

Quality of Real-Time Streaming in Wireless Cellular Networks - Stochastic Modeling and Analysis

Bartłomiej Błaszczyszyn, Miodrag Jovanovic
and Mohamed Kadhém Karray, *Member, IEEE*,

Abstract—We present a new stochastic service model with capacity sharing and interruptions, appropriate for the evaluation of the quality of real-time streaming (e.g. mobile TV) in wireless cellular networks. It takes into account multi-class Markovian process of call arrivals (to capture different radio channel conditions, requested streaming bit-rates and call-durations) and allows for a general resource allocation policy saying which users are temporarily denied the requested fixed streaming bit-rates (put in outage) due to resource constraints. We develop general expressions for the performance characteristics of this model, including the mean outage duration and the mean number of outage incidents for a typical user of a given class, involving only the steady-state of the traffic demand. We propose also a natural class of least-effort-served-first resource allocation policies, which cope with optimality and fairness issues known in wireless networks, and whose performance metrics can be easily calculated using Fourier analysis of Poisson variables. We specify and use our model to analyze the quality of real time streaming in 3GPP Long Term Evolution (LTE) cellular networks. Our results can be used for the dimensioning of these networks.

Index Terms—Real-time streaming, stochastic model, mobile TV, LTE, quality of service, interruptions, outage, deep outage, capacity-sharing, Poisson process.

I. INTRODUCTION

WIRELESS cellular networks offer nowadays possibility to watch TV on mobile devices, which is an example of a real-time content streaming. This type of traffic demand is expected to increase significantly in the future. In order to cope with this process, network operators need to implement in their dimensioning tools efficient methods allowing to predict the quality of this type of service. The quality of real-time streaming (RTS) is principally related to the number and duration of *outage incidents* — (hopefully short) periods when the network cannot deliver to a given user in real-time the requested content of the required quality. In this paper we propose a stochastic model allowing for an analytic evaluation of such metrics. It assumes a traffic demand with different

radio conditions of calls, and can be specified to take into account the parameters of a given wireless cellular technology. We develop expressions for several important performance characteristics of this model, including the mean time spent in outage and the mean number of outage incidents for a typical streaming call in function of its radio conditions. These expressions involve only stationary probabilities of the (free) traffic demand process, which is a vector of independent Poisson random variables describing the number of users in different radio conditions.

We use this model to analyze RTS in a typical cell of a 3GPP Long Term Evolution (LTE) cellular network assuming orthogonal intra-cell user channels with the peak bit-rates (achievable when there are no other users in the same cell) close to the theoretical Shannon's bound in the additive white Gaussian noise (AWGN) channel, with the extra-cell interference treated as noise. These assumptions lead to a radio resource constraint in a multi-rate linear form. Namely, each user experiencing a given signal-to-(extra-cell)-interference-and-noise ratio (SINR) requires a fixed fraction of the normalized radio capacity, related to the ratio between its requested and peak bit-rates. All users of a given configuration (experiencing different SINR values) can be entirely satisfied if and only if the total required capacity is not larger than one.¹

In the above context of a multi-rate linear radio resource constraint, we analyse some natural parametric class of *least-effort-served-first* (LESF) service policies, which assign service to users in order of their increasing radio capacity demand, until the full capacity (possibly with some margin) is reached. The capacity margin may be used to offer some “lower quality” service to users temporarily in outage thus realizing some type of fairness with respect to unequal user radio-channel conditions. This class contains an optimal and a fair policy, the latter being suggested by LTE implementations.

In order to evaluate explicitly the quality of service metrics induced by the LESF policies we relate the mean time spent in outage and the mean number of outage incidents for a typical streaming call in given radio conditions to the distribution

Manuscript received July 30, 2013; revised December 20, 2013 and February 18, 2014; accepted February 18, 2014. The associate editor coordinating the review of this paper and approving it for publication was C.-F. Chiasserini.

This paper reports the results of the research undertaken under CRE-CIFRE thesis co-advancing agreement between Inria and Orange Labs.

B. Błaszczyszyn is with Inria-ENS, 23 Avenue d'Italie, 75214 Paris, France (e-mail: Bartek.Blaszczyszyn@ens.fr).

M. Jovanovic and M. K. Karray are with Orange Labs, 38/40 rue Général Leclerc, 92794 Issy-les-Moulineaux, France (e-mail: {miodrag.jovanovic, mohamed.karray}@orange.com).

Digital Object Identifier 10.1109/TWC.2014.131371

¹Recall that in the case of voice calls and, more generally, constant bit-rate (CBR) calls the multi-rate linear form of the resource constraints has already proved to lead to efficient model evaluation methods, via e.g. Kaufman-Roberts algorithm [1, 2]. Despite some fundamental similarities to CBR service, the RTS gives rise to a new model, due to the fact that the service denials are not definitive for a given call, but have a form of temporal interruptions (outage) periods.

functions of some linear functionals of the Poisson vector describing the steady state of the system. We calculate the Fourier transforms of these functions and use a well-known Fourier transform inversion method to obtain numerical values of the quantities of interest. We also study the mean throughput during a typical streaming call evaluating the expectations of the corresponding non-linear functionals of the Poisson vector describing the steady state of the system via the Monte Carlo method.

Using this approach, we present a thorough study of the quality of RTS with LESF policies in the aforementioned Markovian setting. For completeness we present also some pure-simulation results illustrating the impact of a non Poisson-arrival assumption.

Let us now recollect a few *related works* on the performance evaluation of cellular networks. In early 80's, wireless cellular networks were carrying essentially voice calls, which require constant bit-rates (CBR) and are subject to admission control policies with blocking (at the arrival epoch) to guarantee these rates for calls already in service. An important amount of work has been done to propose efficient call admission policies [3–5]. Policies with admission conditions in the multi-rate linear form have been considered e.g. in [6–8].

Progressively, cellular networks started carrying also calls with variable bit-rates (VBR), used to transmit data files. The available resources are (fairly) shared between such calls and when the traffic demand increases, the file transfer delays increase as well, but (in principle) no call is ever blocked. These delays may be evaluated analytically using multi-rate linear resource constraint in conjunction with multi-class processor sharing models; cf e.g. [8, 9].

Recently, users may access multimedia streaming services through their mobile devices [10]. They are provided via CBR connections, essentially without admission control, but they tolerate temporary interruptions, when network congestions occur. One may distinguish two types of streaming traffic. In *real-time streaming (RTS)* (as e.g. in mobile TV), considered in this paper, the portions of the streaming content emitted during the time when the transmission to a given user is interrupted (is in outage) are definitely lost for him (unless a “secondary”, lower-rate streaming is provided during these periods). In *non-real-time streaming (NRTS)* (like e.g., video-on-demand, YouTube, Dailymotion, etc), a user starts playing back the requested multimedia content after some initial delay, required to deliver and buffer on the user device some initial portion of it. If further transmission is interrupted for some time making the user buffer content drop to zero (buffer starvation) then the play-back is stopped until some new required portion of the content is delivered. Several papers study the effect of the variability of the wireless channel on the performance of a single streaming call; see for e.g. [11], [12]. In [13] VBR transmissions and RTS are considered jointly in some analytical model, however the number and duration of outage periods are not evaluated. In [14] the tradeoff between the start-up delay and the probability of buffer starvation is analyzed in a Markovian queuing framework for NRTS streaming.

We do not consider any cell-load balancing; see [15] for

some recent work on this problem in the video streaming context. Also, [16, 17] consider some admission control policies to guarantee non-dropping of multimedia calls due to caller impatience and/or handoffs.

The remaining part of this paper is organized as follows. In Section II we will present our model for the evaluation of the quality of RTS in wireless cellular networks. Technical proofs of the results presented in this section are postponed to the Appendix, where they are given in a more general context. Section III specifies our model to be compliant with the LTE cellular networks specification and presents numerical results regarding the quality of RTS in these networks.

II. STREAMING IN WIRELESS CELLULAR NETWORKS

In this section we present a new stochastic model of RTS in cellular networks.

A. System assumptions

We consider the following scenario of multi-user streaming in a cellular network.

1) *Network layer*: Geographically distributed users wish to obtain down-link wireless streaming of some (typically video) content, contacting base stations of a network at random times, for random durations, requesting some fixed streaming bit-rates. We consider a uni-cast traffic (as opposed to the broadcast or multi-cast case), i.e.; the content is delivered to all users via private connections. Different classes of users (calls) need to be distinguished, regarding their radio channel conditions, requested streaming bit-rates and mean streaming times. Each user chooses one base station, the one with the smallest path-loss, independently of the configuration of users served by this station. Thus, we do not consider any load-balancing policy.

2) *Data layer — streaming policies*: If a given base station cannot serve all the users present at a given time, it temporarily stops streaming the requested content at the requested rate to users of some classes, according to some given policy (to be described), which is supposed to preserve a maximal subset of served users. We call these (classes of) users with the requested bit-rate temporarily denied *in outage*. The users in outage will not receive the part of the content which is emitted during their outage times (this is the principle of the RTS). We will also consider policies, which offer some “best-effort” streaming bit-rates for some classes of users in outage, thus allowing for example to keep receiving the requested content but of a lower quality. Users, which are (temporarily) denied even this lower quality of service are called *in deep outage*.

3) *Medium access*: In this paper we assume that users are connected to the serving antennas via orthogonal single-input-single-output (SISO) channels allowing for the peak-rate close to the theoretical Shannon’s bound in the additive white Gaussian noise (AWGN) model, with the (extra-cell) interference treated as noise.² We will also comment on how to model multiple-input-multiple-output (MIMO) and broadcast channels.

²Orthogonality of channels is an appropriate assumption for current LTE (Long Term Evolution) norm for cellular networks based on OFDMA, as well as for other multiple access techniques as FDA, TDMA, CDMA assuming perfect in-cell orthogonality, and even HDR neglecting the scheduler gain.

4) *Physical layer*: The quality of channel of a given user depends on the path-loss of the signal with respect to its serving base station, a constant noise, and the interference from other (non-serving) base stations. These three components determine its signal-to-interference-and-noise ratio (SINR). Both path-loss from the serving station and interference account for the distance and random propagation effects (shadowing). Our main motivation for considering a multi-class model is to distinguish users with different SINR values. In other words, even if we assume that all users require the same streaming times and rates, we still need a multi-class model due to (typically) different SINR's values of users in wireless cellular networks.

5) *Performance characteristics*: We will present and analytically evaluate performance of some (realistic) streaming policies in the context described above. We will be particularly interested in the following characteristics:

- fraction of time spent in outage and in deep outage during the typical call of a given class,
- number of outage incidents occurring during this call,
- mean throughput (average bit-rate) during such call, accounting for the requested bit-rates and for the “best-effort” bit-rate obtained during the outage periods.

B. Model description

In what follows we describe a mathematical model of the RTS that is an incarnation of a new, more general, stochastic service model with capacity sharing and interruptions presented and analyzed in the Appendix A. This is a single server model which allows to study the performance of one tagged base station of a multi-cellular network satisfying the above system assumptions. More details on how this model fits the multi-cell scenario will be presented in Section III.

1) *Traffic demand*: Consider $J \geq 1$ classes of calls (or, equivalently, users) characterized by different *requested streaming bit-rates* r_k , *wireless channel conditions* described by the signal-to-(extra-cell)-interference-and-noise ratio SINR_k with respect to the serving base-station³ and *mean requested streaming times* $1/\mu_k$, $k = 1, \dots, J$.

We assume that calls of class $k \in \{1, \dots, J\}$ arrive in time according to a Poisson process with intensity $\lambda_k > 0$ (number of call arrivals per unit of time, per base station) and stay in the system (keep requesting streaming) for independent times, having some *general distribution* with mean $1/\mu_k < \infty$.⁴ Different classes of calls are independent from each other. We denote by $X_k(t)$ the number of calls of a given class requesting streaming from a given BS at time t ; see Section A in the Appendix for a formal definitions of these variables in terms of arrival process and service times. Let $\mathbf{X}(t) = (X_1(t), \dots, X_J(t))$; we call it the (vector of) user configuration at time t . The stationary distribution π of $\mathbf{X}(t)$ coincides with the distribution of the vector $(X_1, \dots, X_J)$ of independent Poisson random variables with

means $\mathbf{E}[X_k] := \rho_k = \lambda_k/\mu_k$, $k = 1, 2, \dots, J$. We call ρ_k the *traffic demand* (per base station) of class k .

2) *Wireless resource constraints*: Users are supposed to be offered the requested streaming rates for the whole requested streaming times. However, due to limited wireless resources, for some configuration of users $\mathbf{X}(t)$, the requested streaming rates $\mathbf{r} = (r_1, \dots, r_J)$ may be not achievable. Following the assumption of orthogonal AWGN SISO wireless channels (with the (extra-cell) interference treated as noise) available for users of a given station, we assume that the requested rates are achievable for all calls present at time t if

$$X_k(t)r_k = \nu_k r_k^{\max}, \quad k = 1, \dots, J, \quad (1)$$

for some non-negative vector $(\nu_1, \dots, \nu_J)$, such that $\sum_{k=1}^J \nu_k \leq 1$, where

$$r_k^{\max} = \gamma W \log(1 + \text{SINR}_k) \quad (2)$$

is the maximal (peak) bit-rate of a user of class k , whose channel conditions are characterized by SINR_k . (The rate $r_k^{\max}$ is available to a user of class k if it is the only user served by the base station.) Here W is the frequency bandwidth and γ (with $0 < \gamma \leq 1$) is a coefficient telling how close a given coding scheme approaches the theoretical Shannon's bound (corresponding to $\gamma = 1$); cf [18, Th .9.1.1].⁵ Note that the assumption (1) corresponds to the situation, when users neither hamper nor assist each other's transmission. They use channels which are perfectly separated in time, frequency or by orthogonal codes, nevertheless sharing these resources.⁶

We can interpret the ratio between the requested and maximal bit-rates $\varphi_k = r_k/r_k^{\max}$ as the *resource demand* of a user of class k . Note that the configuration of users $\mathbf{X}(t)$ can be entirely served if and only if the total resource demand satisfies the constraint

$$\sum_{k=1}^J \varphi_k X_k(t) \leq 1. \quad (3)$$

This is a *multi-rate linear resource constraint*.

3) *Service policy*: If the requested streaming rates are not achievable for a given configuration of users $\mathbf{X}(t)$ present at time t , then some classes of users will be temporarily put in outage at time t , meaning that they will receive some smaller bit-rates (whose values are not guaranteed and may depend on the configuration $\mathbf{X}(t)$). These smaller, “best-effort” bit-rates may drop to 0, in which case we say that users are in deep-outage. Let us recall that the times at which users are in outage and deep outage do *not* alter the original streaming

⁵It was also shown in [19] that the performance of AWGN *multiple input multiple output* (MIMO) channel can be approximated by taking values of $\gamma \geq 1$. Another possibility to consider MIMO channel is to use the exact capacity formula given in [20].

⁶From information theory point of view, the orthogonality assumption is not optimal. In fact, the theoretically optimal performance is offered by the *broadcast channel* model. It is known that in the case of AWGN broadcast channel the rates $\mathbf{r}$ are (theoretically) achievable for the configuration $\mathbf{X}$ if (and only if) there exists a vector $(\nu_1, \dots, \nu_J)$, such that $\sum_{k=1}^J \nu_k \leq 1$ and

$$X_k r_k = W \log \left(1 + \frac{\nu_k}{1/\text{SINR}_k + \sum_{i=1}^{k-1} \nu_i} \right) \quad k = 1, \dots, J,$$

where the classes of users are numbered such that $\text{SINR}_1 \geq \text{SINR}_2 \geq \dots \geq \text{SINR}_J$; cf [21, Eq. 6.29].

³In this paper the interference is always caused only by non-serving base stations.

⁴All the results presented in this paper do not depend on the particular choice of the streaming time distributions. This property is often referred to in the queuing context as the insensitivity property.

times; i.e. the streaming content is not buffered, nor delayed during the outage periods.

We will define now a parametric family of service policies for which *classes with smaller resource demands have higher service priority*. In this regard, in the remaining part of the paper we assume (without loss of generality) that the resource demands of users from different classes are ordered $\varphi_1 < \varphi_2 < \dots < \varphi_J$.

a) *Least-effort-served-first policy*: For a given configuration of users $\mathbf{X} = \mathbf{X}(t)$ requesting streaming at time t , *least-effort-served-first policy with δ -margin (LESF(δ))* (for short) attributes the requested bit-rates to all users in classes $k = 1, \dots, K$, where

$$K = K^\delta(\mathbf{X}) = \max \{k \in \{1, \dots, J\} : \sum_{j=1}^{k-1} \varphi_j X_j + \varphi_k \sum_{j=k}^J X_j \mathbb{1}(\varphi_j \leq \varphi_k(1 + \delta)) \leq 1\}, \quad (4)$$

where $\mathbb{1}_A(x) = 1$ is the indicator function of set A and δ is a constant satisfying $0 \leq \delta \leq \infty$.

Remark 2.1: The LESF(0) policy is *optimal* in the following sense: given constraint (3) and the assumption that the classes with smaller resource demands have higher priority, this policy allows to serve the maximal subset of users present in the system. For the same reason any LESF(δ) policy with $\delta > 0$ is clearly sub-optimal. In order to explain the motivation for considering such policies, one needs to extend the model and explain what actually happens with classes of users which experience outage. In this regard, note that $C = \sum_{j=1}^K \varphi_j X_j \leq 1$ is the actual fraction of the server capacity consumed by the users which are not in outage. The remaining server capacity $1 - C$ (which is not needed to serve users in classes $1, \dots, K$) can be used to offer some “lower quality” service (e.g. streaming with lower video resolution, etc) to the users in classes $K + 1, \dots, J$ which are in outage. Note by (4) that the remaining server capacity under the policy LESF(δ) is at least

$$1 - C \geq \varphi_K \sum_{j=K+1}^J X_j \mathbb{1}(\varphi_j \leq \varphi_K(1 + \delta)).$$

Hence, the server accepting the class K as the least-priority class being “fully” served, leaves enough remaining capacity to be able to make the same effort (allocate service capacity φ_K) for all users in outage in classes whose service demand exceeds φ_K by no more than $\delta \times 100\%$. These latter users will not have “full” required service (since this requires more resources, $\varphi_j > \varphi_K$, for the full service) but only some “lower quality” service (to be specified in what follows). Consequently, one can conclude that policies LESF(δ) with $\delta > 0$, being sub-optimal, ensure some *fairness*, in the sense explained above. Clearly the policy LESF(∞) (i.e., with $\delta = \infty$) is the most fair, in the sense that it reserves enough remaining capacity to offers the “lower quality” service for *all* users in outage (no deep outage). Thus, we will call LESF(∞) the *LESF fair policy*.

b) *Best-effort service for users in outage*: We will specify now a natural model for the “best-effort” streaming bit-rates that can be offered for users in outage in association

with a given LESF(δ) policy. For $k > K = K^\delta(\mathbf{X})$ denote

$$r'_k = r'_k{}^\delta(\mathbf{X}) = r_k^{\max} \frac{1 - \sum_{j=1}^K X_j \varphi_j}{\sum_{j=K+1}^J X_j \mathbb{1}(\varphi_j \leq (1 + \delta)\varphi_K)} \quad (5)$$

if $\varphi_k \leq (1 + \delta)\varphi_K$ and 0 otherwise.

The rates $(r_1, \dots, r_K, r'_{K+1}, \dots, r'_J)$ are achievable for the configuration $\mathbf{X}$ under resource constraint (3). Note that users in classes j such that $\varphi_j > (1 + \delta)\varphi_K$ do not receive any positive bit-rate. We say, they are in *deep outage*. Finally, we remark that the service (5) is “resource fair” among users in outage but not in deep outage.

4) *Performance metrics*: Configuration of users $\mathbf{X}(t)$ evolves in time, it changes at arrival and departure times of users. At each arrival or departure epoch the base station applies the outage policy to the new configuration of users to decide which classes of users receive requested streaming rates and which are in outage (or deep outage).

Let us introduce the following characteristics of the *typical call (user)* of class $k = 1, \dots, J$.

- P_k denotes the *probability of outage at the arrival epoch for class k* . This is the probability that the typical call of this class is put in outage immediately at its arrival epoch.
- D_k denotes the *mean total time spent in outage during the typical call of class k* .
- M_k denotes the *mean number of outage incidents experienced during the typical call of class k* .

More formal definitions of these characteristics, as well as other *system* characteristics (as e.g. the intensity of outage incidents) is given in the Appendix. We also introduce two further characteristics related to the mean *throughput* obtained during the typical call of class $k = 1, \dots, J$.

- Denote by T_k the *mean throughput during the typical call of class k* . This is the mean bit-rate obtained during such a call, taking into account the bit-rate r_k when the call is not in outage and the best-effort bit rate r'_k obtained during the outage periods, averaged over call duration.
- Let T'_k be the *part of the throughput obtained during the outage periods of the typical call of class k* . This is the mean best-effort bit-rate of such call averaged over outage periods.

C. Model evaluation

1) *Results*: We will show how the performance metrics regarding outage incidents and duration, introduced in Section II-B4, can be expressed using probability distribution functions of some *linear* functionals of the random vector $X_1, \dots, X_J$ of independent Poisson random variables with parameters ρ_j , respectively. Recall that these random variables correspond to the number of calls of different classes present in the stationary regime of our streaming model.

Specifically, for given $\delta > 0$, $k = 1, \dots, J$ and $t \geq 0$ denote

$$F_k^\delta(t) := \mathbf{P} \left\{ \sum_{j=1}^k X_j \delta_j \varphi_j \leq t \right\}, \quad (6)$$

where $X_j^{\delta,k} = X_j$ for $j = 1, \dots, k-1$ and $X_k^{\delta,k} = \sum_{j=k}^J X_j \mathbb{1}(\varphi_j \leq \varphi_k(1+\delta))$.

The following results follow from the analysis of a more general model presented in the Appendix.

Proposition 2.2: *The probability of outage at the arrival epoch for user of class k is equal to*

$$P_k = 1 - F_k^\delta(1 - \varphi_k) \quad k = 1, \dots, J. \quad (7)$$

The mean total time spent in outage during the typical call of class k is equal to

$$D_k = \frac{P_k}{\mu_k} = \frac{1 - F_k^\delta(1 - \varphi_k)}{\mu_k} \quad k = 1, \dots, J. \quad (8)$$

The mean number of outage incidents experienced during the typical call of class k (after its arrival) is equal to

$$M_k = \frac{1}{\mu_k} \sum_{j=1}^J \lambda_j \left(F_k^\delta(1 - \varphi_k) - F_k^\delta(1 - \varphi_k - \varphi_j) \right) \quad k = 1, \dots, J. \quad (9)$$

Proof: Note first that the functions $F_k^\delta(t)$ defined in (6) allow one to represent the stationary probability that the configuration of users is in a state in which the LESF(δ) policy serves users of class k

$$F_k^\delta(1) = \mathbf{P} \left\{ \sum_{j=1}^k X_j^{\delta,k} \varphi_j \leq 1 \right\}.$$

In the general model described in the Appendix we denote this state by $\mathcal{F}_k$ and its probability by $\pi(\mathcal{F}_k)$. Thus $\pi(\mathcal{F}_k) = F_k^\delta(1)$. Moreover,

$$1 - F_k^\delta(1 - \varphi_k) = \mathbf{P} \left\{ \sum_{j=1}^k X_j^{\delta,k} \varphi_j > 1 - \varphi_k \right\}$$

is the probability that the steady state configuration of users appended with one user of class k is in the complement $\mathcal{F}_k'$ of the state $\mathcal{F}_k$, i.e., all users of class k are in outage (meaning $k > K^\delta(\mathbf{X}')$, where $\mathbf{X}' = (X_1, \dots, X_k + 1, \dots, X_J)$). Thus the expression (7) follows from Proposition A.3. Similarly (8) follows from Proposition A.4 and (9) follows from Proposition A.5. ■

Regarding the throughput characteristics, we have the following result.

Proposition 2.3: *The mean throughput during the typical call of class k is equal to*

$$T_k = r_k(1 - P_k) + T_k' = r_k F_k^\delta(1 - \varphi_k) + T_k',$$

where

$$T_k' = \mathbf{E} \left[r_k'^\delta(X_1, \dots, X_k + 1, \dots, X_J) \mathbb{1} \left(K^\delta(X_1, \dots, X_k + 1, \dots, X_J) < k \right) \right], \quad (10)$$

with the best-effort rate $r_k'(\cdot)$ given by (5) and the least-priority class $K^\delta(\cdot)$ begin served by the LESF(δ) policy given by (4), is the part of the throughput obtained during the outage periods.

Proof of this proposition is given in the Appendix.

Remark 2.4: Recall from (5) that the variable rates r_k' are obtained by the user of class k when he is in outage, i.e.,

$k > K$. They are non-null, $r_k' > 0$, only if $\varphi_k \leq (1 + \delta)\varphi_K$. In the case of equal requested rates r_k , the intersection of the two conditions $0 < r_k'$ and $k > K$ is equivalent to

$$(1 + \text{SINR}_K)^{1/(1+\delta)} - 1 \leq \text{SINR}_k \leq \text{SINR}_K. \quad (11)$$

2) Remarks on numerical evaluation: In order to be able to use the expressions given in (2.2) we need to evaluate the distribution functions $F_k^\delta(t)$. In what follows we show how this can be done using Laplace transforms. Regarding the throughput in outage T_k' , expressed in (10) as the expectation of a non-linear functional of the vector $(X_1, \dots, X_J)$, we will use Monte Carlo simulations to obtain numerical values for this expectation.

Denote by $\mathcal{L}_k^\delta(\theta) := \int_0^\infty e^{-\theta s} F_k^\delta(s) ds$ the Laplace transform of the function $F_k^\delta(t)$.

Fact 2.5: We have

$$\mathcal{L}_k^\delta(\theta) = \frac{1}{\theta} \exp \left[\sum_{j=1}^k \rho_j^{\delta,k} (e^{-\theta \varphi_j} - 1) \right],$$

where $\rho_j^{\delta,k} = \rho_j$ for $j = 1, \dots, k-1$ and $\rho_k^{\delta,k} = \sum_{j=k}^J \rho_j \mathbb{1}(\varphi_j \leq \varphi_k(1+\delta))$.

Proof: Note that for given $\delta > 0$, $k = 1, \dots, J$ the random variables $X_1^{\delta,k}, \dots, X_k^{\delta,k}$ are independent, of Poisson distribution, with parameters $\rho_1^{\delta,k}, \dots, \rho_k^{\delta,k}$, respectively. The result follows from [22, Proposition 1.2.2] and a general relation $\int_0^\infty e^{-\theta s} F(s) ds = \frac{1}{\theta} \int_0^\infty e^{-\theta s} F(ds)$. ■

The probabilities $F_k^\delta(\cdot)$ may be retrieved from $\mathcal{L}_k^\delta(\cdot)$ using standard techniques. For example [23, with the algorithm implemented by Hollenbeck [24] in Matlab]. In what follows we present a more explicit result based on the Bromwich contour inversion integral. In this regard, denote $\bar{\mathcal{L}}_k^\delta(\theta) = 1/\theta - \mathcal{L}_k^\delta(\theta)$ (which is the Laplace transform of complementary distribution function $1 - F_k^\delta(t)$). Also, denote by $\mathcal{R}(z)$ the real part of the complex number z .

Fact 2.6: We have

$$F_k^\delta(t) = 1 - \frac{2e^{at}}{\pi} \int_0^\infty \mathcal{R} \left(\bar{\mathcal{L}}_k^\delta(a + iu) \right) \cos ut \, du, \quad (12)$$

where $a > 0$ is an arbitrary constant.

Proof: See [25]. ■

Remark 2.7: As shown in [25], the integral in (12) can be numerically evaluated using the trapezoidal rule, with the parameter a allowing to control the approximation error. Specifically, for $n = 0, 1, \dots$ define

$$h_n(t) = h_n(t; a, k, \delta) := \frac{(-1)^n e^{a/2}}{t} \mathcal{R} \left(\bar{\mathcal{L}}_k^\delta \left(\frac{a + 2n\pi i}{2t} \right) \right),$$

$S_n(t) := \frac{h_0(t)}{2} + \sum_{i=1}^n h_i(t)$, and $S(t) = \lim_{n \rightarrow \infty} S_n(t)$. Then $|F_k^\delta(t) - (1 - S(t))| \leq e^{-a}$. Finally, the (alternating) infinite series $S(t)$ can be efficiently approximated using for example the Euler summation rule

$$S(t) \approx \sum_{i=0}^M \binom{M}{i} 2^{-M} S_{N+i}(t)$$

with a typical choice $N = 15$, $M = 11$.

Remark 2.8: The expression (9) for the mean number of outage incidents involves a sum of potentially big number of

terms $F_k^\delta(1 - \varphi_k) - F_k^\delta(1 - \varphi_k - \varphi_j)$, $j = 1, \dots, J$, which are typically small, and which are evaluated via the inversion of the Laplace transform. Consequently the sum may accumulate precision errors. In order to avoid this problem we propose another numerical approach for calculating M_k . It consists in representing M_k equivalently to (9) as

$$M_k = \frac{F_k^\delta(1 - \varphi_k)}{\mu_k} \sum_{j=1}^J \lambda_j b_k(j) \quad k = 1, \dots, J. \quad (13)$$

where

$$b_k(j) = \frac{F_k^\delta(1 - \varphi_k) - F_k^\delta(1 - \varphi_k - \varphi_j)}{F_k^\delta(1 - \varphi_k)} \quad (14)$$

Let k and δ be fixed. Recall the definition of $F_k^\delta(t)$ in (6) and note that the expression (14) may be written as

$$b_k(j) = \frac{\mathbf{P}(X \in \mathcal{F}, X + \epsilon_j \notin \mathcal{F})}{\mathbf{P}(X \in \mathcal{F})}$$

where $\mathcal{F} = \mathcal{F}(k) = \{X \in \mathbb{R}^J : \sum_{j=1}^k X_j^{\delta,k} \varphi_j \leq 1 - \varphi_k\}$. The above expression may be seen as the blocking probability for class j in a classical multi-class Erlang loss system with the admission condition $X \in \mathcal{F}$. Consequently, $b_k(\cdot)$ may be calculated by using the *Kaufman-Roberts algorithm* [1, 2] and plugged into (13). Note that by doing this we still need to calculate $F_k^\delta(1 - \varphi_k)$ however avoid summing of J differences of these functions as in (9).

III. QUALITY OF REAL-TIME STREAMING IN LTE

In this section we will use the model developed in Section II to evaluate the quality of RTS in LTE networks. This single-server (base station) model will be used to study the performance of one tagged base station of a multi-cellular network under the following assumptions:

- We assume a regular hexagonal lattice of base stations on a torus. This allows us to consider the tagged base station of the network as a typical one.
- Homogeneous (in space and time) Poisson arrivals on the torus are marked by i.i.d. (across users and base stations) variables representing their shadowing with respect to different base stations. These variables, together with independent user locations determine their serving (strongest) base stations. A consequence of the independence of users locations and shadowing variables is that the arrivals served by the tagged base station form an independent thinning of the total Poisson arrival process to the torus and thus a Poisson process too. Uniform distribution of user locations and identical distribution of their shadowing variables imply that the intensity of the arrival process to the tagged base station is equal to the total arrival intensity to the torus divided by the number of stations. Moreover, the distribution of the SINR of the typical user of the tagged base station coincides with the distribution of the typical user of the whole network.
- The intensity of arrivals of some particular (SINR)-class to the tagged base station is equal to the total intensity of arrivals to the tagged cell times the probability of the

random SINR of the typical user being in the SINR-interval corresponding to this class.

- We consider the “full interference” scenario, i.e., that all base stations emit the signal with the constant power, regardless of the number of users they serve (this number can be zero). This makes the interference, and hence the service rates, of users of a given base station independent of the service of other base stations (decouples the service processes of different base stations).

A. LTE model and traffic specification

1) *SINR distribution*: Recall that the main motivation for considering a multi-class model was the necessity to distinguish users with different radio conditions, related to different values of the SINR they have with respect to the serving base stations. In order to choose representative values of SINR in a given network and to know what fraction of users experience a given value, we need to know the (*spatial*) *distribution of the SINR* (with respect to the serving base station) experienced in this network (possibly biased by the spatial repartition of arrivals of streaming calls). This distribution can be obtained from real-network measurements, simulations or analytic evaluation of an appropriate spatial, stochastic model.⁷ In this paper we will use the distribution of SINR obtained from the simulation compliant with the 3GPP recommendation in the so-called calibration case (to be explained in what follows). At present, assume simply, that we are given a cumulative distribution function (CDF) of the SINR expressed in dB, $F(x) := \mathbf{P}\{10 \log_{10}(\text{SINR}) \leq x\}$, obtained from either of these methods. In other words, $F(x)$ represents the fraction of mobile users in the given network which experience the SINR (expressed in dB) not larger than x .

Consider a discrete probability mass function

$$p_k := F\left(\frac{x_{k+1} + x_k}{2}\right) - F\left(\frac{x_k + x_{k-1}}{2}\right) \quad k = 1, 2, \dots, J, \quad (15)$$

with $x_0 = -\infty$, $x_{J+1} = \infty$. We define the class $k = 1, \dots, J$ of users as all users having the SINR expressed in dB in the interval $[(x_k + x_{k-1})/2, (x_{k+1} + x_k)/2]$, and approximate their SINR by the common value $\text{SINR}_k = 10^{x_k/10}$. Clearly p_k is the fraction of mobile users in the given network which experience the SINR close to SINR_k . Hence, in the case of a homogeneous streaming traffic (the same requested streaming rates and mean streaming times, which will be our default assumption in the numerical examples) we can assume the intensity of arrivals λ_k of users of class k to be equal to $\lambda_k = p_k \lambda$ where $\lambda = \sum_{i=k}^J \lambda_k$ is the total arrival intensity (per unit of time per serving base station) to be specified together with the CDF F of the SINR.

a) *CDF of the SINR for 3GPP recommendation*: We obtain the CDF F of SINR from the simulation compliant with the 3GPP recommendation in the so-called calibration case, (compare to [29, Figure A.2.2-1(right)]). More precisely, we consider the geometric pattern of BS placed on the 6×6 hexagonal lattice. In the middle of each hexagon there are

⁷For this latter possibility, we refer the reader to a recent paper on Poisson modeling of real cellular networks subject to shadowing [26], as well as to [27], completed in [28], where the distribution of the SINR in Poisson networks is evaluated explicitly.

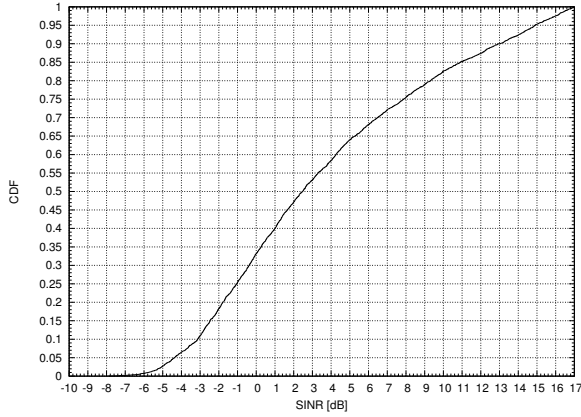


Fig. 1. Cumulative distribution function of the SINR obtained according to 3GPP specification; see Section III-A1a. An abrupt transition of the CDF to 1 at $\text{SINR} = 17\text{dB}$ is due to the cell sectorization: each mobile is interfered by each of the two antennas co-located with its serving antenna on the same site (and serving the different sectors) with the power equal to at least 1% of the power received from the serving BS. Therefore the signal to interference ratio is at most $0.5 \times 10^{-2} = 17\text{dB}$.

three symmetrically oriented BS antennas, which gives a total of 108 BS antennas. The distance between the centers of two neighboring hexagons is 0.5 km. Each BS antenna is characterised by the following horizontal pattern $A(\phi) = -\min(12(\phi/\theta)^2, A_m)$, where ϕ is the angle in degrees, with $\theta = 70^\circ$, $A_m = 20\text{dB}$, and uses transmission power $P = 60\text{dBm}$ (including omnidirectional gain of 14dBi). The distance-loss model (corresponding to the frequency carrier 2GHz) is $L(r) = 128.1 + 37.6 \log_{10}(r)[\text{dB}]$ where r is the distance in km. A supplementary penetration loss of 20dB is added. The shadowing is modeled as a log-normal random variable of mean one and logarithmic standard deviation of deviation 8dB, cf [30]. The noise power equals -95dBm (which corresponds to a system bandwidth of 10MHz, a noise floor of -174dBm/Hz and a noise figure of 9dB). In order to obtain the empirical CDF of the SINR we generate 3600 random user locations uniformly in the network (100 user locations per hexagon on average). Each user is connected to the antenna with the strongest received signal (smallest propagation-loss including distance, shadowing and antenna pattern) and the SINR is calculated. The obtained empirical CDF F of the SINR is shown on Figure 1.

2) *Link characteristics*: 3GPP shows in [31, §A.2] that there is a 25% gap between the practical coding schemes and the Shannon's limit for the AWGN channel. Moreover, some of the transmitted bits are used for signaling, which induces a supplementary capacity loss of about 30% (see [32, §6.8]). This made us assume $\gamma = 0.5 (\approx 0.75(1 - 0.3))$ in (2). The system bandwidth is $W = 10\text{MHz}$.

3) *Streaming traffic*: We assume that all calls require the same streaming rate $r_k = 256\text{ kbit/s}$ and have the same streaming call time distribution. We split them into $J = 100$ user classes characterized by values of the SINR falling into different intervals regularly approximating the SINR domain from $x_1 = -10\text{dB}$ to $x_J = 17\text{dB}$ as explained in Section III-A1. In our performance evaluation we will consider two values of the spatially uniform traffic demand: 900 and 600 Erlang/km².

(All results presented in what follows do not depend on the mean streaming time but only on the traffic demand). Consequently, k th class traffic demand per unit of surface is equal to, respectively, $p_k \times 900$ and $p_k \times 600$ Erlang/km², where p_k are given by (15). Multiplying by the surface served by one base station equal to $\sqrt{3} \cdot (0.5\text{ km})^2 / 6 \approx 0.0722\text{ km}^2$ we obtain the traffic demand per cell, per class, equal to $\rho_k = p_k \times 900 \times 0.0722 \approx p_k \times 64.9$ Erlang and $\rho_k = p_k \times 600 \times 0.0722 \approx p_k \times 43.3$ Erlang, respectively, for the two studied scenarios.

B. Performance evaluation

Assuming the LTE and traffic model described above, we consider now streaming policies $\text{LESF}(\delta)$ defined in Section II-B3. Recall that in doing so, we assume that users are served by the antenna offering the smallest path-loss, and dispose orthogonal down-link channels, with the maximal rates $r_k^{\max}$ depending on the value of the SINR (interference comes from non-serving BS) characterizing class k . Roughly speaking, $\text{LESF}(\delta)$ policy assigns the total requested streaming rate $r_k = 256\text{ kbit/s}$ for the maximal possible subset of classes in the order of decreasing SINR, leaving some capacity margin to offer some “best-effort” streaming rates for (some) users remaining in outage. These streaming rates r'_k given by (5) depend on the current configuration of users and are non-zero for users with SINR within the interval $(1 + \text{SINR}_K)^{1/(1+\delta)} - 1 \leq \text{SINR} \leq \text{SINR}_K$, where SINR_K is the minimal value of SINR for which users are assigned the total requested streaming rate; cf Remark 2.4. In particular, $\text{LESF}(0)$, called the *optimal* policy, leaves no capacity margin for users in outage, while $\text{LESF}(\infty)$, called the *fair* one, offers a “best-effort” streaming rate for all users in outage at the price of assigning the full requested rate 256kbit/s to a smaller number of classes (higher value of SINR_K)⁸. In what follows, we use our results of Section II-C to evaluate performance of these streaming policies in the LTE network model.

1) *Outage time*: Figure 2 shows the mean time of the streaming call spent in outage normalized by call duration, $\mu_k D_k$, evaluated using (8), in function of the SINR value characterizing class k , for the traffic 900 Erlang/km² and different policies $\text{LESF}(\delta)$. Figure 3 shows the analogous results assuming traffic load of 600 Erlang/km². The main observations are as follows:

- All LESF policies exhibit a cut-off behaviour: the fraction of time in outage drops rapidly from 100% to 0% when SINR transgresses some critical values. This cut-off is more strict for the optimal policy.
- For the traffic of 900 Erlang/km², users with $\text{SINR} \geq 3\text{dB}$ are practically never in outage, when the optimal policy is used. The same holds true for users with $\text{SINR} \geq 13\text{dB}$, when the fair policy is used.
- When the traffic drops to 600 Erlang/km², these critical values of SINR decrease by 2dB and 5dB, respectively, for the optimal and the fair policy. Note that the fair policy is more sensitive to higher traffic load.

⁸The LESF fair policy seems to be adopted in some implementations of the LTE.

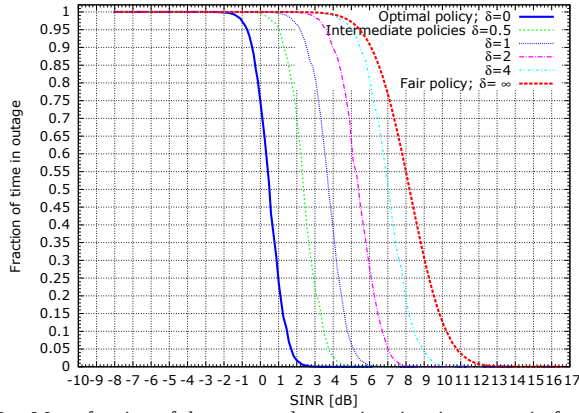


Fig. 2. Mean fraction of the requested streaming time in outage, in function of the user SINR for different policies LESF(δ); traffic 900 Erlang/km².

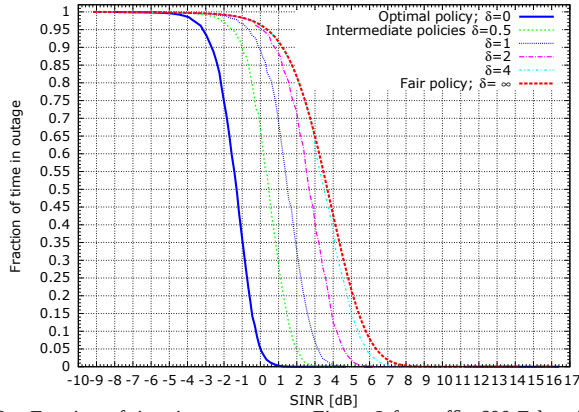


Fig. 3. Fraction of time in outage as on Figure 2 for traffic 600 Erlang/km².

2) *Number of outage incidents*: Figure 4 shows the mean number of outage incidents per streaming call, M_k evaluated using (9), in function of the SINR value characterizing class k , for the traffic 900 Erlang/km² and different policies LESF(δ). (Recall that we assume the same streaming time distribution for all users, and hence $\lambda_j/\mu_k = \rho_j$ making the expression in (9) depend only on the vector of traffic demand per class.) Figure 5 shows the analogous results assuming traffic demand of 600 Erlang/km². The main observations are as follows:

- For all policies, the number of outage incidents (during the service) is non-zero only for users with the SINR close to the critical values revealed by the analysis of the outage times. Users with SINR below these values are constantly in outage while users with SINR above them never in outage.
- More fair policies generate slightly more outage incidents. The worst values are 2 to 2.2 interruptions per call for the optimal policy, depending on the traffic value, and 2.4 to 3 interruptions per call for the fair policy.

Studying outage times and outage incidents we do not see apparent reasons for considering fair policies. This motivates our study of the best-effort service in outage.

3) *The role of the “best effort” service*: Figure 6 shows the fraction of time spent in deep outage in function of the SINR, assuming traffic 900 Erlang/km². These values should be compared to the fraction of time spend in outage (for convenience copied on Figure 6 from Figure 2). Recall, users in outage do not receive the full requested streaming rate (assumed 256kbit/s in our example), however they do receive

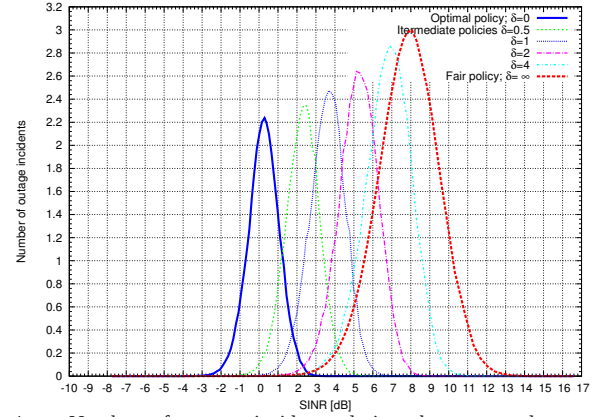


Fig. 4. Number of outage incidents during the requested streaming time, in function of the user SINR for different policies LESF(δ); traffic 900 Erlang/km².

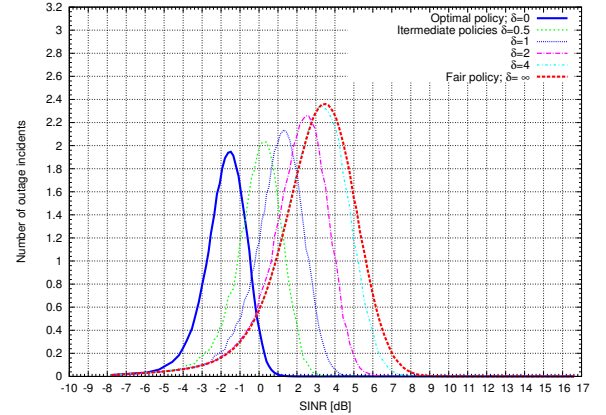


Fig. 5. Number of outage incidents as on Figure 4 for traffic 600 Erlang/km².

some non-null “best effort” rates given by (5), unless they are in deep outage — have SINR too small; cf Remark 2.4. Considering users in outage but not in deep outage as “partially satisfied”, increasing fairness margin δ allows to (at least) partially satisfy users with decreasing SINR values. Obviously the level of the “partial satisfaction” depends on the throughput obtained in outage periods, which is our quantity of interest on Figure 7. It shows also two curves for all policies LESF(δ) assuming traffic 900 Erlang/km². The upper ones represent the mean total throughput realized during the service, normalized to its maximal value; i.e., $T_k/(256\text{kbit/s})$, in function of the SINR value characterizing class k . The fractions of this throughput realized during outage periods, $T'_k/(256\text{kbit/s})$, are represented by the lower curves.

Figures 7 and 6 teach us that the role of the LESF(δ) policies with $\delta > 0$ may be two-fold.

- LESF(δ) policies with small values of δ , e.g. $\delta = 0.5$, improve “temporal homogeneity” of service with respect to the optimal policy, for users having SINR near the critical value. For example, a user having SINR equal to 1dB is served by the optimal policy during 80% of time with the full requested streaming rate (cf. Figure 6). However, for the remaining 20% of time it does not receive any service (deep outage, rate 0b/s). The policy LESF(0.5) offers to such a user 80% of the requested streaming rate during the whole streaming time (cf. Figure 7), with no deep outage periods (cf. Figure 6). The price for this is that a slightly higher SINR is required

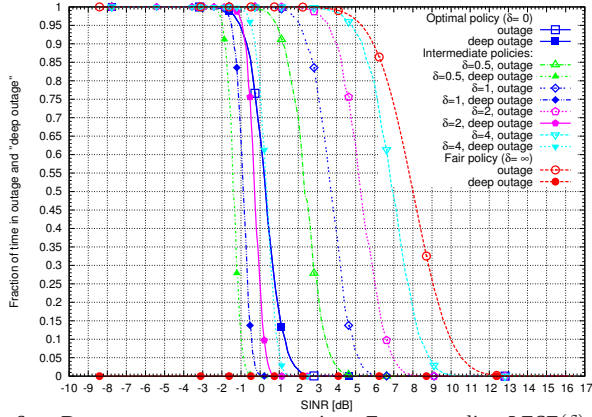


Fig. 6. Deep outage versus outage time. For any policy $\text{LESF}(\delta)$, with $0 < \delta < \infty$, the left curve of a given style represents the fraction of time spent in deep outage. The right curve of a given style recalls the fraction of time spent in outage (already plotted on Figure 2). The optimal policy ($\delta = 0$) does not offer any “best effort” service. The fair policy ($\delta = \infty$) offers this service for all users in outage.

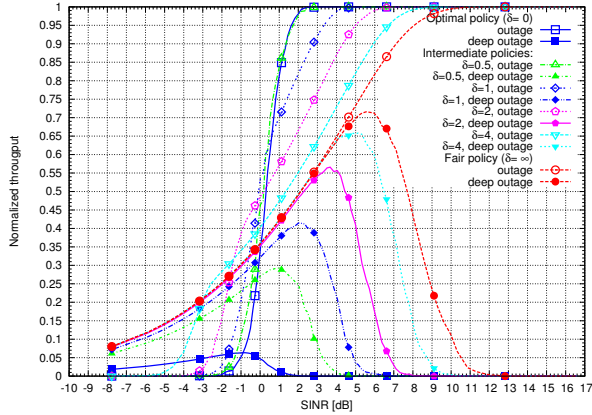


Fig. 7. Mean total throughput normalized to its maximal value 256 kbit/s obtained during the service time (upper curves) and its fraction obtained when a user is in outage (lower curves) for different policies $\text{LESF}(\delta)$ traffic 900 Erlang/km².

to receive the full requested streaming rate (at least 5 dB, instead of 3 dB for the optimal policy).

- The fair policy $\text{LESF}(\infty)$ improves the spatial homogeneity of service. It leaves no user in deep outage, however a much larger $\text{SINR} = 13 \text{ dB}$ is required for not to be in outage (cf. Figure 6). Moreover, the throughput of all users in outage but not in deep outage is substantially reduced e.g. from 80% to 40% for $\text{SINR} = 1 \text{ dB}$, with respect to some intermediate $\text{LESF}(\delta)$ policies (with $0 < \delta < \infty$). These intermediate policies can offer an interesting compromise between the optimality and fairness.

4) *Impact of a non-Poisson-arrivals:* Recall that the performance analysis of the model presented in this paper is insensitive to distribution of the requested streaming times. In this section we will briefly study the impact of a non Poisson-arrival assumption. In this regard we simulate the dynamics of the model with deterministic inter-arrival times (with all other model assumptions as before) and estimate the mean fraction of time in outage $\mu_k D_k$ and mean number of outage incidents M_k for each class k . For the comparison, as well as for the validation of the theoretical work, we perform also the simulation of the model with Poisson arrivals. The results are plotted on Figures 8, 9 and 10, 11. Observe first that the

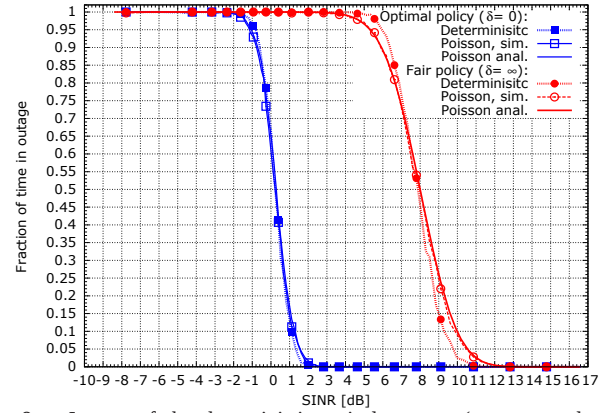


Fig. 8. Impact of the deterministic arrival process (as compared to the Poisson one) on the mean fraction of the requested streaming time in outage, for the optimal and fair policy; traffic 900 Erlang/km².

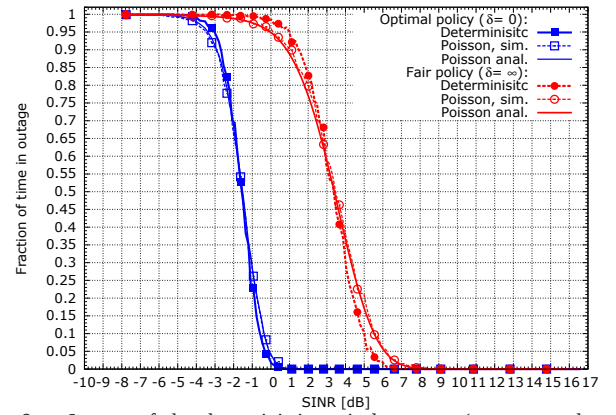


Fig. 9. Impact of the deterministic arrival process (as compared to the Poisson one) on the mean fraction of the requested streaming time in outage, for the optimal and fair policy; traffic 600 Erlang/km².

simulations of the Poisson model confirm the results of the theoretical analysis. Regarding the impact of the deterministic inter-arrival times a (somewhat expected) fact is that the optimal policy remains optimal regarding the fraction of time spent in the outage and the number of outage incidents. Another, less evident, observation is that the deterministic inter-arrivals (more regular than in the Poisson case) do *not* improve the situation for *all* classes of users. In fact, users with small values of the SINR have smaller fraction of time in outage under Poisson arrival assumption than in the deterministic one! This is different from what we can observe for the blocking probability for the classical Erlang’s loss model; cf e.g. [33, Figure 8]. Moreover, the deterministic arrivals increase the number of outage incidents for intermediate values of the SINR and decrease for extreme ones, especially with the fair policy. Concluding these observations one can say however, that the differences between Poisson and deterministic are not very significant and hence we conjecture that the Poisson model can be used to approximate a more realistic arrival traffic.

IV. CONCLUSIONS

In this paper, a real-time streaming (RTS) traffic, as e.g. mobile TV, is analyzed in the context of wireless cellular networks. An adequate stochastic model is proposed to evaluate user performance metrics, such as frequency and number of interruptions during RTS calls in function of user

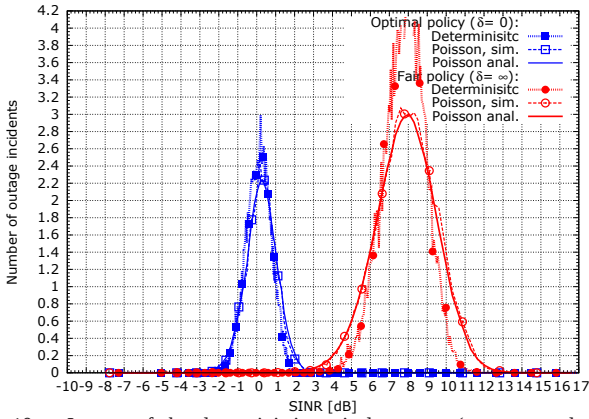


Fig. 10. Impact of the deterministic arrival process (as compared to the Poisson one) on the mean number of outage incidents for the optimal and fair policy; traffic 900 Erlang/km².

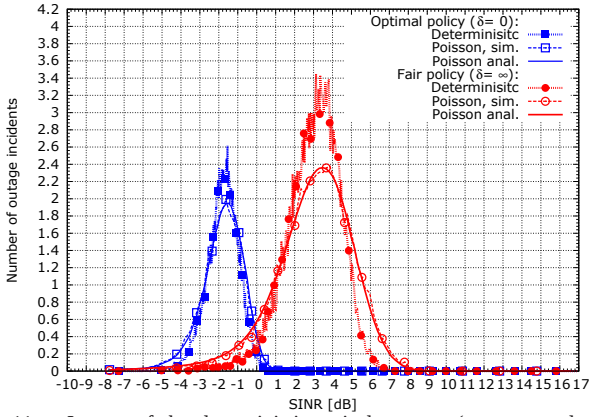


Fig. 11. Impact of the deterministic arrival process (as compared to the Poisson one) on the mean number of outage incidents for the optimal and fair policy; traffic 600 Erlang/km².

radio conditions. Despite some fundamental similarities to the classical Erlang loss model, a new model was required for this type of service, where the service denials are not definitive for a given call, but only temporal – having the form of, hopefully short, interruptions (outage) periods. Our model allows to take into account realistic implementations of the RTS service, e.g. in the LTE networks. In this latter context, several numerical demonstrations are given, presenting the quality of service metrics in function of user radio conditions.

APPENDIX

A GENERAL REAL-TIME STREAMING (RTS) MODEL

In this section we will present a general stochastic model for real-time streaming. An instantiation of this model was used in the main body of the paper to evaluate the real-time streaming in wireless cellular networks. This model comprises Markovian, multi-class process of call arrivals and their independent, arbitrarily distributed streaming times. These calls are served by a server whose service capacity is limited. Depending on numbers of calls of different classes present in the system, the server may not be able to serve some classes of users. If such a congestion occurs, these classes are temporarily denied the service, until the next call arrival or departure, when the situation is reevaluated. These service denial periods, called outage periods, do not alter the call sojourn times in the system. Our model allows for a very general service (outage) policy saying which classes of users are temporarily denied the

service due to insufficient service capacity. We will evaluate key characteristics of this model using the formalism of point processes and their Palm theory, often used in the modern approach to stochastic networking [34]. Specifically, we are interested in the intensity of outage incidents, the mean inter-outage times and the outage durations of a given class, seen from the server perspective, as well as the probability of outage at the arrival epoch, mean total time in outage and mean number of outage incidents experienced by a typical user of a given class. The expressions developed for these characteristics involve only stationary probabilities of the (free) traffic demand process, which in our case is a vector of independent Poisson random variables. Recall that such a representation is possible e.g. for the well known Erlang-B formula, giving the blocking probability in the classical (possibly multi-class) Erlang's loss model. Indeed, our model can be seen as an extension of the classical loss model, where the losses (i.e., service denials) are not definitive for a given call, but only temporal — having the form of outage periods.

A. Traffic demand

Consider $J \geq 1$ classes of users identified with calls. We assume that users of class $k \in \{1, \dots, J\}$ arrive in time according to a Poisson process $N_k = \{T_n^k : n\}$ ⁹ with intensity $\lambda_k > 0$ and stay in the system for independent requested streaming times W_n^k having some general distribution with mean $1/\mu_k < \infty$. All the results presented in what follows do not depend on the particular choice of the streaming time distributions — the property called in the queueing-theoretic context *insensitivity property*. Denote by $\tilde{N}_k = \{(T_n^k, W_n^k) : n\}$ the process of arrival epochs and streaming times (call durations) of users of class k . We assume that $\tilde{N}_k$ are independent across $k = 1, \dots, J$. Denote by $X_k(t) = \sum_n \mathbb{1}_{[T_n^k, T_n^k + W_n^k)}(t)$ the number of users of class k present in the system at time t and let $\mathbf{X}(t) = (X_1(t), \dots, X_J(t))$; we call it the (vector of) user configuration at time t . The stationary distribution π of $\mathbf{X}(t)$ coincides with the distribution of a vector of independent Poisson random variables $(X_1, \dots, X_J)$ with means $\mathbf{E}[X_k] := \rho_k = \lambda_k/\mu_k$, $k = 1, 2, \dots, J$. We call ρ_k the *traffic demand* of class k .

We adopt the usual convention for the numbering of the arrival epochs $T_0^k \leq 0 < T_1^k$. The same convention is used with respect to all point processes denoting some time epochs.

B. Resource constraints and outage policy

For class $k = 1, \dots, J$, let a subset of user configurations $\mathcal{F}_k \subset \mathbb{N}^J$ be given, where $\mathbb{N} = \{0, 1, \dots\}$, such that all X_k users of class k present in the configuration $\mathbf{X} = (X_1, \dots, X_k, \dots, X_J)$ are served if and only if $\mathbf{X} \in \mathcal{F}_k$ and no user of class k is served (we say it is in *outage*) if $\mathbf{X} \notin \mathcal{F}_k$. We call $\mathcal{F}_k$ the *kth class (service) feasibility set*. Denote by $\pi_k = \pi(\mathcal{F}_k)$ the probability that the stationary configuration of users is in *kth class feasibility set*.

⁹The time instants T_n^k are used only in the Appendix and should not be confused with T_k denoting in the main stream of the paper (and in the proof of Proposition 2.3 at the end of the Appendix) the mean throughput of user in class k .

We assume that, upon each arrival or departure of a user, the system updates its decision and, for any class k , it assigns the service to all users of class k if the updated configuration of users is in $\mathcal{F}_k$. All users of any class j for which the updated configuration is in $\mathcal{F}'_k = \bar{\mathbb{N}}^J \setminus \mathcal{F}_k$ will be placed in outage (at least) until the next user arrival or departure.

In what follows we will assume that no user departure can cause outage of any class of users i.e., switch a given configuration from $\mathcal{F}_k$ to $\mathcal{F}'_k$. (However a user departure may make some class j switch from $\mathcal{F}'_j$ to $\mathcal{F}_j$.)

Denote by $\tilde{X}_k(t) := X_k(t) \mathbb{1}_{\mathcal{F}_k}(\mathbf{X}(t))$ the number of users of class k *not in outage* at time t . Denote by $\tilde{\mathbf{X}}(t) = (\tilde{X}_1(t), \dots, \tilde{X}_J(t))$ the configuration of users *not in outage* at time t .

C. Performance metrics

In what follows we will be interested in the following characteristics of the model.

1) *Virtual system metrics*: During its time evolution, the user configuration $\mathbf{X}(t)$ alternates visits in the feasibility set $\mathcal{F}_k$ and its complement $\mathcal{F}'_k$, for each class $k = 1, \dots, J$. We are interested in the expected visit durations in these sets as well as the intensities (frequencies) of the alternations. More formally, for each given $k = 1, \dots, J$, we define the point process $B_k := \{\tau_n^k : n\}$ of exit epochs of $\mathbf{X}(t)$ from $\mathcal{F}_k$; i.e., all epochs t such that $(\mathbf{X}(t-), \mathbf{X}(t)) \in \mathcal{F}_k \times \mathcal{F}'_k$ (with the convention $\tau_0^k \leq 0 < \tau_1^k$). These are epochs when all users of class k present in the system (if any) have their service interrupted.

Denote by $\sigma_n'^k := \sup\{t - \tau_n^k : \mathbf{X}(s) \in \mathcal{F}'_k \forall s \in [\tau_n^k, t)\}$ the duration of the n th visit of the process $\mathbf{X}(t)$ in $\mathcal{F}'_k$ and by $\sigma_n^k := \tau_{n+1}^k - \tau_n^k - \sigma_n'^k$ the duration of the n th visit of the process $\mathbf{X}(t)$ in $\mathcal{F}_k$. We define for each class $k = 1, \dots, J$:

- *The intensity of outage incidents of class k* , i.e., the mean number of outage incidents of this class per unit of time

$$\Lambda_k := \lim_{T \rightarrow \infty} \frac{1}{T} \sum_n \mathbb{1}_{[0, T)}(\tau_n^k).$$

Obviously Λ_k is also the intensity of entrance to the k th class feasibility set $\mathcal{F}_k$.

- *The mean service time between two outage incidents of class k*

$$\bar{\sigma}_k := \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \sigma_n^k.$$

- *The mean outage duration of class k*

$$\bar{\sigma}'_k := \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \sigma_n'^k.$$

Note that the above metrics characterize a “virtual” quality of the service, since some visits in $\mathcal{F}_k$ and $\mathcal{F}'_k$ may occur when there is no k th class user in the system (in the latter case the outage of this class is not experienced by any user).

2) *User metrics*: We adopt now a user point of view on the system. We define for each class $k = 1, \dots, J$:

- *The probability of outage at the arrival epoch for user of class k*

$$P_k = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \mathbb{1}_{\mathcal{F}'_k}(\mathbf{X}(T_n^k)).$$

- *The mean total time in outage of user of class k*

$$D_k = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \int_{[T_n^k, T_n^k + W_n^k)} \mathbb{1}_{\mathcal{F}'_k}(\mathbf{X}(t)) dt.$$

- *The mean number of outage incidents experienced by user of class k after its arrival*

$$M_k = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \sum_m \mathbb{1}_{(T_n^k, T_n^k + W_n^k)}(\tau_m^k).$$

Note that eventual outage experienced at the arrival of a given user is not counted in M_k . The mean total number of outage incidents (including possibly at the arrival epoch) is hence $P_k + M_k$.

For a given class $k = 1, \dots, J$, denote by $\varepsilon_k = (0, \dots, 1, \dots, 0) \in \bar{\mathbb{N}}^J$ the unit vector having its k th component equal to 1. Hence $\mathbf{x} + \varepsilon_k$ represents adding one user of class k to the configuration of users $\mathbf{x} \in \bar{\mathbb{N}}^J$. Denote by $\mathbf{P}$ the probability under which $\{\mathbf{X}(t) : t\}$ is stationary and by $\mathbf{E}$ the corresponding expectation. Recall that $\pi\{\mathbf{x} \in \cdot\} = \mathbf{P}\{\mathbf{X}(t) \in \cdot\}$ is the distribution of the stationary configuration of users $\mathbf{X}(t)$ (it corresponds to independent Poisson variables of mean ρ_k).

D. General results

We present first results regarding the virtual system metrics. These results will be next used to evaluate the user metrics.

Lemma A.1: *The intensity of outage incidents of class k is $\mathbf{P}$ -almost surely equal to*

$$\Lambda_k = \sum_{j=1}^J \lambda_j \pi\{\mathbf{x} \in \mathcal{F}_k, \mathbf{x} + \varepsilon_j \in \mathcal{F}'_k\} \quad k = 1, \dots, J.$$

Proof: Let $N = \sum_{j=1}^J N_j$ be the point process counting the arrival times of users of all classes. By independence, N is the Poisson point process of intensity $\lambda = \sum_{j=1}^J \lambda_j$. Then, by the ergodicity of the process $\{\mathbf{X}(t) : t\}$ and the fact that the exits from $\mathcal{F}_k$ can take place only at some user arrival epoch we have by the Campbell’s formula [cf. e.g. 34, Equation (1.2.19)¹⁰],

$$\begin{aligned} \Lambda_k &= \mathbf{E} \left[\int_{[0,1)} \mathbb{1}_{\mathcal{F}_k \times \mathcal{F}'_k}(\mathbf{X}(t-), \mathbf{X}(t)) N(dt) \right] \\ &= \lambda \mathbf{P}_N^0 \{\mathbf{X}(0-) \in \mathcal{F}_k, \mathbf{X}(0) \in \mathcal{F}'_k\}, \end{aligned}$$

where $\mathbf{P}_N^0$ designates the Palm probability associated to N (which is, roughly speaking, the conditional probability given an arrival at time 0). By PASTA (Poisson Arrivals See Time Averages) property [cf. 34, Equation (3.3.4)] the configuration

¹⁰with $Z_n := (\mathbf{X}(T_n-), \mathbf{X}(T_n))$ and $f(t, z) = \mathbb{1}_{[0,1)}(t) \mathbb{1}_{\mathcal{F}_k \times \mathcal{F}'_k}(z)$

of users $\mathbf{X}(0-)$ under $\mathbf{P}_N^0$ has distribution π . Moreover, $\mathbf{X}(0) = \mathbf{X}(0-) + \varepsilon_\xi$ where $\xi \in \{1, \dots, J\}$ is under $\mathbf{P}_N^0$ independent of $\mathbf{X}(0-)$ and takes value j with probability λ_j/λ . This completes the proof. ■

Lemma A.2: *The mean service time between two outage incidents and the mean outage duration of class k are $\mathbf{P}$ -almost surely equal to, respectively,*

$$\bar{\sigma}_k = \frac{\pi(\mathcal{F}_k)}{\Lambda_k}, \quad \bar{\sigma}'_k := \frac{\pi(\mathcal{F}'_k)}{\Lambda_k} \quad k = 1, \dots, J,$$

where Λ_k is given in Lemma A.1.

Proof: First we prove the expression for $\bar{\sigma}_k$. By ergodicity $\bar{\sigma}_k = \mathbf{E}_{B_k}^0[\sigma_0^k]$ $\mathbf{P}$ -almost surely, where $\mathbf{E}_{B_k}^0$ designates the expectation with respect to the Palm probability associated to B_k , and $\mathbf{E}_{B_k}^0[\tau_0^k] = 1/\Lambda_k$; [see e.g. 34, Equation (1.6.8) and Equation (1.2.27)]. Applying the mean value formula [see 34, Equation (1.3.2)¹¹] we get $\pi(\mathcal{F}_k) = \Lambda_k \mathbf{E}_{B_k}^0[\sigma_0^k]$, which completes the proof of the expression for $\bar{\sigma}_k$. For the other expression, note by the definition of the sequence $\sigma_n^k, \sigma'_n{}^k$ and τ_n^k that $\mathbf{P}$ -almost surely,

$$\bar{\sigma}'_k = \mathbf{E}_{B_k}^0[\sigma_0'^k] = \mathbf{E}_{B_k}^0[\tau_1^k - \sigma_0^k] = \frac{1}{\Lambda_k} - \frac{\pi(\mathcal{F}_k)}{\Lambda_k} = \frac{\pi(\mathcal{F}'_k)}{\Lambda_k},$$

which completes the proof. ■

Proposition A.3: *The probability of outage at the arrival epoch for user of class k is equal to*

$$P_k = \pi\{\mathbf{x} + \varepsilon_k \in \mathcal{F}'_k\} \quad k = 1, \dots, J \quad (16)$$

$\mathbf{P}$ -almost surely.

Proof: By ergodicity we have $P_k = \mathbf{P}_{N_k}^0\{\mathbf{X}(0) \in \mathcal{F}'_k\}$, where $\mathbf{P}_{N_k}^0$ designates the Palm probability associated to N_k (arrival process of the users of class k). By PASTA property the configuration of users $\mathbf{X}(0-)$, just before arrival of the user of class k at time 0, has distribution π . Once the user enters the system, the users configuration becomes $\mathbf{X}(0-) + \varepsilon_k$, whence the result. ■

Proposition A.4: *The mean total time in outage of user of class k is $\mathbf{P}$ -almost surely equal to*

$$D_k = \frac{1}{\mu_k} \pi\{\mathbf{x} + \varepsilon_k \in \mathcal{F}'_k\} \quad k = 1, \dots, J.$$

Proof: Again using the ergodicity of $\{\mathbf{X}(t)\}$ we can write

$$D_k = \mathbf{E}_{N_k}^0 \left[\int_{[0, W_0^k)} \mathbb{1}_{\mathcal{F}'_k}(\mathbf{X}(t)) dt \right].$$

Denote by $\mathbf{Y}(t) := \mathbf{X}(t) - \varepsilon_k \mathbb{1}_{[T_0^k, T_0^k + W_0^k)}(t)$ the process of configurations of users other than the user number 0 of class k (which arrives at time 0 under $\mathbf{E}_{N_k}^0$). By Slivnyak theorem [see e.g. 22, Theorem 1.13] the distribution of the process $\{\mathbf{Y}(t) : t\}$ under $\mathbf{P}_{N_k}^0$ is the same as this of $\{\mathbf{X}(t) : t\}$ under $\mathbf{P}$. Using the fact that W_0^k and $\mathbf{Y}(t)$ are independent under $\mathbf{P}_{N_k}^0$ with $\mathbf{E}_{N_k}^0[W_0^k] = 1/\mu_k$ we obtain

$$\begin{aligned} D_k &= \int_0^\infty \mathbf{E}_{N_k}^0 \left[\mathbb{1}_{[0, W_0^k)}(t) \mathbb{1}_{\mathcal{F}'_k}(\mathbf{Y}(t) + \varepsilon_k) \right] dt \\ &= \frac{1}{\mu_k} \pi\{\mathbf{x} + \varepsilon_k \in \mathcal{F}'_k\}, \end{aligned}$$

¹¹with $Z_k(t) = \mathbb{1}_{\mathcal{F}_k}(\mathbf{X}(t))$

which completes the proof. ■

Proposition A.5: *The mean number of outage incidents experienced by user of class k after its arrival is $\mathbf{P}$ -almost surely equal to*

$$M_k = \frac{1}{\mu_k} \sum_{j=1}^J \lambda_j \pi\{\mathbf{x} + \varepsilon_k \in \mathcal{F}_k, \mathbf{x} + \varepsilon_k + \varepsilon_j \in \mathcal{F}'_k\}, \quad (17)$$

$$k = 1, \dots, J.$$

Proof: Again using the ergodicity of $\{\mathbf{X}(t)\}$ we know that, $\mathbf{P}$ -almost surely,

$$M_k = \mathbf{E}_{N_k}^0 \left[\int_{(0, W_0^k)} B_k(dt) \right].$$

Using the fact that W_0^k and $\mathbf{Y}(t)$ are independent under $\mathbf{P}_{N_k}^0$ with $\mathbf{E}_{N_k}^0[W_0^k] = 1/\mu_k$ we obtain

$$M_k = \mathbf{E}_{N_k}^0[B_k(0, W_0^k)] = \frac{\Lambda_k^*}{\mu_k},$$

where Λ_k^* is the intensity of the point process of exit epochs of $\mathbf{X}(t)$ from $\mathcal{F}_k^* = \{\mathbf{x} : \mathbf{x} + \varepsilon_k \in \mathcal{F}_k\}$. Using Lemma A.1 with $\mathcal{F}_k$ replaced by $\mathcal{F}_k^*$ concludes the proof. ■

We will now prove the result regarding the throughput of the typical call of class k .

Proof of Proposition 2.3: We have

$$\begin{aligned} T_k &= T_k^\delta = \mu_k \mathbf{E}_{N_k}^0 \left[\int_{[0, W_0^k)} r_k \mathbb{1}(\mathbf{X}(t) \in \mathcal{F}_k^\delta) \right. \\ &\quad \left. + r_k'^\delta(\mathbf{X}(t)) \mathbb{1}(\mathbf{X}(t) \notin \mathcal{F}_k^\delta) dt \right]. \end{aligned}$$

It is easy to see, as in the proof of Proposition A.4, that $T_k = r_k \pi\{\mathbf{x} + \varepsilon_k \in \mathcal{F}_k^\delta\} + T'_k$, where

$$T'_k = \mathbf{E} \left[r_k'^\delta(\mathbf{X}(t) + \varepsilon_k) \mathbb{1}((\mathbf{X}(t) + \varepsilon_k) \notin \mathcal{F}_k^\delta) \right] \quad (18)$$

is the part of the throughput obtained by user of class k during its outage time. ■

REFERENCES

- [1] J. Kaufman, "Blocking in a shared resource environment," *IEEE Trans. Commun.*, vol. 29, no. 10, pp. 1474–1481, 1981.
- [2] J. Roberts, "A service system with heterogeneous user requirements," in *Performance of Data Communications Systems and their Applications*, G. Pujolle, Ed., 1981.
- [3] J. Zander, "Distributed co-channel interference control in cellular radio systems," *IEEE Trans. Veh. Technol.*, vol. 41, 1992.
- [4] R. Yates, "A framework for uplink power control in cellular radio systems," *IEEE J. Sel. Areas Commun.*, vol. 13, no. 7, Sept. 1995.
- [5] A. Sampath, P. S. Kumar, and J. Holtzmann, "Power control and resource management for a multimedia CDMA wireless system," in *Proc. 1995 IEEE PIMRC*, vol. 1.
- [6] F. Baccelli, B. Błaszczyszyn, and F. Tournois, "Downlink admission/congestion control and maximal load in CDMA networks," in *Proc. 2003 IEEE Infocom*.
- [7] S.-E. Elayoubi, O. Ben Haddada, and B. Fourestié, "Performance evaluation of frequency planning schemes

- in OFDMA-based networks,” *IEEE Trans. Wireless Commun.*, vol. 7, no. 5-1, pp. 1623–1633, 2008.
- [8] M. K. Karray, “Analytical evaluation of QoS in the downlink of OFDMA wireless cellular networks serving streaming and elastic traffic,” *IEEE Trans. Wireless Commun.*, vol. 9, no. 5, May 2010.
- [9] T. Bonald and A. Proutière, “Wireless downlink data channels: user performance and cell dimensioning,” in *Proc. 2003 Mobicom*.
- [10] F. H. Fitzek, S. Hendrata, P. Seeling, and M. Reisslein, “Video streaming in wireless Internet,” in *Mobile Internet: Enabling Technologies and Services*, ser. Electrical Engineering & Applied Signal Processing, S. Apostolis, Ed. CRC Press, 2004, ch. 11.
- [11] G. Liang and B. Liang, “Effect of delay and buffering on jitter-free streaming over random VBR channels,” *IEEE Trans. Multimedia*, vol. 10, no. 6, pp. 1128–1141, 2008.
- [12] A. ParandehGheibi, M. Médard, S. Shakkottai, and A. Ozdaglar, “Avoiding interruptions — QoE trade-offs in block-coded streaming media applications,” in *Proc. 2010 International Symposium on Information Theory*, pp. 1778–1782.
- [13] L. Rong, S. Elayoubi, and O. Haddada, “Performance evaluation of cellular networks offering TV services,” *IEEE Trans. Veh. Technol.*, vol. 60, no. 2, pp. 644–655, 2011.
- [14] Y. Xu, E. Altman, R. E. Azouzi, M. Haddad, S.-E. Elayoubi, and T. Jiménez, “Probabilistic analysis of buffer starvation in Markovian queues,” in *Proc. 2012 Infocom*.
- [15] D. Bethanabhotla, G. Caire, and M. J. Neely, “Joint transmission scheduling and congestion control for adaptive video streaming in small-cell networks,” *arXiv preprint arXiv:1304.8083*, 2013.
- [16] B. Li, S. T. Chanson, and C. Lin, “Analysis of a hybrid cutoff priority scheme for multiple classes of traffic in multimedia wireless networks,” *Wireless Netw.*, vol. 4, no. 4, pp. 279–290, 1998.
- [17] S. Wu, K. M. Wong, and B. Li, “A dynamic call admission policy with precision QoS guarantee using stochastic control for mobile wireless networks,” *IEEE/ACM Trans. Netw.*, vol. 10, no. 2, pp. 257–271, 2002.
- [18] T. Cover and J. Thomas, *Elements of Information Theory*. New York: John Wiley & Sons, Inc., 2006.
- [19] M. Karray and M. Jovanovic, “Theoretically feasible QoS in a MIMO cellular network compared to the practical LTE performance,” in *Proc. 2012 ICWMC*.
- [20] E. Telatar, “Capacity of multi-antenna Gaussian channels,” *European Trans. Telecommun.*, vol. 10, no. 6, pp. 585–596, Nov. 1999.
- [21] D. Tse and P. Viswanath, *Fundamentals of Wireless Communication*. Cambridge University Press, 2005.
- [22] F. Baccelli and B. Błaszczyszyn, *Stochastic Geometry and Wireless Networks, Volume I — Theory*, ser. Foundations and Trends in Networking. NoW Publishers, 2009, vol. 3, No 3–4.
- [23] F. R. de Hoog, J. H. Knight, and A. N. Stokes, “An improved method for numerical inversion of laplace transforms,” *SIAM J. Scientific and Statistical Computation*, vol. 3, no. 3, pp. 357–366, 1982.
- [24] K. J. Hollenbeck, “Invlap.m: A Matlab function for numerical inversion of Laplace transforms by the de Hoog algorithm,” 1998. [Online]. Available: www.isva.dtu.dk/staff/karl/invlap.htm
- [25] J. Abate and W. Whitt, “Numerical inversion of Laplace transforms of probability distributions,” *ORSA J. on Computing*, vol. 7, no. 1, pp. 38–43, 1995.
- [26] B. Błaszczyszyn, M. K. Karray, and H. P. Keeler, “Using Poisson processes to model lattice cellular networks,” in *Proc. 2013 IEEE INFOCOM*.
- [27] H. Dhillon, R. Ganti, F. Baccelli, and J. Andrews, “Modeling and analysis of k-tier downlink heterogeneous cellular networks,” *IEEE J. Sel. Areas Commun.*, vol. 30, no. 3, pp. 550–560, Apr. 2012.
- [28] H. P. Keeler, B. Błaszczyszyn, and M. K. Karray, “SINR-based coverage probability in cellular networks under multiple connections,” in *Proc. 2013 IEEE ISIT*.
- [29] 3GPP, “Evolved universal terrestrial radio access (E-UTRA); further advancements for E-UTRA — physical layer aspects,” Tech. Rep. 36.814-V900, 2010. [Online]. Available: www.3gpp.org/ftp/Specs/archive/36_series/36.814/
- [30] B. Błaszczyszyn and M. K. Karray, “Linear-regression estimation of the propagation-loss parameters using mobiles’ measurements in wireless cellular network,” in *Proc. 2012 WiOpt*.
- [31] 3GPP, “Evolved universal terrestrial radio access (E-UTRA); radio frequency (RF) system scenarios,” Tech. Rep. 36.942-V830, 2010. [Online]. Available: www.3gpp.org/ftp/Specs/archive/36_series/36.942/
- [32] —, “Evolved universal terrestrial radio access (E-UTRA); physical channels and modulation,” Tech. Rep. 36.211-V910, 2010. [Online]. Available: www.3gpp.org/ftp/Specs/archive/36_series/36.211/
- [33] J. Yu, M. Ameziane *et al.*, “G/g/c simulation model for voip traffic engineering with non-parametric validation,” in *Proc. 2013 International Conference on Digital Telecommunications*, pp. 44–49.
- [34] F. Baccelli and P. Brémaud, *Elements of Queueing Theory; Palm Martingale Calculus and Stochastic Recurrences*. Springer, 2003.



Bartłomiej Błaszczyszyn received his PhD degree and Habilitation qualification in applied mathematics from University of Wrocław (Poland) in 1995 and 2008, respectively. He was an Assistant Professor in the Mathematical Institute, University of Wrocław, and a visiting scientist at the Institute of Stochastics, University of Ulm (Germany) and IEOR Department of Columbia University (USA). He is now a Senior Researcher at Inria (France), and a member of the Computer Science Department of Ecole Normale Supérieure in Paris. His professional interests are in applied probability, in particular in stochastic modeling and performance evaluation of communication networks. He coauthored several publications on this subject in major international journals and conferences, as well as a two-volume book on *Stochastic Geometry and Wireless Networks* NoW Publishers, jointly with F. Baccelli.



Miodrag Jovanovic received the Diploma in engineering from the University of Nis, Nis, Serbia, in 2009 and the Master's degree in electronic systems from Polytech Nantes, France, in 2011. Currently he is preparing a Ph.D. dissertation entitled "Evaluation and Optimization of the Quality Perceived by Mobile Users for New Services in Cellular Networks" under the guidance of M. Karray and B. Błaszczyszyn and carried out at Orange Labs, Issy-Moulineaux and Inria Paris/Rocquencourt, France.



Mohamed Kadhém Karray received his diploma in engineering from Ecole Polytechnique and Ecole Nationale Supérieure des Télécommunications (ENST) in 1991 and 1993, respectively. He prepared a PhD thesis at ENST under the guidance of E. Moulines and B. Błaszczyszyn within 2004-2007. Since 1993 he works at France Télécom R&D (Orange Labs) in France. He co-authored publications in scientific journals and international conferences, and he held two patents on load control in cellular networks. His research activities aim to evaluate the performance of communication networks. His principle tools are probability and stochastic processes, and more specifically information theory, stochastic geometry and queuing theory. In his recent research, he shows how to articulate the tools of these theories to build analytical performance evaluation methods for wireless cellular networks validated with real field measurements. The methods and tools he develops are used by Orange Operator for dimensioning its networks and studying quality of service and energy efficiency.