

On the Multi-Activation Oriented Design of D2D-Aided Caching Networks

Yang Cai and Andreas F. Molisch

Ming Hsieh Department of Electrical Engineering, University of Southern California, Los Angeles, CA, USA

Email: {yangcai, molisch}@usc.edu

Abstract—Caching at the wireless edge has proven to be a promising approach for efficient video distribution, especially when aided by device-to-device communication. A widely explored scheme is to sub-divide a cell into clusters, and allow one pair of users within each cluster to communicate in each time slot. As more devices are raising frequent requests for popular videos, activating multiple links simultaneously can potentially improve the throughput. However, allowing multiple links at the same time requires to solve the problems of avoiding request clashes, i.e., multiple users requesting transmission from the same caching node, as well as interference management. To address these issues, this paper proposes new designs of both the caching policy and the transmission policy (i.e., link scheduling and power control). Furthermore, the duration of each time slot is optimized to improve the throughput. Finally, some numerical results demonstrate the performance gain of the proposed designs.

I. INTRODUCTION

With the rapid proliferation of smartphones, video streaming has become the major consumer for wireless data, and is expected to account for 80% of cellular traffic by 2022 [1]. To satisfy the dramatically increasing demand for video data, general approaches to improve the capacity of wireless systems (e.g., network densification, new spectrum and massive antennas) are possible solutions, however, at a high expense for the additional infrastructure. Conversely, caching at the wireless edge, which exploits the pronounced (asynchronous) content reuse of video by duplicating the video contents and bringing them closer to the users, can address the problem at a much lower cost [2]. Besides, device-to-device (D2D) communication techniques can further improve the spectrum efficiency of content delivery. Therefore, the D2D-aided caching network is a promising paradigm for video sharing, and has attracted much attention, e.g., [3]–[5] and references therein.

To evaluate the performance of a caching network, various metrics have been proposed [6], such as cache hit rate [7], download delay, and energy efficiency. An even more important criterion is the throughput, which characterizes the capacity of a caching network to satisfy users' needs. To optimize throughput, two policies need to be established: (i) the caching policy, i.e., which files are cached by which device, and (ii) the transmission policy, i.e., how the D2D transmissions are scheduled, and with what power. These two policies interact, which motivates a joint optimization of them.

This work was supported in part by the National Science Foundation (NSF) under CNS-1816699 and CCF-1423140.

In the past, a widely investigated D2D transmission policy was as follows. A cell is sub-divided into some clusters, and frequency reuse is employed, such that adjacent clusters do not occupy the same time-frequency resources. Within each cluster, the transmit power is chosen such that the devices can communicate with each other, and only one link is active at a time. For this transmission policy, and under the assumption that each device caches files at random, [2] established the optimum caching policy, i.e., the probability mass function (PMF) according to which the random caching is performed.

In this paper, we will consider an alternative transmission scheme, called multi-activation (MA), which allows multiple links in a cluster to be active simultaneously. Since a common (popular) video can be desired by several users so that more copies of the files are needed to enable clash-free access for each user, which affects the caching policy design. Besides, deciding the combination of active links, and coordinating their transmit powers, can effectively reduce the interferences between the transmissions, leading to higher throughput [8]. Furthermore, the duration of each time slot (called a *transmission session*) will be optimized. Thus the number of simultaneous activations trades off the benefits of *higher data rate on a single link and more links*, which enhances the throughput of the entire network.

The contributions of this work are summarized as follows:

- We propose a new design of caching distribution, which maximizes the number of *non-overlapping* links, and an approximate solution to the problem is provided;
- We formulate the problems of link scheduling and power control to *maximize the minimum data rate*, and develop the corresponding algorithms to solve them;
- We derive a statistics-based model to approximate the D2D-contributed throughput in dynamic scenario, which can be applied to optimize the session duration.

II. SYSTEM MODEL

We consider a cellular communication system with a single base station (BS), and N uniformly distributed users. The video library contains N_f files of interest, and the size of each file is F , i.e., for simplicity we assume that all files have the same size. Each user is equipped with a cache that can store S videos, and a single antenna that enables the communications between any pair of the users (D2D link), or the user and the BS (cellular link). The BS is assumed to connect to the core network, and can provide any videos desired by the users.

A. Request and Service Model

Suppose the number of requests of each user is modeled by a Poisson process with intensity λs^{-1} , and the desired file of each request is independent and identically distributed (i.i.d.), according to the preference distribution $\mathbf{p}_r \in \mathbb{R}^{N_f}$, with its k -th element $p_{r,k} \geq 0$ denoting the probability that file k is needed.¹ Each user keeps at most one request, i.e., if an existing request is not served by the time a new one arises, the recent one will replace the old one.

The videos cached by each user are subject to a common caching policy $\mathbf{p}_c$, and the stored contents do not change over time.² With videos stored in the user devices in addition to the BS, the desired file can be obtained in the following three ways (in order of priority):

- (*Self-supplied*) if the user itself caches the video, the file can be read from the memory (without any delay);
- (*D2D-supplied*) if the video is not stored in its own cache, but can be found in other users, then a D2D link can be constructed to convey the file;
- (*BS-supplied*) if the video is not cached in any user, then the file will be provided by the BS.

Note that this model has been used in a large number of edge-caching papers. Then the probabilities for a request to be self-supplied and D2D-supplied are respectively given by

$$q_s = \sum_{i=1}^{N_f} \mathbf{p}_{r,i} [1 - (1 - \mathbf{p}_{c,i})^S] \quad (1)$$

$$q_d = \sum_{i=1}^{N_f} \mathbf{p}_{r,i} [1 - (1 - \mathbf{p}_{c,i})^{NS}] - q_s. \quad (2)$$

In addition, a user will be called D2D-supplied if it keeps a D2D-supplied request.

B. D2D Communication Model

The D2D communication between the users is modeled as follows. Time is uniformly slotted, and each slot is called a transmission session with duration h . The MA scheme is considered, i.e., multiple D2D links are activated as the transmission session starts, with each user involved in at most one link. To guarantee the quality of transmission, all the deliveries are required to be completed within the session. The spectral band dedicated to D2D communications is non-overlapping with the cellular band, while all the active D2D links use the same band (time-frequency unit) and treat the signals from other links as noise. The path loss of the propagation channel is modeled by $\eta(r) = 10\alpha \log_{10} r + 10 \log_{10} \eta_0$ dB, where α is the path loss exponent, r is the distance the signal travels, and $10 \log_{10} \eta_0$ is a constant. We neglect both small-scale fading and shadowing. The former can be mitigated or eliminated through frequency diversity, while the latter can be either averaged out through temporal diversity, or accounted for by mapping the shadowing variations onto “equivalent”

distance modifications [9]. Based on this model, the signal-to-interference-and-noise-ratio (SINR) for a link $e \in \mathcal{E}$ at the beginning of the transmission session is given by³

$$\gamma_e = \frac{a_e p_e r_e^{-\alpha}}{I_e + \eta_0 \sigma_e^2} \quad (3)$$

where $\mathcal{E}$ is the set of all the potential D2D links; a_e is an indicator variable, which equals 1 if link e is active in the session, and 0 otherwise; p_e and σ_e^2 represent the powers of the transmitted signal and the additive noise, respectively; r_e denotes the length of the link; I_e is the interference caused by other links, given by $I_e = \sum_{s \in \mathcal{E} \setminus \{e\}} a_s p_s r_{s,e}^{-\alpha}$, where $r_{s,e}$ is the distance from the transmitter of link s to the receiver of link e (as mentioned above, the inter-cluster interference is negligible). Furthermore, if a bandwidth of B is allocated for D2D transmission in a cluster, the data rate of the link e is lower bounded as

$$R_e = B \log_2(1 + \gamma_e). \quad (4)$$

C. Problem Formulation

The D2D communication model in the previous section imposes the following constraints on the parameters. First, since each user is associated with at most one active link, the binary vector $\mathbf{a} = [a_e]_{e \in \mathcal{E}}$ must satisfy

$$\sum_{e \in \mathcal{E}(i)} a_e \leq 1 \quad (1 \leq i \leq N). \quad (5)$$

where $\mathcal{E}(i)$ contains all the links involving user i . Second, we require the transmit power for any link not to exceed some maximum power $P_{\max}$, i.e., the vector $\mathbf{p} = [p_e]_{e \in \mathcal{E}}$ satisfies

$$0 \leq p_e \leq P_{\max} \quad (\forall e \in \mathcal{E}). \quad (6)$$

Finally, to ensure the file can be delivered within a session of duration h , a sufficient condition is to require $R_e h \geq F$ for all active links. Then by (3), we have for $\forall e \in \mathcal{E}$,

$$\gamma_e \geq a_e \gamma_0, \text{ with } \gamma_0 \triangleq \left[\exp\left(\frac{F \ln 2}{Bh}\right) - 1 \right]. \quad (7)$$

As long as that the above constraints are not violated, more D2D links should be activated to improve the D2D-contributed throughput, which is given by

$$T = \lim_{n \rightarrow \infty} \frac{F}{nh} \sum_{t=1}^n L(t) \quad (8)$$

where $L(t)$ is the number of activations in session t . Our goal is to maximize T subject to (5) – (7). Intuitively, T is affected by the number of potential D2D links (related to caching policy $\mathbf{p}_c$), the interference management (link scheduling and power control), as well as the threshold of the SINR γ_0 (determined by h). However, the effects of the parameters are interdependent, and it is difficult to derive an explicit expression. To obtain a suboptimal but tractable solution, the problem is decoupled into three parts (see Fig. 1):

¹Note that the requests at different times are i.i.d.; this does *not* mean that the request probabilities for *different files* are identical.

²It can be interpreted by a scheme in which the caches are filled overnight (when wireless transmissions are cheap), and then remain unchanged during the day, when cache refreshes would consume expensive resources.

³We emphasize that this is the SINR at the beginning of the session, when all the links are activated. As some transmissions terminate, the SINR of the rest links will be improved due to the reduced interference.

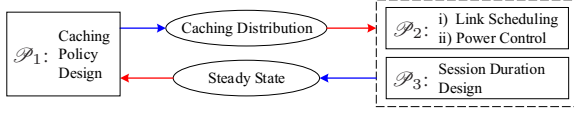


Fig. 1. The iterative design procedure, with the dashed box indicating the network evolution. The caching distribution $\mathbf{p}_c$ is used to optimize the session duration h , leading to a steady state for the number of requests U , which will in turn affect the design of $\mathbf{p}_c$ in the next iteration.

- $\mathcal{P}_1$: How to design the caching policy to provide more D2D links that can be activated simultaneously;
- $\mathcal{P}_2$: How to select a given number of links, and coordinate their transmit powers to improve the data rate;
- $\mathcal{P}_3$: How to decide the transmission session duration to trade off the benefits of higher SINR and more links.

III. CACHING POLICY DESIGN

Under the MA scheme, several users can request files at the same time. When the number of copies for a commonly desired file is less than the number of requests, some requests cannot be satisfied. This motivates to store more copies of popular files, which may reduce the diversity of the cached videos, leading to a tradeoff between preference emphasis and content coverage when designing the caching policy.

A. Expected Number of Non-overlapping Requests

Consider a single transmission session, with the number of requests given by U .⁴ The users not requesting any file will act as potential sources, and the caches on these devices are aggregated into a virtual caching center (VCC) of size $V = (N - U)S$, which is managed by the BS. The requests are called *non-overlapping* if distinct copies for the desired files can be found, so that the caches can transmit in parallel.

Recall that the requested (or cached) videos are i.i.d. and follow the distribution $\mathbf{p}_r$ (or $\mathbf{p}_c$). Therefore, denote the number of requests for file i by u_i , and the corresponding number of copies by v_i , then they both follow the binomial distribution, i.e., $u_i \sim \mathcal{B}(U, \mathbf{p}_{r,i})$ and $v_i \sim \mathcal{B}(V, \mathbf{p}_{c,i})$, with the PMF of a binomial random variable $x \sim \mathcal{B}(n, p)$ given by

$$\mathbb{P}\{x = k\} = \binom{n}{k} p^k (1-p)^{n-k} \triangleq b(k; n, p). \quad (9)$$

Among all the u_i requests, only the non-overlapping ones are qualified for clash-free access, and the number is given by $z_i = \min(u_i, v_i)$. Furthermore, by summing over all possible videos $1 \leq i \leq N_f$ and taking the expectation, we can derive the *expected number of non-overlapping requests* as

$$Z = \mathbb{E}\left\{\sum_{i=1}^{N_f} z_i\right\} = \sum_{i=1}^{N_f} \mathbb{E}\{z_i\}. \quad (10)$$

Given the number of requests U , $Z = Z(\mathbf{p}_c)$ is affected by the caching distribution, and we define the optimal caching policy $\mathbf{p}_c^*$ as the one that maximizes $Z(\mathbf{p}_c)$, i.e.,

$$\mathcal{P}_1: \max_{\mathbf{p}_c \geq 0} Z(\mathbf{p}_c), \text{ s. t. } \sum_{i=1}^{N_f} \mathbf{p}_{c,i} = 1. \quad (11)$$

⁴In fact, the number of requests is a random variable and changes over time. It is assumed to be deterministic and fixed here for tractability. In practice, we use the expected number of requests as a substitute for it.

To derive the relationship between Z and $\mathbf{p}_c$, note that the cumulative distribution function (CDF) for z_i is given by

$$F_{z_i}(s) = F_{u_i}(s) + F_{v_i}(s) - F_{u_i}(s)F_{v_i}(s) \quad (12)$$

where F_{u_i} and F_{v_i} denote the CDFs for u_i and v_i (incomplete beta functions). Since $z_i \in [0, U]$ is non-negative and discrete, its expected value is given by $\sum_{s=0}^U [1 - F_{z_i}(s)]$, i.e.,

$$\mathbb{E}\{z_i\} = U\mathbf{p}_{r,i} - \sum_{k=0}^U g_i(k)b(k; V, \mathbf{p}_{c,i}) \quad (13)$$

where the sequence $\{g_i(k) : 0 \leq k \leq U\}$ is defined as

$$g_i(k) = \sum_{s=k}^U [1 - F_{u_i}(s)] = \sum_{s=1}^{U-k} s b(k+s; U, \mathbf{p}_{r,i}) \quad (14)$$

which is only related to the preference distribution. Substitute (13) into (10), and the objective function is given by

$$Z(\mathbf{p}_c) = U - \sum_{i=1}^{N_f} \sum_{k=0}^U g_i(k)b(k; V, \mathbf{p}_{c,i}). \quad (15)$$

B. Approximated Solution

Although (15) is the accurate expression for Z , it is not straightforward to address $\mathcal{P}_1$ due to the complicated structure of the sequence g_i . To derive a simpler solution, we make the following assumptions/approximations:

- (a) The probability to request or cache any video is small, i.e., $\mathbf{p}_{r,i} \ll 1$ and $\mathbf{p}_{c,i} \ll 1$ for $1 \leq i \leq N_f$;
- (b) The reciprocals of factorials is approximated by a geometrical sequence,⁵ i.e., $(n!)^{-1} \approx aq^n$, where a and q are fitting parameters.⁶

Proposition 1: Assuming that (a) and (b) hold, the sequence g_i can be approximated by a geometrical sequence, given by

$$g_i(k) \approx \left[\frac{aqU\mathbf{p}_{r,i} \exp(-U\mathbf{p}_{r,i})}{(1 - qU\mathbf{p}_{r,i})^2} \right] (qU\mathbf{p}_{r,i})^k \triangleq A_i Q_i^k. \quad (16)$$

Furthermore, the objective function can be approximated by

$$Z(\mathbf{p}_c) \approx U - \sum_{i=1}^{N_f} A_i \exp[-(1 - Q_i)V\mathbf{p}_{c,i}] \quad (17)$$

which is a concave function over $\mathbf{p}_c$.

Proof: The sketch of proof is provided as follows. Assumption (a) implies that u_i and v_i can be approximated by Poisson variables, and we apply approximation (b) to the PMF of the Poisson distribution. Besides, (a) also suggests that (14) can be approximated by the infinite sum, which leads to (16); similar techniques can be used to obtain (17). The concavity can be shown by deriving the Hessian matrix. $\square$

The accuracy of the approximate model will be verified by the numerical results (see Fig. 2(b)). Using (17) to approximate the objective function, (11) is converted to a convex program, which can be efficiently solved.

⁵Although the Stirling's formula is widely used for $n!$ with large n , the terms $(n!)^{-1}$ with small n will dominate in (15) (by assumption (a)). The geometrical model can provide good approximation in such regime.

⁶The fitting result is $a = 1.146$ and $q = 0.618$. However, instead of using this result, we directly employ (16), i.e., $g_i(k) = A_i Q_i^k$, as the model to obtain an approximate sequence for g_i , which yields higher accuracy.

IV. TRANSMISSION POLICY DESIGN

In this section, we consider the problems related to the design of transmission policy, i.e., the interference management and the optimization of the session duration h . First, we assume that the session duration is fixed, and develop the criteria for interference management. In this context, more D2D links should be activated to improve the throughput as long as (7) is not violated. Equivalently, this problem can be expressed as: given a certain number of activations, how to maximize the *minimum SINR* (max-min-SINR) among the active links by link scheduling and power control.

Next, we consider the effect of the session duration h , which determines the number of activations. When h is small, fewer links will be activated, resulting in insufficient use of the spectral resources. On the other hand, for large h , more activations can cause strong interferences between each other, leading to poor throughput. To sum up, the choice of h should trade off the benefits of higher SINR on a single link and more links, in order to optimize the throughput.

A. Link Scheduling and Power Control

Suppose the positions of the users are given, as well as all the potential links $\mathcal{E}$. A given number L_0 of links are to be activated, which are subject to the criteria of max-min-SINR. By introducing an auxiliary parameter γ to represent the minimum SINR, the problem of link selection and power coordination is formulated as

$$\mathcal{P}_2: \max_{\mathbf{a}, \mathbf{p}} \gamma, \text{ s.t. } \gamma_e \geq a_e \gamma, \mathcal{L}, \mathcal{P} \quad (18)$$

where $\mathcal{L}$ (including (5) and $\sum_{e \in \mathcal{E}} a_e = L_0$) and $\mathcal{P}$ (i.e., (6)) denote the constraints on active links and transmit powers, respectively. This is a mixed-integer programming problem, which is difficult to solve in a joint manner. A simpler solution is to address the link scheduling and the power control problem in sequence. The same criterion, i.e., max-min-SINR, is applied when addressing each separate problem.

1) *Link Scheduling*: In this step, we aim to select L_0 links from $\mathcal{E}$ to max-min-SINR, assuming that the transmit powers of all the active links are $P_{\max}$. By substituting $\mathbf{p} = P_{\max} \mathbf{a}$ into (18), the constraint $\gamma_e \geq a_e \gamma$ is equivalent to

$$a_e \left(\frac{r_e^{-\alpha}}{\gamma} - \frac{\eta_0 \sigma_e^2}{P_{\max}} \right) \geq \sum_{s \in \mathcal{E} \setminus \{e\}} \tilde{a}_{es} r_{s,e}^{-\alpha} \quad (19)$$

where $\tilde{\mathbf{a}} = [\tilde{a}_{es}]_{\{e \in \mathcal{E}, s \in \mathcal{E}\}}$ is an auxiliary binary vector that satisfies the following constraints

$$\tilde{a}_{es} \leq a_e, \tilde{a}_{es} \leq a_s, \tilde{a}_{es} \geq a_e + a_s - 1. \quad (20)$$

The problem then becomes to *maximize* γ over $\mathbf{a}$, subject to (19), (20) and $\mathcal{L}$. To address this problem, note that when γ is given, it becomes a feasibility problem of integer linear programming (ILP), which can be approximately solved by the branch-and-bound algorithm. Therefore, we can determine the maximum γ_1^* that makes the constraints feasible by bisection search; the feasible point when $\gamma = \gamma_1^*$ is denoted by $\mathbf{a}^*$, which corresponds to the optimal selection of the links.

Remark 1: Although the accurate formulation for this sub-problem is provided, the solution can cause high computation complexity when the network is large. In that case, a greedy algorithm will be employed, i.e.,

- (i) Start from an empty set $\mathcal{A}$ (active link set);
- (ii) For each link $s \in \mathcal{E} \setminus \mathcal{A}$, calculate the min SINR $\underline{\gamma}_s$ when link s and all links in $\mathcal{A}$ are activated;
- (iii) Add the link with the largest minimum SINR (i.e., $e^* = \arg \max_{s \in \mathcal{E} \setminus \mathcal{A}} \underline{\gamma}_s$) to the set $\mathcal{A}$;
- (iv) Repeat (ii) and (iii) until $\mathcal{A}$ includes L_0 active links.

2) *Power Control*: After the previous step, the active D2D links are decided and collected in $\tilde{\mathcal{E}} = \{e : a_e^* = 1\}$. The remaining problem for this step is to *maximize* γ over the variables $\{\gamma, \mathbf{p}\}$, subject to $\mathcal{P}$ and (7), i.e.,

$$\left(\frac{r_{e_i}^{-\alpha}}{\gamma} \right) \tilde{p}_i - \sum_{j \in \tilde{\mathcal{E}} \setminus \{i\}} r_{e_j, e_i}^{-\alpha} \tilde{p}_j \geq \eta_0 \sigma_{e_i}^2, \quad \forall i, j \in \tilde{\mathcal{E}} \quad (21)$$

where $\tilde{\mathbf{p}} = [\tilde{p}_e]_{e \in \tilde{\mathcal{E}}}$ collects the transmit powers of all the active links (the power of inactive links are set as zero), and e_i denotes the index for the i th link in $\tilde{\mathcal{E}}$.

To address this sub-problem, we notice that given the minimum SINR γ , the objective function becomes a constant, and the above constraints are linear over $\tilde{\mathbf{p}}$. This fact suggests the following solution to the problem: for a given value of γ , checks the feasibility of the linear system corresponding to the constraints (which can be solved efficiently); repeat this procedure and use the bisection search to determine the optimal value for γ^* , as well as the transmit power $\tilde{\mathbf{p}}^*$.

B. Optimal Duration Design

1) *Statistics-Based Method*: Based on the caching policy and the transmission policy proposed in previous sections, we can use a one-dimensional search to determine the optimal session duration h by simulations. However, since this is a dynamic system, and h is continuous-valued, a huge number of experiments will be needed to achieve a reasonable accuracy, and each one is of high computation complexity (for throughput to converge), which is difficult to realize in practice.

Rather than running such simulations, we can gather the statistics of the data rate with much lower cost. More concretely, each experiment considers a single transmission session, where the number of D2D-supplied requests is $W \in [0, N]$, and the number of active links is $L \in [1, L_{\max}]$, with $L_{\max}$ denoting the maximum number of activations in a session; we also assume that the users are equally likely to place the requests. For each pair of W and L , we simulate the SINR for K_0 realizations, with the result of the k th experiments denoted by $\gamma_k(W, L)$. Then for given h , the probability that “the number of activations is L ” can be estimated by

$$\hat{q}_L(W, h) = \sum_{k=1}^{K_0} \frac{\mathbb{I}\{\gamma_k(W, L) \geq \gamma_0, \gamma_k(W, L+1) < \gamma_0\}}{K_0} \quad (22)$$

where the SINR threshold $\gamma_0 = \gamma_0(h)$ is given in (7); and $\mathbb{I}\{\mathcal{A}\}$ equals to 1 if the event $\mathcal{A}$ is true, and 0 otherwise. The

throughput then follows that

$$\hat{T}(h) = \lim_{n \rightarrow \infty} \left[\frac{1}{n} \mathbb{E} \left\{ \sum_{t=1}^n \sum_{L=1}^{L_{\max}} \frac{FL}{h} \hat{q}_L(w(t), h) \right\} \right] \quad (23)$$

where $w(t)$ is the number of D2D-supplied requests at the beginning of session t (which is also the number of D2D-supplied users at that time, since each user keeps only one request), and its distribution is derived in next section.

2) *Stationary Distribution of $w(t)$* : In this section, we assume that the distribution for the number of activations is given by (22), i.e., $\mathbb{P}\{L(t) = L\} = \hat{q}_L(w(t), h)$ where $L(t)$ is the number of activations in session t . Under this assumption, we will derive a stationary distribution for $w(t)$.

Note that the users can be divided into two types: the first type includes all the D2D-supplied users at the beginning of session t , except the ones that will be served in the session (totally $N_I = w(t) - L(t)$ users); and the other users are gathered in the second type (totally $N_{II} = N - N_I$ users). Next we calculate the probabilities for a user in the first and second type to become D2D-supplied in session $t+1$, denoted by q_I and q_{II} , respectively. For any user, if it raises some (possibly more than one) requests during session t , and the last request is D2D-supplied, then it will become D2D-supplied in $t+1$. Besides, for a first-type user, in addition to this situation, if it does not raise any request during session t , then it remains to be D2D-supplied in $t+1$. Thus, the probabilities are $q_I = q_0 + (1 - q_0)q_d$ and $q_{II} = (1 - q_0)q_d$, where $q_0 = \exp(-\lambda h)$ is the probability that a user raises no request during session t , and q_d is given by (2).

Based on above results, the number of D2D-supplied users at session $t+1$ is $w(t+1) = w_I + w_{II}$, where $w_I \sim \mathcal{B}(N_I, q_I)$ and $w_{II} \sim \mathcal{B}(N_{II}, q_{II})$ are binomial random variables. Since N_I and N_{II} are only related to $w(t)$, it implies that $w(t)$ can be modeled by a homogeneous Markov chain (with finite states $\{0, \dots, N\}$), and the state transition matrix $\mathbf{Q}$ is given by

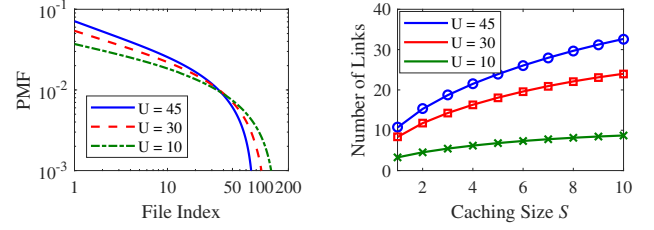
$$Q_{ij} = \sum_{L=1}^{L_{\max}} \hat{q}_L(i, h) \phi_{i,L}(j) \quad (24)$$

where $\phi_{i,L}(n) = b(n; i - L, q_I) \otimes b(n; N - (i - L), q_{II})$ with $b(k; n, p)$ being the PMF of the Binomial distribution (given by (9)) and $\otimes$ denoting the convolution operator.

Finally, we can calculate the stationary distribution π_w for $w(t)$ by solving $\pi_w = \pi_w \mathbf{Q}$, which can be applied to derive the D2D-contributed throughput (23), i.e.,

$$\hat{T}(h) = \sum_{k=0}^N \pi_{w,k} \left[\sum_{L=1}^{L_{\max}} \frac{FL}{h} \hat{q}_L(k, h) \right] \quad (25)$$

where $\pi_{w,k}$ denotes the k th element of π_w . This expression can be calculated efficiently, which enables a one-dimensional search to determine the optimal value of h . Furthermore, the expected number of D2D-supplied requests can be calculated as $\sum_{k=0}^N k \pi_{w,k}$; on the other hand, it is $U q_d$ according to Section III-A (the number of total requests, multiplied by the probability for a request to be D2D-supplied). Consistency of the two results leads to $U = \sum_{k=0}^N k \pi_{w,k} / q_d$, which will be used for the caching policy design in the next iteration.



(a) Caching distributions ($S = 1$) (b) Number of non-overlapping links

Fig. 2. (a) The MA policies under different U ; (b) the simulation (marks) and the analytical (lines) results for the number of non-overlapping links.

V. NUMERICAL RESULTS

Consider a squared cluster of size 100 m $\times$ 100 m, with $N = 100$ users. The users move at a speed of 1.5 m/s, and change their directions randomly every 10 s (they reflect when hitting the boundary). There are $N_f = 500$ videos of interest, with the size of each file equals to $F = 250$ MB. The users request the files at an intensity of $\lambda = 1 \times 10^{-2}$ /s, and the preference model suggested by the UMass Amherst youtube experiment is adopted, which is a Zipf distribution given by $p_{r,i} = i^{-0.6} / (\sum_{j=1}^{N_f} j^{-0.6})$. Each user caches $S = 1$ file, except in the first part of experiment. Drawing inspiration from typical WiFi parameters, a bandwidth of $B = 20$ MHz is allocated for D2D communication in the cluster. The maximum transmit power is set as $P_{\max} = 100$ mW, and all the links have the same noise power of $\sigma_e^2 = -97$ dBm, corresponding to a 3 dB receiver noise figure. The path loss is modeled by $\eta(r) = 36.8 \log_{10}(r) + 37.6$. To account for the typical low efficiency of handset antennas, the antenna gain at each device is assumed to be -3.5 dB. Besides, the BS transmits videos to each destination at a constant data rate $R_b = 200$ kb/s, and $N_b = 10$ users can be served simultaneously.

A. Caching Policy

First, we present the designed caching policy (called MA policy), using different numbers of requests U as the inputs. We note that for $S = 1$ and the network parameters defined above, the iterative optimization procedure of Section III and IV results in $U = 45$; the initial value of the iteration is chosen to be $U = 10$. The number of non-overlapping links is shown, under various caching size $S \in \{1, \dots, 10\}$.

The results are shown in Fig. 2. First, we observe that given a limited caching size, the designed MA policy pays more attention to the popular files for large U , while it increases the diversity of the cached videos as U decreases. This is intuitive since a larger U can lead to a higher probability for a popular file to be desired more than once, and more memories should be allocated to store their copies. Then, Fig. 2(b) shows the number of non-overlapping links obtained from the simulation experiments, as well as the results calculated by (17), which validates the accuracy of the approximated model.

B. Transmission Policy

Next, we evaluate the performance of the proposed transmission policy. A single transmission session is considered,

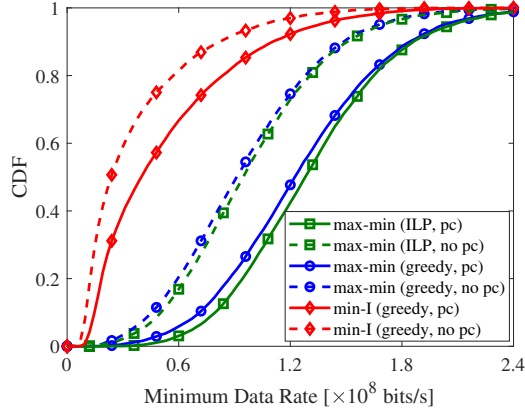


Fig. 3. The CDFs of the minimum data rate under different scheduling strategies, with (solid lines) or without (dashed lines) power control.

and we assume the number of active links is $L_0 = 3$. The number of D2D-supplied requests is set as $U = 45$, and the MA caching policy obtained in the previous section is employed. The proposed max-min approach is compared to the min interference (min-I) scheduling [8], whether with or without power control (in the latter case, all active links transmit signals with power $P_{\max} = 100$ mW).

Fig. 3 depicts the CDFs of the minimum data rate achieved by different strategies. We first focus on the two implementations with the min-max criterion, and find that the ILP-based algorithm performs slightly better than the greedy-based counterpart, however, at the cost of much higher complexity. Next, when selecting the links in a greedy manner, the proposed max-min strategy outperforms the min-I strategy, by 1.26 to 0.52 in the case with power control, and 0.94 to 0.36 without power control (which are the mean values of the minimum data rate, all in $\times 10^8$ bits/s). Finally, according to the result, the link scheduling has a more essential impact on the data rate compared to the power control.

C. The Effect of h

Finally, we show the effect of the session duration h on the local throughput (self-supplied and D2D supplied). The proposed transmission policy is employed, setting the maximum number of activations as $L_{\max} = 5$. We compare different policies, including 1) the MA caching (using $U = 45$ as input), 2) the max-hit caching and 3) the selfish caching.

As we can see from Fig. 4: firstly, the analysis based on the Markov model, combined with the statistics of the data rate (we create $K_0 = 1000$ realizations for each (W, L) pair), provides good approximations to the simulation results.⁷ Second, the proposed MA policy outperforms the max-hit policy at corresponding optimal values of h (which are around $h = 14$ for both policies) by 20%. Finally, the performance curve of TDMA (only one link is activated in each session) also exhibits a maximum as h varies (in this case the optimal

⁷The analytical result for the throughput contributed by self-caching is $\hat{T}_s = FN\lambda q_s$ with q_s given by (1).

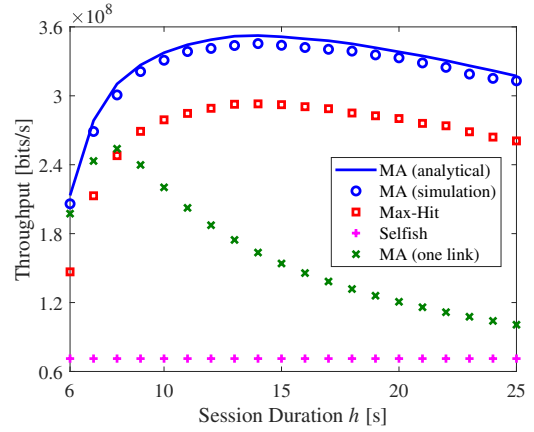


Fig. 4. The comparison of local throughput under various caching policies, and different link activation schemes.

$h = 8$); and compared to it, the MA scheme can achieve 35% of gain in the local throughput. Furthermore, the maximum of the MA policy is broader, making it more robust to achieve.

VI. CONCLUSION

In this paper, we investigated the caching network design for a MA scenario, aiming towards higher throughput. First, we designed a caching policy that allows a maximum number of clash-free access, and provided an efficient approximated solution to the problem. Then, we proposed a link scheduling and power control algorithm that maximizes the minimum SINR given the number of activations. Based on the statistics of the data rate, the transmission session duration was optimized. Finally, numerical results were provided to validate the performance gain of the MA and the proposed designs.

REFERENCES

- [1] Cisco Visual Networking Index, "Global Mobile Data Traffic Forecast Update, 2017–2022," *Cisco White Paper*, Feb. 2019.
- [2] M. Ji, G. Caire, and A. F. Molisch, "Wireless device-to-device caching networks: basic principles and system performance," *IEEE J. Sel. Areas Commun.*, vol. 34, no. 1, pp. 176–189, Jan. 2016.
- [3] R. Amer, M. M. Butt, M. Bennis, and N. Marchetti, "Inter-cluster cooperation for wireless D2D caching networks," *IEEE Trans. Wireless Commun.*, vol. 17, no. 9, pp. 6108–6121, Sep. 2018.
- [4] B. Chen, C. Yang, and A. F. Molisch, "Cache-enabled device-to-device communications: offloading gain and energy cost," *IEEE Trans. Wireless Commun.*, vol. 16, no. 7, pp. 4519–4536, Jul. 2017.
- [5] Y. Wang, X. Tao, X. Zhang, and Y. Gu, "Cooperative caching placement in cache-enabled D2D underlaid cellular network," *IEEE Commun. Lett.*, vol. 21, no. 5, pp. 1151–1154, May 2017.
- [6] M.-C. Lee and A. F. Molisch, "Caching policy and cooperation distance design for base station-assisted wireless D2D caching networks: throughput and energy efficiency optimization and tradeoff," *IEEE Trans. Wireless Commun.*, vol. 17, no. 11, pp. 7500–7514, Nov. 2018.
- [7] Z. Chen, N. Pappas, and M. Kountouris, "Probabilistic caching in wireless D2D networks: cache hit optimal versus throughput optimal," *IEEE Commun. Lett.*, vol. 21, no. 3, pp. 584–587, Mar. 2017.
- [8] L. Zhang, M. Xiao, G. Wu, and S. Li, "Efficient scheduling and power allocation for D2D-assisted wireless caching networks," *IEEE Trans. Commun.*, vol. 64, no. 2, pp. 2438–2452, Jun. 2016.
- [9] H. S. Dhillon and J. G. Andrews, "Downlink rate distribution in heterogeneous cellular networks under generalized cell selection," *IEEE Wireless Commun. Lett.*, vol. 1, no. 3, pp. 42–45, Feb. 2014.