
Learning piecewise Lipschitz functions in changing environments

Maria-Florina Balcan
Carnegie Mellon University

Travis Dick
Carnegie Mellon University

Dravyansh Sharma
Carnegie Mellon University

Abstract

Optimization in the presence of sharp (non-Lipschitz), unpredictable (w.r.t. time and amount) changes is a challenging and largely unexplored problem of great significance. We consider the class of piecewise Lipschitz functions, which is the most general online setting considered in the literature for the problem, and arises naturally in various combinatorial algorithm selection problems where utility functions can have sharp discontinuities. The usual performance metric of ‘static’ regret minimizes the gap between the payoff accumulated and that of the best fixed point for the entire duration, and thus fails to capture changing environments. Shifting regret is a useful alternative, which allows for up to s environment *shifts*. In this work we provide an $O(\sqrt{sdT \log T} + sT^{1-\beta})$ regret bound for β -dispersed functions, where β roughly quantifies the rate at which discontinuities appear in the utility functions in expectation (typically $\beta \geq 1/2$ in problems of practical interest [Balcan et al., 2019, Balcan et al., 2018a]). We also present a lower bound tight up to sub-logarithmic factors. We further obtain improved bounds when selecting from a small pool of experts. We empirically demonstrate a key application of our algorithms to online clustering problems on popular benchmarks.

1 Introduction

Online optimization is well-studied in the online learning community [Cesa-Bianchi and Lugosi, 2006, Hazan et al., 2016]. It consists of a repeated game

with T iterations. At iteration t , the player chooses a point ρ_t from a compact decision set $\mathcal{C} \subset \mathbb{R}^d$; after the choice is committed, a bounded utility function $u_t : \mathcal{C} \rightarrow [0, H]$ is revealed. We treat u_t as a reward function to be maximized, although one may also consider minimizing a loss function. The goal of the player is to minimize the regret, defined as the difference between the online cumulative payoff (i.e. $\sum_{t=1}^T u_t(\rho_t)$) and the cumulative payoff using an optimal offline choice in hindsight. In many real world problems, like online routing [Awerbuch and Kleinberg, 2008, Talebi et al., 2018], detecting spam email/bots [Sculley and Wachman, 2007, Cormack et al., 2008] and ad/content ranking [Wauthier et al., 2013, Combes et al., 2015], it is often inadequate to assume a fixed point will yield good payoff at all times. It is more natural to compute regret against a stronger offline baseline, say one which is allowed to switch the point a few times (say s *shifts*), to accommodate ‘events’ which significantly change the function values for certain time periods. The switching points are neither known in advance nor explicitly stated during the course of the game. This stronger baseline is known as *shifting regret* [Herbster and Warmuth, 1998].

Shifting regret is a particularly relevant metric for online learning problems in the context of algorithm configuration. This is an important family of non-convex optimization problems where the goal is to decide in a data-driven way what algorithm to use from a large family of algorithms for a given problem domain. In the online setting, one has a configurable algorithm such as an algorithm for clustering data [Balcan et al., 2017], and must solve a series of related problems, such as clustering news articles each day for a news reader or clustering drugstore sales information to detect disease outbreaks. For problems of this nature, significant events in the world or changing habits of buyers might require changes in algorithm parameters, and we would like the online algorithms to adapt smoothly.

Related work: We present the first results for shifting regret for non-convex utility functions which potentially have sharp discontinuities. Restricting attention to specific kinds of decision sets $\mathcal{C}$

and utility function classes yields several important problems. If $\mathcal{C}$ is a convex set and utility functions are concave functions (i.e. corresponding loss functions are convex), we get the Online Convex Optimization (OCO) problem [Zinkevich, 2003], which is a generalization of online linear regression [Kivinen and Warmuth, 1997] and prediction with expert advice [Littlestone and Warmuth, 1994]. Algorithms with $O(\sqrt{sT} \log NT)$ regret are known for the case of s shifts for prediction with N experts and OCO on the N -simplex [Herbster and Warmuth, 1998, Cesa-Bianchi et al., 2012] using weight-sharing or regularization. We show how to extend the result to arbitrary compact sets of experts, and more general utility functions where convexity can no longer be exploited. Our key insight is to view the regularization as simultaneously inducing multiplicative weights update with restarts matching all possible shifted expert sequences, which allows us to use the *dispersion* condition introduced in [Balcan et al., 2018a]. Related notions like adaptive regret [Hazan and Seshadhri, 2007], strongly adaptive regret [Daniely et al., 2015, Jun et al., 2017], dynamic regret [Zinkevich, 2003, Jadbabaie et al., 2015] and sparse experts setting [Bousquet and Warmuth, 2002] have also been studied for finite experts.

Intuitively, a sequence of piecewise L -Lipschitz functions is well-*dispersed* if not too many functions are non-Lipschitz in the same region in $\mathcal{C}$. An assumption like this is necessary, since, even for piecewise constant functions, linear regret is unavoidable in the worst case [Cohen-Addad and Kanade, 2017]. Our shifting regret bounds are $O(\sqrt{sdT} \log T + sT^{1-\beta})$ which imply low regret for sufficiently dispersed (large enough β) functions. In a large range of applications, one can show $\beta \geq \frac{1}{2}$ [Balcan et al., 2018a]. This allows us to obtain tight regret bounds modulo sublogarithmic terms, providing a near-optimal characterization of the problem. Our analysis also readily extends to the closely related notion of adaptive regret [Hazan and Seshadhri, 2007]. Note that our setting generalizes the Online Non-Convex Learning (ONCL) problem where all functions are L -Lipschitz throughout [Maillard and Munos, 2010, Yang et al., 2018] for which shifting regret bounds have not been studied.

We demonstrate the effectiveness of our algorithm in solving the algorithm selection problem for a family of clustering algorithms parameterized by different ways to initialize k -means [Balcan et al., 2018b]. We consider the problem of online clustering, but unlike prior work which studies individual data points arriving in an online fashion [Liberty et al., 2016, Rakhlin et al., 2007], we look at complete clustering instances from some distribution(s) presented sequen-

tially. Our experiments provide the first empirical evaluation of online algorithms for piecewise Lipschitz functions — prior work is limited to theoretical analysis [Balcan et al., 2018a] or experiments for the batch setting [Balcan et al., 2018b]. Our results also have applications in non-convex online problems like portfolio optimization [Merton, 1976] and online non-convex SVMs [Ertekin et al., 2010]. More broadly, for applications where one needs to tune hyperparameters that are not ‘nice’, our results imply it is necessary and sufficient to look at dispersion.

2 Problem setup

Consider the following repeated game. At each round $1 \leq t \leq T$ we are required to choose $\rho_t \in \mathcal{C} \subset \mathbb{R}^d$, are presented a piecewise L -Lipschitz function $u_t : \mathcal{C} \rightarrow [0, H]$ and experience reward $u_t(\rho_t)$.

In this work we will study s -shifted regret and (m -sparse, s -shifted) regret notions defined below.

Definition 1. *The s -shifted regret (‘tracking regret’ in [Herbster and Warmuth, 1998]) is given by*

$$\mathbb{E} \left[\max_{\substack{\rho_i^* \in \mathcal{C}, \\ t_0=1 < t_1 < \dots < t_s=T+1}} \sum_{i=1}^s \sum_{t=t_{i-1}}^{t_i-1} (u_t(\rho_i^*) - u_t(\rho_t)) \right]$$

Note that for the i -th phase ($i \in [s]$) given by $[t_{i-1}, t_i - 1]$, the offline algorithm uses the same point ρ_i^* . The usual notion of regret compares the payoff of the online algorithm to the offline strategies that pick a fixed point $\rho^* \in \mathcal{C}$ for all $t \in [T]$ but here we compete against more powerful offline strategies that can use up to s distinct points ρ_i^* by switching the expert $s - 1$ times. For $s = 1$, we retrieve the standard static regret.

Definition 2. *Extend Definition 1 with an additional constraint on the number of distinct experts used, $|\{\rho_i^* \mid 1 \leq i \leq s\}| \leq m$. We call this (m -sparse, s -shifted) regret [Bousquet and Warmuth, 2002].*

This restriction makes sense if we think of the adversary as likely to reuse the same experts again, or the changing environment to experience recurring events with similar payoff distributions.

Without further assumptions, no algorithm achieves sublinear regret, even when the payout functions are piecewise constant [Cohen-Addad and Kanade, 2017]. We will characterize our regret bounds in terms of the ‘dispersion’ [Balcan et al., 2018a, Balcan et al., 2019] of the utility functions, which roughly says that discontinuities are not too concentrated. Several other restrictions can be seen as a special case [Rakhlin et al., 2011, Cohen-Addad and Kanade, 2017].

Definition 3. The sequence of utility functions $u_1, \dots, u_T$ is β -dispersed for the Lipschitz constant L if, for all T and for all $\epsilon \geq T^{-\beta}$, at most $\tilde{O}(\epsilon T)$ functions (the soft- O notation suppresses dependence on quantities beside ϵ, T and β) are not L -Lipschitz in any ball of size ϵ contained in $\mathcal{C}$. Further if the utility functions are obtained from some distribution, the random process generating them is said to be β -dispersed if the above holds in expectation, i.e. if for all T and for all $\epsilon \geq T^{-\beta}$,

$$\mathbb{E} \left[\max_{\rho \in \mathcal{C}} |\{t \mid u_t \text{ not } L\text{-Lipschitz in } \mathcal{B}(\rho, \epsilon)\}| \right] \leq \tilde{O}(\epsilon T)$$

For ‘static’ regret, a continuous version of exponential weight updates gives a tight bound of $\tilde{O}(\sqrt{dT} + T^{1-\beta})$ [Balcan et al., 2018a]. They further show that in several cases of practical interest one can prove dispersion with $\beta = 1/2$ and the algorithm enjoys $\tilde{O}(\sqrt{dT})$ regret. This algorithm may, however, have $\Omega(T)$ s -shifted regret even with a single switch ($s = 2$), and hence is not suited to changing environments (Appendix B).

3 Algorithms with low shifting regret

In this section we describe online algorithms with good shifting regret, but defer the actual regret analysis to Section 4. First we present a discretization based algorithm that simply uses a finite expert algorithm given a discretization of $\mathcal{C}$. This algorithm will give us the reasonable target regret bounds we should shoot for, although the discretization results in exponentially many experts.

Algorithm 1 Discrete Fixed Share Forecaster

Input: β , the dispersion parameter

1. Obtain a $T^{-\beta}$ -discretization $\mathcal{D}$ of $\mathcal{C}$ (i.e. any $c \in \mathcal{C}$ is within $T^{-\beta}$ of some $d \in \mathcal{D}$)
 2. Apply an optimal algorithm for finite experts with points in $\mathcal{D}$ as the experts (e.g. fixed share [Herbster and Warmuth, 1998])
-

We introduce a continuous version of the fixed share algorithm (Algorithm 2). We maintain weights for all points similar to the Exponential Forecaster of [Balcan et al., 2018a] which updates these weights in proportion to their exponentiated scaled utility $e^{\lambda u_t(\cdot)}$ ($\lambda \in (0, 1/H]$ is a *step size parameter* which controls how aggressively the algorithm updates its weights). The main difference is to update the weights with a mixture of the exponential update and a constant additive boost at all points in some proportion α (the *exploration parameter*, optimal value derived in Section 4) which remains fixed for the duration of the

game. This allows the algorithm to balance exploitation (exponential update assigns high weights to points with high past utility) with exploration, which turns out to be critical for success in changing environments. We will show this algorithm has good s -shifted regret in Section 4. It also enjoys good adaptive regret [Hazan and Seshadhri, 2007] (see Appendix D).

Algorithm 2 Fixed Share Exponential Forecaster (Fixed Share EF)

Input: step size parameter $\lambda \in (0, 1/H]$, exploration parameter $\alpha \in [0, 1]$

1. $w_1(\rho) = 1$ for all $\rho \in \mathcal{C}$
 2. For each $t = 1, 2, \dots, T$:
 - i. $W_t := \int_{\mathcal{C}} w_t(\rho) d\rho$
 - ii. Sample ρ with probability proportional to $w_t(\rho)$, i.e. with probability $p_t(\rho) = \frac{w_t(\rho)}{W_t}$
 - iii. Observe $u_t(\cdot)$
 - iv. Let $e_t(\rho) = e^{\lambda u_t(\rho)} w_t(\rho)$. For each $\rho \in \mathcal{C}$, set

$$w_{t+1}(\rho) = (1 - \alpha)e_t(\rho) + \frac{\alpha}{\text{VOL}(\mathcal{C})} \int_{\mathcal{C}} e_t(\rho) d\rho \quad (1)$$
-

Notice that it is not clear how to implement the Algorithm 2 from its description. We cannot store all the weights or sample easily since we have uncountably many points $\rho \in \mathcal{C}$. We will show how to efficiently sample according to p_t without necessarily computing it exactly or storing the exact weights in Section 5.

Algorithm 3 Generalized Share Exponential Forecaster (Generalized Share EF)

Input: step size parameter $\lambda \in (0, 1/H]$, exploration parameter $\alpha \in [0, 1]$, discount rate $\gamma \in [0, 1]$

1. $w_1(\rho) = 1$ for all $\rho \in \mathcal{C}$
 2. For each $t = 1, 2, \dots, T$:
 - i. $W_t := \int_{\mathcal{C}} w_t(\rho) d\rho$
 - ii. Sample ρ with probability proportional to $w_t(\rho)$, i.e. with probability $p_t(\rho) = \frac{w_t(\rho)}{W_t}$
 - iii. Let $e_t(\rho) = e^{\lambda u_t(\rho)} w_t(\rho)$ and $\beta_{i,t} = \frac{e^{-\gamma(t-i)}}{\sum_{j=1}^t e^{-\gamma(t-j)}}$. For each $\rho \in \mathcal{C}$, set $w_{t+1}(\rho)$ to

$$(1 - \alpha)e_t(\rho) + \alpha \left(\int_{\mathcal{C}} e_t(\rho) d\rho \right) \sum_{i=1}^t \beta_{i,t} p_i(\rho)$$
-

As it turns out adding equal weights to all points for exploration does not allow us to exploit recurring environments of the (m -sparse, s -shifted) setting very well. To overcome this, we replace the uniform update with a prior consisting of a weighted mixture of all the previous probability distributions used for sampling

(Algorithm 3). Notice that this includes uniformly random exploration as the first probability distribution $p_1(\cdot)$ is uniformly random, but the weight on this distribution decreases exponentially with time according to *discount rate* γ (more precisely, it decays by a factor $e^{-\gamma}$ with each time step). While exploration in Algorithm 2 is limited to starting afresh, here it includes partial resets to explore again from all past states, with an exponentially discounted rate (cf. Theorems 6, 7).

4 Analysis of algorithms

We will now analyse the algorithms in Section 3. At a high level, the algorithms have been designed to ensure that the optimal solution, and its neighborhood, in hindsight have a large total density. We achieve this by carefully setting the parameters, in particular the *exploration parameter* which controls the rate at which we allow our confidence on ‘good’ experts to change. Lipschitzness and dispersion are then used to ensure that solutions sufficiently close to the optimum are also good on average.

4.1 Regret bounds

In the remainder of this section we will have the following setting. We assume the utility functions $u_t : \mathcal{C} \rightarrow [0, H], t \in [T]$ are piecewise L -Lipschitz and β -dispersed (definition 3), where $\mathcal{C} \subset \mathbb{R}^d$ is contained in a ball of radius R .

Theorem 4. *Let $R^{\text{finite}}(T, s, N)$ denote the s -shifted regret for the finite experts problem on N experts, for the algorithm used in step 2 of Algorithm 1. Then Algorithm 1 enjoys s -shifted regret $R^C(T, s)$ which satisfies*

$$R^C(T, s) \leq R^{\text{finite}}\left(T, s, (3RT^\beta)^d\right) + (sH + L)O(T^{1-\beta}).$$

The proof of Theorem 4 is straightforward using the definition of dispersion and is deferred to Appendix A. This gives us the following target bound for our more efficient algorithms.

Corollary 5. *The s -shifted regret of Algorithm 1 is $O(H\sqrt{sT(d\log(RT^\beta) + \log(T/s))} + (sH + L)T^{1-\beta})$.*

Proof. There are known algorithms e.g. Fixed-Share ([Herbster and Warmuth, 1998]) which obtain $R^{\text{finite}}(T, s, N) \leq O(\sqrt{sT\log(NT/s)})$. Applying Theorem 4 gives the desired upper bound. $\square$

Under the same conditions, we will show the following bounds for our algorithms. In the following statements, we give approximate values for the parameters α, γ and λ under the assumptions $m \ll s, s \ll T$. See proofs in Appendix C for more precise values.

Theorem 6. *The s -shifted regret of Algorithm 2 with $\alpha = s/T$ and $\lambda = \sqrt{s(d\log(RT^\beta) + \log(T/s))}/T/H$ is $O(H\sqrt{sT(d\log(RT^\beta) + \log(T/s))} + (sH + L)T^{1-\beta})$.*

Remark. *The algorithms assume knowledge of s/T , the average number of shifts per time. For unknown s , the strongly adaptive algorithms of [Daniely et al., 2015, Jun et al., 2017] can be used with the same meta-algorithms and substituting continuous exponential forecasters as black-box algorithms.*

Similarly for Algorithm 3 we can show low (m -sparse, s -shifted) regret as well. (In particular this implies s -shifted regret almost as good as Algorithm 2.)

Theorem 7. *The (m -sparse, s -shifted) regret of Algorithm 3 is $O(H\sqrt{T(md\log(RT^\beta) + s\log(mT/s))} + (mH + L)T^{1-\beta})$ for $\alpha = s/T, \gamma = s/mT$ and $\lambda = \sqrt{(md\log(RT^\beta) + s\log(T/s))/T/H}$.*

4.2 Proof sketch and insights

We start with some observations about the weights W_t in Algorithm 2.

Lemma 8 (Algorithm 2). $W_{t+1} = \int_{\mathcal{C}} e^{\lambda u_t(\rho)} w_t(\rho) d\rho$.

The update rule (1) had the uniform exploration term scaled just appropriately so this relation is satisfied. We will now relate W_t with weights resulting from pure exponential updates, i.e. $\alpha = 0$ in Algorithm 2 (also the Exponential Forecaster algorithm of [Balcan et al., 2018a]). The following definition corresponds to weights for running Exponential Forecaster starting at some time τ .

Definition 9. *For any $\rho \in \mathcal{C}$ and $\tau \leq \tau' \in [T]$ define $\tilde{w}(\rho; \tau, \tau')$ to be the weight of expert ρ , and $\tilde{W}(\tau, \tau')$ to be the normalizing constant, if we ran the Exponential Forecaster of [Balcan et al., 2018a] starting from time τ up till time τ' , i.e. $\tilde{w}(\rho; \tau, \tau') := e^{\lambda \sum_{t=\tau}^{\tau'-1} u_t(\rho)}$ and $\tilde{W}(\tau, \tau') := \int_{\mathcal{C}} \tilde{w}(\rho; \tau, \tau') d\rho$.*

We consider Algorithm 4 obtained by a slight modification in the update rule (1) of Fixed Share EF (Algorithm 2) which makes it easier to analyze. Essentially we replace the deterministic α -mixture by a randomized one, so at each turn we either explicitly ‘restart’ with probability α by putting the same weight on each point, or else apply the exponential update. We note that Algorithm 4 is introduced to simplify the proof of Theorem 6, and in particular does *not* result in low regret itself. The issue is that even though the weights are correct in expectation (Lemma 10), their ratio (probability $p_t(\rho)$) is not. In particular, the optimal parameter value of α for Fixed Share EF allows the possibility of pure exponential updates over a long period of time with a constant probability in Algorithm

4, which implies linear regret (see Appendix B, Theorem 21). This also makes the implementation of Fixed Share EF somewhat trickier (Section 5).

Algorithm 4 Random Restarts Exponential Forecaster (Random Restarts EF)

Input: step size parameter $\lambda \in (0, 1/H]$, exploration parameter $\alpha \in [0, 1]$

1. $\hat{w}_1(\rho) = 1$ for all $\rho \in \mathcal{C}$
2. For each $t = 1, 2, \dots, T$:
 - i. $\hat{W}_t := \int_{\mathcal{C}} \hat{w}_t(\rho) d\rho$
 - ii. Sample ρ with probability proportional to $\hat{w}_t(\rho)$, i.e. with probability $p_t(\rho) = \frac{\hat{w}_t(\rho)}{\hat{W}_t}$
 - iii. Sample z_t uniformly in $[0, 1]$ and set

$$\hat{w}_{t+1}(\rho) = \begin{cases} e^{\lambda u_t(\rho)} \hat{w}_t(\rho) & \text{if } z_t < 1 - \alpha \\ \frac{\int_{\mathcal{C}} e^{\lambda u_t(\rho)} \hat{w}_t(\rho) d\rho}{\text{VOL}(\mathcal{C})} & \text{otherwise} \end{cases}$$

The expected weights of Algorithm 4 (over the coin flips used in weight setting) are the same as the actual weights of Algorithm 2 (proof in Appendix C).

Lemma 10. (*Algorithm 2*) For each $t \in [T]$, $w_t(\rho) = \mathbb{E}[\hat{w}_t(\rho)]$ and $W_t = \mathbb{E}[\hat{W}_t]$, where the expectations are over random restarts $\mathbf{z}_t = \{z_1, \dots, z_{t-1}\}$.

The next lemma provides intuition for looking at our algorithm as a weighted superposition of several exponential update subsequences with restarts. This novel insight establishes a tight connection between the algorithms and is crucial for our analysis.

Lemma 11. (*Algorithm 2*) W_{T+1} equals the sum

$$\sum_{s \in [T]} \sum_{t_0=1 < t_1 < \dots < t_s=T+1} \frac{\alpha^{s-1}(1-\alpha)^{T-s}}{\text{VOL}(\mathcal{C})^{s-1}} \prod_{i=1}^s \tilde{W}(t_{i-1}, t_i)$$

Proof Sketch. Each term corresponds to the weight when we pick a number $s \in [T]$ for the number of times we start afresh with a uniformly random point ρ at times $\mathbf{t}_s = \{t_1, \dots, t_{s-1}\}$ and do the regular exponential weighted forecaster in the intermediate periods. We have a weighted sum over all these terms with a factor $\alpha/\text{VOL}(\mathcal{C})$ for each time we restart and $(1-\alpha)$ for each time we continue with the Exponential Forecaster. $\square$

We will now prove Theorem 6. The main idea is to show that the normalized exploration helps the total weights to provide a lower bound for the algorithm payoff. Also the total weights are competitive against the optimal payoff as they contain the exponential updates with the optimal set of switching points in Lemma 11 with a sufficiently large (‘probability’) coefficient.

Proof sketch of Theorem 6. We provide an upper and lower bound to $\frac{W_{T+1}}{W_1}$. The upper bound uses Lemma 8 and helps us lower bound the performance of the algorithm (see Appendix C) as

$$\frac{W_{T+1}}{W_1} \leq \exp\left(\frac{P(\mathcal{A})(e^{H\lambda} - 1)}{H}\right) \quad (2)$$

where $P(\mathcal{A})$ is the expected total payoff for Algorithm 2. We now upper bound the optimal payoff OPT by providing a lower bound for $\frac{W_{T+1}}{W_1}$. By Lemma 11 we have

$$W_{T+1} \geq \frac{\alpha^{s-1}(1-\alpha)^{T-s}}{\text{VOL}(\mathcal{C})^{s-1}} \prod_{i=1}^s \tilde{W}(t_{i-1}^*, t_i^*)$$

by dropping all terms save those that ‘restart’ exactly at the OPT expert switches $t_{0:s}^*$. Now using β -dispersion we can show (full proof in Appendix C)

$$\frac{W_{T+1}}{W_1} \geq \frac{\alpha^{s-1}(1-\alpha)^{T-s}}{(RT^\beta)^{sd}} e^{\lambda(OPT - (sH+L)O(T^{1-\beta}))}$$

Putting together with the upper bound (2), rearranging and optimizing the difference for α and λ concludes the proof. (See Appendix C for a full proof.) $\square$

We now analyze Algorithm 3 for the sparse experts setting. We can adapt proofs of Lemmas 8 and 11 to easily establish Lemmas 12 and 13.

Lemma 12 (Algorithm 3). $W_{t+1} = \int_{\mathcal{C}} e^{\lambda u_t(\rho)} w_t(\rho) d\rho$.

Lemma 13. Let $\pi_t(\rho) = \sum_{i=1}^{t-1} \beta_{i,t} p_i(\rho)$. For Algorithm 3, W_{T+1} can be shown to be equal to the sum

$$\sum_{s \in [T]} \sum_{t_0=1 < \dots < t_s=T+1} \alpha^{s-1}(1-\alpha)^{T-s} \prod_{i=1}^s \tilde{W}(\pi_{t_{i-1}}; t_{i-1}, t_i)$$

where $\tilde{W}(p; \tau, \tau') := \int_{\mathcal{C}} p(\rho) \tilde{w}(\rho; \tau, \tau') d\rho$.

Corollary 14. $W_T \geq \alpha(1-\alpha)^{T-t} \tilde{W}(\pi_t; t, T) W_t$, for all $t < T$.

Proof. Consider the probability of last ‘reset’ (setting $w_t(\rho) = W_t \pi_t(\rho)$) at time t when computing W_{T+1} as the expected weight of a random restart version which matches Algorithm 3 till time t . $\square$

Now to prove Theorem 7, we show that the total weight is competitive with running exponential updates on all partitions (in particular the optimal partition) of $[T]$ into m subsets with s switches, intuitively the property of restarting exploration from all past points crucially allows us to “jump” across intervals where a given expert was inactive (or bad).

Proof sketch of Theorem 7. We provide an upper and lower bound to $\frac{W_{T+1}}{W_1}$ similar to Theorem 6. Using Lemma 12 we can show that inequality 2 holds here as well. By Corollary 14 and Lemma 23 (which relates $\pi_t(\cdot)$ to past weights, proved in Appendix C), and β -dispersion we can show a better lower bound. Putting together the lower and upper bounds, rearranging and optimizing for γ, α, λ concludes the proof. $\square$

5 Efficient implementation

In this section we show that the Fixed Share Exponential Forecaster algorithm (Algorithm 2) can be implemented efficiently when u_t 's are piecewise concave (diminishing returns). In particular we overcome the need to explicitly compute and update $w_t(\rho)$ (there are uncountably infinite ρ in $\mathcal{C}$) by showing that we can sample the points according to $p_t(\rho)$ directly.

The high-level strategy is to show (Lemma 16) that $p_t(\rho)$ is a mixture of t distributions which are Exponential Forecaster distributions from [Balcan et al., 2018a] i.e. $\tilde{p}_i(\rho) := \frac{\tilde{w}(\rho; i, t)}{\tilde{W}(i, t)}$ for each $1 \leq i \leq t$, with proportions $C_{i, t}$. As shown in [Balcan et al., 2018a] these distributions can be approximately sampled from (exactly in the one-dimensional case, $\mathcal{C} \subset \mathbb{R}$), summarized below as Algorithm BDV-18. We need to sample from one of these t distributions with probability $C_{t, i}$ to get the distribution p_t , and we can approximate these coefficients efficiently (or compute exactly in one-dimensional case). The rest of the section discusses how to do these approximations efficiently, and with small extra expected regret. Asymptotically we get the same bound as the exact algorithm. (Formal proofs in Appendix E).

Algorithm BDV-18: Simply *integrate* pieces of the exponentiated utility function, pick a piece with probability proportional to its integral, and *sample from* that piece. [Lovász and Vempala, 2006] show how to efficiently *sample from* and *integrate* logconcave distributions. See [Balcan et al., 2018a] for more details.

The coefficients have a simple form in terms of normalizing constants W_t 's of the rounds so far, so we first express W_{t+1} in terms of W_t 's from previous rounds and some $\tilde{W}(i, j)$'s.

Lemma 15. *In Algorithm 2, for $t \geq 1$,*

$$W_{t+1} = (1 - \alpha)^{t-1} \tilde{W}(1, t+1) + \frac{\alpha}{\text{VOL}(\mathcal{C})} \sum_{i=2}^t \left[(1 - \alpha)^{t-i} W_i \tilde{W}(i, t+1) \right]$$

As indicated above, $p_t(\rho)$ is a mixture of t distributions.

Lemma 16. *In Algorithm 2, for $t \geq 1$, $p_t(\rho) =$*

$\sum_{i=1}^t C_{t, i} \frac{\tilde{w}(\rho; i, t)}{\tilde{W}(i, t)}$. *The coefficients $C_{t, i}$ are given by*

$$C_{t, i} = \begin{cases} 1 & i = t = 1 \\ \alpha & i = t > 1 \\ (1 - \alpha) \frac{W_{t-1}}{W_t} \frac{\tilde{W}(i, t)}{\tilde{W}(i, t-1)} C_{t-1, i} & i < t \end{cases}$$

and $(C_{t, 1}, \dots, C_{t, t})$ lies on the probability simplex Δ^{t-1} .

The observations above allow us to write the algorithms for efficiently implementing Fixed Share EF, for which we obtain formal guarantees in Theorem 17. We present an approximate algorithm (Algorithm 5) with the same expected regret as in Theorem 6 (and also present an exact algorithm, Algorithm 6 in Appendix E, for $d = 1$). We say Algorithm 5 gives a (η, ζ) estimate of Algorithm 2, i.e. with probability at least $1 - \zeta$, its expected payoff is within a factor of e^η of that of Algorithm 2.

Algorithm 5 Fixed Share Exponential Forecaster - efficient approximate implementation

Input: approximation parameter $\eta \in (0, 1)$, confidence parameter $\zeta \in (0, 1)$

1. $W_1 = \text{VOL}(\mathcal{C})$
 2. For each $t = 1, 2, \dots, T$:
 - i. Estimate $C_{t, j}$ using Lemma 16 for each $1 \leq j \leq t$.
 - ii. Sample i with probability $C_{t, i}$.
 - iii. Sample ρ with probability approximately proportional to $\tilde{w}(\rho; i, t)$ by running Algorithm BDV-18 with approximation-confidence parameters $(\eta/3, \zeta/2)$.
 - iv. Estimate W_{t+1} using Lemma 15. Algorithm BDV-18 to get $(\eta/6T, \eta/2T^2)$ estimates for all $\tilde{W}(\tau, \tau')$ and memoize values of $W_i, i \leq t$.
-

Theorem 17. *If utility functions are piecewise concave and L -Lipschitz, we can approximately sample a point ρ with probability $p_{t+1}(\rho)$ in time $\tilde{O}(Kd^4T^4)$ for approximation parameters $\eta = \zeta = 1/\sqrt{T}$ and $\lambda = \sqrt{s(d \ln(RT^\beta) + \ln(T/s))/T/H}$ and enjoy the same regret bound as the exact algorithm. (K is number of discontinuities in u_t 's).*

Note that in this section we concerned ourselves with developing a $\text{poly}(d, T)$ algorithm. For special cases of practical interest, like one-dimensional piecewise constant functions, we can implement much faster $O(K \log KT)$ algorithms as noted in Section 7.

6 Lower bounds

We prove our lower bound for $\mathcal{C} = [0, 1]$ and $H = 1$. Also we will consider functions which are β -dispersed

and 0-Lipschitz (piecewise constant). For such utility functions $u_1, \dots, u_T$ we have shown in Section 4 that the s -shifted regret is $O(\sqrt{sT \log T} + sT^{1-\beta})$. Here we will establish a lower bound of $\Omega(\sqrt{sT} + sT^{1-\beta})$.

We show a range of values of s, β where the stated lower bound is achieved. For $s = 1$, this improves over the lower bound construction of [Balcan et al., 2018a] where the lower bound is shown only for $\beta = 1/2$. In particular our results establish an almost tight characterization of static and dynamic regret under dispersion.

Theorem 18. *For each $\beta > \frac{\log 3s}{\log T}$, there exist utility functions $u_1, \dots, u_T : [0, 1] \rightarrow [0, 1]$ which are β -dispersed, and the s -shifted regret of any online algorithm is $\Omega(\sqrt{sT} + sT^{1-\beta})$.*

Proof. We perform the construction in $\Theta(s)$ phases, each phase accumulating $\Omega(\sqrt{T/s} + T^{1-\beta})$ regret, yielding the desired lower bound.

Let $I_1 = [0, 1]$. In the first phase, for the first $\frac{T-3sT^{1-\beta}}{s}$ functions we have a single discontinuity in the interval $(\frac{1}{2}(1 - \frac{1}{3s}), \frac{1}{2}(1 + \frac{1}{3s})) \subseteq (\frac{1}{3}, \frac{2}{3})$. The functions have payoff 1 before or after (with probability 1/2 each) their discontinuity point, and zero elsewhere. We introduce $3T^{1-\beta}$ functions each for the same discontinuity point, and set the discontinuity points $T^{-\beta}$ apart for β -dispersion. This gives us $\frac{1/3s}{T^{-\beta}} - 1$ potential points inside $[\frac{1}{3}, \frac{2}{3}]$, so we can support $3T^{1-\beta} \left(\frac{1/3s}{T^{-\beta}} - 1 \right) = \frac{T}{s} - 3T^{1-\beta}$ such functions ($\frac{T}{s} - 3T^{1-\beta} > 0$ since $\beta > \frac{\log 3s}{\log T}$). By Lemma 31 (Appendix F) we accumulate $\Omega(\sqrt{\frac{T-3sT^{1-\beta}}{s}}) = \Omega(\sqrt{T/s})$ regret for this part of the phase in expectation. Let I'_1 be the interval from among $[0, \frac{1}{2}(1 - \frac{1}{3s})]$ and $[\frac{1}{2}(1 + \frac{1}{3s}), 1]$ with more payoff in the phase so far. The next function has payoff 1 only at first or second half of I'_1 (with probability 1/2) and zero everywhere else. Any algorithm accumulates expected regret 1/2 on this round. We repeat this in successively halved intervals. β -dispersion is satisfied since we use only $\Theta(T^{1-\beta})$ functions in the interval I' of size greater than $1/3$, and we accumulate an additional $\Omega(T^{1-\beta})$ regret. Notice there is a fixed point used by the optimal adversary for this phase.

Finally we repeat the construction inside the largest interval with no discontinuities at the end of the last phase for the next phase. Note that at the i -th phase the interval size will be $\Theta(\frac{1}{i})$. Indeed at the end of the first round we have ‘unused’ intervals of size $\frac{1}{2}(1 - \frac{1}{3s}), \frac{1}{4}(1 - \frac{1}{3s}), \frac{1}{8}(1 - \frac{1}{3s}), \dots$. At the $i = 2^j$ -th phase, we’ll be repeating inside an interval of size $\frac{1}{2^{j+1}}(1 - \frac{1}{3s}) = \Theta(\frac{1}{i})$. This allows us to run $\Theta(s)$ phases and get the desired lower bound (intervals must be of size at least $\frac{1}{s}$ to support the construction). $\square$

7 Experiments

The simplest demonstration of significance of our algorithm in a changing environment is to consider the 2-shifted regret when a single expert shift occurs. We consider an artificial online optimization problem first, and will then look at applications to online clustering. Let $\mathcal{C} = [0, 1]$. Define utility functions

$$u^{(0)}(\rho) = \begin{cases} 1 & \text{if } \rho < \frac{1}{2} \\ 0 & \text{if } \rho \geq \frac{1}{2} \end{cases} \quad \text{and} \quad u^{(1)}(\rho) = \begin{cases} 0 & \text{if } \rho < \frac{1}{2} \\ 1 & \text{if } \rho \geq \frac{1}{2} \end{cases}$$

Now consider the instance where $u^{(0)}(\rho)$ is presented for the first $T/2$ rounds and $u^{(1)}(\rho)$ is presented for the remaining rounds. We observe constant average regret for the Exponential Forecaster algorithm, while Fixed Share regret decays as $O(1/\sqrt{T})$ (Figure 2). While the example is simple and artificial, it qualitatively captures why Fixed Share dominates Exponential Forecaster here — because the best expert changes and the old expert is no longer competitive. (cf. Appendix B)

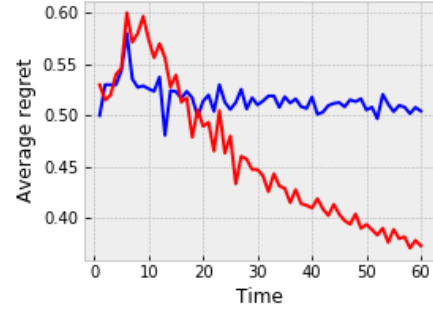


Figure 2: Average 2-shifted regret vs game duration T for a game with single expert shift. Color scheme: **Exponential Forecaster**, **Fixed Share EF**

[Arthur and Vassilvitskii, 2007] proposed k -means++, a celebrated algorithm which shows the importance of initial seed centers in clustering using the k -means algorithm (also called Lloyd’s method). [Balcan et al., 2018b] generalize it to $(\bar{\alpha}, 2)$ -Lloyds++-clustering, which interpolates between random initial seeds (vanilla k -means, $\bar{\alpha} = 0$), k -means++ ($\bar{\alpha} = 2$) and farthest-first traversal ($\bar{\alpha} = \infty$) [Gonzalez, 1985, Dasgupta and Long, 2005] using a single parameter $\bar{\alpha}$. The clustering objective (we use the Hamming distance to the optimal clustering, i.e. the fraction of points assigned to different clusters by the algorithm and the target clustering) is a piecewise constant function of $\bar{\alpha}$, and the best clustering may be obtained for a value of $\bar{\alpha}$ specific to a given problem domain. In an online problem, where clustering instances arrive in a sequential fashion, determining good values of $\bar{\alpha}$ becomes an online optimization problem on piecewise Lipschitz functions.

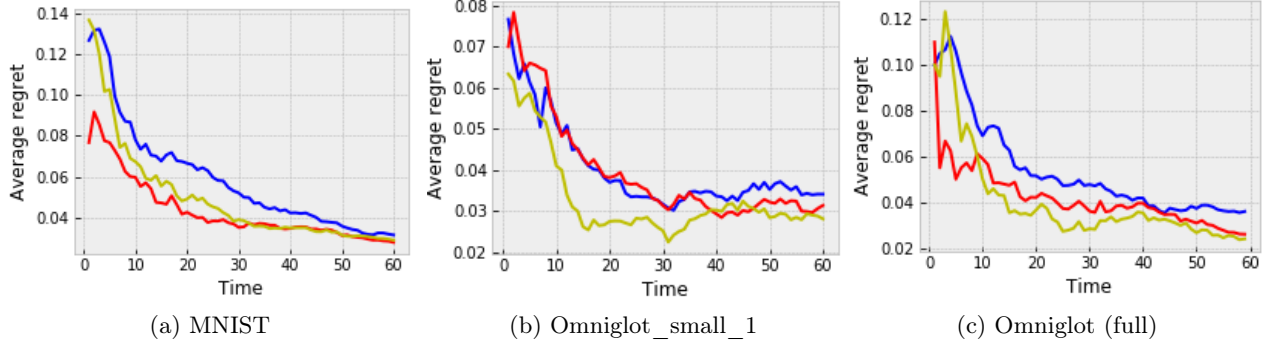


Figure 1: Average 2-shifted regret vs game duration T for online clustering against 2-shifted distributions. Color scheme: **Exponential Forecaster**, **Fixed Share EF**, **Generalized Share EF**

We perform our evaluation on four benchmark datasets to cover a range of examples-set sizes, N and number of clusters, k : *MNIST*, 28×28 binary images of handwritten digits with 60,000 training examples for 10 classes [Deng, 2012]; *Omniglot*, 105×105 binary images of handwritten characters across 30 alphabets with 19,280 examples [Lake et al., 2015]; *Omniglot_small_1*, a “minimal” Omniglot split with only 5 alphabets and 2720 examples.

We consider a sequence of clustering instances drawn from the four datasets and compare our algorithms Fixed Share EF (Algorithm 2) and Generalized Share EF (Algorithm 3) with the Exponential Forecaster algorithm of [Balcan et al., 2018a]. At each time $t \leq T \leq 60$ we sample a subset of the dataset of size 100. For each T , we take uniformly random points from half the classes (even class labels) at times $t = 1, \dots, T/2$ and from the remaining classes (odd class labels) at $T/2 < t \leq T$. We determine the hamming cost of $(\bar{\alpha}, 2)$ -Lloyds++-clustering for $\alpha \in \mathcal{C} = [0, 10]$ which is used as the piecewise constant loss function (or payoff is the fraction of points assigned correctly) for the online optimization game. Notice the Lipschitz constant $L = 0$ since we have piecewise constant utility, and utility function values lie in $[0, 1]$. We set exploration parameter $\alpha = 1/T$ and decay parameter $\gamma = 1/T$ in our algorithms. We plot average 2-shifted regret until time T (i.e. R_T/T) and take average over 20 runs to get smooth curves. (Figure 1). Unlike Figure 2, the optimal clustering parameters before the shift might be relatively competitive to new optimal parameters. So the Exponential Forecaster performance is not terrible, although our algorithms still outperform it noticeably.

We observe that our algorithms have significantly lower regrets (about 15-40% relative for the datasets considered, for $T \geq 40$) compared to the Exponential Forecaster algorithm across all datasets. We also note that the exact advantage of adding exploration to exponential updates varies with datasets and problem

instances. In Appendix G we have compiled further experiments that reaffirm the strengths of our approach against different *changing environments* and also compare against the static setting.

Remark. For the applications considered above, the utility functions are piecewise constant with $d = 1$. For these it is possible to simply maintain the weight on each piece of $\Sigma_t u_t(\rho)$ in $O(K \log Kt)$ time for round t where each $u_t(\cdot)$ has $O(K)$ pieces by using a simple interval tree data structure [Cohen-Addad and Kanade, 2017]. The tree lazily maintains weight for each of $O(Kt)$ pieces, takes $O(\log Kt)$ time each for lazy insertion of $O(K)$ new pieces and allows drawing with probability proportional to weight in $O(\log Kt)$ time. Similarly $O(K \log Kt)$ updates are possible for Algorithm 3. Section 5 addresses the harder problem of polynomial time implementation of Algorithm 2 for arbitrary d .

8 Discussion and open problems

We presented approaches which trade off exploitation with exploration for the online optimization problem to obtain low shifting regret for the case of general non-convex functions with sharp but dispersed discontinuities. Optimizing for the stronger baseline of shifting regret leads to empirically better payout, as we have shown via experiments bearing applications to algorithm configuration. Our focus here is on the full-information setting which corresponds to the entire utility function being revealed at each iteration, and we present almost tight theoretical results for it. Other relevant settings include bandit and semi-bandit feedback where the function value is revealed for only the selected point or a subset of the space containing the point. It would be interesting to obtain low shifting regret in these settings [Auer et al., 2002].

9 Acknowledgements

We thank Ellen Vitercik for helpful feedback. This work was supported in part by NSF grants CCF-1535967, IIS-1618714, IIS-1901403, CCF-1910321, SES-1919453, an Amazon Research Award, a Microsoft Research Faculty Fellowship, a Bloomberg Data Science research grant, and by the generosity of Eric and Wendy Schmidt by recommendation of the Schmidt Futures program. Views expressed in this work do not necessarily reflect those of any funding agency.

References

- [Adamskiy et al., 2012] Adamskiy, D., Koolen, W. M., Chernov, A., and Vovk, V. (2012). A closer look at adaptive regret. In *International Conference on Algorithmic Learning Theory*, pages 290–304. Springer.
- [Arthur and Vassilvitskii, 2007] Arthur, D. and Vassilvitskii, S. (2007). k-means++: The advantages of careful seeding. In *Proceedings of the eighteenth annual ACM-SIAM symposium on Discrete algorithms*, pages 1027–1035. Society for Industrial and Applied Mathematics.
- [Auer et al., 2002] Auer, P., Cesa-Bianchi, N., Freund, Y., and Schapire, R. E. (2002). The nonstochastic multiarmed bandit problem. *SIAM journal on computing*, 32(1):48–77.
- [Awerbuch and Kleinberg, 2008] Awerbuch, B. and Kleinberg, R. (2008). Online linear optimization and adaptive routing. *Journal of Computer and System Sciences*, 74(1):97–114.
- [Balcan et al., 2019] Balcan, M.-F., Dick, T., and Pegden, W. (2019). Semi-bandit Optimization in the Dispersed Setting. *arXiv e-prints*, page arXiv:1904.09014.
- [Balcan et al., 2018a] Balcan, M.-F., Dick, T., and Vitercik, E. (2018a). Dispersion for data-driven algorithm design, online learning, and private optimization. In *2018 IEEE 59th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 603–614. IEEE.
- [Balcan et al., 2017] Balcan, M.-F., Nagarajan, V., Vitercik, E., and White, C. (2017). Learning-theoretic foundations of algorithm configuration for combinatorial partitioning problems. In *Conference on Learning Theory*, pages 213–274.
- [Balcan et al., 2018b] Balcan, M.-F. F., Dick, T., and White, C. (2018b). Data-driven clustering via parameterized lloyd’s families. In *Advances in Neural Information Processing Systems*, pages 10641–10651.
- [Bassily et al., 2014] Bassily, R., Smith, A., and Thakurta, A. (2014). Private empirical risk minimization: Efficient algorithms and tight error bounds. In *2014 IEEE 55th Annual Symposium on Foundations of Computer Science*, pages 464–473. IEEE.
- [Ben-David et al., 2009] Ben-David, S., Pál, D., and Shalev-Shwartz, S. (2009). Agnostic online learning. In *COLT*.
- [Bousquet and Warmuth, 2002] Bousquet, O. and Warmuth, M. K. (2002). Tracking a small set of experts by mixing past posteriors. *Journal of Machine Learning Research*, 3(Nov):363–396.
- [Cesa-Bianchi et al., 2012] Cesa-Bianchi, N., Gaillard, P., Lugosi, G., and Stoltz, G. (2012). Mirror descent meets fixed share (and feels no regret). In *Advances in Neural Information Processing Systems*, pages 980–988.
- [Cesa-Bianchi and Lugosi, 2006] Cesa-Bianchi, N. and Lugosi, G. (2006). *Prediction, learning, and games*. Cambridge university press.
- [Cohen-Addad and Kanade, 2017] Cohen-Addad, V. and Kanade, V. (2017). Online optimization of smoothed piecewise constant functions. In *Artificial Intelligence and Statistics*, pages 412–420.
- [Combes et al., 2015] Combes, R., Magureanu, S., Proutiere, A., and Laroche, C. (2015). Learning to rank: Regret lower bounds and efficient algorithms. *ACM SIGMETRICS Performance Evaluation Review*, 43(1):231–244.
- [Cormack et al., 2008] Cormack, G. V. et al. (2008). Email spam filtering: A systematic review. *Foundations and Trends® in Information Retrieval*, 1(4):335–455.
- [Daniely et al., 2015] Daniely, A., Gonen, A., and Shalev-Shwartz, S. (2015). Strongly adaptive online learning. In *International Conference on Machine Learning*, pages 1405–1411.
- [Dasgupta and Long, 2005] Dasgupta, S. and Long, P. M. (2005). Performance guarantees for hierarchical clustering. *Journal of Computer and System Sciences*, 70(4):555–569.
- [Deng, 2012] Deng, L. (2012). The mnist database of handwritten digit images for machine learning research [best of the web]. *IEEE Signal Processing Magazine*, 29(6):141–142.
- [Ertekin et al., 2010] Ertekin, S., Bottou, L., and Giles, C. L. (2010). Nonconvex online support vector machines. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 33(2):368–381.

- [Gonzalez, 1985] Gonzalez, T. F. (1985). Clustering to minimize the maximum intercluster distance. *Theoretical Computer Science*, 38:293–306.
- [Hazan et al., 2016] Hazan, E. et al. (2016). Introduction to online convex optimization. *Foundations and Trends® in Optimization*, 2(3-4):157–325.
- [Hazan and Seshadhri, 2007] Hazan, E. and Seshadhri, C. (2007). Adaptive algorithms for online decision problems. In *Electronic colloquium on computational complexity (ECCC)*, volume 14.
- [Herbster and Warmuth, 1998] Herbster, M. and Warmuth, M. K. (1998). Tracking the best expert. *Machine learning*, 32(2):151–178.
- [Jadbabaie et al., 2015] Jadbabaie, A., Rakhlin, A., Shahrampour, S., and Sridharan, K. (2015). Online optimization: Competing with dynamic comparators. In *Artificial Intelligence and Statistics*, pages 398–406.
- [Jun et al., 2017] Jun, K.-S., Orabona, F., Wright, S., and Willett, R. (2017). Improved strongly adaptive online learning using coin betting. In *Artificial Intelligence and Statistics*, pages 943–951.
- [Kivinen and Warmuth, 1997] Kivinen, J. and Warmuth, M. K. (1997). Exponentiated gradient versus gradient descent for linear predictors. *information and computation*, 132(1):1–63.
- [Lake et al., 2015] Lake, B. M., Salakhutdinov, R., and Tenenbaum, J. B. (2015). Human-level concept learning through probabilistic program induction. *Science*, 350(6266):1332–1338.
- [Liberty et al., 2016] Liberty, E., Sriharsha, R., and Sviridenko, M. (2016). An algorithm for online k-means clustering. In *2016 Proceedings of the eighteenth workshop on algorithm engineering and experiments (ALENEX)*, pages 81–89. SIAM.
- [Littlestone and Warmuth, 1994] Littlestone, N. and Warmuth, M. K. (1994). The weighted majority algorithm. *Information and computation*, 108(2):212–261.
- [Lovász and Vempala, 2006] Lovász, L. and Vempala, S. (2006). Fast algorithms for logconcave functions: Sampling, rounding, integration and optimization. In *2006 47th Annual IEEE Symposium on Foundations of Computer Science (FOCS'06)*, pages 57–68. IEEE.
- [Maillard and Munos, 2010] Maillard, O.-A. and Munos, R. (2010). Online learning in adversarial lipschitz environments. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 305–320. Springer.
- [Merton, 1976] Merton, R. C. (1976). Option pricing when underlying stock returns are discontinuous. *Journal of financial economics*, 3(1-2):125–144.
- [Rakhlin et al., 2007] Rakhlin, A., Abernethy, J., and Bartlett, P. L. (2007). Online discovery of similarity mappings. In *Proceedings of the 24th international conference on Machine learning*, pages 767–774. ACM.
- [Rakhlin et al., 2011] Rakhlin, A., Sridharan, K., and Tewari, A. (2011). Online learning: Stochastic, constrained, and smoothed adversaries. In *Advances in neural information processing systems*, pages 1764–1772.
- [Sculley and Wachman, 2007] Sculley, D. and Wachman, G. M. (2007). Relaxed online svms for spam filtering. In *Proceedings of the 30th annual international ACM SIGIR conference on Research and development in information retrieval*, pages 415–422. ACM.
- [Talebi et al., 2018] Talebi, M. S., Zou, Z., Combes, R., Proutiere, A., and Johansson, M. (2018). Stochastic online shortest path routing: The value of feedback. *IEEE Transactions on Automatic Control*, 63(4):915–930.
- [Wauthier et al., 2013] Wauthier, F., Jordan, M., and Jojic, N. (2013). Efficient ranking from pairwise comparisons. In *International Conference on Machine Learning*, pages 109–117.
- [Yang et al., 2018] Yang, L., Deng, L., Hajiesmaili, M. H., Tan, C., and Wong, W. S. (2018). An optimal algorithm for online non-convex learning. *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, 2(2):25.
- [Zinkevich, 2003] Zinkevich, M. (2003). Online convex programming and generalized infinitesimal gradient ascent. In *Proceedings of the 20th International Conference on Machine Learning (ICML-03)*, pages 928–936.

Appendix

A Discretization based algorithm

Definition 19 (r -discretization). *An r -discretization or r -net of a bounded set $S \subset \mathbb{R}^d$ is a finite set of points $\mathcal{D}$ such that the Euclidean distance of any point in S is at most r from some point in $\mathcal{D}$.*

Recall that $\mathcal{C} \subset \mathbb{R}^d$ is contained in a ball of radius R . A standard greedy construction gives an r -discretization of size at most $(3R/r)^d$ [Balcan et al., 2018a]. Given the dispersion parameter β , a natural choice is to use a $T^{-\beta}$ -discretization as in Algorithm 1.

Theorem 4. *Let $R^{\text{finite}}(T, s, N)$ denote the s -shifted regret for the finite experts problem on N experts, for the algorithm used in step 2 of Algorithm 1. Then Algorithm 1 enjoys s -shifted regret $R^{\mathcal{C}}(T, s)$ which satisfies*

$$R^{\mathcal{C}}(T, s) \leq R^{\text{finite}}\left(T, s, (3RT^\beta)^d\right) + (sH + L)O(T^{1-\beta}).$$

Proof of Theorem 4. We show we can round the optimal points in $\mathcal{C}$ to points in the $(T^{-\beta})$ -discretization $\mathcal{D}$ with a payoff loss at most $(sH + L)T^{1-\beta}$ in expectation. But in $\mathcal{D}$ we know a way to bound regret by $R^{\text{finite}}(T, s, N)$, where N , the number of points in $\mathcal{D}$, is at most $(\frac{3R}{T^{-\beta}})^d = (3RT^\beta)^d$.

Let $t_{0:s}$ denote the expert switching times in the optimal offline payoff, and ρ_i^* be the point picked by the optimal offline algorithm in $[t_{i-1}, t_i - 1]$. Consider a ball of radius $T^{-\beta}$ around ρ_i^* . It must have some point $\hat{\rho}_i^* \in \mathcal{D}$. We then must have that $\{u_t \mid t \in [t_{i-1}, t_i - 1]\}$ has at most $O(T^{-\beta}T) = O(T^{1-\beta})$ discontinuities due to β -dispersion, which implies

$$\sum_{t=t_{i-1}}^{t_i-1} u_t(\hat{\rho}_i^*) \geq \sum_{t=t_{i-1}}^{t_i-1} u_t(\rho_i^*) - O(T^{1-\beta})H - L(t_i - t_{i-1})T^{-\beta}$$

Let $\hat{\rho}_t = \hat{\rho}_i^*$ for each $t_{i-1} \leq t \leq t_i - 1$. Summing over i gives

$$\begin{aligned} \sum_{t=1}^T u_t(\hat{\rho}_t) &\geq OPT - O(T^{1-\beta})sH - LT^{1-\beta} \\ &= OPT - (sH + L)O(T^{1-\beta}) \end{aligned}$$

Now payoff of this algorithm is bounded above by the payoff of the optimal sequence of experts with s shifts

$$\sum_{t=1}^T u_t(\hat{\rho}_t) \leq OPT^{\text{finite}}$$

Let the finite experts algorithm with shifted regret bounded by $R^{\text{finite}}(T, s, N)$ choose ρ_t at round t . Then,

using the above inequalities,

$$\begin{aligned} \sum_{t=1}^T u_t(\rho_t) &\geq OPT^{\text{finite}} - R^{\text{finite}}(T, s, N) \\ &\geq OPT - (sH + L)O(T^{1-\beta}) - R^{\text{finite}}(T, s, N) \end{aligned}$$

We use this to bound the regret for the continuous case

$$\begin{aligned} R^{\mathcal{C}}(T, s) &= OPT - \sum_{t=1}^T u_t(\rho_t) \\ &\leq R^{\text{finite}}(T, s, N) + (sH + L)O(T^{1-\beta}) \end{aligned}$$

□

B Counterexamples

We will construct problem instances where some sub-optimal algorithms mentioned in the paper suffer high regret.

We first show that the Exponential Forecaster algorithm of [Balcan et al., 2018a] suffers linear s -shifted regret even for $s = 2$. This happens because pure exponential updates may accumulate high weights on well-performing experts and may take a while to adjust weights when these experts suddenly start performing poorly.

Lemma 20. *There exists an instance where Exponential Forecaster algorithm of [Balcan et al., 2018a] suffers linear s -shifted regret.*

Proof. Let $\mathcal{C} = [0, 1]$. Define utility functions

$$u^{(0)}(\rho) = \begin{cases} 1 & \text{if } \rho < \frac{1}{2} \\ 0 & \text{if } \rho \geq \frac{1}{2} \end{cases} \quad \text{and} \quad u^{(1)}(\rho) = \begin{cases} 0 & \text{if } \rho < \frac{1}{2} \\ 1 & \text{if } \rho \geq \frac{1}{2} \end{cases}$$

Now consider the instance where $u^{(0)}(\rho)$ is presented for the first $T/2$ rounds and $u^{(1)}(\rho)$ is presented for the remaining rounds. In the second half, with probability at least $\frac{1}{2}$, the Exponential Forecaster algorithm will select a point from $[0, \frac{1}{2}]$ and accumulate a regret of 1. Thus the expected 2-shifted regret of the algorithm is at least $\frac{T}{2} \cdot \frac{1}{2} = \Omega(T)$. Notice that the construction does not depend on the step size parameter λ . □

We further look at the performance of Random Restarts EF (Algorithm 4), an easy-to-implement algorithm which looks deceptively similar to Algorithm 2, against this adversary. Turns out Random Restarts EF may not restart frequently enough for the optimal value of the exploration parameter, and have sufficiently long chains of pure exponential updates in expectation to suffer high regret.

Theorem 21. *There exists an instance where Random Restarts EF (Algorithm 4) with parameters λ and α as in Theorem 6 suffers linear s -shifted regret.*

Proof. The probability of pure exponential updates from $t = T/4$ through $t = 3T/4$ is at least

$$(1 - \alpha)^{T/2} = \left(1 - \frac{1}{T-1}\right)^{T/2} > \frac{1}{2}$$

for $T > 5$. By Lemma 20, this implies at least $\frac{T}{8}$ regret in this case, and so the expected regret of the algorithm is at least $\frac{T}{16} = \Omega(T)$. $\square$

C Analysis of algorithms

In this section we will provide detailed proofs of lemmas and theorems from Section 4. We will restate them for easy reference.

Lemma 10. (Algorithm 2) *For each $t \in [T]$, $w_t(\rho) = \mathbb{E}[\hat{w}_t(\rho)]$ and $W_t = \mathbb{E}[\hat{W}_t]$, where the expectations are over random restarts $\mathbf{z}_t = \{z_1, \dots, z_{t-1}\}$.*

Proof of Lemma 10. $w_t(\rho) = \mathbb{E}[\hat{w}_t(\rho)]$ implies $W_t = \mathbb{E}[\hat{W}_t]$ by Fubini's theorem (recall $\mathcal{C}$ is closed and bounded). $w_t(\rho) = \mathbb{E}[\hat{w}_t(\rho)]$ follows by simple induction on t . In the base case, $\mathbf{z}_1$ is the empty set and $w_1(\rho) = 1 = \hat{w}_1(\rho) = \mathbb{E}[\hat{w}_1(\rho)]$. For $t > 1$,

$$\begin{aligned} \mathbb{E}[\hat{w}_t(\rho)] &= (1 - \alpha)\mathbb{E}[e^{\lambda u_t(\rho)} \hat{w}_{t-1}(\rho)] + \\ &\quad \frac{\alpha}{\text{VOL}(\mathcal{C})} \mathbb{E} \left[\int_{\mathcal{C}} e^{\lambda u_t(\rho)} \hat{w}_{t-1}(\rho) d\rho \right] \\ &\quad \text{(definition of } \hat{w}_t) \\ &= (1 - \alpha)e^{\lambda u_t(\rho)} \mathbb{E}[\hat{w}_{t-1}(\rho)] + \\ &\quad \frac{\alpha}{\text{VOL}(\mathcal{C})} \int_{\mathcal{C}} e^{\lambda u_t(\rho)} \mathbb{E}[\hat{w}_{t-1}(\rho)] d\rho \\ &\quad \text{(expectation is over } \mathbf{z}_t) \\ &= (1 - \alpha)e^{\lambda u_t(\rho)} w_{t-1}(\rho) \\ &\quad + \frac{\alpha}{\text{VOL}(\mathcal{C})} \int_{\mathcal{C}} e^{\lambda u_t(\rho)} w_{t-1}(\rho) d\rho \\ &\quad \text{(inductive hypothesis)} \\ &= w_t(\rho) \\ &\quad \text{(definition of } w_t) \end{aligned}$$

$\square$

Lemma 11. (Algorithm 2) W_{T+1} equals the sum

$$\sum_{s \in [T]} \sum_{t_0=1 < t_1 < \dots < t_s=T+1} \frac{\alpha^{s-1}(1-\alpha)^{T-s}}{\text{VOL}(\mathcal{C})^{s-1}} \prod_{i=1}^s \tilde{W}(t_{i-1}, t_i)$$

Proof. We use Lemma 10 to note that $W_{T+1} = \mathbb{E}_{\mathbf{z}_T}[\hat{W}_{T+1}] = \mathbb{E}_{s, \mathbf{t}_s}[\hat{W}_{T+1} \mid s, \mathbf{t}_s]$, where $\hat{W}_{T+1} \mid s, \mathbf{t}_s$ is the total weight of Algorithm 4 at time $T+1$ given restarts occur exactly at $\mathbf{t}_s$, and is deterministic since all weights $\hat{w}_{T+1}(\rho)$ are fixed given exact restart times. We will now show by an induction on s ,

$$\hat{w}_{T+1}(\rho) \mid s, \mathbf{t}_s = \tilde{w}(\rho; t_{s-1}, t_s) \prod_{i=1}^{s-1} \frac{\tilde{W}(t_{i-1}, t_i)}{\text{VOL}(\mathcal{C})} \quad (3)$$

In other words, we wish to show that $\hat{w}_{T+1}(\rho) \mid s, \mathbf{t}_s$ (weights of Algorithm 4 at time $T+1$ given restarts occur exactly at $\mathbf{t}_s$) can be expressed as the product of weight $\tilde{w}(\rho; t_{s-1}, t_s)$ at ρ of regular Exponential Forecaster since the last restart times the normalized total weights accumulated over previous runs.

For $s = 1$, we have no restarts and

$$\begin{aligned} \tilde{w}(\rho; t_{s-1}, t_s) \prod_{i=1}^{s-1} \frac{\tilde{W}(t_{i-1}, t_i)}{\text{VOL}(\mathcal{C})} &= \tilde{w}(\rho; t_0, t_1) \prod_{i=1}^0 \frac{\tilde{W}(t_{i-1}, t_i)}{\text{VOL}(\mathcal{C})} \\ &= \tilde{w}(\rho; 1, T+1) \\ &= \hat{w}_{T+1}(\rho) \mid 1, \mathbf{t}_1 \end{aligned}$$

For $s > 1$, the last restart occurs at $t_{s-1} > 1$. By inductive hypothesis for time $t_{s-1} - 1$ until which we've had $s - 2$ restarts,

$$\begin{aligned} \hat{w}_{t_{s-1}-1}(\rho) \mid s, \mathbf{t}_s &= \hat{w}_{t_{s-1}-1}(\rho) \mid s-1, \mathbf{t}_{s-1} \\ &= \tilde{w}(\rho; t_{s-2}, t_{s-1} - 1) \prod_{i=1}^{s-2} \frac{\tilde{W}(t_{i-1}, t_i)}{\text{VOL}(\mathcal{C})} \end{aligned}$$

Due to restart at t_{s-1} ,

$$\begin{aligned} \hat{w}_{t_{s-1}}(\rho) \mid s, \mathbf{t}_s &= \frac{\int_{\mathcal{C}} e^{\lambda u_t(\rho)} \hat{w}_{t_{s-1}-1}(\rho) d\rho}{\text{VOL}(\mathcal{C})} \\ &= \prod_{i=1}^{s-1} \frac{\tilde{W}(t_{i-1}, t_i)}{\text{VOL}(\mathcal{C})} \end{aligned}$$

It's regular exponential updates from this point to t_s , which gives (3).

Integrating (3) to get $\hat{W}_{T+1} \mid s, \mathbf{t}_s$, and noting probability of $s-1$ restarts at $\mathbf{t}_s$ is $\alpha^{s-1}(1-\alpha)^{T-s}$ completes the proof. $\square$

Theorem 6. *The s -shifted regret of Algorithm 2 with $\alpha = s/T$ and $\lambda = \sqrt{s(d \log(RT^\beta) + \log(T/s))}/T/H$ is $O(H \sqrt{sT(d \log(RT^\beta) + \log(T/s))} + (sH + L)T^{1-\beta})$.*

Full proof of Theorem 6. We first provide an upper and lower bound to $\frac{W_{T+1}}{W_1}$.

Upper bound: The proof is similar to the upper bound for exponential weighted forecaster in [Balcan et al., 2018a] and uses Lemma 8 for W_t .

$$\begin{aligned}\frac{W_{t+1}}{W_t} &= \frac{\int_{\mathcal{C}} e^{\lambda u_t(\rho)} w_t(\rho) d\rho}{W_t} \\ &= \int_{\mathcal{C}} e^{\lambda u_t(\rho)} \frac{w_t(\rho)}{W_t} d\rho \\ &= \int_{\mathcal{C}} e^{\lambda u_t(\rho)} p_t(\rho) d\rho\end{aligned}$$

Finally use inequalities $e^{\lambda z} \leq 1 + (e^\lambda - 1)z$ for $z \in [0, 1]$ and $1 + z \leq e^z$ to get

$$\begin{aligned}\frac{W_{t+1}}{W_t} &\leq \int_{\mathcal{C}} p_t(\rho) \left(1 + (e^{H\lambda} - 1) \frac{u_t(\rho)}{H} \right) d\rho \\ &= 1 + (e^{H\lambda} - 1) \frac{P_t}{H} \\ &\leq \exp \left((e^{H\lambda} - 1) \frac{P_t}{H} \right)\end{aligned}$$

where P_t denotes the expected payoff of the algorithm in round t . Let $P(\mathcal{A})$ be the expected total payoff. Then we can write $\frac{W_{T+1}}{W_1}$ as a telescoping product which gives

$$\begin{aligned}\frac{W_{T+1}}{W_1} &= \prod_{t=1}^T \frac{W_{t+1}}{W_t} \leq \exp \left((e^{H\lambda} - 1) \frac{\sum_t P_t}{H} \right) \\ &= \exp \left(\frac{P(\mathcal{A})(e^{H\lambda} - 1)}{H} \right)\end{aligned}\quad (4)$$

Lower bound: Again the proof is similar to [Balcan et al., 2018a] and the major difference is use of Lemma 11.

We first lower bound payoffs of points close to the optimal sequence of experts using dispersion. If the optimal sequence with s shifts has shifts at t_i^* ($1 \leq i \leq s-1$), by β -dispersion for any $\rho_i \in \mathcal{B}(\rho_i^*, w)$

$$\sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho_i) \geq \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho_i^*) - kH - L(t_i^* - t_{i-1}^*)w \quad (5)$$

where $w = T^{-\beta}$ and $k = O(T^{1-\beta})$. Summing both sides over $i \in [s-1]$ helps us relate the lower bound to the payoff OPT of the optimal sequence.

$$\begin{aligned}\sum_{i=1}^s \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho_i) &\geq \sum_{i=1}^s \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho_i^*) - kH - L(t_i^* - t_{i-1}^*)w \\ &= OPT - ksH - LTW\end{aligned}\quad (6)$$

Now to lower bound $\frac{W_{T+1}}{W_1}$, we first lower bound W_{T+1} . We use Lemma 11 and lower bound by picking the term corresponding to times of expert shifts in the optimal

sequence with s -shifted expert.

$$\begin{aligned}W_{T+1} &= \sum_{s \in [T]} \sum_{t_0=1 < t_1 < \dots < t_s=T+1} \frac{\alpha^{s-1}(1-\alpha)^{T-s}}{\text{VOL}(\mathcal{C})^{s-1}} \prod_{i=1}^s \tilde{W}(t_{i-1}, t_i) \\ &\geq \frac{\alpha^{s-1}(1-\alpha)^{T-s}}{\text{VOL}(\mathcal{C})^{s-1}} \prod_{i=1}^s \tilde{W}(t_{i-1}^*, t_i^*)\end{aligned}\quad (7)$$

The product of $\tilde{W}$'s can in turn be lower bounded by restricting attention to points close (i.e. within a ball of radius w centered at optimal expert ρ_i^*) to the optimal sequence. The payoffs of such points was lower-bounded in (5) and (6) in terms of the optimal payoff.

$$\begin{aligned}\prod_{i=1}^s \tilde{W}(t_{i-1}^*, t_i^*) &= \prod_{i=1}^s \int_{\mathcal{C}} \exp \left(\lambda \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho) \right) d\rho \\ &\geq \prod_{i=1}^s \int_{\mathcal{B}(\rho_i^*, w)} \exp \left(\lambda \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho) \right) d\rho \\ &\geq \prod_{i=1}^s \int_{\mathcal{B}(\rho_i^*, w)} \exp \left(\lambda \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho_i^*) - kH - L(t_i^* - t_{i-1}^*)w \right) d\rho \\ &= \text{VOL}(\mathcal{B}(w))^s \cdot \\ &\quad \exp \left(\sum_{i=1}^s \lambda \sum_{t=t_{i-1}^*}^{t_i^*-1} u_t(\rho_i^*) - kH - L(t_i^* - t_{i-1}^*)w \right) \\ &= \text{VOL}(\mathcal{B}(w))^s \exp \left(\lambda(OPT - ksH - LTW) \right)\end{aligned}$$

Plugging into equation (7) we get

$$\begin{aligned}W_{T+1} &\geq \frac{\alpha^{s-1}(1-\alpha)^{T-s} \text{VOL}(\mathcal{B}(w))^s}{\text{VOL}(\mathcal{C})^{s-1}} \\ &\quad \exp \left(\lambda(OPT - ksH - LTW) \right)\end{aligned}$$

Also, $W_1 = \int_{\mathcal{C}} w_1(\rho) d\rho = \text{VOL}(\mathcal{C})$. Thus, using the fact that ratio of volume of balls $\mathcal{B}(w)$ and $\mathcal{B}(R)$ in d -dimensions is $(w/R)^d$, and assuming $\mathcal{C}$ is bounded by some ball $\mathcal{B}(R)$.

$$\begin{aligned}\frac{W_{T+1}}{W_1} &\geq \alpha^{s-1}(1-\alpha)^{T-s} \left(\frac{w}{R} \right)^{sd} \\ &\quad \exp \left(\lambda(OPT - ksH - LTW) \right)\end{aligned}\quad (8)$$

Putting together: Combining upper and lower bounds from (4) and (8) respectively,

$$\begin{aligned}\log(\alpha^{s-1}(1-\alpha)^{T-s}) - sd \log \frac{R}{w} + \\ \lambda(OPT - ksH - LTW) \leq \frac{P(\mathcal{A})(e^{H\lambda} - 1)}{H}\end{aligned}$$

which rearranges to

$$\begin{aligned} OPT - P(\mathcal{A}) \leq & P(\mathcal{A}) \frac{(e^{H\lambda} - 1 - H\lambda)}{H\lambda} + \frac{sd \log(R/w)}{\lambda} \\ & + ksH + LTW - \frac{\log(\alpha^{s-1}(1-\alpha)^{T-s})}{\lambda} \end{aligned}$$

Using $P(\mathcal{A}) \leq HT$ and using $e^z \leq 1 + z + (e-2)z^2$ for $z \in [0, 1]$ we have

$$\begin{aligned} OPT - P(\mathcal{A}) \leq & HT \frac{(e^{H\lambda} - 1 - H\lambda)}{H\lambda} + \frac{sd \log(R/w)}{\lambda} \\ & + ksH + LTW - \frac{\log(\alpha^{s-1}(1-\alpha)^{T-s})}{\lambda} \\ < & H^2T\lambda + \frac{sd \log(R/w)}{\lambda} + ksH + LTW \\ & - \frac{\log(\alpha^{s-1}(1-\alpha)^{T-s})}{\lambda} \end{aligned}$$

Now we tighten the bound, first w.r.t. α then w.r.t. λ . Note $\min_{\alpha} -\log(\alpha^{s-1}(1-\alpha)^{T-s})$ occurs for $\alpha_0 = \frac{s-1}{T-1}$ and

$$\begin{aligned} & -\log(\alpha_0^{s-1}(1-\alpha_0)^{T-s}) \\ & = (T-1) \left[-\frac{s-1}{T-1} \log \frac{s-1}{T-1} - \frac{T-s}{T-1} \log \frac{T-s}{T-1} \right] \\ & \leq (s-1) \log e \frac{T-1}{s-1} \end{aligned}$$

(binary entropy function satisfies $h(x) \leq x \ln(e/x)$ for $x \in [0, 1]$). Finally minimizing over λ gives

$$OPT - P(\mathcal{A}) \leq O(H\sqrt{sT(d \log(R/w) + \log(T/s))} + ksH + LTW)$$

for $\lambda = \sqrt{s(d \log(R/w) + \log(T/s))/T}/H$. Plugging back $w = T^{-\beta}$ and $k = O(T^{1-\beta})$ completes the proof. $\square$

The rest of this section is concerned with the analysis of Algorithm 3 for the sparse experts setting.

Lemma 22. For any $t < T$,

$$w_T(\rho) \geq \alpha(1-\alpha)^{T-t} \pi_t(\rho) \tilde{w}(\rho; t, T) W_t$$

Proof. Follows using the restart algorithm technique used in Lemmas 11 and 13. Consider the probability of last ‘restart’ being at time t . Notice this also implies Corollary 14. $\square$

Lemma 23. Let $\pi_t(\rho) = \sum_{i=1}^t \beta_{i,t} p_i(\rho)$ in Algorithm 3. Then

$$\pi_t(\rho) = \frac{\alpha_{1,t}}{W_1} + \sum_{i=1}^{t-1} \alpha_{i+1,t} \frac{e^{\lambda u_i(\rho)} w_i(\rho)}{W_{i+1}}$$

where

$$\alpha_{i,t} \geq \frac{1-\alpha}{e_t} \left(e^{-\gamma} + \frac{\alpha}{e_t} \right)^{t-i}$$

and $e_t := \sum_{i=1}^t e^{-\gamma(i-1)}$.

Proof. Notice, by definition of weight update in Algorithm 3,

$$\begin{aligned} p_t(\rho) &= (1-\alpha) \frac{e^{\lambda u_{t-1}(\rho)} w_{t-1}(\rho)}{W_t} + \alpha \sum_{i=1}^{t-1} \beta_{i,t-1} p_i(\rho) \\ &= (1-\alpha) \frac{e^{\lambda u_{t-1}(\rho)} w_{t-1}(\rho)}{W_t} + \alpha \pi_{t-1}(\rho) \end{aligned}$$

This gives us a recursive relation for $\alpha_{i,t}$.

$$\alpha_{i,t} = \begin{cases} \beta_{i,t}(1-\alpha) + \alpha \sum_{j=i+1}^t \beta_{j,t} \alpha_{i,j-1} & \text{if } i > 1 \\ \beta_{i,t} + \alpha \sum_{j=i+1}^t \beta_{j,t} \alpha_{i,j-1} & \text{if } i = 1 \end{cases}$$

Thus for each $1 \leq i \leq t$

$$\alpha_{i,t} \geq \beta_{i,t}(1-\alpha) + \alpha \sum_{j=i+1}^t \beta_{j,t} \alpha_{i,j-1}$$

We proceed by induction on $t-i$. For $i = t$,

$$\alpha_{t,t} \geq \beta_{t,t}(1-\alpha) = \frac{1-\alpha}{e_t} \left(e^{-\gamma} + \frac{\alpha}{e_t} \right)^{t-t}$$

For $i < t$, by inductive hypothesis

$$\begin{aligned} \alpha_{i,t} &\geq \beta_{i,t}(1-\alpha) + \sum_{j=i+1}^t \beta_{j,t} \alpha_{i,j-1} \\ &\geq (1-\alpha) \frac{e^{-\gamma(t-i)}}{e_t} \\ &\quad + \alpha \frac{1-\alpha}{e_t} \sum_{j=i+1}^t \frac{e^{-\gamma(t-j)}}{e_t} \left(e^{-\gamma} + \frac{\alpha}{e_t} \right)^{j-1-i} \\ &= \frac{1-\alpha}{e_t} \left(e^{-\gamma} + \frac{\alpha}{e_t} \right)^{t-i} \end{aligned}$$

which completes the induction step. $\square$

Corollary 24. Let $w_t(\rho), W_t$ be as in Algorithm 3 and π_t as in Lemma 13. For each $\tau < \tau' < t$ and any bounded f defined on $\mathcal{C}$.

$$\begin{aligned} \int_{\mathcal{C}} \pi_t(\rho) f(\rho) d\rho &\geq \frac{\alpha(1-\alpha)^{\tau'-\tau}(1-e^{-\gamma})}{(e^{-\gamma} + \alpha(1-e^{-\gamma}))^{\tau'-t}} \frac{W_{\tau}}{W_{\tau'}} \\ &\quad \int_{\mathcal{C}} \pi_{\tau}(\rho) \tilde{w}(\rho; \tau, \tau') f(\rho) d\rho \end{aligned}$$

Proof. By Lemma 23,

$$\begin{aligned}
 \int_{\mathcal{C}} \pi_t(\rho) f(\rho) d\rho &= \int_{\mathcal{C}} \pi_t(\rho) f(\rho) d\rho \\
 &\geq \int_{\mathcal{C}} \alpha_{\tau',t} \frac{e^{\lambda u_{\tau'-1}(\rho)} w_{\tau'-1}(\rho)}{W_{\tau'}} f(\rho) d\rho \\
 &\geq \frac{1-\alpha}{e_t} \left(e^{-\gamma} + \frac{\alpha}{e_t} \right)^{t-\tau'} \frac{1}{W_{\tau'}} \cdot \\
 &\quad \int_{\mathcal{C}} e^{\lambda u_{\tau'-1}(\rho)} w_{\tau'-1}(\rho) f(\rho) d\rho \\
 &\geq \frac{1-\alpha}{e_t} \left(e^{-\gamma} + \frac{\alpha}{e_t} \right)^{t-\tau'} \frac{\alpha(1-\alpha)^{\tau'-1-\tau} W_{\tau}}{W_{\tau'}} \cdot \\
 &\quad \int_{\mathcal{C}} \pi_{\tau}(\rho) \tilde{w}(\rho; \tau, \tau') f(\rho) d\rho
 \end{aligned}$$

where for the last inequality we have used Lemma 22. The lemma then follows by noting

$$\frac{1}{e_t} = \frac{1 - e^{-\gamma}}{1 - e^{-\gamma t}} \geq 1 - e^{-\gamma}$$

where $e_t = \sum_{i=1}^t e^{-\gamma(i-1)}$ as defined in Lemma 23. $\square$

Theorem 7. *The (m -sparse, s -shifted) regret of Algorithm 3 is $O(H\sqrt{T(md\log(RT^\beta) + s\log(mT/s))} + (mH + L)T^{1-\beta})$ for $\alpha = s/T$, $\gamma = s/mT$ and $\lambda = \sqrt{(md\log(RT^\beta) + s\log(T/s))/T/H}$.*

Proof of Theorem 7. Like Theorem 6 we first provide an upper and lower bound to $\frac{W_{T+1}}{W_1}$. The upper bound proof is identical to that of Theorem 6 by replacing Lemma 8 by Lemma 12.

For the lower bound we use Corollaries 14 and 24. Applying corollary 24 repeatedly to collect exponential updates for the times OPT played the same expert lets us use the arguments for Theorem 6 to get Equation ?? . Indeed if $\{(s_i, f_i) \mid 1 \leq i \leq l\}$ are the start and finish times of a particular expert ρ in the OPT sequence, we can use Corollary 14 to write

$$W_{f_{l+1}} \geq \alpha(1-\alpha)^{f_{l+1}-s_l} W_{s_l} \tilde{W}(\pi_{s_l}; s_l, f_l + 1)$$

Applying Corollary 24 repeatedly now gets us

$$\begin{aligned}
 W_{f_{l+1}} &\geq \frac{\alpha^l (1-\alpha)^{\sum_{j=1}^l f_j + 1 - s_j} (1 - e^{-\gamma})^{l-1}}{(e^{-\gamma} + \alpha(1 - e^{-\gamma}))^{\sum_{j=1}^{l-1} f_j + 1 - s_{j+1}}} \cdot \\
 &\quad \frac{\prod_{i=1}^l W_{s_i}}{\prod_{i=1}^{l-1} W_{f_{i+1}}} \int_{\mathcal{C}} \left(\pi_{s_1}(\rho) \prod_{j=1}^l \tilde{w}(\rho; s_j, f_j + 1) \right) d\rho
 \end{aligned}$$

or

$$\begin{aligned}
 \frac{\prod_{i=1}^l W_{f_{i+1}}}{\prod_{i=1}^l W_{s_i}} &\geq \frac{\alpha^l (1-\alpha)^{\sum_{j=1}^l f_j + 1 - s_j} (1 - e^{-\gamma})^{l-1}}{(e^{-\gamma} + \alpha(1 - e^{-\gamma}))^{\sum_{j=1}^{l-1} f_j + 1 - s_{j+1}}} \cdot \\
 &\quad \int_{\mathcal{C}} \left(\pi_{s_1}(\rho) \prod_{j=1}^l \tilde{w}(\rho; s_j, f_j + 1) \right) d\rho
 \end{aligned}$$

Multiplying these inequalities for each of m experts in the optimal sequence gives us $\frac{W_{T+1}}{W_1}$ on the left side. Also note

$$\int_{\mathcal{C}} \pi_t(\rho) f(\rho) d\rho \geq \frac{\alpha_{1,t}}{W_1} \int_{\mathcal{C}} f(\rho) d\rho$$

and, using dispersion as in proof of Theorem 6,

$$\begin{aligned}
 \prod_{\text{experts in OPT}} \int_{\mathcal{C}} \left(\prod_{j=1}^l \tilde{w}(\rho; s_j, f_j + 1) \right) d\rho &\geq \\
 \text{VOL}(\mathcal{B}(T^{-\beta}))^m \exp(\lambda(OPT - (mH + L)O(T^{1-\beta}))) &
 \end{aligned}$$

Putting it all together gives Equation ?? . Combining the lower and upper bounds on $\frac{W_{T+1}}{W_1}$ gives us a bound on $OPT - P(\mathcal{A})$.

$$\begin{aligned}
 OPT - P(\mathcal{A}) &< H^2 T \lambda + \frac{md \log(RT^\beta)}{\lambda} \\
 &\quad + (mH + L)O(T^{1-\beta}) \\
 &\quad - \log \left(\frac{\alpha^s (1-\alpha)^T (1 - e^{-\gamma})^s}{(e^{-\gamma} + \alpha(1 - e^{-\gamma}))^{-mT}} \right) / \lambda
 \end{aligned}$$

We now chose parameters γ, α, λ to get the tightest regret bound. Note that $-\log(\alpha^s (1-\alpha)^T)$ is minimized for $\alpha = \frac{s}{T+s} = \Theta(\frac{s}{T})$ and $-\log((1 - e^{-\gamma})^s (e^{-\gamma} + \alpha(1 - e^{-\gamma}))^{mT})$ is minimized for $\gamma = \log \left(\frac{1+s/mT}{1-s\alpha/mT(1-\alpha)} \right) = \Theta(\frac{s}{mT})$. The corresponding minimum values can be bounded as

$$\begin{aligned}
 -\log(\alpha^s (1-\alpha)^T) &= s \log \frac{T+s}{s} + T \log \left(1 + \frac{s}{T} \right) \\
 &\leq s \log \frac{T+s}{s} + s \\
 &= O \left(s \log \frac{T}{s} \right)
 \end{aligned}$$

using $\log(1+x) \leq x$, and substituting $e^{-\gamma} = \frac{1-s\alpha/mT(1-\alpha)}{1+s/mT}$

$$\begin{aligned}
 &-\log((1 - e^{-\gamma})^s (e^{-\gamma} + \alpha(1 - e^{-\gamma}))^{mT}) \\
 &= -s \log \frac{\frac{s}{mT} \cdot \frac{1}{(1-\alpha)}}{1 + \frac{s}{mT}} - mT \log \frac{1}{1 + \frac{s}{mT}} \\
 &\leq s \log \left((1-\alpha) \left(\frac{mT}{s} + 1 \right) \right) + 1 \\
 &= O \left(s \log \frac{mT}{s} \right)
 \end{aligned}$$

Finally we minimize w.r.t. λ , to obtain the desired regret bound. $\square$

D Adaptive Regret

It is known that the fixed share algorithm obtains good adaptive regret for finite experts and OCO

[Adamskiy et al., 2012]. We show that it is the case here as well.

Definition 25. The τ -adaptive regret (due to [Hazan and Seshadhri, 2007]) is given by

$$\mathbb{E} \left[\max_{\substack{\rho^* \in \mathcal{C}, \\ 1 \leq r < s \leq T, s-r \leq \tau}} \sum_{t=r}^s (u_t(\rho^*) - u_t(\rho_t)) \right]$$

The goal here is to ensure small regret on all intervals of size up to τ simultaneously. Adaptive regret measures how well the algorithm approximates the best expert locally, and it is therefore somewhere between the static regret (measured on all outcomes) and the shifted regret, where the algorithm is compared to a good sequence of experts.

Theorem 26. Algorithm 2 enjoys $O(H\sqrt{\tau(d\log(R/w) + \log \tau)} + (H + L)\tau^{1-\beta})$ τ -adaptive regret for $\lambda = \sqrt{(d\log(R\tau^\beta) + \log(\tau))/\tau}/H$ and $\alpha = 1/\tau$.

Proof sketch of Theorem 26. Apply arguments of Theorem 6 to upper and lower bound W_{s+1}/W_r for any interval $[r, s] \subseteq [1, T]$ of size τ . We get

$$\frac{W_{s+1}}{W_r} \leq \exp \left(\frac{P(\mathcal{A})(e^{H\lambda} - 1)}{H} \right)$$

where $P(\mathcal{A})$ is the expected payoff of the algorithm in $[r, s]$. Also, by Corollary 14 (equivalent for Algorithm 2)

$$\begin{aligned} W_{s+1} &\geq \frac{\alpha(1-\alpha)^{s+1-r}}{\text{VOL}(\mathcal{C})} \tilde{W}(r, s) W_r \\ &= \frac{\alpha(1-\alpha)^\tau}{\text{VOL}(\mathcal{C})} \tilde{W}(r, s) W_r \end{aligned}$$

By dispersion, as in the proof of Theorem 6,

$$\tilde{W}(r, s) \geq \text{VOL}(\mathcal{B}(\tau^{-\beta})) \exp(\lambda(OPT - (H+L)O(\tau^{1-\beta})))$$

Putting the upper and lower bounds together gives us a bound on $OPT - P(\mathcal{A})$, which gives the desired regret bound for $\alpha = \frac{1}{\tau}$. $\square$

E Efficient Sampling

In Section 5 we introduced Algorithm 5 for efficient implementation of Algorithm 2 in $\mathbb{R}^d$. We present proofs of the results in that section, and an exact algorithm for the case $d = 1$.

Lemma 15. In Algorithm 2, for $t \geq 1$,

$$\begin{aligned} W_{t+1} &= (1-\alpha)^{t-1} \tilde{W}(1, t+1) + \\ &\quad \frac{\alpha}{\text{VOL}(\mathcal{C})} \sum_{i=2}^t \left[(1-\alpha)^{t-i} W_i \tilde{W}(i, t+1) \right] \end{aligned}$$

Algorithm 6 Fixed Share Exponential Forecaster - exact algorithm for one dimension

Input: $\lambda \in (0, 1/H]$

1. $W_1 = \text{VOL}(\mathcal{C})$

2. For each $t = 1, 2, \dots, T$:

Estimate $C_{t,j}$ using Lemma 16 for each $1 \leq j \leq t$ using memoized values for weights.

Sample i with probability $C_{t,i}$.

Sample ρ with probability proportional to $\tilde{w}(\rho; i, t)$.

Estimate W_{t+1} using Lemma 15.

Proof of Lemma 15. For $t = 1$, first term is $\tilde{W}(1, 2) = \int_{\mathcal{C}} e^{\lambda u_1(\rho)} d\rho = W_2$ and second term is zero. Also, by Lemma 8, for $t > 1$

$$\begin{aligned} W_{t+1} &= \int_{\mathcal{C}} e^{\lambda u_t(\rho)} w_t(\rho) d\rho \\ &= \int_{\mathcal{C}} e^{\lambda u_t(\rho)} \left[(1-\alpha) e^{\lambda u_{t-1}(\rho)} w_{t-1}(\rho) \right. \\ &\quad \left. + \frac{\alpha}{\text{VOL}(\mathcal{C})} \int_{\mathcal{C}} e^{\lambda u_{t-1}(\rho)} w_{t-1}(\rho) d\rho \right] d\rho \\ &= (1-\alpha) \int_{\mathcal{C}} e^{\lambda(u_t(\rho) + u_{t-1}(\rho))} w_{t-1}(\rho) d\rho \\ &\quad + \frac{\alpha}{\text{VOL}(\mathcal{C})} W_t \int_{\mathcal{C}} e^{\lambda u_t(\rho)} d\rho \end{aligned}$$

Continue substituting $w_j(\rho) = (1-\alpha) e^{\lambda u_j(\rho)} w_{j-1}(\rho) + \frac{\alpha}{\text{VOL}(\mathcal{C})} \int_{\mathcal{C}} e^{\lambda u_j(\rho)} w_{j-1}(\rho) d\rho$ in the first summand until $w_1 = 1$ to get the desired expression. $\square$

Definition 27. For $\alpha \geq 0$ we say $\hat{A}$ is an (α, ζ) -approximation of A if

$$\Pr(e^{-\alpha} A \leq \hat{A} \leq e^{\alpha} A) \geq 1 - \zeta$$

Lemma 28. If $\hat{A}$ is an (α, ζ) -approximation of A and $\hat{B}$ is a (β, ζ') -approximation of B , such that $A, B, \hat{A}, \hat{B}$ are all positive reals

1. $\hat{A}\hat{B}$ is an $(\alpha + \beta, \zeta + \zeta')$ -approximation of AB
2. $p\hat{A} + q\hat{B}$ is a $(\max\{\alpha, \beta\}, \zeta + \zeta')$ -approximation of $pA + qB$ for $p, q \geq 0$

Proof. The results follow from union bound on failure probabilities. $\square$

Corollary 29. For one-dimensional case, we can exactly compute $\tilde{W}(i, j)$, $1 \leq i < j \leq t$, hence W_t at each iteration can be computed in $O(t)$ time using Lemma 15. More generally, if we have a (β, ζ) approximation

for each $\tilde{W}(i, j)$, $1 \leq i < j \leq t$, then by Lemma 15 we can compute a $(t\beta, t^2\zeta)$ -approximation for W_{t+1} .

Proof. Union bound on failure probabilities of all $\tilde{W}(i, j)$, $1 \leq i < j \leq t$ gives we have a β approximation for each with probability at least $1 - t^2\zeta$. This covers failure for all terms in W_i , $2 \leq i \leq t$. Further, by induction, the error for estimates for W_i is at most $(i-1)\beta$. By Lemma 28, the error for W_{t+1} estimates is at most $t\beta$. $\square$

Lemma 16. In Algorithm 2, for $t \geq 1$, $p_t(\rho) = \sum_{i=1}^t C_{t,i} \frac{\tilde{w}(\rho; i, t)}{\tilde{W}(i, t)}$. The coefficients $C_{t,i}$ are given by

$$C_{t,i} = \begin{cases} 1 & i = t = 1 \\ \alpha & i = t > 1 \\ (1 - \alpha) \frac{W_{t-1}}{W_t} \frac{\tilde{W}(i, t)}{\tilde{W}(i, t-1)} C_{t-1,i} & i < t \end{cases}$$

and $(C_{t,1}, \dots, C_{t,t})$ lies on the probability simplex Δ^{t-1} .

Proof of Lemma 16. At each iteration, p_t is obtained by mixing $e^{u_t} p_{t-1}$ with the uniform distribution, i.e. we rescale distributions that p_{t-1} was a mixture of and add one more. Another way to view it is to consider a distribution over the sequences of exponentially updated or randomly chosen points. The final probability distribution is the mixture of a combinatorial number of distributions but a large number of them have a proportional density. $C_{t,i}$ are simply sums of mixture coefficients. This establishes the intuition for the expression for p_t and that the mixing coefficients should sum to 1, but we still need to convince ourselves that the coefficients can be computed efficiently.

We proceed by induction on t . For $t = 1$ (using definitions for $w_2(\rho)$ and $w_2(\rho)$)

$$p_1(\rho) = \frac{w_1(\rho)}{W_1} = \frac{1}{\text{VOL}(\mathcal{C})} = C_{1,1} \frac{\tilde{w}(\rho; 1, 1)}{\tilde{W}(1, 1)}$$

(recall $\tilde{w}(\rho; 1, 1) := 1$ and $\tilde{W}(1, 1) = \int_{\mathcal{C}} \tilde{w}(\rho; 1, 1) d\rho$). For the inductive step, we first express p_{t+1} in terms of p_t

$$\begin{aligned} p_{t+1}(\rho) &= \frac{w_{t+1}(\rho)}{W_{t+1}} \\ &= (1 - \alpha) \frac{e^{\lambda u_t(\rho)} w_t(\rho)}{W_{t+1}} + \frac{\alpha}{\text{VOL}(\mathcal{C})} \\ &= (1 - \alpha) \frac{W_t}{W_{t+1}} \frac{e^{\lambda u_t(\rho)} w_t(\rho)}{W_t} + \frac{\alpha}{\text{VOL}(\mathcal{C})} \\ &= (1 - \alpha) \frac{W_t}{W_{t+1}} e^{\lambda u_t(\rho)} p_t(\rho) + \frac{\alpha}{\text{VOL}(\mathcal{C})} \end{aligned}$$

The lemma is now straightforward to see with induction hypothesis.

$p_{t+1}(\rho)$

$$\begin{aligned} &= (1 - \alpha) \frac{W_t}{W_{t+1}} e^{\lambda u_t(\rho)} \left[\sum_{i=1}^t C_{t,i} \frac{\tilde{w}(\rho; i, t)}{\tilde{W}(i, t)} \right] + \frac{\alpha}{\text{VOL}(\mathcal{C})} \\ &= \sum_{i=1}^t \left[(1 - \alpha) \frac{W_t}{W_{t+1}} C_{t,i} \frac{\tilde{w}(\rho; i, t+1)}{\tilde{W}(i, t)} \right] + \frac{\alpha}{\text{VOL}(\mathcal{C})} \\ &= \sum_{i=1}^t \left[\left((1 - \alpha) \frac{W_t}{W_{t+1}} \frac{\tilde{W}(i, t+1)}{\tilde{W}(i, t)} C_{t,i} \right) \frac{\tilde{w}(\rho; i, t+1)}{\tilde{W}(i, t+1)} \right] \\ &\quad + \frac{C_{t+1,t+1}}{\text{VOL}(\mathcal{C})} \\ &= \sum_{i=1}^t C_{t+1,i} \frac{\tilde{w}(\rho; i, t+1)}{\tilde{W}(i, t+1)} + \frac{C_{t+1,t+1}}{\text{VOL}(\mathcal{C})} \end{aligned}$$

Finally noting

$$C_{t+1,t+1} \frac{\tilde{w}(\rho; t+1, t+1)}{\tilde{W}(t+1, t+1)} = C_{t+1,t+1} \frac{1}{\int_{\mathcal{C}} (1) d\rho} = \frac{C_{t+1,t+1}}{\text{VOL}(\mathcal{C})}$$

completes the proof.

Thus W_t (by Lemma 15) and $C_{t,i}$ can be computed recursively for logconcave utility functions using integration algorithm from [Lovász and Vempala, 2006]. We can compute them efficiently using Dynamic Programming.

Finally it's straightforward to establish that the coefficients for p_t must lie on the probability simplex Δ^{t-1} . All coefficients are positive, which is easily seen from the recursive relation and noting all weights are positive. Also we know

$$p_t(\rho) = \sum_{i=1}^t C_{t,i} \frac{\tilde{w}(\rho; i, t)}{\tilde{W}(i, t)}$$

Since $p_t(\rho)$ is a probability distribution by definition, integrating both sides over $\mathcal{C}$ gives

$$\begin{aligned} \int_{\mathcal{C}} p_t(\rho) d\rho &= \sum_{i=1}^t C_{t,i} \frac{\int_{\mathcal{C}} \tilde{w}(\rho; i, t) d\rho}{\tilde{W}(i, t)} \quad \text{or,} \\ 1 &= \sum_{i=1}^t C_{t,i} \end{aligned}$$

$\square$

Corollary 30. If we have a (β, ζ) approximation for each $\tilde{W}(i, j)$, $1 \leq i < j \leq t$, then by Corollary 29 and Lemma 16 we can compute $\hat{C}_{t+1,i}$ which are $(2t\beta, t^2\zeta)$ -approximation for each $C_{t+1,i}$.

Proof. For $i = t$, we know $C_{t,i}$ exactly by Lemma 16. For $i < t$,

$$C_{t,i} = (1 - \alpha)^{t-i} \frac{W_i}{W_t} \frac{\tilde{W}(i,t)}{\text{VOL}(\mathcal{C})} C_{i,i} \quad (9)$$

In Corollary 29, we show how to compute $((i-1)\beta, (i-1)^2\zeta)$ -approximation for W_i and $((t-1)\beta, (t-1)^2\zeta)$ -approximation for W_t given (β, ζ) approximations for each $\tilde{W}(i, j)$, $1 \leq i < j \leq t$. A similar argument using Lemma 28 shows with failure probability at most $t^2\zeta$, plugging in the approximations in equation 9 has at most $(t+i)\beta$ error. $\square$

Theorem 17. *If utility functions are piecewise concave and L -Lipschitz, we can approximately sample a point ρ with probability $p_{t+1}(\rho)$ in time $\tilde{O}(Kd^4T^4)$ for approximation parameters $\eta = \zeta = 1/\sqrt{T}$ and $\lambda = \sqrt{s(d \ln(RT^\beta) + \ln(T/s))/T}/H$ and enjoy the same regret bound as the exact algorithm. (K is number of discontinuities in u_t 's).*

Proof of Theorem 17. Based on Lemma 16, we can sample a uniformly random number r in $[0, 1]$ and then sample a ρ from one of t distributions (selected based on r) that $p_t(\rho)$ is a mixture of with probability proportional to $C_{t,i}$. The sampling from the exponentials can be done in polynomial time for concave utility functions using sampling algorithm of [Bassily et al., 2014]. At each round we sample from exactly one of t distributions in the sum for p_t in Lemma 16. We compute $(\eta/6T, \zeta/2T^2)$ approximations for $\tilde{W}(i, j)$, $1 \leq i < j \leq T$ in time $O(T^2 K T_f)$ where T_f is the time to integrate a logconcave distribution (at most $\tilde{O}(d^4/\epsilon^2)$ from [Lovász and Vempala, 2006]). These give $(\eta/3, \zeta/2)$ -approximation for $C_{t,i}$'s by corollary 30. Finally we run Algorithm 2 from [Balcan et al., 2018a] with approximation-confidence parameters $(\eta/3, \zeta/2)$. With probability at least $1 - \zeta$, $C_{t,i}$ estimation and ρ sampling according to $\tilde{w}(\rho; i, t)$ succeeds. If $\hat{\mu}$ denotes output distribution of ρ with approximate sampling, and μ denotes the exact distribution per $p_t(\rho)$, then we show $D_\infty(\hat{\mu}, \mu) \leq \eta$. Indeed, for any set of outcomes $E \subset \mathcal{C}$

$$\begin{aligned} \hat{\mu}(E) &= \Pr(\hat{\rho} \in E) = \sum_{i=1}^t \Pr(\hat{\rho} \in E \mid E_{i,t}) \Pr(E_{i,t}) \\ &= \sum_{i=1}^t \hat{\mu}_i(E) \frac{\hat{C}_{t,i}}{\sum_j \hat{C}_{t,j}} \end{aligned}$$

where $E_{i,t}$ denotes the event that $\tilde{w}(\rho; i, t)$ was used for sampling $p_t(\rho)$, and $\hat{\mu}_i$ corresponds to the distribution for approximate sampling of $\tilde{w}(\rho; i, t)$. Noting that we used $\eta/3$ approximation for $\hat{\mu}_i$ and each $\hat{C}_{t,i}$, we have

$$\hat{\mu}(E) \leq \sum_{i=1}^t e^{\eta/3} \mu_i(E) e^{2\eta/3} \frac{C_{t,i}}{\sum_j C_{t,j}} = e^\eta \mu(E)$$

Similarly, $\hat{\mu}(E) \geq e^{-\eta} \mu(E)$ and hence $D_\infty(\hat{\mu}, \mu) \leq \eta$. Finally we can show (cf. Theorem 12 of [Balcan et al., 2018a]) that with probability at least $1 - \zeta$ the expected utility per round of the approximate sampler is at most a $(1 - \eta)$ factor smaller than the expected utility per round of the exact sampler. Together with failure probability of ζ , this implies at most $(\eta + \zeta)HT$ additional regret which results in same asymptotic regret as the exact algorithm for $\eta = \zeta = 1/\sqrt{T}$.

To compute the time complexity, we note from [Lovász and Vempala, 2006] that logconcave functions can be integrated in $\tilde{O}(d^4/\epsilon^2)$ and sampled from in $\tilde{O}(d^3)$ time. The time to integrate dominates the complexity, and the overall complexity can be upper bounded by $O(T^2 K \cdot d^4/(\eta/T)^2) = O(KT^4 d^4)$. Note: The approximate integration and sampling are only needed for multi-dimensional case, for the one-dimensional case we can compute the weights and sample exactly in polynomial time. $\square$

F Lower bounds

We start with a simple lower bound argument for s -shifted regret for prediction with two experts based on a well-known $\Omega(\sqrt{T})$ lower bound argument for static regret. We will then extend it to the continuous setting and use it for the $\Omega(\sqrt{sT})$ part of the lower bound in Theorem 18 in Section 6.

Lemma 31. *For prediction with two experts, there exists a stochastic sequence of losses for which the s -shifted regret of any online learning algorithm satisfies*

$$\mathbb{E}[R_T] \geq \sqrt{sT/8}$$

Proof. Let the two experts predict 0 and 1 respectively at each time $t \in [T]$. The utility at each time t is computed by flipping a coin - with probability $1/2$ we have $u(0) = 1, u(1) = 0$ and with probability $1/2$ it's $u(0) = 0, u(1) = 1$. Expected payoff for any algorithm $\mathcal{A}$ is

$$P(\mathcal{A}, T) = \mathbb{E}\left[\sum_{t=1}^T u_t(\rho_t)\right] = \sum_{t=1}^T \mathbb{E}[u_t(\rho_t)] = \frac{T}{2}$$

since expected payoff is $1/2$ at each t no matter which expert is picked.

To compute shifted regret we need to compare this payoff with the best sequence of experts with $s - 1$ switches. We compare with a weaker adversary $\mathcal{A}'$ which is only allowed to switch up to $s - 1$ times, and

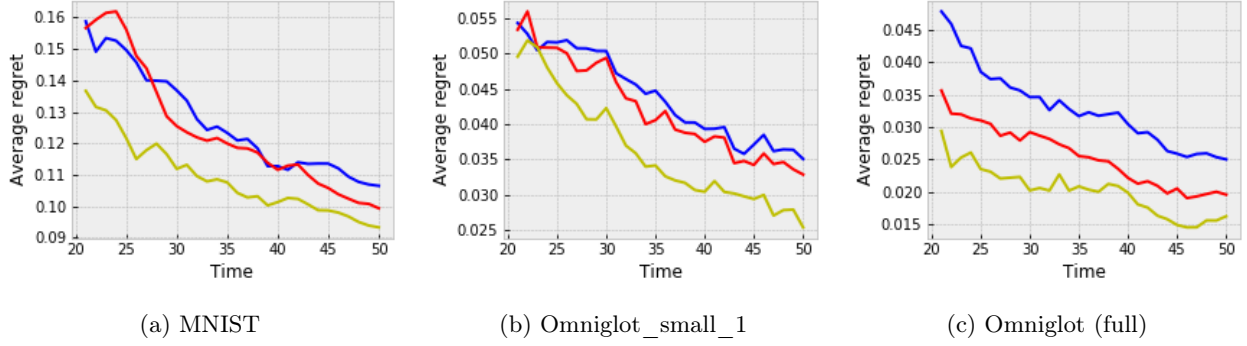


Figure 3: Average k -shifted regret vs game duration T for online clustering against k -shifted distributions. Color scheme: **Exponential Forecaster**, **Fixed Share EF**, **Generalized Share EF**

switches at only a subset of fixed times $t_i = iT/s$ to lower bound the regret.

$$\begin{aligned}
 \mathbb{E}[R_T] &= OPT - P(\mathcal{A}, T) \\
 &\geq P(\mathcal{A}', T) - P(\mathcal{A}, T) \\
 &= \sum_{t=1}^T \mathbb{E}[u_t(\rho'_t)] - \sum_{t=1}^T \mathbb{E}[u_t(\rho_t)] \\
 &= \sum_{i=0}^{s-1} \sum_{t=t_i+1}^{t_{i+1}} \mathbb{E}[u_t(\rho'_t)] - \mathbb{E}[u_t(\rho_t)]
 \end{aligned}$$

Now let $P_{i,j} = \sum_{t=t_i+1}^{t_{i+1}} \mathbb{E}[u_t(j)]$ for $i+1 \in [s]$ and $j \in \{0, 1\}$

$$\begin{aligned}
 \sum_{t=t_i+1}^{t_{i+1}} \mathbb{E}[u_t(\rho'_t)] &= \max_{\rho \in \{0,1\}} \sum_{t=t_i+1}^{t_{i+1}} \mathbb{E}[u_t(\rho)] \\
 &= \frac{1}{2} [P_{i,0} + P_{i,1} + |P_{i,0} - P_{i,1}|] \\
 &= \frac{T}{2s} + |P_{i,0} - T/2s|
 \end{aligned}$$

using $P_{i,0} + P_{i,1} = T/s$. Thus,

$$\begin{aligned}
 \mathbb{E}[R_T] &\geq \sum_{i=0}^{s-1} \left[\left(\frac{T}{2s} + |P_{i,0} - T/2s| \right) - \frac{T}{2s} \right] \\
 &= \sum_{i=0}^{s-1} |P_{i,0} - T/2s|
 \end{aligned}$$

Noting $P_{i,0} = \sum_{t=t_i+1}^{t_{i+1}} \mathbb{E}[u_t(0)] = \sum_{t=t_i+1}^{t_{i+1}} \left(\frac{1+\sigma_t}{2} \right)$ where σ_t are Rademacher variables over $\{-1, 1\}$ and applying Khintchine's inequality (see for example [Ben-David et al., 2009]) we get

$$\mathbb{E}[R_T] \geq \sum_{i=0}^{s-1} \left| \sum_{t=t_i+1}^{t_{i+1}} \sigma_t / 2 \right| \geq \sum_{i=0}^{s-1} \sqrt{T/8s} = \sqrt{sT/8}$$

□

Corollary 32. *We can embed the two-expert setting to get a lower bound for the continuous case.*

Proof. Indeed in Lemma 31 let $\mathcal{C} = [0, 1]$, expert 0 correspond to $\rho_1 = 1/4$, expert 1 corresponds to $\rho_2 = 3/4$ and replace the loss functions by

$$u^{(0)}(\rho) = \begin{cases} 1 & \text{if } \rho < \frac{1}{2} \\ 0 & \text{if } \rho \geq \frac{1}{2} \end{cases} \quad \text{and} \quad u^{(1)}(\rho) = \begin{cases} 0 & \text{if } \rho < \frac{1}{2} \\ 1 & \text{if } \rho \geq \frac{1}{2} \end{cases}$$

We can further generalize this while dispersing the discontinuities somewhat. Instead of having all the discontinuities at $\rho = \frac{1}{2}$, we can have discontinuities dispersed say within an interval $[\frac{1}{3}, \frac{2}{3}]$ and still have $\Omega(\sqrt{sT})$ regret. □

G Experiments

We supplement our results in Section 7 by looking at different changing environments and comparing with performance in the static environment setting. We also look at differences between Generalized and Fixed Share EFs.

G.1 Frequently changing environments

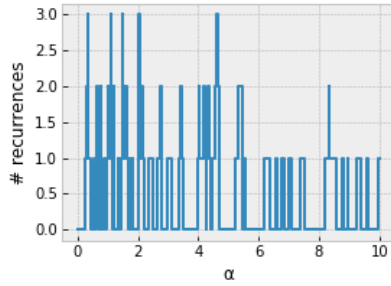
In Section 7 we presented a comparison of our algorithms Fixed Share EF (Algorithm 2) and Generalized Share EF (Algorithm 3) with the Exponential Forecaster algorithm of [Balcan et al., 2018a] for online clustering using well-known datasets. We evaluated the 2-shifted regret for problems where the clustering instance distribution changed exactly once and completely at $T/2$. Here we consider experiments with environments that change more gradually but more frequently.

We consider a sequence of clustering instances drawn from the four datasets. At each time $t \leq T \leq 50$ we sample a subset of the dataset of size 100. For each

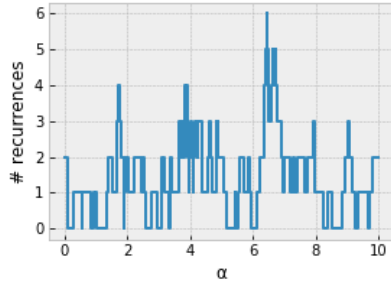
T , we take uniformly random points from all but one classes. The omitted class is changed every T/k rounds, where k is the total number of classes for the dataset. We use parameters $\alpha = \frac{k-1}{T}, \gamma = \frac{1}{T}$ in our algorithms. We determine the hamming cost of $(\bar{\alpha}, 2)$ -Lloyds++-clustering for $\alpha \in \mathcal{C} = [0, 10]$ which is used as the piecewise constant loss function.

We compute the average regret against the best offline algorithm with k shifts. In Figure 3 we plot the average of 20 runs for each dataset. The average regret is higher for all algorithms here since the k -shifted baseline is stronger.

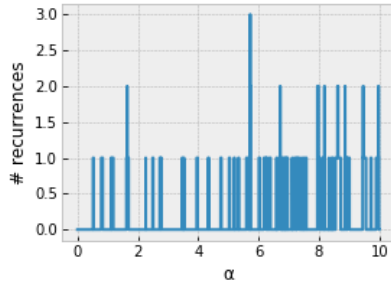
G.2 Generalized vs Fixed Share EFs



(a) MNIST



(b) Omniglot_small_1



(c) Omniglot (full)

Figure 4: Number of recurrences of various values of α in the top decile across all rounds

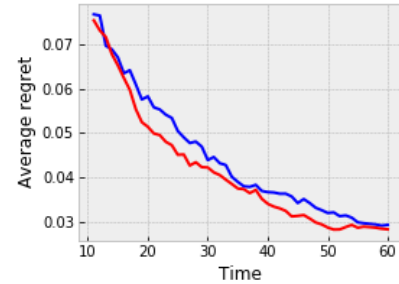
We note that Generalized Share EF performs better on most problem instances. This is because it is better

able to use recurring patterns in *good* values for the parameter that occur non-contiguously, which depends upon the dataset and the problem instance. We verify this hypothesis by a simple experiment (Figure 4).

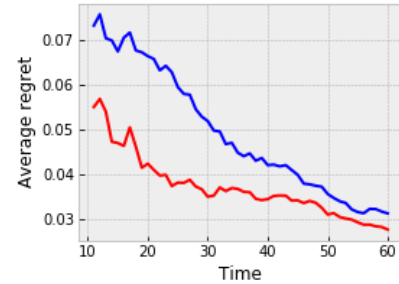
We compute the set of intervals containing the top 10% of the measure of $\alpha \in [0, 10]$ for each t and sum up occurrences of such intervals across all rounds. We observe most recurrences in *Omniglot_small_1* dataset, which explains the large gap between Generalized vs Fixed Share EFs.

G.3 Comparison with static environments

We compare the performance of Fixed Share EF with Exponential Forecaster in static vs dynamic environments on the MNIST dataset. For the changing environment we consider the setting of Section 7, where we present clustering instances for even digits for $t = 1$ through $t = T/2$ and odd digits thereafter. For the static environment we continue to present clustering instances from even labeled digits even after $t = T/2$. We plot the 2-shifted regret in both cases for easier comparison (Figure 5). Note that even though static regret is the more meaningful metric in a static environment, this only changes the baseline and the relative performance of algorithms is unaffected by this choice.



(a) static environment



(b) dynamic environment

Figure 5: Average 2-shifted regret vs game duration T for online clustering against static/dynamic distributions for the MNIST dataset. Color scheme: **Exponential Forecaster**, **Fixed Share EF**

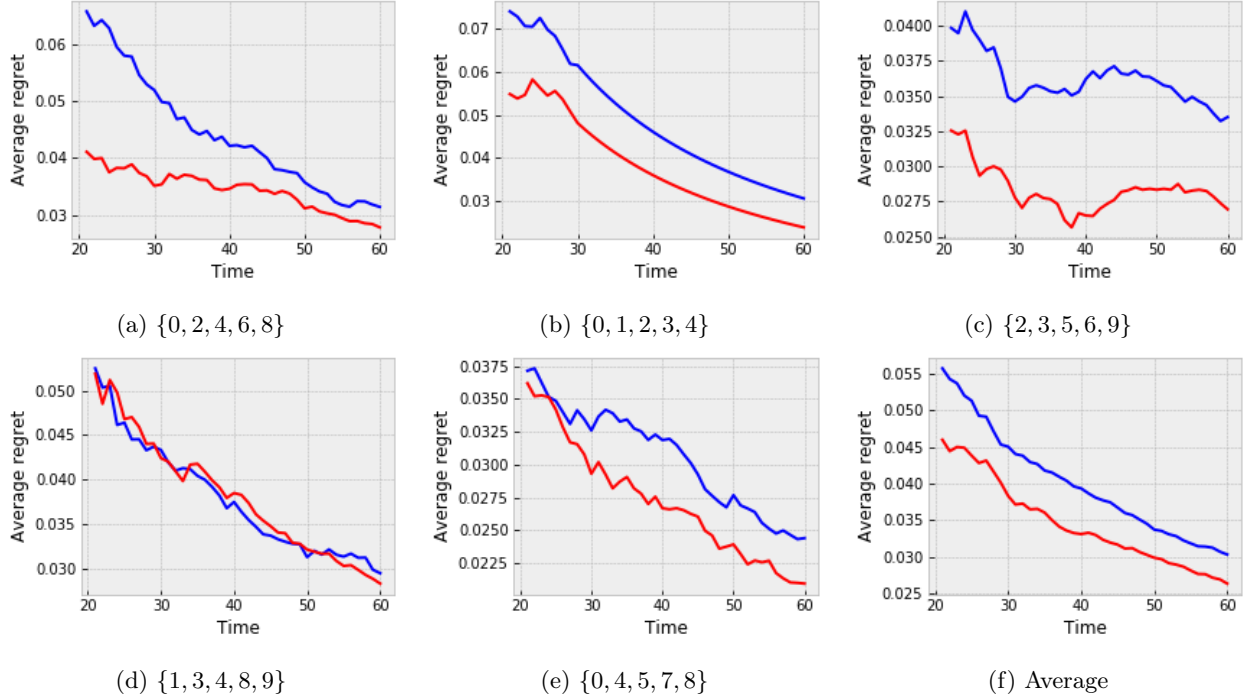


Figure 6: Average 2-shifted regret vs game duration T for online clustering against various dynamic instances for the MNIST dataset. Color scheme: **Exponential Forecaster**, **Fixed Share EF**

Notice that Fixed Share EF is slightly better in the static environment but significantly better in the dynamic environment. It's also worthwhile to note that while the performance of Exponential Forecaster degrades with changing environment, Fixed Share EF actually improves in the dynamic environment since the exploratory updates are more useful.

G.4 Different environments from the same dataset

We look at 2-shifted regret of MNIST clustering instances with the same setting as in Section 7 but with different partitions of clustering classes (i.e. classes used before and after $T/2$). The results are summarized in Figure 6. For each instance we note the set of 5 digits used for drawing uniformly random clustering instances from MNIST till $T/2$, the complement set is used for the remaining rounds. We observe that performance gap between Fixed Share EF and Exponential Forecaster depends not only on the dataset, but also on the clustering instance from the dataset. Across several partitions, Fixed Share EF performs significantly better on average (Figure 6 (f)).