \geq 0} V_\alpha(\mathbf{y}) \\ \text{subject to } y_i &= \left[\lambda_i - \sum_{m=1}^M \phi_m \tilde{\gamma}_m^i F_m^i \right]^+ \text{ for all } i \\ \sum_{i=1}^N \tilde{\gamma}_m^i &= 1 \text{ for all } m = 1, \dots, M. \end{aligned}$$

Note that $w_\alpha(\phi)$ is analogous to $w(\phi)$ defined in Section IV. Again, $\tilde{\gamma}_m^i$ can be interpreted as some long-term fraction of time that user i is served when the channel state is m . Hence, if the channel-rate process is given by $\mathbf{s}(t) = \phi t$, then y_i denotes the long-term growth rate of the backlog of user i . Furthermore, if all queues start empty, then $w_\alpha(\phi)$ is the minimum rate of growth of $V_\alpha(\mathbf{q}(t))$ over all policies. We have the following important lemma.

Lemma 7: For any $\phi \in \mathcal{P}_M$, the following holds:

(a)

$$w_\alpha(\phi) = a(\phi).$$

(b) The optimizer $\mathbf{x}^*$ for $a(\phi)$ and the optimizer $\mathbf{y}^*$ for $w_\alpha(\phi)$ are both unique, and they satisfy $\mathbf{x}^* = \gamma \mathbf{y}^*$ for some $\gamma > 0$. Furthermore, if the optimizer $\mathbf{x}^* \neq 0$, then $\mathbf{x}^*$ and $\mathbf{y}^*$ are the only vectors that satisfy the following conditions: There exist $\mu_m^i \geq 0$ such that $\sum_{i=1}^N \mu_m^i = \phi_m$, $y_i^* = [\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i]^+$, $x_i^* = \gamma y_i^*$ for some $\gamma > 0$, $\sum_{i=1}^N (x_i^*)^{\alpha+1} \leq 1$, and

$$\mu_m^i = 0 \text{ whenever } (x_i^*)^\alpha F_m^i < \max_{1 \leq k \leq N} (x_k^*)^\alpha F_m^k.$$

This lemma is proved by showing that the two problems $a(\phi)$ and $w_\alpha(\phi)$ can be viewed as dual problems of each other. The details of the proof are provided in Appendix-A.

Using part (a) of Lemma 7, we immediately obtain the following:

$$J_\alpha = \inf_{\{\phi \in \mathcal{P}_M | w_\alpha(\phi) > 0\}} \frac{H(\phi|\mathbf{p})}{w_\alpha(\phi)}. \quad (26)$$

Furthermore, according to Proposition 6, the above expression provides a lower bound for the decay rate of the queue-overflow probability $\mathbf{P}_0[V_\alpha(\mathbf{q}_t^B(T)) \geq 1]$ for any $T > 0$. The following lemma shows that J_α is positive, and hence the above bound is nontrivial.

Proposition 8:

$$J_\alpha \geq \frac{1}{N^{\frac{1}{\alpha+1}}} I_{\text{opt}}.$$

Proof: Recall that $J_\alpha = \inf_{\{\phi \in \mathcal{P}_M | w_\alpha(\phi) > 0\}} (H(\phi|\mathbf{p})/w_\alpha(\phi))$ and $I_{\text{opt}} = \inf_{\{\phi \in \mathcal{P}_M | w(\phi) > 0\}} (H(\phi|\mathbf{p})/w(\phi))$.

For all $\mathbf{x} \geq 0$, we have $N^{1/(\alpha+1)} \max_{1 \leq i \leq N} x_i \geq V_\alpha(\mathbf{x})$. Furthermore, since $w(\phi)$ and $w_\alpha(\phi)$ have the same constraint set,

we have $N^{1/(\alpha+1)}w(\phi) \geq w_\alpha(\phi)$, and as a consequence, we have

$$\{\phi | w_\alpha(\phi) > 0\} \subseteq \{\phi | w(\phi) > 0\}. \quad (27)$$

Hence, for any ϕ such that $w_\alpha(\phi) > 0$, we have

$$\frac{H(\phi || \mathbf{p})}{w_\alpha(\phi)} \geq \frac{1}{N^{\frac{1}{\alpha+1}}} \frac{H(\phi || \mathbf{p})}{w(\phi)}.$$

Taking infimum over the corresponding constraint sets and using (27), we then obtain $J_\alpha \geq (1/N^{1/(\alpha+1)})I_{\text{opt}}$. ■

Finally, we can show that the lower bound J_α is tight in the sense that there exists $T > 0$ and a trajectory that overflows at T with cost J_α . We will need the following lemma, which provides a structural property of the fluid sample path when the channel-rate process is linear. Specifically, if the channel-rate process $\mathbf{s}(\cdot)$ is linear, then the queue trajectory $\mathbf{q}(\cdot)$ must also be linear, and its derivative must solve $w_\alpha(\phi)$.

Lemma 9: Consider a fluid sample path $(\mathbf{s}(t), \mathbf{q}(t))$ under the α -algorithm. If $\mathbf{q}(0) = 0$ and $\mathbf{s}(t) = t\phi$ for $t \geq 0$, then the corresponding queue trajectory $\mathbf{q}(t)$ must satisfy the following:

- (a) The queue trajectory is linear, i.e., there exists $\tilde{\mathbf{y}} = [\tilde{y}_i, i = 1, \dots, N] \geq 0$, such that $\mathbf{q}(t) = t\tilde{\mathbf{y}}$ for all $t \geq 0$.
- (b) There must exist $\mu_m^i \geq 0$ such that $\sum_{i=1}^N \mu_m^i = \phi_m$, $\tilde{y}_i = [\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i]^+$, and

$$\mu_m^i = 0 \text{ whenever } \tilde{y}_i^\alpha F_m^i < \max_{1 \leq k \leq N} \tilde{y}_k^\alpha F_m^k.$$

In other words, the queue trajectory $\mathbf{q}(t)$ is consistent with the scheduling rule of the α -algorithm.

- (c) $\mathbf{y}^* = \tilde{\mathbf{y}}$ is the unique minimizer of $w_\alpha(\phi)$.

Proof: Let

$$\Omega(\phi) = \{\lambda | \text{there exists } \mu_m^i \geq 0 \text{ such that}$$

$$\lambda_i \leq \sum_{m=1}^M \mu_m^i F_m^i \text{ for all } i = 1, \dots, N, \text{ and} \\ \sum_{i=1}^N \mu_m^i = \phi_m \text{ for all } m = 1, \dots, M \}.$$

Note that if $\phi = \mathbf{p}$, then $\Omega(\phi)$ corresponds to the capacity region of the system (for stability) [3]. The variables μ_m^i can be viewed as some long-term fraction of time that user i is picked and the channel state is m .

Recall from Proposition 5 that

$$\dot{V}_\alpha(\mathbf{q}(t)) = \left(\sum_{i=1}^N (q_i(t))^{\alpha+1} \right)^{\frac{-\alpha}{\alpha+1}} \\ \times \left[\sum_{i=1}^N (q_i(t))^\alpha \lambda_i - \sum_{m=1}^M \phi_m \max_{1 \leq k \leq N} (q_k(t))^\alpha F_m^k \right].$$

First, consider the case when $\lambda \in \Omega(\phi)$. We will have $\dot{V}_\alpha(\mathbf{q}(t)) \leq 0$ if $\mathbf{q}(t) \neq 0$. Hence, starting from $\mathbf{q}(0) = 0$, we must have $V_\alpha(\mathbf{q}(t)) = 0$ and $\mathbf{q}(t) = 0$ for all $t \geq 0$. Therefore, part (a) holds with $\tilde{y}_i = 0$ for all i . Part (b) then trivially holds. Part (c) follows since the minimizer of $w_\alpha(\phi)$ for this case is $\mathbf{y}^* = 0$.

On the other hand, if $\lambda \notin \Omega(\phi)$, then for all $\mathbf{q}(t) \neq 0$, by setting $\hat{q}_i(t) = q_i(t)/[\sum_{i=1}^N (q_i(t))^{\alpha+1}]^{1/(\alpha+1)}$, we have

$$\dot{V}_\alpha(\mathbf{q}(t)) = \sum_{i=1}^N \hat{q}_i^\alpha(t) \lambda_i - \sum_{m=1}^M \phi_m \max_{1 \leq k \leq N} \hat{q}_k^\alpha(t) F_m^k$$

and $\sum_{i=1}^N \hat{q}_i^{\alpha+1}(t) = 1$. We thus have $\dot{V}_\alpha(\mathbf{q}(t)) \leq a(\phi)$ and $V_\alpha(\mathbf{q}(t)) \leq ta(\phi)$ for all $t \geq 0$. This shows that $ta(\phi)$ upper-bounds the maximum growth of $V_\alpha(\mathbf{q}(t))$. On the other hand, let μ_m^i be the average fraction of time in $[0, t]$ that user i is picked and the channel state is m . Then, $\sum_{i=1}^N \mu_m^i = \phi_m$ for all m , and

$$q_i(t) \geq t \left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+.$$

(The inequality is due to the fact that the queue q_i may be empty at some points in this interval). Hence

$$V_\alpha(\mathbf{q}(t)) \geq tw_\alpha(\phi).$$

However, by Lemma 7, $a(\phi) = w_\alpha(\phi)$. We thus have

$$V_\alpha(\mathbf{q}(t)) = ta(\phi) = tw_\alpha(\phi),$$

i.e., there is only one possible trajectory $V_\alpha(\mathbf{q}(t))$ given that $\mathbf{s}(t) = t\phi$. Furthermore, we have $V_\alpha(\mathbf{q}(t))/t = w_\alpha(\phi)$, i.e., $\mathbf{q}(t)/t$ optimizes $w_\alpha(\phi)$. Since the optimizer of $w_\alpha(\phi)$, denoted by $\tilde{\mathbf{y}}$, is unique, we thus have $\mathbf{q}(t) = t\tilde{\mathbf{y}}$. This shows parts (a) and (c). Part (b) follows from part (b) of Lemma 7. ■

The following result then shows that the lower bound J_α is tight. Recall the definition of $\hat{\Gamma}_T$ in Section V-A.

Proposition 10: There exists T and a fluid sample path in $\hat{\Gamma}_T$ whose cost is equal to J_α .

Proof: Let ϕ^* denote the solution to J_α in (26), i.e., $J_\alpha = H(\phi^* || \mathbf{p})/w_\alpha(\phi^*)$, and let $w^* = w_\alpha(\phi^*) > 0$. (We can show that such a ϕ^* always exists by showing that the infimum in (26) can be taken within a closed subset of the original constraint set.) If we use $\mathbf{s}(t) = t\phi^*$, $t \geq 0$ as the channel-rate process and let the queue process start from $\mathbf{q}(0) = 0$, then $\mathbf{q}(\cdot)$ must follow a linear trajectory according to Lemma 9, i.e.,

$$\mathbf{q}(t) = t\tilde{\mathbf{x}}, \text{ for all } t \geq 0,$$

where $\mathbf{y}^* = \tilde{\mathbf{x}}$ is the minimizer of $w_\alpha(\phi^*)$.

Let $T = 1/w_\alpha(\phi^*)$. Consider such a trajectory over the interval $[0, T]$. Clearly, the cost of this trajectory is equal to J_α . It only remains to show that the trajectory must overflow at T , which is true because $V_\alpha(T\tilde{\mathbf{x}}) = Tw_\alpha(\phi^*) = 1$. ■

Hence, we conclude that the minimum cost to overflow is attained by a simple linear trajectory whose cost is J_α .

VI. ASYMPTOTICAL OPTIMALITY OF α -ALGORITHMS

In this section, we will establish that in the limit as $\alpha \rightarrow \infty$, the α -algorithms asymptotically achieve the largest minimum cost to overflow equal to I_{opt} given in (11). To emphasize the dependence on α , we use $\mathbf{P}_0^\alpha$ to denote the probability distribution conditioned on $\mathbf{Q}(0) = 0$ under the α -algorithm (with a particular value of α). We now show the following.

Proposition 11: For any $T > 0$, the following holds:

$$\lim_{\alpha \rightarrow \infty} \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P}_0^\alpha \left[\max_{1 \leq i \leq N} q_i^B(T) \geq 1 \right] \leq -I_{\text{opt}}.$$

Proof: Since $\max_{1 \leq i \leq N} q_i(T) \geq 1$ implies $V_\alpha(\mathbf{q}(T)) \geq 1$, we must have

$$\mathbf{P}_0^\alpha \left[\max_{1 \leq i \leq N} q_i^B(T) \geq 1 \right] \leq \mathbf{P}_0^\alpha [V_\alpha(\mathbf{q}(T)) \geq 1].$$

Using Proposition 6, for all $T > 0$

$$\begin{aligned} \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P}_0^\alpha \left[\max_{1 \leq i \leq N} q_i^B(T) \geq 1 \right] \\ \leq \limsup_{B \rightarrow \infty} \frac{1}{B} \log \mathbf{P}_0^\alpha [V_\alpha(\mathbf{q}(T)) \geq 1] \leq -J_\alpha. \end{aligned}$$

From Proposition 8, $\lim_{\alpha \rightarrow \infty} J_\alpha \geq I_{\text{opt}}$. The result then follows. ■

Combining Propositions 1 and 11, we conclude that the α -algorithms asymptotically achieve the largest decay rate I_{opt} of the queue-overflow probability over all scheduling policies.

We briefly comment on the behavior of the α -algorithms when α increase. As $\alpha \rightarrow \infty$, the α -algorithm places more and more emphasis on the queue length. For instance, in a two-user system, if $Q_1(t) > Q_2(t)$, $F_{C(t)}^1 > 0$ and $F_{C(t)}^2 > 0$, then all α -algorithms with $\alpha > (\log(F_{C(t)}^2/F_{C(t)}^1)/\log(Q_1(t)/Q_2(t)))$ would serve Q_1 . On the other hand, if $Q_1(t) = Q_2(t)$, then the link with the larger capacity $F_{C(t)}^i$ would be served. Therefore, as $\alpha \rightarrow \infty$, we would expect that the α -algorithm would give more and more preference to the link with the largest queue backlog among all links with nonzero rates. If there are several links that have the same (largest) backlog, the link with the highest rate among them would be served. However, we caution that if we choose $\alpha = \infty$, then the resulting algorithm is the max-queue algorithm, which is not throughput-optimal for general channel models. Therefore, the above intuition does not directly lead to a stable scheduling policy. We will obtain more intuition about this issue when we look at the simulation results in Section VII.

Note that in [7], the authors provide an explanation in the heavy-traffic regime for the conjecture that when $\alpha \rightarrow 0$, the α -algorithm becomes asymptotically optimal in minimizing the average delay. The reason that we have a different regime of asymptotic optimality (i.e., $\alpha \rightarrow \infty$) is because we study a different objective. Although the delay metric in [7] is not clearly defined, the objective appears to be closely related to minimizing the sum of queues, while our goal is to minimize $\max_{1 \leq i \leq N} Q_i$. Hence, in our case, it is more important to serve the queue with the largest backlog, while in [7], it is more important to increase the total service rate in each time-slot.

A. Systems With ON-OFF Channels

Consider the scenario where F_m^i can either take the value 0 or a positive constant C . This scenario corresponds to a wireless system with ON-OFF channels and the ON rates for all users are the same. In this case, for any $\alpha > 0$

$$(q_i)^\alpha F_m^i \leq \max_{1 \leq k \leq N} (q_k)^\alpha F_m^k \Leftrightarrow q_i F_m^i \leq \max_{1 \leq k \leq N} q_k F_m^k.$$

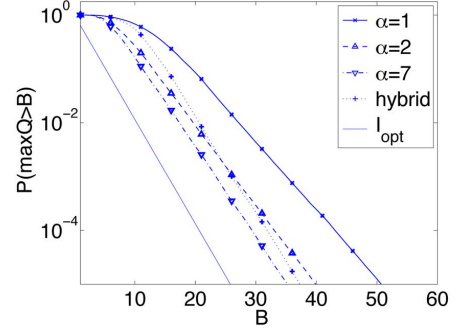


Fig. 1. Case 1: Plot of $\mathbf{P}[\max_{1 \leq i \leq N} Q_i \geq B]$ versus the overflow threshold B for the α -algorithms. Each curve corresponds to a different value of α .

TABLE I
LINK CAPACITIES IN DIFFERENT STATES

F_m^i	$m = 1$	$m = 2$	$m = 3$
$i = 1$	0	3	5
$i = 2$	0	9	0
$i = 3$	0	9	1
$i = 4$	0	9	1

Hence, for any $\alpha \geq 1$, the α -algorithms are equivalent to the max-weight algorithm (i.e., with $\alpha = 1$). Using the result in this paper, we immediately reach the following corollary.

Corollary 12: For the above ON-OFF channel model, the max-weight scheduling algorithm (i.e., $\alpha = 1$) achieves the largest decay rate I_{opt} of the queue-overflow probability over all scheduling policies.

VII. SIMULATION RESULTS

In this section, we will provide simulation results to verify the analytical results in earlier sections. We simulate the following system with four links (i.e., $N = 4$) and three states (i.e., $M = 3$). In each time-slot, one unit of data arrives at each of the links (i.e., $\lambda_1 = \lambda_2 = \lambda_3 = \lambda_4 = 1$). The probabilities of each channel state are denoted as p_1 , p_2 , and p_3 , and will be given shortly. The capacity F_m^i of link i in channel state m is given by Table I. The 95% confidence intervals are very small, and hence they are not shown in the figures.

We first simulate Case 1 when $p_1 = 0.3$, $p_2 = 0.6$ and $p_3 = 0.1$. In Fig. 1, we plot the value of $\mathbf{P}[\max_{1 \leq i \leq N} Q_i \geq B]$ (in log-scale) against the overflow threshold B for the α -algorithms, where each curve corresponds to a different value of α . We have also plotted a line with slope equal to I_{opt} given by (11). Recall that I_{opt} is the maximum decay rate of the queue-overflow probability. We can observe from Fig. 1 that as the value of α increases, the slopes at the tail of the curves (i.e., for large B) approach I_{opt} . Hence, this confirms our analytical result that as the value of α increases, the asymptotic decay rate of the α -algorithms approaches the optimal decay rate I_{opt} .

We have also simulated the exponential rule of [25]. At any time t , if the channel state is m , the exponential rule chooses to serve the link i^* such that

$$i^* = \arg \max_{i=1, \dots, N} \exp \left[\frac{Q_i(t)}{1 + \left(\frac{1}{N} \sum_{k=1}^N Q_k(t) \right)^\eta} \right] F_m^i$$

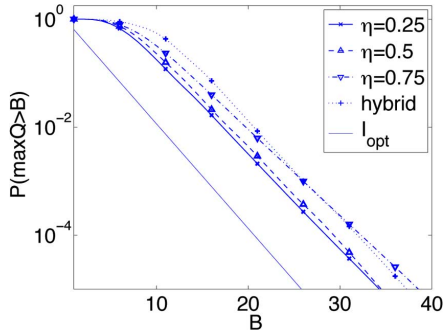


Fig. 2. Case 1: Plot of $P[\max_{1 \leq i \leq N} Q_i \geq B]$ versus the overflow threshold B for the exponential rule. Each curve corresponds to a different value of η .

where η is a constant parameter in $(0,1)$. In Fig. 2, we plot $P[\max_{1 \leq i \leq N} Q_i \geq B]$ against the overflow threshold B for the exponential rule as the parameter η varies. According to the results of [25], the exponential rule achieves the optimal decay rate of the queue-overflow probability for any $0 < \eta < 1$. We observe from Fig. 2 that, for $\eta = 0.25$ and $\eta = 0.5$, the slopes at the tail of the curves indeed become parallel to I_{opt} for large B . For $\eta = 0.75$, such convergence has not occurred even for overflow probability as low as 10^{-5} . Note that one should not conclude from the last curve that the results of [25] are violated: The LDP results of [25] will still kick in eventually, although at a larger value of the threshold B .

The previous set of simulation results raise some important issues on the applicability of large deviation results. Specifically, the results in this paper (and in [25]) are large-buffer asymptotes, i.e., they characterize the behavior of the queue only when the overflow threshold approaches infinity. The results often do not provide much information on what buffer level is large enough for the asymptotic behavior to become dominant. Furthermore, a LDP only specifies the exponential decay rate. The factor in front of exponential term can still vary substantially. Hence, one needs to be careful when comparing the performance predicted by a LDP with the actual performance of the protocol. This point is best illustrated with Case 2 that we simulated. Here, the probability of each channel state is given by $p_1 = 0.35$, $p_2 = 0.5$, and $p_3 = 0.15$. In Fig. 3, we again plot the value of $P[\max_{1 \leq i \leq N} Q_i \geq B]$ against the overflow threshold B for the α -algorithms. We observe from Fig. 3 that as α increases, the slopes at the tail of the curve indeed approaches I_{opt} . However, for small B the curve in fact shifts to the right, indicating that the actual queue-overflow probability $P[\max_{1 \leq i \leq N} Q_i \geq B]$ increases as α increases. Such a shift is more evident for smaller values of B . As B increases, for larger values of α the effect of the steeper slopes eventually dominates, and the queue-overflow probability improves as well.

To better understand this behavior, we introduce a state-space plot as in Fig. 6. The x-axis and the y-axis are the length of any two chosen queues (e.g., Q_1 and Q_3 as in Fig. 6). This state space is divided into regions, each of which corresponds to a fixed scheduling decision. For example, in Region 1, Queue 1 is served irrespective of the channel state (this is the case because the length of Queue 1 is much larger than Queue 3). In Region 2, Queue 1 is served in channel state $m = 3$, and Queue 3 is served in channel state $m = 2$. Finally, in Region 3, Queue 3 is served irrespective of the channel state. We refer to these regions as *decision regions*, and their boundary is determined by the scheduling policy. The

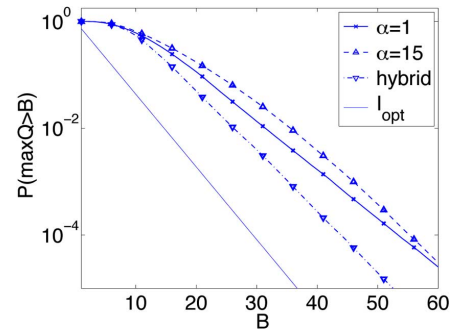


Fig. 3. Case 2: Plot of $P[\max_{1 \leq i \leq N} Q_i \geq B]$ versus the overflow threshold B for the α -algorithm. Each curve corresponds to a different value of α .

dots in the figure are the states that have been visited by the system in the simulation (over some given length of time). A similar state space plot for case 2 is shown in Fig. 8.

Once the probabilities of channel states are given, the capacity region of the system can be determined. For example, Fig. 5 represents the capacity regions of cases 1 and 2, projected to the space of Q_1 and Q_3 . For this system with two active states, we can draw a correlation between the decision regions (e.g., Fig. 6) and the capacity region (e.g., case 1 in Fig. 5). We will refer to Regions 1 and 3 as *max-queue* regions in the sense that the decision is to serve the link with the longest queue, irrespective of the channel state. We refer to Region 2 as the *max-rate* region in the sense that now the decision is to serve the link with the higher rate, depending on which channel state the system is in. The two max-queue regions can be correlated to the points μ_1 and μ_3 of the capacity region, where one user will be served in all states. The max-rate region can be correlated to the point μ_2 of the capacity region. The significance of this correlation is that region 2 contributes to an enlarged capacity region (i.e., the triangular area $\mu_1\mu_2\mu_3$).

For α -algorithms, as the value of α increases, the boundaries between the decision regions all converge to the diagonal line. This convergence has two implications. First, a larger value of α enlarges the two max-queue regions (see Fig. 7). For example, Point A that was in a max-rate region for small α (see Fig. 6) now moves to the max-queue region (see Fig. 7). Note that at Point A, we have $Q_1 > Q_3$. Hence, as the decision boundaries approach the diagonal line, the algorithm places more emphasis on reducing the largest queue. Intuitively, this helps to improve the decay rate of the probability that the largest queue overflows.

However, a second effect of increasing α is that the size of the max-rate region (i.e., Region 2) is reduced. As a result, for smaller values of queue length, it becomes less likely that the system state falls into the max-rate region. Recall that the decision rule in the max-rate region contributes to the improved capacity region (i.e., triangular area $\mu_1\mu_2\mu_3$). Hence, with large values of α , the scheduling algorithm is unlikely to take advantage of the increased capacity at small queue lengths, which leads to a tendency for the queue length to grow. This phenomenon can be observed by the fact that the dots in Fig. 7 now grow along the two boundary lines. It is even more evident in a similar plot for Case 2 (in Fig. 9). After the queue length increases, eventually the width of Region 2 will be sufficiently large so that the system state is more likely to fall into the max-rate region. Only after that, the effect of LDP starts to kick in, and the decay rate of the queue-overflow probability starts to improve.

Although the above discussion is restricted to the dynamics of two queues over two active states, we feel that the above two conflicting trends apply to more general cases. Indeed, the understanding of these two trends help us to interpret the results in Figs. 1 and 3. First, refer to Fig. 6 for Case 1. For small values of α , the queues tend to accumulate around the boundary between Regions 1 and 2. As α increases, the max-queue region (Region 1) enlarges, which helps to reduce the longer queue and push the state space to the origin (Fig. 7). The conflicting effect due to thinning of the max-rate region is not so strong, and the beneficial effect of large α is manifested. Thus, these plots explain why the performance plot in Fig. 1 improves with increasing α . Now, comparing the capacity region for the two cases (Fig. 5), we find that in case 2, the offered load λ is closer to the line $\mu_1\mu_3$. Hence, the triangular section $\mu_1\mu_2\mu_3$ plays a more significant role in reducing the queue length. We would thus expect the effect of thinning of the max-rate region to be relatively stronger than in the previous case. This is exactly what we observe in Figs. 8 and 9. At small values of α (Fig. 8), the queues tend to accumulate relatively more in the max-rate region. Now, as α increases, the stronger effect caused by the thinning of the max-rate region forces the queue length to increase (Fig. 9). As a result, at small values of threshold B , the overflow probability in fact deteriorates.

The above observations motivate us to design a new class of hybrid scheduling policies that have the benefits of both large α (for improving the large-deviation decay rate of the queue-overflow probability) and small α (for having a large max-rate region, which helps to improve the overflow probability at small queue lengths). Essentially, to have good large-deviation decay rates of the queue-overflow probability, we need to use a large α so that the decision boundaries become close to parallel to the diagonal line. However, this may lead to poor performance at small queue lengths due to the thinner max-rate regions. To avoid this, we first use a smaller value of α when the queue length is small and gradually change to large α when queue increases. Note that this does not mean that we can use $\alpha = \infty$ and $\alpha = 0$ for the large α region and the small α region, respectively. The reason is that $\alpha = \infty$ and $\alpha = 0$ will degenerate to the max-queue policy and the max-rate policy, respectively, and neither of them are throughput-optimal policies (see also the discussions before Section VI-A). For example, if we use $\alpha = \infty$, the decision boundaries will be exactly parallel to the diagonal. This means that the max-rate region will not become “thicker” as the queue lengths increase. This may cause instability because the queue state may not be able to stay in the max-rate region for a sufficient fraction of time.

More specifically, the hybrid policy works by modifying the weight function. The scheduling policy still picks the user i for service such that it has the largest value of $w_i(\mathbf{q})F_m^i$. However, the weight of user i , $w_i(\mathbf{q})$, is not equal to q_i^α anymore. Instead, it contains both a term for small α and a term for large α . Specifically, let us assume that we are interested in transitioning from small α to large α when the queue length is around $B_* = 10$. We tested a hybrid policy that uses a combination of $\alpha = 1$ and $\alpha = 15$.⁵ The weight function we used is $w_i(\mathbf{q}) = q_i +$

$([q_i - (K(\mathbf{q})/F_m^i)]^+)^{15}$, where the value $K(\mathbf{q})$ will be specified later. For $q_i < (K(\mathbf{q})/F_m^i)$, the weight function is simply q_i . Hence, the behavior of the scheduling algorithm is similar to $\alpha = 1$. For large q_i , the term $(q_i - (K(\mathbf{q})/F_m^i))^{15}$ dominates. Hence, the behavior of the scheduling algorithm switches to that of $\alpha = 15$. The offset $K(\mathbf{q})/F_m^i$ is chosen to ensure that the decision boundary does not have sudden jumps. Specifically, the value of $K(\mathbf{q})$ is given by

$$K(\mathbf{q}) = \min_{1 \leq i \leq N} \left(B_* F_m^i + [B_* - q_i]^+ \max_{1 \leq k \leq N} F_m^k \right). \quad (28)$$

To understand the intuition behind (28), first consider the case when $q_i > B_*$ for all queues. Then, $K(\mathbf{q}) = B_* \min_{1 \leq i \leq N} F_m^i$. The offset in this case becomes $(B_* \min_{1 \leq i \leq N} F_m^i / F_m^1, \dots, B_*, \dots, B_* \min_{1 \leq i \leq N} F_m^i / F_m^N)$, which is exactly the point where the decision boundary of $\alpha = 1$ meets the threshold boundary $\max_{1 \leq i \leq N} q_i = B_*$. However, if we just use $K(\mathbf{q}) = B_* \min_{1 \leq i \leq N} F_m^i$, the problem is that the transition to large α occurs too early in the case when not all q_i are greater than B_* . For example, consider channel state $m = 2$. In this case, the offset described above becomes $(B_*, B_*/3, B_*/3, B_*/3)$. The projection of this offset value to the space of the queues q_2, q_3 , and q_4 is $(B_*/3, B_*/3, B_*/3)$. As a result, the transition from $\alpha = 1$ to $\alpha = 15$ would occur too early (at $B_*/3$) for q_2, q_3 , or q_4 if q_1 is small. To compensate for this effect, we have introduced the second term in (28). Essentially, if q_1 is small, its channel rates do not play much role in determining the minimum value of (28). In this specific example, if $q_1 = 0$ and $q_2, q_3, q_4 > B_*$, then the offset value is $(3B_*, B_*, B_*, B_*)$. Hence, the transition occurs at the desirable values of q_2, q_3 and q_4 .

We plot the decision boundaries for this hybrid algorithm in Fig. 10. As we can see, the max-rate region is large even for small queue-lengths. In Fig. 3, we also plotted the performance of the hybrid algorithm. Compare with the curve for $\alpha = 15$, we note that the curve for the hybrid algorithm has moved to the left as desired. Also note that the slope of the curve is close to I_{opt} . Hence, this figure confirms that the hybrid algorithm achieves the benefit of both large α and small α .

We find that the same intuitions seem to also apply for the exponential rule [25]. Recall that Fig. 2 plots the value of $P[\max_{1 \leq i \leq N} Q_i \geq B]$ versus the overflow threshold B for the exponential rule when the parameter η varies. A similar figure for Case 2 is given in Fig. 4. To understand why $\eta = 0.5$ seems to produce the best overall performance, we plot the decision boundaries of the exponential rule in Fig. 11. We can see that if the value of η is too small, then the max-rate region (between the decision boundaries) is too narrow, which increases the queue-overflow probability at small threshold values. If the value of η is too large, then the max-rate region is big enough. However, the decision boundaries do not become parallel to the diagonal line until the queue length is very large. Hence, the large-deviation decay rate kicks in only at a larger queue length. A medium value of η (around 0.5) seems to achieve a balance between the above two cases and produces a state-space plot that is similar to our hybrid algorithm (Fig. 10). We have also plotted the performance of the exponential rule and our hybrid

⁵We choose $\alpha = 1$ because we would like to compare with the standard max-weight algorithm, which is an α -algorithm with $\alpha = 1$. The choice of $\alpha = 15$ is somewhat arbitrary. Simulations using $\alpha = 30$ (not shown) resulted in almost identical performance.

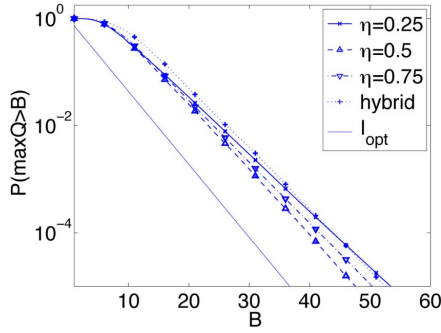


Fig. 4. Case 2: Plot of $P[\max_{1 \leq i \leq N} Q_i \geq B]$ versus the overflow threshold B for the exponential rule. Each curve corresponds to a different value of η .

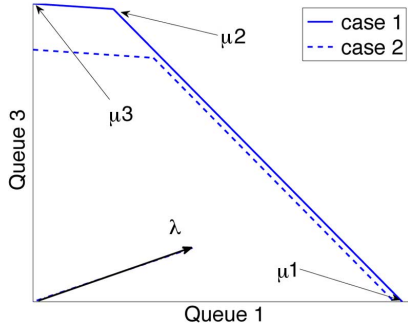


Fig. 5. Shape of the capacity region.

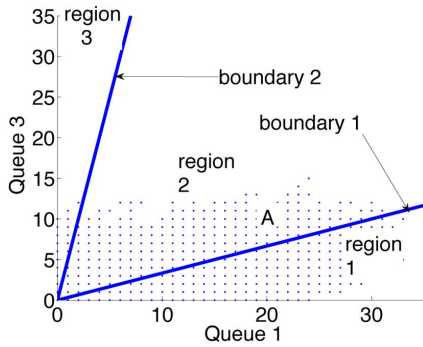


Fig. 6. Case 1: Plot of the state space for $\alpha = 1$.

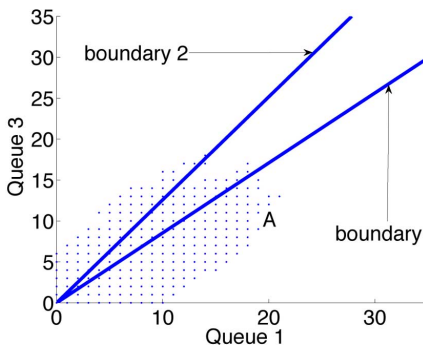


Fig. 7. Case 1: Plot of the state space for $\alpha = 7$.

algorithm in Fig. 4. Their performance appears to be quite comparable. Finally, we plot the performance of the hybrid algorithm for Case 1, and we find that the hybrid algorithm also performs very well, which indicates that the hybrid algorithm is quite robust and seems to work well in all cases.

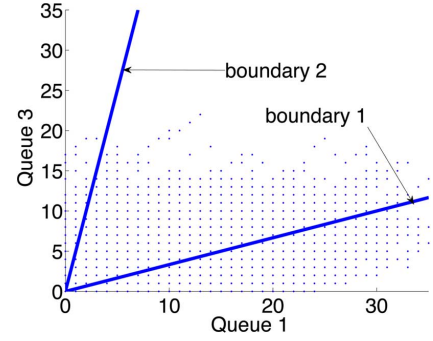


Fig. 8. Case 2: Plot of the state space for $\alpha = 1$.

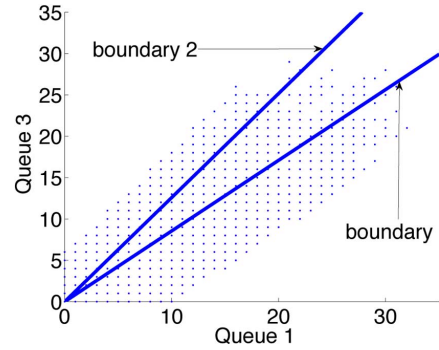


Fig. 9. Case 2: Plot of the state space for $\alpha = 7$.

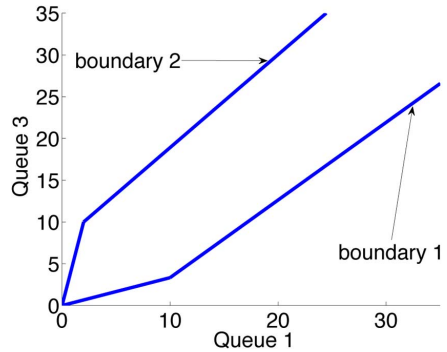


Fig. 10. Plot of the decision boundaries for the hybrid algorithm.

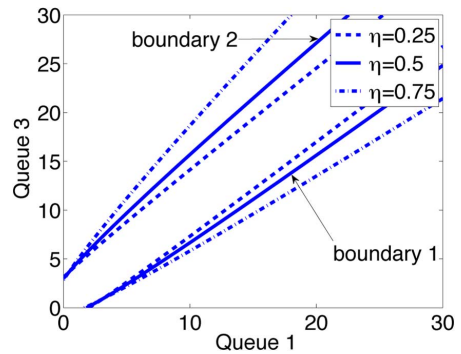


Fig. 11. Plot of the decision boundaries for exponential rule for various values of η .

VIII. CONCLUSION

In this paper, we study wireless scheduling algorithms for the downlink of a single cell that can maximize the asymptotic decay rate of the queue-overflow probability as the overflow

threshold approaches infinity. Specifically, we focus on the class of “ α -algorithms,” which pick the user for service at each time that has the largest product of the transmission rate multiplied by the backlog raised to the power α . We show that when α approaches infinity, the α -algorithms asymptotically achieve the largest decay rate of the queue-overflow probability. A key step in proving this result is to use a Lyapunov function to derive a simple lower bound for the minimum cost to overflow under the α -algorithms. This technique, which is of independent interest, circumvents solving the difficult multidimensional calculus-of-variations problem typical in this type of problem. Finally, using the insight from this result, we design hybrid scheduling algorithms that are both close to optimal in terms of the asymptotic decay rate of the overflow probability and empirically shown to maintain small queue-overflow probabilities over queue-length ranges of practical interest. For future work, we plan to extend the results to more general network and channel models.

A potential limitation of the large-deviations approach used in this work is that although we show optimality in terms of the decay rate, we have not been able to quantify the coefficients before the exponential decay term. Such coefficients may also play an important role, especially when considering small queue values. Unfortunately, they are much more difficult to quantify. The hybrid algorithm in Section VII can be interpreted as an intuitive design engineered to have a better coefficient than the pure α -algorithm.

APPENDIX

A. Proof of Lemma 7

Proof: We first show that $a(\phi)$ and $w_\alpha(\phi)$ are dual problems of each other. Letting $\xi_i = x_i^\alpha$, $i = 1, \dots, N$, $\xi = [\xi_i]$ and introducing the variables $\eta_m \geq \max_{1 \leq i \leq N} \xi_i F_m^i$, the problem $a(\phi)$ can be rewritten as

$$\begin{aligned} a(\phi) = \max_{\xi \geq 0, \eta} & \left[\sum_{i=1}^N \xi_i \lambda_i - \sum_{m=1}^M \phi_m \eta_m \right] \\ \text{subject to} & \sum_{i=1}^N \xi_i^{\frac{\alpha+1}{\alpha}} \leq 1 \\ & \eta_m \geq \xi_i F_m^i \text{ for all } i, m. \end{aligned}$$

This is a convex optimization problem. Introducing the Lagrange multiplier $\mu_m^i \geq 0$ for each of the constraints $\eta_m \geq \xi_i F_m^i$, we obtain the Lagrangian $L(\xi, \mu, \eta) = \sum_{i=1}^N \xi_i \left(\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right) - \sum_{m=1}^M \eta_m \left(\phi_m - \sum_{i=1}^N \mu_m^i \right)$. The dual-objective function is then given by

$$\begin{aligned} D(\mu) = \max_{\xi \geq 0, \eta} & L(\xi, \mu, \eta) \\ \text{subject to} & \sum_{i=1}^N \xi_i^{\frac{\alpha+1}{\alpha}} \leq 1. \end{aligned}$$

Note that if $\sum_{i=1}^N \mu_m^i \neq \phi_m$, then $D(\mu) = +\infty$ since we can set $|\eta_m|$ arbitrarily large. Otherwise, if $\sum_{i=1}^N \mu_m^i = \phi_m$ for all

m , we then have

$$\begin{aligned} D(\mu) = \max_{\xi \geq 0} & \sum_{i=1}^N \xi_i \left(\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right) \\ \text{subject to} & \sum_{i=1}^N \xi_i^{\frac{\alpha+1}{\alpha}} \leq 1. \end{aligned} \quad (29)$$

Clearly, for those i such that $\lambda_i < \sum_{m=1}^M \mu_m^i F_m^i$, the optimal solution for $D(\mu)$ is $\xi_i = 0$. Let $\mathcal{I}$ denote the set of i such that $\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \geq 0$. If $\mathcal{I}$ is an empty set, then $D(\mu) = 0$. If $\mathcal{I}$ is not empty, we can use Holder's inequality that, for any positive p and q such that $1/p + 1/q = 1$, the following holds: $\sum_{i=1}^N a_i b_i \leq \left[\sum_{i=1}^N a_i^p \right]^{1/p} \left[\sum_{i=1}^N b_i^q \right]^{1/q}$, where equality holds if and only if there is a constant γ such that $a_i^p = \gamma b_i^q$ for all i . Hence, for all ξ such that the constraint (29) is satisfied, we have

$$\begin{aligned} \sum_{i \in \mathcal{I}} \xi_i \left(\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right) &= \sum_{i=1}^N \xi_i \left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+ \\ &\leq \left[\sum_{i=1}^N \xi_i^{\frac{\alpha+1}{\alpha}} \right]^{\frac{\alpha}{\alpha+1}} \left[\sum_{i=1}^N \left(\left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+ \right)^{\alpha+1} \right]^{\frac{1}{\alpha+1}} \\ &\leq \left[\sum_{i=1}^N \left(\left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+ \right)^{\alpha+1} \right]^{\frac{1}{\alpha+1}} \end{aligned}$$

where equality holds if and only if

$$\sum_{i=1}^N \xi_i^{\frac{\alpha+1}{\alpha}} = 1 \quad (30)$$

and for some constant $\gamma^{(\alpha+1)/\alpha} > 0$, $\xi_i^{(\alpha+1)/\alpha} = \gamma^{(\alpha+1)/\alpha} \left(\left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+ \right)^{\alpha+1}$, for $i = 1, \dots, N$, or, equivalently, $\xi_i^{1/\alpha} = \gamma^{1/\alpha} \left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+$, for $i = 1, \dots, N$. Such a vector ξ clearly exists when $\mathcal{I}$ is not empty. Hence, if $\sum_{i=1}^N \mu_m^i = \phi_m$ for all m , we have

$D(\mu) = \left[\sum_{i=1}^N \left(\left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+ \right)^{\alpha+1} \right]^{\frac{1}{\alpha+1}}$, which is true even when $\mathcal{I}$ is empty. We can therefore conclude that the dual problem is

$$\begin{aligned} \min_{\mu \geq 0} D(\mu) &= \min_{y \geq 0, \mu \geq 0} \left(\sum_{i=1}^N y_i^{\alpha+1} \right)^{\frac{1}{\alpha+1}} \\ \text{subject to} & y_i = \left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+ \\ & \sum_{i=1}^N \mu_m^i = \phi_m \text{ for all } m. \end{aligned}$$

This is exactly the problem $w_\alpha(\phi)$. Hence, strong duality implies that $a(\phi) = w_\alpha(\phi)$.

The optimizer y of $w_\alpha(\phi)$ must be unique since the objective function in $w_\alpha(\phi)$ is strictly convex in y . Using

the complementary slackness condition, for any optimizer ξ and μ , we must have $\mu_m^i \geq 0$, $\sum_{i=1}^N \mu_m^i = \phi_m$, $\xi_i^{\frac{1}{\alpha}} = \gamma \left[\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i \right]^+$, $\mu_m^i = 0$ if $\xi_i F_m^i < \max_{1 \leq k \leq N} \xi_k F_m^k$

and $\sum_{i=1}^N \xi_i^{(\alpha+1)/\alpha} = 1$ whenever $\xi \neq 0$ by (30). Since $\xi_i = x_i^\alpha$ and $y_i = [\lambda_i - \sum_{m=1}^M \mu_m^i F_m^i]^+$, we must have $\mathbf{x} = \gamma \mathbf{y}$. Furthermore, if $\mathbf{x} \neq 0$, then since $\mathbf{y}$ is unique and $\sum_{i=1}^N x_i^{\alpha+1} = 1$, $\mathbf{x}$ must also be unique. The above set of equations are then exactly the condition in part (b) of the lemma. Conversely, any ξ and μ (or, equivalently, $\mathbf{x}$ and μ) that satisfy the condition must correspond to the maximizer of $a(\phi)$ and $w_\alpha(\phi)$, respectively. Since the optimizers of $a(\phi)$ and $w_\alpha(\phi)$ are both unique, there is at most one $\mathbf{x}$ that satisfies the set of conditions in part (b) of the lemma. ■

ACKNOWLEDGMENT

The authors are grateful for the helpful discussions with Prof. Michael Neely, Prof. Sean Meyn, and Dr. Alexander Stolyar and for the insightful comments provided by the anonymous reviewers.

REFERENCES

- [1] V. J. Venkataramanan and X. Lin, "Structural properties of LDP for queue-length based wireless scheduling algorithms," in *Proc. 45th Annu. Allerton Conf. Commun., Control, Comput.*, Monticello, IL, Sep. 2007.
- [2] X. Lin, N. B. Shroff, and R. Srikant, "A tutorial on cross-layer optimization in wireless networks," *IEEE J. Sel. Areas Commun.*, vol. 24, no. 8, pp. 1452–1463, Aug. 2006.
- [3] L. Tassiulas and A. Ephremides, "Stability properties of constrained queueing systems and scheduling policies for maximum throughput in multihop radio networks," *IEEE Trans. Autom. Control*, vol. 37, no. 12, pp. 1936–1948, Dec. 1992.
- [4] M. J. Neely, E. Modiano, and C. E. Rohrs, "Dynamic power allocation and routing for time-varying wireless networks," in *Proc. IEEE INFOCOM*, San Francisco, CA, Apr. 2003, vol. 1, pp. 745–755.
- [5] A. Eryilmaz, R. Srikant, and J. Perkins, "Stable scheduling policies for fading wireless channels," *IEEE/ACM Trans. Netw.*, vol. 13, no. 2, pp. 411–424, Apr. 2005.
- [6] A. L. Stolyar, "MaxWeight scheduling in a generalized switch: State space collapse and workload minimization in heavy traffic," *Ann. Appl. Probab.*, vol. 14, no. 1, pp. 1–53, 2004.
- [7] D. Shah and D. Wischik, "Optimal scheduling algorithms for input-queued switches," in *Proc. IEEE INFOCOM*, Barcelona, Spain, Apr. 2006.
- [8] S. P. Meyn, "Stability and asymptotic optimality of generalized MaxWeight policies," *SIAM J. Control Optim.*, vol. 47, no. 6, pp. 3259–3294, Jan. 2009.
- [9] M. J. Neely, "Order optimal delay for opportunistic scheduling in multi-user wireless uplinks and downlinks," *IEEE/ACM Trans. Netw.*, vol. 16, no. 5, pp. 1188–1199, Oct. 2008.
- [10] M. J. Neely, "Delay analysis for maximal scheduling in wireless networks with bursty traffic," in *Proc. IEEE INFOCOM*, Phoenix, AZ, April 2008, pp. 6–10.
- [11] A. Schwartz and A. Weiss, *Large Deviations for Performance Analysis: Queues, Communications, and Computing*. London, U.K.: Chapman & Hall, 1995.
- [12] A. Dembo and O. Zeitouni, *Large Deviations Techniques and Applications*, 2nd ed. New York: Springer-Verlag, 1998.
- [13] A. I. Elwalid and D. Mitra, "Effective bandwidth of general Markovian traffic sources and admission control of high speed networks," *IEEE/ACM Trans. Netw.*, vol. 1, no. 3, pp. 329–343, Jun. 1993.
- [14] G. Kesidis, J. Walrand, and C.-S. Chang, "Effective bandwidth for multiclass Markov fluid and other ATM sources," *IEEE/ACM Trans. Netw.*, vol. 1, no. 4, pp. 424–428, Aug. 1993.
- [15] C.-S. Chang and J. A. Thomas, "Effective bandwidth in high-speed digital networks," *IEEE J. Sel. Areas Commun.*, vol. 13, no. 6, pp. 1091–1114, Aug. 1995.
- [16] F. P. Kelly, "Effective bandwidth at multiclass queues," *Queueing Syst.*, vol. 9, pp. 5–16, 1991.
- [17] W. Whitt, "Tail probabilities with statistical multiplexing and effective bandwidth for multi-class queues," *Telecommun. Syst.*, vol. 2, pp. 71–107, 1993.
- [18] A. L. Stolyar and K. Ramanan, "Largest weighted delay first scheduling: Large deviations and optimality," *Ann. Appl. Probab.*, vol. 11, no. 1, pp. 1–48, 2001.
- [19] A. L. Stolyar, "Control of end-to-end delay tails in a multiclass networks: LWDF discipline optimality," *Ann. Appl. Probab.*, vol. 13, no. 3, pp. 1151–1206, 2003.
- [20] D. Wu and R. Negi, "Effective capacity: A wireless link model for support of quality of service," *IEEE Trans. Wireless Commun.*, vol. 2, no. 4, pp. 630–643, Jul. 2003.
- [21] D. Wu and R. Negi, "Downlink scheduling in a cellular network for quality of service assurance," *IEEE Trans. Veh. Technol.*, vol. 53, no. 5, pp. 1547–1557, Sep. 2004.
- [22] D. Wu and R. Negi, "Utilizing multiuser diversity for efficient support of quality of service over a fading channel," *IEEE Trans. Veh. Technol.*, vol. 54, no. 3, pp. 1198–1206, May 2005.
- [23] L. Ying, R. Srikant, A. Eryilmaz, and G. E. Dullerud, "A large deviations analysis of scheduling in wireless networks," *IEEE Trans. Inf. Theory*, vol. 52, no. 11, pp. 5088–5098, Nov. 2006.
- [24] S. Shakkottai, "Effective capacity and QoS for wireless scheduling," *IEEE Trans. Autom. Control*, vol. 53, no. 3, pp. 749–761, Apr. 2008.
- [25] A. L. Stolyar, "Large deviations of queues sharing a randomly time-varying server," *Queueing Syst.*, vol. 59, no. 1, pp. 1–35, May 2008.
- [26] X. Lin, "On characterizing the delay performance of wireless scheduling algorithms," in *Proc. 44th Annu. Allerton Conf. Commun., Control, Comput.*, Monticello, IL, Sep. 2006.
- [27] A. Eryilmaz and R. Srikant, "Scheduling with QoS constraints over Rayleigh fading channels," in *Proc. IEEE Conf. Decision Control*, Dec. 2004, vol. 4, pp. 3447–3452.
- [28] V. J. Venkataramanan and X. Lin, "On wireless scheduling algorithms for minimizing the queue-overflow probability," Purdue Univ., Tech. Rep., 2008 [Online]. Available: <http://min.ecn.purdue.edu/~linx/papers.html>
- [29] D. Bertsimas, I. C. Paschalidis, and J. N. Tsitsiklis, "Asymptotic buffer overflow probabilities in multiclass multiplexers: An optimal control approach," *IEEE Trans. Autom. Control*, vol. 43, no. 3, pp. 315–335, Mar. 1998.



V. J. Venkataramanan received the B.Tech. degree in electrical engineering from the Indian Institute of Technology, Madras, India, in 2006.

He is currently pursuing the Ph.D. degree in electrical and computer engineering at Purdue University, West Lafayette, IN. His research interests are in mathematical modeling and evaluation of communication networks, including applied probability theory and optimization.



Xiaojun Lin (S'02–M'05) received the B.S. from Zhongshan University, Guangzhou, China, in 1994, and the M.S. and Ph.D. degrees from Purdue University, West Lafayette, IN, in 2000 and 2005, respectively.

He is currently an Assistant Professor of electrical and computer engineering at Purdue University. His research interests are resource allocation, optimization, network pricing, routing, congestion control, network as a large system, cross-layer design in wireless networks, and mobile ad hoc and sensor networks.

Dr. Lin received the IEEE INFOCOM 2008 Best Paper Award and 2005 Best Paper of the Year Award from the *Journal of Communications and Networks*. His paper was also one of two runner-up papers for the Best Paper Award at the IEEE INFOCOM 2005. He received the NSF CAREER Award in 2007. He was the Workshop Co-Chair for the IEEE GLOBECOM 2007, the Panel Co-Chair for WICON 2008, and the Technical Program Committee Co-Chair for ACM MobiHoc 2009.