

Throughput Optimal Distributed Power Control of Stochastic Wireless Networks

Yufang Xi, *Student Member, IEEE*, and Edmund M. Yeh, *Member, IEEE*

Abstract—The maximum differential backlog (MDB), or “backpressure” control policy of Tassiulas and Ephremides has been shown to adaptively maximize the stable throughput of multihop wireless networks with random traffic arrivals and queueing. The practical implementation of the MDB policy in wireless networks with mutually interfering links, however, requires the development of distributed optimization algorithms. Within the context of code-division multiple-access (CDMA)-based multihop wireless networks, we develop a set of node-based scaled gradient projection power control algorithms which solves the MDB optimization problem based on the high-signal-to-interference-plus-noise ratio (SINR) approximation of link capacities using low communication overhead. We investigate the impact of the high-SINR approximation and the nonnegligible convergence time required by the power control algorithms on the throughput region achievable by the iterative MDB policy. We show that the policy can achieve at least the stability region induced by the high-SINR capacity region.

Index Terms—Distributed optimization, multihop wireless networks, stochastic control.

I. INTRODUCTION

THE optimal control of multihop wireless networks is a major research and design challenge due, in part, to the interference between nodes, the time-varying nature of the communication channels, the energy limitation of mobile nodes, and the lack of centralized coordination. This problem is further complicated by the fact that data traffic in wireless networks often arrive at random instants into network buffers. Although a complete solution to the optimal control problem is still elusive, a major advance is made in the seminal work of Tassiulas and Ephremides [1]. In this work, the authors consider a stochastic multihop wireless network with random traffic arrivals and queueing, where the activation of links satisfies specified constraints reflecting, for instance, channel interference. For this network, the authors characterize the stability region, i.e., the set of all end-to-end demands that the network can support. Moreover, they obtain a *throughput optimal* routing and link activation policy which stabilizes the network whenever the arrival rates are in the interior of the stability region, without a

priori knowledge of arrival statistics. The throughput optimal policy operates on the Maximum differential backlog (MDB) principle, which essentially seeks to achieve load-balancing in the network. The MDB policy (sometimes called the “backpressure algorithm”) has been extended to multihop networks with general capacity constraints in [2] and has been combined with congestion control mechanisms in [3], [4].

While the MDB policy represents a remarkable achievement, there remains a significant difficulty in applying the policy to wireless networks. The mutual interference between wireless links imply that the evaluation of the MDB policy involves a centralized network optimization. It is shown in [5] that when there are finitely many feasible schedules, the centralized optimization can be approximated by a randomized algorithm with linear complexity while preserving throughput optimality. In wireless networks with limited transmission range and scarce battery resources, however, any centralized algorithm is undesirable. The call for more distributed scheduling algorithms with guaranteed throughput has given rise to two main lines of research.

One approach is to adopt simple physical and MAC layer models and apply computationally efficient scheduling rules in a distributed manner [6]–[10]. It is shown in [6] that Maximal Greedy Scheduling can achieve a guaranteed fraction of the maximum throughput region. This result is generalized in [7], [8] to multihop networks where the end-to-end paths are given and fixed. Despite its simplicity, the distributed scheduling considered in the above work applies to only a limited class of networks. Moreover, the simplicity is gained at the expense of throughput optimality [11]. Recently, the work in [10] developed a distributed implementation of Tassiulas’s linear complexity randomized algorithm [5] in networks with primary interference constraints. The scheme, which involves distributed randomized matching, schedule selection, and improvement, is shown to achieve a throughput region arbitrarily close to the optimal region.

Another line of research develops distributed power control and rate allocation algorithms for implementing the MDB policy in interference-limited networks with the aim of preserving the throughput optimality. Thus far, distributed MDB control has been investigated only for networks with relatively simple physical layer models. For example, Neely [12] studies a cell-partitioned network model where different cells do not interfere with each other so that scheduling can be decentralized to each cell. In general, the MDB policy for interference-limited networks has to choose from a continuum of rate vectors. Therefore, the randomized technique developed in [5], [10] for simple physical layer models is not readily applicable in this context. So far, the question of how the MDB policy can

Manuscript received September 18, 2006; revised May 02, 2009; approved by IEEE/ACM TRANSACTIONS ON NETWORKING Editor N. B. Shroff. This research is supported in part by the Army Research Office (ARO), Young Investigator Program (YIP) under Grant DAAD19-03-1-0229, and by the National Science Foundation (NSF) under Grant CCR-0313183.

The authors are with the Department of Electrical Engineering, Yale University, New Haven, CT 06520 USA (e-mail: yufang.xi@yale.edu; edmund.yeh@yale.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.2009.2035919

be efficiently applied in general interference-limited wireless networks remains unanswered.

In this paper, we consider the implementation of the MDB algorithm within interference-limited CDMA wireless networks, where all links can potentially be active at the same time, and transmission on any given link can interfere with transmissions on all other active links. We consider a control strategy where transmission powers and link service rates are dynamically adjusted according to the queue state in order to maximize network throughput. A major difficulty in realizing this goal is the fact that the physical-layer capacity region of the interference-limited code-division multiple-access (CDMA) network is in general nonconvex. To achieve all arrival rates in the stability region induced by the nonconvex physical-layer capacity region typically requires the use of scheduling, which is often not amenable to distributed and efficient implementation, especially for medium- to large-scale networks. In light of this, we focus on achieving the stability region induced by the approximate physical-layer capacity region for the high-signal-to-interference-plus-noise ratio (SINR) regime, which turns out to be a convex subset of the full capacity region. It turns out that this latter problem is amenable to distributed and efficient implementation through power control algorithms.

We present two main sets of results for achieving throughput optimality with respect to the high-SINR stability region. First, we develop a set of node-based scaled gradient projection power control algorithms which solves the MDB optimization using the high-SINR approximation of the link capacities. Specifically, as the transmission powers are iteratively updated by the algorithms, the high-SINR capacity expressions (which are always feasible since they underestimate the actual link capacities) are used. However, the network queues are still served using the actual link capacities. We show that the power control algorithms which use the high-SINR approximation can be implemented in a distributed manner using low communication overhead. On the other hand, since the power control algorithms typically require nonnegligible time to converge, the optimal high-SINR capacities for any given queue state can only be found iteratively over time. In the second result, we develop a new geometric approach for analyzing the expected Lyapunov drift, and show that the iterative power control policy with convergence time is guaranteed to achieve the stability region induced by the high-SINR capacities, as long as the second moments of traffic processes are bounded. Combining these two results, we conclude that our iterative power control algorithms represent a distributed and nearly optimal solution to the problem of throughput optimal control of CDMA wireless networks with random traffic arrivals in the high-SINR regime.

To the best of our knowledge, the iterative distributed power control algorithms developed in this work are the first node-based scheme for implementing the MDB policy in interference-limited wireless networks. Concurrent with our study, iterative implementation of the MDB policy has also been studied independently by Giannoulis *et al.* [13]. They investigate distributed power control algorithms for CDMA networks which are iterated once for every update of the queue state. In their scheme, the power and rate allocation algorithms are iterated only once in a slot, after which the queue state is re-sampled.

In contrast, our algorithms re-sample the queue state only after the power and rate allocation has converged to the optimum for the previous queue state. This paper is organized as follows. Section II introduces the model of stochastic multihop wireless networks and reviews the throughput optimal MDB (Backpressure) policy. In Section III, we introduce the interference-limited physical-layer model and present node-based power control algorithms that achieve the optimal rate allocation based on the high-SINR approximation of link capacities. In Section IV, we prove that the iterative MDB policy proposed in Section III is guaranteed to achieve the stability region induced by the high-SINR capacity region, even in the presence of nonnegligible convergence time.

II. NETWORK MODEL AND THROUGHPUT OPTIMAL CONTROL

A. Model of Stochastic Multihop Wireless Networks

Consider a wireless network represented by a directed and connected graph $\mathcal{G} = (\mathcal{N}, \mathcal{E})$. Each node $i \in \mathcal{N}$ models a wireless transceiver. An edge $(i, j) \in \mathcal{E}$ represents a unidirectional radio channel from node i to j . For convenience, let $\mathcal{O}(i) \triangleq \{j : (i, j) \in \mathcal{E}\}$ and $\mathcal{I}(i) \triangleq \{j : (j, i) \in \mathcal{E}\}$ denote the sets of node i 's next-hop and previous-hop neighbors, respectively. Let the vector $\mathbf{h} = (h_{ij})_{(i,j) \in \mathcal{E}}$ represent the (constant) channel gains on all links.

Denote the transmission power used on link (i, j) at (continuous) time τ by $P_{ij}(\tau) \geq 0$, and the instantaneous service rate of link (i, j) by $R_{ij}(\tau) \geq 0$. A feasible service rate vector $\mathbf{R}(\tau) = (R_{ij}(\tau))_{(i,j) \in \mathcal{E}}$ must belong to a given *instantaneous capacity region* $\mathcal{C}(\mathbf{P}(\tau))$ reflecting the physical-layer coding mechanism. Under individual power constraints $\hat{P}_i, i \in \mathcal{N}$, let

$$\Pi = \left\{ \mathbf{P}(\tau) \in \mathbb{R}_+^{|\mathcal{E}|} : \sum_{j \in \mathcal{O}(i)} P_{ij}(\tau) \leq \hat{P}_i, \forall i \in \mathcal{N} \right\}$$

be the set of feasible power allocations and

$$\mathcal{C}(\Pi) \triangleq \text{conv} \left(\bigcup_{\mathbf{P} \in \Pi} \mathcal{C}(\mathbf{P}) \right)$$

be the *long-term capacity region*. Here, the convex hull operation $\text{conv}(\cdot)$ indicates the possibility of time sharing among different feasible power configurations $\mathbf{P} \in \Pi$.

Let the data traffic in the network be classified according to their destinations. Traffic of type $k \in \mathcal{K}$ is destined for a set of nodes $\mathcal{N}_k \subset \mathcal{N}$ (when type- k traffic reaches any node in $\mathcal{N}_k$, it exits the network), where $\mathcal{K}$ is the set of all traffic types. Let $T > 0$ be a given time slot length. Let the number of bits of type- k traffic entering the network at node i from time tT to $(t+1)T$ be a nonnegative random variable $B_i^k[t]$. Assume that for all $t \in \mathbb{Z}_+$, $B_i^k[t]$ are independent and identically distributed with finite first and second moments and $\mathbb{P}(B_i^k[t] = 0) > 0$. Let $a_i^k \triangleq \frac{1}{T} \mathbb{E} B_i^k[t]$ denote the average (exogenous) arrival rate of type- k traffic at node i . Furthermore, assume all arrival processes $\{B_i^k[t]\}_{t=1}^\infty, i \in \mathcal{N}, k \in \mathcal{K}$ are mutually independent.

Assume node $i \in \mathcal{N}$ provides a (separate) infinite buffer i^k for each type k of traffic that is not destined for i . Denote the unfinished work in i^k at time τ by $U_i^k(\tau)$. We focus on the queue

states sampled at slot boundaries $\tau = tT$, $t \in \mathbb{Z}_+$. Let $U_i^k[t]$ denote the instantaneous backlog at the beginning of the t th slot, i.e., $U_i^k[t] = U_i^k(tT)$. Over the t th slot, link (i, j) serves i^k at average rate $R_{ij}^k[t] = \frac{1}{T} \int_{tT}^{(t+1)T} R_{ij}^k(\tau) d\tau$. Thus, we have the following queueing dynamics:

$$U_i^k[t+1] \leq \left(U_i^k[t] - T \sum_{j \in \mathcal{O}(i)} R_{ij}^k[t] + T \sum_{m \in \mathcal{I}(i)} R_{mi}^k[t] + B_i^k[t] \right)^+. \quad (1)$$

Here $(x)^+$ denotes $\max\{x, 0\}$, and the inequality comes from the fact that in general, since certain queues may be empty, the actual endogenous arrivals are less than or equal to the nominal amount $T \sum_{m \in \mathcal{I}(i)} R_{mi}^k[t]$. Finally, the aggregate service rate on link (i, j) over the t th slot is $TR_{ij}[t] = \sum_{k \in \mathcal{K}} TR_{ij}^k[t]$.

B. Stability Region and Throughput Optimal Policy

Given the wireless network model, we now define notions of stability and throughput optimal control policies.

Definition 1: [2] The queue i^k is *stable* if $g_i^k(\xi) \triangleq \limsup_{n \rightarrow \infty} \frac{1}{n} \sum_{t=1}^n \mathbb{P}[U_i^k[t] > \xi] \rightarrow 0$ as $\xi \rightarrow \infty$. Input processes $\left\{ B[t] = \left(B_i^k[t] \right)_{i \in \mathcal{N}, k \in \mathcal{K}} \right\}_{t=1}^\infty$ are *stabilizable* if there exist service processes $\{R_{ij}^k[t]\}$ for all $(i, j) \in \mathcal{E}$ and $k \in \mathcal{K}$ such that for every $t \in \mathbb{Z}_+$, $\mathbf{R}[t] \in \mathcal{C}(\Pi)$ and the resulting queueing processes are all stable.

Definition 2: The *stability region* Λ of a wireless multihop network is the closure of the set of the average arrival rate vectors $\mathbf{a}$ of all stabilizable input processes.

For a general wireless multihop network, its stability region has a simple characterization in terms of supporting multicommodity rates that are feasible for a given capacity region.

Theorem 1: [2] The stability region Λ of the wireless multihop network with transmission power constraint Π is the set of all average rate vectors $(a_i^k) \geq \mathbf{0}$ such that there exists a multicommodity flow vector (f_{ij}^k) satisfying

$$\begin{aligned} f_{ij}^k &\geq 0, \quad \forall (i, j) \in \mathcal{E} \quad \text{and} \quad k \in \mathcal{K} \\ a_i^k &\leq \sum_{j \in \mathcal{O}(i)} f_{ij}^k - \sum_{m \in \mathcal{I}(i)} f_{mi}^k \quad \forall i \in \mathcal{N}, \quad k \in \mathcal{K} \\ \sum_{k \in \mathcal{K}} f_{ij}^k &\leq C_{ij}, \quad \forall (i, j) \in \mathcal{E} \quad \text{where} \quad (C_{ij})_{(i,j) \in \mathcal{E}} \in \mathcal{C}(\Pi). \end{aligned}$$

We refer to the stability region characterized in Theorem 1 as the stability region induced by the capacity region $\mathcal{C}(\Pi)$ and denote it by $\Lambda(\mathcal{C}(\Pi))$. The following Maximum Differential Backlog (MDB) policy has been shown to be *throughput optimal* [1], [2] in the sense that it stabilizes all input processes with average rate vectors belonging to the interior of $\Lambda(\mathcal{C}(\Pi))$, without knowledge of arrival statistics. The policy can be described as follows:

- 1) At slot t , find traffic type $k_{ij}^*[t]$ having the *maximum differential backlog* over link (i, j) for all $(i, j) \in \mathcal{E}$. That is, $k_{ij}^*[t] = \arg \max_{k \in \mathcal{K}} \{U_i^k[t] - U_j^k[t]\}$, where $U_j^k[t] \equiv 0$

if $j \in \mathcal{N}_k$. Let $b_{ij}^*[t] = \max\{0, U_i^{k^*}[t] - U_j^{k^*}[t]\}$, where $k^* \equiv k_{ij}^*[t]$.

- 2) Find the rate vector $\mathbf{R}^*[t]$ which solves

$$\max_{\mathbf{R} \in \mathcal{C}(\Pi)} \sum_{(i,j) \in \mathcal{E}} b_{ij}^*[t] \cdot R_{ij}. \quad (2)$$

- 3) The service rate provided by link (i, j) to queue i^k is determined by

$$R_{ij}^k[t] = \begin{cases} R_{ij}^*[t], & \text{if } k = k_{ij}^*[t] \\ 0, & \text{otherwise.} \end{cases}$$

For wired networks, the above MDB policy can be implemented in a fully distributed manner. In wireless networks, however, the capacity of a link is usually affected by interference from other links. Therefore, solving (2) in general requires centralized computation. Thus far, distributed solutions for (2) are available only for relatively simple physical layer models [12].

In the following, we develop efficient distributed MDB control algorithms for interference-limited CDMA networks with random traffic. Throughout the rest of the paper, we assume all nodes have synchronized clocks so that the boundaries of time slots at all nodes are aligned. This assumption guarantees that the MDB values in (2) are taken at the same instant across all links. The study of MDB policy based on asynchronously sampled queue state will be a subject of future work.

III. DISTRIBUTED MAXIMUM DIFFERENTIAL BACKLOG CONTROL

A. Throughput Optimal Power Control

We study a wireless network using direct-sequence spread-spectrum CDMA. The received signal-to-interference-plus-noise ratio (SINR) per channel code symbol of link (i, j) is given by

$$\text{SINR}_{ij} = \frac{K h_{ij} P_{ij}}{h_{ij}(P_i - P_{ij}) + \sum_{m \neq i} h_{mj} P_m + N_j},$$

where K is the processing gain, $P_m = \sum_{k \in \mathcal{O}(m)} P_{mk}$ is the total transmission power of node m , and N_j represents the noise power of receiver j .

Assume the receiver of every link decodes its own signal against the interference from other links as Gaussian noise. The information-theoretic capacity of link (i, j) is given by

$$\bar{C}_{ij}(\mathbf{P}) = R_s \log \left(1 + \frac{K h_{ij} P_{ij}}{h_{ij}(P_i - P_{ij}) + \sum_{m \neq i} h_{mj} P_m + N_j} \right). \quad (3)$$

For convenience, we normalize the channel symbol rate R_s to be one for subsequent analysis. We also take $\log(\cdot)$ to be the natural logarithm to simplify differentiation operations.

For a given power configuration $\mathbf{P}$, let

$$\bar{\mathcal{C}}(\mathbf{P}) \equiv \left\{ \mathbf{R} \in \mathbb{R}_+^{|\mathcal{E}|} : R_{ij} \leq \bar{C}_{ij}(\mathbf{P}), \forall (i, j) \in \mathcal{E} \right\}$$

denote the corresponding capacity region. It is known that function $\bar{C}_{ij}(\mathbf{P})$ is not concave in $\mathbf{P}$, and therefore that the achiev-

able capacity region $\bigcup_{\mathbf{P} \in \Pi} \bar{\mathcal{C}}(\mathbf{P})$ under the individual power constraints Π is not convex. Thus, scheduling and time-sharing are generally needed to achieve the long-term capacity region

$$\bar{\mathcal{C}}(\Pi) = \text{conv} \left(\bigcup_{\mathbf{P} \in \Pi} \bar{\mathcal{C}}(\mathbf{P}) \right).$$

Scheduling and time-sharing, however, are often not amenable to distributed implementation, especially for medium- to large-scale networks. As we explain next, however, the achievable capacity region based on the *high-SINR approximation* of link capacities is *convex*. This region is a subset of the actual capacity region $\bigcup_{\mathbf{P} \in \Pi} \bar{\mathcal{C}}(\mathbf{P})$. The gap between the actual capacity region and the approximate high-SINR region is typically small in CDMA networks.

In most CDMA systems, due to the large multiplication factor K , the SINR *per symbol*

$$\frac{Kh_{ij}P_{ij}}{h_{ij}(P_i - P_{ij}) + \sum_{m \neq i} h_{mj}P_m + N_j}$$

is typically high [14]. Therefore, the link capacity in (3) is well approximated by [15], [16]

$$\underline{\mathcal{C}}_{ij}(\mathbf{P}) = \log \left(\frac{Kh_{ij}P_{ij}}{h_{ij} \sum_{k \neq j} P_{ik} + \sum_{m \neq i} h_{mj} \sum_{k \in \mathcal{O}(m)} P_{mk} + N_j} \right).$$

We call the region

$$\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P}) = \left\{ \mathbf{R} \in \mathbb{R}_+^{|\mathcal{E}|} : R_{ij} \leq \underline{\mathcal{C}}_{ij}(\mathbf{P}), \mathbf{P} \in \Pi \right\}$$

the *achievable high-SINR capacity region* under the individual node power constraints Π . Notice that $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P})$ is an underestimate of $\bigcup_{\mathbf{P} \in \Pi} \bar{\mathcal{C}}(\mathbf{P})$, since for any power configuration $\mathbf{P}$, $\underline{\mathcal{C}}_{ij}(\mathbf{P}) < \bar{\mathcal{C}}_{ij}(\mathbf{P})$, and thus $\underline{\mathcal{C}}(\mathbf{P}) \subset \bar{\mathcal{C}}(\mathbf{P})$ and $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P}) \subset \bigcup_{\mathbf{P} \in \Pi} \bar{\mathcal{C}}(\mathbf{P})$. The gap between $\underline{\mathcal{C}}(\mathbf{P})$ and $\bar{\mathcal{C}}(\mathbf{P})$ is characterized by

$$\bar{\mathcal{C}}_{ij}(\mathbf{P}) - \underline{\mathcal{C}}_{ij}(\mathbf{P}) = \log \left(1 + \frac{1}{\text{SINR}_{ij}} \right)$$

which is close to zero if SINR_{ij} is large. Thus, $\underline{\mathcal{C}}(\mathbf{P})$ is a close approximation of $\bar{\mathcal{C}}(\mathbf{P})$ in the high-SINR regime.

Note that $\underline{\mathcal{C}}(\mathbf{P})$ is an *analytical* approximation of the actual region $\bar{\mathcal{C}}(\mathbf{P})$. When $\mathbf{P}$ is the instantaneous link power configuration, the *physically-feasible* capacity region is still given by $\bar{\mathcal{C}}(\mathbf{P})$. We wish to analytically approximate the true capacities ($\bar{\mathcal{C}}_{ij}(\mathbf{P})$) by ($\underline{\mathcal{C}}_{ij}(\mathbf{P})$) for the following reason.

With a change of variables $S_i = \ln P_i$, $\hat{S}_i = \ln \hat{P}_i$, and $S_{ij} = \ln P_{ij}$, the high-SINR capacity function becomes

$$\underline{\mathcal{C}}_{ij}(\mathbf{S}) = \log(Kh_{ij}) + S_{ij} - \log \left(h_{ij} \sum_{k \neq j} e^{S_{ik}} + \sum_{m \neq i} h_{mj} \sum_{k \in \mathcal{O}(m)} e^{S_{mk}} + N_j \right) \quad (4)$$

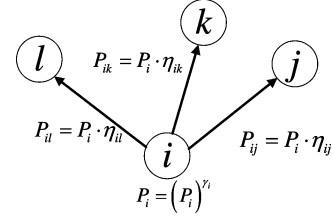


Fig. 1. Transmission powers in terms of the power control and power allocation variables.

which is known to be concave in $\mathbf{S}$ [15], [16]. It follows that the high-SINR capacity region $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P})$ is *convex*. Thus, $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P}) = \text{conv} \left(\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P}) \right) \equiv \underline{\mathcal{C}}(\Pi)$. Therefore, any point in the long-term high-SINR capacity region $\underline{\mathcal{C}}(\Pi)$ can be achieved by using an appropriate power configuration, without the need for scheduling and time sharing among different power configurations.

In a given time slot, the MBD optimization problem with capacities approximated by the high-SINR expression (4) is given by (2) with the long-term capacity region being $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P})$.¹ It can be rewritten as the following concave maximization problem:

$$\begin{aligned} & \text{maximize} && \sum_{(i,j) \in \mathcal{E}} b_{ij}^* R_{ij} \\ & \text{subject to} && R_{ij} = \underline{\mathcal{C}}_{ij}(\mathbf{S}), \forall (i,j) \in \mathcal{E} \\ & && \sum_{j \in \mathcal{O}(i)} e^{S_{ij}} \leq \hat{P}_i, \quad \forall i \in \mathcal{N}. \end{aligned} \quad (5)$$

Without loss of generality, we assume the MDB $b_{ij}^* > 0$ for all (i,j) [otherwise, we can simply exclude those links having $b_{ij}^* = 0$ from the objective function in (5)]. Next, we will develop distributed power control algorithms which iteratively solve the optimization in (5).

B. Power Adjustment Variables

First we introduce a set of node-based control variables for adjusting the transmission powers on all links. They are

$$\begin{aligned} \text{Power allocation variables:} & \quad \eta_{ik} \triangleq \frac{P_{ik}}{P_i}, (i,k) \in \mathcal{E} \\ \text{Power control variables:} & \quad \gamma_i \triangleq \frac{S_i}{\hat{S}_i}, \quad i \in \mathcal{N}. \end{aligned}$$

These variables are illustrated in Fig. 1. With appropriate scaling, we can always let $\hat{P}_i > 1$ for all $i \in \mathcal{N}$ so that $\hat{S}_i > 0$. Therefore, we have the equivalent High-SINR Power Control (HSPC) problem shown in equation (6) at the bottom of the next page.

C. Conditions For Optimality

To solve the HSPC problem in (6), we compute the gradients of the objective function, denoted by F , with respect to the

¹Notice that even if $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P})$ is not convex, restricting the feasible set of the optimization in (2) to $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P})$ does not lose any optimality. This is because the objective function is linear in the link rates, and so the maximum over $\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P})$ is equal to the maximum over $\underline{\mathcal{C}}(\Pi) = \text{conv} \left(\bigcup_{\mathbf{P} \in \Pi} \underline{\mathcal{C}}(\mathbf{P}) \right)$.

power allocation variables and the power control variables, respectively. For all $i \in \mathcal{N}$ and $j \in \mathcal{O}(i)$

$$\frac{\partial F}{\partial \eta_{ij}} = P_i \left[- \sum_{m \in \mathcal{N}} \sum_{k \in \mathcal{O}(m)} b_{mk}^* \frac{h_{ik}}{IN_{mk}} + \delta \eta_{ij} \right]$$

where the *power allocation marginal gain indicator* is

$$\delta \eta_{ij} \triangleq b_{ij}^* \left(\frac{1}{P_{ij}} + \frac{h_{ij}}{IN_{ij}} \right). \quad (7)$$

For all $i \in \mathcal{N}$

$$\frac{\partial F}{\partial \gamma_i} = \hat{S}_i \cdot \delta \gamma_i$$

where the *power control marginal gain indicator* is

$$\delta \gamma_i \triangleq P_i \left[- \sum_{m \in \mathcal{N}} \sum_{k \in \mathcal{O}(m)} \frac{b_{mk}^* h_{ik}}{IN_{mk}} + \sum_{k \in \mathcal{O}(i)} \delta \eta_{ik} \cdot \eta_{ik} \right]. \quad (8)$$

The term IN_{ij} appearing above is short-hand notation for the overall interference-plus-noise power at the receiver end of link (i, j) , that is

$$IN_{ij} = h_{ij} \sum_{k \neq j} e^{S_{ik}} + \sum_{m \neq i} h_{mj} \sum_{k \in \mathcal{O}(m)} e^{S_{mk}} + N_j.$$

The marginal gain indicators fully characterize the optimality conditions as follows.

Theorem 2: A feasible set of transmission power variables $\{\eta_{ik}\}_{(i,k) \in \mathcal{E}}$ and $\{\gamma_i\}_{i \in \mathcal{N}}$ is the solution of the HSPC problem (6) if and only if the following conditions hold. For all $i \in \mathcal{N}$, there exists a constant ν_i such that

$$\delta \eta_{ik} = \nu_i, \quad \forall k \in \mathcal{O}(i) \quad (9)$$

$$\delta \gamma_i = 0, \quad \text{if } \gamma_i < 1 \quad (10)$$

$$\delta \gamma_i \geq 0, \quad \text{if } \gamma_i = 1 \quad (11)$$

Here, all $\eta_{ik} > 0$ since $b_{ik}^* > 0$ by assumption.

For the detailed proof of Theorem 2, see [17]. Due to the distributed form of the optimality conditions, every node can check the conditions with respect to its controlled variables locally, and adjust them towards the optimum. In Section III-D, we present a set of distributed algorithms that achieve the globally optimal power configuration for the HSPC problem.

D. Distributed Power Control Algorithms

We design scaled gradient projection algorithms which iteratively update the nodes' power allocation variables and power control variables in a distributed manner, so as to asymptotically converge to the optimal solution of (6). At each iteration, the variables are updated in the positive gradient direction, scaled by a positive definite matrix. When an update leads to a point outside the feasible set, the point is projected back into the feasible set [18].

1) *Power Allocation Algorithm (PA)*: At the k th iteration at node i , the current local power allocation vector $\boldsymbol{\eta}_i^k = (\eta_{ij}^k)_{j \in \mathcal{O}(i)}$ is updated by

$$\boldsymbol{\eta}_i^{k+1} = PA(\boldsymbol{\eta}_i^k) = \left[\boldsymbol{\eta}_i^k + (Q_i^k)^{-1} \cdot \delta \boldsymbol{\eta}_i^k \right]_{Q_i^k}^+.$$

Here, $\delta \boldsymbol{\eta}_i^k = (\delta \eta_{ij}^k)_{j \in \mathcal{O}(i)}$ and the matrix Q_i^k is symmetric, positive definite on the subspace $\{\mathbf{v}_i : \sum_{j \in \mathcal{O}(i)} v_{ij} = 0\}$. Finally, $[\cdot]_{Q_i^k}^+$ denotes the projection on the feasible set of $\boldsymbol{\eta}_i$ relative to the norm induced by Q_i^k .²

Suppose each node j can measure the value of SINR_{ij} for any of its incoming links. Before an iteration of PA, node i collects the current SINR_{ij} 's via feedback from its next-hop neighbors j . Node i can then readily compute all $\delta \eta_{ij}$'s according to

$$\delta \eta_{ij} = b_{ij}^* \left(\frac{1}{P_{ij}} + \frac{h_{ij}}{IN_{ij}} \right) = \frac{b_{ij}^*}{P_{ij}} \left(1 + \frac{\text{SINR}_{ij}}{K} \right).$$

Note that since the calculation of $\delta \eta_{ij}$ involves only locally obtainable measures, the PA algorithm does not require global exchange of control messages.³

2) *Power Control Algorithm (PC)*: After a phase for exchanging control messages (which will be discussed below), every node i is able to calculate its power control marginal gain indicator $\delta \gamma_i$. From a network-wide viewpoint, the power control vector $\boldsymbol{\gamma}^k = (\gamma_i^k)_{i \in \mathcal{N}}$ is updated by

$$\boldsymbol{\gamma}^{k+1} = PC(\boldsymbol{\gamma}^k) = \left[\boldsymbol{\gamma}^k + (V^k)^{-1} \cdot \delta \boldsymbol{\gamma}^k \right]_{V^k}^+.$$

Here, the scaling matrix V^k is symmetric and positive definite. Note that PC becomes amenable to distributed implementation if and only if V^k is diagonal.

²In general, $[\tilde{\mathbf{x}}]_{Q_i^k}^+ \equiv \arg \min_{\mathbf{x} \in \mathcal{F}} (\mathbf{x} - \tilde{\mathbf{x}})' \cdot Q_i^k \cdot (\mathbf{x} - \tilde{\mathbf{x}})$, where $\mathcal{F}$ is the feasible set of $\mathbf{x}$.

³It is assumed here and for the following that control messages are exchanged between nodes on a communication channel separate from the main channel where the network traffic is carried.

$$\begin{aligned} & \text{maximize} && \sum_{(i,j)} b_{ij}^* \log \frac{K h_{ij} (\hat{P}_i)^{\gamma_i} \eta_{ij}}{h_{ij} (\hat{P}_i)^{\gamma_i} (1 - \eta_{ij}) + \sum_{m \neq i} h_{mj} (\hat{P}_m)^{\gamma_m} + N_j} \\ & \text{subject to} && \eta_{ij} \geq 0, \quad \forall (i, j) \in \mathcal{E}, \\ & && \sum_{j \in \mathcal{O}(i)} \eta_{ij} = 1, \quad \gamma_i \leq 1, \quad \forall i \in \mathcal{N}. \end{aligned} \quad (6)$$

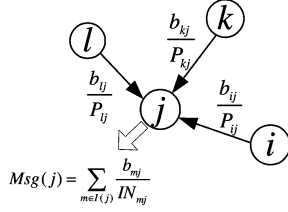


Fig. 2. Information exchange protocol for power control algorithm.

We now derive an efficient protocol which allows each node to calculate its own $\delta\gamma_i$ given limited control messaging. We first reorder the summations on the RHS of (8) as

$$\delta\gamma_i = P_i \left[\sum_{j \neq i} \left\{ -h_{ij} \sum_{m \in \mathcal{I}(j)} \frac{b_{mj}^*}{IN_{mj}} \right\} + \sum_{j \in \mathcal{O}(i)} \delta\eta_{ij} \cdot \eta_{ij} \right]. \quad (12)$$

With reference to the above expression, we propose the following procedure for computing $\delta\gamma_i$.

Power Control Message Exchange Protocol: Let each node j assemble the measures $\frac{b_{mj}^*}{IN_{mj}}$ from all its incoming links (m, j) . For this purpose, an upstream neighbor m needs to inform j of the value b_{mj}^*/P_{mj} . Since node j can measure both $SINR_{mj}$ and h_{mj} , it can calculate

$$\frac{b_{mj}^*}{IN_{mj}} = \frac{b_{mj}^*}{P_{mj}} \frac{SINR_{mj}}{h_{mj}K}.$$

After obtaining the measures from all incoming links, node j sums them up to form the power control message

$$Msg(j) \triangleq \sum_{m \in \mathcal{I}(j)} \frac{b_{mj}^*}{IN_{mj}}.$$

It then broadcasts $Msg(j)$ to the whole network. The process for control messaging is illustrated by Fig. 2, where the solid arrows represent local message communication and the hollow arrow signifies the broadcasting of the message.

Upon obtaining $Msg(j)$ from node $j \neq i$, node i processes it according to the following rule. If j is a next-hop neighbor of i , it multiplies the message with h_{ij} and subtracts the product from the local measure $\delta\eta_{ij} \cdot \eta_{ij}$. Otherwise, it multiplies $Msg(j)$ with $-h_{ij}$. Finally, node i adds up the results derived from processing all other nodes' messages, and this sum multiplied by P_i equals $\delta\gamma_i$. Note that in a symmetric duplex channel, $h_{ij} \approx h_{ji}$, and node i may use its own measure of h_{ji} in place of h_{ij} . Otherwise, it will need channel feedback from node j to calculate h_{ij} . To summarize, the protocol requires *only one message from each node* to be broadcast to the whole network. Moreover in practice, node i can ignore the messages generated by distant nodes, since they contribute very little to $\delta\gamma_i$ due to the negligible multiplicative factor h_{ij} on $Msg(j)$ when i and j are far apart [cf. (12)].

3) **Convergence of Algorithms:** We now formally state the central convergence result for the *PA* and *PC* algorithms discussed above.

Theorem 3: From any feasible initial transmission power configuration $\{\eta_i^0\}$ and γ^0 , there exist appropriate scaling matrices $\{Q_i^k\}$ and V^k such that the sequences generated by the

algorithms $PA(\cdot)$ and $PC(\cdot)$ converge, i.e., $\eta_i^k \rightarrow \eta_i^*$ for all i , and $\gamma^k \rightarrow \gamma^*$ as $k \rightarrow \infty$. Furthermore, $\{\eta_i^*\}$ and γ^* constitute a set of jointly optimal solution to the HSPC problem (6).

In the *PA* and *PC* algorithms, the scaling matrices are chosen to be appropriate diagonal matrices which approximate the relevant Hessians such that the objective value is increased by every iteration until the optimum is achieved. This allows the scaled gradient projection algorithms to approximate constrained Newton algorithms, which are known to have fast convergence rates. Furthermore, the scaling matrices are shown to guarantee convergence from all initial conditions. These features are crucial for the applicability of these algorithms to large networks which lack the ability of centrally scheduling and synchronizing node operations. Finally, we note that the scaling matrices can be easily calculated at each node using very limited control messaging. The detailed derivation and the full proof of Theorem 3 can be found in [17].

Also note that convergence of the algorithms does not require any particular order of running *PA* and *PC* algorithms at different nodes. Any node i only needs to update its own variables η_i and γ_i using *PA* and *PC* until its local variables satisfy the optimality conditions (9)–(11).

IV. STABILITY REGION OF ITERATIVE MAXIMUM DIFFERENTIAL BACKLOG POLICY

The MDB policy in Section II-B is implemented through the node-based power control scheme as follows.

- 1) At the beginning of the t th slot, i.e., at time $\tau = tT$, queue state $U_i^k[t]$ is sampled by every node i for each type k of traffic.⁴ Every node i then computes the maximum differential backlog (MDB) on each of its outgoing links (i, j) for the current slot. Denote the MDB by $b_{ij}^*[t] = \max\{0, \max_{k \in \mathcal{K}} \{U_i^k[t] - U_j^k[t]\}\}$, and if $b_{ij}^*[t] > 0$, denote the type of traffic attaining the MDB by $k_{ij}^*[t] = \arg \max_{k \in \mathcal{K}} \{U_i^k[t] - U_j^k[t]\}$.
- 2) At any time $\tau \in [tT, (t+1)T)$, the link capacities are determined by the instantaneous transmission powers $\mathbf{P}(\tau) = (P_{ij}(\tau))$ according to the *actual* capacity formula $\bar{C}_{ij}(\mathbf{P}(\tau))$. The instantaneous service rates for different queues are allocated as follows:

$$\bar{R}_{ij}^k(\tau) = \begin{cases} \bar{C}_{ij}(\mathbf{P}(\tau)), & \text{if } k = k_{ij}^*[t] \\ 0, & \text{otherwise.} \end{cases} \quad (13)$$

- 3) The link transmission powers, and thus the link capacities and service rates, are iterated in time. Specifically at any $\tau \in [tT, (t+1)T)$, each node i iterates the *PA* or *PC* algorithms to update its local power η_i or γ_i . Algorithm parameters such as $\delta\eta_i$ and $\delta\gamma_i$ are calculated using the MDB values $(b_{mn}^*[t])_{(m,n) \in \mathcal{E}}$ at the beginning of the slot. All calculations and iterations are performed using the *high-SINR* capacity formula $\underline{C}_{ij}(\mathbf{P}(\tau))$, as described in Sections III-B–D.
- 4) At the beginning of the $t+1$ th slot, the queue state is re-sampled and the MDBs are recalculated as in step (1). Steps (2)–(3) are carried out using the new queue state.

⁴For $j \in \mathcal{N}_k$, $U_j^k[t]$ is always taken to be 0.

The above implementation of the MDB policy deviates from the original policy in two important respects. First, the high-SINR link capacity formula (4) is used in the iterated optimization of the power configuration $\mathbf{P}$, while the network queues are served at the actual link capacities given by (3). Second, since the PA and PC algorithms require a certain number of iterations before reaching a close neighborhood of an optimum to the problem in (6), the transmission rates (which are computed according to the actual link capacity formulas) corresponding to the optimal high-SINR power configuration cannot be applied instantaneously. Rather, the optimal high-SINR transmission powers can only be found iteratively over time. At any moment in the convergence interval, the queues are served at rates which are iteratively updated towards the transmission rates corresponding to the optimal power configuration for the queue state at the beginning of the slot. The power configuration obtained at the end of a convergence period are optimal only for the queue state some time ago. The effect of using lagging optimal service rates is studied in the context of $N \times N$ packet switches by Neely *et al.* [19]⁵ and in a queueing network with Poisson arrivals and exponential service rates by Tassiulas and Ephremides [20]. In [19], [20], however, the process of finding the optimal service rates is not iterative. It is assumed that once the (outdated) queue state information becomes available, the optimal service rates are obtained instantaneously.

In the following, we will analyze the impact of both the high-SINR approximation and the convergence time on the stability region achievable by the proposed iterative MDB policy. We show that the policy can stabilize any arrival processes (with finite second moments) whose average arrival rates are within the stability region induced by the high-SINR capacity region. To accomplish this, we develop a new geometric approach for computing the expected Lyapunov drift of the queue state.

A. Transient Optimal Rates

Without loss of generality, assume that the convergence time of the distributed power control algorithms in Section III-D is the length of a time slot T ,⁶ i.e., at time $\tau = (t+1)T$, the optimal high-SINR transmission power vector for $\mathbf{U}[t]$ is obtained. For ease of analysis and without loss of generality, we further scale time so that $T = 1$.

Recall from Section II that $R_{ij}^k(\tau)$ denotes the instantaneous service rate on link (i, j) for type k traffic. The service rate vector $(R_{ij}^k(\tau))$ is understood to vary continuously over time τ according to the iterations of the PA and PC algorithms. The total service rate (in bits/slot) provided by (i, j) to type- k traffic over the t th slot (from time tT to $(t+1)T$, where $T = 1$) is given by

$$R_{ij}^k[t] = \int_t^{t+1} R_{ij}^k(\tau) d\tau.$$

⁵In [19], the current queue state is taken to be the state of the Markov chain used for stability analysis. As we show below, however, the Markov state should consist of the current queue state as well as the previous queue state.

⁶In practice, the gradient projection algorithms can only find an approximate optimal solution within a finite period of time. In this work, we make the idealization that the exact optimum can be achieved after the convergence period T . Such an assumption simplifies the following analysis while its loss of precision is small when we take T sufficiently large.

Instead of studying the service rates $(R_{ij}^k(\tau))$ or $(R_{ij}^k[t])$, in this section we focus on *net service rates*. First define the instantaneous net service rate of queue i^k by⁷

$$\tilde{R}_i^k(\tau) = \sum_{j \in \mathcal{O}(i)} R_{ij}^k(\tau) - \sum_{m \in \mathcal{I}(i)} R_{mi}^k(\tau).$$

Let $\tilde{\mathcal{C}}(\mathbf{P})$ be the set of net service rate vectors induced by service rate vectors (R_{ij}^k) such that $\sum_k R_{ij}^k \leq \bar{C}_{ij}(\mathbf{P})$. This represents the feasible set of instantaneous net service rates when the power configuration is $\mathbf{P}$. Next, let $\tilde{\mathcal{C}}(\mathbf{P})$ be the set of $(\tilde{R}_i^k)$ induced by (R_{ij}^k) such that $\sum_k R_{ij}^k \leq \underline{C}_{ij}(\mathbf{P})$. That is, $\tilde{\mathcal{C}}(\mathbf{P})$ is the feasible set of instantaneous net service rates under the high-SINR approximation.

Working with net service rates considerably simplifies our subsequent analysis. The net service rates (in bits/slot) over the t th slot, induced by service rate vector $\mathbf{R}[t]$, are given by

$$\tilde{R}_i^k[t] = \sum_{j \in \mathcal{O}(i)} R_{ij}^k[t] - \sum_{m \in \mathcal{I}(i)} R_{mi}^k[t].$$

Let $\tilde{\mathcal{C}}(\Pi)$ be the set of $(\tilde{R}_i^k[t])$ induced by $\mathbf{R}[t] \in \bar{\mathcal{C}}(\Pi)$, and let $\tilde{\mathcal{C}}(\Pi)$ be the set of $\tilde{R}_i^k[t]$ induced by $\mathbf{R}[t] \in \underline{\mathcal{C}}(\Pi)$. It is straightforward to verify that both $\tilde{\mathcal{C}}(\Pi)$ and $\tilde{\mathcal{C}}(\Pi)$ are compact and convex. By Theorem 1 of [3], the intersection of $\tilde{\mathcal{C}}(\Pi)$ and $\tilde{\mathcal{C}}(\Pi)$ with the positive orthant are the stability regions induced by the actual capacity region $\bar{\mathcal{C}}(\Pi)$ and by the high-SINR capacity region $\underline{\mathcal{C}}(\Pi)$, respectively.

Using the net service rates $\tilde{\mathbf{R}}[t] \equiv (\tilde{R}_i^k[t])$, we can re-write the queueing dynamics in (1) as follows:

$$\mathbf{U}[t+1] \leq (\mathbf{U}[t] - \tilde{\mathbf{R}}[t] + \mathbf{B}[t])^+. \quad (14)$$

In the following, we use $\tilde{\mathbf{R}}[t]$ to denote the net service rate vector induced by the actual service rates allocated according to (13). Correspondingly, we let $\underline{\tilde{\mathbf{R}}}[t]$ denote the net service rates calculated using the high-SINR capacity formula, i.e., $\underline{\tilde{\mathbf{R}}}[t]$ is induced by the service rates allocated according to

$$\underline{R}_{ij}^k = \begin{cases} \underline{C}_{ij}(\mathbf{P}), & \text{if } k = k_{ij}^*[t] \\ 0, & \text{otherwise.} \end{cases} \quad (15)$$

It is easy to verify that for any $\mathbf{P} \in \Pi$ and any net service rate vector $\tilde{\mathbf{R}} \in \tilde{\mathcal{C}}(\mathbf{P})$

$$\sum_{(i,j) \in \mathcal{E}} b_{ij}^*[t] \cdot \bar{C}_{ij}(\mathbf{P}) \geq \mathbf{U}[t]' \cdot \tilde{\mathbf{R}}$$

with equality if and only if $\tilde{\mathbf{R}}$ is induced by the service rates allocated according to (13). Similarly, for any $\tilde{\mathbf{R}} \in \tilde{\mathcal{C}}(\mathbf{P})$, we have

$$\sum_{(i,j) \in \mathcal{E}} b_{ij}^*[t] \cdot \underline{C}_{ij}(\mathbf{P}) \geq \mathbf{U}[t]' \cdot \tilde{\mathbf{R}}$$

with equality if and only if $\tilde{\mathbf{R}}$ is induced by the service rates allocated according to (15). Therefore, if $\mathbf{P}^*$ is the optimal power

⁷Net service rates can be negative, as when a queue's endogenous incoming rate is higher than its outgoing rate.

configuration for (5), then the net service rates induced by the service rates allocated according to (15) maximize $\mathbf{U}[t]' \cdot \tilde{\mathbf{R}}$ over all $\tilde{\mathbf{R}} \in \tilde{\mathcal{C}}(\Pi)$. Denote the maximizing $\tilde{\mathbf{R}} \in \tilde{\mathcal{C}}(\Pi)$ by $\tilde{\mathbf{R}}^*(\mathbf{U}[t])$. In the rest of the paper, we denote $\tilde{\mathcal{C}}(\Pi)$ by $\mathcal{C}$ for brevity. Expressed using the simplified notation

$$\tilde{\mathbf{R}}^*(\mathbf{U}[t]) = \arg \max_{\tilde{\mathbf{R}} \in \mathcal{C}} \mathbf{U}[t]' \cdot \tilde{\mathbf{R}}.$$

In the following, we will simply refer to $\tilde{\mathbf{R}}$ as the service rate vector and $\tilde{\mathbf{R}}^*(\mathbf{U}[t])$ as the *optimal rate allocation* for queue state $\mathbf{U}[t]$.

Recall our discussion of the distributed MDB control algorithms in the last section. Due to the iterative nature of the algorithms, the optimal power vector and the optimal rate allocation for a given queue state can be found only when the algorithms converge. Therefore in practice, the rate vector solving (2) for $(b_{ij}^*[t])$ cannot be applied instantly at the beginning of the t th slot. The service rates $\tilde{\mathbf{R}}(\tau)$, $\tau \in \mathbb{R}_+$, are always in transience, shifting from the previous optimum to the next optimum. Thus, the instantaneous rate vector at time $\tau = t$ is $\tilde{\mathbf{R}}(t) = \tilde{\mathbf{R}}^*(\mathbf{U}[t-1])$, and at time $\tau = t+1$, $\tilde{\mathbf{R}}(t+1) = \tilde{\mathbf{R}}^*(\mathbf{U}[t])$.

B. Lyapunov Drift Criterion

Following the previous model, the process $\{(\mathbf{U}[t], \mathbf{U}[t-1])\}_{t=1}^{\infty}$ forms a Markov chain. The state $(\mathbf{U}[t], \mathbf{U}[t-1]) \triangleq \mathbf{W}[t]$ lies in the state space $\mathcal{W} = \mathbb{R}_+^M \times \mathbb{R}_+^M$ where M is the total number of queues. As an extension of Foster's criterion for a recurrent Markov chain [21], the following condition is used in studying the stability of stochastic queueing systems [1], [19].

Lemma 1: [1] If there exist a (Lyapunov) function $V : \mathcal{W} \mapsto \mathbb{R}_+$, a compact subset $\mathcal{W}_0 \subset \mathcal{W}$, and a positive constant ε_0 such that for all $\mathbf{w} \in \mathcal{W}_0$

$$\mathbb{E}[V(\mathbf{W}[t+1]) - V(\mathbf{W}[t]) | \mathbf{W}[t] = \mathbf{w}] < \infty \quad (16)$$

and for all $\mathbf{w} \notin \mathcal{W}_0$

$$\mathbb{E}[V(\mathbf{W}[t+1]) - V(\mathbf{W}[t]) | \mathbf{W}[t] = \mathbf{w}] \leq -\varepsilon_0 \quad (17)$$

then the Markov chain $\{\mathbf{W}[t]\}$ is recurrent.⁸ Hence, the queueing system is stable in the sense of Definition 1.

We use the Lyapunov function from [20]

$$\begin{aligned} V(\mathbf{W}[t]) &= \sum_{k \in \mathcal{K}} \sum_{i \in \mathcal{N}} U_i^k[t]^2 + (U_i^k[t] - U_i^k[t-1])^2 \\ &= \|\mathbf{U}[t]\|^2 + \|\mathbf{U}[t] - \mathbf{U}[t-1]\|^2 \end{aligned}$$

where $\|\cdot\|$ denotes the L^2 norm.

⁸For a Markov chain with continuous state space to be recurrent, the following condition usually is required in addition to those in Lemma 1: there exists a subset of states which can be visited from any other state (in a finite number of steps) with positive probability. For $\{\mathbf{W}[t]\}$ studied here, the zero state constitutes such a subset because by assumption $\mathbb{P}(B_i^k[t] = 0) > 0$ for all queues i^k .

Now let service rates be allocated according to (13). Using (14), we derive the following upper bound on the expected one-step Lyapunov drift conditioned on $\mathbf{W}[t] = (\mathbf{u}_t, \mathbf{u}_{t-1})$:

$$\begin{aligned} \mathbb{E}[V(\mathbf{W}[t+1]) - V(\mathbf{W}[t]) | \mathbf{W}[t] = (\mathbf{u}_t, \mathbf{u}_{t-1})] \\ \leq 2\mathbf{u}_t' (\mathbf{a} - \tilde{\mathbf{R}}[t]) + 2(|\mathbf{z}| + \|\tilde{\mathbf{R}}[t]\|^2) \\ - \|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 \end{aligned}$$

where $\mathbf{z} \triangleq (\mathbb{E}(B_i^k[t])^2)_{i \in \mathcal{N}, k \in \mathcal{K}}$ is the vector of second moments of the random arrival rates and $|\cdot|$ denotes the L^1 norm. The detailed derivation of the above inequality is left to Appendix A.

Recall that the distributed power control algorithms in Section III-D increase the objective value $\sum_{(i,j)} b_{ij}^*[t] \cdot \underline{C}_{ij}(\mathbf{P}(\tau))$ with every iteration from time t to $t+1$. Therefore, given $\mathbf{W}[t] = (\mathbf{u}_t, \mathbf{u}_{t-1})$, we have

$$\begin{aligned} \mathbf{u}_t' \cdot \tilde{\mathbf{R}}[t] &\stackrel{(a)}{=} \int_t^{t+1} \sum_{(i,j)} b_{ij}^*[t] \cdot \bar{C}_{ij}(\mathbf{P}(\tau)) d\tau \\ &\stackrel{(b)}{>} \int_t^{t+1} \sum_{(i,j)} b_{ij}^*[t] \cdot \underline{C}_{ij}(\mathbf{P}(\tau)) d\tau \\ &\stackrel{(c)}{\geq} \sum_{(i,j)} b_{ij}^*[t] \cdot \underline{C}_{ij}(\mathbf{P}(t)) \\ &\geq \mathbf{u}_t' \cdot \tilde{\mathbf{R}}(t) = \mathbf{u}_t' \cdot \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1}). \end{aligned}$$

Equation (a) holds because the iterative MDB policy always allocates actual service rates according to (13). Inequality (b) is obvious since $\bar{C}_{ij}(\mathbf{P}(\tau)) > \underline{C}_{ij}(\mathbf{P}(\tau))$ for all (i, j) . Inequality (c) follows from the fact that $\sum_{(i,j)} b_{ij}^*[t] \cdot \underline{C}_{ij}(\mathbf{P}(\tau))$ is increasing over $[t, t+1)$.

Note that because the second moment vector $\mathbf{z}$ is assumed to be finite and the actual service rates $\tilde{\mathbf{R}}[t]$ must be bounded, we can find a finite constant λ such that $2(|\mathbf{z}| + \|\tilde{\mathbf{R}}[t]\|^2) \leq \lambda$. Thus, the conditional expected Lyapunov drift is upper bounded by

$$2\mathbf{u}_t' \cdot (\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})) - \|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 + \lambda.$$

Using the above Lyapunov function and the upper bound for the expected Lyapunov drift, we show the following main result.

Theorem 4: The iterative MDB policy with convergence time can stabilize all arrival processes whose average arrival rate vector $\mathbf{a}$ is in the interior of the stability region $\Lambda(\tilde{\mathcal{C}}(\Pi))$ induced by the high-SINR capacity region $\tilde{\mathcal{C}}(\Pi)$.

Guided by the Lyapunov drift criterion, the proof aims to find an $\varepsilon_0 > 0$ and a compact set $\mathcal{W}_0$ (which may depend on ε_0) which satisfy the conditions (16)–(17) for any average arrival rate vector $\mathbf{a} \in \text{int } \Lambda(\tilde{\mathcal{C}}(\Pi))$. As we have explained, $\Lambda(\tilde{\mathcal{C}}(\Pi))$ coincides with $\tilde{\mathcal{C}}(\Pi)$ (which we simply denote by $\mathcal{C}$) in the positive orthant. So from now on, we assume $\mathbf{a} \in \text{int } \mathcal{C}$. Note that condition (16) is always satisfied since the first and second moments of arrival rates as well as the service rate vector are bounded. Now consider the compact region characterized by

$$\mathcal{W}_0 = \{\mathbf{w} \in \mathbb{R}_+^M \times \mathbb{R}_+^M : V(\mathbf{w}) \leq \Omega\}. \quad (18)$$

Given $\varepsilon_0 > 0$, we need to specify a finite Ω and show that when $\mathbf{w}[t] = (\mathbf{u}_t, \mathbf{u}_{t-1}) \notin \mathcal{W}_0$

$$2\mathbf{u}'_t \cdot (\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})) - \|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 + \lambda \leq -\varepsilon_0. \quad (19)$$

Towards this objective, we devise a geometric method to relate the position of $\mathbf{u}_t$ and $\mathbf{u}_{t-1}$ in the state space to the value of the inner product $\mathbf{u}'_t \cdot [\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})]$. In order to reveal the insight underlying this approach, we first develop the methodology in $\mathbb{R}^2$. The generalization to higher dimensions as well as the proof for Theorem 4 can be found in Appendices B and C.

C. Geometric Analysis

In this section, we analyze vectors of arrival rates, service rates, and queue states geometrically. In view of condition (19), we characterize a neighborhood around $\mathbf{u}_t$ which has the following properties: if $\mathbf{u}_{t-1}$ lies in the neighborhood, then the first term $2\mathbf{u}'_t \cdot (\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1}))$ is substantially negative ($\leq -\lambda - \varepsilon_0$); if $\mathbf{u}_{t-1}$ lies outside the neighborhood (meaning that $\|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2$ is relatively large), then the second term $-\|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2$ is sufficiently negative for (19) to hold.

We assume an average arrival rate vector $\mathbf{a} \in \text{int } \mathcal{C}$. There must exist a point $\bar{\mathbf{a}} \in \text{bd } \mathcal{C}$, and a positive constant ε such that $\mathbf{a} + \varepsilon \cdot \mathbf{1} \leq \bar{\mathbf{a}}$. Therefore the point $\mathbf{e} = \mathbf{a} + \frac{\varepsilon}{2} \cdot \mathbf{1}$ is also in the interior of $\mathcal{C}$.

Given the current queue state vector $\mathbf{u}_t \geq \mathbf{0}$, the hyperplane $\mathcal{B}_e(\mathbf{u}_t) \triangleq \{\mathbf{x} : \mathbf{u}'_t \cdot \mathbf{x} = \mathbf{u}'_t \cdot \mathbf{e}\}$ is perpendicular to $\mathbf{u}_t$ and crosses the point $\mathbf{e}$. The intersection of halfspace $\mathcal{H}_e^+(\mathbf{u}_t) \triangleq \{\mathbf{x} : \mathbf{u}'_t \cdot \mathbf{x} \geq \mathbf{u}'_t \cdot \mathbf{e}\}$ with $\mathcal{C}$, denoted by $\mathcal{C}_e^+(\mathbf{u}_t)$, is closed and convex with nonempty interior [22].

Lemma 2: For $\mathbf{y} \in \mathcal{C}_e^+(\mathbf{u}_t)$, $\mathbf{u}'_t \cdot [\mathbf{a} - \mathbf{y}] \leq -\frac{\varepsilon}{2}\|\mathbf{u}_t\|$.

Proof: Since $\mathbf{y} \in \mathcal{H}_e^+(\mathbf{u}_t)$, by definition $\mathbf{u}'_t \cdot \mathbf{y} \geq \mathbf{u}'_t \cdot \mathbf{e}$. Thus

$$\begin{aligned} \mathbf{u}'_t \cdot [\mathbf{a} - \mathbf{y}] &\leq \mathbf{u}'_t \cdot [\mathbf{a} - \mathbf{e}] \\ &= -\frac{\varepsilon}{2}\|\mathbf{u}_t\| \leq -\frac{\varepsilon}{2}\|\mathbf{u}_t\|. \end{aligned}$$

The last inequality follows from $\|\mathbf{u}_t\| \geq \|\mathbf{u}_t\|$ since $\mathbf{u}_t \geq \mathbf{0}$. $\square$

1) *Two-Dimensional Heuristic:* Assume there are two queues in the network and index them by 1 and 2. In this subsection, all vectors, hyperplanes, surfaces, etc., are in $\mathbb{R}^2$. The hyperplane $\mathcal{B}_e(\mathbf{u}_t)$ must intersect $\text{bd } \mathcal{C}$ at two different points, as illustrated in Fig. 3. Let the two points be $\mathbf{f}_1$ and $\mathbf{f}_2$, where $\mathbf{f}_1$ is the upper-left one. Denote the hyperplane (which is a line in $\mathbb{R}^2$) tangent⁹ to $\mathcal{C}$ at $\mathbf{f}_1$ by $\mathcal{B}_{\mathbf{f}_1}(\mathbf{n}_1)$, where $\mathbf{n}_1$ is the unit normal vector of the tangent line. Specifically, we require $\mathbf{n}_1$ to be pointing outward from $\mathcal{C}$. Since $\mathcal{C}$ is not confined in $\mathbb{R}_+^2$, $\mathbf{f}_1$ is not necessarily nonnegative, and neither is $\mathbf{n}_1$. If there exist multiple tangent lines at $\mathbf{f}_1$, take $\mathbf{n}_1$ to be any one of them. Let the unit normal vector at $\mathbf{f}_2$ be $\mathbf{n}_2$, defined in the same manner. Let

$$\theta_1(\vec{\mathbf{u}}_t) = \arccos(\mathbf{n}'_1 \cdot \vec{\mathbf{u}}_t), \quad \theta_2(\vec{\mathbf{u}}_t) = \arccos(\mathbf{n}'_2 \cdot \vec{\mathbf{u}}_t)$$

⁹The tangent hyperplane contains $\mathbf{f}_1$ and defines a halfspace containing $\mathcal{C}$.

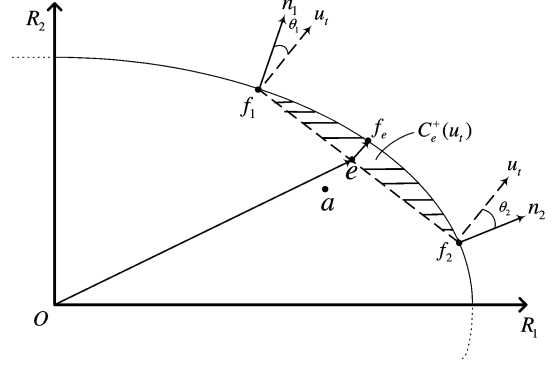


Fig. 3. The geometry when $\mathcal{B}_e(\mathbf{u}_t)$ intersects $\text{bd } \mathcal{C}$ at two different points in $\mathbb{R}_+^2$.

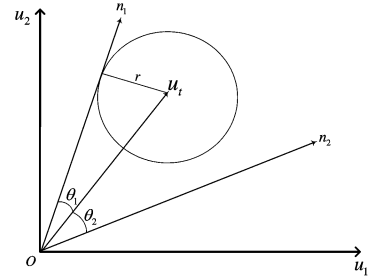


Fig. 4. The geometry of $\mathbf{u}_{t-1}$ lying in the neighborhood of $\mathbf{u}_t$, where $r = \|\mathbf{u}_t\| \cdot \alpha(\vec{\mathbf{u}}_t)$.

where $\vec{\mathbf{u}}_t$ stands for the normalized vector of $\mathbf{u}_t$. Since $\mathbf{e} \in \text{int } \mathcal{C}$, $\mathbf{n}_1$ and $\mathbf{n}_2$ can never be parallel to $\vec{\mathbf{u}}_t$. Thus

$$\mathbf{n}'_1 \cdot \vec{\mathbf{u}}_t < 1, \quad \mathbf{n}'_2 \cdot \vec{\mathbf{u}}_t < 1$$

and $\theta_1(\vec{\mathbf{u}}_t) > 0$, $\theta_2(\vec{\mathbf{u}}_t) > 0$. Moreover, $\theta_1(\vec{\mathbf{u}}_t)$ and $\theta_2(\vec{\mathbf{u}}_t)$ are bounded away from zero for all $\mathbf{u}_t$. To see this, we make use of Fig. 3 again. The point $\mathbf{f}_e$ is on the boundary and the vector $\mathbf{f}_e - \mathbf{e}$ is parallel to $\mathbf{u}_t$. By simple geometry, the convexity of the capacity region implies $\theta_1(\mathbf{u}_t) \geq \arctan(\|\mathbf{f}_e - \mathbf{e}\|/\|\mathbf{f}_1 - \mathbf{e}\|)$. Because $\mathbf{e}$ is an interior point, $\|\mathbf{f}_e - \mathbf{e}\| \geq \xi > 0$. Moreover, $\|\mathbf{f}_1 - \mathbf{e}\| \leq D < \infty$ since $\mathcal{C}$ is a bounded region. Therefore, $\theta_1(\vec{\mathbf{u}}_t) \geq \arctan(\xi/D) > 0$.

The same is true for $\theta_2(\vec{\mathbf{u}}_t)$. Thus, we can construct a nonempty cone emanating from the origin sweeping from the direction of vector $\mathbf{u}_t$ clockwise by $\theta_2(\vec{\mathbf{u}}_t)$ and counterclockwise by $\theta_1(\vec{\mathbf{u}}_t)$. Such a cone always contains $\mathbf{u}_t$ in its strict interior. This is illustrated in Fig. 4.

We consider the following two cases. First, if $\|\mathbf{u}_t - \mathbf{u}_{t-1}\|/\|\mathbf{u}_t\| \leq \sin[\min\{\theta_1(\vec{\mathbf{u}}_t), \theta_2(\vec{\mathbf{u}}_t), \pi/2\}] \triangleq \alpha(\vec{\mathbf{u}}_t)$, then the pair of points $(\mathbf{u}_t, \mathbf{u}_{t-1})$ both lie in the cone described above. In this case, $\mathbf{u}_{t-1}$ is said to be in the *neighborhood* of $\mathbf{u}_t$. See Fig. 4.

Let α be the infimum of $\alpha(\vec{\mathbf{u}}_t)$ over all nonnegative unit vector $\vec{\mathbf{u}}_t$. Because all $\theta_1(\vec{\mathbf{u}}_t)$ and $\theta_2(\vec{\mathbf{u}}_t)$ are strictly positive, α must be strictly positive. If $\|\mathbf{u}_t - \mathbf{u}_{t-1}\|/\|\mathbf{u}_t\| \leq \alpha$, $\mathbf{u}_{t-1}$ is also in the cone with $\mathbf{u}_t$. In this case, the hyperplane of normal

vector $\mathbf{u}_{t-1}$ tangent to the capacity region $\mathcal{C}$ touches $\text{bd } \mathcal{C}$ at $\tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})$ somewhere between $\mathbf{f}_1$ and $\mathbf{f}_2$, i.e., $\tilde{\mathbf{R}}^*(\mathbf{u}_{t-1}) \in \mathcal{C}_e^+(\mathbf{u}_t)$. By Lemma 2, the inner product $\mathbf{u}_t' \cdot [\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})] \leq -\frac{\varepsilon}{2} \|\mathbf{u}_t\|$. Then for all $\mathbf{w}[t]$ such that $V(\mathbf{w}[t]) > (1 + \alpha^2) \cdot (\varepsilon_0 + \lambda)^2 / \varepsilon^2 \equiv \Omega_1$, $\|\mathbf{u}_t\| > (\varepsilon_0 + \lambda) / \varepsilon$, and therefore

$$\begin{aligned} 2\mathbf{u}_t' \cdot (\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})) - \|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 + \lambda \\ \leq 2\mathbf{u}_t' \cdot (\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})) + \lambda < -\varepsilon_0 \end{aligned}$$

which is the desired condition (19).

If $\|\mathbf{u}_t - \mathbf{u}_{t-1}\| / \|\mathbf{u}_t\| > \alpha$ and assume $\|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 = \omega$, then

$$\begin{aligned} 2\mathbf{u}_t' \cdot (\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})) - \|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 + \lambda \\ \leq 2\|\mathbf{u}_t\| \|\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})\| - \omega + \lambda \\ < 2\sqrt{\omega/\alpha^2} \sqrt{\lambda/2} - \omega + \lambda \\ = \sqrt{2\omega\lambda}/\alpha - \omega + \lambda. \end{aligned}$$

Define

$$\omega_2 = \inf \left\{ \omega > 0 : \sqrt{2\omega\lambda}/\alpha - \omega + \lambda \leq -\varepsilon_0 \right\}. \quad (20)$$

Then for all $\mathbf{w}[t]$ such that $V(\mathbf{w}[t]) > (1 + 1/\alpha^2)\omega_2 \equiv \Omega_2$, $\|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 > \omega_2$ and (19) holds.

Combining the above two cases and letting $\Omega = \max\{\Omega_1, \Omega_2\}$, we see that the region specified in (18) satisfies Lemma 1 and Theorem 4 follows.

V. NUMERICAL EXPERIMENTS

To assess the practical performance of the node-based distributed MDB policy in stochastic wireless networks, we conduct the following simulation to compare the total backlogs resulting from the same arrival processes under different MDB schemes.

Our scheme iteratively adjusts the transmission powers during a slot to find the optimal rates (under the high-SINR approximation) for the queue state at the beginning of a slot. As a consequence, the MDB optimization is done with delayed queue state information. The transmission rates keep changing with time, and the optimal rates are achieved only at the end (beginning) of the current (next) slot. Recently, Giannoulis *et al.* [13] proposed another distributed power control algorithm to implement the MDB policy in CDMA networks. Instead of converging to the optimal solution for the current MDB problem (also under the high-SINR approximation), their scheme updates the link powers based on the present queue state only once in a slot. The new queue state at the beginning of the next slot is used for the subsequent iteration. To highlight the above difference, we refer to our method as “iterative MDB with convergence”, and the method studied in [13] as “iterative MDB without convergence”.

For a single run of the experiment, we use a network with N nodes uniformly distributed in a disc of unit radius. Nodes i and j share a link if their distance $d(i, j)$ is less than $2.5/\sqrt{N}$, so that the average number of a node's neighbors remains constant with N . The path gain is modeled as $h_{ij} = d(i, j)^{-4}$. The processing

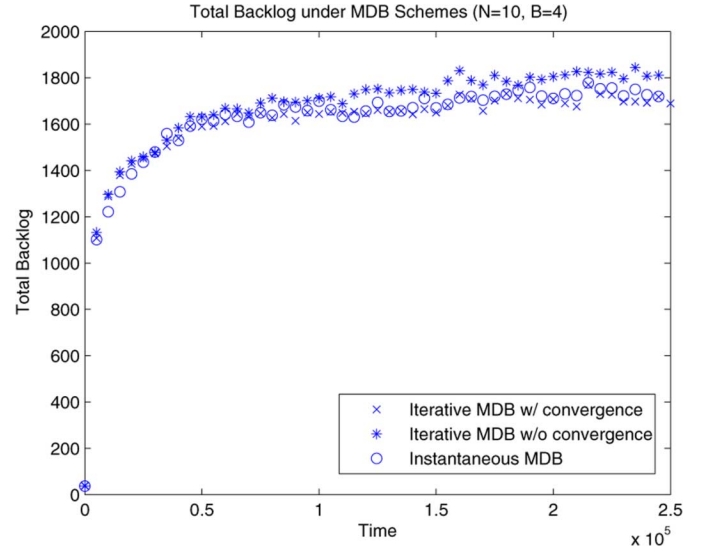


Fig. 5. Total backlogs under three MDB schemes ($N = 10$, $B = 4$).

gain of the CDMA system is $K = 10^5$. All nodes are subject to the common total power constraint $\hat{P}_i = 100$ and AWGN with power $N_i = 0.1$.

Each node is the source node of one session with the destination chosen from the other $N - 1$ nodes uniformly at random. At the beginning of every slot, the new arrivals of all N sessions are independent Poisson random variables with the same parameter B . At any instant, queues are served at actual link capacities determined by the instantaneous power configuration. As an approximation, we assume the iterative MDB scheme converges after 50 iterations of the *PA* and *PC* algorithms. The convergence time is taken to be the length of a slot, as in Section IV. The network performance is investigated under each of the MDB schemes with the same set of arrival processes. The total backlog in the network is recorded after every slot. Fig. 5 shows the backlog curves generated by the three schemes after averaging 10 independent runs with the parameters $N = 10$ and $B = 4$. Each point on a given curve represents the average total backlog sampled at boundaries of time slots. The three schemes all manage to stabilize the network queues in the long run (we show the trajectories up to 2.5×10^5 time slots in Fig. 5). However, the iterative MDB scheme with convergence and the instantaneous MDB scheme result in lower queue occupancy, hence lower delay, than the iterative MDB scheme without convergence.

For a closer look at the performance of our iterative MDB scheme, we show in Fig. 6 the average trajectory of the MDB objective value $\sum_{(i,j)} b_{ij}^*(\tau) R_{ij}(\tau)$ generated by our scheme. Note that the objective value is computed using the high-SINR capacity formula. In Fig. 6, one iteration involves every node updating its power allocation and power control variables using *PA* and *PC* once, respectively. The MDB values $b_{ij}^*(\tau)$ at the τ th iteration are computed based on the last sample of queue state $\mathbf{U}[t]$ where $tT < \tau \leq (t+1)T$. In this simulation, the length of a time slot T is taken to be 50, i.e., the queue state is sampled every 50 iterations. It can be observed from the plot that our MDB scheme constantly improves the objective value within a

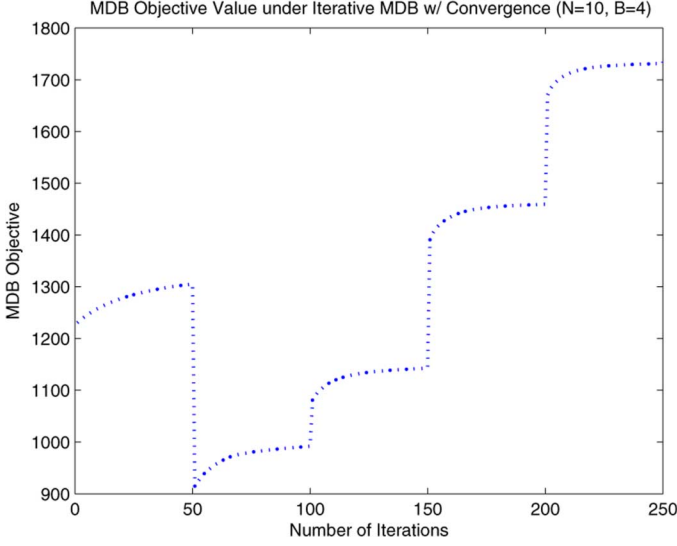


Fig. 6. MDB objective value under iterative MDB with convergence ($N = 10$ and $B = 10$).

time slot and readjusts the link rates toward the new optimum once the queue state is updated after each slot.

VI. CONCLUSION

In this work, we studied the distributed implementation of the Maximum Differential Backlog (backpressure) algorithm within interference-limited CDMA wireless networks with random traffic arrivals. In the first half of the paper, we developed a set of node-based iterative power control algorithms for solving the MDB optimization problem using the high-SINR approximation of the link capacities. Our algorithms are based on the scaled gradient projection method. We showed that the algorithms can be implemented in a distributed manner using low communication overhead, without the need for time sharing and scheduling. Although the high-SINR capacity approximation is used in optimizing the transmission powers, the actual link capacities are still used to service the network queues.

Since the power control algorithms require nonnegligible time to converge, even the optimal high-SINR capacities for any given queue state can only be found iteratively over time. In the second half of the paper, we investigated the impact of both the high-SINR approximation and the convergence time on the stability region achievable by the iterative MDB policy. Using a new geometric approach to analyze the expected Lyapunov drift, we proved that the iterative MDB policy can stabilize all arrival processes whose average arrival rates are within the stability region induced by the high-SINR capacity region, as long as the second moments of arrival process are bounded.

The stability region induced by the high-SINR capacity region is a subset of the stability region induced by the actual capacity region. The gap between the two, however, is negligible in the high-SINR regime. Thus, our iterative power control algorithms represent a distributed and close-to-optimal solution to the problem of throughput optimal control of CDMA wireless networks with random traffic arrivals in the high-SINR regime.

APPENDIX

Derivation of Lyapunov Drift: By definition, the difference of Lyapunov values $V(\mathbf{W}[t+1])$ and $V(\mathbf{W}[t])$ can be written as

$$\begin{aligned} V(\mathbf{W}[t+1]) - V(\mathbf{W}[t]) &= \|\mathbf{U}[t+1]\|^2 - \|\mathbf{U}[t]\|^2 + \|\mathbf{U}[t+1] - \mathbf{U}[t]\|^2 \\ &\quad - \|\mathbf{U}[t] - \mathbf{U}[t-1]\|^2 \\ &= 2\mathbf{U}[t+1]' \cdot (\mathbf{U}[t+1] - \mathbf{U}[t]) \\ &\quad - \|\mathbf{U}[t] - \mathbf{U}[t-1]\|^2. \end{aligned}$$

Using relation (14), we have

$$\begin{aligned} &\mathbf{U}[t+1]' \cdot (\mathbf{U}[t+1] - \mathbf{U}[t]) \\ &\leq \left((\mathbf{U}[t] - \tilde{\mathbf{R}}[t] + \mathbf{B}[t])^+ \right)' \cdot \left((\mathbf{U}[t] - \tilde{\mathbf{R}}[t] + \mathbf{B}[t])^+ - \mathbf{U}[t] \right) \\ &\leq (\mathbf{U}[t] - \tilde{\mathbf{R}}[t] + \mathbf{B}[t])' \cdot (\mathbf{B}[t] - \tilde{\mathbf{R}}[t]) \\ &\leq \mathbf{U}[t]' \cdot (\mathbf{B}[t] - \tilde{\mathbf{R}}[t]) + \|\mathbf{B}[t]\|^2 + \|\tilde{\mathbf{R}}[t]\|^2. \end{aligned}$$

Therefore, we finally obtain

$$\begin{aligned} V(\mathbf{W}[t+1]) - V(\mathbf{W}[t]) &\leq 2\mathbf{U}[t]' \cdot (\mathbf{B}[t] - \tilde{\mathbf{R}}[t]) + 2 \left(\|\mathbf{B}[t]\|^2 + \|\tilde{\mathbf{R}}[t]\|^2 \right) \\ &\quad - \|\mathbf{U}[t] - \mathbf{U}[t-1]\|^2. \end{aligned}$$

Geometric Analysis in $\mathbb{R}^M$: We now generalize our geometric analysis in Section IV-C to M -dimensional space. We retain the notation from Section IV-C.

Analogous to the argument used in the two-dimensional case, we focus on characterizing the neighborhood of $\mathbf{u}_t$.

Lemma 3: For any $\mathbf{u}_t \geq \mathbf{0}$, there exists a region $\mathcal{K}(\mathbf{u}_t) \subset \mathbb{R}_+^M$ such that

- 1) $\mathbf{u}_t \in \mathcal{K}(\mathbf{u}_t)$;
- 2) $\mathcal{K}(\mathbf{u}_t)$ has nonempty and convex interior relative to any one-dimensional affine space containing $\mathbf{u}_t$;
- 3) for all $\mathbf{u}_{t-1} \in \mathcal{K}(\mathbf{u}_t)$, the optimal rate vector $\tilde{\mathbf{R}}^*(\mathbf{u}_{t-1})$ with respect to $\mathbf{u}_{t-1}$ is in $\mathcal{C}_{\mathcal{E}}^+(\mathbf{u}_t)$.

Note that $\mathcal{K}(\mathbf{u}_t)$ is the M -dimensional analogue of the circle $\mathcal{S}(\mathbf{u}_t, r)$ of radius r around $\mathbf{u}_t$ in Fig. 4. To facilitate the proof, define the set of feasible unit incremental vectors around a non-negative unit vector $\bar{\mathbf{u}}$ as

$$\begin{aligned} \Delta_{\bar{\mathbf{u}}} &\triangleq \left\{ \Delta = (\Delta_1, \dots, \Delta_M) : \right. \\ &\quad \left. \|\Delta\| = 1, \text{ and } \Delta_i^k \geq 0 \text{ if } \bar{u}_i^k = 0 \right\}. \end{aligned}$$

Proof of Lemma 3: Each $\Delta \in \Delta_{\bar{\mathbf{u}}}$ spans a one-dimensional affine space containing $\mathbf{u}_t$. It is sufficient to show that given any $\Delta \in \Delta_{\bar{\mathbf{u}}}$, there exists $\bar{\delta} > 0$ such that for all $\delta \in [0, \bar{\delta}]$ and $\mathbf{f} \in \mathcal{C}$ satisfying

$$(\mathbf{u}_t + \delta\Delta)' \cdot \mathbf{f} \geq (\mathbf{u}_t + \delta\Delta)' \cdot \mathbf{R}, \quad \forall \mathbf{R} \in \mathcal{C} \quad (21)$$

we have $\mathbf{f} \in \mathcal{C}_e^+(\mathbf{u}_t)$.

We prove the claim by construction. We make use of the dominant point $\bar{\mathbf{a}}$ of $\mathbf{a}$ such that $\mathbf{a} + \varepsilon \cdot \mathbf{1} \leq \bar{\mathbf{a}}$ (also $\mathbf{e} + \varepsilon/2 \cdot \mathbf{1} \leq \bar{\mathbf{a}}$). Define the parameter

$$d(\Delta) \triangleq \max_{\mathbf{R} \in \mathcal{C}} \Delta' \cdot (\mathbf{R} - \bar{\mathbf{a}}) \quad (22)$$

which is at least zero (by setting $\mathbf{R} = \bar{\mathbf{a}}$ in the objective function). It is possibly equal to zero, and must be bounded from above, because Δ is a unit vector and the optimization region $\mathcal{C}$ is compact.

Now consider

$$\bar{\delta} = \frac{\varepsilon \|\mathbf{u}_t\|}{2d(\Delta)}$$

which by the above analysis is positive. Because $\mathcal{C}$ is convex and compact, for any $\delta \in [0, \bar{\delta}]$ there exists at least one $\mathbf{f}$ satisfying (21). Picking any one such $\mathbf{f}$ and specifically letting $\mathbf{R} = \bar{\mathbf{a}}$ on the RHS of (21), we have

$$(\mathbf{u}_t + \delta \Delta)' \cdot \mathbf{f} \geq (\mathbf{u}_t + \delta \Delta)' \cdot \bar{\mathbf{a}}.$$

By using the inequality

$$\mathbf{u}_t' \cdot \bar{\mathbf{a}} \geq \mathbf{u}_t' \cdot \mathbf{e} + \frac{\varepsilon}{2} \|\mathbf{u}_t\| \geq \mathbf{u}_t' \cdot \mathbf{e} + \frac{\varepsilon}{2} \|\mathbf{u}_t\|$$

we have

$$\begin{aligned} \mathbf{u}_t' \cdot \mathbf{f} &\geq \mathbf{u}_t' \cdot \bar{\mathbf{a}} - \delta \Delta' \cdot (\mathbf{f} - \bar{\mathbf{a}}) \\ &\geq \mathbf{u}_t' \cdot \mathbf{e} + \frac{\varepsilon}{2} \|\mathbf{u}_t\| - \delta \Delta' \cdot (\mathbf{f} - \bar{\mathbf{a}}) \\ &\geq \mathbf{u}_t' \cdot \mathbf{e} + \frac{\varepsilon}{2} \|\mathbf{u}_t\| - \bar{\delta} \max_{\mathbf{R} \in \mathcal{C}} \Delta' \cdot (\mathbf{R} - \bar{\mathbf{a}}) \\ &= \mathbf{u}_t' \cdot \mathbf{e} + \frac{\varepsilon}{2} \|\mathbf{u}_t\| - \frac{\varepsilon \|\mathbf{u}_t\|}{2d(\Delta)} \cdot d(\Delta) \\ &= \mathbf{u}_t' \cdot \mathbf{e}. \end{aligned}$$

Thus, we can conclude that $\mathbf{f} \in \mathcal{C}_e^+(\mathbf{u}_t)$. Since $\mathbf{f}$ is chosen arbitrarily, the claim at the beginning of the proof is proved.

Finally, define $\mathcal{K}(\mathbf{u}_t)$ as

$$\left\{ \mathbf{u}_{t-1} \in \mathbb{R}_+^M : \|\mathbf{u}_{t-1} - \mathbf{u}_t\| \leq \frac{\varepsilon \|\mathbf{u}_t\|}{2d(\frac{\mathbf{u}_{t-1} - \mathbf{u}_t}{\|\mathbf{u}_{t-1} - \mathbf{u}_t\|})} \right\} \quad (23)$$

where $d(\cdot)$ is defined as in (22). To accommodate the special case of $\mathbf{u}_{t-1} = \mathbf{u}_t$, we define $d(\mathbf{0}) = 0$. It is easily verified that the so-constructed $\mathcal{K}(\mathbf{u}_t)$ is a valid neighborhood of $\mathbf{u}_t$, as required by the lemma. $\square$

Proof for Theorem 4: If

$$\frac{\|\mathbf{u}_{t-1} - \mathbf{u}_t\|}{\|\mathbf{u}_t\|} \leq \frac{\varepsilon}{2 \sup_{\vec{\mathbf{u}} \geq \mathbf{0}} d(\vec{\mathbf{u}})} \equiv \alpha$$

then $\mathbf{u}_{t-1} \in \mathcal{K}(\mathbf{u}_t)$ where $\mathcal{K}(\mathbf{u}_t)$ is defined in (23). In this case, for all $\mathbf{w}[t]$ such that $V(\mathbf{w}[t]) > (1 + \alpha^2) \cdot (\varepsilon_0 + \lambda)^2 / \varepsilon^2 \equiv \Omega_1$, $\|\mathbf{u}_t\| > (\varepsilon_0 + \lambda) / \varepsilon$, and therefore

$$\begin{aligned} &2\mathbf{u}_t' \cdot \left(\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1}) \right) - \|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 + \lambda \\ &\leq 2\mathbf{u}_t' \cdot \left(\mathbf{a} - \tilde{\mathbf{R}}^*(\mathbf{u}_{t-1}) \right) + \lambda \\ &< -\varepsilon_0 \end{aligned}$$

which is the desired condition (19).

If $\|\mathbf{u}_t - \mathbf{u}_{t-1}\| / \|\mathbf{u}_t\| > \alpha$, define ω_2 as in (20), then for all $\mathbf{w}[t]$ such that $V(\mathbf{w}[t]) > (1 + 1/\alpha^2) \omega_2 \equiv \Omega_2$, $\|\mathbf{u}_t - \mathbf{u}_{t-1}\|^2 > \omega_2$ and (19) holds.

Combining the above two cases and letting $\Omega = \max\{\Omega_1, \Omega_2\}$, we see that the region specified in (18) satisfies Lemma 1 and therefore the queueing system is stable under any average arrival rate vector $\mathbf{a} \in \text{int } \mathcal{C}$. $\square$

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. Autom. Control*, vol. 37, no. 12, pp. 1936–1948, Dec. 1992.
- [2] M. Neely, E. Modiano, and C. Rohrs, "Dynamic power allocation and routing for time varying wireless networks," in *Proc. IEEE INFOCOM*, Mar. 2003, vol. 1, pp. 745–755.
- [3] M. Neely, E. Modiano, and C. Rohrs, "Dynamic power allocation and routing for time-varying wireless networks," *IEEE J. Sel. Areas Commun.*, vol. 23, no. 1, pp. 89–103, Jan. 2005.
- [4] A. Eryilmaz and R. Srikant, "Fair resource allocation in wireless networks using queue-length-based scheduling and congestion control," in *Proc. IEEE INFOCOM*, Mar. 2005, vol. 3, pp. 1794–1803.
- [5] L. Tassiulas, "Linear complexity algorithms for maximum throughput in radio networks and input queued switches," in *Proc. IEEE INFOCOM*, Mar. 1998, vol. 2, pp. 533–539.
- [6] P. Chaporkar, K. Kar, and S. Sarkar, "Throughput guarantees through maximal scheduling in wireless networks," presented at the 2005 Allerton Conf. Commun., Contr., Computing, Sep. 2005.
- [7] X. Wu, R. Srikant, and J. R. Perkins, "Queue-length stability of maximal greedy schedules in wireless networks," in *Proc. Workshop Inf. Theory Appl.*, Feb. 2006.
- [8] X. Wu and R. Srikant, "Regulated maximal matching: A distributed scheduling algorithm for multi-hop wireless networks with node-exclusive spectrum sharing," in *Proc. IEEE Conf. Dec. Contr.*, Dec. 2005, pp. 5342–5347.
- [9] L. Bui, A. Eryilmaz, R. Srikant, and X. Wu, "Joint asynchronous congestion control and distributed scheduling for multi-hop wireless networks," in *Proc. IEEE INFOCOM*, Apr. 2006, pp. 1–12.
- [10] E. Modiano, D. Shah, and G. Zussman, "Maximizing throughput in wireless networks via gossiping," in *Proc. Joint Int. Conf. Measur. Model. Comput. Syst.*, 2006, pp. 27–38.
- [11] X. Lin and N. Shroff, "The impact of imperfect scheduling on cross-layer rate control in wireless networks," in *Proc. IEEE INFOCOM*, Mar. 2005, vol. 3, pp. 1804–1814.
- [12] M. J. Neely, "Energy optimal control for time varying wireless networks," in *Proc. IEEE INFOCOM*, Mar. 2005, vol. 1, pp. 572–583.
- [13] A. Giannoulis, K. Tsoukatos, and L. Tassiulas, "Lightweight cross-layer control algorithms for fairness and energy efficiency in CDMA ad-hoc networks," in *Proc. IEEE WiOpt*, Apr. 2006, pp. 1–8.
- [14] D. Tse and P. Viswanath, *Fundamentals of Wireless Communication*. Cambridge, U.K.: Cambridge Univ. Press, 2004.
- [15] M. Johansson, L. Xiao, and S. Boyd, "Simultaneous routing and power allocation in CDMA wireless data networks," in *Proc. IEEE Int. Conf. Commun.*, May 2003, vol. 1, pp. 51–55.
- [16] M. Chiang, "To layer or not to layer: Balancing transport and physical layers in wireless multihop networks," in *Proc. IEEE INFOCOM*, Mar. 2004, vol. 4, pp. 2525–2536.
- [17] Y. Xi, "Distributed resource allocation in communication networks," Ph.D. dissertation, Yale University, New Haven, CT, 2008.

- [18] D. P. Bertsekas, *Nonlinear Programming*, second ed. New York: Athena Scientific, 1999.
- [19] M. J. Neely, E. Modiano, and C. Rohrs, "Tradeoffs in delay guarantees and computation complexity for $N \times N$ packet switches," in *Proc. Conf. Inf. Sci. Syst.*, Princeton, NJ, Mar. 2002.
- [20] L. Tassiulas and A. Ephremides, "Throughput properties of a queueing network with distributed dynamic routing and flow control," *Adv. Appl. Prob.*, vol. 28, pp. 285–307, Mar. 1996.
- [21] S. Asmussen, *Applied Probability and Queues*. New York: Wiley, 1987.
- [22] H. Eggleston, *Convexity*. Cambridge, U.K.: Cambridge Univ. Press, 1977.

Yufang Xi (S'06) received the B.Sci. degree in electronic engineering from Tsinghua University, Beijing, China, in 2003 and the M.Sci., M.Phil., and Ph.D. degrees in electrical engineering from Yale University, New Haven, CT, in 2005, 2006, and 2008, respectively.

His research interests include cross-layer optimization and distributed resource allocation algorithms for wireless networks.

Edmund M. Yeh (M'09) received the B.S. degree in electrical engineering with Distinction from Stanford University, Stanford, CA, in 1994, the M.Phil. degree in engineering from the University of Cambridge, Cambridge, U.K., in 1995, and the Ph.D. degree in electrical engineering and computer science from the Massachusetts Institute of Technology (MIT), Cambridge, in 2001.

Since July 2001, he has been on the faculty at Yale University, New Haven, CT, where he is currently an Associate Professor of Electrical Engineering and Computer Science. He has been visiting faculty at MIT, Princeton University, Princeton, NJ; University of California at Berkeley; and Swiss Federal Institute of Technology, Lausanne, Switzerland.

Dr. Yeh is a member of Phi Beta Kappa, and Tau Beta Pi. He is a recipient of the Army Research Office (ARO) Young Investigator Program (YIP) Award (2003), the Winston Churchill Scholarship (1994), the National Science Foundation and Office of Naval Research Fellowships (1994) for graduate study, the Frederick E. Terman Award from Stanford University (1994), and the Barry M. Goldwater Scholarship from the United States Congress (1993).