

Heavy-Traffic Analysis of QoE Optimality for On-Demand Video Streams Over Fading Channels

Ping-Chun Hsieh, *Student Member, IEEE*, and I-Hong Hou, *Member, IEEE*

Abstract—This paper proposes online scheduling policies to optimize quality of experience (QoE) for video-on-demand applications in wireless networks. We consider wireless systems where an access point (AP) transmits video content to clients over fading channels. The QoE of each flow is measured by its duration of video playback interruption. We are specifically interested in systems operating in the heavy-traffic regime. We first consider a special case of ON-OFF channels plus constant-bit-rate videos and establish a scheduling policy that achieves every point in the capacity region under heavy-traffic conditions. This policy is then extended for more general fading channels and variable-bit-rate videos, and we prove that it remains optimal under some mild conditions. We then formulate a network utility maximization problem based on the QoE of each flow. We demonstrate that our policies achieve the optimal overall utility when their parameters are chosen properly. Finally, we compare our policies against three popular policies. Simulation and experimental results validate that the proposed policies indeed outperform existing policies.

Index Terms—Heavy traffic; Quality of experience; Video streaming; Wireless networks; Scheduling.

I. INTRODUCTION

In recent years, video streaming applications have been demanding more and more resource in wireless networks. According to the latest report from Cisco [2], video-centric services are projected to increase 13-fold and occupy nearly three-fourths of global mobile data traffic in the near future. However, upgrade of wireless network capacity can hardly catch up with the explosive mobile data traffic. Therefore, better scheduling algorithms are required for service providers to serve more customers without sacrificing user satisfaction.

From the perspective of service providers, scheduling policies are conventionally designed to meet the performance requirement of a general wireless network, such as average throughput, latency, delay jitter, etc. However, these statistics fail to directly characterize real user experience in enjoying video service. Hence, various research works have been carried out to study quality of experience (QoE) in video streaming applications. Much effort has been dedicated to quantifying subjective user experience and constructing analytical models based on different experiment setups, such as [3]–[5]. In general, QoE can be affected by several factors, such as playback smoothness, mean video quality, temporal variation in quality, etc. Among these elements, playback interruption has been shown to be the dominant factor of QoE performance

[4], [5]. Therefore, we define QoE by the duration of video interruption during playback process for wireless on-demand video streams.

Playback interruption of a single video stream has been studied extensively in recent literature. In [6], the probability of interruption-free video playback is analyzed for variable bit-rate video over wireless channels with variable data rate. By modelling a playback buffer as a M/D/1 queue, [7] provides a bound on interruption probability for media streams over Markovian channels. Likewise, Xu *et al.* [8] and Anttonen *et al.* [9] provide explicit results of the distribution of video interruption based on different queueing models. Similarly, [10] presents an online algorithm to adaptively control playback buffer underflow and overflow based on large deviation theory. Moreover, by applying diffusion approximation to a G/G/1 queue, Luan *et al.* [11] characterize the dynamics of a video playback buffer under a threshold-based buffer management scheme. The common focus of these works is to provide an indicator to make the best tradeoff between initial prefetching delay and playback buffer emptiness. However, these results only work for a single video stream and thus can hardly be applied to a wireless network.

Regarding scheduling for QoE of multiple video streams, [12] and [13] provide a flow-level analytical framework to study the effect of flow dynamics on playback interruption and average throughput. [14] proposes an online algorithm based on Proportional-Fair scheduling to achieve fairness among video users while maintaining required throughput. In [15], a modified version of Proportional-Fair scheduling has been presented to reduce video inter-frame delay for wireless LTE networks. To offer better average rate guarantees, Bhatia *et al.* [16] design a scheduling policy which exploits slow-fading variation of wireless channels. In a multi-cast wireless network, [17] proves by dynamic programming that a Max-Weight like policy is throughput optimal. To improve video-rate-based QoE, Li *et al.* [18] design a scheduling policy based on the head-of-line packet delay and packet deadlines. In [19], a resource allocation algorithm is proposed to maximize video-rate-based utility while maintaining fairness in term of buffer level. In [20], Li *et al.* propose a joint rate control and scheduling algorithm to optimize rate-based network utility for scalable videos in a multi-cast wireless network. In [21], Anand and de Veciana propose a scheduling policy that achieves asymptotically optimal delay-based QoE. In [22], Joseph and de Veciana consider a more comprehensive QoE metric and propose the NOVA algorithm to asymptotically optimize QoE for a wireless network. Based on a similar decomposition approach as [22], Xiao *et al.* [23] propose an

The authors are with the Department of Electrical and Computer Engineering, Texas A&M University, College Station, TX 77843, USA (e-mail: {pingchun.hsieh, ihou}@tamu.edu).

Part of this work has been presented in IEEE INFOCOM 2016 [1].

online algorithm to optimize QoE specifically for OFDMA wireless networks. One common feature of the above policies is that they aim to optimize QoE in the sense of long-term average performance, such as average video quality and average playback interruption. Moreover, [24] considers the short-term QoE performance by studying the diffusion limit. However, it only considers constant-bit-rate video streams in wireless networks where the channel qualities are static.

In this paper, we address QoE optimality and propose online scheduling policies for wireless on-demand video streams. We are particularly interested in QoE performance under heavy-traffic conditions. Different from the prior efforts of [12]–[23], this paper addresses not only long-term average performance but also *short-term QoE performance* by studying diffusion limits. Different from the study of diffusion limit for a single video stream without scheduling in [11], we study diffusion limits for a network of multiple video streams and propose online scheduling policies. In [11], diffusion approximation for a G/G/1 queue is directly applicable since the arrival and departure processes of a playback buffer are given and completely independent of the buffer management scheme. By contrast, for a network of video streams, the arrival process of each video buffer is controlled by the scheduling policy, and hence it is not immediately clear how to apply diffusion limits. Despite this challenge, we are able to characterize the capacity region for QoE and show that the proposed policies achieve the whole capacity region based on diffusion limits. Instead of assuming static channel qualities and constant-bit-rate videos as in [24], we study the dynamic behavior of playback process of variable-bit-rate videos over time-varying channels. This paper can be summarized as follows:

- For ease of presentation, we start from a special case of constant-bit-rate videos over ON-OFF channels. We first consider long-term average playback interruption by studying the stability region and providing a polynomial-time algorithm for checking the stabilizability when channels are independent (Section III).
- We study short-term QoE performance of playback interruption by applying diffusion limits. Specifically, we start by deriving a lower bound for total playback interruption for all scheduling policies and providing necessary conditions for the capacity region for QoE (Section IV-C). We then propose an online scheduling policy and explicitly characterize the distribution of playback interruption at any given time under the proposed policy (Section IV-D). We thereby show that the proposed policy achieves every interior point in the capacity region for QoE.
- The policy and the heavy-traffic analysis are then generalized for general i.i.d. fading channels and variable-bit-rate videos. The proposed policies are proved to remain optimal under some mild assumptions.
- We formulate a network utility maximization problem based on the QoE of each client. We show that our policy can achieve the optimal network utility by selecting proper parameters.
- We compare the proposed policies against three popular policies and show by simulation that our policy surpasses

other three policies by a large margin in short-term QoE, despite that the long-term average duration of video interruption approaches zero asymptotically for all these policies. Moreover, we also implement the proposed policy on a software-defined wireless testbed and demonstrate the performance via real video streaming applications.

Through analysis and extensive evaluation of the proposed policies, we thereby demonstrate the essential difference between QoE and the conventional QoS metrics for wireless networks.

The rest of the paper is organized as follows. Section II describes the network model for analyzing QoE of on-demand video streams. Section III discusses the stability region and an algorithm for checking stabilizability for ON-OFF channels. In Section IV, we present an online scheduling policy for ON-OFF channels plus constant-bit-rate videos and prove that it is optimal in heavy traffic. We then extend the policy for fading channels and variable-bit-rate videos in Section V. In Section VI, we formulate a network utility maximization problem based on QoE. Simulation and experimental results of the proposed policies as well as the counterparts are shown in Section VII and VIII, respectively. Finally, Section IX concludes the paper.

II. SYSTEM MODEL

We consider a time-slotted wireless network consisting of a wireless access point (AP) and a group of N mobile clients denoted by $S_{tot} = \{1, 2, \dots, N\}$. Each client is downloading an on-demand video which has been pre-stored by video service providers. The video content is partitioned into packets and streamed to clients via the AP and the wireless links. On the AP's side, we assume that the AP always has packets at hand for transmission to each video client. In other words, the throughput for the AP to acquire video content from video providers is assumed to be much larger than that between the AP and the mobile clients. We also assume that there is no network coding mechanism involved in the system. Thus, in each time slot, the AP can transmit data to at most one client.

In a wireless network, the quality of a wireless channel usually changes with time. To capture the variation of wireless channels, we model the wireless link of each client n as a discrete-time random process $r_n(t)$, which is i.i.d. across time and takes only non-negative integer values in a finite *data rate space* denoted by $\mathcal{R}$. If the AP schedules a transmission for client n at time t , it can deliver exactly $r_n(t)$ bits. For example, the IEEE 802.11a standard has a maximum physical data rate of 54 Mbit/s, and the data rate can also be adaptively reduced by applying different modulation and coding, depending on channel conditions. In this model, we make no assumption about the relationship between different wireless links. Therefore, unless stated otherwise, the channels of different clients are *not* required to be independent.

On the mobile clients' side, the received packets are first decoded and queued in a playback buffer, whose size is assumed to be infinite. For each client n , let $A_n(t)$ denote the total amount of received video content in bits until time t , and $B_n(t)$ be the amount of video content in bits stored in

the playback buffer at time t . Next, we consider the playback process of *variable-bit-rate* videos. Specifically, each client n plays one video frame every k_n slots, and different frames of the same video may contain different number of bits. Suppose for client n , the j -th frame contains $F_n(j)$ bits for every $j \geq 1$. We assume that $F_n(j)$ are i.i.d. across j with mean q_n^* and variance $\sigma_{q,n}^2$. We define $q_n := q_n^*/k_n$ to be the average video playback rate per slot, which reflects the desired video resolution. Moreover, we assume that $F_n(j)$ is upper bounded by $q_{n,\max}$ for every frame j . Let $C_n(J)$ be the sum of the frame size up to the J -th frame, i.e. $C_n(J) := \sum_{j=1}^J F_n(j)$. Let $S_n(t)$ be the total number of frames that the client n plays up to t . Therefore, the total bits played up to time t is $C_n(S_n(t))$. At time slot t , the number of bits in the playback buffer of client n is

$$B_n(t) = A_n(t) - C_n(S_n(t)), \quad (1)$$

where we assume that $B_n(0) = 0$ for every client.

When the client is about to play a new video frame from the playback buffer, playback interruption might occur if there is no enough video content in the playback buffer. To check this condition, it is equivalent to check if the buffer $B_n(t)$ becomes negative after the video frame about to play is taken out of the buffer. Let $D_n(t)$ be the total number of slots in which video is interrupted by time t . Given $D_n(t-1)$, $\lfloor \frac{t-D_n(t-1)}{k_n} \rfloor$ is the total number of video frames that client n would play up to time t if video interruption does not happen at time t . Hence, $C_n(\lfloor \frac{t-D_n(t-1)}{k_n} \rfloor)$ represents the total number of bits played up to t if video interruption does not happen at time t . Then, $D_n(t)$ can be updated recursively by

$$D_n(t) = \begin{cases} D_n(t-1) + 1, & \text{if } A_n(t) - C_n(\lfloor \frac{t-D_n(t-1)}{k_n} \rfloor) < 0 \\ D_n(t-1), & \text{otherwise} \end{cases} \quad (2)$$

From (2), we know that in each slot, $D_n(t)$ either stays unchanged or increases by 1. Define $B_n^*(t) := q_{n,\max} \lfloor \frac{B_n(t)}{q_{n,\max}} \rfloor$ to be the quantized version of $B_n(t)$ and let $e_n(t) := B_n^*(t) - B_n(t)$ be the quantization error of $B_n(t)$ with respect to $q_{n,\max}$. Note that $e_n(t)$ is bounded, i.e. $|e_n(t)| < q_{n,\max}$, for all $t \geq 0$. After the above manipulation, we know that $B_n^*(t) = 0$ if $D_n(t+1) - D_n(t) = 1$. Therefore, we have

$$B_n^*(t) = A_n(t) - C_n(S_n(t)) + e_n(t) \quad (3)$$

$$= (A_n(t) - q_n t) + q_n D_n(t) \quad (4)$$

$$- [C_n(S_n(t)) - q_n(t - D_n(t)) - e_n(t)]. \quad (5)$$

Define

$$Y_n(t) := C_n(S_n(t)) - q_n(t - D_n(t)) - e_n(t). \quad (6)$$

Since $t - D_n(t)$ is the total playback time with no interruption, then on average client n should already play $q_n(t - D_n(t))$ bits by time t . Since $C_n(S_n(t))$ is the total number of bits played up to time t , $Y_n(t)$ therefore reflects whether the amount of played video content matches the average playback rate of the variable-bit-rate videos. We also define that

$$X_n(t) := A_n(t) - q_n t. \quad (7)$$

Since on average client n plays q_n bits per time slot, then $X_n(t)$ reflects whether the amount of received data matches the average video playback rate. Now, we can summarize the basic properties as follows.

$$B_n^*(t) = X_n(t) - Y_n(t) + q_n D_n(t) \geq 0 \quad (8)$$

$$[D_n(t+1) - D_n(t)] \in \{0, 1\}, \quad D_n(0) = 0, \quad (9)$$

$$B_n^*(t)[D_n(t+1) - D_n(t)] = 0 \quad (10)$$

Based on (8)-(10), we can further connect $D_n(t)$ with $X_n(t)$ as follows.

Theorem 1 *Given $X_n(t)$ and $Y_n(t)$, there exists a unique pair of $B_n^*(t)$ and $D_n(t)$ that satisfies (8)-(10). Moreover, we have*

$$D_n(t) = \sup_{0 \leq \tau \leq t} (\max\{0, -\frac{X_n(\tau) - Y_n(\tau)}{q_n}\}). \quad (11)$$

Proof This is a direct result of Theorem 6.1 in [25]. $\square$

Remark 1 We know that $X_n(t)$ reflects whether the total amount of received data $A_n(t)$ matches the total number of played bits $q_n t$. Moreover, $Y_n(t)$ captures the fluctuation in video frame size. $X_n(t)$ and $Y_n(t)$ are sufficient to determine the buffer status and video interruption. Therefore, it is not surprising that there exists a unique pair of $B_n^*(t)$ and $D_n(t)$ that satisfies (8)-(11) as stated in Theorem 1.

Remark 2 To intuitively understand (11), consider an example where a client n receives a_n bits from the AP in every time slot and the video has a constant bit rate. Then, $X_n(t) = (a_n - q_n)t$, and $|Y_n(t)| \leq q_n k_n + |e_n(t)|$ is bounded.

- If $a_n > q_n$, then $X_n(t)$ grows linearly with t and by (11) we know $D_n(t)$ remains 0 for all t . In other words, there is no video interruption at all if the amount of received data per slot is always greater than the playback rate.
- If $a_n < q_n$, then $X_n(t)$ is always non-positive and decreases linearly with t . By (11) we know $D_n(t)$ grows almost linearly with t with some minor fluctuation due to the bounded term $Y_n(t)$. This corresponds to the fact that the video gets continually interrupted when the amount of received data per slot is always less than the required playback rate.
- If $a_n = q_n$, then both $X_n(t)$ and $Y_n(t)$ are 0 for all t . Thus, $D_n(t) = 0$ for all t .

In this paper, QoE of each video stream is measured by its total duration of playback interruption. One usual way to assess QoE of a wireless network is through long-term average performance which is formally defined as follows.

Definition 1 *A video streaming system is said to be stabilizable if there exists a scheduling policy η such that $\limsup_{t \rightarrow \infty} \frac{D_n(t)}{t} = 0$ for all n , almost surely. Moreover, η is a stabilizing scheduling policy for QoE.* $\square$

In other words, a wireless video-streaming system is stabilizable if the total duration of video interruption grows sublinearly after some finite time. We first consider the long-term average behavior of $Y_n(t)$. Since the frame sizes are i.i.d. with mean q_n^* , then by the Strong Law of Large Numbers for i.i.d. random variables, it can be easily shown that

TABLE I
MAIN NOTATIONS USED IN THE PAPER.

Notation	Description
q_n^*	average frame size of client n
k_n	interval between two consecutive frames (if no interruption)
q_n	q_n^*/k_n , i.e. video playback rate of client n
$A_n(t)$	total number of bits received by client n up to time t
$S_n(t)$	total number of frames that client n plays up to t
$C_n(J)$	sum of the frame size of client n up to frame J
$X_n(t)$	$A_n(t) - q_n t$ (reflects if the arrival rate matches q_n)
$Y_n(t)$	$C_n(S_n(t)) - q_n(t - D_n(t)) - e_n(t)$, where $e_n(t)$ denotes a small bounded error
$D_n(t)$	total amount of video interruption seen by client n up to t
$\hat{X}_n(t), \hat{Y}_n(t)$	diffusion limits of $X_n(t)$ and $Y_n(t)$
$\hat{D}_n(t)$	diffusion limits of $D_n(t)$
$X(t)$	sum of $X_n(t)$ of all client n
$\hat{X}(t)$	diffusion limit of $X(t)$
$\hat{D}(t)$	$\sup_{0 \leq \tau \leq t} (\max\{0, -\hat{X}(\tau)\})$
$r_n(t)$	data rate of client n at time t
$R(t, S)$	$\max\{r_n(t) : n \in S\}$
$R(t)$	highest data rate among all the clients at time t
$Z_n(t)$	$C_n(\lfloor \frac{t}{k_n} \rfloor) - q_n t$
$Z(t)$	sum of $Z_n(t)$ of all client n
$\hat{Z}_n(t), \hat{Z}(t)$	diffusion limits of $Z_n(t)$ and $Z(t)$

$\lim_{t \rightarrow \infty} \frac{Y_n(t)}{t} = 0$, almost surely, regardless of the scheduling policy. By (8), this implies that the fluctuation in video frame size does not affect the long-term average video interruption. Then, it can be easily shown that $\limsup_{t \rightarrow \infty} \frac{D_n(t)}{t} = 0$ if and only if the long-term average throughput of each client is at least q_n [24]. If $\liminf_{t \rightarrow \infty} \frac{A_n(t)}{t} < q_n$, then $X_n(t)$ will go to negative infinity as $t \rightarrow \infty$, almost surely. By (8), since $B_n^*(t)$ is always nonnegative, $\liminf_{t \rightarrow \infty} \frac{A_n(t)}{t} < q_n$ implies that $D_n(t)$ goes to infinity as $t \rightarrow \infty$, almost surely. Therefore, studying whether a system is stabilizable is equivalent to studying the capacity region of achievable throughput. However, this definition fails to characterize the behavior of the system in the heavy-traffic regime.

To fully characterize the growth of playback interruption with time, we study the dynamic behavior by using *diffusion limits* in the following parts of the paper. Moreover, in Section VII, we will compare the proposed policy with other popular scheduling policies that are all stabilizing for QoE but are rather different in short-term performance.

To study the behavior of video interruption in the heavy-traffic regime, we consider the diffusion limit of $D_n(t)$, which is defined as

$$\hat{D}_n(t) := \lim_{k \rightarrow \infty} \frac{D_n(kt)}{\sqrt{k}}, \quad 0 \leq t \leq 1. \quad (12)$$

Similarly, we define

$$\hat{X}_n(t) := \lim_{k \rightarrow \infty} \frac{X_n(kt)}{\sqrt{k}}, \quad (13)$$

$$\hat{Y}_n(t) := \lim_{k \rightarrow \infty} \frac{Y_n(kt)}{\sqrt{k}}, \quad (14)$$

$$\hat{B}_n^*(t) := \lim_{k \rightarrow \infty} \frac{B_n^*(kt)}{\sqrt{k}}. \quad (15)$$

Given the properties in (7)–(10), we then have the following useful results.

Theorem 2 Given $\hat{X}_n(t)$ and $\hat{Y}_n(t)$, there exists a unique pair of $(\hat{B}_n^*(t), \hat{D}_n(t))$ that satisfies

$$\hat{B}_n^*(t) = \hat{X}_n(t) - \hat{Y}_n(t) + q_n \hat{D}_n(t) \geq 0 \quad (16)$$

$$\frac{d\hat{D}_n(t)}{dt} \geq 0, \quad \hat{D}_n(0) = 0 \quad (17)$$

$$\hat{B}_n^*(t) \frac{d\hat{D}_n(t)}{dt} = 0. \quad (18)$$

Moreover, $\hat{D}_n(t)$ can be expressed as

$$\hat{D}_n(t) = \sup_{0 \leq \tau \leq t} (\max\{0, -\frac{\hat{X}_n(t) - \hat{Y}_n(t)}{q_n}\}) \quad (19)$$

Proof This is a direct result of Theorem 1 in [24]. $\square$

Remark 3 (19) can be viewed as the diffusion limit version of (11).

Theorem 2 suggests a general recipe on how to study video interruption in the diffusion limit:

- Characterize $\hat{X}_n(t)$ based on the channel model and the scheduling policy.
- Characterize $\hat{Y}_n(t)$ based on the dynamics of video bit rate.
- Combine $\hat{X}_n(t)$ and $\hat{Y}_n(t)$ to derive $\hat{D}_n(t)$ based on (19).

For ease of presentation, we start from a special case of constant-bit-rate videos over ON-OFF channels in Section III and IV, and then extend the analysis for fading channels and variable-bit-rate videos in Section V.

In order to distinguish the analysis on $\limsup_{t \rightarrow \infty} \frac{D_n(t)}{t}$ and that on $\hat{D}_n(t)$, we use *stability region* to denote the set of stabilizable systems, and *capacity region* to denote the set of achievable vectors of $\hat{D}_n(t)$. A more formal definition of capacity region is introduced in Section IV.

For convenience, we summarize the main notations used in this paper in Table I.

III. STABILITY REGION FOR ON-OFF CHANNELS

We first consider a special case of ON-OFF channels, where transmission rate of each client can only be either zero or a positive value r^* , and therefore $\mathcal{R} = \{0, r^*\}$.

Recall that the AP can transmit data to at most one client in each time slot. In this case, the stability region for ON-OFF channels has been shown to be associated with a set of necessary and sufficient conditions [26], [27]. We summarize the results as follows.

Lemma 1 [26, Theorem 1] Let W_n be the event that client n has an ON channel, i.e., $r_n(t) = r^*$. A video streaming system with ON-OFF channels is stabilizable if and only if the video playback rates $\{q_n\}$ satisfy the following equations:

$$\Pr \left[\bigcup_{n \in S} W_n \right] \geq \frac{1}{r^*} \sum_{n \in S} q_n, \quad \forall S \subseteq S_{tot}. \quad (20)$$

Remark 4 For a given subset S , (20) indicates that the total demand of the subset S should not exceed the maximum total channel resource of the subset S .

The above condition requires checking (20) for all subsets, which can be intractable. However, for the special case where the channel conditions of different clients are independent, we can derive a polynomial-time algorithm to check whether a system is stabilizable. The algorithm is described in the following theorem.

Theorem 3 *Let p_n be the probability that client n has an ON channel, and $p_n > 0, \forall n$. Suppose that the clients are sorted based on $\frac{q_n}{p_n}$ in descending order, i.e. $\frac{q_1}{p_1} \geq \frac{q_2}{p_2} \geq \dots \geq \frac{q_N}{p_N}$. Denote by S_k the subset $\{1, \dots, k\}$ of all clients. Then, the system is stabilizable if and only if*

$$1 - \prod_{n \in S_k} (1 - p_n) \geq \frac{1}{r^*} \sum_{n \in S_k} q_n, \quad 1 \leq k \leq N. \quad (21)$$

Moreover, the complexity of checking this condition is $O(N \log N)$. $\square$

Proof The proof can be found in Theorem 2 of [1]. $\square$

IV. HEAVY-TRAFFIC ANALYSIS FOR ON-OFF CHANNELS AND CONSTANT-BIT-RATE VIDEOS

We are particularly interested in the situation where the set of video playback rates $\{q_n\}$ is on the boundary of the stability region, that is, under the *heavy-traffic* condition.

A. Heavy-Traffic Conditions for ON-OFF Channels

Recall that W_n denotes the event that client n has ON channel. In this section, we assume that,

$$\Pr \left[\bigcup_{n=1}^N W_n \right] = \frac{1}{r^*} \sum_{n=1}^N q_n, \quad (22)$$

while for any subset $S \subset \{1, \dots, N\}$,

$$\Pr \left[\bigcup_{n \in S} W_n \right] > \frac{1}{r^*} \sum_{n \in S} q_n. \quad (23)$$

The constraint (23) corresponds to the *complete resource pooling* condition in [28]. The complete resource pooling condition guarantees that there is enough overlap in the channel resources of different clients. This technical condition enables us to characterize the system using one-dimensional Brownian motion as in [28].

B. Constant-Bit-Rate Videos

For constant-bit-rate videos, all the frames of the same video have exactly the same size q_n^* , for each client n . From Remark 2, we know that $|Y_n(t)| \leq q_n k_n + |e_n(t)|$ is bounded, regardless of the scheduling policy. Therefore, by the definition of diffusion limit, we have $\hat{Y}_n(t) = 0$, for all t , under any scheduling policy. By Theorem 2, we have

$$\hat{D}_n(t) = \sup_{0 \leq \tau \leq t} (\max\{0, -\frac{\hat{X}_n(t)}{q_n}\}). \quad (24)$$

This implies that for constant-bit-rate video streams, $\hat{X}_n(t)$ fully characterizes the behavior of video playback interruption in the diffusion limit. Therefore, we can focus on characterizing $\hat{X}_n(t)$ in the rest of this section.

C. A Lower-Bound of Capacity Region for QoE

We derive fundamental properties of $D_n(t)$ with ON-OFF channels under the heavy-traffic conditions. Let us define a random process

$$X(t) := \sum_{n=1}^N X_n(t) = \sum_{n=1}^N (A_n(t) - q_n t). \quad (25)$$

Let $\Delta X(t+1) := X(t+1) - X(t)$ be the amount of change in $X(t)$, for all $t \geq 0$. Regardless of the scheduling policy, the AP can deliver exactly r^* bits to some client n if at least one client in S_{tot} has an ON channel. Let $\gamma := \Pr \left[\bigcup_{n=1}^N W_n \right]$ be the probability of the event that at least one client has an ON channel. Then, in each time slot t , $\sum_{n=1}^N A_n(t)$ increases by r^* with probability γ and stays the same with probability $1 - \gamma$, regardless of the scheduling policy. Therefore, we have

$$\Delta X(t+1) = \begin{cases} -\sum_{n=1}^N q_n, & \text{with probability } 1 - \gamma \\ r^* - \sum_{n=1}^N q_n, & \text{with probability } \gamma \end{cases} \quad (26)$$

Since the channels are i.i.d. across time, the equations in (26) hold regardless of time and thus $\Delta X(t)$ is i.i.d. across all time slots. Due to the heavy-traffic assumption given by (22), we further have

$$E[\Delta X(t)] = \gamma(r^* - \sum_{n=1}^N q_n) + (1 - \gamma)(-\sum_{n=1}^N q_n) = 0$$

$$\text{Var}[\Delta X(t)] = \gamma(r^* - \sum_{n=1}^N q_n)^2 + (1 - \gamma)(\sum_{n=1}^N q_n)^2 = \gamma(1 - \gamma)(r^*)^2.$$

By the *functional central limit theorem* for i.i.d. random variables [25], we have the following important properties regarding the diffusion limit of $X(t)$.

Theorem 4 *Let $\hat{X}(t) := \lim_{k \rightarrow \infty} \frac{X(k t)}{\sqrt{k}}$. Then $\hat{X}(t)$ is a driftless Brownian motion with variance σ_x^2 , where $\sigma_x = r^* \sqrt{\gamma(1 - \gamma)}$. Moreover, given $\hat{X}(\tau)$, for any $\tau, t \geq 0$ with $\tau + t \leq 1$, $\hat{X}(\tau + t) - \hat{X}(\tau)$ is a Gaussian random variable with zero mean and variance $\sigma_x^2 t$. $\square$*

Remark 5 By Theorem 4, we are able to fully characterize the distribution of $\hat{X}(t)$ while the original process $X(t)$ can be difficult to analyze. This manifests the benefit of taking the diffusion limit of the original process.

Similar to (24), we define

$$\hat{D}(t) := \sup_{0 \leq \tau \leq t} (\max\{0, -\hat{X}(\tau)\}) \quad (27)$$

Since $\hat{X}(t)$ is a Brownian motion, we can thereby derive the distribution and important statistics of $\hat{D}(t)$ based on the following lemma.

Lemma 2 [29, Section 1.6] *Let $\Phi(x)$ be the cumulative distribution function (CDF) of a standard Gaussian random variable with zero mean and unit variance. The CDF of $\hat{D}(t)$ is given by*

$$\Pr[\hat{D}(t) \leq x] = \Phi\left(\frac{x}{\sqrt{\sigma_x^2 t}}\right) - \Phi\left(\frac{-x}{\sqrt{\sigma_x^2 t}}\right)$$

for all $x \geq 0$, $t \geq 0$. The expected value of $\hat{D}(t)$ is given by

$$E[\hat{D}(t)] = \int_0^\infty x \sqrt{\frac{2}{\pi \sigma_x^2 t}} \exp\left(-\frac{x^2}{2\sigma_x^2 t}\right) dx = \sqrt{\frac{2t\sigma_x^2}{\pi}}. \quad \square$$

Given the characteristics of $\hat{D}(t)$, we obtain a lower bound for dynamics of video interruption seen by the clients. We first introduce the concept of stochastic ordering as follows [25].

Definition 2 Let $\hat{D}^{\eta_1}(t)$ and $\hat{D}^{\eta_2}(t)$ be two real-valued random processes under policies η_1 and η_2 , respectively. We say that $\hat{D}^{\eta_1}(t) \leq_{st} \hat{D}^{\eta_2}(t)$ if

$$Pr[\hat{D}^{\eta_1}(t) \geq x] \leq Pr[\hat{D}^{\eta_2}(t) \geq x], \quad (28)$$

for all $x \in \mathbb{R}$ and for any $t \in [0, 1]$. $\square$

Then, we further build the relationship between $\hat{D}(t)$ and $\hat{D}_n(t)$ by using of $\hat{X}(t)$ and $\hat{X}_n(t)$:

$$\hat{D}(t) = \sup_{0 \leq \tau \leq t} (\max\{0, -\hat{X}(t)\}) \quad (29)$$

$$= \sup_{0 \leq \tau \leq t} (\max\{0, -\sum_{n=1}^N \hat{X}_n(\tau)\}) \quad (30)$$

$$\leq_{st} \sum_{n=1}^N \sup_{0 \leq \tau \leq t} (\max\{0, -\hat{X}_n(t)\}) = \sum_{n=1}^N q_n \hat{D}_n(t). \quad (31)$$

Motivated by (29)-(31), we define the *capacity region* for QoE in terms of the diffusion limits $\hat{D}(t)$ and $\hat{D}_n(t)$ as follows.

Definition 3 A N -tuple vector $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]$ is said to be feasible if there exists a scheduling policy such that

$$\hat{D}_n(t) \leq_{st} \frac{\lambda_n}{q_n} \hat{D}(t), \quad n = 1, 2, \dots, N. \quad (32)$$

Then, the capacity region for QoE, denoted by Λ , is defined as the set of all feasible vectors λ . $\square$

Remark 6 Note that the capacity region for QoE is defined in terms of the diffusion limits of video interruption time, instead of the arrival rate region considered in many studies on long-term average throughput (such as [30]). Regarding the capacity analysis of long-term average throughput in classical queueing theory, it has been widely known that the capacity region can be characterized by using stationary randomized policies. By contrast, in the capacity analysis for QoE studied in this paper, it is not immediately clear how to characterize the capacity region based on diffusion limits or how to achieve the whole capacity region. This also manifests the difference between our study on QoE and the conventional studies on long-term average throughput.

Theorem 5 A feasible vector $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]$ with $\lambda_n \geq 0$, for all n , must satisfy $\sum_{n=1}^N \lambda_n \geq 1$. $\square$

Proof By Definition 3, if λ is feasible, then $q_n \hat{D}_n(t) \leq_{st} \lambda_n \hat{D}(t)$, for every client n . Therefore, $\sum_{n=1}^N q_n \hat{D}_n(t) \leq_{st} \sum_{n=1}^N \lambda_n \hat{D}(t)$. By (29)-(31), we have $\sum_{n=1}^N \lambda_n \geq 1$. $\square$

D. Scheduling Policy

Now, we introduce a scheduling policy for constant-bit-rate videos over ON-OFF channels and show that it achieves every point in the interior of the capacity region for QoE.

Joint Channel-Deficit Policy (JCD):

In each time slot, the AP schedules the client n with the smallest value of $w_n X_n(t)$ among those clients with $r_n(t) = r^*$, where w_n is a predetermined weight factor. $\square$

To prove that JCD policy achieves every point in the interior of the capacity region for QoE, we first establish the *state space collapse* property to characterize the diffusion limit $\hat{X}_n(t)$ of each individual client.

Theorem 6 Let w_n be the weight for client n which is predetermined by the AP. For any pair of clients n, m in S_{tot} , we have $w_n \hat{X}_n(t) = w_m \hat{X}_m(t)$. Moreover, we can obtain that

$$\hat{X}_n(t) = \frac{\frac{1}{w_n}}{\sum_{m=1}^N \frac{1}{w_m}} \hat{X}(t) = \beta_n \hat{X}(t), \quad (33)$$

where $\beta_n := \frac{\frac{1}{w_n}}{\sum_{m=1}^N \frac{1}{w_m}}$ and $\sum_{n=1}^N \beta_n = 1$. $\square$

Proof To show state-space collapse, we start from a fluid system induced by $X_n(t)$. Next, we consider a Lyapunov function and show the random process of the fluid system is positive recurrent. The complete proof can be found in Theorem 5 of [1]. $\square$

Based on Theorem 6, we show that the JCD policy achieves every point in the capacity region by choosing proper $\{w_n\}$.

Theorem 7 Given any vector $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]$ which satisfies $\lambda_n > 0$, $\forall n$, and $\sum_{n=1}^N \lambda_n \geq 1$, JCD policy achieves $\hat{D}_n(t) = \frac{\beta_n}{q_n} \hat{D}(t) \leq_{st} \frac{\lambda_n}{q_n} \hat{D}(t)$ with $\beta_n := \frac{\frac{1}{w_n}}{\sum_{m=1}^N \frac{1}{w_m}}$ and $\sum_{n=1}^N \beta_n = 1$, and thus the vector λ is feasible. Moreover,

$$E[\hat{D}_n(t)] = \sqrt{\frac{2t\sigma_x^2}{\pi}} \frac{\beta_n}{q_n} = \sqrt{\frac{2t\sigma_x^2}{\pi}} \frac{\frac{1}{q_n w_n}}{\sum_{m=1}^N \frac{1}{w_m}}. \quad (34)$$

Proof From (24) and (33), we know that

$$\hat{D}_n(t) = \sup_{0 \leq \tau \leq t} (\max\{0, -\frac{\hat{X}_n(t)}{q_n}\}) \quad (35)$$

$$= \sup_{0 \leq \tau \leq t} (\max\{0, -\frac{\beta_n \hat{X}(t)}{q_n}\}) = \frac{\beta_n}{q_n} \hat{D}(t). \quad (36)$$

By assigning $w_n = \frac{1}{\lambda_n}$ for all n , we have $\beta_n = \frac{\frac{1}{w_n}}{\sum_{m=1}^N \frac{1}{w_m}} = \frac{\lambda_n}{\sum_{m=1}^N \lambda_m} \leq \lambda_n$, where the last inequality holds since $\sum_{m=1}^N \lambda_m \geq 1$. Therefore, we conclude that $\hat{D}_n(t) = \frac{\beta_n}{q_n} \hat{D}(t) \leq_{st} \frac{\lambda_n}{q_n} \hat{D}(t)$. Moreover, (34) follows directly from (36) and Lemma 2. $\square$

Remark 7 From an engineering point of view, by choosing w_n for each client, we can control the total playback interruption seen by each client and hence differentiate levels of service among all clients. In real applications, $\{w_n\}$ can be determined by a proper pricing scheme given by service providers. One simple example is paid VIP membership:

choose a proper pair $\mu_1, \mu_2 > 0$ with $\mu_1 > \mu_2$, assign $w_n = \mu_1$ if client n is a VIP member; otherwise, assign $w_n = \mu_2$. This scheme enables differentiated service in QoE.

Here, we do not consider the boundary points which have $\lambda_n = 0$ for some n . In practice, we can get as close as possible to the boundary points by technically assigning extremely large w_n to our policy. Theorem 7 also characterizes a sufficient condition for the capacity region.

Theorem 8 *Given a vector $[\lambda_1, \lambda_2, \dots]$ with $\lambda_n > 0, \forall n$, the vector is feasible if and only if $\sum_{n=1}^N \lambda_n \geq 1$. $\square$*

V. HEAVY-TRAFFIC ANALYSIS FOR GENERAL FADING CHANNELS AND VARIABLE-BIT-RATE VIDEOS

In this section, we further relax the assumptions that channels are ON-OFF and videos have constant bit rates in every frame, and study the playback process with general i.i.d. fading channels and variable-bit-rate videos.

A. General Fading Channels and Heavy-Traffic Conditions

We consider general i.i.d. fading channels where the data rate space $\mathcal{R}$ can consist of any number of different rates, as described in Section II. Recall that the AP can transmit data to at most one client in each time slot. Unlike the case of ON-OFF channels, the stability region cannot be determined by a simple set of conditions as those in Lemma 1. Instead, we impose the following conditions to simplify the analysis.

Recall that q_n^* and q_n are the mean frame size and the mean video consumption rate in bits per slot, respectively. Let $R(t, S) := \max\{r_n(t) : n \in S\}$ and $R(t)$ be the shorthand for $R(t, S_{tot})$. In other words, $R(t)$ represents the highest data rate at time t among all the clients. We assume that

$$E[R(t)] = \sum_{n=1}^N q_n, \quad (37)$$

and

$$E[R(t) \cdot \mathbb{I}\{R(t, S) = R(t)\}] > \sum_{n \in S} q_n, \quad (38)$$

for all $S \subsetneq S_{tot}$, where $\mathbb{I}\{\cdot\}$ is an indicator function. (37) serves as the heavy-traffic condition for general fading channels, i.e. the condition where the maximum system-wide channel resource exactly matches the total video playback rate. Besides, similar to (23), (38) corresponds to the complete resource pooling condition for general fading channels. We also note that these conditions reduce to (22) and (23) when $\mathcal{R} = \{0, r^*\}$, i.e. ON-OFF channels. It is easy to check these two conditions are sufficient for a system to be stabilizable. Further, it characterizes a portion of the boundary of the stability region, as it is not possible to increase q_n for any client n without making the system unstabilizable.

Similar to Section IV, we define $X(t) := \sum_{n=1}^N X_n(t) = \sum_{n=1}^N (A_n(t) - q_n t)$ and $\Delta X(t) := X(t) - X(t-1)$. Under the heavy-traffic condition given by (37), $\forall t > 0$ we have

$$E[\Delta X(t)] = \sum_{n=1}^N E[A_n(t) - A_n(t-1)] - \sum_{n=1}^N q_n \quad (39)$$

$$\leq E[R(t)] - \sum_{n=1}^N q_n \leq 0, \quad (40)$$

regardless of the scheduling policy. To obtain a lower bound of capacity region as in Section IV, we first consider a special class of policies, denoted by Π^* , which only schedule clients with the largest $r_n(t)$ at any time $t > 0$. This class of policies still need to determine which client to schedule when there are multiple clients with the same largest $r_n(t)$. Let $X_\pi(t)$ denote the random process of $X(t)$ under a scheduling policy $\pi \in \Pi^*$. Then, given any scheduling policy $\pi \in \Pi^*$,

$$X_\pi(t) \geq X_\eta(t), \quad \forall t \geq 0, \quad (41)$$

for every sample path, for any scheduling policy η . Let $\Delta X_\pi(t+1) := X_\pi(t+1) - X_\pi(t)$. Since $R(t)$ is i.i.d. across all time slots, $\Delta X_\pi(t)$ is also i.i.d. for all $t > 0$. Moreover,

$$E[\Delta X_\pi(t)] = E[R(t)] - \sum_{n=1}^N q_n = 0 \quad (42)$$

$$\text{Var}[\Delta X_\pi(t)] = \text{Var}[R(t)] = E[(R(t))^2] - \left(\sum_{n=1}^N q_n\right)^2 \quad (43)$$

Then, by the functional central limit theorem for i.i.d. random variables [25], we summarize the fundamental properties of the diffusion limit of $X_\pi(t)$ as follows.

Theorem 9 *For any scheduling policy $\pi \in \Pi^*$, define $\hat{X}_\pi(t) := \lim_{k \rightarrow \infty} \frac{X_\pi(kt)}{\sqrt{k}}$ and $\sigma^2 := E[(R(t))^2] - \left(\sum_{n=1}^N q_n\right)^2$. Then, $\hat{X}_\pi(t)$ is a driftless Brownian motion with variance σ^2 . Furthermore, given $\hat{X}_\pi(\tau)$, for any $\tau, t \geq 0$ with $\tau + t \leq 1$, $\hat{X}_\pi(\tau + t) - \hat{X}_\pi(\tau)$ is a Gaussian random variable with zero mean and variance $\sigma^2 t$. $\square$*

By Theorem 9, we know that $\hat{X}_\pi(t)$ has the same behavior for all scheduling policy $\pi \in \Pi^*$. For simplicity, we use $\hat{X}^*(t)$ to denote the diffusion limit of $X_\pi(t)$ for any policy $\pi \in \Pi^*$.

B. Variable-Bit-Rate Videos

Next, we study the dynamics of variable-bit-rate videos to characterize the behavior of $\hat{Y}_n(t)$. Define

$$Z_n(t) := C_n(\lfloor \frac{t}{k_n} \rfloor) - q_n t. \quad (44)$$

Note that $C_n(\lfloor \frac{t}{k_n} \rfloor)$ is the total number of frames played up to t if $D_n(t) = 0$. Similar to $Y_n(t)$, $Z_n(t)$ aims to capture the dynamics of video frame size but without taking video interruption into account. We consider $Z_n(t)$ for two reasons:

- $Z_n(t)$ and $Y_n(t)$ behave the same in the diffusion limit. This will be shown later in Lemma 3.

- Using $Z_n(t)$ instead of $Y_n(t)$ greatly simplifies the design and implementation of scheduling policy. This will become more clear in Section V-D (see Remark 9).

Consider the diffusion limits of $Z_n(t)$ and $C_n(t)$ as $\hat{Z}_n(t) := \lim_{k \rightarrow \infty} \frac{Z_n(kt)}{\sqrt{k}}$ and $\hat{C}_n(t) := \lim_{k \rightarrow \infty} \frac{C_n(kt) - q_n kt}{\sqrt{k}}$. We state a useful property in the following lemma.

Lemma 3 For any client n , at any t , we have

$$\hat{Y}_n(t) = \hat{Z}_n(t) = \hat{C}_n\left(\frac{t}{k_n}\right), \quad (45)$$

where $\hat{C}_n(t)$ is a driftless Brownian motion with variance $\sigma_{q,n}^2$. Therefore, both $\hat{Z}_n(t)$ and $\hat{Y}_n(t)$ are driftless Brownian motion with variance $\sigma_{z,n}^2 = \frac{\sigma_{q,n}^2}{k_n}$, regardless of scheduling policy. $\square$

Proof The proof is provided in Appendix A. $\square$

By Lemma 3, we are able to fully characterize the distribution of $\hat{Y}_n(t)$ and $\hat{Z}_n(t)$, regardless of the scheduling policy.

C. A Lower Bound of Capacity Region

By Lemma 3, (19) can be written as $\hat{D}_n(t) = \sup_{0 \leq \tau \leq t} (\max\{0, -\frac{\hat{X}_n(t) - \hat{Z}_n(t)}{q_n}\})$. Define

$$Z(t) := \sum_{n=1}^N Z_n(t) = \sum_{n=1}^N C_n(\lfloor \frac{t}{k_n} \rfloor) - \sum_{n=1}^N q_n t. \quad (46)$$

Since $\hat{Z}_n(t)$ is a driftless Brownian motion and $\hat{Z}_n(t)$ are independent among different n , then $\hat{Z}(t) := \lim_{k \rightarrow \infty} \frac{Z(kt)}{\sqrt{k}}$ is also a driftless Brownian motion with variance $\sigma_z^2 := \sum_{n=1}^N \sigma_{z,n}^2$. Similar to (27), we define

$$\hat{D}^*(t) := \sup_{0 \leq \tau \leq t} (\max\{0, -(\hat{X}^*(\tau) - \hat{Z}(\tau))\}). \quad (47)$$

Note that $\hat{X}^*(\tau) + (-\hat{Z}(\tau))$ is the sum of two independent driftless Brownian motion and therefore is also a driftless Brownian motion with variance $(\sigma^2 + \sigma_z^2)$. By using a similar argument as (29)-(31), we further have

$$\hat{D}^*(t) \leq_{st} \sum_{n=1}^N q_n \hat{D}_n(t), \quad (48)$$

under any such scheduling policy.

Definition 4 For a system with fading channels and variable-bit-rate videos, a vector $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]$ is said to be feasible if there exists a scheduling policy such that

$$\hat{D}_n(t) \leq_{st} \frac{\lambda_n}{q_n} \hat{D}^*(t), \quad n = 1, 2, \dots, N. \quad (49)$$

Then, the capacity region for QoE is defined as the set of all feasible vectors λ . $\square$

Again, we obtain a lower bound of capacity region as follows.

Theorem 10 For a system with fading channels and variable-bit-rate videos, a feasible vector $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]$ with $\lambda_n \geq 0$, for all n , must satisfy $\sum_{n=1}^N \lambda_n \geq 1$. $\square$

Proof Similar to the proof of Theorem 5, this can be proved by using (48) and Definition 4. $\square$

D. Scheduling Policy

For fading channels and variable-bit-rate videos, we propose the following extended version of the JCD policy.

Highest Data Rate Policy For Variable-Bit-Rate Videos (HDR-VBR):

In each time slot t , the AP schedules a client with the largest $r_n(t)$ and break ties by choosing the one with the smallest $w_n(X_n(t) - Z_n(t))$, where w_n is a predetermined weight factor. $\square$

Remark 8 Note that $X_n(t) - Z_n(t) = (A_n(t) - q_n t) - (C_n(\lfloor \frac{t}{k_n} \rfloor) - q_n t) = A_n(t) - C_n(\lfloor \frac{t}{k_n} \rfloor)$, which reflects the difference between the total received video content and the total video content that should have been played if there is no video interruption at all. Hence, $X_n(t) - Z_n(t)$ still loosely reflects the status of the playback buffer of client n .

Remark 9 Under the HDR-VBR policy, the AP requires the information of $A_n(t)$ and $C_n(\lfloor \frac{t}{k_n} \rfloor)$. In wireless networks, $A_n(t)$ can be obtained by collecting ACKs from the clients. For $C_n(\lfloor \frac{t}{k_n} \rfloor)$, since the AP has the video files, the AP can simply refer to the accumulative size of the frames that should have been played up to current time t . Hence, the HDR-VBR policy can be easily implemented on the AP.

Remark 10 For the special case of constant-bit-rate videos, the HDR-VBR policy degenerates to the HDR policy studied in [1]. Under the HDR policy, the AP schedules a client with the largest $r_n(t)$ and break ties by choosing the one with the smallest $w_n X_n(t)$ in each time slot t .

Theorem 11 Let w_n be the predetermined weight for client n . For variable-bit-rate video, under the HDR-VBR policy and conditions (37) and (38), we have $w_n(\hat{X}_n(t) - \hat{Z}_n(t)) = w_m(\hat{X}_m(t) - \hat{Z}_m(t))$, for any pair of clients n, m . $\square$

Proof The proof is provided in Appendix B. $\square$

Given the state-space collapse property, the HDR-VBR can achieve every interior point in the capacity region. The key results are summarized in the following theorem.

Theorem 12 Given any vector $\lambda = [\lambda_1, \lambda_2, \dots, \lambda_N]$ which satisfies $\lambda_n > 0$, $\forall n$, and $\sum_{n=1}^N \lambda_n \geq 1$, HDR-VBR policy can achieve $\hat{D}_n(t) = \frac{\frac{1}{q_n} \frac{w_n}{\sum_{m=1}^N \frac{1}{w_m}}}{\frac{1}{q_n} \frac{w_n}{\sum_{m=1}^N \frac{1}{w_m}}} \hat{D}^*(t) \leq_{st} \frac{\lambda_n}{q_n} \hat{D}^{**}(t)$ by assigning $w_n = \frac{1}{\lambda_n}$ for all n . Moreover, we have

$$E[\hat{D}_n(t)] = \sqrt{\frac{2t(\sigma^2 + \sigma_z^2)}{\pi}} \frac{\frac{1}{q_n w_n}}{\sum_{m=1}^N \frac{1}{w_m}}. \quad (50)$$

Remark 11 In (50), we see that video interruption arises from two factors: σ^2 due to the randomness in fading channels and σ_z^2 due to the randomness in variable-bit-rate videos.

Based on Theorems 10 and 12, we characterize the capacity region for fading channels and variable-bit-rate videos.

Theorem 13 For a system with fading channels and variable-bit-rate videos, a vector $[\lambda_1, \lambda_2, \dots]$ with $\lambda_n > 0$, $\forall n$ is feasible if and only if $\sum_{n=1}^N \lambda_n \geq 1$. $\square$

VI. NETWORK UTILITY MAXIMIZATION FOR QoE

In this section, we propose a network utility maximization (NUM) problem for QoE, and obtain tractable solutions for special cases. Given $\hat{D}_n(t)$ and T , we assume that each client n suffers from some *penalty* $f_n(E[\hat{D}_n(T)])$, where $f_n(\cdot)$ is an increasing, differentiable, and convex function. Note that the expectation $E[\hat{D}_n(T)]$ is taken over all sample paths for $t \in [0, T]$. Here we use $E[\hat{D}_n(T)]$ to approximate the short-term playback interruption $D_n(T)$. We then aim to minimize the total penalty in the system, which can be expressed as $\sum_n f_n(E[\hat{D}_n(T)])$.

By Theorems 8 and 13, we have $\sum_n q_n E[\hat{D}_n(T)] \geq E[\hat{D}^*(T)]$ for ON-OFF channels and $\sum_n q_n E[\hat{D}_n(T)] \geq E[\hat{D}^*(T)]$ for fading channels and variable-bit-rate videos. Further, if we have the additional condition of $E[\hat{D}_n(T)] > 0$, the JCD policy and the HDR-VBR policy can achieve any set of $\{E[\hat{D}_n(T)]\}$ by properly assigning the weight $\{w_n\}$ to each client. Since the formulation of the NUM problem is the same for ON-OFF channels and fading channels plus variable-bit-rate videos, we consider the more general case in the rest of this section.

Below, we study an example of NUM problem which aims to minimize the sum of polynomial penalty functions.

NUM with Polynomial Penalty Functions:

Given the distribution of $\mathbf{r}(t)$, $T > 0$, $\alpha \geq 1$, and a vector $[\delta_1, \delta_2, \dots]$ with $\delta_n > 0$, for all n ,

$$\text{Min. } \sum_{n=1}^N \delta_n \cdot (E[\hat{D}_n(T)])^\alpha$$

$$\text{s.t. } q_1 E[\hat{D}_1(T)] + \dots + q_N E[\hat{D}_N(T)] \geq E[\hat{D}^*(T)]. \quad \square$$

To minimize the total penalty, we define a function L_1 with a Lagrange multiplier μ_1 as

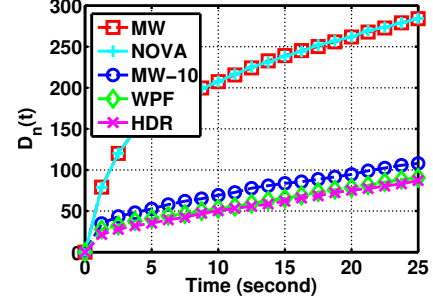
$$L_1 = \sum_{n=1}^N \delta_n (E[\hat{D}_n(T)])^\alpha - \mu_1 \left(\sum_{n=1}^N q_n E[\hat{D}_n(T)] - E[\hat{D}^*(T)] \right).$$

Next, we take the partial derivative of L_1 with respect to each $E[\hat{D}_n(T)]$ and set them to zero, i.e. $\frac{\partial L_1}{\partial E[\hat{D}_n(T)]} = \delta_n \alpha (E[\hat{D}_n(T)])^{\alpha-1} - \mu_1 q_n = 0$, $\forall n$. If $\alpha > 1$, an optimal solution occurs when $\frac{E[\hat{D}_n(T)]}{E[\hat{D}_m(T)]} = \left(\frac{q_n / \delta_n}{q_m / \delta_m} \right)^{\frac{1}{\alpha-1}}$, for any pair n, m . From (50), it is equivalent to have $\frac{\beta_n}{\beta_m} = \frac{q_n^{\frac{\alpha}{\alpha-1}} \cdot \delta_n^{\frac{-1}{\alpha-1}}}{q_m^{\frac{\alpha}{\alpha-1}} \cdot \delta_m^{\frac{-1}{\alpha-1}}}$.

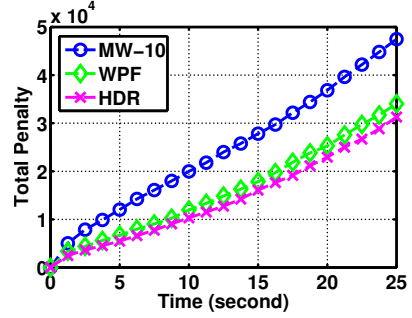
Then, we can simply assign $w_n = \delta_n^{\frac{1}{\alpha-1}} q_n^{\frac{-\alpha}{\alpha-1}}$ for each client so that HDR-VBR achieves an optimal solution. If $\alpha = 1$, the problem degenerates to a linear program. An optimal solution is obtained by assigning almost all the video interruption time to a client with the smallest $\frac{\delta_n}{q_n}$. Without loss of generality, we may assume that $\frac{\delta_1}{q_1} \leq \frac{\delta_2}{q_2} \leq \dots \leq \frac{\delta_N}{q_N}$. Then, we just assign $w_1 = 1$ and let w_n be extremely large for the other clients.

VII. SIMULATION RESULTS

We evaluate the proposed policies through ns-2 simulation. Following the IEEE 802.11a standard, we simulate a wireless network that allows data transmission at 54, 48, 36, 18, and 6



(a) Average duration of video interruption.



(b) Total penalty of the network.

Fig. 1. Comparisons of the five policies in a fully symmetric system with constant-bit-rate videos.

Mbit/s. The time to transmit a packet and to receive an ACK is set to be 500 μ s, which is short enough so that the channel quality stays almost the same in a time slot. We thereby obtain the corresponding packet size for each data rate: 2340, 2080, 1560, 750, 220 bytes. The frame rate of each video stream is 30 frames per second, and thus each client plays one frame every 33.3 milliseconds (equivalent to about 66 time slots). All the results presented in this section are the average of 50 simulation trials. We compare HDR-VBR policy against four policies: HDR policy, Max-Weight policy (MW), weighted Proportional-Fair policy (WPF), and NOVA algorithm. In MW policy, the AP schedules the one with the largest $(-r_n X_n(t))$ and breaks tie by choosing the one with the largest $(-X_n(t))$. To further explore the difference between HDR-VBR and MW, we also consider the Max-Weight- α policy (MW- α), which schedules the client with the largest $r_n (\max(0, -X_n(t)))^{\frac{1}{\alpha}}$. When $\alpha > 1$, the instantaneous data rate becomes more influential than $X_n(t)$. In the following simulations, we assign $\alpha = 10$. For WPF policy, the scheduled client at time t is the one that maximizes $q_n (r_n(t) / A_n(t-1))$ [31]. For NOVA, we choose the same objective function as that in [22] with a slight change in the initial condition (b_{i0} in [22]) to fit in our simulation scenario.

A. Constant-Bit-Rate Videos

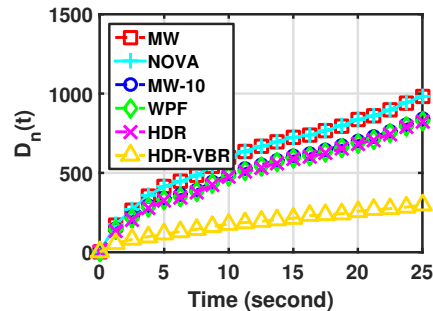
For the special case of constant-bit-rate videos, HDR-VBR is exactly the same as the HDR policy. Therefore, we merge

the results of HDR and HDR-VBR for constant-bit-rate videos. We consider a fully symmetric system of 20 clients under the heavy-traffic condition given by (37) and (38). The channel distribution of each client is the same and evenly distributed, i.e. the probability of each data rate is 0.2. Under this setting, $E[R(t)] \approx 2340$. Therefore, q_n is chosen to be 117 byte/slot for every client. We study a quadratic QoE objective function given by $\sum_n \delta_n (E[\hat{D}_n(T)])^2$, where $\delta_n = 1$ for all n . Since the system is fully-symmetric, we choose $w_n = 1$ for all the clients. Fig. 1 shows the results of the symmetric system. In Fig. 1(a), HDR has the smallest $D_n(t)$ among all the policies, while MW and NOVA perform rather poorly. As expected, MW-10 policy has a moderate $D_n(t)$ since MW-10 serves as an intermediate between MW and HDR. Moreover, it is noticeable that WPF has similar performance to HDR. The main reason is that in the symmetric case, $A_n(t)$ of each client grows almost at the same speed and thus maximizing $q_n(r_n(t)/A_n(t-1))$ is equivalent to maximizing $r_n(t)$ in each slot. Moreover, Fig. 1(b) shows the total penalty of MW-10, WPF, and HDR policy to further compare the difference between these three policies.

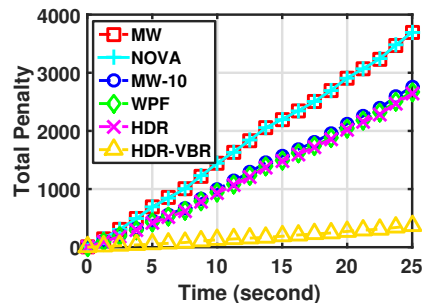
B. Variable-Bit-Rate Videos

Following the same simulation setup as that of constant-bit-rate videos, we evaluate the HDR-VBR policy against HDR as well as other popular policies with variable-bit-rate videos. First, we consider a fully-symmetric system as that in Section VII-A but with variable-bit-rate videos. We assume that the frame size of each video is uniformly distributed between 100 bytes and 15344 bytes with average frame size of 7722 bytes. This corresponds to an average playback rate of 117 bytes per slot for every client. Fig. 2(a) shows the average playback interruption of the fully-symmetric system. As expected, the HDR-VBR has the smallest $D_n(t)$ among all the policies. More importantly, with variable-bit-rate videos, the HDR-VBR policy outperforms the HDR policy since the original HDR does not include the information about variable playback rates. Besides, Fig. 2(b) demonstrates the total penalty incurred by playback interruption. It is also noticeable that the total penalty is larger in Fig. 2(b) than that in Fig. 1(b) due to the fluctuation in video playback rates.

Next, we turn to the asymmetric case. We divide the clients equally into two classes. We assign $\delta_n = 10$ to Class 1 and $\delta_n = 1$ to Class 2, with the result that Class 1 dominates the overall QoE performance. In addition, we assume that the two classes have the same evenly-distributed channel but different playback rates. Suppose the clients in Class 1 and Class 2 watch videos with resolution of 480p and 720p, respectively. According to the recommended bitrates for YouTube videos in [32], we choose $q_n = 156$ and 78 bytes/slot for 720p and 480p videos, respectively. To include randomness in frame sizes, we assume that each client in Class 1 plays a video with frame size uniformly distributed between a minimum 100 bytes and maximum 10196 bytes with average playback rate of 78 bytes per slot. Similarly, each client in Class 2 plays a video with frame size uniformly distributed between 100 bytes and 20492 bytes with average playback rate of 156 bytes per slot. By



(a) Average duration of video interruption.



(b) Total penalty of the network.

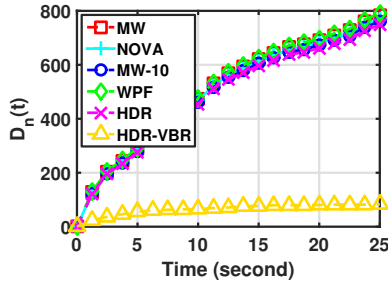
Fig. 2. Comparisons of the six policies in a fully symmetric system with variable-bit-rate videos.

TABLE II
INFORMATION OF THE VIDEOS IN THE EXPERIMENTS.

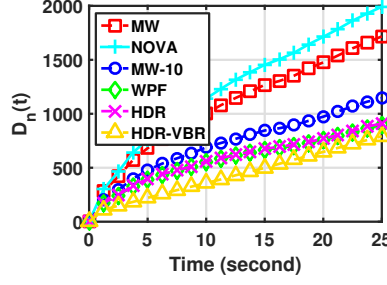
Video #	Video Name	Avg. Bitrate (Mbps)	Frame Rate
1	The Simpsons (Official Trailer)	2.04	24
2	Serenity (Official Trailer)	2.25	24
3	Toy Story 3 (Official Trailer)	0.71	30
4	Angry Birds (Official Trailer)	0.65	24
5	The Simpsons (Official Trailer, HD)	3.42	24

the discussion in Section VI, we assign $w_n = 40$ to Class 1 and $w_n = 1$ to Class 2 to optimize the network utility under the HDR-VBR policy. Fig. 3 shows the playback interruption and the total penalty of the asymmetric system. Clearly, the HDR-VBR still achieves the smallest playback interruption for both clients in Class 1 and Class 2 by taking the fluctuation in video frame size into account. Moreover, the HDR-VBR policy intelligently allocates $D_n(t)$ among the two classes by assigning proper weights w_n so that it can achieve the smallest total penalty.

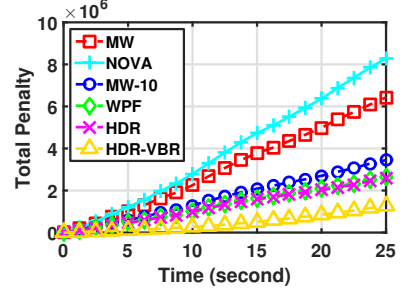
From simulation, we note that all of the five policies are stabilizing since the duration of video interruption grows sub-linearly. However, the short-term performance of these policies are rather different in the heavy-traffic regime. Therefore, diffusion limit indeed provides more detailed information on the playback process.



(a) Average playback interruption: Class 1.



(b) Average playback interruption: Class 2.



(c) Total penalty of the network.

Fig. 3. Comparisons of the five policies in a system with the same channel distribution but heterogeneous playback rates.

TABLE III
EXPERIMENTAL RESULTS UNDER HEAVY-TRAFFIC CONDITION.

Policy	Trial	$D_1(t)$ (sec)	$D_2(t)$ (sec)	$D_3(t)$ (sec)	$D_4(t)$ (sec)	$D_5(t)$ (sec)	RX Thruput (Mbps)
HDR-VBR	#1	3.12	9.14	4.06	4.11	0.00	8.39
	#2	3.41	9.29	4.27	4.09	0.00	8.40
	#3	2.72	9.01	3.80	3.46	0.00	8.55
	#4	2.36	7.48	3.88	3.83	0.00	8.57
	#5	0.29	7.32	1.24	1.06	0.00	9.47
	#6	0.19	7.33	1.21	1.02	0.00	9.36
WPF	#1	9.78	9.61	10.36	8.96	0.00	8.47
	#2	9.36	9.46	9.99	8.62	0.00	8.51
	#3	9.26	9.38	9.85	8.46	0.00	8.59
	#4	6.56	7.59	6.77	5.21	0.00	9.46
	#5	8.33	8.59	8.61	7.54	0.00	8.85
	#6	7.10	8.17	7.30	6.00	0.00	9.19

TABLE IV
EXPERIMENTAL RESULTS UNDER NON-HEAVY-TRAFFIC CONDITION.

Policy	Trial	$D_1(t)$ (sec)	$D_2(t)$ (sec)	$D_3(t)$ (sec)	$D_4(t)$ (sec)	$D_5(t)$ (sec)	RX Thruput (Mbps)
HDR-VBR	#1	0.00	4.02	0.00	0.00	0.00	11.93
	#2	0.00	3.59	0.00	0.00	0.00	11.98
	#3	0.00	3.56	0.00	0.00	0.00	12.30
	#4	0.00	3.48	0.00	0.00	0.00	12.43
	#5	0.00	3.64	0.00	0.00	0.00	12.38
	#6	0.00	3.53	0.00	0.00	0.00	12.44
WPF	#1	2.54	4.00	1.42	1.57	0.00	11.88
	#2	1.99	3.37	0.72	1.03	0.00	12.62
	#3	2.03	3.41	0.78	1.08	0.00	12.55
	#4	2.15	3.50	0.91	1.22	0.00	12.48
	#5	1.97	3.37	0.71	1.03	0.00	12.54
	#6	2.07	3.43	0.81	1.12	0.00	12.41

VIII. EXPERIMENTAL RESULTS WITH REAL VIDEOS

We further evaluate the performance of the proposed policies with real videos on a software-defined wireless testbed. The experiments are done on WiMAC, which is a FPGA-based wireless platform introduced by Yau *et al.* in [33]. With the clean separation between software and hardware, WiMAC allows us to quickly prototype the proposed scheduling policies completely in the software domain. We consider five on-demand videos streams from an AP to five different clients using real video files. Table II shows the basic information

about the videos played in the experiments. For simplicity, we consider symmetric wireless channels for all the clients by having the five clients co-located at the same station. For the MAC and PHY specification, the AP and the clients follow the IEEE 802.11a standard. Suppose the AP aims to minimize a linear objective function as $\sum_n \delta_n E[D_n(t)]$ with $\delta_1 = \delta_5 = 10$ and $\delta_2 = \delta_3 = \delta_4 = 1$. This implies that client 1 and client 5 should expect better QoE than the other three clients. By the results in Section VI, since client 2 has the smallest $\frac{\delta_n}{q_n}$, we choose $w_n = 1$ for client 2 and $w_n = 1000$ for the rest of the clients. We compare the HDR-VBR policy with the WPF policy, which has already been shown to have better performance than MW and NOVA policy in simulations.

A. Heavy-Traffic Case

We first run experiments under heavy-traffic condition, i.e. the total receiver throughput is about the same as total video playback rate. By using a fixed 16-QAM modulation and coding rate 1/2, the total receiver throughput is about 9 Mbps. Figure 4 shows the experimental results under the HDR-VBR and WPF policy. Compared to the WPF policy, the HDR-VBR indeed significantly reduces playback interruption by keeping track of the variation in video bit rates. Moreover, Table III provides the accumulated playback interruption at 40 second and the average receiver throughput of each experiment trial. From Table III, we know that HDR-VBR performs much better than WPF while the receiver throughput under both policies are almost the same.

B. Non-Heavy-Traffic Case

We increase the receiver throughput by using 64-QAM modulation and coding rate 3/4. Under this setting, the receiver throughput is about 12 Mbps, which corresponds to roughly 75% system load. Figure 5 shows the video interruption of each client and the total penalty under the HDR-VBR and WPF policy. Due to the increase in throughput, both policies achieve much less video interruption than in the heavy-traffic case. In this case, HDR-VBR achieves zero video interruption for client 1, 3, 4, and 5 while still maintaining about the same amount of video interruption for client 2 as that under WPF. Table IV further verifies this result through multiple experimental trials.

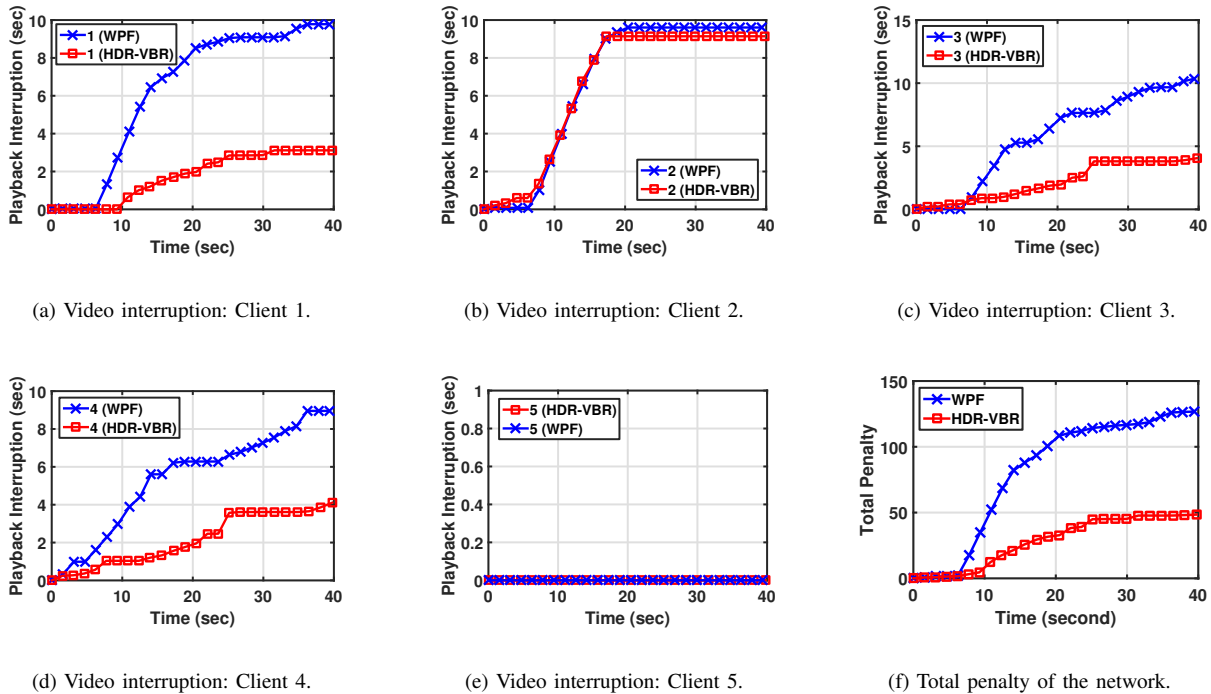


Fig. 4. Experimental results with five real video streams under heavy-traffic condition.

Hence, the experiments demonstrate that HDR-VBR outperforms its counterparts with real videos under both heavy-traffic and non-heavy-traffic conditions.

IX. CONCLUSIONS

In this paper, we study dynamic behavior of QoE in heavy traffic by using diffusion approximation. We characterize the capacity region for QoE and propose online scheduling policies to optimize QoE. Simulation and experimental results show that the proposed policies outperform existing popular policies. In the future, we intend to further study the effect of adaptive video bit rates and user engagement on QoE.

ACKNOWLEDGEMENTS

This material is based upon work supported in part by NSF under contract number CNS-1719384, the U.S. Army Research Laboratory and the U.S. Army Research Office under contract/Grant Number W911NF-15-1-0279, Office of Naval Research under Contract N00014-18-1-2048, and NPRP Grant 8-1531-2- 651 of Qatar National Research Fund (a member of Qatar Foundation).

REFERENCES

- [1] P.C. Hsieh and I.H. Hou, "Heavy-traffic analysis of QoE optimality for on-demand video streams over fading channels," in *Proc. of IEEE INFOCOM*, April 2016, pp. 1–9.
- [2] Cisco, "Cisco visual networking index: Global mobile data traffic forecast update, 2014-2019," Feb. 2015.
- [3] N. Staelens *et al.*, "Assessing quality of experience of IPTV and video on demand services in real-life environments," *IEEE Transactions on Broadcasting*, vol. 56, no. 4, pp. 458–466, 2010.
- [4] R. Mok *et al.*, "Measuring the quality of experience of HTTP video streaming," in *Proc. of IFIP/IEEE International Symposium on Integrated Network Management (IM)*, 2011, pp. 485–492.
- [5] F. Dobrian *et al.*, "Understanding the impact of video quality on user engagement," in *Proc. of ACM SIGCOMM*, 2011, pp. 362–373.
- [6] G. Liang and B. Liang, "Effect of delay and buffering on jitter-free streaming over random VBR channels," *IEEE Transactions on Multimedia*, vol. 10, no. 6, pp. 1128–1141, Oct 2008.
- [7] A. ParandehGheibi *et al.*, "Avoiding interruptions - a QoE reliability function for streaming media applications," *IEEE Journal on Selected Areas in Communications*, vol. 29, no. 5, pp. 1064–1074, 2011.
- [8] Y. Xu *et al.*, "Analysis of buffer starvation with application to objective QoE optimization of streaming services," *IEEE Transactions on Multimedia*, vol. 16, no. 3, pp. 813–827, 2014.
- [9] A. Anttonen and A. Mammela, "Interruption probability of wireless video streaming with limited video lengths," *IEEE Transactions on Multimedia*, vol. 16, no. 4, pp. 1176–1180, June 2014.
- [10] J. Yang *et al.*, "Online buffer fullness estimation aided adaptive media playout for video streaming," *IEEE Transactions on Multimedia*, vol. 13, no. 5, pp. 1141–1153, Oct 2011.
- [11] T. Luan *et al.*, "Impact of network dynamics on user's video quality: Analytical framework and QoS provision," *IEEE Transactions on Multimedia*, vol. 12, no. 1, pp. 64–78, Jan 2010.
- [12] Y. Xu *et al.*, "Impact of flow-level dynamics on QoE of video streaming in wireless networks," in *Proc. of IEEE INFOCOM*, 2013, pp. 2715–2723.
- [13] —, "Modeling buffer starvations of video streaming in cellular networks with large-scale measurement of user behavior," *IEEE Transactions on Mobile Computing*, vol. 16, no. 8, pp. 2228–2245, 2017.
- [14] D. De Vleeschauwer *et al.*, "Optimization of HTTP adaptive streaming over mobile cellular networks," in *Proc. of IEEE INFOCOM*, April 2013, pp. 898–997.
- [15] P. Chandur and K. Sivalingam, "Quality of experience aware video scheduling in LTE networks," in *Proc. of NCC*, Feb 2014, pp. 1–6.
- [16] R. Bhatia *et al.*, "Improving mobile video streaming with link aware scheduling and client caches," in *Proc. of IEEE INFOCOM*, April 2014, pp. 100–108.
- [17] K.S. Kim *et al.*, "Scheduling multicast traffic with deadlines in wireless networks," in *Proc. of IEEE INFOCOM*, April 2014, pp. 2193–2201.
- [18] M. Li *et al.*, "QoE-aware scheduling for video streaming in 802.11 n/ac-based high user density networks," in *IEEE Vehicular Technology Conference (VTC Spring)*, 2016, pp. 1–5.

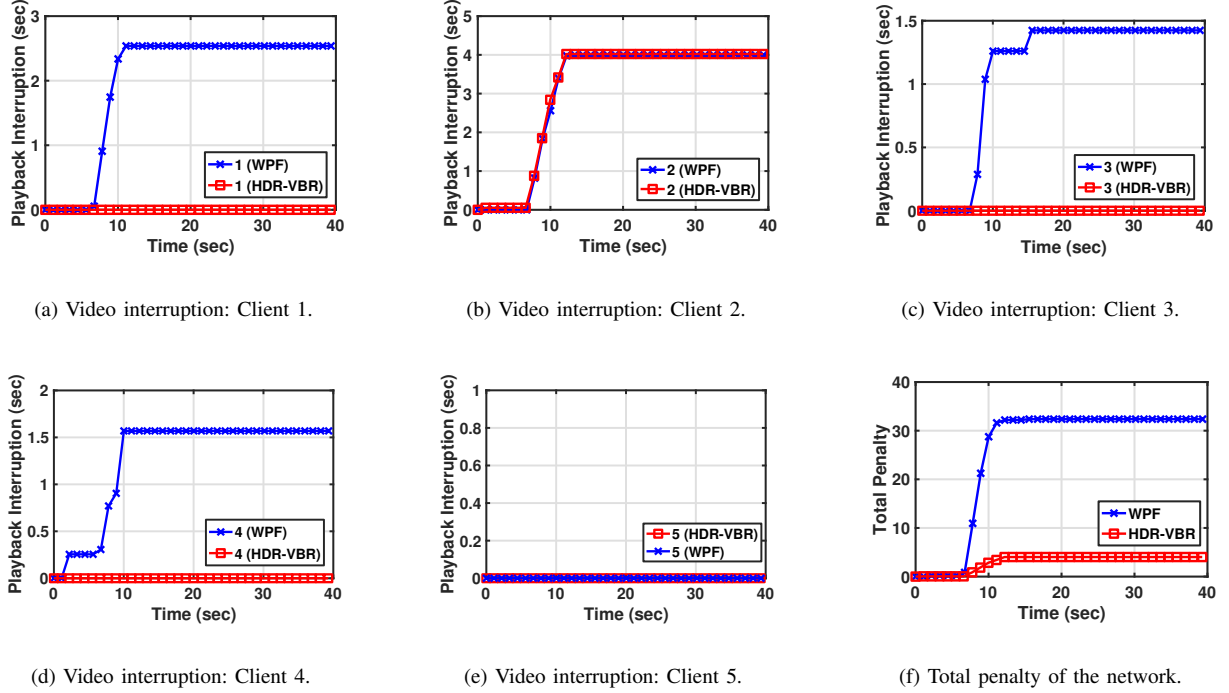


Fig. 5. Experimental results with five real video streams under non-heavy-traffic condition.

- [19] S. Cicalo *et al.*, “Improving QoE and fairness in HTTP adaptive streaming over LTE network,” *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 26, no. 12, pp. 2284–2298, 2016.
- [20] C. Li *et al.*, “Joint dynamic rate control and transmission scheduling for scalable video multirate multicast over wireless networks,” *IEEE Transactions on Multimedia*, vol. 20, no. 2, pp. 361–378, 2018.
- [21] A. Anand and G. de Veciana, “Measurement-based scheduler for multi-class QoE optimization in wireless networks,” in *Proc. of IEEE INFOCOM*, 2017, pp. 1–9.
- [22] V. Joseph and G. de Veciana, “NOVA: QoE-driven optimization of DASH-based video delivery in networks,” in *Proc. of IEEE INFOCOM*, April 2014, pp. 82–90.
- [23] K. Xiao *et al.*, “QoE-driven resource allocation for DASH over OFDMA networks,” in *Proc. of GLOBECOM*, 2016, pp. 1–6.
- [24] I.H. Hou and P.C. Hsieh, “QoE-optimal scheduling for on-demand video streams over unreliable wireless networks,” in *Proc. of ACM MobiHoc*, 2015, pp. 207–216.
- [25] H. Chen and D. Yao, *Fundamentals of Queueing Networks: Performance, Asymptotics and Optimization*. Springer-Verlag, 2001.
- [26] L. Tassiulas and A. Ephremides, “Dynamic server allocation to parallel queues with randomly varying connectivity,” *IEEE Transactions on Information Theory*, vol. 39, no. 2, pp. 466–478, 1993.
- [27] M. Neely, “Delay analysis for Max Weight opportunistic scheduling in wireless systems,” *IEEE Transactions on Automatic Control*, vol. 54, no. 9, pp. 2137–2150, Sept 2009.
- [28] J. Harrison and M. Lopez, “Heavy traffic resource pooling in parallel-server systems,” *Queueing Systems*, vol. 33, no. 4, pp. 339–368, 1999.
- [29] J.M. Harrison, *Brownian motion and stochastic flow systems*. Wiley New York, 1985.
- [30] M.J. Neely *et al.*, “Dynamic power allocation and routing for time-varying wireless networks,” *IEEE Journal on Selected Areas in Communications*, vol. 23, no. 1, pp. 89–103, 2005.
- [31] H. Kushner and P. Whiting, “Convergence of proportional-fair sharing algorithms under general conditions,” *IEEE Transactions on Wireless Communications*, vol. 3, no. 4, pp. 1250–1259, July 2004.
- [32] “Live encoder settings, bitrates and resolutions.” [Online]. Available: <https://support.google.com/youtube/answer/2853702?hl=EN>
- [33] S. Yau *et al.*, “WiMAC: Rapid Implementation Platform for User Definable MAC Protocols Through Separation,” in *Proc. of ACM SIGCOMM (demo paper)*, 2015, pp. 109–110.

APPENDIX A PROOF OF LEMMA3

Proof Define $\tilde{C}_n(j) := C_n(j) - jq_n^*$ and $\Delta\tilde{C}_n(j+1) := \tilde{C}_n(j+1) - \tilde{C}_n(j)$. By the i.i.d. assumption on frame size, $\Delta\tilde{C}_n(j)$ is i.i.d. across all time slots. Moreover,

$$E[\Delta\tilde{C}_n(j)] = E[F_n(j) - q_n^*] = E[F_n(j)] - q_n^* = 0,$$

$$\text{Var}[\Delta\tilde{C}_n(j)] = \text{Var}[F_n(j) - q_n^*] = \text{Var}[F_n(j)] = \sigma_{q,n}^2.$$

By the functional central limit theorem for i.i.d. random variables, we know $\hat{C}_n(t)$ is a driftless Brownian motion with variance $\sigma_{q,n}^2$. Hence, $\hat{C}_n(\frac{t}{k_n})$ is a driftless Brownian motion with variance $\sigma_{q,n}^2/k_n$. Next, consider $\hat{Z}_n(t)$ as

$$\hat{Z}_n(t) = \lim_{k \rightarrow \infty} \frac{C_n(\lfloor \frac{kt}{k_n} \rfloor) - q_n kt}{\sqrt{k}} \quad (51)$$

$$= \lim_{k \rightarrow \infty} \frac{C_n(k \lfloor \frac{kt}{k_n} \rfloor) - q_n^* k \frac{t}{k_n}}{\sqrt{k}} = \hat{C}_n\left(\frac{t}{k_n}\right). \quad (52)$$

Finally, we consider $\hat{Y}_n(t)$ as

$$\hat{Y}_n(t) = \lim_{k \rightarrow \infty} \frac{C_n(S_n(kt)) - q_n(kt - D_n(kt)) + e_n(kt)}{\sqrt{k}} \quad (53)$$

$$= \lim_{k \rightarrow \infty} \frac{C_n(\lfloor \frac{kt - D_n(kt)}{k_n} \rfloor) - q_n(kt - D_n(kt)) + e_n(kt)}{\sqrt{k}} \quad (54)$$

$$= \lim_{k \rightarrow \infty} \frac{C_n(k \lfloor \frac{kt - D_n(kt)}{k_n} \rfloor) - q_n(k(t - \frac{D_n(kt)}{k_n})) + e_n(kt)}{\sqrt{k}}. \quad (55)$$

By the Random Time-Change Theorem (Theorem 5.3 in [25]) and (52)-(55), we conclude $\hat{Y}_n(t) = \hat{C}_n\left(\frac{t}{k_n}\right) = \hat{Z}_n(t)$. $\square$

APPENDIX B PROOF OF THEOREM 11

We prove the state space collapse property by introducing a fluid system. First, define

$$Q_n(t) = -w_n(X_n(t) - Z_n(t)) + \frac{\sum_{m=1}^N (X_m(t) - Z_m(t))}{\sum_{m=1}^N \frac{1}{w_m}}.$$

Two important facts of the above definition:

- At any time t , the client n with the largest $Q_n(t)$ also has the largest $-w_n(X_n(t) - Z_n(t))$.
- Since $\frac{\sum_{m=1}^N (X_m(t) - Z_m(t))}{\sum_{m=1}^N \frac{1}{w_m}}$ is the weighted average of $w_n(X_n(t) - Z_n(t))$, we have $\max_{1 \leq m \leq N} Q_m \geq 0$.

Next, we study the fluid limit of $Q_n(t)$ defined as

$$\bar{Q}_n(t) := \lim_{k \rightarrow \infty} \frac{Q_n(kt)}{k} \quad (56)$$

Similarly, define $\bar{X}_n(t) := \lim_{k \rightarrow \infty} \frac{X_n(kt)}{k}$ and $\bar{Z}_n(t) := \lim_{k \rightarrow \infty} \frac{Z_n(kt)}{k}$ to be the fluid limits of $X_n(t)$ and $Z_n(t)$, respectively. Since $Z_n(t) := C_n(\lfloor \frac{t}{k_n} \rfloor) - q_n t$, we thus have

$$\bar{Z}_n(t) = \lim_{k \rightarrow \infty} \frac{C_n(\lfloor \frac{kt}{k_n} \rfloor) - q_n kt}{k} = \lim_{k \rightarrow \infty} \frac{C_n(k \lfloor \frac{kt}{k_n} \rfloor)}{k} - q_n t \quad (57)$$

Since $\lim_{k \rightarrow \infty} \frac{\lfloor \frac{kt}{k_n} \rfloor}{k} = \frac{t}{k_n}$, then by the Random Time-Change Theorem (Theorem 5.3 in [25]), we have $\lim_{k \rightarrow \infty} \frac{C_n(k \lfloor \frac{kt}{k_n} \rfloor)}{k} = q_n^* \frac{t}{k_n} = q_n t$. Hence, $\bar{Z}_n(t) = 0$, for any $t \geq 0$, for every client n . Thus, $\bar{Q}_n(t)$ can be written as

$$\bar{Q}_n(t) = -w_n(\bar{X}_n(t) - \bar{Z}_n(t)) + \frac{\sum_{m=1}^N (\bar{X}_m(t) - \bar{Z}_m(t))}{\sum_{m=1}^N \frac{1}{w_m}} \quad (58)$$

$$= -w_n \bar{X}_n(t) + \frac{\sum_{m=1}^N \bar{X}_m(t)}{\sum_{m=1}^N \frac{1}{w_m}}. \quad (59)$$

The rest of the proof is to show that the random process $\{Q_n(t)\}$ is positive recurrent for all n . Define a Lyapunov function

$$L_Q(t) = \sum_{n=1}^N \frac{1}{2w_n} [\bar{Q}_n(t)]^2. \quad (60)$$

We again assume that fluid limits of $Q_m(t)$ are sorted in descending order, i.e. $\bar{Q}_1(t) \geq \bar{Q}_2(t) \geq \dots \geq \bar{Q}_N(t)$. Let U_n be the event that $r_n(t)$ equals $R(t)$ at some given time t . Since $X_n(t) = A_n(t) - q_n t$, under the HDR-VBR policy we have

$$\frac{d\bar{X}_n(t)}{dt} = E \left[R(t) \cdot \mathbb{I} \left\{ \left(\bigcap_{k=1}^{n-1} U_k^c \right) \cap U_n \right\} \right] - q_n, \quad (61)$$

where $\left\{ \left(\bigcap_{k=1}^{n-1} U_k^c \right) \cap U_n \right\}$ represents the event that client n is the only client in $\{1, 2, \dots, n\}$ which has the largest

transmission rate among all clients. Now, let $\tilde{r}_n := E \left[R(t) \cdot \mathbb{I} \left\{ \left(\bigcap_{k=1}^{n-1} U_k^c \right) \cap U_n \right\} \right]$. Then, we define $h_k := \sum_{j=1}^k (\tilde{r}_j - q_j) = E \left[R(t) \cdot \mathbb{I} \left\{ \bigcup_{j=1}^k U_j \right\} \right] - \sum_{j=1}^k q_j$, where $\left\{ \bigcup_{j=1}^k U_j \right\}$ represents the event that at least one client in $\{1, 2, \dots, k\}$ has the largest transmission rate among all clients. By using the conditions in (37) and (38), we obtain that

$$\begin{cases} h_k > 0, & \text{if } k = 1, 2, \dots, N-1 \\ h_k = 0, & \text{if } k = N \end{cases} \quad (62)$$

where the last equality holds since $\sum_{m=1}^N (\tilde{r}_m - q_m) = h_N$ and should be zero. For convenience, let $h_0 = 0$. Thus,

$$\frac{d\bar{Q}_n(t)}{dt} = -w_n(\tilde{r}_n - q_n) + \frac{\sum_{m=1}^N (\tilde{r}_m - q_m)}{\sum_{m=1}^N \frac{1}{w_m}} = -w_n(\tilde{r}_n - q_n),$$

Finally, the Lyapunov drift is given by

$$\begin{aligned} \frac{dL_Q}{dt} &= - \sum_{n=1}^N (\tilde{r}_n - q_n) \cdot \bar{Q}_n(t) \\ &= - \sum_{n=1}^N (h_n - h_{n-1}) \cdot \bar{Q}_n(t) \\ &= - \left[\left(\sum_{n=1}^{N-1} h_n (\bar{Q}_n(t) - \bar{Q}_{n+1}(t)) \right) + h_N \bar{Q}_N(t) \right] \leq 0. \end{aligned}$$

Note that the drift is zero only if $\bar{Q}_1(t) = \bar{Q}_2(t) = \dots = \bar{Q}_N(t) = 0$. Hence, the random process $\{Q_n(t)\}$ is positive recurrent. Therefore, $\hat{Q}_n(t) = 0$, for every n . This result implies that $w_n(\hat{X}_n(t) - \hat{Z}_n(t)) = w_m(\hat{X}_m(t) - \hat{Z}_m(t))$, for every pair n, m . $\square$



Ping-Chun Hsieh received the B.S. in Electrical Engineering and M.S. in Electronics Engineering from National Taiwan University in 2011 and 2013, respectively. He is currently a PhD student in the Department of Electrical and Computer Engineering at Texas A&M University, College Station, Texas, USA. His research interests include wireless networks, networked transportation systems, and software-defined wireless platforms. He received the Best Paper Award in ACM MobiHoc 2017.



I-Hong Hou (S'10-M'12) received the B.S. in Electrical Engineering from National Taiwan University in 2004, and his M.S. and Ph.D. in Computer Science from University of Illinois, Urbana-Champaign in 2008 and 2011, respectively.

In 2012, he joined the department of Electrical and Computer Engineering at the Texas A&M University, where he is currently an assistant professor. His research interests include wireless networks, wireless sensor networks, real-time systems, distributed systems, and vehicular ad hoc networks.

Dr. Hou received the Best Paper Award in ACM MobiHoc 2017, the Best Student Paper Award in WiOpt 2017, and the C.W. Gear Outstanding Graduate Student Award from the University of Illinois at Urbana-Champaign.