

Throughput Analysis of TCP-Friendly Rate Control in Mobile Hotspots

Sangheon Pack, *Member, IEEE*, Xuemin (Sherman) Shen, *Senior Member, IEEE*,
Jon W. Mark, *Life Fellow, IEEE*, and Lin Cai, *Member, IEEE*

Abstract—By integrating wireless wide area networks (WWANs) and wireless local area networks (WLANs), mobile hotspot technologies enable seamless Internet multimedia services to users on-board a vehicle. In this paper, we investigate the performance of TCP-Friendly Rate Control (TFRC) protocol supporting multimedia services in mobile hotspots. To quantify the throughput of TFRC flows in mobile hotspots, we first develop discrete-time queuing models for the WWAN link and the WLAN link. We then derive the steady state TFRC throughput using an iterative algorithm. Analytical and extensive simulation results reveal how the end-to-end TFRC throughput is affected by the number of users in a mobile hotspot, the vehicle velocity, the WWAN/WLAN link bandwidth, the retransmission limit, and the buffer size. It is found that the WWAN channel profile and link bandwidth have significant impacts on the TFRC throughput, and therefore suitable resource allocation and admission control are indispensable for the quality and efficiency of multimedia services in mobile hotspots.

Index Terms—Mobile hotspot, TCP-friendly rate control, TFRC throughput, performance analysis.

I. INTRODUCTION

WITH the advances in wireless access technologies and the ever-increasing demand of anywhere, anytime Internet services, wireless local area network (WLAN)-based hotspot services in public areas (e.g., convention centers, cafes, airports, shopping malls, etc.) are being proliferated. In addition, the extension of hotspot services to moving vehicles (e.g., subways, buses, trains, vessels, airplanes, etc.) is gaining more attention [1], [2]. The hotspot service in a mobile platform is referred to as *mobile hotspot* [3], which is a novel concept to provide ubiquitous and always best connected (ABC) services in future wireless/mobile networks.

Figure 1 shows a typical network architecture for mobile hotspots consisting of a wireless wide area network (WWAN) and a WLAN. Within a moving vehicle, a WLAN is used to

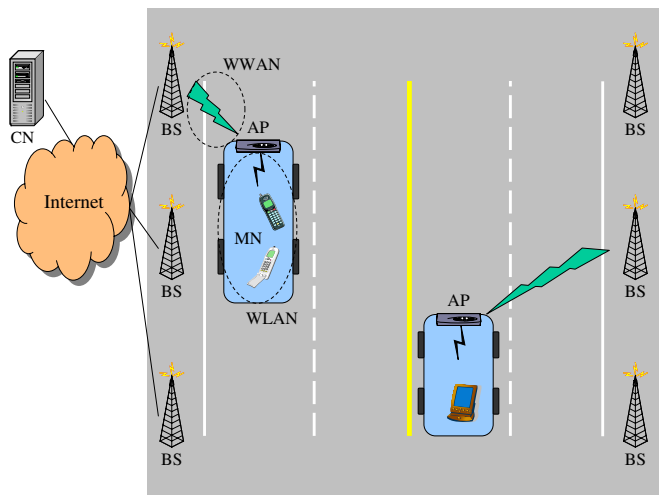


Fig. 1. Mobile hotspot architecture.

connect a number of mobile nodes (MNs) to an access point (AP). Meanwhile, a WWAN is employed for the connection between the AP and the base station (BS), which is in turn connected to the Internet through a wireline link. The WWAN can be IP-based cellular systems or IEEE 802.16 Worldwide Interoperability for Microwave Access (WiMAX) networks [4]. Packets sent from a correspondent node (CN) to an MN in the mobile hotspot are first routed to the BS through the Internet, and then transmitted to the MN over an integrated WWAN-WLAN link. This mobile hotspot architecture has the following advantages. First, the aggregated traffic at the AP is transmitted to the BS through an antenna mounted on top of the vehicle, which has better communication channels to the BS as compared to a channel between the BS and the MNs inside the vehicle. Second, the AP has less energy constraint than the MNs. Using the AP to relay the data to the BS far away can significantly save the energy consumption of MNs. Third, the AP has better knowledge of the mobility and location of the vehicle, and therefore handoff management of aggregated traffic at the AP can be simpler.

Multimedia streaming is expected to be a promising application in mobile hotspots [2]. Since multimedia streaming traffic is normally long-lived and requires high data rate, flow and congestion control is important for both network stability and users' perceived quality of service (QoS). In addition, fairness among multimedia flows and the currently dominant TCP controlled flows should be considered in mobile hotspots.

Manuscript received June 29, 2006; revised October 1, 2006; accepted December 6, 2006. The associate editor coordinating the review of this paper and approving it for publication was W. Liao. This work has been supported in part by a Strategic Grant from the Natural Sciences and Engineering Research Council (NSERC) of Canada under Grant No. STPGP 257682, and in part by the Korea Research Foundation Grant No. M01-2005-000-10073-0.

S. Pack is with the School of Electrical Engineering, Korea University, Seoul, Korea (e-mail: shpack@ieee.org).

X. S. Shen and J. W. Mark are with the Centre for Wireless Communications, Department of Electrical and Computer Engineering, University of Waterloo, Waterloo, Ontario, Canada (e-mail: {xshen, jw-mark}@bbcr.uwaterloo.ca).

L. Cai is with the Department of Electrical and Computer Engineering, University of Victoria, Victoria, British Columbia, Canada (e-mail: cai@ece.uvic.ca).

Digital Object Identifier 10.1109/TWC.2008.060426.

TCP-friendliness can be achieved when the long-term throughput of a non-TCP controlled flow does not exceed that of any TCP flows under the same circumstance [5].

In the literature, a large number of rate control protocols with TCP-friendliness have been proposed [6]. A representative TCP-friendly protocol is the TCP-Friendly Rate Control (TFRC) protocol [7]. A TFRC sender sets the sending rate according to a TCP response function [8], which is a function of round-trip time (RTT), retransmission timeout, packet size, and packet loss event rate. To determine the TFRC sending rate, the TFRC receiver sends feedback messages to the TFRC sender at least once per round-trip time or whenever a packet loss event is detected. Since TFRC is less aggressive in probing for available bandwidth and more moderate in responding to transient network congestion, the TFRC flow throughput is much smoother than that of a TCP flow. In addition, for time-sensitive applications, the TFRC sender does not need to retransmit lost packets to avoid an excessive delay due to end-to-end retransmissions.

Originally, TFRC has been designed for wired networks and its performance in wired networks has been extensively investigated. Unlike wired communications, wireless communications are characterized by limited bandwidth, high bit error rate, time-varying and location-dependent channel conditions, etc. Several studies to quantify and improve the performance of TFRC in wireless networks have been reported. Borri *et al.* [9] evaluate the performance of TFRC in an experimental IEEE 802.11g WLAN testbed. To overcome throughput degradation and unfairness problems, they propose a tuning scheme using a timer and backoff parameters. Li *et al.* [10] investigate TFRC over wireless ad-hoc networks. To avoid setting an inaccurate sending rate, the sender determines the sending rate depending on measurement results and the model value for round-trip time in IEEE 802.11 wireless ad-hoc networks. Nahm *et al.* [11] examine the steady state TCP behavior and utility of TFRC in IEEE 802.11 multi-hop networks, and propose a fractional window scheme which resembles the stop-and-go protocol. Chen and Zakhori [12] propose a multiple TFRC (MULTFRC) scheme to fully utilize the wireless link, and conduct actual experiments for video streaming over 1xRTT CDMA wireless data networks. Shen *et al.* [13] develop a discrete-time Markov model for TFRC over wireless fading channels and evaluate the TFRC performance under different situations. However, to the best of our knowledge, no studies on the performance of TFRC in mobile hotspots have been reported.

In this paper, we consider multimedia streaming applications over heterogeneous wireless links and evaluate the throughput of TFRC in mobile hotspots by developing discrete-time queuing models for the WWAN link and the WLAN link. Based on the queuing models, the packet loss probability and RTT of a TFRC flow are derived, and the TFRC throughput in steady state is obtained using an iterative algorithm. We further investigate how the end-to-end TFRC throughput is affected by the number of MNs in the mobile hotspot, the vehicle velocity, the WWAN/WLAN link bandwidth, the retransmission limit, and the buffer size. Extensive simulation results are given to validate the analytical results.

Our major contributions are as follows. We propose a

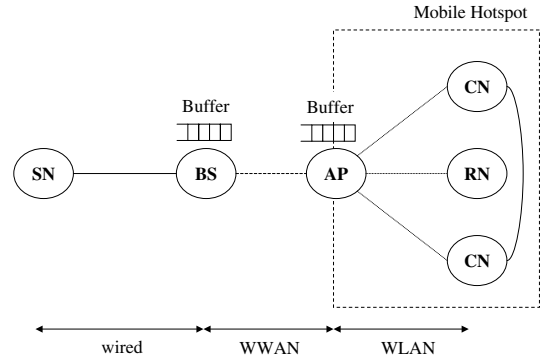


Fig. 2. System model.

novel analytical framework to evaluate the TFRC throughput in mobile hotspots. Specifically, a dedicated WWAN link with a truncated automatic retransmission request (ARQ) scheme over a Rayleigh fading channel is considered. For the WLAN link, an unsaturated infrastructure-based WLAN using the IEEE 802.11 distributed coordination function (DCF) is considered. Based on the numerical results from the analytical model, we present valuable observations that provide important guidelines for resource management and admission control in mobile hotspots. To the best of our knowledge, this is the first research work to quantify the performance of multimedia applications in mobile hotspots.

The remainder of this paper is organized as follows. In Section II, the system model for the WWAN and WLAN links is described. In Section III, the throughput model for the TFRC flow in mobile hotspots is developed and an iterative algorithm for calculating the steady state TFRC throughput is presented. Various analytical and simulation results are given in Section IV, followed by concluding remarks in Section V.

II. SYSTEM DESCRIPTION

Figure 2 shows the system model for analyzing TFRC throughput in mobile hotspots. We consider a TFRC flow between a TFRC sender node (SN) and a receiver node (RN). The multimedia flow is controlled by TFRC, and we focus on the downlink transmission (i.e., from the SN to the RN) for multimedia streaming applications. The model can be extended for interactive multimedia applications. Since the link layer fragmentation is not considered here, the term *packet* is used for a protocol data unit (PDU) in the data link layer as well as a PDU in the transport layer.

With an emphasis on the packet losses in the wireless links, it is assumed that packet losses in the wired domain is negligible and the transmission delay in the wired domain between the SN and the BS, t_{wired} , is constant. For downlink transmissions in the WWAN (i.e., from the BS to the AP), a dedicated channel is allocated, so transmission errors are mainly due to channel fading. In addition, a truncated ARQ scheme is used for reliable transmission and an instantaneous feedback channel is assumed [13]. Hence, the feedback reception time for the ARQ scheme is zero. The WWAN uplink transmissions (i.e., from the AP to the BS) of TFRC acknowledgements (*acks*) are assumed error-free. This assumption can be justified

because the size of *acks* is relatively small and an occasional loss of *ack* can be recovered by the subsequent *acks*.

The infrastructure-based WLAN is shared by N mobile nodes (MNs), i.e., the RN and $N - 1$ contending nodes, and an AP which is connected to the BS through the WWAN link. In an infrastructure-based WLAN, all traffic from and to all MNs will go through the AP. Generally, the AP is the bottleneck and the MNs are unsaturated, i.e., they do not always have data to send. Since the WLAN is used for the internal connection within a vehicle, transmission failures can be assumed due to collisions only and a collision is caused by simultaneous transmissions from two MNs or the AP and an MN [14]. The WLAN operates with the IEEE 802.11 distributed coordination function (DCF) [15] in a basic access mode, since the request-to-send (RTS)/clear-to-send (CTS) exchange is not very useful for infrastructure-based WLANs and it is disabled in most products available in the current market [16].

Since the load of upstream traffic (from MNs to the BS) is typically much less than that of downstream traffic (from the BS to MNs) in multimedia streaming applications, the downlink from the BS to the AP and that from the AP to MNs are bottlenecks. The BS and the AP have buffers of sizes $B - 1$ and $Q - 1$, respectively.

In the following subsections, we introduce two discrete-time queuing models for the WWAN and WLAN links. Table I summarizes the notations for the queuing models.

A. WWAN Link Model

For the downlink transmission from the BS to the AP, a non-line-of-sight (NLOS) frequency-nonselective (flat) multipath fading channel with packet transmission rate (in packets/seconds) much higher than the maximum Doppler frequency (Hz) is considered. Given the modulation scheme, the dynamics of the fading channel can be characterized at the packet level. However, the performance analysis of high-level protocols becomes quite complex. As an alternative to this problem, a widely adopted two-state Markov channel model is used to approximate the error process at the packet level [17]. The discrete-time two-state Markov channel model has a *good* state and a *bad* state: packet loss probability is one in the *bad* state and zero in the *good* state. The duration of a time slot, D , corresponds to the packet transmission time over the WWAN link. Let v and f_c be the velocity of the vehicle and the carrier frequency, respectively. The Doppler frequency f_d is given by $f_c v / v_c$, where v_c is the speed of light¹. Let F be the fading margin. Then the average transmission error probability is

$$\pi_e = 1 - e^{-1/F}. \quad (1)$$

Let $\mathbf{P} = \begin{bmatrix} p_{bb} & p_{bg} \\ p_{gb} & p_{gg} \end{bmatrix}$ be the WWAN link state transition matrix. The state transition probabilities are

$$p_{bb} = 1 - \frac{Q(\theta, \rho\theta) - Q(\rho\theta, \theta)}{e^{1/F} - 1}$$

¹We assume that the distance between the vehicle and the BS is much larger than the distance between the BS and the moving track of the vehicle. Therefore, the angle between the moving direction and the communication link can be neglected to simplify the analysis.

and

$$p_{gg} = \frac{1 - \pi_e(2 - p_{bb})}{1 - \pi_e}$$

where $\theta = \sqrt{\frac{2/F}{1-\rho^2}}$ and $\rho = J_0(2\pi f_d D)$ [17]. ρ is the Gaussian correlation coefficient between two samples of the complex amplitude of a fading channel with Doppler frequency f_d , which are sampled D seconds apart. $J_0(\cdot)$ is the Bessel function of the first kind and zero order and $Q(\cdot, \cdot)$ is the Marcum Q function.

In the truncated ARQ scheme, if a sender experiences a transmission failure due to a bad link condition, it retransmits the packet at the next time slot. A packet is removed from the BS buffer after being successfully delivered or after l attempts (including the first transmission). Therefore, the packet loss probability in the WWAN link can be computed as

$$\varepsilon_W = \pi_e p_{bb}^{l-1}. \quad (2)$$

The arrival process of TFRC packets at the BS buffer can be approximated by a stationary Bernoulli process, and simulation results show that the Bernoulli model is acceptable [13]. Therefore, the WWAN downlink transmission at the BS buffer is modeled as a discrete-time $M/M/1/K$ queue with the time slot length D . Let λ_T (packets/sec) be the sum of downlink transmission rates of the N MNs in a mobile hotspot. Then, the arrival rate to the BS buffer in a given time slot, λ_B (packets/slot), is computed as

$$\lambda_B = \lambda_T D. \quad (3)$$

On the other hand, the service rate, μ_B (packets/slot), of the BS buffer can be obtained by deriving the average service time as (4) at next page. In (4), $p_{gb}p_{bb}^{i-2}p_{bg}$ is the probability that a packet transmission is successful at the i th transmission attempt ($2 < i < l$) when the WWAN link state at the last transmission of the previous packet is g , whereas $p_{bb}^{i-1}p_{bg}$ is the probability that a packet transmission is successful at the i th transmission attempt ($2 < i < l$) when the WWAN link state at the last transmission of the previous packet is b . Therefore, the first and second terms on the right-hand side of (4) represent the average service times when the previous packet departs successfully and unsuccessfully, respectively.

B. WLAN Link Model

Bianchi [18] has proposed an analytical model for IEEE 802.11 DCF WLANs, which considers a saturated condition where each node always has data to send. However, the saturated condition is not realistic for infrastructure-based WLANs. Therefore, we extend the queuing model in [14] for an infrastructure-based WLAN. The downlink transmission at the AP and the uplink transmission at the MN are modeled as a discrete-time $M/M/1/K$ queue and a discrete-time $M/M/1$ queue, respectively [19], where a slot length corresponds to the length of a backoff slot δ . In IEEE 802.11 DCF, a packet is dropped if the packet transmission fails after $m + 1$ attempts. Therefore, for WLAN downlink transmissions from the AP, the packet loss probability due to collisions is given by

$$\varepsilon_L^A = p_A^{m+1} \quad (5)$$

TABLE I
SUMMARY OF IMPORTANT NOTATIONS.

symbol	explanation
$\lambda_M(\lambda_M^U)$	Downstream (Upstream) transmission rate to (from) an MN
μ_M	Service rate at the MN queue
N	The number of MNs in a mobile hotspot
D	Time slot length in the WWAN link
λ_T	The sum of downlink transmission rates of N MNs
λ_B	Traffic arrival rate at the BS downlink buffer in a time slot
μ_B	Service rate of the BS downlink
δ	Time slot length in the WLAN link (i.e., a backoff slot length)
λ_A	Traffic arrival rate at the AP buffer in a time slot
μ_A	Service rate of the AP downlink
ε_W	Packet loss probability in the WWAN link
$\varepsilon_L^A(\varepsilon_L^M)$	Packet loss probability in the WLAN downlink (uplink) transmission
$P_B(\bar{P}_Q)$	Blocking probability at the BS (AP) downlink buffer
$p_A(p_M)$	Conditional collision probability when the AP (MN) transmits a packet
B	AP system size (=buffer size + 1)
$\bar{Q}$	BS system size (=buffer size + 1)
$\bar{C}_A(\bar{C}_M)$	Average number of collisions during a packet transmission by the AP (MN)
$\bar{C}_A Succ$	Average number of collisions during a successful packet transmission by the AP
$\bar{BO}_A(\bar{BO}_M)$	Average number of backoff slots during a packet transmission by the AP (MN)
$\bar{BO}_A Succ$	Average number of backoff slots during a successful packet transmission by the AP
$T_C(T_S)$	Numbers of time slots for a collided (successful) transmission
m	The maximum number of retransmissions
m'	The number of contention window sizes
θ_S	Average service time for a successful WWAN downlink transmission (in numbers of time slots)
θ_S^U	Average service time for a successful WWAN uplink transmission (in seconds)
η_S	Average service time for a successful WLAN downlink transmission (in numbers of time slots)
η_S^U	Average service time for a successful WLAN uplink transmission (in numbers of time slots)
Q_B	Queuing delay at the BS buffer (in downlink transmission)
Q_A	Queuing delay at the AP buffer (in downlink transmission)
$\tau_A(\tau_M)$	Probability that the AP (MN) with a non-empty queue transmits a packet in a randomly chosen slot
$\sigma_A(\sigma_M)$	Probability that the queue of the AP (MN) is not empty

$$\begin{aligned}
\frac{1}{\mu_B} &= (1 - \pi_e)(p_{gg} + 2p_{gb}p_{bg} + 3p_{gb}p_{bb}p_{bg} + \dots + (l-1)p_{gb}p_{bb}^{l-3}p_{bg} + lp_{gb}p_{bb}^{l-2}) \\
&+ \pi_e(p_{bg} + 2p_{bb}p_{bg} + 3p_{bb}p_{bb}p_{bg} + \dots + (l-1)p_{bb}p_{bb}^{l-3}p_{bg} + lp_{bb}p_{bb}^{l-2}) \\
&= (1 - \pi_e) \left(1 + \frac{p_{gb}(1 - p_{bb}^{l-1})}{1 - p_{bb}} \right) + \pi_e \left(\frac{1 - p_{bb}^{l-1}}{1 - p_{bb}} + p_{bb}^{l-1} \right)
\end{aligned} \tag{4}$$

where p_A is the collision probability when the AP transmits a packet. Similarly, for WLAN uplink transmissions from unsaturated MNs, the packet loss probability due to collisions, ε_L^M , can be computed as p_M^{m+1} , where p_M is the collision probability when an MN transmits a packet.

For downstream traffic, the arrival rate at the AP, λ_A (packet/slot), equals the departure rate of the successfully transmitted packets over the WWAN downlink. Therefore, λ_A can be approximated as

$$\lambda_A = (1 - P_B)(1 - \varepsilon_W)\lambda_T\delta. \tag{6}$$

where P_B is the blocking probability at the BS downlink buffer. On the other hand, the service rate at the AP downlink, μ_A , can be obtained from the average service time at the AP, which includes the time for backoffs and that for successful or collided transmissions by the AP. In addition, during the interval between two packets serviced by the AP, MNs may transmit some packets successfully or experience collisions. We assume a collision occurs by simultaneous transmissions of two stations (i.e., the AP and an MN or two MNs). Therefore, the average number of collisions during the AP service time (i.e., $1/\mu_A$) is one half of the total number of collisions

during the corresponding period because a collision can be counted twice in the service times of the two stations. Let λ_M^U be the upstream transmission rate of an MN. Consequently, the average service time of the AP is given by

$$\begin{aligned}
\frac{1}{\mu_A} &= \left(N \frac{\lambda_M^U}{\mu_A} (1 - \varepsilon_L^M) + (1 - \varepsilon_L^A) \right) T_S \\
&+ \frac{1}{2} \left(N \frac{\lambda_M^U}{\mu_A} \bar{C}_M + \bar{C}_A \right) T_C + \bar{BO}_A,
\end{aligned} \tag{7}$$

where $\left(N \frac{\lambda_M^U}{\mu_A} (1 - \varepsilon_L^M) + (1 - \varepsilon_L^A) \right)$ is the average number of successful transmissions during $1/\mu_A$ and $\left(N \frac{\lambda_M^U}{\mu_A} \bar{C}_M + \bar{C}_A \right)$ is the total number of collisions during $1/\mu_A$. T_S and T_C are the numbers of time slots for successful and collided transmissions, respectively. $\bar{BO}_A$ is the average number of backoff slots during a packet transmission by the AP, and $\bar{C}_A$ and $\bar{C}_M$ are the average numbers of collisions during a packet transmission by the AP and the MN, respectively.

The probability of i collisions ($0 \leq i \leq m$) for a successful transmission by the AP is $p_A^i(1 - p_A)$. For a failed AP transmission, $m + 1$ collisions occur and its probability is p_A^{m+1} . Then, $\bar{C}_A$ and $\bar{C}_M$ can be computed as (8) and (9),

$$\overline{C_A} = 1 \cdot p_A(1 - p_A) + 2 \cdot p_A^2(1 - p_A) + \dots + m \cdot p_A^m(1 - p_A) + (m + 1) \cdot p_A^{m+1} = \frac{p_A(1 - p_A^{m+1})}{1 - p_A}. \quad (8)$$

respectively.

In the AP transmission, the probability that the i th backoff ($0 \leq i \leq m$) is triggered is p_A^i . Therefore, $\overline{BO_A}$ is given by

$$\begin{aligned} \overline{BO_A} &= \frac{W_0 - 1}{2} \cdot 1 + \dots + \frac{W_m - 1}{2} \cdot p_A^m \\ &= \sum_{i=0}^m \frac{W_i - 1}{2} p_A^i \end{aligned} \quad (10)$$

where $W_i = 2^i W$ and W is the minimum contention window size. T_C and T_S are given by $T_S = (DIFS + H + P + \sigma + SIFS + ACK + \sigma)/\delta$ and $T_C = (DIFS + H + P + SIFS + ACK)/\delta$, where σ is the propagation delay, and $DIFS$ and $SIFS$ represent the DCF inter-frame space and small inter-frame space, respectively. H , P , ACK are the transmission times of the header, payload, and ACK frame, respectively.

The service rate of the MN queue, μ_M , can be obtained as follows. During the service time of the MN, the average number of successful transmissions by other MNs, except the tagged MN, is $\left((N - 1) \frac{\lambda_M^U}{\mu_M}\right) (1 - \varepsilon_L^M)$ and the average number of successful transmissions by the AP is $\frac{\lambda_A(1 - P_Q)}{\mu_M} (1 - \varepsilon_L^A)$, where P_Q is the blocking probability at the AP downlink buffer. On the other hand, the numbers of collisions by other MNs and the AP are $(N - 1) \frac{\lambda_M^U}{\mu_M} \overline{C_M}$ and $\frac{\lambda_A(1 - P_Q)}{\mu_M} \overline{C_A}$, respectively. Therefore, the average service time at the MN queue can be computed as (11).

Similar to $\overline{BO_A}$, the average number of backoff slots with respect to the MN can be computed as

$$\begin{aligned} \overline{BO_M} &= \frac{W_0 - 1}{2} \cdot 1 + \dots + \frac{W_m - 1}{2} \cdot p_M^m \\ &= \sum_{i=0}^m \frac{W_i - 1}{2} p_M^i. \end{aligned} \quad (12)$$

In the unsaturated infrastructure-based WLAN, the AP and the MN transmit a packet only when their queues are non-empty. Let σ_A and σ_M be the probabilities that the AP and the MN queues are not empty, respectively. They are given by

$$\sigma_A = \frac{\rho_A(1 - \rho_A^Q)}{(1 - \rho_A^{Q+1})} \quad \text{and} \quad \sigma_M = \rho_M$$

where $\rho_A = \lambda_A/\mu_A$ and $\rho_M = \lambda_M^U/\mu_M$. From [20], the probability that the AP with a non-empty queue transmits a packet in a randomly chosen slot is given by (13) where m is the retransmission limit and m' is the number of contention window sizes (i.e., the maximum contention window size is $2^{m'}$). Similarly, the probability of a packet transmission by a non-empty MN is given by (14).

Then, the probabilities that the AP and the MN transmit a packet in a randomly chosen slot are $\sigma_A \tau_A$ and $\sigma_M \tau_M$, respectively. p_A and p_M are given by

$$p_A = 1 - (1 - \sigma_M \tau_M)^N \quad (15)$$

and

$$p_M = 1 - (1 - \sigma_M \tau_M)^{N-1} (1 - \sigma_A \tau_A). \quad (16)$$

Finally, (7), (11), (15), and (16) can be solved numerically to obtain μ_A , μ_M , p_A , and p_M .

III. TFRC THROUGHPUT ANALYSIS

Since TFRC is a rate-based protocol, the TFRC sender produces a smooth flow with rate λ (packets/sec) that is determined by

$$\lambda = \frac{s}{r \sqrt{\frac{2p}{3}} + 3p(t_{RTO}(1 + 32p^2) \sqrt{\frac{3p}{8}})} \quad (17)$$

where r is the round-trip time, t_{RTO} is the retransmission timeout value, s is the packet size, and p is the packet loss event rate. p is measured by the TFRC receiver, while r and t_{RTO} are estimated and calculated by the TFRC sender. Initially, the TFRC sender sets its sending rate to one packet per second and doubles the rate every RTT until a packet loss occurs. Thereafter, the sending rate is determined by (17). Therefore, we need to determine λ , r , and p to quantify the throughput of the TFRC flow.

First, a packet in the TFRC flow can be lost due to overflows at the BS/AP buffer or transmission errors in the WWAN/WLAN link. Therefore, the packet loss rate can be obtained from

$$p = 1 - (1 - P_B)(1 - P_Q)(1 - \varepsilon_W)(1 - \varepsilon_L^A). \quad (18)$$

where ε_W and ε_L^A have been obtained in Section II. From queuing theory, P_B and P_Q can be computed by

$$P_B = \frac{(1 - \rho_B)\rho_B^B}{1 - \rho_B^{B+1}} \quad \text{and} \quad P_Q = \frac{(1 - \rho_A)\rho_A^Q}{1 - \rho_A^{Q+1}} \quad (19)$$

where $\rho_B = \lambda_B/\mu_B$ and $\rho_A = \lambda_A/\mu_A$. The arrival rates and service rates at the BS and the AP have been derived in Section II.

On the other hand, the round-trip time r can be expressed as

$$r = 2t_{\text{wired}} + t_{\text{wireless}}^{\text{down}} + t_{\text{wireless}}^{\text{up}} \quad (20)$$

where t_{wired} is the transmission latency in the wired link (i.e., from the SN to the BS). $t_{\text{wireless}}^{\text{down}}$ and $t_{\text{wireless}}^{\text{up}}$ are the latencies for downlink and uplink transmissions in the integrated WWAN-WLAN link, respectively.

The term $t_{\text{wireless}}^{\text{down}}$ is given by

$$t_{\text{wireless}}^{\text{down}} = D \cdot (Q_B + \theta_S) + \delta \cdot (Q_A + \eta_S) \quad (21)$$

where D and δ are time slot lengths in the queuing models for the WWAN and WLAN links, respectively. θ_S and η_S are the average service times for a successful WWAN downlink transmission and for a successful WLAN downlink transmission, respectively. Q_B and Q_A are the queueing delays at the BS buffer and the AP buffer for downlink transmission, respectively. By the $M/M/1/K$ queuing model, Q_B and Q_A are given by

$$Q_B = \frac{1}{\lambda_B(1 - P_B)} \left(\frac{\rho_B}{1 - \rho_B} - \frac{\rho_B(B\rho_B^B + 1)}{1 - \rho_B^{B+1}} \right) \quad (22)$$

$$\overline{C_M} = 1 \cdot p_M(1 - p_M) + 2 \cdot p_M^2(1 - p_M) + \dots + m \cdot p_M^m(1 - p_M) + (m + 1) \cdot p_M^{m+1} = \frac{p_M(1 - p_M^{m+1})}{1 - p_M}. \quad (9)$$

$$\begin{aligned} \frac{1}{\mu_M} &= \left(\left((N - 1) \frac{\lambda_M^U}{\mu_M} + 1 \right) (1 - \varepsilon_L^M) + \frac{\lambda_A(1 - P_Q)}{\mu_M} (1 - \varepsilon_L^A) \right) T_S \\ &+ \frac{1}{2} \left(\left((N - 1) \frac{\lambda_M^U}{\mu_M} + 1 \right) \overline{C_M} + \frac{\lambda_A(1 - P_Q)}{\mu_M} \overline{C_A} \right) T_C + \overline{BO_M}. \end{aligned} \quad (11)$$

$$\tau_A = \begin{cases} \frac{2(1-2p_A)(1-p_A^{m+1})}{W(1-(2p_A)^{m+1})(1-p_A)+(1-2p_A)(1-p_A^{m+1})} & m \leq m' \\ \frac{2(1-2p_A)(1-p_A^{m+1})}{W(1-(2p_A)^{m'+1})(1-p_A)+(1-2p_A)(1-p_A^{m+1})+W2^{m'}p_A^{m'+1}(1-2p_A)(1-p_A^{m-m'})} & m > m' \end{cases} \quad (13)$$

$$\tau_M = \begin{cases} \frac{2(1-2p_M)(1-p_M^{m+1})}{W(1-(2p_M)^{m+1})(1-p_M)+(1-2p_M)(1-p_M^{m+1})} & m \leq m' \\ \frac{2(1-2p_M)(1-p_M^{m+1})}{W(1-(2p_M)^{m'+1})(1-p_M)+(1-2p_M)(1-p_M^{m+1})+W2^{m'}p_M^{m'+1}(1-2p_M)(1-p_M^{m-m'})} & m > m' \end{cases} \quad (14)$$

and

$$Q_A = \frac{1}{\lambda_A(1 - P_Q)} \left(\frac{\rho_A}{1 - \rho_A} - \frac{\rho_A(Q\rho_A^Q + 1)}{1 - \rho_A^{Q+1}} \right). \quad (23)$$

The derivations of θ_S and η_S are similar to those of $1/\mu_B$ and $1/\mu_A$, respectively, except that the successful packet transmission is assumed. θ_S is then computed as (24) where $1 - p_{gb}p_{bb}^{l-1}$ and $1 - p_{bb}^l$ are the probabilities that a packet transmission is successful when the WWAN link states at the last transmission of the previous packet are g and b , respectively. On the other hand, η_S can be derived as

$$\begin{aligned} \eta_S &= (N\lambda_M^U\eta_S(1 - \varepsilon_L^M) + 1) T_S \\ &+ \frac{1}{2} (N\lambda_M^U\eta_S\overline{C_M} + \overline{C_A}|Succ) T_C + \overline{BO_A}|Succ \end{aligned} \quad (25)$$

where $\overline{C_A}|Succ$ and $\overline{BO_A}|Succ$ are the average numbers of collisions and backoffs experienced by the AP when a packet is successfully transmitted over the WLAN link, respectively. They are respectively given by (26) and (27) where $1 - p_A^{m+1}$ is the probability of a successful transmission by the AP.

The queuing delay for the upstream transmission is negligible. Therefore, $t_{wireless}^{up}$ is given by the transmission latency in the WWAN-WLAN link as

$$t_{wireless}^{up} = \delta \cdot \eta_S^U + \theta_S^U. \quad (28)$$

where θ_S^U and η_S^U are the average service times for a successful WWAN uplink transmission and for a successful WLAN uplink transmission, respectively. η_S^U can be derived as (29). On the other hand, since an ideal WWAN uplink channel is assumed, θ_S^U is simply given by P_{ack}/W_{up} , where P_{ack} and W_{up} are the TFRC *ack* size and the WWAN uplink bandwidth, respectively.

We can derive the TFRC throughput in steady state using an iterative algorithm given in Algorithm 1. We consider two cases: 1) there is only one TFRC flow and the sending rates of other flows are fixed, i.e., constant bit rate (CBR) flows, (referred to as *single TFRC flow case*) and 2) each MN has a TFRC flow and therefore there are N TFRC flows in a mobile hotspot (referred to as *multiple TFRC flows case*). For the

single TFRC flow case, let λ'_M be the downlink transmission rate of the tagged MN with a TFRC flow. The sending rate λ_{Init} is initialized to 1.0 and then λ_T can be determined as $\lambda_T = \lambda_{Init} + (N - 1)\lambda_F$, where λ_F is the constant sending rate (in packets/sec) of an MN except the tagged MN with a TFRC flow. In the sequel, r and p are computed using λ_T , (18) and (20). t_{RTO} is set to $4r$ [7] and a new TFRC sending rate λ'_M is calculated in line 5 of Algorithm 1. In lines 6-12, using a sufficiently small value ϵ , λ'_M is repeatedly calculated until it converges². On the other hand, for the multiple TFRC flows case, λ_T is set according to $\lambda_T \leftarrow N\lambda_{Init}$ in lines 2 and 8, and λ'_M is repeatedly computed using (17). Consequently, the throughput of a TFRC flow, T , can be computed as $T = \lambda^*(1 - p^*)$ [12], where λ^* and p^* are the TFRC sending rate and the packet loss rate in steady state, respectively.

Algorithm 1 The iterative algorithm.

- 1: $\lambda_{Init} \leftarrow 1$;
 - 2: $\lambda_T \leftarrow \lambda_{Init} + (N - 1)\lambda_F$;
 - 3: Calculate r and p using λ_T , (18), and (20);
 - 4: $t_{RTO} \leftarrow 4 \times r$;
 - 5: Calculate λ'_M using (17);
 - 6: **while** $|\lambda'_M - \lambda_{Init}| \geq \epsilon$ **do**
 - 7: $\lambda_{Init} = (\lambda'_M + \lambda_{Init})/2$;
 - 8: $\lambda_T \leftarrow \lambda_{Init} + (N - 1)\lambda_F$;
 - 9: Calculate r and p using λ_T , (18), and (20);
 - 10: $t_{RTO} \leftarrow 4 \times r$;
 - 11: Calculate λ'_M using (17);
 - 12: **end while**
-

IV. NUMERICAL RESULTS

In the simulation, the network topology is the same as that shown in Figure 2. The following parameters are used unless otherwise explicitly stated. The carrier frequency (f_c) of the WWAN link is 900 MHz. The payload sizes of a

²TFRC adjusts the sending rate as a monotonic decreasing function of the packet loss rate whereas the packet loss rate is a monotonic non-decreasing function of the sending rate. Therefore, the system will converge to the steady state, which is the crossing point of two functions for the sending rate and the packet loss rate [13].

$$\begin{aligned}
\theta_S &= (1 - \pi_e) \left(\frac{1}{1 - p_{gb}p_{bb}^{l-1}} (p_{gg} + 2p_{gb}p_{bg} + 3p_{gb}p_{bb}p_{bg} + \dots + lp_{gb}p_{bb}^{l-2}p_{bg}) \right) \\
&+ \pi_e \left(\frac{1}{1 - p_{bb}^l} (p_{bg} + 2p_{bb}p_{bg} + 3p_{bb}p_{bb}p_{bg} + \dots + lp_{bb}p_{bb}^{l-2}p_{bg}) \right) \\
&= \frac{1 - \pi_e}{1 - p_{gb}p_{bb}^{l-1}} \left(1 + p_{gb} \left(\frac{1 - p_{bb}^{l-1}}{1 - p_{bb}} - lp_{bb}^{l-1} \right) \right) + \frac{\pi_e}{1 - p_{bb}^l} \left(\frac{1 - p_{bb}^l}{1 - p_{bb}} - lp_{bb}^l \right)
\end{aligned} \quad (24)$$

$$\overline{C_A} | Succ = 1 \cdot \frac{p_A(1 - p_A)}{1 - p_A^{m+1}} + 2 \cdot \frac{p_A^2(1 - p_A)}{1 - p_A^{m+1}} + \dots + m \cdot \frac{p_A^m(1 - p_A)}{1 - p_A^{m+1}} = \frac{1}{1 - p_A^{m+1}} \left(\frac{p_A(1 - p_A^m)}{1 - p_A} - mp_A^{m+1} \right) \quad (26)$$

$$\begin{aligned}
\overline{BO_A} | Succ &= \frac{W_0 - 1}{2} \cdot \frac{(1 - p_A^{m+1})}{1 - p_A^{m+1}} + \frac{W_1 - 1}{2} \cdot \frac{(p_A - p_A^{m+1})}{1 - p_A^{m+1}} + \dots + \frac{W_m - 1}{2} \cdot \frac{(p_A^m - p_A^{m+1})}{1 - p_A^{m+1}} \\
&= \sum_{i=0}^m \frac{W_i - 1}{2} \frac{(p_A^i - p_A^{m+1})}{1 - p_A^{m+1}}
\end{aligned} \quad (27)$$

$$\begin{aligned}
\eta_S^U &= ((N - 1)\lambda_M^U \eta_S^U (1 - \varepsilon_L^M) + \lambda_A(1 - P_Q)\eta_S^U (1 - \varepsilon_L^A) + 1) T_S \\
&+ \frac{1}{2} ((N - 1)\lambda_M^U \eta_S^U \overline{C_M} + \overline{C_M} + \lambda_A(1 - P_Q)\eta_S^U \overline{C_A}) T_C + \overline{BO_M}.
\end{aligned} \quad (29)$$

data packet and an *ack* packet are fixed to 250 bytes and 50 bytes, respectively. The default downlink and uplink WWAN bandwidths are 400 Kbps and 80 Kbps, respectively. Hence, the transmission time of a packet (i.e., data or *ack* packet) over the WWAN link is 5 msec. The default velocity (v) and fading margin (F) are 20 m/s and 10 dB, respectively. The default values of retransmission limits (i.e., l and m) for the WWAN and WLAN links are 3 and 5, respectively. The number of MNs within a vehicle is varied from 1 to 40 and its default value is 20. t_{wired} is set to 20 msec. The default system sizes of the AP buffer (Q) and the BS buffer (B) are 10 (packets). The parameters for WLAN follow those of the IEEE 802.11b specification in [20] and the data rate is 11 Mbps. To validate analytical results, simulations are performed using the *ns-2* simulator [21].

A. Effects of v and N

Figure 3 shows the effect of velocity (v) on the TFRC throughput T for the multiple TFRC flows case. Note that the impact of v on the physical layer, e.g., synchronization error, is not considered. It can be seen that, given the same level of fading margin, T increases as v increases. This observation can be explained as follows. When v is high, the WWAN link's coherence time is short and therefore the burstiness of transmission errors in the WWAN link is not severe. With non-bursty transmission errors, packet losses can be effectively reduced by the truncated ARQ scheme. Consequently, the reduced packet loss rate at high velocity results in a higher TFRC throughput. When v exceeds a certain value, around 20 m/s, T remains fairly constant. Figure 3 also shows that T decreases as the number of MNs sharing the WLAN increases. This is because a large N leads to a higher packet loss rate due to more channel collisions in

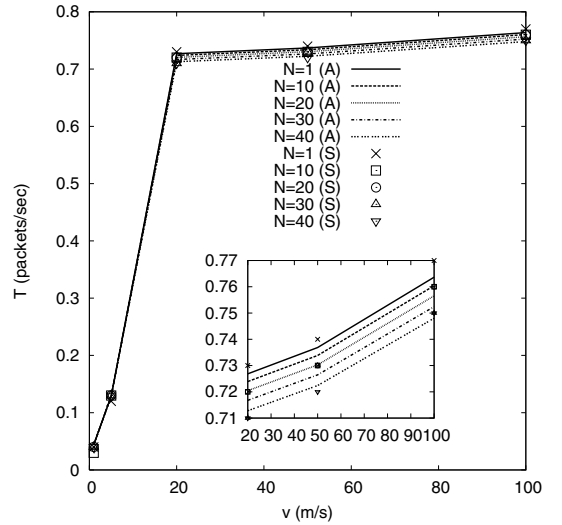


Fig. 3. T vs. v : multiple TFRC flows (A: analytical, S: simulation).

WLAN and a longer queuing delay. However, the effect of N on T is not significant, especially when v is low. This is because each TFRC flow adapts to the network conditions. In other words, when N is large, the TFRC flows will reduce their sending rates if higher packet loss rates are observed. In short, the flow and congestion control mechanisms of TFRC can intelligently adjust the sending rate to effectively mitigate network congestion, and efficiently utilize network resources independent of N .

The effects of v and N for the single TFRC flow case are illustrated in Figure 4. We consider two values for λ_F : 4 packet/sec (i.e., 8 Kbps) and 8 packets/sec (i.e., 16 Kbps). Similar to Figure 3, the TFRC throughput increases with the

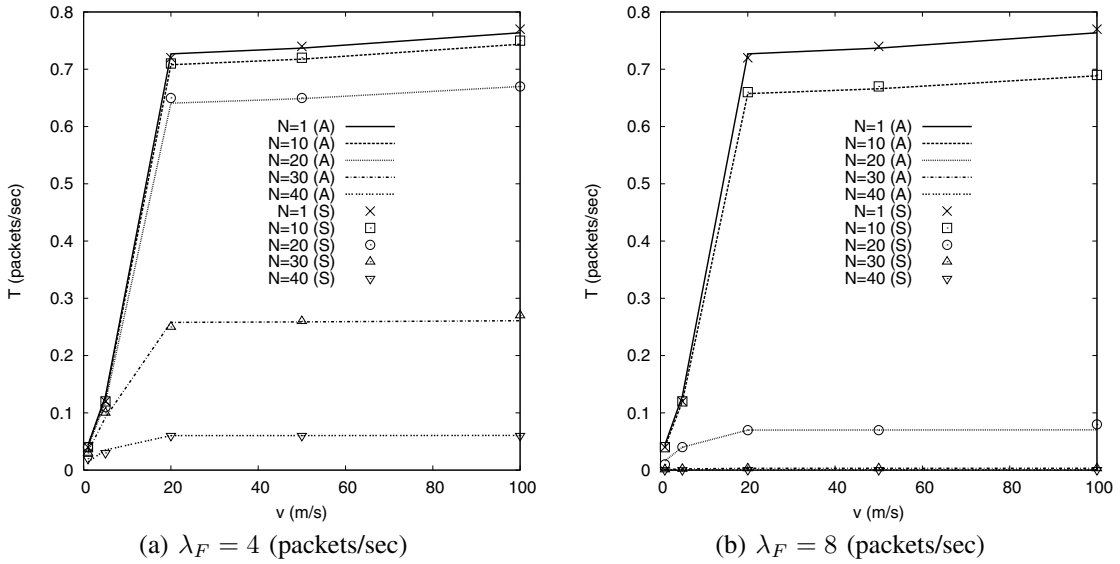


Fig. 4. T vs. v : single TFRC flow (A: analytical, S: simulation).

increase of v , and the increasing rate becomes stable when v exceeds a certain point. On the other hand, the effect of N is clear in the single TFRC case, especially when λ_F is high. From both Figures 3 and 4, the analytical results are consistent with the simulation results.

The TFRC throughputs in single and multiple TFRC cases are compared in Figure 5. For the multiple TFRC case, the TFRC throughput is not highly sensitive to N due to TFRC flow and congestion control mechanisms. However, in the single TFRC case, the TFRC throughput drastically decreases as N increases. Especially, for $\lambda_F = 8$ (packet/sec), the TFRC throughput reduces below 0.001 (packets/sec) when N is larger than 30, which indicates that a TFRC flow cannot be effectively supported. This is because the high traffic load incurs more packet losses (due to buffer overflow and channel collisions in WLAN) and a longer queuing delay, and thus degrades the TFRC throughput. Therefore, it can be shown that an admission control algorithm to limit the number of MNs in a mobile hotspot should be deployed to provide satisfactory quality of services.

B. Effects of l and m

Retransmissions up to $l - 1$ and m times are deployed in the WWAN and WLAN links, respectively, and they affect the packet loss probability and the round-trip time. Figure 6 shows the effects of l and m on the TFRC throughput. As shown in Figure 6(a), when l increases, a higher TFRC throughput can be obtained, since the packet loss rate can be significantly reduced for a large l . However, a larger l may not be preferable for delay-sensitive multimedia applications because it will increase the end-to-end delay. Therefore, an optimal l should be determined by considering the tradeoff between latency and throughput. On the other hand, the effect of m is not significant as shown in Figure 6(b). Since the WWAN is most likely to be the bottleneck and TFRC flows can adjust the sending rate to mitigate congestion in WLAN, packet losses due to collisions are rare for the multiple TFRC case.

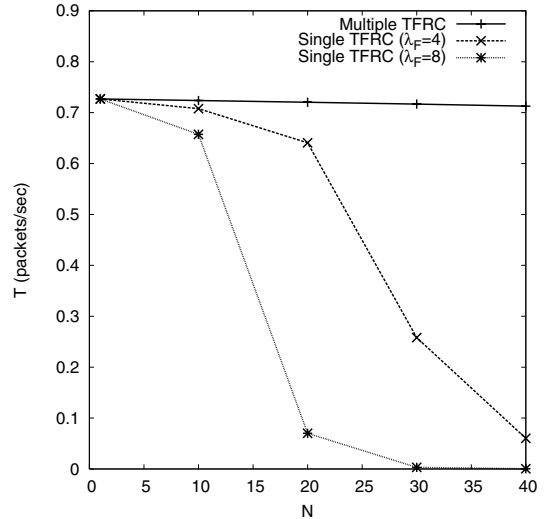


Fig. 5. Effect of N : multiple TFRC flows vs. single TFRC flow.

C. Effects of B and Q

The BS and AP buffer sizes determine the queuing delay and the packet loss rate. Table II shows the TFRC throughput when different B and Q are employed. It can be seen that the TFRC throughput remains almost the same value when Q varies and B is fixed. That is, the effect of Q on the TFRC throughput is quite limited. As mentioned before, for the multiple TFRC flows case, network congestion in WLAN is not significant due to TFRC's flow and congestion control mechanisms. Also, the WLAN bandwidth is much larger than the WWAN bandwidth. Therefore, the AP is unsaturated, i.e., $\rho_A < 1.0$. Consequently, under a lightly loaded WLAN condition, the variation of Q has no significant impact on the TFRC throughput. On the other hand, throughput degradation can be observed when B is reduced from 20 (or 10) to 5. In addition, the degradation is clearer when N is large.

Tables III and IV demonstrate the effects of B and Q for the single TFRC flow case. It can be seen that the effect of Q is

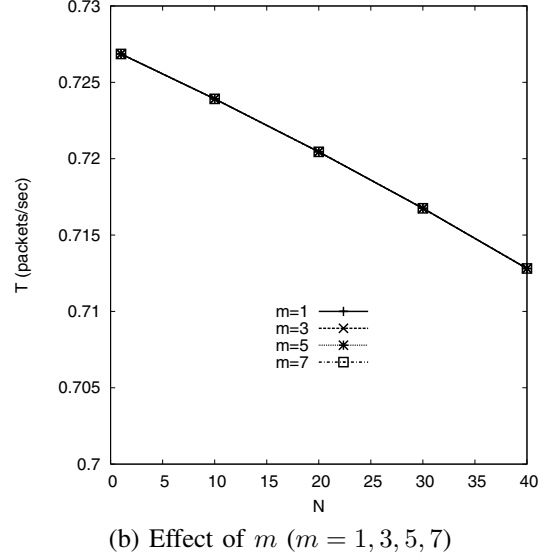
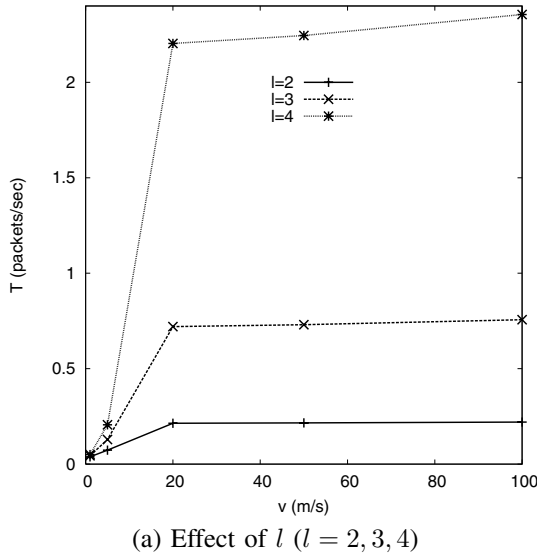


Fig. 6. Effect of retransmission limits: multiple TFRC flows.

TABLE II

TFRC THROUGHPUT FOR DIFFERENT (B, Q) : MULTIPLE TFRC FLOWS.

N	(5,10)	(10,10)	(20,10)	(10,5)	(10,20)
1	0.7239	0.7269	0.7269	0.7269	0.7269
10	0.7239	0.7239	0.7239	0.7239	0.7239
20	0.7194	0.7204	0.7204	0.7204	0.7204
30	0.7098	0.7168	0.7168	0.7168	0.7168
40	0.6894	0.7128	0.7128	0.7128	0.7128

TABLE III

TFRC THROUGHPUT FOR DIFFERENT (B, Q) : SINGLE TFRC FLOW
($\lambda_F = 4$).

N	(5,10)	(10,10)	(20,10)	(10,5)	(10,20)
1	0.7269	0.7269	0.7269	0.7269	0.7269
10	0.6271	0.7078	0.7078	0.7078	0.7078
20	0.2147	0.6407	0.6743	0.6406	0.6407
30	0.0682	0.2580	0.5941	0.2579	0.2580
40	0.0216	0.0603	0.1506	0.0603	0.0603

not significant regardless of λ_F . For $\lambda_F = 4$ (packets/sec), significant throughput degradation is observed for a small value of B . This is because a small size of BS buffer induces more packet losses due to buffer overflow. On the other hand, for $\lambda_F = 8$ (packets/sec), a large B gives a higher TFRC throughput when N is less than 20. However, when N is equal to or larger than 20, a small size of buffer can increase the TFRC throughput. Obviously, a large size of buffer can reduce the packet loss rate while it increases the round-trip time due to queuing delay. Since the WWAN link is not congested when the number of CBR flows is small, the increase in the round-trip time due to a large buffer is not significant and therefore a large buffer is better to improve the TFRC throughput. On the contrary, when there are many CBR flows (i.e., N is large), increasing the BS buffer size cannot significantly reduce the packet loss rate, while the queuing delay will be substantially prolonged. Therefore, it is not desirable to use a larger BS buffer. From these observations, it can be concluded that the dimensioning of the BS buffer size is critical to improve the throughput of the TFRC flow in mobile hotspots. In addition, it can be shown that the load of non-responsive traffic has a significant impact on the TFRC throughput in mobile hotspots.

D. Effect of WWAN/WLAN Bandwidth

Figure 7 shows the effect of the WWAN downlink bandwidth with different velocities. For $v = 5$ m/s, the TFRC throughput decreases as the allocated WWAN bandwidth increases. This counter-intuitive observation can be explained as follows. In the WWAN link, a time slot equals a packet

TABLE IV
TFRC THROUGHPUT FOR DIFFERENT (B, Q) : SINGLE TFRC FLOW
($\lambda_F = 8$).

N	(5,10)	(10,10)	(20,10)	(10,5)	(10,20)
1	0.7269	0.7269	0.7269	0.7269	0.7269
10	0.2422	0.6574	0.6785	0.6573	0.6574
20	0.0244	0.0700	0.1888	0.0700	0.0700
30	0.0024	0.0032	0.0024	0.0032	0.0032
40	0.0005	0.0004	0.0003	0.0004	0.0004

transmission time over the WWAN link. Therefore, with higher WWAN bandwidth, the transmission time of a packet is smaller and thus the time slot duration is short. When v is low, the channel coherence time in the WWAN link is long and a shorter time slot under the long channel coherence time leads to longer burst of transmission errors, which may not be recovered by the truncated ARQ scheme. Thus, more packet losses are observed by the transport layer. Consequently, at a low velocity, even though the round-trip time can be shortened by a large WWAN bandwidth, the TFRC throughput is degraded due to the high packet loss rate. This undesirable situation can be improved by exploiting diversity, e.g., retransmitting the corrupted packets at another band/subcarrier in a FDMA/OFDM system, using delayed retransmission in a TDMA system, etc., which are beyond the scope of this paper.

On the other hand, for $v = 20$ m/s and $v = 100$ m/s, it can be observed that the TFRC throughput can be improved

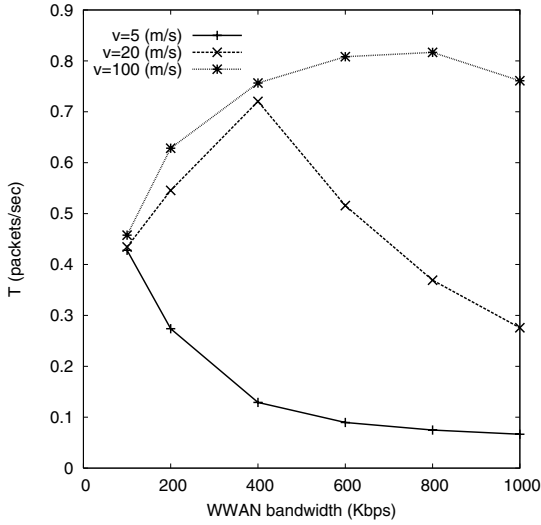


Fig. 7. Effect of WWAN bandwidth: multiple TFRC flows.

by allocating more bandwidth, since the channel coherence time is not too long. Another interesting result is that there is an optimal WWAN bandwidth to maximize the TFRC throughput. Therefore, when the allocated bandwidth is larger than the optimal value, the TFRC throughput is reduced due to the burstiness of transmission errors. In addition, the optimal WWAN bandwidth increases with v , i.e., given the wireless channel profile, it is preferable to allocate more WWAN bandwidth to maximize the TFRC throughput according to v .

To investigate the effect of the WLAN bandwidth, we consider a IEEE 802.11a [22] WLAN supporting a high data rate of 54 Mbps. As shown in Figure 8, the TFRC throughput can be improved when IEEE 802.11a is used. However, the improvement is not significant even though the bandwidth of IEEE 802.11a is much larger than that of IEEE 802.11b. Since the bottleneck of mobile hotspots is the WWAN link in general, the WLAN load is not heavy to require more bandwidth. Only if larger bandwidth is allocated in the WWAN such that the bottleneck is shifted to WLAN, the gain of higher data rate WLANs with IEEE 802.11a/g will be significant. From Figure 8, it can be seen that limiting the number of unresponsive flows is more important than increasing the WLAN bandwidth for improving the TFRC throughput and this result demonstrates the necessity of admission control in mobile hotspots.

V. CONCLUSION

In this paper, we have studied the TFRC performance in mobile hotspots. Specifically, the throughput model has been developed and the TFRC throughput in steady state has been exploited. Analytical and simulation results demonstrate that the TFRC throughput is mainly affected by the WWAN channel condition and bandwidth. In addition, since the WWAN-WLAN link is shared by multiple MNs, the traffic load has a significant impact on the TFRC throughput in mobile hotspots. Therefore, a suitable bandwidth allocation method and an admission control algorithm are necessary to support multimedia applications with QoS guarantee in mobile

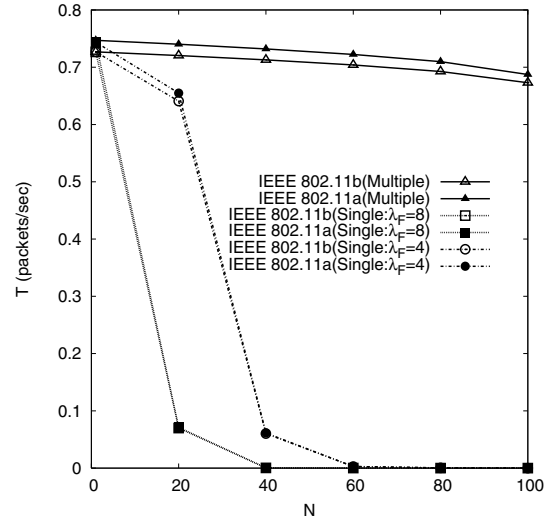


Fig. 8. Effect of WLAN bandwidth.

hotspots. The analytical and simulation results presented in this paper can be used as a guideline for effective admission control, which will be investigated in our future research work.

REFERENCES

- [1] J. Ott and D. Kutscher, "Drive-thru Internet: IEEE 802.11b for automobile users," in *Proc. IEEE INFOCOM 2004*, March 2004.
- [2] V. Mancuso and G. Bianchi, "Streaming for vehicular users via elastic proxy buffer management," *IEEE Commun. Mag.*, vol. 42, no. 11, pp. 144-152, Nov. 2004.
- [3] D. Ho and S. Valaee, "Information raining and optimal link-layer design for mobile hotspots," *IEEE Trans. Mobile Comput.*, vol. 4, no. 3, pp. 271-284, May/June 2005.
- [4] WiMAX Forum: <http://www.wimaxforum.org/home/>.
- [5] S. Floyd and K. Fall, "Promoting the use of end-to-end congestion control in the Internet," *IEEE/ACM Trans. Networking*, vol. 7, no. 4, pp. 458-472, Aug. 1999.
- [6] J. Widmer, R. Denda, M. Mauve, "A survey on TCP-friendly congestion control," *IEEE Network*, vol. 15, no. 3, pp. 28-37, May-June 2001.
- [7] S. Floyd, M. Handley, J. Padhye, and J. Widmer, "Equation-based congestion control for unicast applications," in *Proc. ACM SIGCOMM 2000*, Aug. 2000.
- [8] J. Padhye, V. Firoio, D. Towsley, and J. Kurose, "Modeling TCP Reno performance: a simple model and its empirical validation," *IEEE/ACM Trans. Networking*, vol. 8, no. 2, pp. 133-145, April 2000.
- [9] M. Borri, M. Casoni, and M. Merani, "Performance and TCP-fairness of TFRC in an 802.11g WLAN: experiments and tuning," in *Proc. Int. Symposium on Wireless Communication Systems (ISWCS) 2004*, Sept. 2004.
- [10] M. Li, C. Lee, E. Agu, M. Claypool, and R. Kinicki, "Performance enhancement of TFRC in wireless ad hoc networks," in *Proc. Int. Conference on Distributed Multimedia Systems (DMS) 2004*, Sept. 2004.
- [11] K. Nahm, A. Helmy, and C. Kuo, "On interaction between MAC and transport layers for multimedia streaming in 802.11 ad hoc networks," in *Proc. SPIE Int. Symposium on Information Technologies and Communications (ITCOM)*, Oct. 2004.
- [12] M. Chen and A. Zakhor, "Rate control for video streaming over wireless," in *Proc. IEEE INFOCOM 2004*, March 2004.
- [13] H. Shen, L. Cai, and X. Shen, "Performance analysis of TFRC over wireless link with truncated link level ARQ," *IEEE Trans. Wireless Commun.*, vol. 5, no. 6, pp. 1479-1487, June 2006.
- [14] L. X. Cai, X. Shen, J. W. Mark, L. Cai and Y. Xiao, "Voice capacity analysis of WLAN with unbalanced traffic," *IEEE Trans. Veh. Technol.*, vol. 55, no. 3, pp. 752-761, May 2006.
- [15] IEEE 802.11b, Part 11, "Wireless LAN medium access control (MAC) and physical layer (PHY) specification: high-speed physical layer extension in the 2.4 GHz band," IEEE, Sept. 1999.
- [16] J. Kim, S. Kim, S. Choi, and D. Qiao, "CARA: collision-aware rate adaptation for IEEE 802.11 WLANs," in *Proc. IEEE INFOCOM 2006*, April 2006.

- [17] M. Zorzi, R. Rao, and L. Milstein, "ARQ error control for fading mobile radio channels," *IEEE Trans. Veh. Technol.*, vol. 46, no. 2, pp. 445-455, May 1997.
- [18] G. Bianchi, "Performance analysis of the IEEE 802.11 distributed coordination function," *IEEE J. Sel. Areas Commun.*, vol. 18, no. 3, pp. 535-547, March 2000.
- [19] H. Zhai, Y. Kwon, and Y. Fang, "Performance analysis of IEEE 802.11 MAC protocols in wireless LANs," *Wiley Wireless Commun. Mobile Comput.*, vol. 4, no. 8, pp. 917-931, Dec. 2004.
- [20] H. Wu, Y. Peng, K. Long, S. Cheng, and J. Ma, "Performance of reliable transport protocol over IEEE 802.11 wireless LAN: analysis and enhancement," in *Proc. IEEE INFOCOM 2002*, June 2002.
- [21] ns-2, "The Network Simulator ns-2," <http://www.isi.edu/nsnam/ns>, 2005.
- [22] IEEE 802.11a, Part 11, "Wireless LAN, medium access control (MAC) and physical layer (PHY) specifications: high-speed physical layer in the 5GHz band," *IEEE*, Sept. 1999.



Sangheon Pack (S'03-M'05) received B.S. (2000, magna cum laude) and Ph. D. (2005) degrees from Seoul National University, both in computer engineering. Since March 2007, he has been an Assistant Professor in the School of Electrical Engineering, Korea University, Korea. From July 2006 to February 2007, he was a Postdoctoral Fellow in Seoul National University, Korea. From July 2005 to June 2006, he was a Postdoctoral Fellow in the Broad-band Communications Research (BBCR) Group at University of Waterloo, Canada. From 2002 to 2005,

he was a recipient of the Korea Foundation for Advanced Studies (KFAS) Computer Science and Information Technology Scholarship. He has been also a member of Samsung Frontier Membership (SFM) from 1999. He received a student travel grant award for the IFIP Personal Wireless Conference (PWC) 2003. He was a visiting researcher in Fraunhofer FOKUS, Germany in 2003. He was listed in *Marquis Who's Who of Emerging Leader*, 2007. His research interests include mobility management, multimedia transmission, vehicular networks, and QoS provision issues in the next-generation wireless/mobile networks. He is a member of the ACM.



Xuemin (Sherman) Shen (M'97-SM'02) received the B.Sc. (1982) degree from Dalian Maritime University (China) and the M.Sc. (1987) and Ph.D. degrees (1990) from Rutgers University, New Jersey (USA), all in electrical engineering. Dr. Shen is with the Department of Electrical and Computer Engineering, University of Waterloo, Canada, where he is a Professor and the Associate Chair for Graduate Studies. Dr. Shen's research focuses on mobility and resource management in interconnected wireless/wireline networks, UWB wireless communications systems, wireless security, and ad hoc and sensor networks. He is a coauthor of two books, and has published more than 200 papers and book chapters in wireless communications and networks, control and filtering.

Dr. Shen serves as the Technical Program Committee Chair for IEEE Globecom'07, IEEE WCNC'07 Network Symposium, Qshine'05, IEEE Broadnet'05, WirelessCom'05, IFIP Networking'05, ISPAN'04, IEEE Globecom'03 Symposium on Next Generation Networks and Internet. He also serves as Associate Editor for *IEEE Transactions on Wireless Communications*, *IEEE Transactions on Vehicular Technology*, *ACM Wireless Networks*, *Computer Networks*, *Wireless Communications and Mobile Computing* (Wiley), etc., and the Guest Editor for *IEEE Journal on Selected Areas in Communications*, *IEEE Wireless Communications*, and *IEEE Communications Magazine*. Dr. Shen received the Premier's Research Excellence Award (PREA) from the Province of Ontario, Canada for demonstrated excellence of scientific and academic contributions in 2003, and the Outstanding Performance Award from the University of Waterloo, for outstanding contribution in teaching, scholarship and service in 2002. Dr. Shen is a registered Professional Engineer of Ontario, Canada.



Jon W. Mark (M'62-SM'80-F'88-LF'03) received the Ph.D. degree in electrical engineering from McMaster University, Canada in 1970. Upon graduation, he joined the Department of Electrical Engineering (now Electrical and Computer Engineering) at the University of Waterloo, became a full Professor in 1978, and served as Department Chairman from July 1984 to June 1990. In 1996, he established the Centre for Wireless Communications (CWC) at the University of Waterloo and has since been serving as the founding Director. He was on sabbatical

leave at the IBM Thomas Watson Research Center, Yorktown Heights, NY, as a Visiting Research Scientist (1976-77); at AT&T Bell Laboratories, Murray Hill, NJ, as a Resident Consultant (1982-83); at the Laboratoire MASI, Université Pierre et Marie Curie, Paris, France, as an Invited Professor (1990-91); and at the Department of Electrical Engineering, National University of Singapore, as a Visiting Professor (1994-95). His current research interests are in wireless communications and wireless/wireline interworking, particularly in the areas of resource management, mobility management, and end-to-end information delivery with QoS provisioning.

Dr. Mark is a co-author of the textbook *Wireless Communications and Networking* (Prentice-Hall, 2003). Dr. Mark is a Life Fellow of the IEEE and has served as a member of a number of editorial boards, including *IEEE Transactions on Communications*, *ACM/Baltzer Wireless Networks*, *Telecommunication Systems*, etc. He was a member of the Inter-Society Steering Committee of the *IEEE/ACM Transactions on Networking* from 1992-2003, and a member of the IEEE COMSOC Awards Committee during the period 1995-1998.



Lin Cai (S'00-M'06) received the MASc and PhD degrees (with Outstanding Achievement in Graduate Studies Award) in electrical and computer engineering from the University of Waterloo, Waterloo, Canada, in 2002 and 2005, respectively. Since July 2005, she has been an Assistant Professor in the Department of Electrical and Computer Engineering at the University of Victoria, Victoria, Canada. Her research interests span several areas in wireless communications and networking, with a focus on network protocol and architecture design supporting emerging multimedia traffic over wireless networks. She serves as the Associate Editor for *IEEE Transactions on Vehicular Technology* (2007-), *EURASIP Journal on Wireless Communications and Networking* (2006-), and *International Journal of Sensor Networks* (2006-).