

Online Caching with Optimistic Learning

Mhaisen, Naram; Iosifidis, George; Leith, Douglas

DOI

[10.23919/IFIPNetworking55013.2022.9829806](https://doi.org/10.23919/IFIPNetworking55013.2022.9829806)

Publication date

2022

Document Version

Final published version

Published in

Proceedings of the 2022 IFIP Networking Conference (IFIP Networking)

Citation (APA)

Mhaisen, N., Iosifidis, G., & Leith, D. (2022). Online Caching with Optimistic Learning. In *Proceedings of the 2022 IFIP Networking Conference (IFIP Networking)* (pp. 1-9). Article 982980 (2022 IFIP Networking Conference, IFIP Networking 2022). IEEE. <https://doi.org/10.23919/IFIPNetworking55013.2022.9829806>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

Online Caching with Optimistic Learning

Naram Mhaisen*, George Iosifidis*, Douglas Leith†

*Software Technology Group, Delft University of Technology, Netherlands.

†School of Computer Science and Statistics, Trinity College Dublin, Ireland.

Abstract—The design of effective online caching policies is an increasingly important problem for content distribution networks, online social networks and edge computing services, among other areas. This paper proposes a new algorithmic toolbox for tackling this problem through the lens of *optimistic* online learning. We build upon the Follow-the-Regularized-Leader (FTRL) framework which is developed further here to include predictions for the file requests, and we design online caching algorithms for bipartite networks with fixed-size caches or elastic leased caches subject to time-average budget constraints. The predictions are provided by a content recommendation system that influences the users viewing activity, and hence can naturally reduce the caching network’s uncertainty about future requests. We prove that the proposed optimistic learning caching policies can achieve *sub-zero* performance loss (regret) for perfect predictions, and maintain the best achievable regret bound $O(\sqrt{T})$ even for arbitrary-bad predictions. The performance of the proposed algorithms is evaluated with detailed trace-driven numerical tests.

I. INTRODUCTION

Motivation. The quest for efficient data caching policies spans more than 50 years and remains today one of the most important research areas for wireless and wired communication systems [1]. Caching was first studied in computer systems where the aim was to decide which files to store in fast-accessible memory segments (*paging*) [2]. Its scope was later expanded due to the explosion of Internet web traffic [3] and the advent of content distribution networks (CDNs) [4], and was recently revisited as a technique to improve the operation of wireless networks through edge caches [5] and on-device caching [6]. A common challenge in these systems is to design an online policy that decides which files to store at a cache, without knowing the future file requests, so as to maximize the cache *hits* or some other more general cache-related performance metric.

There is a range of online caching policies that tackle this problem under different assumptions about the request arrivals. Policies such as the LRU and LRU are widely-deployed, yet their performance deteriorates when the file popularity is non-stationary, i.e., the requests are drawn from a time-varying probability distribution [7]–[9]. This motivated modeling non-stationary request patterns [10], [11] and optimizing accordingly the caching decisions [12], [13]. Another line of work relies on techniques such as reinforcement learning to estimate the request probabilities and make caching decisions [14], [15]; but typically these solutions do not scale nor offer optimality bounds. Caching was studied as an online learning problem in [16], [17] for a single-cache system; and in its more

general form in [18] that proposed an online gradient descent (OGD) caching policy. Interesting follow-up works include sub-modular policies [19], online mirror-descent policies [20], and the characterization of their performance limits [21]. The advantage of these online learning-based caching policies is that they are scalable, do not require training data, and their performance bounds are *robust* to any possible request pattern.

An aspect that has not been studied, however, is whether predictions about future requests can improve the performance of these learning-based caching policies. This is important in modern caching systems where most often the users receive content viewing recommendations from a recommendation system (*rec-sys*). For instance, recommendations are a standard feature in streaming platforms such as YouTube and Netflix [22]; but also in online social network platforms such as Facebook and Twitter, which moderate the users’ viewing feeds [23]. Not surprisingly, the interplay between recommendations and caching has attracted substantial attention and recent works devised static joint policies aiming, e.g., to increase the cache hit rate or reduce the routing costs by recommending to users already-cached files [24], [25].

Changing vantage point, one can observe that since recommendations bias the users towards viewing certain content files, they can effectively serve as predictions of the forthcoming requests. This prediction information, if properly leveraged, can hugely improve the efficacy of caching policies, transforming their design from an online learning to an online optimization problem. Nevertheless, the caching policy needs to adapt to the accuracy of recommendations and the users propensity to follow them – which is typically unknown and potentially time-varying. Otherwise, the caching performance might as well deteriorate by following these misleading request *hints*. The goal of this work is to tackle exactly this challenging new problem and *propose online learning-based caching policies which leverage predictions (of unknown quality) to achieve robust performance bounds*.

Contributions. Our approach is based on the theory of Online Convex Optimization (OCO) that was introduced in [26] and has since been applied in different problems [27]. The basic premise of OCO is that a learner (here the caching system) selects in each slot t a decision vector x_t from a convex set $\mathcal{X}$, without knowing the t -slot convex performance function $f_t(x)$, that change with time. The learner’s goal is to minimize the growth rate of *regret* $R_T = \sum_{t=1}^T f_t(x^*) - f_t(x_t)$, where $x^* = \arg \max_{x \in \mathcal{X}} \sum_{t=1}^T f_t(x)$ is the benchmark solution designed with hindsight. The online caching problem fits squarely in this setup, where $f_t(x)$ depends on the users requests and is

unknown when the caching is decided. And previous works [18], [20], [21] have proved that OCO-based policies achieve $R_T = O(\sqrt{T})$, thus ensuring $\lim_{T \rightarrow \infty} R_T/T = 0$.

Different from these studies, we extend the learning model to include predictions that are available through the content recommendations. Improving the regret of learning policies via predictions is a relatively new area in machine learning research. For instance [28] used predictions $\tilde{c}_t$ for the function gradient $c_t = \nabla f_t(x_t)$ with guaranteed quality, i.e., $c_t^\top \tilde{c}_t \geq a\|c_t\|^2$, to reduce R_T from $O(\sqrt{T})$ to $O(\log T)$; and [29] enhanced this result by allowing some predictions to fail the quality condition. A different line of works uses regularizing functions which enable the learner to adapt to the predictions' quality [30], [31]. This approach is more promising for the caching problem where the recommendations might be inaccurate, or followed by the users for only arbitrary time windows.

Our approach relies on the Follow-The-Regularized-Leader (FTRL) algorithm [32] which we extend with predictions that offer *optimism* by reducing the uncertainty about the next-slot functions. We first design a policy (OFTRL) for the bipartite caching model [5], which generalizes the standard single cache case. Theorem 1 proves that R_T is proportional to prediction errors ($\|c_t - \tilde{c}_t\|^2, \forall t$) diminishing to zero for perfect predictions; while still meeting the best achievable bound $O(\sqrt{T})$ [18], [21] even if all predictions fail. We continue with the *elastic* caching problem [33], where the system resizes the caches at each slot based, e.g., on volatile storage leasing costs [33]–[35]. The aim is to maximize the performance subject to a long-term budget constraint. This places the problem in the realm of constrained-OCO [36]–[39]. Using a new saddle point analysis with predictions, we devise Theorem 2 which reveals how $R_T^{(e)}$ and the budget violation $V_T^{(e)}$ depend on the caches and prediction errors, and how we can prioritize one metric over the other while achieving sublinear growth rates for both.

The above algorithms make no assumption about the predictions accuracy, which might be high or low, or even alternate between these extremes (e.g., as user behavior changes) in any unpredictable and frequent fashion. However, in many cases, a rec-sys exhibits *consistent* performance, namely its recommendations are of similar quality within a certain time window; either accurately due to recently trained model, or poorly due to e.g., distributional shift, see [40] and references therein. Our final contribution is a meta-learning caching framework that utilizes such consistent behavior in order to achieve *negative* regret while maintaining sublinear regret when the consistency fails, see Theorem 3.

In summary, the contributions of this work are the following:

- Introduces an online learning framework for bipartite and elastic caching networks that leverages predictions to achieve a constant *zero* regret for perfect recommendations and a sub-linear $O(\sqrt{T})$ regret for arbitrary bad recommendations.
- Introduces a meta-learning framework that can achieve *negative* regret by leveraging consistently-performing rec-sys.
- Evaluates the policies using various request models and real datasets [41] and compares them with key benchmarks.

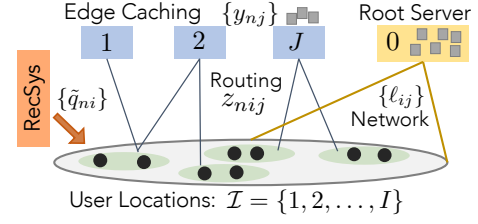


Fig. 1. **System Model.** A network of $\mathcal{J}$ caches serves file requests from a set $\mathcal{I}$ of user locations. Unserved requests are routed to the Root Server. Caching decisions are aided via the recommendations provided by the rec-sys.

The work presents conceptual innovations, i.e., using recommendations as predictions for caching, and using different online caching algorithms in a meta-learning algorithm; as well as technical contributions such as the new optimistic FTRL algorithm with budget constraints (Theorem 2). While we focus on data caching, the proposed algorithms can be directly applied to caching of services on edge systems.

II. SYSTEM MODEL AND PROBLEM STATEMENT

A. Model Preliminaries

Network. The caching network includes a set of edge caches $\mathcal{J} = \{1, 2, \dots, J\}$ and a root cache indexed with 0, Fig. 1. The file requests emanate from a set of non-overlapping user locations $\mathcal{I} = \{1, 2, \dots, I\}$. The connectivity between $\mathcal{I}$ and $\mathcal{J}$ is modeled with parameters $\ell = (\ell_{ij} \in \{0, 1\} : i \in \mathcal{I}, j \in \mathcal{J})$, where $\ell_{ij} = 1$ if cache j can be reached from location i . The root cache is within the range of all users in $\mathcal{I}$. This is a general non-capacitated bipartite model [42] that encompasses as a special case the celebrated femtocaching model [5], and can be used both for wired and wireless networks.

Requests. The system operation is time slotted, $t = 1, 2, \dots, T$. Users submit requests for obtaining files from a library $\mathcal{N}$ of N files with unit size; we note that the analysis can be readily extended to files with different sizes. Parameter $q_{ni}^t \in \{0, 1\}$ indicates the submission of a request for file $n \in \mathcal{N}$ by a user at location $i \in \mathcal{I}$ in the beginning of slot t . At each slot we assume there is one request; i.e., the caching decisions are updated after every request, as in LFU and LRU policies, [43], [44]. Hence, the request process comprises successive vectors $q_t = (q_{ni}^t \in \{0, 1\} : n \in \mathcal{N}, i \in \mathcal{I})$ from the set:

$$\mathcal{Q} = \left\{ q \in \{0, 1\}^{N \cdot I} \mid \sum_{n \in \mathcal{N}} \sum_{i \in \mathcal{I}} q_{ni} = 1 \right\}.$$

We make no assumptions for the request pattern; it might follow a fixed or time-varying distribution that is unknown to the system; and can be even selected strategically by an *adversary* aiming to degrade the caching operation. If a policy performs satisfactory under this model, it is ensured to achieve (at least) the same performance for other request models.

Recommendations. There is a recommender system (*rec-sys*) that suggests up to K_i files to each user $i \in \mathcal{I}$, see [22] for the case of Netflix. User i requests a recommended file with a certain probability that captures the user's propensity to follow one of the recommendations. Unlike prior works that consider these probabilities fixed [24], [45], we model them as unknown

and possibly time-varying. A key point in our approach is that the content recommendations, if properly leveraged, can serve as predictions for the next-slot user requests which are otherwise unknown. We denote with $\tilde{q}_t$ the prediction for the request q_t that the system will receive at the beginning of slot t , and we assume that $\tilde{q}_t$ is available at the end of slot $t-1$, i.e., when the rec-sys provides its recommendations.

Caching. Each cache $j \in \mathcal{J}$ stores up to $C_j < N$ files, while the root cache stores the entire library, i.e., $C_0 \geq N$. We also define $C = \max_{j \in \mathcal{J}} C_j$. Following the standard femtocaching model [5], we perform caching using the *Maximum Distance Separable* (MDS) codes, where files are split into a fixed number of F chunks, which include redundancy chunks. A user can decode the file if it receives any F -sized subset of its chunks. For large values of F , the MDS model allows us to use continuous caching variables.¹ Hence, we define the variable $y_{nj}^t \in [0, 1]$ which denotes the portion of F chunks of file $n \in \mathcal{N}$ stored at cache $j \in \mathcal{J}$, and we introduce the t -slot caching vector $y_t = (y_{nj}^t : n \in \mathcal{N}, j \in \mathcal{J})$ that belongs to set:

$$\mathcal{Y} = \left\{ y \in [0, 1]^{N \cdot J} \mid \sum_{n \in \mathcal{N}} y_{nj} \leq C_j, j \in \mathcal{J} \right\}.$$

Routing. Since each user location $i \in \mathcal{I}$ may be connected to multiple caches, we need to introduce routing variables. Let z_{nij}^t denote the portion of request q_{ni}^t served by cache j . In the MDS caching model the requests can be simultaneously routed from multiple caches and, naturally, we restrict² the amount of chunks not to exceed F . Hence, the t -slot routing vector $z_t = (z_{nij}^t \in [0, 1] : n \in \mathcal{N}, i \in \mathcal{I}, j \in \mathcal{J})$ is drawn from:

$$\mathcal{Z} = \left\{ z \in [0, 1]^{N \cdot J \cdot I} \mid \sum_{j \in \mathcal{J}} z_{nij} \leq 1, n \in \mathcal{N}, i \in \mathcal{I} \right\}.$$

Requests that are not (fully) served by the edge caches $\mathcal{J}$ are served by the root server that provides the missing chunks. This decision needs not to be explicitly modeled as it is directly determined by the routing vector z_t .

B. Problem Statement

Cache Utility & Predictions. We use parameters $w_{nij} \in [0, w]$ to model the system utility when delivering a chunk of file $n \in \mathcal{N}$ to location $i \in \mathcal{I}$ from cache $j \in \mathcal{J}$, instead of using the root server. This utility model can be used to capture bandwidth or delay savings, and other edge-caching gains in wired or wireless networks. The caching benefits can in general differ for each cache and user location, and may vary with time as it is explained in the sequel. Note that the cache-hit maximization problem is a special case of this more general setting [1]. To streamline presentation we introduce vector $x_t = (y_t, z_t) \in \mathbb{R}^m$, with $m = NIJ + NJ$, and define the system utility in slot t as:

$$f_t(x_t) = \sum_{n \in \mathcal{N}} \sum_{i \in \mathcal{I}} \sum_{j \in \mathcal{J}} w_{nij} q_{ni}^t z_{nij}^t, \quad (1)$$

¹Large files are composed of thousands of chunks and hence the rounding operation induces practically negligible errors [42]

²This practical constraint is called the *inelastic* model and compounds the problem, cf. [21] for a comparison with the simpler elastic model.

and we denote its gradient $c_{t+1} = \nabla f_{t+1}(x_{t+1})$. As it will become clear, our analysis holds also for non-linear concave functions $f_t(x)$; this generalization is useful in case, e.g., we wish to enforce fairness in the dispersion of caching gains across the user locations [35].

The main challenge in online caching is the following: at the end of each slot t where we need to decide the cache configuration, the utility function f_{t+1} is not available. Indeed, this function depends on the next-slot request q_{t+1} that is revealed only after y_{t+1} is fixed³, see [17], [18], [21]. Besides, this is also the timing of the LRU/LFU policies [43], [44]. However, the recommendations provided to users can be used to form a predicted request $\tilde{q}_{t+1}$. For example, the caching system can set $\tilde{q}_{\hat{n}\hat{i}}^{t+1} = 1$ and $\tilde{q}_{ni}^{t+1} = 0, \forall (n, i) \neq (\hat{n}, \hat{i})$, where $(\hat{n}, \hat{i})$ is the request with the highest predicted probability⁴. Then, we can use $\tilde{q}_{t+1}$ to create a prediction for the next slot function $\tilde{f}_{t+1}$, or for its gradient $\tilde{c}_{t+1}$, which suffices to solve the caching problem, as we will see.

Benchmark. In such learning problems, it is important to understand the learning objective that our learning algorithm aims to achieve. If we had access to an oracle for the requests $\{q_t\}_{t=1}^T$ (and utility parameters) we could devise the utility-maximizing static caching and routing policy $x^* = (y^*, z^*)$, by solving the following convex optimization problem:

$$\mathbb{P}_1 : \max_x \sum_{t=1}^T f_t(x) \quad (2)$$

$$\text{s.t. } z_{nij} \leq y_{nj} \ell_{ij}, \quad i \in \mathcal{I}, j \in \mathcal{J}, n \in \mathcal{N}, \quad (3)$$

$$z \in \mathcal{Z}, \quad y \in \mathcal{Y}, \quad (4)$$

where constraints (3) ensure the routing decisions for each requested file use only the edge caches that store enough chunks of that file. And let us define the set of constraints $\mathcal{X} = \{\mathcal{Y} \times \mathcal{Z} \cap \{(3)\}\}$, which is compact and convex.

This hypothetical solution x^* can be designed only with *hindsight* and is the benchmark for evaluating our online learning policy π . To that end, we use the metric of regret:

$$R_T(\pi) = \sup_{\{f_t\}_{t=1}^T} \left[\sum_{t=1}^T f_t(x^*) - \sum_{t=1}^T f_t(x_t) \right], \quad (5)$$

which quantifies the performance gap of π from x^* , for any possible sequence of requests or, equivalently, functions $\{f_t\}_t$. Our goal is to find a policy that achieves sublinear regret, $R_T(\pi) = o(T)$, thus ensuring the average performance gap will diminish as T grows. This policy, similar to other online policies, decides x_{t+1} at the end of each slot t using the previous utility functions $\{f_\tau\}_{\tau=1}^t$ and the next-slot prediction $\tilde{f}_{t+1}$ which is made available through the rec-sys.

III. OPTIMISTIC BIPARTITE CACHING

Unlike recent caching solutions that rely on Online Gradient Descent (OGD) [18] or on the Follow-the-Perturbed-Leader

³In our case, since the routing is directly shaped by the caching, this restriction affects also z_{t+1} .

⁴Note that our caching policy is orthogonal to the mechanism that maps the recommendations to predictions.

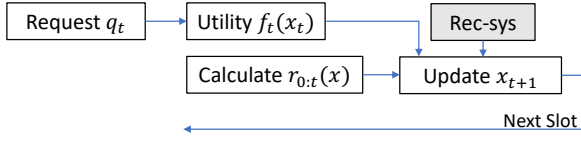


Fig. 2. A decision step for OBC. When a request q_t arrives, the file is routed based on the current cache configuration. The caches are updated using the observed utility $f_t(x_t)$ and the new prediction from the recommender.

(FTPL) policy [21], our approach draws from the *Follow-The-Regularized-Leader* (FTRL) policy, cf. [46]. A key element in our proposal is the *optimism* emanating from the availability of predictions, namely the content recommendations that are offered to users by the rec-sys in each slot.

Let us begin by defining the proximal regularizers⁵:

$$r_0(x) = \mathbf{I}_{\mathcal{X}}(x), \quad r_t(x) = \frac{\sigma_t}{2} \|x - x_t\|^2, \quad t \geq 1 \quad (6)$$

where $\|\cdot\|$ is the Euclidean norm, and $\mathbf{I}_{\mathcal{X}}(x) = 0$ if $x \in \mathcal{X}$ and ∞ otherwise. We apply the following regularizing parameters:

$$\sigma_t = \sigma \left(\sqrt{h_{1:t}} - \sqrt{h_{1:t-1}} \right), \quad \sigma_1 = \sigma \sqrt{h_1}, \quad h_t = \|c_t - \tilde{c}_t\|^2 \quad (7)$$

where $\sigma \geq 0$, $c_t = \nabla f_t(x_t)$, and we used the shorthand notation $h_{1:t} = \sum_{i=1}^t h_i$ for the aggregate prediction errors during the first t slots. The basic step of the algorithm is the update:

$$x_{t+1} = \arg \min_{x \in \mathbf{R}^m} \left\{ r_{0:t}(x) - (c_{1:t} + \tilde{c}_{t+1})^\top x \right\}, \quad (8)$$

which calculates the decision vector based on past observations $c_{1:t}$, the aggregate regularizer $r_{0:t}(x)$ and the prediction $\tilde{c}_{t+1}$ (see Fig. 2). The update employs the negative gradients as it concerns a maximization problem, cf. [46]. Henceforth, we refer to (8) as the *optimistic* FTRL (OFTRL) update.

Policy π_{obc} is outlined in Algorithm OBC. In each iteration, OBC solves a convex optimization problem, (8), involving a projection on the feasible set $\mathcal{X}$ (via $r_0(x)$). For the latter, one can rely on fast-projection algorithms specialized for caching, e.g., see [18]; while it is possible to obtain a closed-form solution for the OFTRL update for linear functions. We quantify next the performance of Algorithm OBC.

Theorem 1. *Algorithm OBC ensures the regret bound:*

$$R_T \leq 2\sqrt{2(1+JC)} \sqrt{\sum_{t=1}^T \|c_t - \tilde{c}_t\|^2}$$

Proof. We start from [31, Theorem 1] which proved that a proximal OFTRL update with regularizer $r_{0:t}(x)$ that is 1-strongly-convex w.r.t. some norm $\|\cdot\|_{(t)}$ yields regret:

$$R_T \leq r_{1:T}(x^*) + \sum_{t=1}^T \|c_t - \tilde{c}_t\|_{(t),*}^2, \quad \forall x^* \in \mathcal{X}. \quad (9)$$

Now, $r_{1:t}$ is 1-strongly-convex w.r.t. norm $\|x\|_{(t)} = \sqrt{\sigma_{1:t}} \|x\|$ which has dual norm $\|x\|_{(t),*} = \|x\|/\sqrt{\sigma_{1:t}}$. Using the

⁵A proximal regularizer is one that induces a proximal mapping for the objective function; see [47, Ch. 6.1] for the formal definition.

Algorithm OBC: Optimistic Bipartite Caching (π_{obc})

```

1 Input:  $\{\ell_{ij}\}_{(i,j)}; \{C_j\}_j; \mathcal{N}; x_1 \in \mathcal{X}; \sigma = 2/D_{\mathcal{X}}$ .
2 Output:  $x_t = (y_t, z_t), \forall t$ .
3 for  $t = 1, 2, \dots$  do
4   Route request  $q_t$  according to configuration  $x_t$ 
5   Observe system utility  $f_t(x_t)$ 
6   Observe the new prediction  $\tilde{c}_{t+1}$ 
7   Update the regularizer  $r_{0:t}(x)$  using (6)-(7)
8   Calculate the new policy  $x_{t+1}$  using (8)
end

```

regularization parameter (7), we get $\sigma_{1:t} = \sigma \sqrt{h_{1:t}}$, and replacing all the above into (9) we get:

$$R_T \leq \frac{\sigma}{2} \sqrt{h_{1:T}} D_{\mathcal{X}}^2 + \sum_{t=1}^T \frac{h_t}{\sigma \sqrt{h_{1:t}}} \quad (10)$$

where we upper bounded each $\|x - x_t\|$ term with the Euclidean diameter of $\mathcal{X}$ denoted by $D_{\mathcal{X}}$. Namely $\forall x, x_t \in \mathcal{X}$ holds:

$$\begin{aligned} \|x - x_t\|^2 &= \sum_{n,j} (y_{nj} - y_{nj}^t)^2 + \sum_{n,i,j} (z_{nij} - z_{nij}^t)^2 \\ &\stackrel{(a)}{\leq} \sum_{n,j} |y_{nj} - y_{nj}^t| + \sum_{n,i,j} |z_{nij} - z_{nij}^t| \stackrel{(b)}{\leq} 2(JC + 1) \triangleq D_{\mathcal{X}}^2 \end{aligned}$$

where (a) holds as $y_{nj}, z_{nij} \in [0, 1]$; (b) holds by the triangle inequality and definitions of $\mathcal{Y}, \mathcal{Z}, \mathcal{Q}$. Finally, using [48, Lem. 3.5] to bound the second regret term in (10) $\sum_{t=1}^T h_t / \sqrt{h_{1:t}} \leq 2\sqrt{h_{1:T}}$ and setting $\sigma = 2/D_{\mathcal{X}}$ we arrive at the result. $\square$

Discussion. Theorem (1) shows that the regret does not depend on the library size N and is also modulated by the quality of the content recommendations; accurate predictions tighten the bound, and in the case of perfect prediction, i.e., when users follow the recommendations, we get a negative regret $R_T \leq 0, \forall T$, which is much stronger than the sub-linear growth rates in other works [18], [49]. On the other hand, for worst-case prediction, it is $\|c_t - \tilde{c}_t\|^2 \leq 2w^2$, thus $R_T \leq 4w\sqrt{2(JC+1)}\sqrt{T} = O(\sqrt{T})$; i.e., the regret is at most a constant factor worse than the regret of those policies that do not incorporate predictions⁶, regardless of the predictions' quality. Thus, OBC offers an efficient and safe approach for incorporating predictions in cases where we are uncertain about their accuracy, e.g., either due to the quality of the rec-sys or the behaviour of users.

Another key point is that the *utility parameters might vary with time* as well. Indeed, replacing $w_t = (w_{nij}^t \leq w, n \in \mathcal{N}, i \in \mathcal{I}, j \in \mathcal{J})$ in $f_t(x_t)$ does not affect the analysis nor the bound. This is important when the caching system employs a wireless network where the link capacities vary, or when the caching utility changes. Similarly, for edge computing and caching services, the utility of each computation or service might vary substantially across users and time. Parameters w_t can be even unknown to the caching system when x_t is decided, exactly as it is with q_t , and they can be predicted either using the rec-sys or other side information (e.g., channel measurements).

⁶The factor is 2 compared to the “any-time” version of the bound that do not use predictions, and $2\sqrt{2}$ compared to those that assume a known T .

IV. OPTIMISTIC CACHING IN ELASTIC NETWORKS

We extend our analysis to *elastic* caching networks where the caches can be resized dynamically. Such architectures are important for two reasons. Firstly, there is a growing number of small-size content providers that implement their services by leasing storage on demand from infrastructure providers [50]; and secondly, CDNs often resize their caches responding to the time-varying user needs and operating expenditures [51].

We introduce the t -slot price vector $s_t = (s_j^t \leq s, j \in \mathcal{J})$, where s_j^t is the leasing price per unit of storage at cache j in slot t , and s its maximum value. In the general case, these prices may change arbitrarily over time, e.g., because the provider has a dynamic pricing scheme or the electricity cost changes [34], [35]; hence the caching system has access only to s_t at each slot t . We denote with B_T the budget the system intends to spend during a period of T slots for leasing cache capacity. The objective is to maximize the caching gains while satisfying the constraint:

$$\sum_{t=1}^T g_t(x_t) = \sum_{t=1}^T \sum_{j \in \mathcal{J}} \sum_{n \in \mathcal{N}} s_j^t y_{nj}^t - B_T \leq 0. \quad (11)$$

In particular, the new benchmark problem in this case is:

$$\mathbb{P}_2: \max_{x \in \mathcal{X}} \sum_{t=1}^T f_t(x) \quad \text{s.t.} \quad (3), \quad \sum_{t=1}^T g_t(x) \leq 0, \quad (12)$$

which differs from $\mathbb{P}_1$ due to the leasing constraint.

Indeed, in this case the regret is defined as:

$$R_T^{(e)}(\pi) = \sup_{\{f_t, g_t\}_{t=1}^T} \left[\sum_{t=1}^T f_t(x^*) - \sum_{t=1}^T f_t(x_t) \right], \quad (13)$$

where $x^* \in \mathcal{X}_e \triangleq \{x \in \mathcal{X} \mid (3), g_t(x) \leq 0, \forall t\}$, i.e., x^* is a feasible point of $\mathbb{P}_2$ with the additional restriction to satisfy $g_t(x) \leq 0$ in every slot. In the definition of $\mathcal{X}$, C now denotes the maximum leasable space. Learning problems with time-varying constraints are notoriously hard to tackle, see impossibility result in [52], and hence require such additional restrictions on the selected benchmarks. We refer the reader to [36] for a related discussion, and to [37], [38] for more competitive benchmarks. These ideas are directly applicable to our OFTRL framework. For instance, the analysis follows directly for the K -slot benchmark of [37] where $\sum_{\tau=t}^{t+K} g_\tau(x^*) \leq 0, \forall t$, instead of $g_t(x^*) \leq 0, \forall t$. Finally, apart from $R_T^{(e)}$, we need also to ensure sublinear growth rate for the budget violation:

$$V_T^{(e)} = \sum_{t=1}^T g_t(x_t).$$

To tackle this new problem we follow a saddle point analysis, which is new in the context of OFTRL.

We first define a Lagrangian-type function by relaxing the budget constraint and introducing the dual variable $\lambda \geq 0$:

$$\mathcal{L}_t(x, \lambda) = \frac{\sigma_t}{2} \|x - x_t\|^2 - f_t(x_t) + \lambda g_t(x_t) - \frac{\lambda^2}{a_t}. \quad (14)$$

Algorithm OEC: Optimistic Elastic Caching (π_{oec})

1 **Input:** $\{\ell_{ij}\}_{(i,j)}, \{C_j\}_j, \mathcal{N}, \lambda_1=0, x_1 \in \mathcal{X}_e, a_t = at^{-\beta}$
2 **Output:** $x_t = (y_t, z_t), \forall t$.
3 **for** $t = 1, 2, \dots$ **do**
4 Route request q_t according to configuration x_t
5 Observe system utility $f_t(x_t)$ and cost $g_t(x_t)$
6 Update the budget parameter λ_{t+1} using (15)
7 Update the regularizer $r_{0:t}(x)$ using (6)-(7)
8 Observe prediction $\tilde{c}_{t+1}$ and price s_{t+1}
9 Calculate the new policy x_{t+1} using (16)
end

The last term is a non-proximal regularizer for the dual variable; and we use $a_t = at^{-\beta}$, where parameter $\beta \in [0, 1)$ can be used to prioritize either $R_T^{(e)}$ or $V_T^{(e)}$. The main ingredients of policy π_{oec} are the saddle-point iterations:

$$\lambda_{t+1} = \arg \max_{\lambda \geq 0} \left\{ -\frac{\lambda^2}{a_{t+1}} + \lambda \sum_{i=1}^t g_i(x_i) \right\}, \quad (15)$$

$$x_{t+1} = \arg \min_{x \in \mathbb{R}^m} \left\{ r_{0:t}(x) + \left(\sum_{i=1}^{t+1} \lambda_i s_i - c_{1:t} - \tilde{c}_{t+1} \right)^\top x \right\} \quad (16)$$

and its implementation is outlined in Algorithm OEC. Note that we use the same regularizer for the primal variables x_t , while λ_t modulates the caching decisions by serving as a *shadow price* for the average budget expenditure.

The performance of Algorithm OEC is characterized next.

Theorem 2. *Algorithm OEC ensures the bounds:*

$$R_T^{(e)} \leq 2D\mathcal{X} \sqrt{\sum_{t=1}^T \|c_t - \tilde{c}_t\|^2} + \frac{a(sJC)^2}{2(1-\beta)} T^{1-\beta}$$

$$V_T^{(e)} \leq \sqrt{\frac{4D\mathcal{X}T^\beta}{a}} \sqrt{\sum_{t=1}^T \|c_t - \tilde{c}_t\|^2} + \frac{T(sJC)^2}{1-\beta} - \frac{2R_T^{(e)}T^\beta}{a}$$

Proof. Observe that the update in (16) is similar to (8) but applied to the Lagrangian in (14) instead of just the utility, and the known prices when x_{t+1} is decided represent perfect prediction for $g_t(x)$. Using Theorem 1 with $c_t - \lambda_t s_t$ instead of c_t , and $\tilde{c}_t - \lambda_t s_t$ instead of $\tilde{c}_t$, we can write:

$$\sum_{t=1}^T \left(f_t(x^*) - f_t(x_t) + \lambda_t g_t(x_t) - \lambda_t g_t(x^*) \right) \leq 2D\mathcal{X} \sqrt{h_{1:T}},$$

and rearrange to obtain:

$$R_T^{(e)} \leq 2D\mathcal{X} \sqrt{h_{1:T}} + \sum_{t=1}^T \lambda_t g_t(x^*) - \sum_{t=1}^T \lambda_t g_t(x_t). \quad (17)$$

For the dual update (15), we can use the non-proximal-FTRL bound [46, Theorem 1] to write:

$$-\sum_{t=1}^T \lambda_t g_t(x_t) + \lambda \sum_{t=1}^T g_t(x_t) \leq \frac{\lambda^2}{a_T} + \frac{1}{2} \sum_{t=1}^T a_t g_t^2(x_t). \quad (18)$$

Since $g_t(x^*) \leq 0, \forall t$ and combining (17), (18) we get:

$$R_T^{(e)} \leq 2D_{\mathcal{X}} \sqrt{h_{1:T}} - \lambda \sum_{t=1}^T g_t(x_t) + \frac{\lambda^2}{a_T} + \frac{1}{2} \sum_{t=1}^T a_t g_t^2(x_t). \quad (19)$$

Setting $\lambda=0$, using the identity $\sum_{t=1}^T a t^{-\beta} \leq a T^{1-\beta}/(1-\beta)$ and the bound $g_t(x_t) \leq sJC$, we prove the $R_T^{(e)}$ bound. Using:

$$\frac{a_T}{2} \left[\sum_{t=1}^T g_t(x_t) \right]^2 = \sup_{\lambda \geq 0} \left[\sum_{t=1}^T g_t(x_t) \lambda - \frac{\lambda^2}{2a_T} \right],$$

we can replace this term to (19) and write:

$$\frac{a_T}{2} (V_T^{(e)})^2 \leq 2D_{\mathcal{X}} \sqrt{h_{1:T}} + \frac{a(sJC)^2}{2-2\beta} T^{1-\beta} - R_T^{(e)}.$$

Rearranging and taking the square root yields $V_T^{(e)}$ bound. $\square$

Discussion. The worst-case bounds in Theorem 2 arise when the predictions are failing. In that case, we have $\|c_t - \tilde{c}_t\|^2 \leq 2w^2$ and use the bound $-R_T^{(e)} = O(T)$ for the last term of $V_T^{(e)}$, to obtain $R_T^{(e)} = O(T^\kappa)$, with $\kappa = \max\{1/2, 1-\beta\}$ while $V_T^{(e)} = O(T^\phi)$, with $\phi = \frac{1+\beta}{2}$. Hence, for $\beta=1/2$ we achieve the desired sublinear rates $R_T^{(e)} = O(\sqrt{T})$, $V_T^{(e)} = O(T^{3/4})$. However, when the rec-sys manages to predict accurately the user preferences, the performance of π_{oec} improves substantially as the first terms in each bound are eliminated. Thus, for bounded T , we practically halve the regret and violation bounds.

It is also interesting to observe the tension between $V_T^{(e)}$ and $R_T^{(e)}$, which is evident from the $V_T^{(e)}$ bound and the condition $-R_T^{(e)} = O(T)$. The latter refers to the upper bound of the *negative* regret, thus when it is consistently satisfied (i.e., for all T), we obtain an even better result: π_{oec} *outperforms* the benchmark. Another likely case is when $-R_T^{(e)} = O(\sqrt{T})$, i.e., the policy does not outperform the benchmark at a rate larger than $\sqrt{T}$. Then, Theorem 2 yields $R_T^{(e)} = O(T^\kappa)$ with $\kappa = \max\{1/2, 1-\beta\}$ while $V_T^{(e)} = O(T^\phi)$ with $\phi = \max\{1/2, 1/4 + \beta/2\}$. Hence, for $\beta=1/2$ the rates are reduced to $R_T^{(e)} = O(\sqrt{T})$, $V_T^{(e)} = O(\sqrt{T})$.

V. CACHING WITH NON-VOLATILE PREDICTIONS

We now introduce a different approach on modeling recommendations as predictions, which, in cases of consistent prediction performance, delivers better regret. Namely, we model the problem of online caching using the *experts model*, see [27]. The first expert represents a robust learner (referred to as *pessimistic*) and proposes an FTRL-based caching policy without any predictions. The second expert represents an *optimistic* learner and implements a policy that always caches the file predicted to be requested. To streamline the presentation, we present the results using a single cache scenario (hence using only y below), but it will become clear that this method can be readily extended to caching networks.

Formally, the pessimistic expert proposes caching actions $\{y_t^{(p)}\}_t$ according to step (8), but with setting $\tilde{c}_t=0$ for the

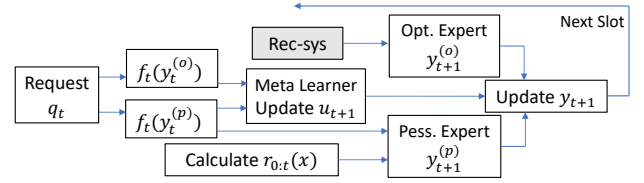


Fig. 3. A decision step for XC. Experts' utilities are used to update the weights u . The new caching decisions are then the combination of the experts' proposals. The optimistic decisions are updated based only on predictions from rec-sys. The pessimistic decisions are updated based only on past requests.

regularization parameter σ_t in (7). Its regret w.r.t the optimal-in-hindsight caching configuration $y^* = \arg\max_{y \in \mathcal{Y}} c_{1:T}^\top y$ is denoted with $R_T^{(p)}$. On the other hand, the optimistic expert proposes actions $\{y_t^{(o)}\}_t$ according to the linear program:

$$y_{t+1}^{(o)} = \arg\max_{y \in \mathcal{Y}} \tilde{c}_{t+1}^\top y, \quad (20)$$

and we denote its regret with $R_T^{(o)}$. The optimistic expert represents a high-risk high-reward policy; $R_T^{(o)}$ is linear in the worst case predictions and negative linear for perfect predictions. In contrast, the pessimistic expert is more robust as it is not affected by bad predictions, but guarantees only a sub-linear regret. We aim to have the best of both worlds and design an algorithm that, in the best case, is able to obtain negative regret, while being worse only by a constant factor than the pessimistic expert in the general case.

Unlike π_{obc} and π_{ec} , the predictions are not appended to the FTRL step itself but rather treated independently through the optimistic expert. The challenge is to meta-learn which of the two experts to rely upon. To that end, we will be using Online Gradient Ascent (OGA) to learn how to combine the experts' proposed caching vectors $y_t^{(p)}$ and $y_t^{(o)}$. The decisions of the meta-learner are then these combination weights $u_t = (u_t^{(p)}, u_t^{(o)})$, drawn from the 2-dimensional simplex set Δ , (see Fig. 3). The weights are learned through the OGA step:

$$u_{t+1} = \mathcal{P}_\Delta \{u_t + \delta_t l_t\}, \quad (21)$$

where $\mathcal{P}$ is the projection operator, δ_t is the OGA learning rate and $l_t = (l_t^{(p)}, l_t^{(o)})$ is the t -slot performance vector for the experts, i.e., $l_t^{(p)} = c_t^\top y_t^{(p)}$, and $l_t^{(o)} = c_t^\top y_t^{(o)}$. The caching decision is the convex combination of experts' proposals:

$$y_{t+1} = u_{t+1}^{(p)} y_{t+1}^{(p)} + u_{t+1}^{(o)} y_{t+1}^{(o)}. \quad (22)$$

Thus, y_{t+1} is still a feasible caching policy. The steps are shown in Algorithm XC, and the following theorem bounds the regret of the *caching decisions* $\{y_t\}_t$.

Theorem 3. Algorithm XC ensures the bound $R_T^{(xc)} =$

$$\sum_{t=1}^T c_t^\top (y^* - y_t) \leq 2w\sqrt{2T} + A, \quad A \in [-wT, 2w\sqrt{2CT}]$$

Proof. First, we relate the regret of the combined caching decisions to that of the expert selection,

Algorithm XC: Experts Caching (π_{xc})

```

1 Input:  $C$ ;  $y_1 \in \mathcal{Y}$ ;  $\sigma = 2/D_{\mathcal{Y}}$ .
2 Output:  $y_t, \forall t$ .
3 for  $t = 1, 2, \dots$  do
4   Serve request  $q_t$  according to configuration  $y_t$ 
5   Observe utilities  $f_t(y_t^{(p)}), f_t(y_t^{(o)})$ 
6   Update  $r_{0:t}(x)$  using (6)-(7) with  $\tilde{c}_{t+1} = 0$ 
7   Calculate pessimistic expert's proposal  $y_{t+1}^{(p)}$  as in (8)
8   Observe the prediction  $\tilde{q}_{t+1}$ , and calculate  $\tilde{c}_{t+1}$ 
9   Calculate optimistic expert's proposal  $y_{t+1}^{(o)}$  using (20)
10  Calculate the new weights  $u_{t+1}$  using (21)
11  Calculate the new policy  $y_{t+1}$  using (22)
end

```

$$\begin{aligned}
R_T^{(xc)} &= \sum_{t=1}^T c_t^\top y^* - c_t^\top (u_t^{(p)} y_t^{(p)} + u_t^{(o)} y_t^{(o)}) = \sum_{t=1}^T c_t^\top y^* - l_t^\top u_t \\
&= \sum_{t=1}^T c_t^\top y^* - l_t^\top u^* + l_t^\top u^* - l_t^\top u_t \\
&= R_T^{(u)} + \min \left\{ R_T^{(p)}, R_T^{(o)} \right\}, \tag{23}
\end{aligned}$$

where $R_T^{(u)}$ is the regret for the *expert selection weights* u : $R_T^{(u)} = \sum_{t=1}^T l_t^\top u^* - l_t^\top u_t$. (23) holds because $u^* = \arg\max_{u \in \Delta} l_{1:t}^\top u = \max\{l_t^{(o)}, l_t^{(p)}\}$. Thus, we have that

$$l_{1:t}^\top u^* = \max \left\{ \sum_{t=1}^T c_t^\top y_t^{(p)}, \sum_{t=1}^T c_t^\top y_t^{(o)} \right\}. \tag{24}$$

Now, we write the expressions for the terms in (23). $R_T^{(p)}$ can be bounded in the same manner as Theorem 1 with prediction vectors $\tilde{c}_t = 0$, and substituting an upper bound w for $\|c_t\|$:

$$R_T^{(p)} \leq 2w\mathcal{D}_{\mathcal{Y}}\sqrt{T} \leq 2w\sqrt{2CT}. \tag{25}$$

$R_T^{(o)}$ is hard to calculate as it depends on both, prediction $\{\tilde{c}_t\}_t$, and the relationship between $c_{1:t}$ and c_t . However, we can easily deduce lower and upper bounds. Since c_t and $\tilde{c}_t$ represent the utility of one request to a file, each term of the optimistic regret can be maximally w . Hence, we have that $R_T^{(o)} \in [-wT, wT]$, and:

$$\min \left\{ R_T^{(p)}, R_T^{(o)} \right\} \in [-wT, 2w\sqrt{2CT}]. \tag{26}$$

For $R_T^{(u)}$, we use the OGA bound [53, Thm. 4.14] tailored to the simplex decision set. Using $\eta_t^{(u)} = \frac{1}{w\sqrt{t}}$, and $\|l_t\| \leq \sqrt{2}w$:

$$R_T^{(u)} \leq 2w\sqrt{2T} \tag{27}$$

Substituting (26) and (27) in (23) gives the bound. $\square$

Discussion. The regret in Theorem 3 can now be *strictly* negative for perfect predictions, which is tighter than π_{obc} . In general, however, the regret bound can be worse than that of π_{obc} . Namely, $R_T^{(xc)}$ is bounded by $2w\sqrt{2T} + 2\sqrt{2}w\sqrt{CT}$, while in the single cache case R_T is bounded by $2\sqrt{2C}\sqrt{\sum_t \|c_t - \tilde{c}_t\|^2}$, which might be better or worse than $R_T^{(xc)}$, depending on the number of steps with accurate

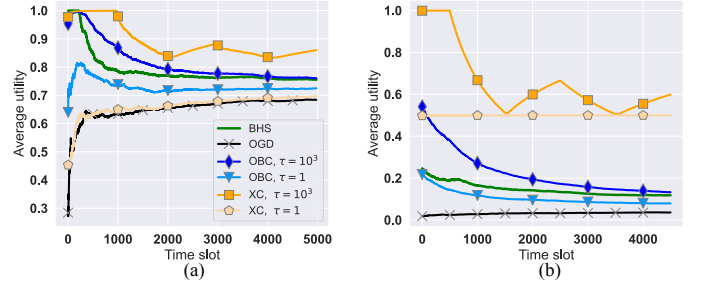


Fig. 4. Utility in the single cache model under different prediction quality levels in (a) Zipf requests with $\zeta = 1.2$, (b) YouTube request traces [41].

predictions. In all cases, $R_T^{(xc)} = O(\sqrt{T})$. An additional point to highlight is that the first term, $2w\sqrt{2T}$, is a worst-case bound for finding the best expert, i.e., $R_T^{(u)}$. In request sequences where the best expert is easily identifiable, e.g., due to consistent predictions which make l_t similar, it would be a loose upper bound and its actual value is constant (negligible) compared to the second term $\min \left\{ R_T^{(p)}, R_T^{(o)} \right\}$. This second term is the regret of the best expert, and falls in a range that depends on predictions' quality at each step. Thus, if the optimistic expert is better than the best-in-hindsight solution, this min term will be negative for some request sequences.

VI. PERFORMANCE EVALUATION

We evaluate π_{obc} , π_{oec} and π_{xc} under different request patterns and predictions modes; and we benchmark them against x^* and the OGD policy [18] that outperforms other state-of-the-art policies [43], [44]. We observe that when reasonable predictions are available, the proposed policies have an advantage, and under noisy predictions, they still reduce the regret at the same rate with OGD, as proven in the Theorems. First, we compare π_{obc} and π_{xc} against OGD [18] in the single cache case. We then study π_{obc} for the bipartite model and π_{oec} with the presence of budget constraints. We consider two requests scenarios, stationary Zipf requests (with parameter $\zeta = 1.2$) and an actual trace from the dataset in [41]. Predictions alternate between accurate and adversarial (i.e., requesting the recommended file vs. any other file, respectively), for τ time step in each mode. While low values of τ represent an unstable performance, the high value of τ is meant to approximate the consistent performance of practical rec-sys. We also experiment with random accuracies where at each t , the prediction is accurate with probability ρ .

Single Cache Scenarios. We set $w = 1$ to study the cache hit rate scenario. Fig. 4.a shows the performance of π_{obc} and π_{xc} for a library of $N = 10^4$ files and cache size $C = 100$. At each slot, we plot the attained average utility, $\frac{1}{T} \sum_{t=1}^T f_t(x_t)$, for each policy and the best static cache configuration *until that slot*, i.e., we find the best in hindsight⁷ for each t .

In the simulated requests case (Fig. 4.a), π_{obc} achieves negative regret through the experiment for $\tau = 10^3$ and a regret that is 57.1% better than that of the OGD for $\tau = 1$. Such an advantage for the former is due to having more time steps

⁷unlike [18] that calculates x^* for the largest t , we use x_t^* for each t . Thus, the gap between any policy and BHS at t is the policy's average regret R_t/t .

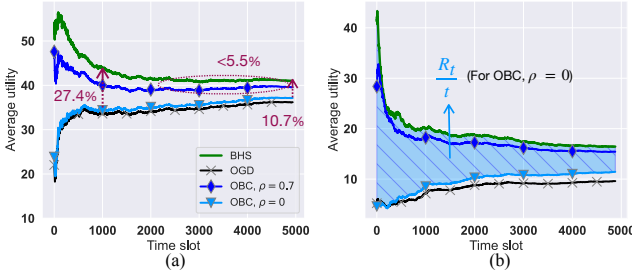


Fig. 5. Attained utility in the bipartite model under different prediction quality levels in (a) Zipf requests with $\zeta = 1.2$, (b) YouTube request traces [41].

with accurate predictions. π_{xc} also maintains negative regret that even outperforms π_{obc} when $\tau = 10^3$. This is because the stable performance of experts allows the policy to efficiently find the best expert and stick with it within each time window. However, a value of $\tau = 1$ induces frequent switching between the two experts in π_{xc} : the performance of the optimistic expert alternate between 0 and 1, while that of pessimistic expert is in the range $(0.6, 0.7)$. Hence, π_{xc} is inclined to place some weight on the opt. expert at one step, only to retract and suffer a greater loss at the following one had it stayed with the full weight on the pess. expert. Due to the additional regret caused by such frequent switching, π_{obc} performs better when $\tau = 1$.

For the trace used in Fig. 4.b, π_{obc} maintains the advantage over OGD in both prediction modes. Regarding π_{xc} , the alternation of the performance of the opt. expert (when $\tau = 1$) no longer induces a switching between the experts since even when the opt. expert performs poorly (gets 0 reward), there is a high chance, especially initially, that the pess. perform similarly⁸. Hence, finding that the opt. expert is better is still easy (due to differences in their utility). Thus, in this trace, π_{xc} performs well with both τ values.

Bipartite Networks. We consider next a bipartite graph with 3 caches and 4 user locations, where the first two locations are connected with caches 1 and 2, and the rest are connected to caches 2 and 3. The utility vector is $w_n = (1, 2, 100)$, $\forall i, j$, thus an efficient policy places popular files on cache 3. This is the setup used in [18] that we adopt here to make a fair comparison. For the stationary scenario, we consider a library of $N = 500$ files and $C = 50$. For the traces scenario, files with at least 10 requests are considered, forming a library of $N = 456$ files, and we keep $C = 50$. In this experiment, we assume that at each time step, the user follows the recommendation with probability ρ . The location of each request is selected uniformly at random. Similar to the single-cache case, we plot the average utility of the online policies and the best static configuration *until each t*.

Scenario 1 in Fig. 5.a shows the effect of good predictions as OBC maintains utility within 5.32% of BHS's utility after $t = 2.5k$. Even when the recommendations are not followed, OBC preserves the sublinear regret, achieving a gap of 27.4% and 10.36% for $t = 1k$ and $t = 5k$, respectively. Akin patterns appear in the second scenario (Fig. 5.b) but with lower utilities

⁸The poor performance of the pess. expert here is due to more uniform, unpredictable, request vectors.

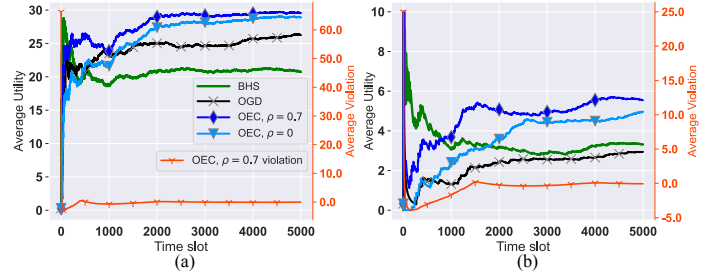


Fig. 6. Attained utility and constraints violations for OEC, OGD with (a) Zipf requests with $\zeta = 1.5$ and (b) YouTube request traces [41].

across all policies due to the more spread requests. Recall that the area between a policy and BHS is the average regret.

Next, we consider the case of budget constraint and evaluate π_{oec} for scenario 1, Fig. 6.a, and scenario 2, Fig. 6.b. The prices at each slot are generated uniformly at random in the normalized range $[0, 1]$, and the available budget is generated randomly $b_t = \mathcal{N}(0.5, 0.05) \times 10$, i.e., enough for approximately 10 files. Such tight budgets magnify the role of dual variables and allow us to test the constraint satisfaction. The benchmark x^* is computed once for the *full time horizon*, and its utility is plotted for each t . In both scenarios, we note that the constraint violation for all policies is approximately similar, fluctuating during the first few slots and then stabilizing at zero. Hence, we plot it for one case.

Concluding, we find that π_{oec} can even outperform the benchmark since it is allowed to violate the constraints at some time slots, provided that the constraints are eventually satisfied, which occurs either due to strict satisfaction or due to having an ample subsidy at some slots. Moreover, in the first scenario (Fig. 6.a), the good predictions enable OEC to outperform x^* by 42.5% after observing all requests ($T = 5K$). OGD, and OEC with noisy predictions attain utility units improvement of 26.5%, 39.3%, respectively, over the BHS. In the second scenario (Fig. 6.b), the good forecast enables a utility gain of 67.1% compared to, -11.3%, and 49.7% for OGD and OEC with noisy prediction, respectively. The algorithms scale for very large libraries $\mathcal{N}$, and the only bottleneck is finding x^* which involves the horizon T , see also [18], [21]; this is not required in real systems. The code for the presented policies and experiments is available at [54].

VII. CONCLUSIONS

The problem of online caching is timely with applications that extend beyond content delivery to edge computing [1]. This work proposes a new suite of caching policies that leverage predictions obtained from content-viewing recommendations to achieve negative regret w.r.t to an ideal (unknown) benchmark. As recommender systems permeate online content platforms, such policies can play an essential role in optimizing the caching efficacy. We identified and built upon this new connection. The framework is scalable and robust to the quality of recommendations, improves previously known caching regret bounds [18], [20], [21], and opens new directions. Among them, the design of optimistic policies for uncoded caching is perhaps the most promising.

VIII. ACKNOWLEDGMENT

This publication has emanated from research conducted with the financial support of the European Commission through Grant No. 101017109 (DAEMON).

REFERENCES

- [1] G. S. Paschos, G. Iosifidis, M. Tao, D. Towsley, and G. Caire, "The Role of Caching in Future Communication Systems and Networks," *IEEE J. Select. Areas Commun.*, vol. 36, no. 6, 2018.
- [2] L. A. Belady, "A study of Replacement Algorithms for a Virtual-Storage Computer," *IBM Systems Journal*, vol. 5, no. 2, 1966.
- [3] C. Aggarwal, J. L. Wolf, and P. S. Yu, "Caching on the World Wide Web," *Trans. Knowledge Data Eng.*, vol. 11, no. 1, 1999.
- [4] J. Kangasharju, J. Roberts, and K. Ross, "Object Replication Strategies in Content Distribution Networks," *Computer Communications*, vol. 25, no. 4, 2002.
- [5] K. Shanmugam *et al.*, "Femtocaching: Wireless Content Delivery Through Distributed Caching Helpers," *IEEE Trans. Inform. Theory*, vol. 59, no. 12, 2013.
- [6] N. Golrezaei, A. F. Molisch, A. G. Dimakis, and G. Caire, "Femto-caching and Device-to-device Collaboration: A New Architecture for Wireless Video Distribution," *IEEE Commun. Mag.*, vol. 51, no. 4, 2013.
- [7] D. D. Sleator and R. E. Tarjan, "Amortized Efficiency of List Update and Paging Rules," *Commun. ACM*, vol. 28, no. 2, 1985.
- [8] P. R. Jelenković and X. Kang, "Characterizing the Miss Sequence of the LRU Cache," *SIGMETRICS Perform. Eval. Rev.*, vol. 36, no. 2, 2008.
- [9] D. Lee *et al.*, "On the Existence of a Spectrum of Policies That Subsumes the Least Recently Used (LRU) and Least Frequently Used (LFU) Policies," *SIGMETRICS Perform. Eval. Rev.*, vol. 27, no. 1, 1999.
- [10] S. Traverso, M. Ahmed, M. Garetto, P. Giaccone, E. Leonardi, and S. Niccolini, "Temporal Locality in Today's Content Caching: Why It Matters and How to Model It," *SIGCOMM Comput. Commun. Rev.*, vol. 43, no. 5, 2013.
- [11] F. Olmos, B. Kauffmann, A. Simonian, and Y. Carlinet, "Catalog dynamics: Impact of content publishing and perishing on the performance of a LRU cache," in *Proc. of ITC*, 2014.
- [12] M. Leconte, G. Paschos, L. Gkatzikis, M. Draief, S. Vassilaras, and S. Chouvardas, "Placing Dynamic Content in Caches with Small Population," in *Proc. of IEEE INFOCOM*, 2016.
- [13] S.-E. Elayoubi and J. Roberts, "Performance and Cost Effectiveness of Caching in Mobile Access Networks," in *Proc. of ICN*, 2015.
- [14] S. O. Somuyiwa, A. György, and D. Gündüz, "A Reinforcement-Learning Approach to Proactive Caching in Wireless Networks," *IEEE J. Select. Areas Commun.*, vol. 36, no. 6, 2018.
- [15] A. Sadeghi, F. Sheikholeslami, and G. B. Giannakis, "Optimal and Scalable Caching for 5G Using Reinforcement Learning of Space-Time Popularities," *IEEE J. Select. Areas Commun.*, vol. 12, no. 1, 2018.
- [16] S. Geulen, B. Vöcking, and M. Winkler, "Regret Minimization for Online Buffering Problems Using the Weighted Majority Algorithm," in *Proc. of COLT*, 2010.
- [17] T. Lykouris and S. Vassilvitskii, "Competitive Caching with Machine Learned Advice," in *Proc. of ICML*, 2018.
- [18] G. S. Paschos, A. Destounis, L. Vigneri, and G. Iosifidis, "Learning to Cache With No Regrets," in *Proc. of IEEE INFOCOM*, 2019.
- [19] Y. Li, T. Si Salem, G. Neglia, and S. Ioannidis, "Online Caching Networks with Adversarial Guarantees," *Proc. ACM Meas. Anal. Comput. Syst.*, vol. 5, no. 3, 2021.
- [20] T. Si Salem, G. Neglia, and S. Ioannidis, "No-Regret Caching via Online Mirror Descent," in *Proc. of ICC*, 2021.
- [21] R. Bhattacharjee, S. Banerjee, and A. Sinha, "Fundamental Limits on the Regret of Online Network-Caching," *Proc. ACM Meas. Anal. Comput. Syst.*, vol. 4, no. 2, 2020.
- [22] C. A. Gomez-Urbe and N. Hunt, "The Netflix Recommender System: Algorithms, Business Value, and Innovation," *ACM Trans. Manage. Inf. Syst.*, vol. 6, no. 4, 2016.
- [23] X. Amatriain, "Building Industrial-Scale Real-World Recommender Systems," in *Proc. of RecSys*, 2012.
- [24] L. E. Chatzileftheriou, M. Karaliopoulos, and I. Koutsopoulos, "Jointly Optimizing Content Caching and Recommendations in Small Cell Networks," *IEEE Trans. Mobile Comput.*, vol. 18, no. 1, 2019.
- [25] Y. Fu, Q. Yu, T. Q. S. Quek, and W. Wen, "Revenue Maximization for Content-Oriented Wireless Caching Networks (CWCNs) With Repair and Recommendation Considerations," *IEEE Trans. Wireless Commun.*, vol. 20, no. 1, 2021.
- [26] M. Zinkevich, "Online Convex Programming and Generalized Infinitesimal Gradient Ascent," in *Proc. of ICML*, 2003.
- [27] E. Hazan, "Introduction to Online Convex Optimization," 2019. [Online]. Available: <https://arxiv.org/abs/1909.05207>
- [28] O. Dekel, A. Flajolet, N. Haghtalab, and P. Jaillet, "Online Learning with a Hint," in *Proc. of NeurIPS*, 2017.
- [29] A. Bhaskara, A. Cutkosky, R. Kumar, and M. Purohit, "Online Learning with Imperfect Hints," in *Proc. of ICML*, 2020.
- [30] A. Rakhlin and K. Sridharan, "Optimization, Learning, and Games with Predictable Sequences," in *Proc. of NeurIPS*, 2013.
- [31] M. Mohri and S. Yang, "Accelerating Online Convex Optimization via Adaptive Prediction," in *Proc. of AISTATS*, 2016.
- [32] S. Shalev-Shwartz and Y. Singer, "A Primal-Dual Perspective of Online Learning Algorithms," *Mach. Learn.*, vol. 69, no. 2-3, 2007.
- [33] E. Nygren, R. K. Sitaraman, and J. Sun, "The Akamai Network: A Platform for High-Performance Internet Applications," *SIGOPS Oper. Syst. Rev.*, vol. 44, no. 3, 2010.
- [34] A. Sadeghi, F. Sheikholeslami, A. G. Marques, and G. B. Giannakis, "Reinforcement Learning for Adaptive Caching With Dynamic Storage Pricing," *IEEE J. Select. Areas Commun.*, vol. 37, no. 10, 2019.
- [35] J. Kwak, G. Paschos, and G. Iosifidis, "Dynamic Cache Rental and Content Caching in Elastic Wireless CDNs," in *Proc. of WiOpt*, 2018.
- [36] T. Chen, Q. Ling, and G. B. Giannakis, "An Online Convex Optimization Approach to Proactive Network Resource Allocation," *IEEE Trans. Signal Processing*, vol. 65, no. 24, 2017.
- [37] N. Liakopoulos *et al.*, "Cautious Regret Minimization: Online Optimization with Long-Term Budget Constraints," in *Proc. of ICML*, 2019.
- [38] V. Valls, G. Iosifidis, D. Leith, and L. Tassiulas, "Online Convex Optimization with Perturbed Constraints: Optimal Rates against Stronger Benchmarks," in *Proc. of AISTATS*, 2020.
- [39] X. Yi, X. Li, L. Xie, and K. H. Johansson, "Distributed Online Convex Optimization With Time-Varying Coupled Inequality Constraints," *IEEE Trans. Signal Processing*, vol. 68, 2020.
- [40] Y. Zhang, F. Feng, C. Wang, X. He, M. Wang, Y. Li, and Y. Zhang, "How to Retrain Recommender System? A Sequential Meta-Learning Method," in *Proc. of SIGIR*, 2020.
- [41] M. Zink, K. Suh, Y. Gu, and J. Kurose, "Characteristics of YouTube Network Traffic at a Campus Network - Measurements, Models, and Implications," *Comput. Netw.*, vol. 53, no. 4, 2009.
- [42] G. Paschos, G. Iosifidis, and G. Caire, "Cache Optimization Models and Algorithms," *Fnt in Commun. and Inf. Theory*, vol. 16, no. 3-4, 2020.
- [43] A. Giovanidis and A. Avranas, "Spatial Multi-LRU Caching for Wireless Networks with Coverage Overlaps," *SIGMETRICS Perform. Eval. Rev.*, vol. 44, no. 1, 2016.
- [44] E. Leonardi and G. Neglia, "Implicit Coordination of Caches in Small Cell Networks Under Unknown Popularity Profiles," *IEEE J. Select. Areas Commun.*, vol. 36, no. 6, 2018.
- [45] D. Tsigikar and T. Spyropoulos, "User-centric Optimization of Caching and Recomm. in Edge Cache Networks," in *Proc. of WoWMoM*, 2020.
- [46] H. B. McMahan, "A Survey of Algorithms and Analysis for Adaptive Online Learning," *J. Mach. Learn. Res.*, vol. 18, no. 1, 2017.
- [47] A. Beck, *First-Order Methods in Optimization*. SIAM-Society for Industrial and Applied Mathematics, 2017.
- [48] P. Auer, N. Cesa-Bianchi, and C. Gentile, "Adaptive and Self-Confident On-Line Learning Algorithms," *Journal of Computer and System Sciences*, vol. 64, no. 1, 2002.
- [49] G. S. Paschos, A. Destounis, and G. Iosifidis, "Online Convex Optimization for Caching Networks," *IEEE/ACM Trans. Networking*, vol. 28, no. 2, 2020.
- [50] "Amazon Elastic CDN Service - ElastiCache." [Online]. Available: <https://aws.amazon.com/elasticache/>
- [51] Juniper Networks, "The Elastic CDN Solution," Solution Brief, Dec. 2014. [Online]. Available: <https://www.juniper.net/assets/kr/kr/local/pdf/solutionbriefs/3510532-en.pdf>
- [52] S. Mannor, J. N. Tsitsiklis, and J. Y. Yu, "Online Learning with Sample Path Constraints," *J. Mach. Learn. Res.*, vol. 10, 2009.
- [53] F. Orabona, "A Modern Introduction to Online Learning," 2019. [Online]. Available: <https://arxiv.org/abs/1912.13213>
- [54] N. Mhaisen, "online-caching." [Online]. Available: <https://github.com/Naram-m/online-caching>