

Delay Analysis of Maximum Weight Scheduling in Wireless Ad Hoc Networks

Long Bao Le, Krishna Jagannathan, and Eytan Modiano

Abstract—This paper studies delay properties of the well-known maximum weight scheduling algorithm in wireless ad hoc networks. We consider wireless networks with either one-hop or multi-hop flows. Specifically, this paper shows that the maximum weight scheduling algorithm achieves order optimal delay for wireless ad hoc networks with single-hop traffic flows if the number of activated links in one typical schedule is of the same order as the number of links in the network. This condition would be satisfied for most practical wireless networks. This result holds for both i.i.d and Markov modulated arrival processes with two states. For the multi-hop flow case, we also derive tight backlog bounds in the order sense.

Index Terms—Maximum weight scheduling, backlog/delay bounds, capacity region, order optimal delay

I. INTRODUCTION

Wireless scheduling has been known to be a key problem for throughput/capacity optimization in wireless networks. The well-known maximum weight scheduling algorithm has been proposed by Tassiulas in his seminal paper [1] where he proved its throughput optimality. Latter developments in this area include extension of this maximum weight scheduling algorithm to wireless networks with rate/power control [2], [3], network control when offered traffic is outside the capacity region [4], and other scheduling policies with lower-complexity [5]-[8]. While most existing works in the area of stochastic network control focused on throughput performance of optimal and suboptimal scheduling policies, delay properties of most scheduling policies proposed for wireless ad hoc networks remain unknown. In this paper, we study backlog/delay properties of the *maximum weight* scheduling algorithm in wireless ad hoc networks.

There are some recent works which investigated backlog/delay bounds for the supoptimal *maximal scheduling* algorithm in wireless ad hoc networks and *maximum weight* scheduling algorithm in the downlink/uplink of cellular networks. Specifically, in [13] Neely showed that maximal scheduling achieves delay scaling of $O(1/(1-\rho))$ for traffic inside the reduced stability region derived in [8]. This reduced stability region can be as small as $1/I$ of the capacity region, where I is the maximum number of links in any link interference set which do not interfere with one another. In [14], [15], Neely also proved the “order optimal” delay for the maximum weight scheduling algorithm in the wireless cellular uplink/downlink with ON/OFF wireless links. Note that the capacity region in the cellular setting can be explicitly

described which significantly eases the backlog/delay analysis. Average backlog bounds were derived for maximum weight scheduling in several works [2], [4], [10], [9]. These backlog bounds were obtained by bounding the maximum transmission rates and the number of arrivals in each time slot which are, therefore, not tight in general. There are some other works which investigate exponents for the tails of queue backlogs in the wireless cellular setting [11], [12].

In this paper, we consider a wireless ad hoc network with either one-hop and multi-hop traffic flows. We show that average delay for the case of one-hop traffic flows scales as $O(1/(1-\rho))$ if we can construct a set of distinct schedules to cover the network where the number of activated links in each of these schedules is of the same order as the number of network links. This condition would be satisfied for most practical large-scale wireless networks. This delay scaling holds for both i.i.d and Markov modulated traffic arrival processes with at most two states. These results are stated in Propositions 4 and 5 of the paper. To the best of our knowledge, these are the first delay optimal results for the maximum weight scheduling algorithm in wireless ad hoc networks. For wireless ad hoc networks with multihop traffic flows, we also derive a tight backlog bound which scales as $O(N/(1-\rho))$ where N is the number of wireless nodes.

The remaining of this paper is organized as follows. Delay analysis for single-hop traffic flows is presented in section II. In section III, we derive backlog bounds for wireless networks with multihop traffic flows.

II. ANALYSIS OF SINGLE-HOP FLOW CASE

A. System Models and Assumptions

We model a wireless ad hoc network as directed graph $G = (V, E)$ where V is the set of wireless nodes and E is the set of wireless links. Suppose the cardinalities of V and E are N and L , respectively. We consider single-hop traffic flows in this section. Data from all flows traversing a particular link l is buffered at the corresponding transmitter of the link. Assume time is slotted with fixed-size slot intervals. For now, traffic arriving to source nodes of single-hop flows is assumed to be independent and identically distributed (i.i.d) over time.

Assume that packets arriving during time slot t can only be transmitted from time slot $t+1$ at the earliest. Let denote by $A_l(t)$ the number of packets arriving at link l in time slot t and $\mu_l(t)$ the number of packet transmitted on link l in time slot t . For simplicity, assume that $\mu_l(t) = 1$ if link l is scheduled in time slot t , otherwise $\mu_l(t) = 0$. In the remaining of this paper, we will use $\vec{r}$ to describe a column vector with elements r_l denoting quantities such as queue length, scheduled links,

This work was supported by NSERC Postdoctoral Fellowship and by ARO MURI grant number W911NF-08-1-0238. The authors are with Communications and Networking Research Group, Massachusetts Institute of Technology, Cambridge, MA, USA. Emails: {longble, krishnaj, modiano}@mit.edu.
978-1-4244-2734-5/09/\$25.00 ©2009 IEEE

etc. The queue evolution for the flow at link l can be written as follows:

$$Q_l(t+1) = Q_l(t) - \mu_l(t) + A_l(t). \quad (1)$$

Assume that only backlogged links are scheduled, it can be verified that this queue evolution equation holds for arbitrary $Q_l(t)$ and $\mu_l(t)$. Regarding the scheduling, we consider the well-known maximum weight scheduling algorithm which is known to achieve the capacity region [1]. The maximum weight scheduling algorithm determines the optimal schedule $\vec{\mu}^*(t)$ based on the link queue backlogs as follows:

$$\vec{\mu}^*(t) = \operatorname{argmax}_{\vec{\mu}(t) \in S} \sum_{l=1}^L Q_l(t) \mu_l(t) \quad (2)$$

where S denotes the set of all possible feasible schedules according some interference constraints. In this section, we are going to derive the delay bound for this scheduling policy assuming that arrival traffic is strictly within the capacity region so that the maximum weight scheduling algorithm will stabilize the network [1]-[3].

B. Backlog/Delay Analysis for i.i.d Arrival Traffic

In this subsection, we obtain a delay bound for the aforementioned scheduling scheme using the Lyapunov drift technique [1]-[3]. Traffic arriving to a transmitting buffer of wireless link l is assumed to be i.i.d over time with average arrival rate λ_l . In the following, we use a result which was stated in [13].

Lemma 1: (Theorem 1 from [13]) Let $\vec{Q}(t)$ be the queue backlog vector in time slot t and $L(\vec{Q}(t))$ be a Lyapunov function. Also, define a one-step Lyapunov drift as follows:

$$\Delta(t) \triangleq E \left\{ L(\vec{Q}(t+1)) - L(\vec{Q}(t)) \right\} \quad (3)$$

where the expectation $E(\cdot)$ is taken over the randomness of queue backlogs $\vec{Q}(t)$ and system dynamics given the queue backlogs $\vec{Q}(t)$. If the Lyapunov drift satisfies

$$\Delta(t) \leq E \{g(t)\} - E \{f(t)\} \quad (4)$$

then we have

$$\limsup_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=0}^{t-1} E \{f(t)\} \leq \limsup_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=0}^{t-1} E \{g(t)\}. \quad (5)$$

Now, consider the following quadratic Lyapunov function

$$L(\vec{Q}(t)) \triangleq \sum_{l=1}^L Q_l^2(t). \quad (6)$$

We have the following result for the Lyapunov drift.

Proposition 1: The Lyapunov drift satisfies the following relation for any time slot t :

$$\Delta(t) = E \{B(t)\} + 2 \sum_{l=1}^L E \{Q_l(t) (A_l(t) - \mu_l(t))\} \quad (7)$$

where

$$B(t) \triangleq [A_l(t) - \mu_l(t)]^2 = A_l(t)^2 + 2A_l(t)\mu_l(t) + \mu_l(t)^2. \quad (8)$$

The proof of this proposition follows directly by applying the queue evolution relation in (1), and is omitted for brevity. As shown in [1], the capacity region coincides with the convex hull of all possible feasible schedules. Let $S = \{\vec{R}_i\}$ be the set of all possible schedules where one particular schedule $\vec{R}_i$ is a column vector of dimension L with the l -th element equal 1 if link l is scheduled and equal 0 otherwise. For any arrival rate vector $\vec{\lambda}$ strictly inside the capacity region, we have

$$\vec{\lambda} \leq \sum_{i=1}^{|S|} \beta_i \vec{R}_i \quad (9)$$

where $|S|$ denotes the cardinality of set S , $\sum_{i=1}^{|S|} \beta_i < 1$ and " $<, \leq$ " denotes both a regular inequality and an element-wise inequality. We have the following relation

$$\sum_{l=1}^N Q_l \lambda_l = \vec{Q}^{tr} \vec{\lambda} \leq \sum_{i=1}^{|S|} \beta_i \vec{Q}^{tr} \vec{R}_i \leq \vec{Q}^{tr} \vec{\mu}^* \sum_{i=1}^{|S|} \beta_i < \vec{Q}^{tr} \vec{\mu}^* \quad (10)$$

where $[\cdot]^{tr}$ denotes the vector transposition and $\vec{\mu}^*$ is the optimal scheduled vector given backlog vector $\vec{Q}$. It can be verified that these results hold by using the relations in (9) and (2). Now, we state a bound on the total queue backlogs for i.i.d. arrival traffic in the following proposition.

Proposition 2: Assume that the arrival rate vector $\vec{\lambda}$ is strictly inside the capacity region so that there exists a vector $\vec{\epsilon}$ such that $\vec{\lambda} + \vec{\epsilon}$ is inside the capacity region where $\vec{\epsilon}$ is a vector with all elements equal to ϵ . Also, assume that all arrival processes on all wireless links have bounded second moments. Then, the network is stable and the total average queue backlog can be bounded as

$$\sum_{l=1}^L \bar{Q}_l \leq \frac{\lambda_{tot} + \sum_{l=1}^L E \{A_l(t)^2\} - 2 \sum_{l=1}^L \lambda_l^2}{2\epsilon} \quad (11)$$

where $\lambda_{tot} \triangleq \sum_{l=1}^L \lambda_l$ is the total link arrival rates.

Proof: Using (10) for backlog vector $\vec{Q}(t)$, we have

$$[\vec{Q}(t)]^{tr} [\vec{\lambda} + \vec{\epsilon}] \leq [\vec{Q}(t)]^{tr} \vec{\mu}^*(t). \quad (12)$$

Hence,

$$[\vec{Q}(t)]^{tr} \vec{\lambda} - [\vec{Q}(t)]^{tr} \vec{\mu}^*(t) \leq -[\vec{Q}(t)]^{tr} \vec{\epsilon}. \quad (13)$$

Note that the second term of (7) can be written as

$$\begin{aligned} & 2 \sum_{l=1}^L E \{Q_l(t) (A_l(t) - \mu_l(t))\} \\ &= 2 \sum_{l=1}^L E \{Q_l(t) (\lambda_l - \mu_l(t))\} \\ &= 2E \left\{ [\vec{Q}(t)]^{tr} (\vec{\lambda} - \vec{\mu}(t)) \right\}. \end{aligned} \quad (14)$$

Using (13) and (14) in (7) with $\vec{\mu}(t)$ representing an optimal scheduled vector, we have

$$\Delta(t) \leq E \{B(t)\} - 2 [\vec{Q}(t)]^{tr} \vec{\epsilon}$$

$$= E\{B(t)\} - 2\epsilon \sum_{l=1}^L E\{Q_l(t)\}. \quad (15)$$

Using the result in Lemma 1 in (15), we have

$$\limsup_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=0}^{t-1} \sum_{l=1}^L E\{Q_l(\tau)\} \leq \frac{\bar{B}}{2\epsilon} \quad (16)$$

where $\bar{B} \triangleq \limsup_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=0}^{t-1} E\{B(\tau)\}$.

From (8), using the fact that $\mu_l(t) \leq 1$ and arrival processes have bounded second moments, we have $\bar{B} < \infty$. Therefore, the queueing network is strongly stable. Because it evolves according to an ergodic Markov chain with countable state space, the limiting time averages of queue backlogs equal to the corresponding steady state averages. To calculate $\bar{B}$, we note that under the stability condition we have $\limsup_{t \rightarrow \infty} \frac{1}{t} \sum_{\tau=0}^{t-1} \mu_l(\tau) = \lambda_l$. Also, note that $\mu_l(t)^2 = \mu_l(t)$ because $\mu_l(t) = \{0, 1\}$ depending on whether link l is scheduled in time slot t or not. As a consequence, $\bar{B}$ can be written as

$$\begin{aligned} \bar{B} &= \sum_{l=1}^L E\{A_l(t)^2\} - 2 \sum_{l=1}^L \lambda_l^2 + \sum_{l=1}^L \lambda_l \\ &= \lambda_{tot} + \sum_{l=1}^L E\{A_l(t)^2\} - 2 \sum_{l=1}^L \lambda_l^2. \end{aligned} \quad (17)$$

Because time average limits of queue backlogs are equal to their steady state averages, using (17) the inequality (16) can be rewritten as

$$\sum_{l=1}^L \bar{Q}_l \leq \frac{\lambda_{tot} + \sum_{l=1}^L E\{A_l(t)^2\} - 2 \sum_{l=1}^L \lambda_l^2}{2\epsilon}. \quad (18)$$

Hence, the proposition is proved. ■

1) *Delay Bound:* Applying Little's law to (11), we can obtain a delay bound as follows:

$$\bar{W} \leq \frac{1 + \frac{1}{\lambda_{tot}} \sum_{l=1}^L [E\{A_l(t)^2\} - 2\lambda_l^2]}{2\epsilon}. \quad (19)$$

Now, in order to understand the scaling of this delay bound, we need to determine the relationship between the "traffic loading factor" ρ and the parameter ϵ . Let us denote by Λ the capacity region. Assume that the arrival rate vector $\vec{\lambda} = (\lambda_1, \lambda_2, \dots, \lambda_L)^{tr}$ is strictly inside the capacity region Λ , then there exists a loading factor $\rho < 1$ such that

$$\vec{\lambda} \in \rho\Lambda. \quad (20)$$

In the following, we state a delay bound by choosing a straightforward loading factor ρ as a function of ϵ .

Proposition 3: If arrival rate is in ρ -scaled capacity region as described in (20), the average total delay can be bounded as

$$\bar{W} \leq \frac{L \left[1 + \frac{1}{\lambda_{tot}} \sum_{l=1}^L [E\{A_l(t)^2\} - 2\lambda_l^2] \right]}{2(1-\rho)}. \quad (21)$$

In the special case where the arrival process on each wireless link is Poisson, we have

$$\bar{W} \leq \frac{L(1 - \frac{1}{2\lambda_{tot}} \sum_{l=1}^L \lambda_l^2)}{1-\rho}. \quad (22)$$

Proof: The proof follows by using the fact that we can choose $\vec{\epsilon} = \frac{(1-\rho)}{L} \mathbf{1}_L$ where $\mathbf{1}_L$ is an all-one vector with dimension L such that $\vec{\lambda} + \vec{\epsilon} \in \Lambda$ for any $\vec{\lambda} \in \rho\Lambda$. Specifically, by plugging $\epsilon = \frac{(1-\rho)}{L}$ into the delay bound in (19), we can obtain (21). Now, we show that $\vec{\lambda} + \vec{\epsilon} \in \Lambda$ for $\epsilon = \frac{(1-\rho)}{L}$. Note that for any $\vec{\lambda} \in \rho\Lambda$, we can write $\vec{\lambda} = \rho \sum_{i=1}^{|S|} \beta_i \vec{R}_i$ where $\sum_{i=1}^{|S|} \beta_i < 1$. Define $\vec{e}_i$ be a vector of dimension L with all zeros except a one at the i -th position. It can be easily seen that any $\vec{e}_i$ ($i = 1, 2, \dots, L$) represents a feasible schedule (with link i being activated). Also, note that $\sum_{i=1}^L \vec{e}_i = \mathbf{1}_L$. Hence, we have the following result

$$\frac{1-\rho}{L} \mathbf{1}_L + \rho \sum_{i=1}^{|S|} \beta_i \vec{R}_i = \frac{1-\rho}{L} \sum_{i=1}^L \vec{e}_i + \rho \sum_{i=1}^{|S|} \beta_i \vec{R}_i \in \Lambda. \quad (23)$$

When the arrival processes are Poisson, we have $E\{A_l(t)^2\} = \lambda_l + \lambda_l^2$. Using this relationship in the delay bound (21), we obtain (22). ■

Note that the term $\frac{1}{\lambda_{tot}} \sum_{l \in E} E\{A_l(t)^2\}$ is typically $O(1)$ for any traffic satisfying $A_l(t) \leq A_{\max}$. In fact, in such cases we have $\frac{1}{\lambda_{tot}} \sum_{l \in E} E\{A_l(t)^2\} \leq A_{\max}$. Hence, the delay bound stated in Proposition 3 is typically $O(L/(1-\rho))$.

C. Tighter Delay Bound

In the following, we state a tighter delay bound under specific assumptions which can be achieved by exploiting underlying interference constraints and network topology.

Proposition 4: Assume that the arrival rate is in the ρ -scaled capacity region as described in (20). Also, assume that we can find a set of feasible schedules, namely $\Psi = \{\vec{s}_i, i = 1, 2, \dots, T\}$, satisfying the following assumptions

- For any schedule $\vec{s}_i \in \Psi$, if link l is activated in $\vec{s}_i$ then link l is not activated in any other $\vec{s}_j \in \Psi$ for $j \neq i$ (i.e., any link should belong to one and only one schedule in the set Ψ).
- Let E^* be the set of links activated by all schedules in Ψ , then $E^* = E$ where recall that E is the set of all network links (i.e., the union of activated links by all schedules covers the whole network).

Let K_i denote the number of activated links in schedule $\vec{s}_i$ and $K_{\min} = \min_i K_i$. Then, we have the following delay bound

$$\bar{W} \leq \frac{1 + \frac{1}{\lambda_{tot}} \sum_{l=1}^L [E\{A_l(t)^2\} - 2\lambda_l^2]}{2K_{\min}(1-\rho)/L}. \quad (24)$$

Before proving this proposition, we note that for wireless networks such that $K_{\min} = O(L)$, proposition 4 implies that the network delay typically scales as $O(1/(1-\rho))$. This condition would hold if the network topology is sufficiently *sparse and uniform* so that the most balanced set of schedules Ψ (i.e., almost all schedules in Ψ have the same number of activated links in the order sense) satisfies $K_{\min} = O(L)$. Note that this condition would be satisfied for most practical wireless networks because a typical schedule would activate most links in the network. We will provide one such network example after the proof.

Proof: The proof for this proposition follows the same line as that for proposition 3. However, a tighter delay bound

is achieved in this proposition by constructing $\vec{\epsilon}$ from the set of schedules Ψ each of which has at least $K_{\min}$ activated links. Now, consider the following linear combination of feasible schedules whose outcome lies inside the capacity region

$$(1 - \rho) \sum_{i=1}^T \frac{K_i}{\sum_j K_j} \vec{s}_i + \rho \sum_{i=1}^{|S|} \beta_i \vec{R}_i \in \Lambda. \quad (25)$$

Therefore, the result stated in proposition 4 follows by plugging $\epsilon = (1 - \rho) \frac{K_{\min}}{L}$ into the delay bound (19). ■

In the following, we provide a simple example where the assumptions of the proposition hold.

Example: Consider a grid network and one-hop (primary) interference model for the sake of simplicity as being shown in Fig. 1. In this figure, we also show how to construct a set of feasible schedules Ψ that covers the whole network graph (again, each schedule has the same link pattern). To analyze its delay bound, assume that the size in one dimension of the grid network is H links, then it can be verified that $L = 2H(H+1)$. From the constructed set of schedules Ψ as shown in this figure, we have $K_{\min} = (H+1) \lfloor H/2 \rfloor$. Therefore, using the result in proposition 4, the delay can be bounded as

$$\begin{aligned} \bar{W} &\leq \frac{2H(H+1) \left[1 + \frac{1}{\lambda_{\text{tot}}} \sum_{l=1}^L [E\{A_l(t)^2\} - 2\lambda_l^2] \right]}{2(H+1) \lfloor H/2 \rfloor (1 - \rho)} \\ &\approx \frac{2 \left[1 + \frac{1}{\lambda_{\text{tot}}} \sum_{l=1}^L [E\{A_l(t)^2\} - 2\lambda_l^2] \right]}{(1 - \rho)} \end{aligned}$$

which scales typically as $O(1/(1 - \rho))$.

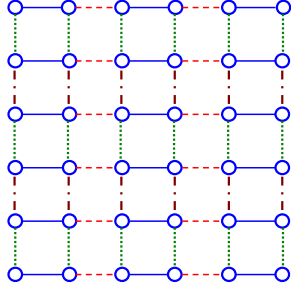


Fig. 1. Grid networks with one-hop (primary) interference model.

D. Analysis for Time-Correlated Arrivals with Two States

Here, assume that arrival process $A_l(t)$ for links l is either i.i.d. or modulated by a discrete time stationary and ergodic Markov chain $Z_l(t)$ having two states (i.e., states 1 and 2). Let σ_l and δ_l be transition probabilities from state 1 to state 2 and from state 2 to state 1, respectively. For each link l , define the conditional average arrival rates $\lambda_l^{(m)}$ as follows:

$$\lambda_l^{(m)} \triangleq E\{A_l(t) | Z_l(t) = m\}.$$

Now, let denote by $E_1 \subset E$ as the set of links with time-correlated arrivals where $\lambda_l^{(1)} \neq \lambda_l^{(2)}$. Also, assume that arrival traffic to any other links in $E_2 = E - E_1$ is either i.i.d. or time-correlated with two states satisfying $\lambda_l^{(1)} = \lambda_l^{(2)}$. Assume that the modulating Markov chains of all arrival processes are stationary so that for all links l we have $E\{A_l(t)\} = \lambda_l$ for

all time t . In order to obtain delay bound for this case, we will use one result proved in [13] which is stated in the following lemma.

Lemma 2: (from section V.A of [13]) Define $C_l = E\{A_l(t-1)A_l(t)\}$ for $l \in E_1$ and $C_l = 0$ for $l \in E_2$. For all link l , we have

$$E\{Q_l(t)A_l(t)\} \leq E\{Q_l(t)\} \lambda_l + \frac{C_l}{\delta_l + \sigma_l}.$$

Now, we state delay bounds for the case of time-correlated arrivals in the following proposition.

Proposition 5: If the arrival traffic is within the ρ -scaled capacity region and the assumptions in proposition 4 are satisfied, then the network is stable and the average delay can be bounded as

$$\bar{W} \leq \frac{\tilde{B} + \tilde{C}}{2K_{\min}(1 - \rho)/L}. \quad (26)$$

where

$$\tilde{B} \triangleq 1 + \frac{1}{\lambda_{\text{tot}}} \sum_{l \in E} E\{A_l(t)^2\}, \tilde{C} \triangleq \frac{1}{\lambda_{\text{tot}}} \sum_{l \in E_1} \frac{C_l}{\sigma_l + \delta_l}. \quad (27)$$

The proof follows by using results in Lemma 2 and Proposition 1 so it is omitted for brevity. The term $\frac{1}{\lambda_{\text{tot}}} \sum_{l \in E} E\{A_l(t)^2\}$ is typically $O(1)$ for any traffic satisfying $A_l(t) \leq A_{\max}$. In fact, in such cases we have $\tilde{B} \leq 1 + A_{\max}$. It is not very difficult to see that for $l \in E_1$, we have $C_l \leq \lambda_l \lambda_{\max}$ where $\lambda_{\max} < 1$ is the maximum conditional rate over all links and states. Hence, we can obtain the following delay bound

$$\bar{W} \leq \frac{1 + A_{\max} + \max_{l \in E_1} \{\lambda_{\max}/(\sigma_l + \delta_l)\}}{2K_{\min}(1 - \rho)/L} \quad (28)$$

which scales as $O(1/(1 - \rho))$ for $K_{\min} = O(L)$.

III. ANALYSIS OF MULTIHOP FLOW CASE

A. System Models and Assumptions

We consider the same network model as section II. We assume that there is set of multihop flows F where flow $f \in F$ has a fixed route from a source node $s(f)$ to a destination node $d(f)$. We denote the set of links and nodes on the route of flow f as $L(f)$ and $R(f)$, respectively. For simplicity, we assume that packet arrivals to source nodes of all flows are i.i.d stochastic processes.

We denote the queue length of flow f at node n at the beginning of time slot t as $Q_n^f(t)$ and the number of packets arriving at the source node of flow f as $A_{s(f)}^f(t)$. Note that data packets of any flow are delivered to the higher layer upon reaching the destination node, so $Q_{d(f)}^f(t) = 0$. In addition, let $\mu_n^f(t)$ be the number of packets of flow f transmitted from node n along link (n, m) of its route which is buffered at node m if $m \neq d(f)$. Again, we assume that $\mu_n^f(t) = 1$ if we activate link (n, m) on the route of flow f and $\mu_n^f(t) = 0$, otherwise. Given the routes for all flows, the maximum weight scheduling algorithm is used for data delivery [1]. Specifically, the scheduling is performed in every time slot as follows:

- Each link (n, m) finds the maximum differential backlogs as follows:

$$w_{nm}(t) = \max_{f:(n,m) \in L(f)} \{Q_n^f(t) - Q_m^f(t)\}. \quad (29)$$

- Based on calculated link weights, a maximum weight schedule is found as

$$\vec{\mu}^*(t) = \operatorname{argmax}_{\vec{\mu}(t) \in S} \sum_{(n,m)} w_{nm}(t) \mu_{nm}(t). \quad (30)$$

- For any scheduled link, one packet is transmitted from the buffer of the flow achieving the maximum differential backlog.

The queue evolutions can be written as

$$Q_n^f(t+1) = Q_n^f(t) - \mu_n^f(t) + \pi_n^f(t) \quad (31)$$

where this equation holds because $\mu_n^f(t) = 1$ only if $Q_n^f(t) \geq 1$ (i.e., we do not schedule links with empty queues). Also, $\pi_n^f(t)$ is the number of packets arriving to queue $Q_n^f(t)$ in time slot t which can be written as

$$\pi_n^f(t) = \begin{cases} A_n^f(t), & \text{if } n = s(f) \\ \mu_{n-1}^f(t), & \text{otherwise.} \end{cases} \quad (32)$$

B. Backlog Bound

The queue backlog bound is stated in the following proposition. The tighter backlog bound will be developed after that.

Proposition 6: Assume that arrival processes for all flows are i.i.d over slots and have bounded second moments. Also, assume that arrival rate lies within the ρ -scaled capacity region. Then,

- 1) The total average queue backlogs can be bounded as

$$\sum_{f=1}^{|F|} \sum_{n \in R(f)} \bar{Q}_n^f(t) \leq \frac{L\theta_{\max}L_{\max}(N+D)}{2(1-\rho)}.$$

where ρ is a loading factor, $L_{\max} \triangleq \max_{f \in F} |L(f)|$, $\theta_{\max}$ is the maximum number of flows traversing any link in the network, N is total number of network nodes, and

$$D \triangleq \sum_{f=1}^{|F|} E \left\{ \left(A_{s(f)}^f(t) \right)^2 - 2\lambda_f^2 + \lambda_f \right\}. \quad (33)$$

- 2) For Poisson arrival process, the backlog can be bounded as

$$\sum_{f=1}^{|F|} \sum_{n \in R(f)} \bar{Q}_n^f(t) \leq \frac{L\theta_{\max}L_{\max} \left(N + 2\lambda_{tot} - \sum_{f \in F} \lambda_f^2 \right)}{2(1-\rho)}$$

where $\lambda_{tot} = \sum_{f \in F} \lambda_f$.

Proof: Consider the following Lyapunov function

$$L(\vec{Q}(t)) \triangleq \sum_{f=1}^{|F|} \sum_{n \in R(f)} (Q_n^f(t))^2 \quad (34)$$

where again the expectation $E(\cdot)$ is taken over the randomness of queue backlogs $\vec{Q}(t)$ and system dynamics given the queue backlogs $\vec{Q}(t)$. The Lyapunov drift can be written as follows:

$$\Delta(t) = E \left\{ \sum_{f=1}^{|F|} \sum_{n \in R(f)} \left[(Q_n^f(t+1))^2 - (Q_n^f(t))^2 \right] \right\}$$

978-1-4244-2734-5/09/\$25.00 ©2009 IEEE

$$= E \left\{ \sum_{f=1}^{|F|} \sum_{n \in R(f)} [\pi_n^f(t) - \mu_n^f(t)]^2 \right\} \quad (35)$$

$$+ 2E \left\{ \sum_{f=1}^{|F|} \sum_{n \in R(f)} Q_n^f(t) [\pi_n^f(t) - \mu_n^f(t)] \right\} \quad (36)$$

Now, consider (35) we have

$$E \left\{ \sum_{f=1}^{|F|} \sum_{n \in R(f)} [\pi_n^f(t) - \mu_n^f(t)]^2 \right\} \leq N + B_1(t) \quad (37)$$

where $B_1(t) \triangleq \sum_{f=1}^{|F|} E \left\{ [A_{s(f)}^f(t) - \mu_{s(f)}^f(t)]^2 \right\}$. This inequality holds because we have $\sum_{f \in F, s(f) \neq n} [\pi_n^f(t) - \mu_n^f(t)]^2 \leq 1$. Now, using (32), we can manipulate (36) as follows:

$$\begin{aligned} & 2E \left\{ \sum_{f=1}^{|F|} \sum_{n \in R(f)} Q_n^f(t) [\pi_n^f(t) - \mu_n^f(t)] \right\} \\ &= -2 \sum_{(n,m) \in E} \mu_{nm}^*(t) w_{nm}^* + 2 \sum_{f=1}^{|F|} Q_{s(f)}^f(t) \lambda_f \end{aligned} \quad (38)$$

where $\mu_{nm}^*(t)$ corresponds to the maximum weight schedule with queue backlogs $\vec{Q}(t)$ and $w_{nm}^* = \max_{f:(n,m) \in L(f)} [Q_n^f(t) - Q_m^f(t)]^+$. Note that we have written down $[Q_n^f(t) - Q_m^f(t)]^+$ instead of $[Q_n^f(t) - Q_m^f(t)]$ because links with negative weight will not be scheduled by the maximum weight scheduling algorithm. Note that we can rewrite $\sum_{f=1}^{|F|} Q_{s(f)}^f(t) \lambda_f$ as follows:

$$\begin{aligned} 2 \sum_{f=1}^{|F|} Q_{s(f)}^f(t) \lambda_f &= 2 \sum_{f=1}^{|F|} \sum_{(n,m) \in L(f)} \lambda_f [Q_n^f(t) - Q_m^f(t)] \\ &\leq 2 \sum_{(n,m) \in E} \left(\sum_{f:(n,m) \in L(f)} \lambda_f \right) \\ &\quad \times \max_{f:(n,m) \in L(f)} [Q_n^f(t) - Q_m^f(t)]^+ \end{aligned} \quad (39)$$

Suppose that $\vec{\lambda} = (\lambda_f)_{f \in F}$ is strictly inside the capacity region. Then, there exists a vector $\vec{\epsilon}$ such that $\vec{\lambda} + \vec{\epsilon}$ is still inside the capacity region. This implies that there exists a vector of link rates $(\mu_{nm})_{(n,m) \in E} \in Co(S)$ such that [1]

$$\sum_{f:(n,m) \in L(f)} (\lambda_f + \epsilon) = \sum_{f:(n,m) \in L(f)} \lambda_f + \theta_{nm} \epsilon \leq \mu_{nm} \quad (40)$$

where $Co(S)$ represents the convex hull of all feasible link schedules and θ_{nm} denotes the number of flows traversing link (n, m) . Using (38), (39), (40), we can rewrite (36) as

$$\begin{aligned} (36) &\leq 2 \sum_{(n,m) \in E} \left(-\mu_{nm}^*(t) + \sum_{f:(n,m) \in L(f)} \lambda_f \right) w_{nm}^* \\ &\leq 2 \sum_{(n,m) \in E} (-\mu_{nm}^*(t) + \mu_{nm} - \theta_{nm} \epsilon) w_{nm}^* \\ &\leq -2\epsilon \sum_{(n,m) \in E} \theta_{nm} w_{nm}^*. \end{aligned} \quad (41)$$

Now, we have the following

$$\begin{aligned} \sum_{(n,m) \in L(f)} Q_n^f(t) &\leq |L(f)| \sum_{(n,m) \in L(f)} [Q_n^f(t) - Q_m^f(t)]^+ \\ &\leq |L(f)| \sum_{(n,m) \in L(f)} w_{nm}^*. \end{aligned} \quad (42)$$

Note that a similar bounding technique for $\sum_{(n,m) \in L(f)} Q_n^f(t)$ used in the above inequality has been employed in [16] to prove the network stability. We instead aim at obtaining a tight upper-bound of queue backlogs as a function of the loading factor ρ , so the Lyapunov drift analysis in this paper is very different from that in [16]. Using (42), we have

$$\begin{aligned} \sum_{f=1}^{|F|} \sum_{(n,m) \in L(f)} Q_n^f(t) &\leq \sum_{f=1}^{|F|} |L(f)| \sum_{(n,m) \in L(f)} w_{nm}^* \\ &\leq L_{\max} \sum_{f=1}^{|F|} \sum_{(n,m) \in L(f)} w_{nm}^* = L_{\max} \sum_{(n,m) \in E} \theta_{n,m} w_{nm}^*. \end{aligned} \quad (43)$$

where $L_{\max} = \max_{f \in F} |L(f)|$. Using (37), (41) and (43), the Lyapunov drift can be bounded as

$$\Delta(t) \leq N + B_1(t) - \frac{2\epsilon}{L_{\max}} \sum_{f=1}^{|F|} \sum_{n \in R(f)} Q_n^f(t). \quad (44)$$

If the arrival processes for all flows have bounded second moments then $B_1(t)$ is bounded. Under this condition, the Lyapunov drift will be negative when the total backlog becomes large enough. Hence, the network is stable and the time average limits of queue backlogs are equal to their steady-state averages. Also, it is not very difficult to see that the time average limit of $B_1(t)$ is equal to D . Therefore, by using Lemma 1 and substituting time average limits of queue backlogs by their steady-state averages, we have

$$\sum_{f=1}^{|F|} \sum_{n \in R(f)} \bar{Q}_n^f \leq \frac{L_{\max}(N + D)}{2\epsilon}. \quad (45)$$

Now, to understand the scaling of this backlog bound, we need to find ϵ as a function of the loading factor ρ as before. Suppose arrival rate vector $\vec{\lambda}$ is inside the ρ -scaled capacity region, and let $\vec{\lambda}'$ be a vector with (n, m) -th elements equal $\lambda_{nm} = \sum_{f: (n,m) \in L(f)} \lambda_f$. Then, we have

$$\vec{\lambda}' = \rho \sum_{i=1}^{|S|} \beta_i \vec{R}_i \quad (46)$$

for some non-negative β_i such that $\sum_{i=1}^{|S|} \beta_i < 1$. Let $\theta_{\max} = \max_{(n,m) \in E} \theta_{nm}$ where recall that θ_{nm} is the number of flows traversing link (n, m) . We will show that $\vec{\lambda}^{(1)} = \vec{\lambda} + \vec{\epsilon} \in \Lambda$ for $\epsilon = (1 - \rho)/(L\theta_{\max})$. Now, let us construct a vector $\vec{\lambda}^{(2)}$ with its (n, m) -th elements equal to $\bar{\lambda}_{nm}^{(2)} = \sum_{f: (n,m) \in L(f)} \lambda_f^{(1)}$. Then, we have

$$\vec{\lambda}^{(2)} = \vec{\lambda}' + \left(\frac{1 - \rho}{L\theta_{\max}} \theta_{nm} \right)_{(n,m) \in E}$$

$$\leq \rho \sum_{i=1}^{|S|} \beta_i \vec{R}_i + \frac{1 - \rho}{L} \sum_{(n,m) \in E} \vec{\epsilon}_{nm} \in \Lambda \quad (47)$$

where $\left(\frac{1 - \rho}{L\theta_{\max}} \theta_{nm} \right)_{(n,m) \in E}$ denotes a vector with (n, m) -th elements equal to the quantity inside the bracket.

Substitute $\epsilon = (1 - \rho)/(L\theta_{\max})$ into (45), we can obtain part 1) of the proposition. To prove part 2) of the proposition, we need some manipulation of D for the Poisson arrival process.

Specifically, for Poisson process we have $E \left\{ \left(A_{s(f)}^f(t) \right)^2 \right\} = \lambda_f + \lambda_f^2$. Substitute this into (33), we have $D = 2\lambda_{tot} - \sum_{f=1}^{|F|} \lambda_f^2$. Plug this into the bound in part 1), we obtain part 2) of the proposition. ■

It can be shown that the average backlogs derived in this proposition scale as $O(LN/(1 - \rho))$ for $L_{\max} = O(1)$. In fact, using the similar idea as that of proposition 4, a tighter backlog bound can be obtained. Specifically, if the arrival processes satisfy assumptions in proposition 6 and the assumptions of proposition 4 hold, then the total queue backlogs can be bounded as

$$\sum_{f=1}^{|F|} \sum_{n \in R(f)} \bar{Q}_n^f \leq \frac{\theta_{\max} L_{\max} (N + D)}{2K_{\min}(1 - \rho)/L}. \quad (48)$$

This backlog bound typically scales as $O(N/(1 - \rho))$ for $L_{\max} = O(1)$ and $K_{\min} = O(L)$.

REFERENCES

- [1] L. Tassiulas and A. Ephremides, "Stability properties of constrained queueing systems and scheduling policies for maximum throughput in multihop radio networks," *IEEE Trans. Automatic Control*, vol. 37, no. 12, pp. 1936-1948, Dec. 1992.
- [2] M. Neely, E. Modiano, and C. Rohrs, "Power allocation and routing in multibeam satellites with time-varying channels," *IEEE/ACM Trans. Networking*, vol. 11, no. 1, pp. 138-152, Feb. 2003.
- [3] M. Neely, E. Modiano, and C. Rohrs, "Dynamic power allocation and routing for time varying wireless networks," *IEEE INFOCOM 2003*.
- [4] M. Neely, E. Modiano, and C. Li, "Fairness and optimal stochastic control for heterogeneous networks," *IEEE INFOCOM 2005*.
- [5] E. Modiano, D. Shah, and G. Zussman, "Maximizing throughput in wireless networks via gossiping," *ACM SIGMETRICS 2006*.
- [6] S. Sanghavi, L. Bui, and R. Srikant, "Distributed link scheduling with constant overhead," *ACM SIGMETRICS 2007*, June 2007.
- [7] X. Lin and N. Shroff, "The impact of imperfect scheduling on cross-layer rate control in wireless networks," *IEEE INFOCOM 2005*.
- [8] P. Charporkar, K. Kar, and S. Sarkar, "Throughput guarantees through maximal scheduling in wireless networks," *Allerton 2005*, Sept. 2005.
- [9] Y. Yi, A. Proutiere, and M. Chiang, "Complexity of wireless scheduling: Impact and tradeoffs," *ACM Mobihoc*, May 2008.
- [10] G. Gupta and N. B. Shroff, "Delay analysis of scheduling policies in wireless networks," *Asilomar Conference on Signals, Systems, and Computers*, Oct. 2008.
- [11] A. Stolyar, "Large deviations of queues under QoS scheduling algorithms," *Allerton 2006*, Sept. 2006.
- [12] V. J. Venkataramanan and X. Lin, "Structural properties of LDP for queue-length based wireless scheduling algorithms," *Allerton 2007*, Sept. 2007.
- [13] M. Neely, "Delay analysis for maximal scheduling in wireless networks with bursty traffic," *IEEE INFOCOM 2008*.
- [14] M. Neely, "Order optimal delay for opportunistic scheduling in multi-user wireless uplinks and downlinks," *Allerton 2006*, Sept. 2006.
- [15] M. Neely, "Delay analysis for max weight opportunistic scheduling in wireless systems," *Allerton 2008*, Sept. 2008.
- [16] L. Bui, R. Srikant, and A. Stolyar, "Novel architecture and algorithms for delay reduction in back-pressure scheduling and routing," *IEEE INFOCOM 2009*.