

Offloading Cellular Traffic Through Opportunistic Communications: Analysis and Optimization

Vincenzo Sciancalepore, *Student Member, IEEE*, Domenico Giustiniano, *Member, IEEE*,
Albert Banchs, *Senior Member, IEEE*, and Andreea Hossmann-Picu

Abstract—Offloading traffic through opportunistic communications has been recently proposed as a way to relieve the current overload of cellular networks. Opportunistic communication can occur when mobile device users are (temporarily) in each other's proximity, such that the devices can establish a local peer-to-peer connection (e.g., via WLAN or Bluetooth). Since opportunistic communication is based on the spontaneous mobility of the participants, it is inherently unreliable. This poses a serious challenge to the design of any cellular offloading solutions, that must meet the applications' requirements. In this paper, we address this challenge from an *optimization analysis* perspective, in contrast to the existing *heuristic* solutions. We first model the dissemination of content (injected through the cellular interface) in an opportunistic network with heterogeneous node mobility. Then, based on this model, we derive the optimal content injection strategy, which minimizes the load of the cellular network while meeting the applications' constraints. Finally, we propose an adaptive algorithm based on control theory that implements this optimal strategy without requiring any data on the mobility patterns or the mobile nodes' contact rates. The proposed approach is extensively evaluated with both a heterogeneous mobility model as well as real-world contact traces, showing that it substantially outperforms previous approaches proposed in the literature.

Index Terms—Cellular traffic offloading, content dissemination, D2D communication, epidemic dissemination, opportunistic communication, optimization analysis, control theory.

I. INTRODUCTION

FOLLOWING the huge popularization of smartphones and the ensuing explosion of mobile data traffic [1], cellular networks are currently overloaded and this is foreseen to worsen in the near future [2]. A recent promising approach to alleviate this problem is to offload cellular traffic through op-

portunistic communications [3]. The key idea is to inject mobile application content to a small subset of the interested users through the cellular network and let these users opportunistically spread the content to others interested upon meeting them. By exploiting opportunistic communications in this way, such an approach has the potential to substantially relieve the load of the cellular infrastructure. Among other mobile applications, this can be used for news [4], road traffic updates [5], social data [6] or streaming content [7]. Indeed, as shown by our performance evaluation results, the load of the cellular network can be reduced between 50% and 95%, depending on the application.

Opportunistic networking exploits the daily mobility of users, which enables intermittent *contacts* whenever two mobile devices are in each other's proximity. These contacts are used to transport data through the opportunistic network, which may introduce substantial delays. However, the type of content concerned by cellular offloading may not always be entirely delay-tolerant. In many applications, it is indeed critical that the content reach all users before a given deadline, lest it lose its relevance or its usability. Therefore, the design of opportunistic-based cellular offloading techniques faces serious challenges from the intermittent availability of transmission opportunities and the high dynamics of the mobile contacts. In order to find the best trade-off between the *load* of the cellular network and the *delay* until the content reaches the interested users, any opportunistic-based offloading design must answer crucial questions such as, *how many* copies of the content to inject, to *which users* and *when*.

While a number of techniques have been proposed in the literature to offload cellular traffic through opportunistic communications, previous approaches are either based on heuristics (and hence do not ensure that the load of the cellular network is minimized) [5]–[7] or fail to provide delay guarantees [4], [7]. In contrast to the above approaches, in this paper, we propose the HYPE (HYbrid oPportunistic and cELLular) technique, which *minimizes* the load of the cellular network while meeting the constraint in terms of *delay guarantees*. To our best knowledge, we are the first to provide such features. The key contributions of our work are as follows:

- 1) Building on the foundations of *epidemic analysis* [8], we propose a model to understand the fundamental trade-offs and evaluate the performance of a hybrid opportunistic and cellular communication approach. Our model reveals that content tends to disseminate faster through opportunistic contacts when a sufficient, but not excessive, number of nodes have already received the content; in

Manuscript received May 20, 2013; revised November 15, 2013; accepted May 17, 2015. Date of publication July 2, 2015; date of current version December 15, 2015. This work has been sponsored by the HyCloud project, supported by Microsoft Innovation Cluster for Embedded Software (ICES), and by the EU H2020-ICT-2014-2 Flex5Gware project, no. 671563.

V. Sciancalepore is with IMDEA Networks Institute, Madrid 28918, Spain, with University Carlos III of Madrid, Madrid 28911, Spain, and also with Politecnico di Milano, Milano 28911, Italy (e-mail: vincenzo.sciancalepore@imdea.org).

D. Giustiniano is with IMDEA Networks Institute, Madrid 28911, Spain (e-mail: domenico.giustiniano@imdea.org).

A. Banchs is with IMDEA Networks Institute, Madrid 28918, Spain, and also with University Carlos III of Madrid, Madrid 28911, Spain (e-mail: albert.banchs@imdea.org).

A. Hossmann-Picu is with Swisscom, Berne 3050, Switzerland (e-mail: hossmann@iam.unibe.ch).

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/JSAC.2015.2452472

contrast, dissemination is slower when either few users have the content or few users are missing it.

- 2) Based on our model, we derive the optimal strategy for injecting content through the cellular network. In line with our previous findings, this strategy uses the cellular network when low speed of opportunistic propagation is statistically expected, and lets the opportunistic network spread the content the rest of the time.
- 3) We design an adaptive algorithm, based on control theory, that implements the optimal strategy for injecting content through the cellular network. The key strengths of this algorithm over previous approaches are that it adapts to the current network conditions without monitoring the nodes' mobility and that it incurs very low signaling overhead and complexity. Both features are essential features for a practical implementation.

The rest of the paper is structured as follows. After thoroughly reviewing related work in Section II, we outline the basic design guidelines of our approach and theoretically analyze its performance in Section III. Based on this analysis, in Section IV, we then derive the optimal strategy and present our adaptive algorithm, which implements this optimal strategy. The algorithm's performance is extensively evaluated in Section V, using mobility models as well as experimental contact traces. Finally, Section VI closes the paper with some final remarks.

II. RELATED WORK

The problem of the unsustainable increase in cellular network traffic and how to offload some of it has become more and more popular. Two types of solutions can be distinguished, on the basis of the outlet chosen for part of the cellular traffic: (i) offloading through additional (new or existing) infrastructure, and (ii) offloading through ad hoc communication. Our proposal, HYPE, falls into the second category.

In the first category, many solutions [9], [10] are aiming to exploit the relatively large number of existing WLAN access points, as well as cellular diversity. A different approach, based on new infrastructure, is introduced in [11], in the context of vehicular networks. In that paper, the authors advocate the deployment of fixed roadside infrastructure units and study the performance of the system in offloading traffic information from the cellular network.

In the second category, along with our study, an increasing body of work is investigating the use of infrastructure-free opportunistic networking as a complement for the cellular infrastructure. In particular, the studies in [4]–[7], [12] propose solutions based on this idea.

In [4], the authors propose to push updates of dynamic content from the infrastructure to subscribers, which then disseminate the content epidemically. The distribution of content updates over a mobile social network is shown to be scalable, and different rate allocation schemes are investigated to maximize the data dissemination speed. A substantial difference between this work and HYPE is that [4] does not minimize the load incurred in the cellular network and does not provide any delay guarantees, which are central objectives in our approach. Moreover, the solution introduced in [4] results in higher re-

source consumption for the “most central” users (i.e., highest contact rates) and/or the “most social” users.

Han *et al.* investigate, in [6], which initial subset of users (who receive the content through the cellular) will lead to the greatest infection ratio. A heuristic algorithm is proposed, that uses the history of user mobility of the previous day to identify a target set of users for the cellular deliveries. HYPE differs significantly from this, in the following aspects: (i) the solution in [6] is heuristic and thus does not guarantee optimal performance, (ii) it requires to know the mobility patterns of all users, which may not be realistic in most scenarios, and (iii) it only investigates *which* users to choose, but not *how many* of them.

In [7], an architecture is implemented to stream video content to a group of smartphones users within proximity of each other, using both the cellular infrastructure and WLAN ad-hoc communication. The decision of who will download the content from the cellular network is based on the phones' download rates. In contrast to our work, the focus of [7] is on the implementation rather than the model and the algorithm. Indeed, the algorithm proposed is a simple heuristic, which does not guarantee optimal performance.

Another study where opportunistic networking is used to offload the mobile infrastructure is [12]. Here, some chosen users, named “helpers,” participate in the offloading, and incentives for these users are provided by using a micro-payment scheme. Alternatively, the operator can offer the participants a reduced cost for the service or better quality of service. Thus, the focus of [12] is on incentives, which is out of the scope of our work.

Most similar to HYPE is the Push-and-Track solution, presented in [5]. There, a subset of users initially receive content from a content provider and subsequently propagate it epidemically. Upon reception of the content, every node sends an acknowledgment to the provider, which may decide to re-inject extra copies to other users. Upon reaching the content deadline, the system enters into a “panic zone” and pushes the content to all nodes that have not yet received it. The most prominent difference between this approach and ours is that Push-and-Track relies on a heuristic to choose *when* to feed more content copies into the opportunistic network, which does not guarantee that the load on the cellular network is minimized. In contrast, we build on analytical results to guarantee that performance is optimal. An additional drawback of Push-and-Track is that it incurs a very high signaling overhead, which compromises the scalability with the number of subscribed users. Results in Section V confirm that our theory-driven algorithm outperforms the heuristics proposed in [5] both in terms of cellular load and of signaling overhead.

Finally, from a different perspective, HYPE is also related to content dissemination solutions in purely opportunistic networks [13], [14]. However, most of these studies focus on finding the best ways to collaborate or contribute to the dissemination, under various constraints (e.g., limited “public” buffer space). Evaluation is usually based on the delay incurred to obtain desired content or the equivalent metric of average content freshness over time. In contrast, our metric is the load incurred in the cellular network. However, when developing our initial model, we do use a similar modeling method as in purely opportunistic dissemination (e.g., [13], [15]).

Like all the previous works on offloading cellular networks through opportunistic communications [4]–[6], with our approach all the transmissions over the cellular network are unicast. There are several key reasons that limit the usage of multicast messages in a cellular network. First, multicast cannot be easily combined with opportunistic transmissions, as this would require that the Content Server is aware of the cell of each node and can dynamically select the subset of nodes at each cell that receives the multicast message, which is not possible with current cellular multicast approaches. Second, in urban scenarios users will likely be associated to different base stations (there are hundreds/thousands of them in the city, each covering some sector, and in dense urban areas femtocells have started to be deployed). Thus, there is a low probability that users subscribed to a specific content are associated to the same cells at the same time, and hence multicast may collapse to unicast. Finally, transmissions with multicast would occur at the lowest rate to preserve users in the edge of the cell, which degrades the resulting performance.

III. THE HYPE APPROACH

In this section, we present the basic design guidelines of the HYPE (HYbrid oPportunistic and cEllular) approach. HYPE is a hybrid cellular and opportunistic communications approach that delivers content to a set of users by (i) sending the content through the cellular network to an initial subset of the users (which we will call *seed nodes*), and (ii) letting these initial users or seed nodes share the content opportunistically with the other nodes. We aim at designing HYPE so as to combine the cellular and opportunistic communication paradigms in a way that retains the key strengths of each paradigm, while overcoming their drawbacks.

HYPE consists of two main building blocks: (i) the *Content Server*, and (ii) the *Mobile Applications*. The Content Server runs inside the network infrastructure, while the Mobile Applications run in mobile devices that are equipped with cellular connectivity, as well as able to directly communicate with each other via short range connections (e.g., via WLAN or Bluetooth). The Content Server monitors the Mobile Applications and, based on the feedback received from them, delivers the content through the cellular network to a selected subset of Mobile Applications (the seed nodes). When two mobile devices are within transmission range of each other, the corresponding Mobile Applications opportunistically exchange the content by using local (short-range) communications.

A. Objectives

The fundamental challenge of the HYPE approach is the design of the algorithm that decides *which* mobile devices and *when* they should receive the content through the cellular network. The rest of this paper is devoted to the design of such an algorithm. The key objectives in the design are:

- (i) *Maximum Traffic Offload*: Our fundamental objective is to maximize the traffic offloaded and thus reduce the load of the cellular network as much as possible. This is beneficial both for the operators (who may otherwise need to

upgrade their network, if the cellular infrastructure is not capable of coping with current demand), as well as for the users (who must pay for cellular usage, either directly or by seeing their data rate reduced).

- (ii) *Guaranteed delay*: Most types of content have an expiration time, arising either from the content's usefulness to the user (e.g., road traffic information), its validity after an update (e.g., daily news) or its play-out time (e.g., streaming). Therefore, a key requirement for our approach is that the content reaches all the interested users before its deadline.
- (iii) *Fairness among users*: In order to make sure that all users benefit from HYPE, it is important to guarantee a good level of fairness both in terms of cellular usage (for which users have to pay), as well as in terms of opportunistic communications (which may increase the energy consumption of the device).¹
- (iv) *Reduced signaling overhead*: The signaling overhead between the Content Server and the Mobile Applications needs to be low. This is important for two reasons: first, to ensure that HYPE *scales* with the number of mobile devices (otherwise the signaling traffic would overload the cellular network); second, to avoid using the cellular interface for small control packets (which is highly energy inefficient due to the significant tail consumption after a cellular transmission [16]).

The above objectives involve some trade-offs, making it very challenging to satisfy all of them simultaneously. For instance, to maximize the traffic offload, we may consider a greedy approach, where the Content Server sends the content to users with the highest contact rates; however this would (i) deteriorate the fairness among users, and (ii) increase the signaling overhead to gather data on user mobility patterns. Another approach may instead minimize the signaling overhead by injecting content as long as there is enough bandwidth available, avoiding thus any signaling; however, this will not maximize the traffic offload. In the following, we set the basic design guidelines of an approach that satisfies all these objectives.

B. Basic Design Guidelines

In order to satisfy the above objectives, a key decision of HYPE is how to deliver a certain piece of content (hereafter referred to as *data chunk*) through the cellular network. In particular, this decision involves the selection of the nodes to which the data chunk is delivered via cellular, as well as the times when to perform these deliveries.

In HYPE, a data chunk is initially delivered to one or more users through the cellular network; additional copies may be injected later if needed. The decision of when to inject another copy of the chunk is driven by the number of users that have already received it. As long as the deadline has not expired, any user with a copy of the chunk will opportunistically transmit it to all the users it meets, that do not have the chunk. Finally,

¹ Indeed, an important drawback of certain existing solutions is that they tend to over-exploit the users with high contact rates [4], [6], thus discouraging the participation of such users.

upon reaching the deadline of the content, the remaining users that have not yet received the chunk, download it from the cellular network²; this ensures that the *delay guarantees* are met and thus we satisfy objective (ii) from Section III-A.

In order to provide a good level of fairness among users, which is objective (iii), HYPE selects each of the seed nodes uniformly at random. Over the long term, this ensures that, on the one hand, all users have the same load in terms of cellular usage and, on the other hand, they also share fairly well the load incurred in opportunistic communications. This is confirmed by the simulation results presented in Section V, which show that HYPE provides a good level of fairness while paying a small price in terms of performance.³

The approach sketched above meets objectives (ii) and (iii). In the following, we first present a model for the opportunistic dissemination of content injected by a cellular network. Based on this model, in Section IV we derive the optimal strategy for the delivery of a single data chunk, that minimizes the load of the cellular network fulfilling objective (i), and then we design an algorithm to implement this strategy, that incurs very low signaling overhead thus also satisfying objective (iv).

C. Model

In order to derive the optimal strategy, with the above approach, for the delivery of data chunks through the cellular network, we need to determine:

- The *total number of copies* of the data chunk to be delivered by the cellular network. This is not trivial: for example, an overly conservative approach, that delivers too few copies before the deadline, may have the side-effect of overloading the cellular network with a large number of copies when the deadline expires.
- The *optimal instants* for their delivery. The decision of when to deliver a copy of a data chunk through the cellular network is based on the current status of the network, which is given by the number of users that already have the chunk.

In the following, we model the opportunistic dissemination of content injected by a cellular network and analyze the load of the cellular network as a function of the strategy followed. Then, based on this analysis, in Section IV we obtain the optimal strategy, that minimizes the load of the cellular network for a given content deadline.

Let $\mathcal{N}$ be a set of mobile nodes subscribed to the same content, with $N = |\mathcal{N}|$ the size of this set (total number of nodes). All nodes have access to the cellular network. Any two nodes also have the ability to setup pairwise bi-directional wireless links, when they are in each other's communication range (in *contact*). Thus, opportunistic communication happens via the store-carry-forward method, through the sequences of intermittent contacts established by node mobility.

²An added advantage of this architecture is that the mobile nodes only need to keep the data chunks for forwarding until their deadline and no longer. The burden on the mobile nodes' buffers is thus kept very low.

³This is also supported by the results of [5], which show that the difference in terms of performance between the random selection and other strategies is very small.

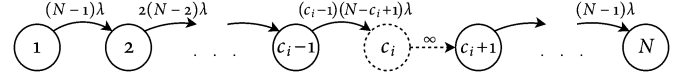


Fig. 1. Markov chain for HYPE communication, assuming **homogeneous node mobility**. Transitions can be caused either by (i) a contact between two nodes, or (ii) injection of the chunk to one node through the cellular network (instantaneous transition, represented with ∞ rate in the figure).

At time 0, a data chunk is injected in the (opportunistic) network, i.e., copies of the chunk are pushed via the cellular interface to a small subset of $\mathcal{N}$, the seed nodes. Throughout the model description, we follow the epidemic dissemination of this chunk of content. We denote by $M(t)$ the number of mobile nodes holding the chunk at time t (we refer to such nodes as “infected”). The delivery deadline assigned to a data chunk is given by T_c (its value depends on the mobile application's requirements).

1) *Opportunistic Communication*: In the opportunistic phase of HYPE, data are exchanged only upon contacts in the network $\mathcal{N}$, therefore a mobility model based on contact patterns is sufficient for our analysis.

We assume every pair of nodes (x, y) in the network $\mathcal{N}$ meets independently of other pairs, at exponentially distributed time intervals⁴ with rate $\beta_{xy} \geq 0$. Then, the opportunistic network $\mathcal{N}$ can be represented as a weighted contact graph using the $N \times N$ matrix $\mathbf{B} = \{\beta_{xy}\}$. We further assume that the inter-contact rates β_{xy} are samples of a generic probability distribution $F(\beta) : (0, \infty) \rightarrow [0, 1]$ with known expectation μ_β (various distribution types for $F(\beta)$ and their effects on aggregated inter-contact times are investigated in [20]). Additionally, we assume that the duration of a contact is negligible in comparison to the time between two consecutive contacts, and that the transmission of a single chunk is instantaneous in both the cellular and the opportunistic network.

The assumptions of exponential inter-contact and negligible contact duration are the norm in analytical work dealing with opportunistic networks [21]–[23]. Studies based on looser assumptions (generic inter-contact models, non-zero contact duration) have, so far, only resulted in broad, qualitative conclusions (e.g., infinite vs. finite delay), while we aim at obtaining more concrete, quantitative results. In addition, all our simulations feature non-zero contact duration and some of them also have non-exponential inter-contact times, thus testing the applicability of our results outside the domain of these assumptions.

Epidemic dissemination in opportunistic networks is typically described with a pure-birth Markov chain, similar to the one in Fig. 1 (slightly adapted from, e.g., [21]). This type of chain only models the *number of copies* of a chunk in the network $\mathcal{N}$ at any point in time, regardless of the specific nodes carrying those copies. This is only possible when considering node mobility to be entirely homogeneous (i.e., all node pairs

⁴Though all pairwise inter-contact rates may not always be exactly exponential (preliminary studies of traces [17] suggested that this is true for subsets of node pairs only), the most in-depth and recent studies [18], [19] conclude that inter-contact time intervals do feature an exponential tail. This is supported by the recent results of Passarella *et al.* [20], which show that the non-exponential aggregated inter-contacts discovered in the preliminary trace studies [17] can, in fact, be the result of exponentially distributed pairwise inter-contacts with different rates.

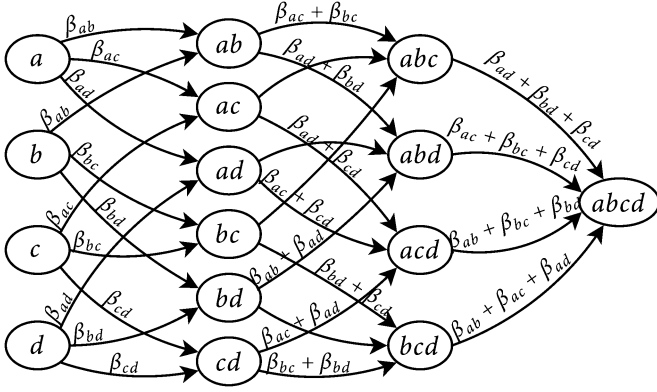


Fig. 2. Markov chain for epidemic spreading, assuming **heterogeneous node mobility**. HYPE specific transitions (i.e., chunk injection by cellular) are left out for clarity. This Markov chain is very complex and intractable for large scenarios; in Theorem 1 we can then reduce it to an equivalent Markov chain that is much simpler and for which we can derive a closed-form solution.

meet at a unique rate: $\beta_{xy} = \lambda$ for all $x, y \in \mathcal{N}$), which allows all nodes to be treated as equivalent.

However, as stated in the beginning of this subsection, we consider node mobility to be heterogeneous, with node pairs meeting at different rates β_{xy} with $x, y \in \mathcal{N}$. In this case, not only the number of spread copies must be modeled, but also the *specific nodes carrying those copies*. This results in more complex Markov chains, as illustrated in Fig. 2 for a 4-node network $\mathcal{N} = \{a, b, c, d\}$.

Transition rates in Markov chains like the one shown in Fig. 2 depend on the nodes “infected” in each of the departure and the arriving states. For example, in Fig. 2, the transition between state a and state ab can happen if node a meets node b . Therefore, the transition time between these two states is exponential with rate given by the meeting rate of the (a, b) node pair, β_{ab} . Similarly, the transition between state ab and state abc can happen if node a meets node c , or if node b meets node c (whichever meeting happens first). Thus, the transition time for this transition is the minimum of two exponential variables with rates β_{ac} and β_{bc} . Since inter-contact times are exponential, this minimum is also exponential with rate $\beta_{ac} + \beta_{bc}$, as shown in Fig. 2.

2) *Cellular Communication*: The decision to deliver a copy of the chunk through the cellular network is based on the current dissemination level, i.e., the number of nodes that already have the chunk. We say that the HYPE process or its associated Markov chain (similar to Fig. 2) is at **level** i , when i mobile nodes are infected, i.e., $M(t) = i$. Each level i corresponds to a set of $\binom{N}{i}$ states $\{\mathbf{K}_1^i, \mathbf{K}_2^i, \dots, \mathbf{K}_{\binom{N}{i}}^i\}$ in the Markov chain. For instance, in our 4-node network from Fig. 2, the HYPE process is at level 3, when the chain is in any of the states $\mathbf{K}_1^3 = abc$, $\mathbf{K}_2^3 = abd$, $\mathbf{K}_3^3 = acd$ or $\mathbf{K}_4^3 = bcd$.

The strategy to transmit copies of the chunk over the cellular network is given by the levels at which we inject a copy. We denote these levels by $C = \{c_1, c_2, \dots, c_d\}$: as soon as we reach one of these levels $c_i \in C$ before the deadline T_c , a copy of the chunk is sent to a randomly chosen node. With this, the transitions in the HYPE Markov chain can be caused either by: (i) a contact between two nodes (one infected, the other uninfected),

which occurs at rates indicated in the previous subsection, or (ii) the injection of the chunk to one node through the cellular network. The latter corresponds to an instantaneous transition (since the chain instantly “jumps” to a state of the next dissemination level), and is represented in Fig. 1 with ∞ rate.⁵ Finally, upon reaching the deadline T_c , the chunk is sent through the cellular network to those nodes that do not have the content by that time.

D. Analysis

Based on the above model, in the following, we analyze the load of the cellular network (which is the metric that we want to minimize) as a function of the strategy followed to inject content (which is given by $C = \{c_1, c_2, \dots, c_d\}$). The cellular network load corresponds to the number of copies delivered through the cellular network, which we denote by D . Let $p_i(t) = \mathbb{P}[M(t) = i]$ denote the probability of being at level i at time t . Then, D is given by:

$$D = \sum_{i=1}^N (d_i + d_i^*) p_i(T_c) \quad (1)$$

where d_i is the number of deliveries through the cellular network that take place until level i is reached ($d_i = |\{1, 2, \dots, i\} \cap C|$) and d_i^* is the number of copies delivered upon reaching the deadline T_c , if it expires at level i ($d_i^* = N - i$).

In order to compute $p_i(T_c)$, we first analyze the case $C = \{c_1\}$,⁶ i.e., when we only inject one copy of the data chunk at the beginning and do not inject any other until we reach the deadline. Let $p_1^{c_1}(T_c)$ denote the probability that, in this case, the system is at level i at time T_c . In order to compute $p_i^{c_1}(T_c)$, we model the transient solution of our Markov chain as shown in the following theorem. (The formal proofs of the theorems are provided in the Appendix.)

Theorem 1: According to the HYPE Markov chain for heterogeneous mobility (similar to Fig. 2), the process $\{M(t), t \geq 0\}$ is described by the following system of differential equations:

$$\begin{cases} \frac{d}{dt} p_1^{c_1}(t) = -\lambda_1 p_1^{c_1}(t), & i = 1 \\ \frac{d}{dt} p_i^{c_1}(t) = -\lambda_i p_i^{c_1}(t) + \lambda_{i-1} p_{i-1}^{c_1}(t), & 1 < i < N \\ \frac{d}{dt} p_N^{c_1}(t) = \lambda_{N-1} p_{N-1}^{c_1}(t), & i = N \end{cases} \quad (2)$$

where $\lambda_i = i(N - i)\mu_\beta$. (Recall that μ_β is the known expectation of the generic probability distribution $F(\beta) : (0, \infty) \rightarrow [0, 1]$, from which the inter-contact rates describing our network are drawn: $\{\beta_{xy}\} = \mathbf{B}$.)

Theorem 1 has effectively reduced our complicated Markov chain for heterogeneous mobility back to a simpler Markov chain, like the one in Fig. 1 (the λ factor being replaced by μ_β). In the simpler chain, each state represents a level of chunk dissemination (i.e., number of nodes holding a copy of the chunk). This is possible, as shown in the proof, thanks to the fact that our heterogeneous contact rates β_{xy} are all drawn from the same

⁵Note that, for clarity, the Markov chain of Fig. 2 does not model transitions caused by chunk injection through the cellular network. This type of transition would be the same as in Fig. 1 (i.e., ∞ rate).

⁶Note that c_1 must necessarily be equal to 0.

distribution, $F(\beta) : (0, \infty) \rightarrow [0, 1]$, which means that all the states of a certain dissemination level i : $\{\mathbf{K}_1^i, \mathbf{K}_2^i, \dots, \mathbf{K}_{\binom{N}{i}}^i\}$ are, in fact, statistically equivalent.

Applying the Laplace transform to the above differential equations, and taking into account that $p_i^{c1}(0) = \delta_{i1}$, leads to

$$\begin{cases} sP_1^{c1}(s) = -\lambda_1 P_1^{c1}(s) + 1, & i = 1 \\ sP_i^{c1}(s) = -\lambda_i P_i^{c1}(s) + \lambda_{i-1} P_{i-1}^{c1}(s), & 1 < i < N \\ sP_N^{c1}(s) = \lambda_{N-1} P_{N-1}^{c1}(s), & i = N \end{cases} \quad (3)$$

from which

$$\begin{cases} P_i^{c1}(s) = \frac{1}{s+\lambda_i} \prod_{j=1}^{i-1} \frac{\lambda_j}{s+\lambda_j}, & i < N \\ P_N^{c1}(s) = \frac{1}{s} \prod_{j=1}^{N-1} \frac{\lambda_j}{s+\lambda_j}, & i = N \end{cases} \quad (4)$$

In case we deliver the data chunk through the cellular network at the levels $C = \{c_1, c_2, \dots, c_d\}$, then the transitions corresponding to those levels are instantaneous, and the Laplace transforms of the probabilities $P_i(s)$ are computed as:

$$P_i(s) = \begin{cases} \frac{1}{s+\lambda_i} \prod_{j \in S_{i-1}} \frac{\lambda_j}{s+\lambda_j}, & i < N, i \notin C \\ 0, & i < N, i \in C \\ \frac{1}{s} \prod_{j \in S_{N-1}} \frac{\lambda_j}{s+\lambda_j}, & i = N \end{cases} \quad (5)$$

where S_{i-1} is the set of levels up to level $i-1$, without including those that belong to set C , i.e., $S_{i-1} = \{1, 2, \dots, i-1\} \setminus (\{1, 2, \dots, i-1\} \cap C)$. For the levels $i \in C$, we simply have $P_i^C(s) = 0$, since we will never be at these levels.

From Eq. (5), we can obtain a closed-form expression for the probabilities $p_i(t)$ as follows. The polynomial $P_i(s)$ is characterized by first and second order poles which have all negative real values. Let $\{s = -\lambda_n\}$ be the poles of $P_i(s)$. Then, $p_i(t)$ for $i < N, i \notin C$ is computed as:

$$p_i(t) = \left(\prod_{j \in S_{i-1}} \lambda_j \right) \sum_{\{s = -\lambda_n\}} \text{Res} \left(\frac{e^{st}}{\prod_{j \in S_i} (\lambda_j + s)} \right) \quad (6)$$

where Res indicates the residue, which is given by:

$$\begin{aligned} & \text{Res}_{s=-\lambda_n} \left(\frac{e^{st}}{\prod_{j \in S_i} (\lambda_j + s)} \right) \\ &= \begin{cases} \frac{e^{-\lambda_n t}}{\prod_{\substack{j \in S_i \\ j \neq n}} (\lambda_j - \lambda_n)}, & -\lambda_n \text{ is a 1st order pole} \\ \frac{e^{-\lambda_n t} \left[t - \sum_{\substack{j \in S_i \\ \lambda_j \neq \lambda_n}} \frac{1}{(\lambda_j - \lambda_n)} \right]}{\prod_{\substack{j \in S_i \\ \lambda_j \neq \lambda_n}} (\lambda_j - \lambda_n)}, & -\lambda_n \text{ is a 2nd order pole} \end{cases} \end{aligned}$$

Additionally, for $i < N, i \in C$ we have $p_i(t) = 0$, and for $i = N$, $p_N(t) = 1 - \sum_{i=1}^{N-1} p_k(t)$.

By evaluating $p_i(t)$ at time $t = T_c$ and applying Eq. (1), we can compute the average number of deliveries over the cellular network, D .

IV. OPTIMAL STRATEGY AND ADAPTIVE ALGORITHM

In this section, we first leverage on the above model to determine the optimal strategy for the delivery of data chunk, and then we design an adaptive algorithm to implement this strategy.

A. Optimal Strategy Analysis

Our goal is to find the best strategy $C = \{c_1, c_2, \dots, c_d\}$ for injecting chunk copies over the cellular network, that minimizes the total load D of the cellular network while meeting the content's deadline T_c . To solve this optimization problem, we proceed along the following two steps:

- 1) We show that the optimal strategy is to deliver the content through the cellular network only at the beginning and at the end of the data chunk's *period*, and never in-between. The data chunk's *period* is defined as the interval between $t = 0$ (when we first start distributing the content) and $t = T_c$ (when the content's deadline expires).
- 2) We obtain the optimal number of copies of the chunk to be delivered at the beginning of the period such that the average load of the cellular network, D , is minimized.

The following theorem addresses the first step.

Theorem 2: In the optimal strategy, the data chunk is delivered through the cellular network to d seed nodes at time $t = 0$, and to the nodes that do not have the content by the deadline at time $t = T_c$.

According to Theorem 2, the optimal strategy is to: (i) deliver a number of copies through the cellular network at the beginning of the period, (ii) wait until the deadline *without* delivering any additional copy, and (iii) deliver a copy of the chunk to the mobile nodes missing the content at the end of the period.

The intuition behind this result is as follows. When few users have the content, information spreads slowly, since it is unlikely that a meeting between two nodes involves one of the few that have already the content. Similarly, information spreads slowly when many users have the content, as a meeting involving a node that does not yet have the content is improbable.

The strategy given by Theorem 2 avoids the above situations by delivering a number of chunk copies through cellular communication at the beginning (when few users have the content) and at the end (where few users miss the content). As a result, the strategy lets the content disseminate through opportunistic communication when the expected speed of dissemination is higher, which allows to minimize the average load of the cellular network.

The second challenge in deriving the optimal strategy is to compute the optimal number of copies of the chunk to be delivered at the beginning of the period, which we denote by d . To that end, the following proposition defines the notion of gain and computes it.

Proposition 1: Let us define G_d as the gain resulting from sending the $(d+1)^{\text{th}}$ chunk of chunk copy at the beginning

of the period (i.e., $G_d = D_d - D_{d+1}$, where D_{d+1} and D_d are the values of D when we deliver $d + 1$ and d copies at the beginning, respectively). Then, G_d can be computed from the following equation:

$$G_d = \sum_{j=d}^{N-1} \frac{\lambda_j}{\lambda_d} p_j^d(T_c) - 1, \quad (7)$$

Building on the above notion of G_d , the following theorem provides the optimal point of operation:

Theorem 3: The optimal value of d is the one that satisfies $G_d = 0$.

The rationale behind the above theorem is as follows. When $G_d > 0$, by sending one additional copy at the beginning, we save more than one copy at the end of the period and hence obtain a gain. Conversely, when $G_d < 0$, we do not benefit from increasing d . The proof shows that G_d is a strictly decreasing function of d , which implies that, to find the optimal point of operation, we need to increase d as long as $G_d > 0$ and stop when we reach $G_d = 0$ (after this point, $G_d < 0$ and further increasing d yields a loss).

B. Adaptive Algorithm for Optimal Delivery

While the previous section addressed the delivery of a single data chunk, in this section we focus on the delivery of the entire content, e.g., a flow of road traffic updates, news feeds or a streaming sequence. We consider that the distribution of content in mobile applications is typically performed by independently delivering different pieces of content in a sequence of data chunks. For instance, a streaming content of 800 MB may be divided into a sequence of chunks of 1 MB. When delivering chunks in sequence, we need to adapt to the system dynamics. For instance, inter-contact time statistics may vary depending on the time of the day [24], which means that the optimal d value obtained by Theorem 3 needs to be adapted accordingly. Similarly, the number of mobile nodes N subscribed to the content may change with time, e.g., based on the content popularity.

To address the above issues, we design an adaptive algorithm based on control theory, that adjusts the number of chunk copies d delivered at the beginning of each period to the behavior observed in previous rounds (hereafter we refer to the sequence of periods as rounds). For instance, in the example above we would have a total of 800 rounds. In the following, we first present the basic design guidelines of our adaptive algorithm. Building on these guidelines, we then design our system based on control theory. Finally, we conduct an analysis of the system to guarantee its stability and ensure good response times.

C. Adaptive Algorithm Basics

In order to devise an adaptive algorithm that drives the system to optimality, we first need to identify *which variable* we should monitor and *what value* this variable should take in optimal operation. To do this, we build on the results of the previous section to design an algorithm that: (i) monitors *how many additional infected nodes* we would have at the end of a round, if we injected one extra copy at the beginning of that round; and (ii) drives the system to optimality by increasing

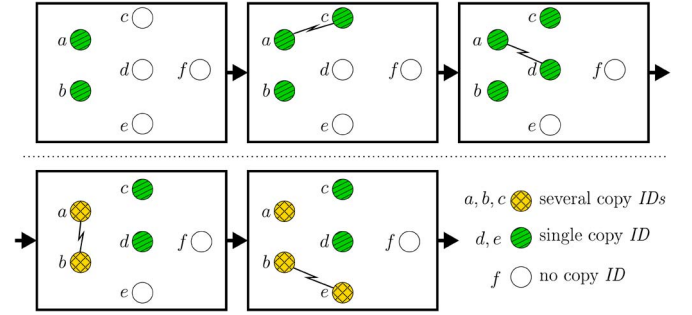


Fig. 3. Example of chunk dissemination in optimal operation. Node a and b receive a copy of the chunk from the Content Server ($d = 2$). At the end of the round, there are two nodes with a single copy ID, that is, $s = 2$.

or decreasing d depending on whether this number is above or below its optimal value.

To efficiently monitor the number of additional infected nodes, we apply the following reasoning. According to Theorem 3, in optimal operation, one extra delivery at the beginning of a round leads to one additional infected node at the end of that round. If we focus on a single copy of the chunk delivered over the cellular network and consider it as the extra delivery, the nodes that would receive the content due to this one extra delivery are those that *received this specific copy and could not have received the chunk from any other source*. Since this holds for each of the d copies delivered over the cellular network, in optimal operation there are on average a total of d nodes at the end of the round, which received the chunk from one source and could not have received it from any other source. Our algorithm focuses on this aggregate behavior of the d deliveries rather than on a single copy, as this provides more accurate information about the epidemic dissemination of the data chunk.

Based on this, each round of the adaptive algorithm proceeds as follows (see Fig. 3 for an example):

- 1) Initially, copies of the data chunk are transmitted to a random set of d seed nodes over the cellular network. Each of the copies is marked with a different *ID* that uniquely identifies the source of the copy.
- 2) When a node that does not have the chunk receives it from another node opportunistically, it records the *ID* of the copy received.
- 3) If two nodes that have copies with different *IDs* meet, they mark this event, to record that they could have potentially received the chunk from different sources.⁷ We say that such nodes have “several copy *IDs*,” while those that keep only one *ID* have a “single copy *ID*.”
- 4) If a node who does not have any copy or has a single copy *ID* meets with another node who recorded the “several copy *IDs*” event, the first node also marks its copy with the “several copy *IDs*” mark.
- 5) At the end of the round, the nodes whose chunk comes from a single source (i.e., no “several copy *IDs*” mark) send a signal to the Content Server.

⁷Note that, for a node with “several copy *IDs*,” we only mark the event and do not keep the *IDs* of the copies, since (i) we are only interested in signaling the number of nodes with a single copy *ID*, and (ii) this leads to more efficient operation, requiring fewer communications and less protocol overhead.

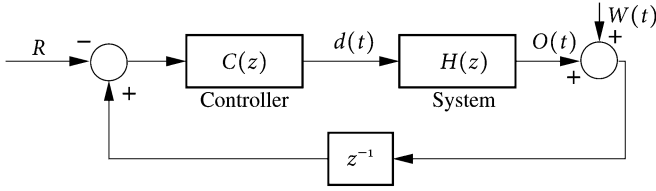


Fig. 4. Our system is composed by two modules: the controlled system $H(z)$, that models the behavior of HYPE, and the PI controller $C(z)$, that drives the controlled system to the optimal point of operation.

By running the above algorithm, we count the number of nodes whose copy of the chunk comes from a single source (i.e., who have a single copy ID at the end of the round), which we denote by s . As argued at the beginning of this section, in optimal operation this number is equal to the number of seed nodes. This implies that, at this operating point, *the number of data chunks injected through the cellular network at the beginning of the round is equal, in expectation, to the number of signals received at the end of the round*, i.e., $s = d$.

A key feature of the above algorithm design is that it does not require gathering any complex statistics on the network, such as the behavior of the mobile nodes, their mobility or social patterns, or their contact rates. Instead, we just need to keep track of the number of chunks injected at the beginning of each round and the signals received at the end, and this is sufficient to drive the system to optimal operation. As a result, the proposed algorithm involves very reduced signaling overhead, which fulfills one of the objectives that we had identified in Section III-A, namely objective (iv).

D. System Design

Based on the above design guidelines, our adaptive algorithm should (i) monitor the number of signals received at the end of each round, and (ii) drive the system to the point of operation where this value is equal to the number of copies injected at the beginning of the round. To do this, in this paper we rely on control theory, which provides the theoretical basis for monitoring a given variable (the *output signal* in control theory terminology) and driving it to some desired value (the *reference signal*).

Following a control theoretic design, we propose the system depicted in Fig. 4. This system is composed from a controller $C(z)$, which is the adaptive algorithm that controls the chunk delivery, and the controlled system $H(z)$, which represents the HYPE network. Furthermore, the component z^{-1} provides the delay in the feedback-loop (to account for the fact that the d value used in the current round is computed from the behavior observed in the previous round). For the controller, we have decided to use a Proportional-Integral (PI), because of its simplicity and the fact that it guarantees zero error in the steady-state. The z transform of the PI controller is given by:

$$C(z) = K_p + \frac{K_i}{z - 1} \quad (8)$$

where K_p and K_i are the parameters of the controller.

Here, the variable that we want to optimally adjust is the number of deliveries at the beginning of the round (i.e., d). Following classical control theory [25], this variable is the *control*

signal provided by the controller. In each round, the controller monitors the system behavior (and in particular the output signal, which we will define later), given the value d that is currently used. Based on this behavior, it decides whether to increase or decrease d in the next round, in order to drive the output signal to the reference signal.

A key aspect of the system design is the definition of the output and reference signals. On the one hand, we need to enforce that by driving the output signal to the reference signal, we bring the system to the optimal point of operation. On the other hand, we also need to ensure that the reference signal is a constant value that does not depend on variable parameters, such as the number of nodes or the contact rates.

Following the arguments exposed in Section IV-C, we design the output signal $O(t)$ and the reference signal R of our controller as follows:

$$\begin{cases} O(t) = s(t) - d(t) \\ R = 0 \end{cases} \quad (9)$$

where $d(t)$ is the number of deliveries at the beginning of a given round t , and $s(t)$ is the number of signals received at the end of this round. Note that, with the above output and reference signals, by driving $O(t)$ to R we bring the system to the point of operation given by $s = d$, which, as discussed previously, corresponds to the optimal point of operation. Following classical control theory, we represent the randomness of the system by adding some noise $W(t)$ to the output signal, as shown by Fig. 4.

E. Control Theoretic Analysis

The behavior of the proposed system (in terms of stability and response time) depends on the parameters of the controller $C(z)$, namely K_p and K_i . In the following, we conduct a control theoretic analysis of the system and, based on this analysis, calculate the setting of these parameters. Note that this analysis guarantees that the algorithm quickly converges to the desired point of operation and remains stable at that point.

In order to analyze our system from a control theoretic standpoint, we need to characterize the HYPE network with a transfer function $H(z)$ that takes d as input and provides $s - d$ as output. In order to derive $H(z)$, we proceed as follows. According to the definition given in Proposition 1, G_d is the gain resulting from sending an extra copy of the chunk. In one round, by sending one extra copy of the chunk at the beginning, there are on average s/d additional nodes that have the chunk at the end. Indeed, s is the total number of nodes that receive the chunk from only one of the d initial seed nodes, which means that on average each seed node contributes with s/d to this number. This yields to:

$$G_d = s/d - 1, \quad (10)$$

from which:

$$s - d = G_d d. \quad (11)$$

The above provides a nonlinear relationship between d and $s - d$, since G_d (given by Eq. (7)) is a non-linear function of d .

To express this relationship as a transfer function $H(z)$, we linearize it at the optimal point of operation.⁸ Then, we study the linearized model and ensure its stability through appropriate choice of parameters. Note that the stability of the linearized model guarantees that our system is locally stable.⁹

To obtain the linearized model, we approximate the perturbations suffered by $s - d$ at the optimal point of operation, $\Delta(s - d)$, as a linear function of the perturbations suffered by d , Δd ,

$$\Delta(s - d) \approx \frac{\partial(s - d)}{\partial d} \Delta d, \quad (12)$$

which gives the following transfer function for the linearized system:

$$H(z) = \frac{\partial(s - d)}{\partial d}. \quad (13)$$

Combining the above with Eq. (11), we obtain the following expression for $H(z)$:

$$H(z) = \frac{\partial(s - d)}{\partial d} = G_d + d \frac{\partial G_d}{\partial d}. \quad (14)$$

Evaluating $H(z)$ at the optimal point of operation ($G_d = 0$) yields:

$$H(z) = d \frac{\partial G_d}{\partial d}. \quad (15)$$

To calculate the above derivative, we approximate λ_i (given by $\lambda_i = i(N - i)\mu_\beta$) by its first order Taylor polynomial evaluated at level $i = \hat{d}$, where $\hat{d}$ is the average value of i at time T_c (i.e., the average number of nodes that have the chunk at the deadline). Since the Taylor polynomial provides an accurate approximation for small perturbations around $\hat{d}$, and the number of nodes that have the chunk at time T_c is distributed around this value, we argue that this approximation leads to accurate results. The first order Taylor polynomial for λ_i at $i = \hat{d}$ is:

$$\lambda_i \approx \lambda_{\hat{d}} - (i - \hat{d})(2\hat{d} - N)\mu_\beta. \quad (16)$$

Substituting this into Eq. (7) yields

$$\begin{aligned} G_d &= \frac{1}{\lambda_d} \sum_{i=1}^N p_i^d(T_c) \left(\lambda_{\hat{d}} - (i - \hat{d})(2\hat{d} - N)\mu_\beta \right) - 1 \\ &= \frac{\lambda_{\hat{d}}}{\lambda_d} - 1 = \frac{\hat{d}(N - \hat{d})\mu_\beta}{d(N - d)\mu_\beta} - 1 \end{aligned} \quad (17)$$

Since at the optimal point of operation we have $G_d = 0$, this implies that (at this operating point) $d = \hat{d}$. Moreover, from Theorem 3 we have that, when operating at the optimal point, if we deliver one additional copy at the beginning (i.e., increase d by one unit), this leads to one additional node with the chunk at

the end (i.e., $\hat{d}$ also increases by one unit). Therefore, at the optimal operating point we also have $\partial \hat{d} / \partial d = 1$. Accounting for all of this when performing the partial derivative of G_d yields:

$$\frac{\partial G_d}{\partial d} = \frac{2(2d - N)}{d(N - d)}, \quad (18)$$

from which:

$$H(z) = d \frac{\partial G_d}{\partial d} = -\frac{2(N - 2d)}{N - d}. \quad (19)$$

Having obtained the transfer function of our HYPE network, we finally address the configuration of the controller parameters K_p and K_i , that will ensure a good trade-off between our system's stability and response time. To this end, we apply the Ziegler-Nichols rules [28], which have been designed for this purpose. According to these rules, we first obtain the K_p value that leads to instability when $K_i = 0$; this value is denoted by K_u . We also calculate the oscillation time T_i under these conditions. Once the K_u and T_i values have been derived, K_p and K_i are configured as follows:

$$K_p = 0.4K_u, \quad K_i = \frac{K_p}{0.85T_i}. \quad (20)$$

Let us start by computing K_u , i.e., the K_p value that ensures stability when $K_i = 0$. From control theory [25], we have that the system is stable as long as the absolute value of the closed-loop gain is smaller than 1. The closed-loop transfer function $T(z)$ of the system depicted in Fig. 4 is given by:

$$T(z) = \frac{-H(z)C(z)}{1 - z^{-1}H(z)C(z)}. \quad (21)$$

To ensure that the closed-loop gain of the above transfer function is smaller than 1, we need to impose $|H(z)C(z)| < 1$. Doing this for $K_i = 0$ yields:

$$|H(z)C(z)| = \left| -\frac{2(N - 2d)}{N - d} K_p \right| < 1. \quad (22)$$

The above inequality gives the following upper bound for K_p , at which the system turns unstable:

$$K_p < \frac{N - d}{2(N - 2d)}. \quad (23)$$

We want to ensure that the system is stable independently of N and d , that is, the above inequality holds for any N and d values. Since the smallest possible value that the right-hand side of Eq. (23) can take is $1/2$ (when $d \rightarrow 0$), the system is guaranteed to be stable as long as $K_p < 1/2$, and may turn unstable when K_p exceeds this value. Accordingly, we set $K_u = 1/2$. Furthermore, when the system becomes unstable, the control signal d may change its sign up to every round, yielding an oscillation period of two rounds, which gives $T_i = 2$. With these K_u and T_i values, we set K_p and K_i following Eq. (20),

$$K_p = \frac{0.4}{2}, \quad K_i = \frac{0.4}{2 \cdot 2 \cdot 0.85}, \quad (24)$$

which terminates the configuration of the PI controller.

⁸This linearization provides a good approximation of the behavior of the system when it suffers small perturbations around the stable point of operation [26]. Note that the approximation only affects the transient analysis and not the analysis of the stable point of operation at which the system is brought by the algorithm.

⁹A similar approach was used in [27] to analyze the Random Early Detection (RED) scheme from a control theoretic standpoint.

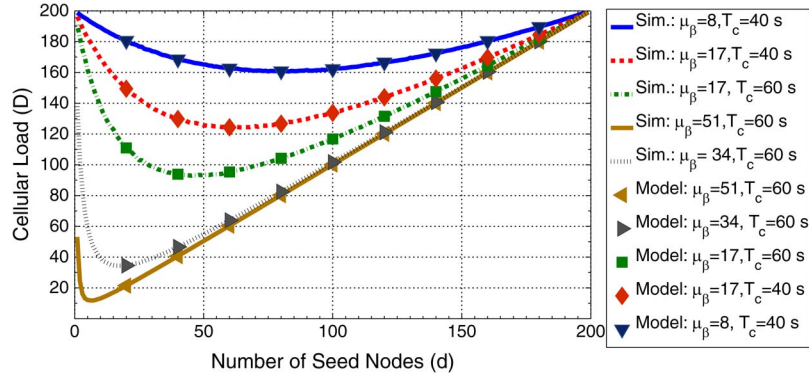


Fig. 5. The analytical model provides very accurate results for different settings (μ_β is given in contacts/pair/day).

While the Ziegler-Nichols rules aim at providing a good trade-off between stability and response time, they are heuristic in nature and thus do not guarantee the stability of the system. The following theorem proves that the system is stable with the proposed configuration.

Theorem 4: The HYPE control system is stable for $K_p = 0.2$ and $K_i = 0.4/3.4$.

V. PERFORMANCE EVALUATION

In this section, we evaluate HYPE for a wide range of scenarios, including several instances of a heterogeneous mobility model, as well as real-world mobility traces. We show that:

- The analytical model provides very accurate results.
- The optimal strategy for data chunk delivery effectively minimizes the load incurred in the cellular network.
- The proposed adaptive algorithm is stable and quickly converges to optimal operation.
- HYPE outperforms previously proposed heuristics in terms of the cellular load, signaling load and fairness among users.

From the four design objectives introduced in Section III-A, our evaluation focuses on the traffic offload, fairness and signaling overhead. Note that, since the delay guarantees are satisfied by design, we meet the objective on the delay.

a) Simulation Setting: To evaluate the performance of HYPE, we use both real mobility traces and a heterogeneous mobility model. For the evaluation with real mobility traces, we select the contact traces collected in the Hagggle project for 4 days during Infocom 2006 [24], and the GPS location traces of San Francisco taxicabs,¹⁰ collected through the Cabspotting project [29]. The number of users for the Infocom 2006 and San Francisco traces are 78 and 536, respectively.

As for the heterogeneous mobility model, we generate contacts as follows. For any given node pair (x, y) , the pairwise inter-contact times are exponentially distributed with rate β_{xy} . The pairwise contact rates, β_{xy} , are drawn from a Pareto distribution¹¹ with mean μ_β (which determines the average frequency of the contacts) and standard deviation σ (which determines the

level of heterogeneity). To account for sparser scenarios, we also run some experiments where a node pair has a probability $p > 0$ of never meeting, i.e., $\beta_{xy} = 0$ (otherwise the inter-contact rate for the pair β_{xy} is drawn as above). In addition, we generate contact durations δ from a Pareto distribution with parameter $\alpha = 2$, as observed in [30]. Following the findings in [17], we choose the average contact rate μ_β and the average contact duration $\mathbb{E}[\delta]$ values such that $1/(\mu_\beta \cdot \mathbb{E}[\delta])$ is between 100 and 1000.

In all the simulations, we set the throughput of the cellular communication to one mobile node equal to 600 kb/s [31] and the bandwidth of opportunistic communication to 20 Mb/s. All the results given in this section are provided with 95% confidence intervals below 0.1%.

b) Baseline Scenarios: For the heterogeneous mobility model, we use the following four baseline scenarios:

- **streaming:** $N = 100$, mean contact rate $\mu_\beta = 13$ contacts/pair/day [24] and $\sigma = 0.58 \cdot \mu_\beta$, Pareto-distributed contact duration $\mathbb{E}[\delta] = 66.46$ s, $T_c = 120$ s [32] and chunk size $L = 1$ MB,
- **road traffic update:** $N = 1000$, mean contact rate $\mu_\beta = 1.2$ contacts/pair/day and $\sigma = 1.5 \cdot \mu_\beta$, Pareto-distributed contact duration $\mathbb{E}[\delta] = 72$ s, $T_c = 600$ s, $L = 1$ MB [5],
- **news feed:** $N = 100$, mean contact rate $\mu_\beta = 0.69$ contacts/pair/day [24] and $\sigma = 2 \cdot \mu_\beta$, Pareto-distributed contact duration $\mathbb{E}[\delta] = 125$ s, $T_c = 3600$ s [32], $L = 0.5$ MB,
- **social data:** $N = 50$, mean contact rate $\mu_\beta = 3.5$ contacts/pair/day [24] and $\sigma = \mu_\beta$, Pareto-distributed contact duration $\mathbb{E}[\delta] = 164$ s, $T_c = 900$ s, $L = 4$ KB.

A. Validation of the Model

In order to validate the analysis conducted in Section III, we evaluate the total load incurred in the cellular network (D) as a function of the strategy followed (which is given by the number of copies of the data chunk delivered at the beginning of a round, d). The results obtained are depicted in Fig. 5 for a scenario with $N = 200$, $\sigma = 0.04$ contacts/pair/day, and different values of T_c (in seconds) and μ_β (in contacts/pair/day). We observe that the analytical results follow very closely those resulting from simulations, which validates the accuracy of our analysis. We further observe that, as pointed out in Section IV, performance degrades for smaller and larger values of d , since when either too few or too many nodes have the content,

¹⁰We assume two taxicabs are in contact when they are within 100 meters of each other.

¹¹Under these conditions, the tail of the aggregate inter-contact times decays as a power law with exponential cut-off [20], as observed in traces, in [19].

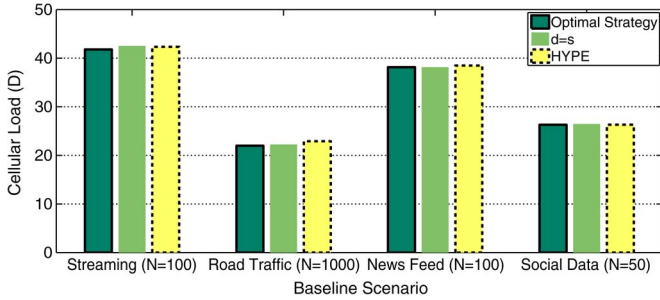


Fig. 6. Validation of the optimal strategy for the four baseline scenarios.

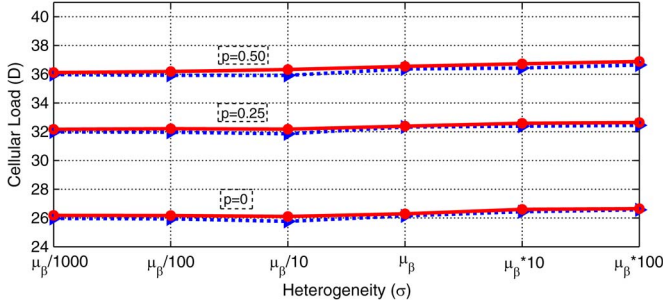


Fig. 7. Cellular load D as a function of the level of heterogeneity (σ) and network sparsity (p).

information spreads slowly. The figure finally shows that—given $\mu_b = 17$ —a smaller T_c (of 40 s) causes a higher load of cellular network than a larger one (of 60 s).

B. Performance Gain and Validation of the Optimal Strategy

We next evaluate the performance gains that can be achieved by opportunistic communications in the four baseline scenarios identified earlier and validate the optimal strategy to confirm that it achieves the highest possible gains. Fig. 6 gives the performance obtained for the four baseline scenarios with: (i) the optimal d value provided by Theorem 3, labeled *Optimal Strategy*, (ii) the strategy proposed in Section IV-C for the design of the adaptive algorithm, labeled $d = s$, and (iii) the adaptive algorithm implemented by HYPE, labeled *HYPE*. For each strategy, the figure shows the absolute average load of the cellular network in number of chunk copies per round (D).

The results obtained show that the proposed approach can reduce very substantially the load of the cellular network (with offloaded traffic ranging from almost 50% in the social data scenario to more than 95% in the road traffic one). The tests also show that the adaptive algorithm implemented by HYPE is very effective in minimizing this load, as it performs practically as the benchmarks given by the optimal and $d = s$ strategies.

C. Impact of Heterogeneity and Sparsity

To understand the impact of the heterogeneity of pairwise contact rates β_{xy} on the proposed approach, Fig. 7 depicts the total cellular load D for the streaming scenario, with varying σ 's. The effect of network sparsity is also shown by using different values for the probability p that a pair of nodes never meet.

We note that HYPE achieves a performance very close to the optimal, which confirms that the HYPE design also works for

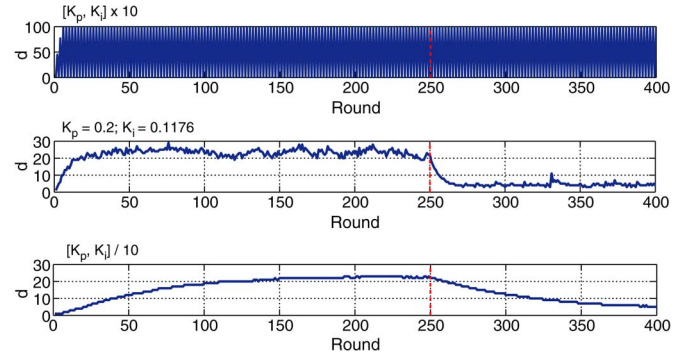


Fig. 8. Evolution of the control signal d over time for different K_p, K_i settings. Our selection of parameters is stable and reacts quickly.

heterogeneous settings, as well as sparse ones. In the sparsest tested scenario ($p = 0.50$), D increases by $\approx 38\%$ as compared to $p = 0$, as a result of the slower dissemination caused by the decreasing number of connections (i.e., larger p). Furthermore, for all tested p values, the cellular load D is mostly insensitive to variations of σ both for HYPE and the optimal strategy, which is in line with Theorem 1.

D. Stability and Response Time

Based on the control theoretic analysis conducted in Section IV, the parameters K_p, K_i of the PI controller have been chosen to guarantee stability and ensure a good response time. In order to assess the effectiveness of this configuration, we evaluate its performance for the streaming baseline scenario and compare it against different choices for the values of parameters K_p, K_i . In Fig. 8, we show the evolution of the control signal d over time for our setting $K_p = 0.2, K_i = 0.1176$, as well as a setting of these parameters ten times larger, labeled $[K_p, K_i] \times 10$ and ten times smaller, labeled $[K_p, K_i]/10$. In the test, μ_b increases from 13 contacts/pair/day to 40 contacts/pair/day after 250 rounds. (For instance, this could be the result of an increase in the number of contacts at rush hour). Results show that our setting is stable and reacts quickly, while a larger setting of K_p, K_i is highly unstable and a smaller setting reacts very slowly. This confirms the choice of parameters made for our controller. We also conducted a similar experiment in which we varied N (which could be for instance the result of a change in content popularity) and observed a similar behavior (not shown in the figure for space reasons).

The above experiment shows the response of the controller to a drastic change. In order to confirm that this response is sufficiently quick to follow the variations of the opportunistic contacts in a realistic environment, we consider the San Francisco real traces and study the temporal evolution of the cellular load D . To provide a benchmark, we compare HYPE against an optimal strategy that selects the best d value every ten rounds.¹² The results, for a content deadline $T_c = 600$ s, are plotted in Fig. 9. These results confirm that HYPE reacts rapidly to dynamic conditions: as there are fewer number of contacts during night

¹²For the optimal strategy, we make an exhaustive search over all possible d values every ten rounds and select the best one. Note that such a strategy cannot be used in practice and is only considered for comparison purposes.

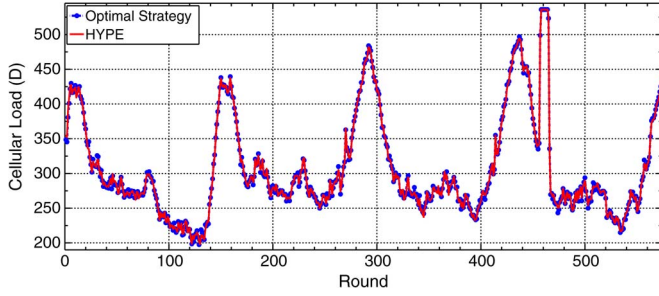


Fig. 9. San Francisco real traces ($N = 536$): Temporal evolution of the cellular load D for HYPE and optimal strategy, with deadline $T_c = 600$ sec.

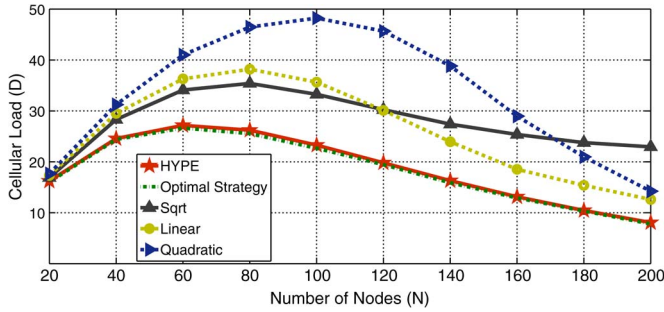


Fig. 10. Comparison with Push-and-Track heuristics [5].

time, HYPE needs to inject more content through the cellular network (up to $D \approx N$), while the higher number of contacts during day time greatly reduces the network load ($D \approx 220$).

E. When to Deliver: HYPE Strategy Versus Other Approaches

One of our key findings in Section IV is that performance is optimized when all the deliveries over the cellular network take place at the beginning and at the end of the period. To validate this result, Fig. 10 compares the performance of HYPE against the Push-and-Track heuristics proposed in [5] (namely, *Sqrt*, *Linear* and *Quadratic*), which distribute the deliveries along the period. Results are given for the social data scenario with a varying number of subscribed users N . We observe from the figure that HYPE substantially outperforms all other approaches (the cellular load is even halved, in some cases), and performs very closely to the *Optimal Strategy* benchmark.

In addition to the above experiment, conducted with a mobility model, we also compare the performance of HYPE against the other approaches with real mobility traces. The results, depicted in Fig. 11, show that HYPE closely follows the performance of the benchmark given by the optimal strategy and outperforms previous heuristics. For the San Francisco real traces, HYPE can offload about 20% more traffic than the previous heuristics. For the Infocom 2006 traces, the employed strategy has a smaller impact on cellular load performance, which yields to a smaller gain (up to about 12%).¹³

¹³Indeed, by conducting experiments with the Infocom 2006 traces for many different strategies (unreported here for space reasons), we observed that performance was relatively similar for all of them, which shows that the impact of the specific strategy followed is limited for this case.

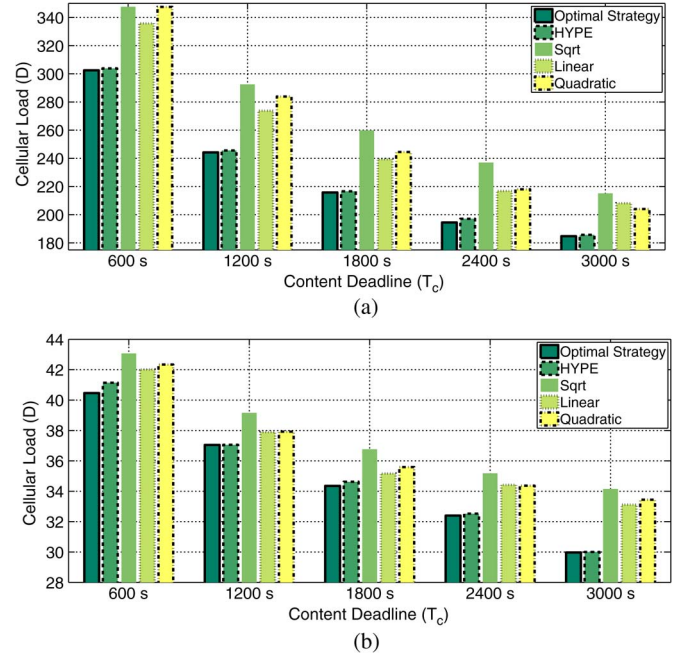


Fig. 11. Tests using real mobility traces for different deadlines T_c . HYPE performs closely to the benchmark provided by the optimal strategy and substantially outperforms previous heuristics. (a) San Francisco real traces. ($N = 536$) (b) Infocom 2006 real traces ($N = 78$).

F. Which Seed Nodes: Comparison to Other Selection Methods

One of the key decisions in the HYPE design is to randomly select a node when transmitting content over the cellular network. In order to gain insight into the impact of this design decision, we compare HYPE against the heuristic approach proposed in [6] to select the seed nodes in the opportunistic network. Unlike HYPE, [6] requires full knowledge of the pairwise contact rates to identify the target set of users, which involves a much higher level of complexity. Note that, since [6] does not provide an algorithm to compute the number of copies d of the data chunk, we apply the HYPE strategy to compute d also for [6].

Fig. 12 shows the performance of both approaches in terms of cellular traffic load (D) and fairness for different values of heterogeneity σ , for the social data scenario. To measure fairness, we apply the Jain's Fairness Index (*JFI*) to the total number of cellular and opportunistic communications involving a node.¹⁴ The results show that HYPE provides a much higher level of fairness than [6] with negligible loss in terms of cellular traffic load. Therefore, HYPE does not only feature a simpler implementation than [6], as it does not need to know the individual inter-contact rates, but also provides a much better trade-off between fairness and cellular load performance. The results of this and the previous section are particularly relevant as the algorithms of [5] and [6] are the only existing proposals in the literature to offload cellular networks with opportunistic communications while providing deterministic delay guarantees.

¹⁴For instance, a node that (i) receives the content through the cellular network and (ii) sends it to n nodes in range during the period, will have a total number of communications equal to $n + 1$.

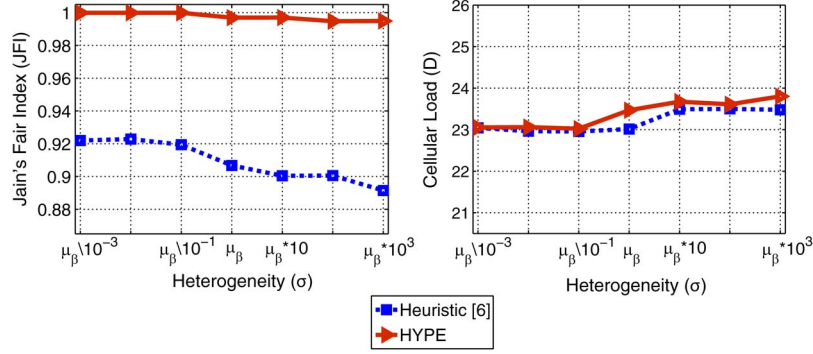


Fig. 12. HYPE versus the heuristic solution of [6], varying the heterogeneity (σ). HYPE provides a much better trade-off between fairness and cellular load performance.

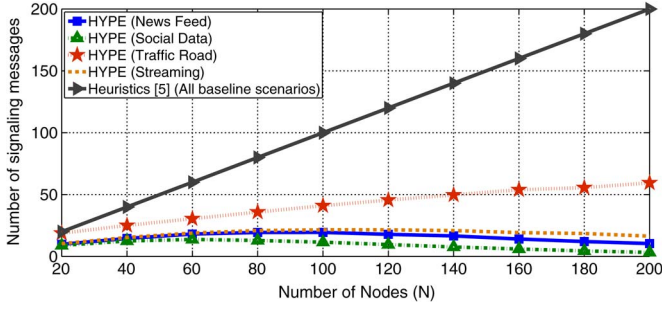


Fig. 13. HYPE signaling load as a function of the number of nodes N for the four baseline scenarios. HYPE outperforms the previous heuristics of [5] and scales very efficiently with the number of nodes.

G. Signaling Load

In order to gain insight into the scalability of our design, we analyze the number of uplink signaling messages sent over cellular network. In HYPE, such messages are the signals sent by the nodes with a single copy ID to the Content Server at the end of the period. Fig. 13 shows the signaling load of HYPE as a function of the number of users for each of the four baseline scenarios, and compares it with the Push-and-Track heuristics (labeled as *Heuristics* [5] in the figure), which require an uplink signal each time a node receives the chunk.¹⁵ In contrast, HYPE scales very efficiently with the number of users: the more users are subscribed to the content, the lower the signaling load per user.

Summarizing the results of the performance evaluation conducted in this section, we have shown that our analytical model is very accurate, that the optimal strategy proposed does indeed minimize the load incurred in the cellular network, that the adaptive algorithm is stable and quick to converge to optimality and that HYPE outperforms existing heuristics in crucial aspects: cellular load, signaling load and fairness among users.

VI. CONCLUSION

In this paper we have presented HYPE, a novel approach to offload cellular traffic through opportunistic communications.

¹⁵We have not compared the signaling messages of HYPE against the approach of [6] since that approach requires gathering data from nodes' mobility patterns. Even though [6] does not explain the signaling mechanism employed, we expect that the need to collect the mobility patterns involves a substantially higher signaling overhead than [5].

To design HYPE, we have developed a theoretical model to analyze the performance of opportunistic dissemination when data can be selectively injected through a cellular network. Based on this model, we have derived the optimal strategy that minimizes the total amount of data injected through the cellular network while meeting delay guarantees. To implement the optimal strategy obtained from the analysis, HYPE runs an adaptive algorithm that adjusts the data delivery over the cellular network to the current network conditions. By building on control theory, we guarantee that this algorithm is stable and quickly adapts to dynamic conditions. The algorithm incurs very low signaling overhead and does not need to monitor the contacts between nodes nor to gather complex statistics, which are important requirements for a practical implementation.

APPENDIX

Theorem 1: According to the HYPE Markov chain for heterogeneous mobility (similar to Fig. 2), the process $\{M(t), t \geq 0\}$ is described by the following system of differential equations:

$$\begin{cases} \frac{d}{dt}p_1^{c_1}(t) = -\lambda_1 p_1^{c_1}(t), & i = 1 \\ \frac{d}{dt}p_i^{c_1}(t) = -\lambda_i p_i^{c_1}(t) + \lambda_{i-1} p_{i-1}^{c_1}(t), & 1 < i < N \\ \frac{d}{dt}p_N^{c_1}(t) = \lambda_{N-1} p_{N-1}^{c_1}(t), & i = N \end{cases} \quad (25)$$

where $\lambda_i = i(N-i)\mu_\beta$. (Recall that μ_β is the known expectation of the generic probability distribution $F(\beta) : (0, \infty) \rightarrow [0, 1]$, from which the inter-contact rates describing our network are drawn: $\{\beta_{xy}\} = \mathbf{B}$.)

Proof: Recall that we denoted by $\{\mathbf{K}_1^i, \mathbf{K}_2^i, \dots, \mathbf{K}_{(N)}^i\}$ the set of $\binom{N}{i}$ states in the Markov chain corresponding to level i . Also, $\mathbf{B} = \{\beta_{xy}\}$ is our network. Then, assuming the Markov chain starts in initial state $\mathbf{K}_m^1$ for $1 \leq m \leq N$ (i.e., the chunk was initially injected to a node m), the probability of still being in this state after a small time interval dt is:

$$\mathbb{P}[\mathbf{K}_m^1 \text{ at } t + dt | \mathbf{B}] = \mathbb{P}[\mathbf{K}_m^1 \text{ at } t | \mathbf{B}] \cdot \left(1 - \sum_{\substack{x \in \mathbf{K}_m^1 \\ y \notin \mathbf{K}_m^1}} \beta_{xy} dt\right) \quad (26)$$

Then, averaging over all states of this dissemination level, the probability of still being at dissemination level 1 after a small

time interval dt is:

$$\mathbb{P}[1 \text{ at } t + dt | \mathbf{B}] = \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t + dt | \mathbf{B}] \quad (27)$$

Finally, considering all possible network realizations:

$$\begin{aligned} p_1^{c_1}(t + dt) &= \mathbb{P}[1 \text{ at } t + dt] \\ &= \int_{\mathbf{B}} \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t + dt | \mathbf{B}] \mathbb{P}[\mathbf{B}] d\mathbf{B}, \end{aligned} \quad (28)$$

where $\mathbb{P}[\mathbf{B}]$ is given by the generic distribution, $F(\beta) : (0, \infty) \rightarrow [0, 1]$, which determines the inter-contact rates of our network. Combining this last equation with Eq. (26) and using basic probability theory, we obtain Eq. (29) below:

$$\begin{aligned} p_1^{c_1}(t + dt) &= \int_{\mathbf{B}} \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t | \mathbf{B}] \left(1 - \sum_{\substack{x \in \mathbf{K}_m^1 \\ y \notin \mathbf{K}_m^1}} \beta_{xy} dt\right) \mathbb{P}[\mathbf{B}] d\mathbf{B} \\ &= \int_{\mathbf{B}} \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t | \mathbf{B}] \mathbb{P}[\mathbf{B}] d\mathbf{B} \\ &\quad - \int_{\mathbf{B}} \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t | \mathbf{B}] \sum_{\substack{x \in \mathbf{K}_m^1 \\ y \notin \mathbf{K}_m^1}} \beta_{xy} dt \cdot \mathbb{P}[\mathbf{B}] d\mathbf{B} \\ &= p_1^{c_1}(t) - \sum_{m=1}^N \int_{\mathbf{B}} \sum_{\substack{x \in \mathbf{K}_m^1 \\ y \notin \mathbf{K}_m^1}} \beta_{xy} dt \\ &\quad \cdot \mathbb{P}[\mathbf{B} | \mathbf{K}_m^1 \text{ at } t] \mathbb{P}[\mathbf{K}_m^1 \text{ at } t] d\mathbf{B} \\ &= p_1^{c_1}(t) - \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t] \mathbb{E}[X | \mathbf{K}_m^1 \text{ at } t] dt, \end{aligned} \quad (29)$$

where $X = \sum_{\substack{x \in \mathbf{K}_m^1 \\ y \notin \mathbf{K}_m^1}} \beta_{xy}$ (that is a sum of $N - 1$ terms).

Since our network's inter-contact rates forming the matrix $\mathbf{B}$ are independent and identically distributed (with generic distribution $F(\beta) : (0, \infty) \rightarrow [0, 1]$ of mean μ_β), the terms of the sum forming X are distributed according to $F(\beta)$, regardless of the specific node combination $\mathbf{K}_m^1$. Hence, $\mathbb{E}[X | \mathbf{K}_m^1 \text{ at } t] = \mathbb{E}[X]$ and Eq. (29) becomes:

$$p_1^{c_1}(t + dt) = p_1^{c_1}(t) - \mathbb{E}[X] dt \cdot \sum_{m=1}^N \mathbb{P}[\mathbf{K}_m^1 \text{ at } t] \quad (30)$$

$$= p_1^{c_1}(t) - (N - 1)\mu_\beta dt \cdot p_1^{c_1}(t). \quad (31)$$

Thus, we obtain as desired:

$$\frac{d}{dt} p_1^{c_1}(t) = -(N - 1)\mu_\beta p_1^{c_1}(t) \quad (32)$$

The remaining two differential equations are obtained by the same process. $\square$

Theorem 2: In the optimal strategy, the data chunk is delivered through the cellular network to d seed nodes at time $t = 0$, and to the nodes that do not have the content by the deadline $t = T_c$.

Proof: The proof goes by contradiction: we first assume that in the optimal strategy the data chunk is transmitted to some mobile node at time $t \neq \{0, T_c\}$ and then we find an alternative strategy that provides a better performance.

If the chunk is transmitted to some mobile node at $t \neq \{0, T_c\}$, this means that $C \neq \{1, \dots, d\}$ and hence there exists some missing value smaller than c_d in C . Indeed, if $C = \{1, \dots, d\}$, all the first d states are instantaneous states and the data chunk is transmitted to d nodes at the beginning of the round.

Let us denote the largest value in C (c_d) by k and the largest value that is missing by $k - l$. Let D_k further denote the value of D for the optimal configuration $C_k = \{c_1, \dots, c_d\}$ (where $c_d = k$), D_{k-l} the value of D for the configuration $C_{k-l} = \{c_1, \dots, c_{d-1}, k - l\}$ and D_{k+1} the value of D for the configuration $C_{k+1} = \{c_1, \dots, c_{d-1}, k + 1\}$.¹⁶ In the following, we show that either D_{k-l} or D_{k+1} , or both, are smaller than D_k , which contradicts the initial assumption that the configuration $\{c_1, \dots, c_d\}$ is optimal.

If we compare the state probabilities for the configurations C_k and C_{k+1} , we have that

$$\begin{cases} P_i^{C_k}(s) = P_i^{C_{k+1}}(s), & i < k \\ P_i^{C_k}(s) = \frac{\lambda_{k+1}(s + \lambda_k)}{\lambda_k(s + \lambda_{k+1})} P_i^{C_{k+1}}(s), & i > k + 1. \end{cases} \quad (33)$$

From the above, we have that the following holds for $i > k + 1$,

$$\begin{aligned} P_i^{C_k}(s) - P_i^{C_{k+1}}(s) &= \left(\frac{\lambda_{k+1}}{s + \lambda_{k+1}} - \frac{\lambda_k}{s + \lambda_k} \right) \prod_{j \in S_{i-1}^{C_k} \setminus k} \frac{\lambda_j}{s + \lambda_j} \\ &= \frac{\lambda_{k+1} - \lambda_k}{\lambda_{k+1} \lambda_k} s P_i^{C_{d-1}}(s), \end{aligned} \quad (34)$$

where $S_{i-1}^{C_k} = \{1, 2, \dots, i-1\} \setminus (\{1, 2, \dots, i-1\} \cap \{c_1, \dots, c_d\})$ and $P_i^{C_{d-1}}(s)$ is state i 's probability for the configuration $C_{d-1} = \{c_1, \dots, c_{d-1}\}$.

By doing the inverse Laplace transform of the above, we have that

$$\begin{aligned} p_i^{C_k}(t) - p_i^{C_{k+1}}(t) &= \frac{\lambda_{k+1} - \lambda_k}{\lambda_k \lambda_{k+1}} \frac{dP_i^{C_{d-1}}(t)}{dt} \\ &= \frac{\lambda_{k+1} - \lambda_k}{\lambda_k \lambda_{k+1}} \left(-\lambda_i P_i^{C_{d-1}}(t) + \lambda_{i-1} P_{i-1}^{C_{d-1}}(t) \right). \end{aligned} \quad (35)$$

Furthermore, we also have

$$P_{k+1}^{C_k}(s) - P_k^{C_{k+1}}(s) = \frac{\lambda_k - \lambda_{k+1}}{\lambda_k} P_{k+1}^{C_{d-1}}(s), \quad (36)$$

and, hence,

$$p_{k+1}^{C_k}(t) - p_k^{C_{k+1}}(t) = \frac{\lambda_k - \lambda_{k+1}}{\lambda_k} P_{k+1}^{C_{d-1}}(t). \quad (37)$$

Combining Eqs. (35) and (37) with Eq. (1), and taking into account that for $i > k + 1$ it holds that $d_i + d_i^* = d_{i+1} + d_{i+1}^* + 1$,

¹⁶Without loss of generality we assume that $c_d \neq N$, as it can be easily shown that a configuration with $c_d = N$ is not optimal.

we obtain

$$D_k - D_{k+1} = (\lambda_k - \lambda_{k+1}) \sum_{i=k+1}^{N-1} \frac{1}{\lambda_k \lambda_{k+1}} p_i^{C_{d-1}}(T_c). \quad (38)$$

Following a similar approach for the configurations C_k and C_{k-l} , we obtain

$$D_k - D_{k-l} = (\lambda_k - \lambda_{k-l}) \sum_{i=k+1}^{N-1} \frac{1}{\lambda_k \lambda_{k-l}} p_i^{C_{d-1}}(T_c). \quad (39)$$

Since it holds that either $\lambda_k - \lambda_{k-l}$ or $\lambda_k - \lambda_{k+1}$ is greater than zero, we have that at least one of the two alternative configurations (C_{k+1} or C_{k-l}) provides a D value smaller than C_k . This contradicts the assumption that in the optimal strategy the data chunk is transmitted to some node at time $t \neq \{0, T_c\}$, which proves the theorem. $\square$

Proposition 1: Let us define G_d as the gain resulting from sending the $(d+1)^{th}$ chunk of chunk copy at the beginning of the period (i.e., $G_d = D_d - D_{d+1}$, where D_{d+1} and D_d are the values of D when we deliver $d+1$ and d copies at the beginning, respectively). Then, G_d can be computed from the following equation:

$$G_d = \sum_{j=d}^{N-1} \frac{\lambda_j}{\lambda_d} p_j^d(T_c) - 1. \quad (40)$$

Proof: G_d can be expressed as:

$$G_d = \sum_{j=d}^{N-1} (N-j) \left(p_j^d(T_c) - p_j^{d+1}(T_c) \right) - 1. \quad (41)$$

The term $p_j^d(T_c) - p_j^{d+1}(T_c)$ is calculated as follows. From Eqs. (5), we have

$$P_j^d(s) - P_j^{d+1}(s) = -\frac{s P_j^d(s)}{\lambda_d}. \quad (42)$$

Making the inverse Laplace transform of the above for $j > d$ yields

$$\begin{aligned} p_j^d(T_c) - p_j^{d+1}(T_c) &= -\frac{1}{\lambda_d} \left. \frac{dp_j^d(t)}{dt} \right|_{T_c} \\ &= \frac{1}{\lambda_d} \left(\lambda_j p_j^d(T_c) - \lambda_{j-1} p_{j-1}^d(T_c) \right). \end{aligned} \quad (43)$$

Note that the above equation also holds for $j = d$ since in this case $p_{j-1}^d(t) = 0$ and $p_j^{d+1}(t) = 0$. Combining it with Eq. (41) leads to

$$G_d = \sum_{j=d}^{N-1} \frac{\lambda_j}{\lambda_d} p_j^d(T_c) - 1. \quad (44)$$

$\square$

Theorem 3: The optimal value of d is the one that satisfies $G_d = 0$.

Proof: As long as $G_d > 0$, we benefit from increasing d , since by sending one additional chunk at the beginning, we save more than one chunk at the end. Conversely, if $G_d < 0$ we do

not benefit. It can be seen that $G_1 > 0$ and $G_N < 0$. Furthermore, it can also be seen that G_d strictly decreases with d :

$$\begin{aligned} G_{d+1} - G_d &= \sum_{j=d}^{N-1} \frac{\lambda_j}{\lambda_{d+1}} p_j^{d+1}(T_c) - \frac{\lambda_j}{\lambda_d} p_j^d(T_c) \\ &= \sum_{j=d}^{N-1} \frac{p_j^d(T_c) \lambda_j (\lambda_{j+1} - \lambda_j - (\lambda_{d+1} - \lambda_d))}{\lambda_d \lambda_{d+1}} < 0. \end{aligned} \quad (45)$$

From the above, it follows that the value of d that minimizes D is the one that satisfies $G_d = 0$, since up to this value we benefit from increasing d and after this value we stop benefiting, which proves the theorem. $\square$

Theorem 4: The HYPE control system is stable for $K_p = 0.2$ and $K_i = 0.4/3.4$.

Proof: The closed-loop transfer function of our system is

$$T(z) = \frac{-z(z-1)HK_p - zHK_i}{z^2 + (-HK_p - 1)z + H(K_p - K_i)} \quad (46)$$

where

$$H = -2 \quad (47)$$

A sufficient condition for stability is that the poles of the above polynomial fall within the unit circle $|z| < 1$. This can be ensured by choosing coefficients $\{a_1, a_2\}$ of the characteristic polynomial that belong to the stability triangle [25]:

$$\begin{cases} a_2 < 1 \\ a_1 < a_2 + 1 \\ a_1 > -1 - a_2 \end{cases} \quad (48)$$

In the transfer function of Eq. (46), the coefficients of the characteristic polynomial are $a_1 = -HK_p - 1$ and $a_2 = H(K_p - K_i)$. From Eqs. (24) and (47), we have $HK_p = -0.4$ and $HK_i = -0.4/(0.85 \cdot 2)$, from which $a_1 = -0.6$ and $a_2 = -0.16$. It can be easily seen that these $\{a_1, a_2\}$ values satisfy Eq. (48), which proves the theorem. $\square$

ACKNOWLEDGMENT

We thank the anonymous reviewers for their inspiring and insightful feedbacks.

REFERENCES

- [1] Gartner Says Worldwide Smartphone Sales Soared in Fourth Quarter of 2011 With 47 Percent Growth, Press Release, Gartner, Inc., Stamford, CT, USA, 2012. [Online]. Available: <http://www.gartner.com/it/page.jsp?id=1924314>
- [2] Cisco Visual Networking Index: Global Mobile Data Traffic Forecast Update, 2011–2016, White Paper, Cisco, San Jose, CA, USA, 2012. [Online]. Available: http://www.cisco.com/en/US/solutions/collateral/ns341/ns525/ns537/ns705/ns827/white_paper_c11-520862.html
- [3] L. Pelusi, A. Passarella, and M. Conti, "Opportunistic networking: Data forwarding in disconnected mobile ad hoc networks," *IEEE Commun. Mag.*, vol. 44, no. 11, pp. 134–141, Nov. 2006.
- [4] S. Ioannidis, A. Chaintreau, and L. Massoulie, "Optimal and scalable distribution of content updates over a mobile social network," in *Proc. IEEE INFOCOM*, 2009, pp. 1422–1430.

- [5] J. Whitbeck, M. Amorim, Y. Lopez, J. Leguay, and V. Conan, "Relieving the wireless infrastructure: When opportunistic networks meet guaranteed delays," in *Proc. IEEE WoWMoM*, 2011, pp. 1–10.
- [6] B. Han *et al.*, "Cellular traffic offloading through opportunistic communications: A case study," in *Proc. ACM CHANTS*, 2010, pp. 31–38.
- [7] L. Keller *et al.*, "MicroCast: Cooperative video streaming on smartphones," in *Proc. ACM MobiSys*, 2012, pp. 57–70.
- [8] D. Daley and J. Gani, *Epidemic Modelling: An Introduction*. Cambridge, U.K.: Cambridge Univ. Press, 2005.
- [9] K. Lee, J. Lee, Y. Yi, I. Rhee, and S. Chong, "Mobile data offloading: How much can WiFi deliver?" in *Proc. ACM CoNEXT*, 2010, p. 26.
- [10] C.-L. Tsao and R. Sivakumar, "On effectively exploiting multiple wireless interfaces in mobile hosts," in *Proc. ACM CoNEXT*, 2009, pp. 337–348.
- [11] F. Malandrino, C. Casetti, C. Chiasserini, and M. Fiore, "Offloading cellular networks through ITS content download," in *Proc. IEEE SECON*, Jun. 2012, pp. 263–271.
- [12] Y. Li *et al.*, "Multiple mobile data offloading through delay tolerant networks," in *Proc. ACM CHANTS*, 2011, pp. 43–48.
- [13] O. Helgason, F. Legendre, V. Lenders, M. May, and G. Karlsson, "Performance of opportunistic content distribution under different levels of cooperation," in *Proc. Eur. Wireless*, 2010, pp. 903–910.
- [14] L. Hu, J.-Y. Le Boudec, and M. Vojnović, "Optimal channel choice for collaborative ad-hoc dissemination," in *Proc. IEEE INFOCOM*, 2010, pp. 1–9.
- [15] P. Sermpezis and T. Spyropoulos, "Delay analysis of epidemic schemes in sparse and dense heterogeneous contact environments," *Eurecom*, Biot, France, Tech. Rep. 3868, Jul. 2012.
- [16] N. Balasubramanian, A. Balasubramanian, and A. Venkataramani, "Energy consumption in mobile phones: A measurement study and implications for network applications," in *Proc. ACM IMC*, 2009, pp. 280–293.
- [17] V. Conan, J. Leguay, and T. Friedman, "Characterizing pairwise inter-contact patterns in delay tolerant networks," in *Proc. 1st Int. Conf. Auton. Comput. Commun. Syst.*, 2007, p. 19.
- [18] H. Cai and D. Y. Eun, "Crossing over the bounded domain: From exponential to power-law intermeeting time in mobile ad hoc networks," *IEEE/ACM Trans. Netw.*, vol. 17, no. 5, pp. 1578–1591, Oct. 2009.
- [19] T. Karagiannis, J.-Y. Le Boudec, and M. Vojnović, "Power law and exponential decay of intercontact times between mobile devices," *IEEE Trans. Mobile Comput.*, vol. 9, no. 10, pp. 1377–1390, Oct. 2010.
- [20] A. Passarella and M. Conti, "Characterising aggregate inter-contact times in heterogeneous opportunistic networks," in *Proc. 10th Int. IFIP TC 6 Conf. NETWORKING*, 2011, pp. 301–313.
- [21] R. Groenevelt, P. Nain, and G. Koole, "The message delay in mobile ad hoc networks," *Perform. Eval.*, vol. 62, no. 1–4, pp. 210–228, Oct. 2005.
- [22] G. Neglia and X. Zhang, "Optimal delay-power tradeoff in sparse delay tolerant networks: A preliminary study," in *Proc. ACM CHANTS*, 2006, pp. 237–244.
- [23] A. Picu and T. Spyropoulos, "Forecasting DTN performance under heterogeneous mobility: The case of limited replication," in *Proc. IEEE SECON*, 2012, pp. 569–577.
- [24] A. Chaintreau *et al.*, "Impact of human mobility on opportunistic forwarding algorithms," *IEEE Trans. Mobile Comput.*, vol. 6, no. 6, pp. 606–620, Jun. 2007.
- [25] K. J. Åström and B. Wittenmark, *Computer Controlled Systems: Theory and Design*. Upper Saddle River, NJ, USA: Prentice-Hall, 1990.
- [26] R. E. Kalman, *Topics in Mathematical System Theory*. New York, NY, USA: McGraw-Hill, 1969.
- [27] C. Hollot, V. Misra, D. Towsley, and W.-B. Gong, "A control theoretic analysis of red," in *Proc. IEEE INFOCOM*, 2001, pp. 1510–1519.
- [28] G. F. Franklin, J. D. Powell, and M. L. Workman, *Digital Control of Dynamic Systems*, 2nd ed. Reading, MA, USA: Addison-Wesley, 1990.
- [29] M. Piorkowski, N. Sarafjanovic-Djukic, and M. Grossglauser, *CRAWDAD Data Set Epi/Mobility (v.2009-02-24)*, 2009.
- [30] A. Chaintreau *et al.*, "Pocket switched networks: Real-World mobility and its consequences for opportunistic forwarding," *Comput. Lab., Univ. Cambridge*, Cambridge, U.K., Tech. Rep. UCAM-CL-TR-617, 2005.

- [31] A. Balasubramanian, R. Mahajan, and A. Venkataramani, "Augmenting mobile 3g using WiFi," in *Proc. ACM MobiSys*, 2010, pp. 209–222.
- [32] A. Sentinelli, G. Marfia, M. Gerla, L. Kleinrock, and S. Tewari, "Will IPTV ride the peer-to-peer stream?" *IEEE Commun. Mag.*, vol. 45, no. 6, pp. 86–92, Jun. 2007.



Vincenzo Sciancalepore (S'11) received the B.Sc. degree in computer engineering from the University Politecnico di Bari, Italy, in 2008, the M.Sc. degree in telecommunications engineering from the University Politecnico di Milano in 2011, and the M.Sc. degree in telematics engineering from the University Carlos III of Madrid in 2012. Currently, he is working at IMDEA Networks and enrolled as a Ph.D. student in a double Ph.D. programme between University Carlos III of Madrid and Politecnico di Milano, focusing his activity on inter-cell coordinated scheduling for LTE-Advanced networks and device-to-device communication.



Domenico Giustiniano (S'06–M'13) received the Ph.D. degree in telecommunication engineering from the University of Rome Tor Vergata in 2008. He is a Research Assistant Professor at IMDEA Networks Institute and leader of the Pervasive Wireless Systems group. Before joining IMDEA, he was a Senior Researcher and Lecturer in the Communication System Group, ETH Zurich (2012–2013). Prior to that, he was a Post-Doctoral Researcher at Disney Research Zurich (2010–2012) and Telefonica Research Barcelona (2008–2010). He devotes most of

his research to the system design and implementation of emerging wireless technologies. He is an author of more than 50 international papers and an inventor of four patents.



Albert Banchs (M'04–SM'12) received the M.Sc. and Ph.D. degrees from the Polytechnic University of Catalonia in 1997 and 2002, respectively. He was at ICSI in 1997, at Telefonica I+D in 1998, and at NEC Europe Ltd. from 1998 to 2003. He has been with the University Carlos III of Madrid since 2003. Since 2009, he also has a double affiliation as Deputy Director of the IMDEA Networks Institute. He is Editor of IEEE TRANSACTIONS ON WIRELESS COMMUNICATIONS, and of the Green Communications and Networking Series of IEEE JOURNAL OF

SELECTED AREAS IN COMMUNICATIONS (JSAC). He is currently coordinating the iJOIN European project. His research interests include the performance evaluation and algorithm design in wireless and wired networks.



Andreea Hosmann-Picu received the master's degree in telecommunications from INSA Lyon, in 2007, the M.Sc., in computer science from the Université de Lyon, in 2007, and the Ph.D. degree in information technology from ETH Zurich, in 2014. She is a Senior Data Scientist at Swisscom (Switzerland) Ltd. Prior to that, she was a Senior Researcher and Lecturer in the Communications and Distributed Systems group, University of Bern (2014–2015). She won grants from the "Georges Besse" foundation, the University of Illinois at Urbana-Champaign, and the French regions of Upper Normandy and Rhone-Alpes.