

A Low-Complexity QoS-Aware Proportional Fair Multicarrier Scheduling Algorithm for OFDM Systems

Zhen Kong, Yu-Kwong Kwok, *Senior Member, IEEE*, and Jiangzhou Wang, *Senior Member, IEEE*

Abstract—Orthogonal frequency-division multiplexing (OFDM) systems are the major cellular platforms for supporting ubiquitous high-speed mobile applications. However, a number of research challenges remain to be tackled. One of the most important challenges is the design of a judicious packet scheduler that will make efficient use of the spectrum bandwidth. Due to the multicarrier nature of the OFDM systems, the applicability and performance of traditional wireless packet-scheduling algorithms, which are usually designed for single-carrier systems, are largely unknown. In this paper, we propose a new quality-of-service (QoS)-aware proportional fairness (QPF) packet-scheduling policy with low complexity for the downlink of multiuser OFDM systems to allocate radio resources among users. Our proposed algorithm is based on a cross-layer design in that the scheduler is aware of both the channel (i.e., physical layer) and the queue state (i.e., data link layer) information to achieve proportional fairness while maximizing each user's packet-level QoS performance. The simulation results show that the proposed QPF algorithm is efficient in terms of average system throughput, packet-dropping probability, and packet delay, while maintaining adequate fairness among users with relatively low scheduling overhead.

Index Terms—Cross-layer design, orthogonal frequency-division multiplexing (OFDM) systems, packet scheduling, proportional fairness (PF), quality of service (QoS).

I. INTRODUCTION

WITH the rapid development and integration of Internet and wireless communication networks, the forthcoming fourth-generation (4G) mobile systems are envisioned to support the outburst and popularity of high-speed multimedia packet-based applications. Such services usually exhibit a large variety of quality-of-service (QoS) requirements, such as transmission rate, delay, and packet-dropping ratio, which are difficult to be satisfied in a wireless environment due to the limited radio resource, the time-varying channel condition, and

the resource contention among multiple users. Recently, the orthogonal frequency-division multiplexing (OFDM) system has been identified as a promising wireless interface solution for 4G systems due to its high capacity to combat channel fading and support for a high data rate (HDR) [1]. Moreover, because it can provide a high degree of flexibility for resource control on different subcarriers, an adaptive packet scheduling for a multiuser OFDM system is widely considered as an important strategy to improve system performance.

In fact, the design of an efficient packet scheduler for a wireless system is a difficult task that typically involves a large number of conflicting requirements, which must be analyzed and weighted before a balanced solution can be implemented [2]. On one hand, the scheduler must be efficient in utilizing the radio resource since the wireless spectrum is the most precious resource in wireless communication systems. On the other hand, the services should fairly be scheduled so as to guarantee a certain level of service to users with low average channel conditions. Among the various fairness criteria, the proportional fairness (PF) scheduling is widely considered as a good solution because it provides an attractive tradeoff between the maximum average throughput and the user fairness by exploiting the temporal diversity and game-theoretic equilibrium in a multiuser environment [3]. Specifically, when the resource allocation is said to be proportionally fair, it is then impossible to enhance the throughput of any particular user by $x\%$ without decreasing some other users' total throughput by $x\%$ [4]. Under this consideration, some scheduling algorithms are proposed to achieve PF in wireless systems. For example, a well-known PF scheduling algorithm was implemented in [5] for its HDR system.

Unfortunately, there are some problems that have not been well tackled if we want to adopt PF in OFDM systems. First, for instance, the PF scheduling method in [5] can only be used in a single-carrier (SC) situation, and only one user is allowed to transmit at a time. Although a PF scheduler was proposed for the OFDM system in [6], it was based on the PF criterion derived from the SC system without considering the PF criteria for a multicarrier (MC) system. The optimal criteria for PF scheduling in MC systems has been studied in [7], but its practical implementation is very complex. Thus, there is a pressing need in developing an alternative PF scheduling algorithm for practical implementation in multiuser OFDM systems. Second, a typical PF scheduler only considers the performances of average system throughput and fairness. Such a PF scheduler

Manuscript received January 31, 2008; revised July 27, 2008. First published November 21, 2008; current version published May 11, 2009. The review of this paper was coordinated by Dr. P. Lin.

Z. Kong was with the Department of Electrical and Electronic Engineering, University of Hong Kong, Hong Kong. He is now with the Department of Electrical and Computer Engineering, Colorado State University, Fort Collins, CO 80523-1373 USA (e-mail: zkong@eee.hku.hk; zhen.kong@colostate.edu).

Y.-K. Kwok is with the Department of Electrical and Computer Engineering, Colorado State University, Fort Collins, CO 80523-1373 USA (e-mail: ricky.kwok@colostate.edu).

J. Wang is with the Department of Electronics, University of Kent, CT2 7NT Kent, U.K. (e-mail: j.z.wang@kent.ac.uk).

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/TVT.2008.2009874

usually does not consider other QoS performance metrics, such as packet-dropping probability (PDP) and packet delay. To meet the QoS requirements of the current wireless multimedia services, we have to improve the packet level QoS in the context of PF scheduling. Furthermore, these research efforts are usually based on the assumption of bit stream transmission [1] or infinite queue length [8] without considering realistic queueing conditions, which makes the proposed schemes impractical in real-life wireless packet transmission systems. If we want to meet the packet-level QoS performance specifications, the data link layer queueing performance should also be analyzed on top of the channel conditions so as to allocate the radio resources in a cross-layer manner.

Some work that is related to the packet scheduling in OFDM systems takes the queue conditions into consideration as well as the channel dynamics. For instance, the average delay performance in a multiuser OFDM was studied in [8]. However, the queue size is assumed to be infinite, which is clearly impractical. In [9], a unified packet-scheduling method is proposed by using a unified priority function to take each user's packet delay, packet type, and channel condition into consideration. In each subcarrier, the user with the largest priority can transmit its packets, but the fairness performance is not considered. The upper and lower bounds to the buffer overflow probability and mean packet waiting time in a single-user MC system were theoretically derived in [10] through a cross-layer analysis, but how to approach these bounds through a practical packet-scheduling method was not analyzed. In [11], the single-user OFDM scheduling framework is proposed as a Markov decision process, and the authors proposed both optimal and suboptimal algorithms to minimize the power consumption and satisfy the long-term average delay. Our model is different since we consider the throughput instead of the power and assume a multiuser scheduling instead of a single user. A utility-based subcarrier assignment method was proposed in [12] to balance the efficiency and fairness of resource allocation, whereas the packet-level QoS performance metrics were not considered. A cross-layer packet-scheduling algorithm for multiuser OFDM systems was developed in [13] to improve the performance on throughput and packet-dropping rate, but only one packet from one user can be scheduled at any time slot. By contrast, in our model, several packets from different users can simultaneously be scheduled as so to further exploit the multiuser diversity.

In this paper, we focus on an efficient cross-layer approach to allocate the subcarrier and schedule the packets in the down-link of multiuser OFDM systems. Specifically, we consider to achieve PF while improving the individual user's packet level QoS performance. In physical layer OFDM transmission, the data rate is determined by the wireless channel condition and the adaptive modulation level. At data link layer, the relationship among packet arrival rate, queue length, and packet transmission rate is set up to analyze the packet-level QoS performance. Thus, the packet scheduling and the corresponding subcarrier allocation are related not only to the channel state information (CSI, such as SNR) observed in the physical layer but also to the queue state information (QSI, such as queue length and packet arrival rate) obtained at the data link layer. We first discuss the PF scheduling criteria in the MC system and pro-

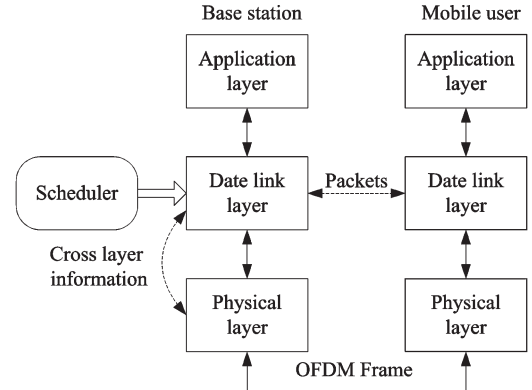


Fig. 1. Cross-layer scheduling architecture.

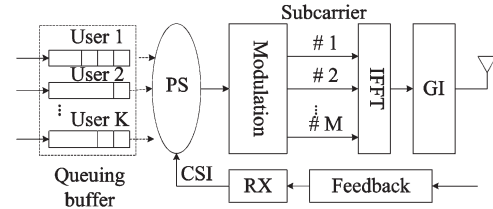


Fig. 2. Cross-layer OFDM packet scheduling system model.

pose an efficient greedy PF method to realize PF in an OFDM environment. After that, we analyze the average PDP and the average packet delay for a wireless multimedia traffic by investigating the queueing performance, and present a QoS-aware PF (QPF) algorithm to achieve the PF among users while improving these users' QoS. The cross-layer QPF algorithm consists of two steps. In the first step, the greedy PF method is used to achieve PF, and then a subcarrier reassignment procedure is utilized to improve the packet-level QoS performances according to the queue analysis. The simulation results show that the proposed QPF algorithms with low complexity can achieve good performances in terms of average system throughput, fairness, and packet-level QoS performances.

The remainder of this paper is organized as follows. In Section II, we describe the cross-layer system model. In Section III, an efficient MC PF scheduling algorithm for a multiuser OFDM system is presented. After that, the cross-layer QPF scheduling with packet-level QoS improvement is described and explained in Section IV. Then, we present the simulation results in Section V. Finally, we give some concluding remarks in Section VI.

II. CROSS-LAYER MODEL

The cross-layer scheduling architecture is depicted in Fig. 1. The OFDM frames are transmitted through the physical layers between the base station (BS) and mobile user (MS). The data link layer maintains the queue buffers that contain the packets delivered from the higher application layer. The packet scheduler collects the cross-layer information (CSI and QSI) and selects different users' packets for transmission according to the information obtained. Based on this architecture, our proposed cross-layer model is depicted in Fig. 2. We assume a total of K MSs in the OFDM system with M subcarriers, and the total

bandwidth is W . In the BS, the incoming packets of each user arrive from some higher layers and are then buffered in its own first-in–first-out (FIFO) queue with a finite space of B packets waiting to be scheduled. Similar to that in [8] and [14], we assume that the packet length is fixed to be N bits/packet. At the beginning of each time slot, the packet scheduler selects some packets in the queues for transmission according to CSI and QSI so as to meet the QoS requirements. The selected packets are then forwarded to the OFDM transmitter. Then, they are adaptively modulated at the corresponding mode related to the CSI and distributed on different subcarriers. After inverse fast Fourier transformation (IFFT) and guard interval (GI) insertion, the scheduled packets from the different users form an OFDM symbol and are then sent to the MSs via the downlink channels. In our model, we group S OFDM symbols into one OFDM frame, and one frame transmission time is assumed to be T_0 seconds, which can be expressed as

$$\begin{aligned} T_0 &= S \cdot (\text{length of one OFDM symbol}) \\ &= S \cdot \frac{M}{W} \cdot (1 + 0.25) \end{aligned} \quad (1)$$

where $S \cdot (M/W)$ is the length of data, and $0.25 \cdot S \cdot (M/W)$ is the length of GI. Thus, each OFDM frame is time slotted, and T_0 is the length of one time slot. In this paper, T_0 also corresponds to one scheduling time interval.

At the physical layer, we consider an L -path Rayleigh fading channel. We assume that the perfect CSI is sent to the BS through feedback channels from all the MSs. By using adaptive modulation, at time slot t , the achievable transmission rate (in bits per symbol) for the k th user on the m th subcarrier $c_{k,m}(t)$ can usually be decided by the current channel SNR and the required bit error rate (BER). For instance, if an adaptive uncoded multilevel quadrature amplitude modulation (MQAM) and an ideal phase detection are used as in [8] and [15], it can be expressed as

$$c_{k,m}(t) = \log_2 \left(1 + \frac{-1.5}{\ln(5 \cdot P_{\text{ber}})} \cdot \gamma_{k,m}(t) \right) \quad (2)$$

where $\gamma_{k,m}(t)$ is the SNR for the k th user's m th subcarrier signal at time instant t , and P_{ber} is the required BER for this transmission. Let Δf be the frequency spacing between two adjacent subcarriers, and let $x_{k,m}(t)$ be the channel assignment status that takes the value of 1 or 0 when the m th subcarrier is or is not occupied by the k th user. Then, the k th user's instantaneous data rate is defined as

$$r_k(t) = \sum_{m=1}^M r_{k,m}(t) \cdot x_{k,m}(t) \quad (3)$$

where $r_{k,m}(t) = \lfloor c_{k,m}(t) \rfloor \cdot \Delta f$ b/s is the k th user's data rate on the m th subcarrier.

Furthermore, to analyze the QoS performances for wireless multimedia services, we assume that the k th user belongs to a QoS profile with the parameters $\{\mu_k, d_k\}$, where μ_k and d_k are the maximum allowable long-term average PDP and the packet delay, respectively.

III. PF PACKET SCHEDULING

In this section, we discuss the PF scheduling for OFDM systems and present an efficient solution.

A. Problem Formulation for PF Scheduling in OFDM Systems

As discussed before, the PF criterion for an SC system is not suitable if we want to transmit multiple users in OFDM systems with multiple carriers. Then, based on the PF definition for MC systems in [7], we can formulate the PF scheduling in the OFDM systems as

$$P = \arg \max_{C_k, k \in U} \prod_{k \in U} \left(1 + \frac{\sum_{m \in C_k} r_{k,m}(t)}{(t_c - 1) \bar{R}_k(t)} \right) \quad (4)$$

subject to

$$(I) \bigcup_{k \in U} C_k \subseteq C \quad (II) C_i \cap C_j = \emptyset, \quad i \neq j \quad \forall i, j \in U$$

where $C = \{1, 2, \dots, M\}$ and $U = \{1, 2, \dots, K\}$ denote the subcarrier index set and the user index set, respectively. C_k is the set of subcarriers that are allocated to the k th user. $\bar{R}_k(t)$ is the average data rate at time slot t of user i . Here, the constraint I means that each user selects its subcarriers from C , and the constraint II shows that each subcarrier can only be assigned to one user. $\bar{R}_k(t)$ can be approximated by a moving average value with average window size t_c time slots

$$\bar{R}_k(t+1) = \left(1 - \frac{1}{t_c} \right) \cdot \bar{R}_k(t) + \frac{1}{t_c} \cdot r_k(t). \quad (5)$$

When introducing the channel assignment status $x_{k,m}(t)$, the preceding MC PF problem can be reformulated by

$$\max_x \prod_{k=1}^K \left(1 + \frac{\sum_{m=1}^M x_{k,m}(t) \cdot r_{k,m}(t)}{(t_c - 1) \bar{R}_k(t)} \right) \quad (6)$$

subject to

$$(I) \sum_{k=1}^K x_{k,m}(t) = 1, \quad x_{k,m} = \{0, 1\}$$

where $x = [x_{1,1}, \dots, x_{1,M}, \dots, x_{K,1}, \dots, x_{K,M}]^T$. The optimization in (6) is a nonlinear integer programming problem, and the algorithm complexity for the optimal solution is prohibitively high. Specifically, for a system with K users and M subcarriers, in the worst case, a total number of M^K times of iterations are needed to find the solution. Such an algorithm will consume much time and memory space, which is not suitable for practical systems. To make the scheduling method amenable for practical implementation, a simplification for the foregoing generic MC PF scheduler is needed. In this paper, we present a fast and efficient greedy method to solve this problem.

B. MC PF Scheduling Algorithm in OFDM Systems

The objective function in (6) is a multiple-product function. To simplify the algorithm implementation, we first convert the

product into a summation by taking a logarithm function on the objective. Then, the equivalent problem can be described as

$$\max_x \sum_{k=1}^K \log \left(1 + \frac{\sum_{m=1}^M x_{k,m}(t) \cdot r_{k,m}(t)}{(t_c - 1)R_k(t)} \right) \quad (7)$$

subject to

$$(I) \sum_{k=1}^K x_{k,m}(t) = 1, \quad x_{k,m} = \{0, 1\}.$$

To be specific, we define $PF(t)$ as the system PF value at time slot t , i.e.,

$$\begin{aligned} PF(t) &= \sum_{k=1}^K \log \left(1 + \frac{\sum_{m=1}^M x_{k,m}(t) \cdot r_{k,m}(t)}{(t_c - 1)R_k(t)} \right) \\ &= \sum_{k=1}^K \log \left(1 + \frac{r_k(t)}{(t_c - 1)R_k(t)} \right). \end{aligned} \quad (8)$$

Notice that if one subcarrier is assigned to a different user, then the resulting PF value will be different. For example, if before the assignment of the m th subcarrier the user rate for every user $k \in U$ is $r_k(t)$, then when this subcarrier is allocated to the k' th user, the new PF value is then given by

$$\begin{aligned} PF(t, k') &= \log \left(1 + \frac{r_{k'}(t) + r_{k',m}(t)}{(t_c - 1)R_{k'}(t)} \right) \\ &\quad + \sum_{k \neq k'} \log \left(1 + \frac{r_k(t)}{(t_c - 1)R_k(t)} \right). \end{aligned} \quad (9)$$

Thus, for the m th subcarrier, if the largest system PF value $PF(t, k)$ is obtained when the k th user gets it, then it should be assigned to this k th user. With this strategy, the subcarrier can be allocated so as to get the highest PF value. Consequently, we can assign all of the subcarriers one by one in this greedy manner, and the resulting PF value will be the maximal value when all the subcarriers are allocated. Furthermore, because the scheduler needs to compare K users to allocate a subcarrier, the total computational complexity is KM , which is efficient compared to the number of comparisons M^K for the original problem in (7). Thus, the greedy MC PF scheduling method can be formalized as follows in Algorithm 1.

Algorithm 1 The Greedy MC PF scheduling algorithm

- 1: Initialization: Let $r_k(t) = 0$ and $C_k = \emptyset$ for all $k \in U$;
- 2: **for** subcarrier $m = 1$ to M **do**
- 3: Calculate $r_{k,m}(t)$ for all $k \in U$;
- 4: **for** every user $k \in U$ **do**
- 5: Calculate the PF value $PF(t, k)$ if user k occupies subcarrier m according to (8);
- 6: **end for**
- 7: The user getting the largest $PF(t, k)$ is assigned with this subcarrier, i.e., $\tilde{k} = \arg \max_k (PF(t, k))$;
- 8: Update $r_{\tilde{k}}(t) = r_{\tilde{k}}(t) + r_{\tilde{k},m}(t)$;

- 9: Update the set of subcarriers allocated to this user, $C_{\tilde{k}} = C_{\tilde{k}} + \{m\}$;
 - 10: **end for**
 - 11: Update $\bar{R}_k(t+1)$ for all $k \in U$;
 - 12: Transmit each user's packets on the assigned subcarriers with the corresponding rate.
-

IV. PF SCHEDULING WITH PACKET-LEVEL QoS IMPROVEMENT

As discussed in Section I, the PF packet scheduler previously presented does not consider the queue state, coupled with the assumption of infinite incoming packets. When considering the queue state and bursty traffic, the scheduler should not serve any empty queue because doing that will waste radio resource. In addition, if a packet delay exceeds its delay limit, then this packet should be dropped from its queue. Furthermore, the packet-level QoS performance metrics, such as packet delay and PDP, should also be considered. In this section, we analyze the QoS performance and propose a QPF scheduling algorithm in the OFDM systems.

A. Packet-Level QoS Performance Analysis

To analyze the packet-level QoS performance, we first model the queuing service. Assuming that at the beginning of time slot t the queue length is $Q_k(t)$, the scheduler serves the k th user at rate $r_k(t)$ according to CSI and QSI. Furthermore, there are $\pi_k(t)$ packets to be dropped because their packet delays exceed the limit. Then, if there are $\nu_k(t)$ packets arriving during this time slot and regardless of the queue limit, the queue length at the end of this time slot can be expressed as

$$U_k(t+1) = Q_k(t) - \frac{r_k(t) \cdot T_0}{N} + \nu_k(t) - \pi_k(t) \quad (10)$$

Because the queue has a capacity limit of B packets, the actual queue length $Q_k(t+1)$ at the beginning of time slot $t+1$ should be verified by

$$\begin{aligned} Q_k(t+1) &= \min \{B, U_k(t+1)\} \\ &= \min \left\{ B, Q_k(t) - \frac{r_k(t) \cdot T_0}{N} + \nu_k(t) - \pi_k(t) \right\}. \end{aligned} \quad (11)$$

The number of dropped packets due to the overflow at the end of time slot t can then be evaluated as

$$\begin{aligned} D_k(t+1) &= \max \{0, U_k(t+1) - B\} \\ &= \max \left\{ 0, Q_k(t) - \frac{r_k(t) \cdot T_0}{N} + \nu_k(t) - \pi_k(t) - B \right\}. \end{aligned} \quad (12)$$

Furthermore, to avoid serving the empty queues, the scheduler should control the service rate so that

$$r_k(t) \leq \frac{Q_k(t) \cdot N}{T_0}. \quad (13)$$

Then, the average packet delay and the average PDP can be analyzed based on this model. Similar to the average data rate definition in (5) and regardless of the queue limit, let the average queue length $\overline{U_k(t)}$ over the average window size t_c be

$$\overline{U_k(t+1)} = \left(1 - \frac{1}{t_c}\right) \cdot \overline{U_k(t)} + \frac{1}{t_c} \cdot U_k(t+1). \quad (14)$$

Subsequently, at the beginning of time slot t , when given $\overline{U_k(t)}$, the actual queue length $Q_k(t)$, the dropped packets $\pi_k(t)$ due to deadline missed, and the average packet incoming rate $E\{\nu_k(t)\}$, the predicted average queue length (regardless of the queue limit) over the average window size at the end of time slot t can be expressed as

$$\begin{aligned} \hat{U}_k(t+1) &= E_{\nu_k(t)} \left\{ \overline{U_k(t+1)} \right\} \\ &= E_{\nu_k(t)} \left\{ \left(1 - \frac{1}{t_c}\right) \overline{U_k(t)} + \frac{1}{t_c} U_k(t+1) \right\} \\ &= \left(1 - \frac{1}{t_c}\right) \overline{U_k(t)} \\ &\quad + \frac{E_{\nu_k(t)} \left\{ Q_k(t) - \frac{r_k(t)T_0}{N} + \nu_k(t) - \pi_k(t) \right\}}{t_c} \\ &= \left(1 - \frac{1}{t_c}\right) \overline{U_k(t)} \\ &\quad + \frac{Q_k(t) - \frac{r_k(t)T_0}{N} + E\{\nu_k(t)\} - \pi_k(t)}{t_c} \\ &= H(r_k(t)) \end{aligned} \quad (15)$$

where $E_{\nu_k(t)}\{\cdot\}$ is the expectation for expression $\{\cdot\}$ with respect to $\nu_k(t)$, which can be obtained according to the incoming traffic characteristics.

Now, if we schedule the k th user with data rate $r_k(t)$, we can predict the average number of dropped packets at the end of time slot t due to overflow and express it as

$$\begin{aligned} \hat{D}_k(t+1) &= \max \left\{ 0, \hat{U}_k(t+1) - B \right\} \\ &= \max \left\{ 0, H(r_k(t)) - B \right\} = F(r_k(t)). \end{aligned} \quad (16)$$

Letting $\mu_k(t+1)$ be the estimate of the k th user's average PDP at the end of time slot t , then

$$\mu_k(t+1) = \frac{\hat{D}_k(t+1) + \pi_k(t)}{E\{\nu_k(t)\}}. \quad (17)$$

Thus, at the beginning of time slot t , we can get the required data rate $r_k(t)$ if we want to meet the PDP requirement μ_k when given $\hat{U}_k(t)$, $\pi_k(t)$, $Q_k(t)$, and queue size B . Furthermore, we have

$$\mu_k(t+1) = \frac{\hat{D}_k(t+1) + \pi_k(t)}{E\{\nu_k(t)\}} \leq \mu_k. \quad (18)$$

Then, we have

$$F(r_k(t)) \leq E\{\nu_k(t)\} \cdot \mu_k - \pi_k(t). \quad (19)$$

Because $F(r_k(t)) - \pi_k(t)$ is a nonincreasing function related to $r_k(t)$, we can express the required rate by

$$r_k(t) \geq F^{-1}(E\{\nu_k(t)\} \cdot \mu_k - \pi_k(t)) = \alpha_k(t). \quad (20)$$

Then, the estimate of the average packet delay for the k th user can be expressed by Little's law [16] as

$$\hat{d}_k(t+1) = \frac{\hat{Q}_k(t+1)}{E\{\nu_k(t)\}} \quad (21)$$

where $\hat{Q}_k(t+1)$ is the estimated actual queue length at the end of time slot $t+1$, and

$$\begin{aligned} \hat{Q}_k(t+1) &= \min \left\{ B, \hat{U}_k(t+1) \right\} \\ &= \min \left\{ B, H(r_k(t)) \right\} = G(r_k(t)). \end{aligned} \quad (22)$$

Because $G(r_k(t))$ is also a nonincreasing function related to $r_k(t)$, to meet the packet delay requirement d_k , we have

$$\hat{d}_k(t+1) = \frac{\hat{Q}_k(t+1)}{E\{\nu_k(t)\}} = \frac{G(r_k(t))}{E\{\nu_k(t)\}} \leq d_k \quad (23)$$

$$r_k(t) \geq G^{-1}(E\{\nu_k(t)\} \cdot d_k) = \beta_k(t). \quad (24)$$

Thus, when combining (20) and (24), if we have

$$r_k(t) \geq \max \left\{ \alpha_k(t), \beta_k(t) \right\} \quad (25)$$

then the average PDP and the average packet delay performance could be linked together in the system.

B. Formulation

With the foregoing packet-level QoS analysis, we can reformulate the MC PF scheduling problem in Section III as an MC QPF scheduling problem

$$\max_x \sum_{k=1}^K \log \left(1 + \frac{\sum_{m=1}^M x_{k,m}(t) \cdot r_{k,m}(t)}{(t_c - 1)R_k(t)} \right) \quad (26)$$

subject to

$$(I) \quad \sum_{k=1}^K x_{k,m}(t) = 1, \quad x_{k,m} = \{0, 1\}$$

$$(II) \quad r_k(t) \leq \frac{Q_k(t) \cdot N}{T_0}$$

$$(III) \quad r_k(t) \geq \max \left\{ \alpha_k(t), \beta_k(t) \right\}$$

where constraint *II* implies that the scheduler does not serve the empty queue, and constraint *III* is used for improving the average packet delay and the average PDP performance metrics, as previously analyzed.

C. Cross-Layer Design for QPF Scheduling in OFDM Systems

To solve the QPF scheduling problem, we present a two-step cross-layer QPF algorithm. First, we modify the greedy

PF scheduling algorithm described in Section III to ensure that constraint *II* in (26) is satisfied. Then, the subcarrier reassignment procedure is implemented to improve the QoS performance according to constraint *III*.

Our modified greedy PF scheduling algorithm is described as follows. When the m th subcarrier is allocated to the k th user with data rate $r_{k,m}(t)$ to get the largest PF value, this user's queue length is updated by $Q_k(t+1) = Q_k(t) - ((r_{k,m}(t) \cdot T_0)/N)$. If $Q_k(t) \leq 0$. Then, the queue is empty, and consequently, in the following scheduling steps, this user will not be served so as not to waste resource. This allocation process is formalized in Algorithm 2.

Algorithm 2 The modified greedy MC PF scheduling algorithm

- 1: Initialization: Let $r_k(t) = 0$ and $C_k = \emptyset$ for all $k \in U$;
 - 2: **for** subcarrier $m = 1$ to M **do**
 - 3: Calculate $r_{k,m}(t)$ for all $k \in U$ according to (2);
 - 4: **for** every user $k \in U$ **do**
 - 5: Calculate the PF value $PF(t, k)$ if user k occupies subcarrier m according to (8);
 - 6: **end for**
 - 7: The user getting the largest $PF(t, k)$ is assigned with this subcarrier, i.e., $\tilde{k} = \arg \max_k (PF(t, k))$;
 - 8: Update $r_{\tilde{k}}(t) = r_{\tilde{k}}(t) + r_{\tilde{k},m}(t)$;
 - 9: Update $C_{\tilde{k}} = C_{\tilde{k}} + \{m\}$;
 - 10: Update this user's queue length by $Q_{\tilde{k}}(t) = Q_{\tilde{k}}(t) - ((r_{\tilde{k},m}(t) \cdot T_0)/N)$;
 - 11: **if** $Q_{\tilde{k}}(t) \leq 0$ **then**
 - 12: Remove this user from the user set, and $U = U - \{\tilde{k}\}$;
 - 13: **end if**
 - 14: **end for**
 - 15: Update $R_k(t+1)$ for all users;
 - 16: Transmit each user's packets on the assigned subcarriers with the corresponding rate.
-

In the subcarrier reassignment process, we assume that we get the set of subcarriers allocated to the k th user C_k , i.e., the corresponding $r_k(t)$ and $r_{k,m}(t)$, after performing the previously modified greedy PF scheduling algorithm at time slot t . Let Ψ be the user set in which each user's packet-level QoS performance constraint *III* is satisfied, that is

$$\Psi = \{k : r_k(t) \geq \max \{\alpha_k(t), \beta_k(t)\}\} \subseteq U \quad (27)$$

and let $\Psi^c = U - \Psi$ be the user set where each user's QoS performance requirement is not satisfied.

If $\Psi = U$, and each user's QoS requirement is satisfied, then the allocation is already optimal; thus, the packets can be transmitted at the allocated subcarriers with an assigned data rate. Otherwise, the subcarrier reassignment procedure is needed. Define $S_\Psi = \{m : m \in S_k, k \in \Psi\}$ as the subcarriers allocated to the users whose rate limits are satisfied. In this case, some subcarriers in S_Ψ should successively be reassigned to the users in Ψ^c . To control the subcarrier reassignment while maintaining good PF performance, we define a new

swap PF value $PF(i, j, m)$ when the m th subcarrier previously belonging to the i th user ($i \in \Psi$) is assigned to the j th user ($j \in \Psi^c$), i.e.,

$$\begin{aligned} PF(i, j, m) = & \log \left(1 + \frac{r_i(t) - r_{i,m}(t)}{(t_c - 1)R_i(t)} \right) \\ & + \log \left(1 + \frac{r_j(t) + r_{j,m}(t)}{(t_c - 1)R_j(t)} \right) \\ & + \sum_{k \neq i, j} \log \left(1 + \frac{r_k(t)}{(t_c - 1)R_k(t)} \right). \end{aligned} \quad (28)$$

If the i th user can still satisfy the QoS constraint *III*, and the corresponding $PF(i, j, m)$ is the largest among all of the reassignment attempts, then we can reassign this m th subcarrier to the j th user. Then the reallocation process continues until its rate requirement is satisfied or there are not enough subcarriers that could be assigned to this user if the users in the system are experiencing deep fading. Specifically, the subcarrier reassignment algorithm is formulated in Algorithm 3.

Algorithm 3 Subcarrier reassignment algorithm

- 1: **for** every subcarrier $m \in S_\Psi$ **do**
 - 2: **if** m has been allocated to user $i \in \Psi$ and $r_i(m) - r_{i,m}(t) \geq \max \{\alpha_i, \beta_i\}$ **then**
 - 3: Calculate swap PF value for user i and every $j \in \Psi^c$ according to (27);
 - 4: Assign subcarrier m with the maximal $PF(i, j, m)$ to the user j , i.e., $j = \arg \max_{j \in \Psi^c} PF(i, j, m)$;
 - 5: Update $r_i(t) = r_i(t) - r_{i,m}(t)$, $r_j(t) = r_j(t) + r_{j,m}(t)$;
 - 6: Update $S_\Psi = S_\Psi - \{m\}$;
 - 7: Update $C_j = C_j + \{m\}$, $C_i = C_i - \{m\}$;
 - 8: **if** $r_j \geq \max \{\alpha_j, \beta_j\}$ **then**
 - 9: $\Psi^c = \Psi^c - \{j\}$;
 - 10: **end if**
 - 11: **end if**
 - 12: **end for**
-

D. Qualitative Comparison

Based on the preceding analysis, we can give a general qualitative comparison of our QPF method with several related approaches, such as an SC PF scheduler extended for OFDM systems [6], the original MC PF scheduler [7], the greedy MC PF scheduler proposed in Section III, a unified method [9], a utility-based MC scheduler [12], and a MR scheduler [17], which chooses the user with the largest SNR to occupy the subcarrier. As shown in Table I, the QPF method is the only approach to take efficiency, fairness, and QoS into consideration together. Here, the proposed QPF algorithm consists of two parts, i.e., the modified greedy MC PF algorithm (Algorithm 2) and the subcarrier reassignment algorithm (Algorithm 3). For a system with K users and M subcarriers, the complexity of Algorithm 2 is KM . The worst-case computational complexity

TABLE I
QUALITATIVE COMPARISON OF THE SEVEN SCHEDULING ALGORITHMS

	Efficiency	Fairness	QoS	Complexity
QPF	Y	Y	Y	$2KM$
Greedy MC-PF	Y	Y	N	KM
MC-PF	Y	Y	N	M^K
Utility MC	Y	Y	N	$(K-1)^2(M+1)\log_2(M)$
SC-PF	Y	Y	N	KM
MR	Y	N	N	KM
Unified	Y	N	Y	KM

TABLE II
SIMULATION PARAMETERS FOR THE OFDM SYSTEM

Parameter	Value
System bandwidth	20 MHz
Number of subcarriers	256
OFDM symbol time	$16\mu s$
Symbols per OFDM frame	100
OFDM frame time(T_0)	$1600\mu s$
BER requirement	10^{-5}
Number of multi-path	6
Max Doppler frequency	30
Power delay profile	e^{-2l}
AWGN power density	-80dBW/Hz
Packet length	100bit/packet
Queue size	200
Average window size	$10 \cdot T_0$

for Algorithm 3 is also KM . Thus, the total computational complexity of the QPF is at most $2KM$, which is still at the same level with the MR and SC PF methods. On the other hand, the computational complexities for the original MC PF scheduler and the utility MC scheduler are M^K and $(K-1)^2(M+1)\log_2 M$, respectively, which are relatively higher than the QPF method. Thus, the complexity performance for the QPF is also acceptable even when adding the additional QoS improvement.

V. SIMULATION RESULTS

In this section, we present the performance results of the proposed QPF scheduler.

A. Simulation Configuration

Similar to that in [8] and [14], the simulation parameters for the OFDM system are presented in Table II. The entire system bandwidth is 20 MHz, which is divided into 256 subcarriers. Thus, the OFDM symbol time is $16\mu s$, in which $3.2\mu s$ is the GI. We group 100 OFDM symbols into an OFDM frame, and thus, the OFDM frame time is $1600\mu s$, i.e., $T_0 = 1600\mu s$.

The wireless channel is modeled as a six-path frequency-selective Rayleigh fading channel. Each path is simulated by Clark's fading model and suffers from different Rayleigh fading with the maximum Doppler frequency of 30 Hz, which corresponds to the average speed of 4.5 km/h for the carrier frequency of 5 GHz. We also assume that the power delay profile is exponentially decaying with e^{-2l} , where l is the multipath index. Moreover, the additive white Gaussian noise (AWGN) power density is -80 dBW/Hz. Furthermore, the channel that corresponds to the different users has the same but independent statistics. We assume that the queue space is 200 packets,

TABLE III
SIMULATION PARAMETERS FOR TRAFFIC

Parameter	Value
Mean ON period for voice traffic	1.00s
Mean OFF period for voice traffic	1.35s
Constant data rate for voice traffic	64kbps
Delay limit for voice traffic	20ms
Packet dropping probability for voice traffic	1%
Delay limit for video traffic	75ms
Packet dropping probability for video traffic	1%
Delay limit for data traffic	0.5s
Packet dropping probability for data traffic	5%

and that the packet length is fixed to 100 bits/packet. The maximum BER requirement is 10^{-5} .

Three classes of packetized traffic are considered, i.e., video, voice, and data, and the corresponding traffic parameters are presented in Table III. The video traffic is modeled by an eight-state Markov-modulated Poisson process (MMPP) [14]. The average duration in each state is 40 ms, which is equivalent to the length of one video frame with a frame rate of 25 frames/s. The voice traffic is generated according to a two-state ON-OFF model [14], [18], where the mean ON period is 1 s, and the mean OFF period is 1.35 s. In the ON state, the voice packet arrives at a constant rate of 64 kb/s, whereas in the OFF state, no packet is generated. The data packet is just modeled as a Poisson process. As in [18], for the video traffic, the delay limit is 75 ms, and the PDP is 1%. For the voice traffic, the delay limit is 50 ms, and the PDP is 1%. For the data traffic, it usually corresponds to the best-effort application. Here, we borrow the QoS requirement for the WWW model in [9] as its QoS profile, i.e., the delay limit is 0.5 s, and the PDP is 5%.

Using this simulation model, we compare the performance of the QPF algorithm with respect to four algorithms with similar scheduling complexities, namely, the MR scheduler [17], the unified scheduler [9], the SC PF scheduler [6], and the greedy MC PF scheduler without QoS constraint proposed in Section III. All of the simulations were run in Matlab 7.3 on Windows XP with a Pentium-4 2.8-GHz CPU and 512-MB RAM. Each test case was repeated for 100 trials. From the simulations, we have calculated the sample mean and the 95% confidence interval. To be noticed, for most simulation points, the size of the confidence interval has been very small in comparison with the sample mean. When these confidence intervals are included in the figures, they are almost unnoticeable. Furthermore, plotting the confidence intervals deteriorates the readability of the figures. Thus, we only plot the confidence interval in Fig. 3.

B. Performance Results for Video Traffic

First, we present the performance results for the video traffic. Fig. 3 shows the average system throughput comparisons of the five algorithms. We define the average system throughput as the average transmitted bit per second in the system. As can be seen, the proposed QPF algorithm achieves the highest throughput among all the algorithms. The MR and the unified method also achieve a high throughput performance because their scheduling policies are maximizing the system

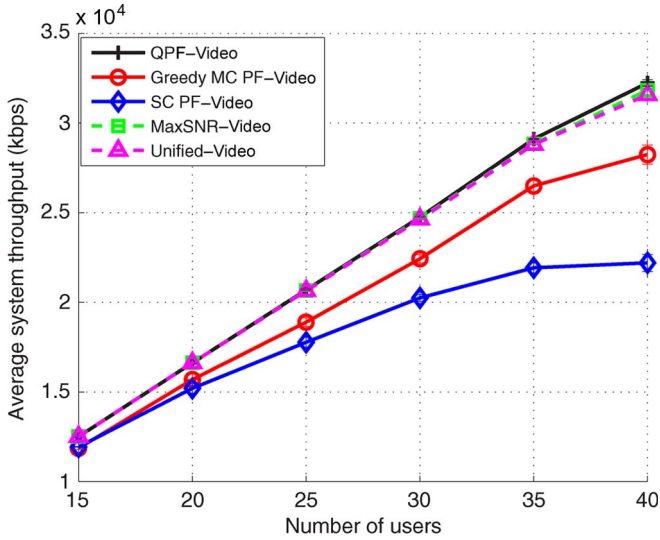


Fig. 3. Average system throughput versus number of users with video traffic.

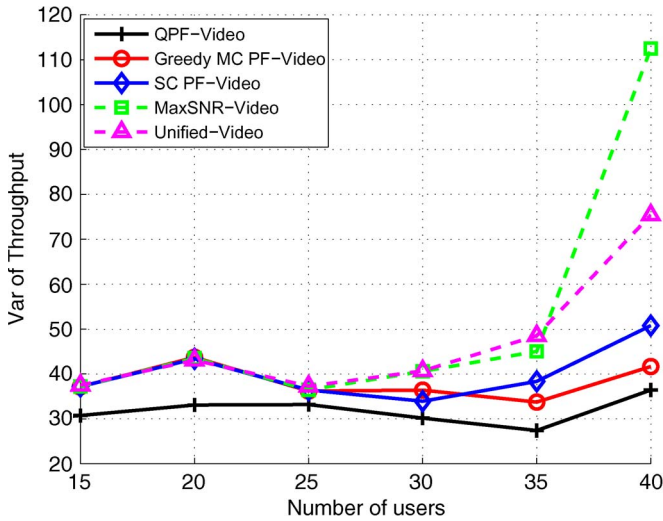


Fig. 4. Variance of different users' throughput versus number of users with video traffic.

throughput. However, their performance results are poorer than that for the QPF algorithm, particularly in the case of a higher packet arrival rate. For example, when there are 40 users, the average system throughput for QPF is 32.1 Mb/s, with a 95% confidence interval of (31.8 M, 32.4 M). Its throughput is 5% higher than that of the MR and the unified algorithms, because when the users increase, the dropped and delayed packets also increase if there is no QoS control. Because our approach controls the packet-level QoS, the numbers of dropped packets and delayed packets are smaller than the other algorithms. Similarly, the performance for the QPF is also better than that of the greedy MC PF and SC PF algorithms. It should be noticed that the greedy MC PF scheduler achieves a higher throughput than that of the SC PF algorithm because the former uses the optimal PF criterion for packet transmission in MC systems.

The fairness performance is indicated by the variance of the different users' average throughput, as shown in Fig. 4. It can be observed that the variances for the MR and unified methods are both much larger than the other algorithms. This means that the

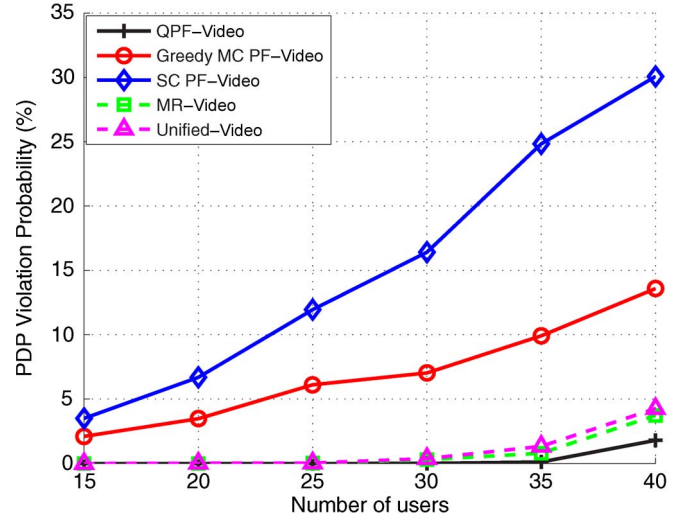


Fig. 5. PDP violation probability versus number of users with video traffic.

throughput differences among the users are much higher when these two methods are used because they do not perform fairness control in the algorithm design. When observing other PF-based algorithms, their fairness performance results are much better due to their capabilities to balance efficiency and fairness for resource scheduling. Furthermore, when the traffic density is low, e.g., $K \leq 30$, the throughput variances among the users for all five algorithms are similar to each other, because under these conditions, an individual user may obtain enough resources to satisfy its throughput and QoS requirements. Thus, fairness can relatively easily be achieved. When $K > 30$, the variances significantly increase with the user number due to the resource competition among users. However, the QPF can still maintain the best fairness performance. This can be explained by the fact that during the process of QoS control that was introduced by the QPF, some transmission chances are switched from the QoS-satisfied users to some other unsatisfied users. Thus, the packets that would likely be dropped due to the deadlines missed or overflow when using other algorithms without QoS consideration are also successfully selected for transmission by the QPF. Then, when using the QPF algorithm, the radio resource is allocated in a fairer manner among users. Whereas the throughput variance for the MR algorithm significantly increases when there are 40 users. This is because the MR just select the users with good channel conditions; thus, some users with bad channel conditions will not get services when there are a large number of users. Consequently, their packets will be dropped if the queue space limit or delay limit is reached. Then, the fairness performance deteriorates.

The PDP performance is shown in Fig. 5 in terms of the PDP violation probability. Similar to the definition used in [19], a PDP violation occurs when the calculated average PDP at the end of one scheduling slot for a particular user is not satisfied with the predefined requirement. We then define the PDP violation probability as the ratio of the number of occurred PDP violations over all the scheduling time slots for all the users to the total number of calculated PDP over these time slots for these users. We can see that our proposed algorithm considerably improves the packet-dropping performance, particularly

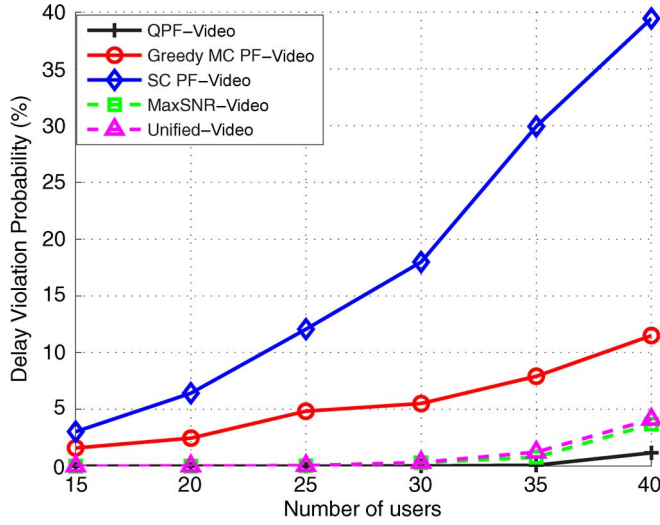


Fig. 6. Packet delay violation probability versus number of users with video traffic.

when the traffic density is high. We also find that the SC PF method has the worst performance. For example, when there are 40 users, its PDP violation is almost 30%. This result confirms our analysis that the PF scheduling algorithm designed for an SC system cannot directly be used for the multimedia traffic in an OFDM system due to its lower throughput performance as well as its inability to adapt to the queue state. Similarly, the average delay performance is shown in Fig. 6 in terms of the packet delay violation probability in the system, which is defined as the ratio of the number of occurred delay violations over all the time slots for all the users to the total number of calculated packet delay over these time slots for these users. Here, the packet delay violation probability of the proposed QPF method is the smallest among all the algorithms because it can control the packet-level QoS and schedule the packets according to the CSI as well as QSI.

C. Performance Results for Heterogeneous Traffic

To demonstrate the performance of the proposed scheduling scheme in a heterogeneous traffic environment, we simulate a system in which the voice, video, and data traffic flows are transmitted at the same time, where the video traffic occupies 40% of the total traffic, and the voice and data traffic each occupies 30%. In Fig. 7, the average system throughputs of all these algorithms increase with the increasing number of users due to the utilization of multiuser diversity. More importantly, the proposed QPF algorithm has the best performance. The fairness performance is shown in Fig. 8. Because the packet incoming rates for the different types of traffic are different, the variance of throughput in terms of heterogeneous traffic is larger than that in terms of homogeneous video traffic. Thus, the fairness performance for the QPF decreases when using the heterogeneous traffic. However, the proposed QPF algorithm still outperforms the other algorithms. Furthermore, the PDP violation probability in Fig. 9 and the delay violation probability in Fig. 10 for QPF are also smaller than that of the other algorithms. Thus, the QPF algorithm achieves the best

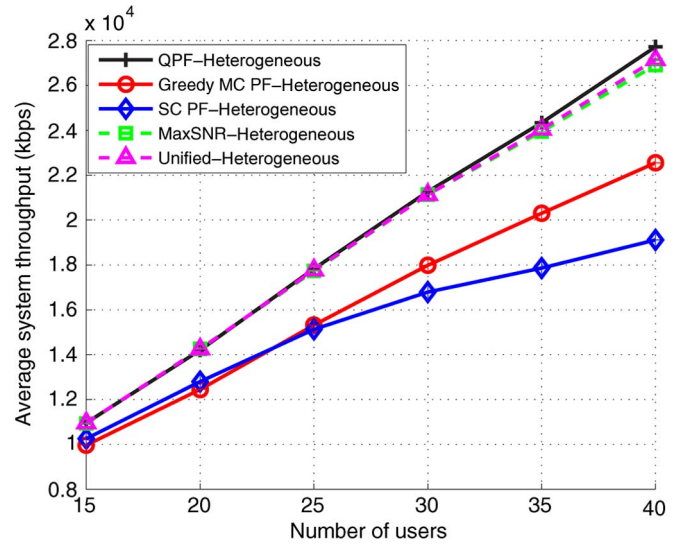


Fig. 7. Average system throughput versus number of users with heterogeneous traffic.

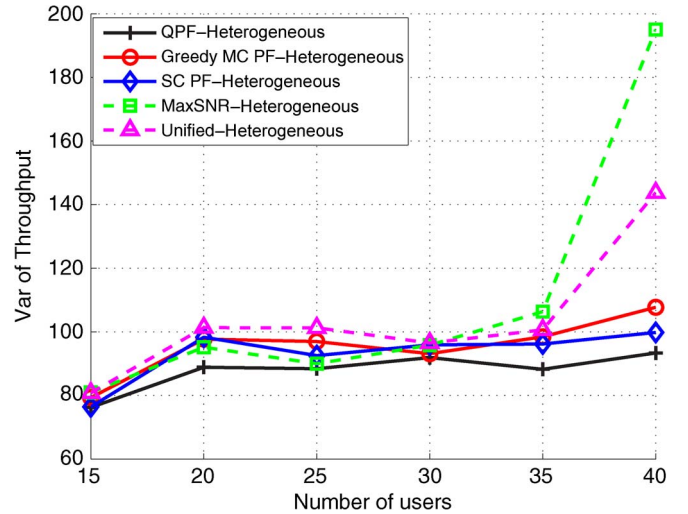


Fig. 8. Variance of different users' throughput versus number of users with heterogeneous traffic.

performance in terms of average system throughput, packet delay, and PDP with a relatively low computational complexity.

D. Impact of Average Window Size t_c

As described in [5], for the SC PF algorithm, t_c is related to the maximum time for which a user can be starved during poor channel conditions. Thus, the scheduler usually chooses a large t_c , e.g., $1000T_0$, to wait long for a user's channel condition to improve so as to consequently increase the total throughput. However, in MC systems, this requirement can be relaxed due to the high capacity to combat the channel fading for OFDM systems. Furthermore, the QPF algorithm has QoS control and can schedule users to improve their QoS performance. Thus, the impact of the average window size here is not as important as that in an SC system. In Table IV, we show the simulation results for the QPF algorithm with 40 video users in terms of different t_c values. We can see that the performance results with

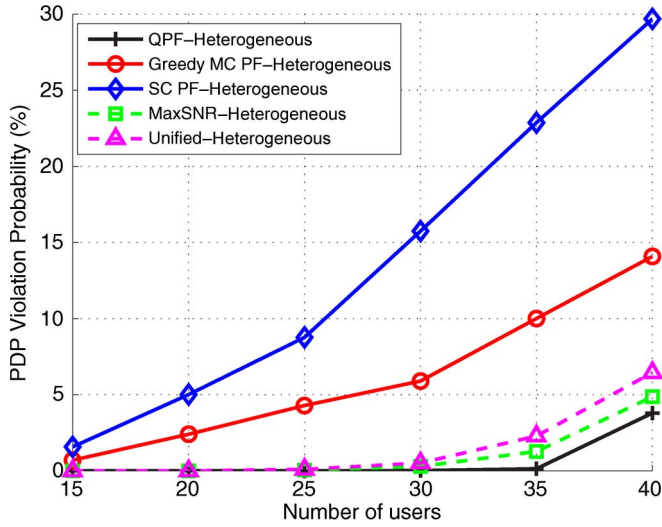


Fig. 9. PDP violation probability versus number of users with heterogeneous traffic.

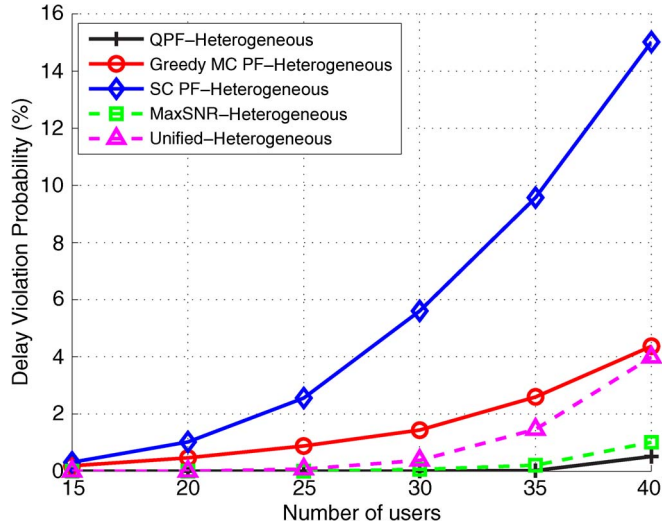


Fig. 10. Packet delay violation probability versus number of users with heterogeneous traffic.

TABLE IV
IMPACT OF t_c ON THE QPF ALGORITHM

t_c	Throughput	Variance of throughput	Delay violation probability	PDP violation probability
$5T_0$	32.67Mbps	33.76	1.36%	2.05%
$5T_0$	32.10Mbps	38.18	1.15%	1.53%
$50T_0$	33.10Mbps	32.44	1.08%	1.21%
$100T_0$	32.95Mbps	38.13	1.38%	1.65%
$1000T_0$	33.13Mbps	40.77	1.28%	1.51%

different window sizes are similar. Thus, we choose a shorter average window size, i.e., $t_c = 10T_0$, in the above simulations.

VI. CONCLUSION

In this paper, a cross-layer packet-scheduling algorithm called QPF has been proposed for the multimedia services in a multiuser OFDM system based on the different users' channel conditions, as well as their queue states. The design objective

of the QPF is to achieve PF in the system while improving the different user's packet-level QoS performances. With the consideration that the traditional SC PF method is not suitable for the OFDM systems, we have proposed an efficient greedy method based on the MC PF criterion with a relatively low computational complexity to allocate subcarriers to different users in order to achieve PF. Then, based on the analysis of the packet-level QoS performance, a subcarrier reassignment procedure has been proposed to improve the QoS performance. Through mathematical analysis and simulation, we find that MC PF scheduling is more suitable than the SC PF method in OFDM systems. Furthermore, we find that the benefits of introducing QoS control into PF scheduling for OFDM transmission are threefold. First, the occurrences of packet overflow and deadlines missed significantly decrease, and thus, the delay and PDP performances are improved. Second, with the decrease of failed packets, the throughput performance also correspondingly increases. Third, the radio resource can be allocated to users in a fairer manner when more users become satisfied with their QoS performances. Therefore, the cross-layer QPF packet-scheduling method with low complexity can be used in the OFDM system to achieve good performances in terms of throughput, average PDP, packet delay, and fairness.

REFERENCES

- [1] Z. Shen, J. G. Andrews, and B. L. Evans, "Adaptive resource allocation in multiuser OFDM systems with proportional fairness," *IEEE Trans. Wireless Commun.*, vol. 4, no. 6, pp. 2726–2737, Nov. 2005.
- [2] I. F. Akyildiz, D. A. Levine, and I. Joe, "A slotted CDMA protocol with BER scheduling for wireless multimedia networks," *IEEE/ACM Trans. Netw.*, vol. 7, no. 2, pp. 146–158, Apr. 1999.
- [3] F. Kelly, "Charging and rate control for elastic traffic," *Eur. Trans. Telecommun.*, vol. 8, no. 1, pp. 33–37, Jan. 1997.
- [4] V. K. N. Lau, "Proportional fair space time scheduling for wireless communications," *IEEE Trans. Commun.*, vol. 53, no. 4, pp. 1353–1360, Apr. 2005.
- [5] A. Jalali, R. Padovani, and R. Pankai, "Data throughput of CDMA-HDR a high efficiency-high data rate personal communication wireless system," in *Proc. 51st IEEE VTC—Spring*, Tokyo, Japan, Jan. 2001, vol. 3, pp. 1854–1858.
- [6] S. Yoon, Y. Cao, C. B. Chae, and H. Lee, "System level performance of OFDMA forward link with proportional fair scheduling," in *Proc. 15th IEEE Int. Symp. PIMRC*, Barcelona, Spain, Sep. 2004, pp. 1384–1388.
- [7] H. Kim and Y. Han, "A proportional fair scheduling for multicarrier transmission systems," *IEEE Commun. Lett.*, vol. 9, no. 3, pp. 210–212, Mar. 2005.
- [8] G. Song, Y. G. Li, L. J. Cimini, Jr., and H. Zheng, "Joint channel-aware and queue-aware data scheduling in multiple shared wireless channels," in *Proc. IEEE WCNC*, Atlanta, GA, Mar. 2004, pp. 1939–1944.
- [9] Y. Ofuji, S. Abeta, and M. Sawahashi, "Unified fast packet-scheduling method considering fluctuation in frequency domain in forward link for OFCDM broadband packet wireless access," in *Proc. 60th IEEE VTC—Fall*, Los Angeles, CA, Sep. 2004, vol. 4, pp. 2724–2729.
- [10] L. Yang and M.-S. Alouini, "Cross layer analysis of buffered adaptive multicarrier transmission," in *Proc. 60th IEEE VTC—Fall*, Los Angeles, CA, Sep. 2004, vol. 7, pp. 4767–4771.
- [11] M. J. Hossain, D. V. Djonin, and V. K. Bhargava, "Delay limited optimal and suboptimal power and bit loading algorithms for OFDM systems over correlated fading," in *Proc. IEEE GLOBECOM*, St. Louis, MO, Dec. 2005, pp. 2787–2792.
- [12] G. Song and Y. Li, "Cross-layer optimization for OFDM wireless networks—Part II: Algorithm development," *IEEE Trans. Wireless Commun.*, vol. 4, no. 2, pp. 625–634, Mar. 2005.
- [13] J. Huang and Z. Niu, "Buffer-aware and traffic-dependent packet-scheduling in wireless OFDM networks," in *Proc. IEEE WCNC*, Hong Kong, Mar. 2007, pp. 1556–1560.

- [14] J. Cai, X. Shen, and J. W. Mark, "Downlink resource management for packet transmission in OFDM wireless communication systems," *IEEE Trans. Wireless Commun.*, vol. 4, no. 4, pp. 1688–1703, Jul. 2005.
- [15] A. J. Goldsmith and S. G. Chua, "Variable-rate variable-power MQAM for fading channels," *IEEE Trans. Commun.*, vol. 45, no. 10, pp. 1218–1230, Oct. 1997.
- [16] B. D. Bunday, *An Introduction to Queueing Theory*. New York: Halsted, 1996.
- [17] B. S. Tsybakov, "File transmission over wireless fast fading downlink," *IEEE Trans. Inf. Theory*, vol. 48, no. 8, pp. 2323–2337, Aug. 2002.
- [18] H. T. Cheng and W. Zhuang, "Joint power–frequency–time resource allocation in clustered wireless mesh networks," *IEEE Netw.*, vol. 22, no. 1, pp. 45–51, Jan./Feb. 2008.
- [19] H. Kwon, W. Lee, and B. Lee, "Low-overhead resource allocation with load balancing in multi-cell OFDMA systems," in *Proc. IEEE 61st VTC—Spring*, Stockholm, Sweden, May 2005, vol. 5, pp. 3063–3067.



Yu-Kwong Kwok (SM'03) received the B.Sc. degree in computer engineering from the University of Hong Kong (HKU), Hong Kong, in 1991 and the M.Phil. and Ph.D. degrees in computer science from the Hong Kong University of Science and Technology (HKUST), Kowloon, Hong Kong, in 1994 and 1997, respectively.

In August 1998, he was a Visiting Scholar for one year with the Parallel Processing Laboratory, School of Electrical and Computer Engineering, Purdue University, West Lafayette, IN. He recently served as

a Visiting Associate Professor with the Department of Electrical Engineering Systems, University of Southern California, Los Angeles, from August 2004 to July 2005, on sabbatical leave from HKU. He is currently an Associate Professor with the Department of Electrical and Computer Engineering, Colorado State University, Fort Collins, on leave from the Department of Electrical and Electronic Engineering, HKU. His research interests include distributed computing systems, wireless networking, and mobile computing.

Dr. Kwok is a Senior Member of the Association for Computing Machinery (ACM). He is also a Member of the IEEE Computer Society and the IEEE Communications Society. He received the Outstanding Young Researcher Award from HKU in November 2004.



Jiangzhou Wang (M'91–SM'94) received the B.S. and M.S. degrees from Xidian University, Xian, China, in 1983 and 1985, respectively, and the Ph.D. degree (with Greatest Distinction) from the University of Ghent, Ghent, Belgium, in 1990, all in electrical engineering.

He is currently a Professor and Chair of the Department of Electronics, the University of Kent, Kent, U.K. From 1995 to 2005, he was with the University of Hong Kong, where he is still serving as an Honorary Professor. From 1992 to 1995, he

was a Senior System Engineer at Rockwell International Corporation (now Conexant). From 1990 to 1992, he was a Postdoctoral Fellow with the University of California at San Diego. He has held a Visiting Professor position with NTT DoCoMo, Japan. He was a Technical Chairman of the IEEE Workshop on 3G Mobile Communications in December 2000. He has been a technical committee member and session chair of a number of international conferences. He has published over 140 papers, including more than 40 IEEE TRANSACTIONS/JOURNAL papers in the areas of wireless mobile and spread spectrum communications. He has written/edited three books: *High Speed Wireless Communications* (Cambridge Univ. Press, 2008), *Broadband Wireless Communications* (Boston, MA: Kluwer, 2001), and *Advances in 3G Enhanced Technologies for Wireless Communications* (Norwood, MA: Artech House, 2002), respectively. The last book has been translated into Chinese.

Dr. Wang is an Editor for the IEEE TRANSACTIONS ON COMMUNICATIONS and has been a Guest Editor four times for the IEEE JOURNAL ON SELECTED AREAS IN COMMUNICATIONS (Wideband CDMA, August 2000 and January 2001, Advances in Multicarrier CDMA, June 2006, and Wireless Video Transmission, January 2010). He holds one U.S. patent in the GSM system. He is an Evaluation Expert of the European Commission Framework Program 7 (FP7). He is listed in *Who's Who in the World*.



Zhen Kong received the B.Eng. degree in communication engineering and the M.Eng. degree in communication and information systems from Huazhong University of Science and Technology (HUST), Wuhan, China, in 1999 and 2002, respectively, and the Ph.D. degree in electrical and electronics engineering from the University of Hong Kong (HKU), Hong Kong, in 2008.

From 1999 to 2004, he was a Lecturer on R&D of 3G WCDMA core networks with the Department of Electronics and Information Engineering, HUST.

From March 2007 to June 2007, he was a Research Intern with the National Research Institute in Informatics and Control (INRIA), Rennes, France. He is currently a Postdoctoral Fellow with the Department of Electrical and Computer Engineering, Colorado State University, Fort Collins. His research interests include wireless packet scheduling, game-theoretic analysis for selfish behaviors in noncooperative wireless networks, and price competition in wireless environments.