
Corralling Stochastic Bandit Algorithms

Raman Arora
Johns Hopkins University
arora@cs.jhu.edu

Teodor V. Marinov
Johns Hopkins University
tmarino2@jhu.edu

Mehryar Mohri
Courant Institute and Google Research
mohri@google.com

Abstract

We study the problem of corralling stochastic bandit algorithms, that is combining multiple bandit algorithms designed for a stochastic environment, with the goal of devising a corralling algorithm that performs almost as well as the best base algorithm. We give two general algorithms for this setting, which we show benefit from favorable regret guarantees. We show that the regret of the corralling algorithms is no worse than that of the best algorithm containing the arm with the highest reward, and depends on the gap between the highest reward and other rewards.

1 Introduction

We study the problem of *corralling* multi-armed bandit algorithms in a stochastic environment. This consists of selecting, at each round, one out of a fixed collection of bandit algorithms and playing the action returned by that algorithm. Note that the corralling algorithm does not directly select an arm, but only a base algorithm. It never requires knowledge of the action set of each base algorithm. The objective of the corralling algorithm is to achieve a large cumulative reward or a small pseudo-regret, over the course of its interactions with the environment. This problem was first introduced and studied by Agarwal et al. (2016). Here, we are guided by the same motivation but consider the stochastic setting and seek more favorable guarantees. Thus, we assume that the reward, for each arm, is drawn from an unknown distribution.

In the simplest setting of our study, we assume that each base bandit algorithm has access to a distinct set of arms. This scenario appears in several applica-

tions. As an example, consider the online contractual display ads allocation problem (BasuMallick, 2020): when users visit a website, say some page of the online site of a national newspaper, an ads allocation algorithm chooses an ad to display at each specific slot with the goal of achieving the largest value. This could be an ad for a clothing item, which could be meant for the banner of the online front page of that newspaper. To do so, the ads allocation algorithm chooses one out of a large set of advertisers, each a clothing brand or company in this case, which have signed a contract with the ads allocation company. Each clothing company has its own marketing strategy and thus its own bandit algorithm with its own separate set of clothing items or arms. There is no sharing of information between these companies which are typically competitors. Furthermore, the ads allocation algorithm is not provided with any detailed information about the base bandits algorithms of these companies, since that is proprietary information private to each company. The allocation algorithm cannot choose a specific arm or clothing item, it can only choose a base advertiser. The number of ads or arms can be very large. The number of advertisers can also be relatively large in practice, depending on the domain. The number of times the ads allocation is run is in the order of millions or even billions per day, depending on the category of items.

A similar problem arises with online mortgage broker companies offering loans to new applicants. The mortgage broker algorithm must choose a bank, each with different mortgage products. The broker brings a new application exclusively to one of the banks, as part of the contract, which also entitles them to incentives. The bank's online algorithm can be a bandit algorithm proposing a product, and the details of the algorithm are not accessible to the broker; for instance, the bank's credit rate and incentives may depend on the financial and credit history of the applicant. The number of mortgage products is typically fairly large, and the number of online loan requests per day is in the order of several thousands. Other instances of this problem appear when an algorithm can only select one of multiple bandit algorithms and, for privacy or reg-

ulatory reasons, it cannot directly select an arm or receive detailed information about the base algorithms.

In the most general setting we study, there may be an arbitrary sharing of arms between the bandit algorithms. We will only assume that only one algorithm has access to the arm with maximal expected reward, which implies a positive gap between the expected reward of the best arm of any algorithm and that of the best algorithm. This is because we seek to devise a corraling algorithm with favorable gap-dependent pseudo-regret guarantees.

Related work. The previous work most closely related to this study is the seminal contribution by Agarwal et al. (2016) who initiated the general problem of corraling bandit algorithms. The authors gave a general algorithm for this problem, which is an instance of the generic Mirror Descent algorithm with an appropriate mirror map (LOG-BARRIER-OMD), (Foster et al., 2016; Wei and Luo, 2018), and which includes a carefully constructed non-decreasing step-size schedule, also used by Bubeck et al. (2017). The algorithm of Agarwal et al. (2016), however, cannot in general achieve regret bounds better than $\tilde{O}(\sqrt{T})$ in the time horizon, unless optimistic instance-dependent regret bounds are known for the corralled algorithms. Prior to their work, Arora et al. (2012) presented an algorithm for learning deterministic Markov decision processes (MDPs) with adversarial rewards, using an algorithm for corraling bandit linear optimization algorithms. In an even earlier work, Maillard and Munos (2011) attempted to corral EXP3 algorithms (Auer et al., 2002b) with a top algorithm that is a slightly modified version of EXP4. The resulting regret bounds are in $\tilde{O}(T^{2/3})$.

Our work can also be viewed as selecting the best algorithm for a given unknown environment and, in this way, is similar in spirit to the literature solving the *best of both worlds* problem (Audibert and Bubeck, 2009; Bubeck and Slivkins, 2012; Seldin and Slivkins, 2014; Auer and Chiang, 2016; Seldin and Lugosi, 2017; Wei and Luo, 2018; Zimmert and Seldin, 2018; Zimmert et al., 2019) and the model selection problem for linear bandit (Foster et al., 2019; Chatterji et al., 2019).

Very recently, Pacchiano et al. (2020) also considered the problem of corraling stochastic bandit algorithms. The authors seek to treat the problem of model selection, where multiple algorithms might share the best arm. More precisely, the authors consider a setting in which there are K stochastic contextual bandit algorithms and try to minimize the regret with respect to the best overall policy belonging to any of the bandit algorithms. They propose two corraling algorithms, one based on the work of (Agarwal et al., 2016) and

one based on EXP3.P (Auer et al., 2002b). The main novelty in their work is a smoothing technique for each of the base algorithms, which avoids having to restart the base algorithms throughout the T rounds, as was proposed in (Agarwal et al., 2016). The proposed regret bounds are of the order $\tilde{\Theta}(\sqrt{T})$. We expect that the smoothing technique is also applicable to one of the corraling algorithms we propose. Since Pacchiano et al. (2020) allow for algorithms with shared best arms, their main results do not discuss the optimistic setting in which there is a gap between the reward of the optimal policy and all other competing policies, and do not achieve the *optimistic guarantees* we provide. Further, they show a min-max lower bound which states that even if one of the base algorithms is optimistic and contains the best arm, there is still no hope to achieve regret better than $\tilde{\Omega}(\sqrt{T})$ if the best arm is shared by an algorithm with regret $\tilde{\Omega}(\sqrt{T})$. We view their contributions as complementary to ours.

In general, some caution is needed when designing a corraling algorithm, since aggressive strategies may discard or disregard a base learner that admits an arm with the best mean reward if it performs poorly in the initial rounds. Furthermore, as noted by Agarwal et al. (2016), additional assumptions are required on each of the base learners if one hopes to achieve non-trivial corraling guarantees.

Contributions. We first motivate our key assumption that all of the corralled algorithms must have favorable regret guarantees during all rounds. To do so, in Section 3, we show that if one does not assume anytime regret guarantees, then even when corraling simple stochastic bandit algorithms, each with $o(\sqrt{T})$ regret, any corraling strategy will have to incur $\Omega(\sqrt{T})$ regret. Therefore, for the rest of the paper we assume that each base learner admits anytime guarantees. In Section 4 and Section 5, we present two general corraling algorithms whose pseudo-regret guarantees admit a dependency on the gaps between base learners, that is their best arms, and only poly-logarithmic dependence on time horizon. These bounds are syntactically similar to the instance-dependent guarantees for the stochastic multi-armed bandit problem (Auer et al., 2002a). Thus, our corraling algorithm performs almost as well as the best base learner, if it were to be used on its own, modulo gap-dependent terms and logarithmic factors. The algorithm in Section 4 uses the standard UCB ideas combined with a boosting technique, which runs multiple copies of the same base learner. In Section 4.1, we show that simply using UCB-style corraling without boosting can incur linear regret. If, additionally, we assume that each of the base learners satisfy the stability condition adopted in (Agarwal et al., 2016), then, in Section 5 we show that

it suffices to run a single copy of each base learner by using a corraling approach based on OMD. We show that UCB-I (Auer et al., 2002a) can be made to satisfy the stability condition, as long as the confidence bound is rescaled and changed by an additive factor. In Section 6, to further examine the properties of our algorithms, we report the results of experiments with our algorithms for synthetic datasets. Finally, while our main motivation is not model selection, in Section 7, we briefly discuss some related matters and show that our algorithms can help recover several known results in that area.

2 Preliminaries

We consider the problem of corraling K stochastic multi-armed bandit algorithms $\mathcal{A}_1, \dots, \mathcal{A}_K$, which we often refer to as *base algorithms* (base learners). At each round t , a *corraling algorithm* selects a base algorithm $\mathcal{A}_{i_t}$, which plays action a_{i_t, j_t} . The corraling algorithm is not informed of the identity of this action but it does observe its reward $r_t(a_{i_t, j_t})$. The top algorithm then updates its decision rule and provides feedback to each of the base learners $\mathcal{A}_i$. We note that the feedback may be just the empty set, in which case the base learners do not update their state. We will also assume access to the parameters controlling the behavior of each $\mathcal{A}_i$ such as the step size for mirror descent-type algorithms, or the confidence bounds for UCB-type algorithms. Our goal is to minimize the cumulative pseudo-regret of the corraling algorithm as defined in Equation 1:¹

$$\mathbb{E}[R(T)] = T\mu_{1,1} - \mathbb{E}\left[\sum_{t=1}^T r_t(a_{i_t, j_t})\right], \quad (1)$$

where $\mu_{1,1}$ is the mean reward of the best arm.

Notation. We denote by e_i the i th standard basis vector, by $\mathbf{0}_K \in \mathbb{R}^K$ the vector of all 0s, and by $\mathbf{1}_K \in \mathbb{R}^K$ the vector of all 1s. For two vectors $x, y \in \mathbb{R}^K$, $x \odot y$ denotes their Hadamard product. We also denote the line segment between x and y as $[x, y]$. $w_{t,i}$ denotes the i -th entry of a vector $w_t \in \mathbb{R}^K$. Δ^{K-1} denotes the probability simplex in $\mathbb{R}^K$, $D_\Psi(x, y)$ the Bregman divergence induced by the potential Ψ , whose conjugate function we denote by Ψ^* . We use $\mathbf{I}_C$ to denote the indicator function of a set C . For any $k \in \mathbb{N}$, we use the shorthand $[k] := \{1, 2, \dots, k\}$.

For the base algorithms $\mathcal{A}_1, \dots, \mathcal{A}_K$, let $T_i(t)$ be the number of times algorithm $\mathcal{A}_i$ has been played until time t . Let $T_{i,j}(t)$ be the number of times action j has been proposed by algorithm $\mathcal{A}_i$ until time t . Let

¹For conciseness, from now on, we will simply write *regret* instead of *pseudo-regret*.

$[k_i]$ denote the set of arms or action set of algorithm $\mathcal{A}_i$. We denote the reward of arm j in the action set of algorithm i at time t as $r_t(a_{i,j})$ and denote its mean reward by $\mu_{i,j}$. We also use a_{i,j_t} to denote the arm proposed by algorithm $\mathcal{A}_i$ during time t . Further, the algorithm played at time t is denoted as i_t , its action played at time t is a_{i_t, j_t} and the reward for that action is $r_t(a_{i_t, j_t})$ with mean μ_{i_t, j_t} . Let i^* denote the index of the base algorithm that contains the arm with the highest mean reward. Without loss of generality, we will assume that $i^* = 1$. Similarly, we assume that $a_{i,1}$ is the arm with highest reward in algorithm $\mathcal{A}_i$. We assume that the best arm of the best algorithm has a gap to the best arm of every other algorithm. We denote the gap between the best arm of $\mathcal{A}_1$ and the best arm of $\mathcal{A}_i$ as Δ_i : $\Delta_i = \mu_{i^*,1} - \mu_{i,1} > 0$ for $i \neq i^*$. Further, we denote the intra-algorithm gaps by $\Delta_{i,j} = \mu_{i,1} - \mu_{i,j}$. We denote by $\bar{R}_i(t)$ an upper bound on the regret of algorithm $\mathcal{A}_i$ at time t and by $R_i(t)$ the actual regret of $\mathcal{A}_i$, so that $\mathbb{E}[R_i(t)]$ is the expected regret of algorithm $\mathcal{A}_i$ at time t . The asymptotic notations $\tilde{\Omega}$ and $\tilde{O}$ are equal to Ω and O up to poly-logarithmic factors.

3 Lower bounds without anytime regret guarantees

We begin by showing a simple and yet instructive lower bound that helps guide our intuition regarding the information needed from the base algorithms $\{\mathcal{A}_i\}_{i=1}^K$ in the design of a corraling algorithm. Our lower bound is based on corraling base algorithms that only admit a fixed-time horizon regret bound and do not enjoy anytime regret guarantees. We further assume that the corraling strategy cannot simulate anytime regret guarantees on the base algorithms, say by using the so-called doubling trick. This result suggests that the base algorithms must admit a strong regret guarantee during every round of the game.

The key idea behind our construction is the following. Suppose one of the corralled algorithms, $\mathcal{A}_i$, incurs a linear regret over the first $R_i(T)$ rounds. In that case, the corraling algorithm is unable to distinguish between $\mathcal{A}_i$ and another algorithm that mimics the linear regret behavior of $\mathcal{A}_i$ throughout all T rounds, unless the corraling algorithm plays $\mathcal{A}_i$ at least $R_i(T)$ times. The successive elimination algorithm (Even-Dar et al., 2002) benefits from gap-dependent bounds and can have the behavior just described for a base algorithm. Thus, our lower bound is presented for successive elimination base algorithms, all with regret $O(T^{1/4})$. It shows that, with constant probability, no corraling strategy can achieve a more favorable regret than $\tilde{\Omega}(\sqrt{T})$ in that case.

Theorem 3.1. *Let the corralled algorithms be instances of successive elimination defined by a parameter α . With probability $1/4$ over the random sampling of α , any corraling strategy will incur regret at least $\tilde{\Omega}(\sqrt{T})$, while the gap, Δ , between the best and second best reward is such that $\Delta > \omega(T^{-1/4})$ and all algorithms have a regret bound of $\tilde{O}(1/\Delta)$.*

This theorem shows that, even when corraling natural algorithms that benefit from asymptotically better regret bounds, corraling can incur $\tilde{\Omega}(\sqrt{T})$ regret. It can be further proven (Theorem B.3, Appendix B) that, even if the worst case upper bounds on the regret of the base algorithms were known, achieving an optimistic regret guarantee for corraling would not be possible, unless some additional assumptions were made.

4 UCB-style corraling algorithm

The negative result of Section 3 hinge on the fact that the base algorithms do not admit anytime regret guarantees. Therefore, we assume, for the rest of the paper, that the base algorithms, $\{\mathcal{A}_i\}$, satisfy the following:

$$\mathbb{E} \left[t\mu_{i,1} - \sum_{s=1}^t r_s(a_{i_s,j_s}) \right] \leq \bar{R}_i(t), \quad (2)$$

for any time $t \in [T]$. For UCB-type algorithms, such bounds can be derived from the fact that the expected number of pulls, $T_{i,j}(t)$, of a suboptimal arm j , is bounded as $\mathbb{E}[T_{i,j}(t)] \leq c \frac{\log(t)}{(\Delta_{i,j})^2}$, for some time and gap-independent constant c (e.g., Bubeck (2010)), and take the following form, $\bar{R}_i(t) \leq c' \sqrt{k_i t \log(t)}$, for some constant c' .

Suppose that the bound in Equation 2 holds with probability $1 - \delta_t$. Note that such bounds are available for some UCB-type algorithms (Audibert et al., 2009). We can then adopt the optimism in the face of uncertainty principle for each $\mu_{i,1}$ by overestimating it with $\frac{1}{t} \sum_{s=1}^t r_s(a_{i_s,j_s}) + \frac{1}{t} \bar{R}_i(t)$. As long as this occurs with high enough probability, we can construct an upper confidence bound for $\mu_{i,1}$ and use it in a UCB-type algorithm. Unfortunately, the upper confidence bounds required for UCB-type algorithms to work need to hold with high enough probability, which is not readily available from Equation 2 or from probabilistic bounds on the pseudo-regret of anytime stochastic bandit algorithms. In fact, as discussed in Section 4.1, we expect it to be impossible to corral any-time stochastic MAB algorithms with a standard UCB-type strategy. However, a simple boosting technique, in which we run $2 \log(1/\delta)$ copies of each algorithm $\mathcal{A}_i$, gives the following high probability version of the bound in Equation 2.

Algorithm 1 UCB-C

Input: Stochastic bandit algorithms $\mathcal{A}_1, \dots, \mathcal{A}_K$

Output: Sequence of algorithms $(i_t)_{t=1}^T$.

```

1:  $t = 1$ 
2: for  $i = 1, \dots, K$ 
3:    $\mathbb{A}_i = \emptyset$  % contains all copies of  $\mathcal{A}_i$ 
4:   for  $s = 1, \dots, \lceil 2 \log(T) \rceil$ 
5:     Initialize  $\mathcal{A}_i(s)$  as a copy of  $\mathcal{A}_i$ ,  $\hat{\mu}_i(s) = 0$ 
6:     Append  $(\mathcal{A}_i(s), \hat{\mu}_i(s))$  to  $\mathbb{A}_i$ 
7:   for  $i = 1, \dots, K$ 
8:     Foreach  $(\mathcal{A}_i(s), \hat{\mu}_i(s)) \in \mathbb{A}_i$ , play  $\mathcal{A}_i(s)$ , update empirical mean  $\hat{\mu}_i(s)$ ,  $t = t + 2 \log(T)$ 
9:      $\hat{\mu}_{med_i} = \text{Median}(\{\hat{\mu}_i(s)\}_{s=1}^{\lceil 2 \log(T) \rceil})$ 
10:  while  $t \leq T$ 
11:     $b_\ell(t) = \frac{\sqrt{2\bar{R}_{med_\ell}(T_{med_\ell}(t))} + \sqrt{2T_{med_\ell}(t) \log(t)}}{T_{med_\ell}(t)}, \forall \ell \in [K]$ 
12:     $i = \arg\max_{\ell \in [K]} \{\hat{\mu}_{med_\ell} + b_\ell(t)\}$ 
13:    Foreach  $(\mathcal{A}_i(s), \hat{\mu}_i(s)) \in \mathbb{A}_i$ , play  $\mathcal{A}_i(s)$ , update empirical mean  $\hat{\mu}_i(s)$ ,  $t = t + 2 \log(T)$ 
14:     $\hat{\mu}_{med_i} = \text{Median}(\{\hat{\mu}_i(s)\}_{s=1}^{\lceil 2 \log(T) \rceil})$ 

```

Lemma 4.1. *Suppose we run $2 \log(1/\delta)$ copies of algorithm $\mathcal{A}_i$ which satisfies Equation 2. If $\mathcal{A}_{med_i}$ is the algorithm with median cumulative reward at time t , then $\mathbb{P}[t\mu_{i,1} - \sum_{s=1}^t r_s(a_{med_i,j_s}) \geq 2\bar{R}_i(t)] \leq \delta$.*

We consider the following variant of the standard UCB algorithm for corraling. We initialize $2 \log(T)$ copies of each base algorithm $\mathcal{A}_i$. Each $\mathcal{A}_i$ is associated with the median empirical average reward of its copies. At each round, the corraling algorithm picks the $\mathcal{A}_i$ with the highest sum of median empirical average reward and an upper confidence bound based on Lemma 4.1. The pseudocode is given in Algorithm 1. The algorithm admits the following regret guarantees.

Theorem 4.2. *Suppose that algorithms $\mathcal{A}_1, \dots, \mathcal{A}_K$ satisfy the following regret bound $\mathbb{E}[R_i(t)] \leq \sqrt{\alpha k_i t \log(t)}$, respectively for $i \in [K]$. Algorithm 1 selects a sequence of algorithms $i_1, \dots, i_T$ which take actions $a_{i_1,j_1}, \dots, a_{i_T,j_T}$, respectively, such that*

$$\mathbb{E}[R(T)] \leq O \left(\sum_{i \neq i^*} \frac{k_i \log(T)^2}{\Delta_i} + \log(T) \mathbb{E}[R_{i^*}(T)] \right),$$

$$\mathbb{E}[R(T)] \leq O \left(\log(T) \sqrt{KT \log(T) \max_{i \in [K]} (k_i)} \right).$$

We note that both the optimistic and the worst case regret bounds above involve an additional factor that depends on the number of arms, k_i , of the base algorithm $\mathcal{A}_i$. This dependence reflects the complexity of the decision space of algorithm $\mathcal{A}_i$. We conjecture that a complexity-free bound is not possible, in general. To see this, consider a setting where each

$\mathcal{A}_i$, for $i \neq i^*$, only plays arms with equal means $\mu_i = \mu_{1,1} - \Delta_i$. Standard stochastic bandit regret lower bounds, e.g. (Garivier et al., 2018b), state that any strategy on the combined set of arms of all algorithms will incur regret at least $\Omega(\sum_{i \neq i^*} k_i \log(T) / \Delta_i)$. The $\log(T)$ factor in front of the regret of the best algorithm comes from the fact that we are running $\Omega(\log(T))$ copies of it.

4.1 Discussion regarding tightness of bounds

A natural question is if it is possible to achieve bounds that do not have a $\log(T)^2$ scaling. After all, for the simpler stochastic MAB problem, regret upper bounds only scale as $O(\log(T))$ in terms of the time horizon. As already mentioned, the extra logarithmic factor comes from the boosting technique, or, more precisely, the need for exponentially fast concentration of the true regret to its expected value, when using a UCB-type corraling strategy. We now show that, in the absence of such strong concentration guarantees, if only a single copy of each of the base algorithms in Algorithm 1 is run, then linear regret is unavoidable.

Theorem 4.3. *There exist instances $\mathcal{A}_1$ and $\mathcal{A}_2$ of UCB-I and a reward distribution, such that, if Algorithm 1 runs a single copy of $\mathcal{A}_1$ and $\mathcal{A}_2$, then $\mathbb{E}[R(T)] \geq \tilde{\Omega}(\Delta_2 T)$.*

Further, for any algorithm $\mathcal{A}_1$ such that $\mathbb{P}[R_1(t) \geq \frac{1}{2}\Delta_{1,2}\tau] \geq \frac{1}{\tau^c}$, there exists a reward distribution such that if Algorithm 1 runs a single copy of $\mathcal{A}_1$ and $\mathcal{A}_2$, then $\mathbb{E}[R(T)] \geq \tilde{\Omega}((\Delta_{1,2})^c \Delta_2 T)$.

The proof of the above theorem and further discussion can be found in Appendix C.1. The requirement that the regret of the best algorithm satisfies $\mathbb{P}[R_1(t) \geq \frac{1}{2}\Delta_{1,2}\tau] \geq \frac{1}{\tau^c}$ in Theorem 4.3 is equivalent to the condition that the regret of the base algorithms admit only a polynomial concentration. Results in (Salomon and Audibert, 2011) suggest that there cannot be a tighter bound on the tail of the regret for anytime algorithms. It is therefore unclear if the $\log(T)^2$ rate can be improved upon or if there exists a matching information-theoretic lower bound.

5 Corraling using Tsallis-INF

In this section, we consider an alternative approach, based on the work of Agarwal et al. (2016), which avoids running multiple copies of base algorithms. Since the approach is based on the OMD framework, which is naturally suited to losses instead of rewards, for the rest of the section we switch to losses.

We design a corraling algorithm that maintains a probability distribution $w \in \Delta^{K-1}$ over the base al-

gorithms, $\{\mathcal{A}_i\}_{i=1}^K$. At each round, the corraling algorithm samples $i_t \sim w$. Next, $\mathcal{A}_{i_t}$ plays a_{i_t, j_t} and the corraling algorithm observes the loss $\ell_t(a_{i_t, j_t})$. The corraling algorithm updates its distribution over the base algorithms using the observed loss and provides an unbiased estimate $\hat{\ell}_t(a_{i_t, j_t})$ of $\ell_t(a_{i_t, j_t})$ to algorithm $\mathcal{A}_i$: the feedback provided to $\mathcal{A}_i$ is $\hat{\ell}_t(a_{i_t, j_t}) = \frac{\ell_t(a_{i_t, j_t})}{w_{t, i_t}}$, and for all $a_{i_t, j_t} \neq a_{i_t, j_t}$, $\hat{\ell}_t(a_{i_t, j_t}) = 0$. Notice that $\hat{\ell}_t \in \mathbb{R}^K$, as opposed to $\ell_t \in [0, 1]^{\prod_i k_i}$. Essentially, the loss fed to $\mathcal{A}_i$, with probability $w_{t, i}$, is the true loss rescaled by the probability $w_{t, i}$ to observe the loss, and is equal to 0 with probability $1 - w_{t, i}$.

The change of environment induced by the rescaling of the observed losses is analyzed in Agarwal et al. (2016). Following Agarwal et al. (2016), we denote the environment of the original losses $(\ell_t)_t$ as $\mathcal{E}$ and that of the rescaled losses $(\hat{\ell}_t)_t$ as $\mathcal{E}'$. Therefore, in environment $\mathcal{E}$, algorithm $\mathcal{A}_i$ observes $\ell_t(a_{i_t, j_t})$ and in environment $\mathcal{E}'$, $\mathcal{A}_i$ observes $\hat{\ell}_t(a_{i_t, j_t})$. A few important remarks are in order. As in (Agarwal et al., 2016), we need to assume that the base algorithms admit a *stability property* under the change of environment. In particular, if $w_{s, i} \geq \frac{1}{\rho_t}$ for all $s \leq t$ and some $\rho_t \in \mathbb{R}$, then $\mathbb{E}[R_i(t)]$ under environment $\mathcal{E}'$ is bounded by $\mathbb{E}[\sqrt{\rho_t} R_i(t)]$. For completeness, we provide the definition of stability by Agarwal et al. (2016).

Definition 5.1. *Let $\gamma \in (0, 1]$ and let $R: \mathbb{N} \rightarrow \mathbb{R}_+$ be a non-decreasing function. An algorithm $\mathcal{A}$ with action space A is $(\gamma, R(\cdot))$ -stable with respect to an environment $\mathcal{E}$ if its regret under $\mathcal{E}$ is $R(T)$ and its regret under $\mathcal{E}'$ induced by the importance weighting is $\max_{a \in A} \mathbb{E} \left[\sum_{t=1}^T \hat{\ell}_t(a_{i_t, j_t}) - \ell_t(a) \right] \leq \mathbb{E}[(\rho_T)^\gamma R(T)]$.*

We show that UCB-I (Auer et al., 2002a) satisfies the stability property above with $\gamma = \frac{1}{2}$. The techniques used in the proof are also applicable to other UCB-type algorithms. Other algorithms for stochastic bandits like Thompson sampling and OMD/FTRL variants have been shown to be $1/2$ -stable in (Agarwal et al., 2016).

The corraling algorithm of Agarwal et al. (2016) is based on Online Mirror Descent (OMD), where a key idea is to increase the step size whenever the probability of selecting some algorithm $\mathcal{A}_i$ becomes smaller than some threshold. This induces a negative regret term which, coupled with a careful choice of step size (dependent on regret upper bounds of the base algorithms), provides regret bounds that scale as a function of the regret of the best base algorithm.

Unfortunately, the analysis of the corraling algorithm always leads to at least a regret bound of $\tilde{\Omega}(\sqrt{T})$ and also requires knowledge of the regret

bound of the best algorithm. Since our goal is to obtain instance-dependent regret bounds, we cannot appeal to this type of OMD approach. Instead, we draw inspiration from the recent work of [Zimmert and Seldin \(2018\)](#), who use a Follow-the-Regularized-Leader (FTRL) type of algorithm to design an algorithm that is simultaneously optimal for both stochastic and adversarially generated losses, without requiring knowledge of instance-dependent parameters such as the sub-optimality gaps to the loss of the best arm. The overall intuition for our algorithm is as follows. We use the FTRL-type algorithm proposed by [Zimmert and Seldin \(2018\)](#) until the probability to sample some arm falls below a threshold. Next, we run an OMD step with an increasing step size schedule which contributes a negative regret term. After the OMD step, we resume the normal step size schedule and updates from the FTRL algorithm. After carefully choosing the initial step size rate, which can be done in an instance-independent way, the accumulated negative regret terms are enough to compensate for the increased regret due to the change of environment.

5.1 Algorithm and the main result

We now describe our corralling algorithm in more detail. The potential function Ψ_t used in all of the updates is defined by $\Psi_t(w) = -4 \sum_{i \in [K]} \frac{1}{\eta_{t,i}} (\sqrt{w_i} - \frac{1}{2} w_i)$, where $\eta_t = [\eta_{t,1}, \eta_{t,2}, \dots, \eta_{t,K}]$ is the step-size schedule during time t . The algorithm proceeds in epochs and begins by running each base algorithm for $\log(T) + 1$ rounds. Each epoch is twice as large as the preceding, so that the number of epochs is bounded by $\log_2(T)$, and the step size schedule remains non-increasing throughout the epochs, except when an OMD step is taken. The algorithm also maintains a set of thresholds, $\rho_1, \rho_2, \dots, \rho_n$, where $n = O(\log(T))$. These thresholds are used to determine if the algorithm executes an OMD step, while increasing the step size:

$$\begin{aligned} w_{t+1} &= \operatorname{argmin}_{w \in \Delta^{K-1}} \langle \hat{\ell}_t, w \rangle + D_{\Psi_t}(w, w_t), \\ \eta_{t+1,i} &= \beta \eta_{t,i} \text{ (for } i : w_{t,i} \leq 1/\rho_{s_i}), \\ w_{t+2} &= \operatorname{argmin}_{w \in \Delta^{K-1}} \langle \hat{\ell}_{t+1}, w \rangle + D_{\Psi_{t+1}}(w, w_{t+1}), \rho_{s_i} = 2\rho_{s_i} \end{aligned} \quad (3)$$

or the algorithm takes an FTRL step

$$w_{t+1} = \operatorname{argmin}_{w \in \Delta^{K-1}} \langle \hat{L}_t, w \rangle + \Psi_{t+1}(w), \quad (4)$$

where $\hat{L}_t = \hat{L}_{t-1} + \hat{\ell}_t$, unless otherwise specified by the algorithm. We note that the algorithm can only increase the step size during the OMD step. For technical reasons, we require an FTRL step after each OMD step. Further, we require that the second step of

Algorithm 2 Corralling with Tsallis-INF

Input: Mult. constant β , thresholds $\{\rho_i\}_{i=1}^n$, initial step size η , epochs $\{\tau_i\}_{i=1}^m$, algorithms $\{\mathcal{A}_i\}_{i=1}^K$.

Output: Algorithm selection sequence $(i_t)_{t=1}^T$.

```

1: Initialize  $t = 1$ ,  $w_1 = \text{Unif}(\Delta^{K-1})$ ,  $\eta_1 = \eta$ 
2: Initialize current threshold list  $\theta \in [n]^K$  to  $\mathbf{1}$ 
3: while  $t \leq K \log(T) + K$ 
4:   for  $i \in [K]$ 
5:      $\mathcal{A}_i$  plays  $a_{i,j_t}$ ,  $\hat{L}_{1,i} += \ell_t(a_{i,j_t})$ ,  $t += 1$ 
6:  $t = 2$ ,  $w_2 = \nabla \Phi_2(-\hat{L}_1)$ ,  $1/\eta_{t+1}^2 = 1/\eta_t^2 + 1$ 
7: while  $j \leq m$ 
8:   for  $t \in \tau_j$ 
9:      $\mathcal{R}_t = \emptyset$ ,  $\hat{\ell}_t = \text{PLAY-ROUND}(w_t)$ 
10:    if  $t$  is first round of  $\tau_j$  and  $\exists w_{t,i} \leq \frac{1}{\rho_1}$ 
11:      for  $i : w_{t,i} \leq \frac{1}{\rho_1}$ 
12:         $\theta_i = \min\{s \in [n] : w_{t,i} > \frac{1}{\rho_s}\}$ ,  $\mathcal{R}_t = \mathcal{R}_t \cup \{i\}$ .
13:         $(w_{t+3}, \hat{L}_{t+2}) = \text{NRS}(w_t, \hat{\ell}_t, \eta_t, \mathcal{R}_t, \hat{L}_{t-1})$ ,  $t = t + 2$ ,  $\hat{\ell}_t = \text{PLAY-ROUND}(w_t)$ 
14:      if  $\exists i : w_{t,i} \leq \frac{1}{\rho_{\theta_i}}$  and prior step was not NRS
15:        for  $i : w_{t,i} \leq \frac{1}{\rho_{\theta_i}}$ 
16:           $\theta_i += 1$ ,  $\mathcal{R}_t = \mathcal{R}_t \cup \{i\}$ .
17:         $(w_{t+3}, \hat{L}_{t+2}) = \text{NRS}(w_t, \hat{\ell}_t, \eta_t, \mathcal{R}_t, \hat{L}_{t-1})$ ,  $t = t + 2$ ,  $\hat{\ell}_t = \text{PLAY-ROUND}(w_t)$ 
18:      else
19:         $1/\eta_{t+1}^2 = 1/\eta_t^2 + 1$ ,  $w_{t+1} = \nabla \Phi_{t+1}(-\hat{L}_t)$ 
```

each epoch be an OMD step if there exists at least one $w_{t,i} \leq \frac{1}{\rho_1}$. The algorithm also can enter an OMD step during an epoch if at least one $w_{t,i}$ becomes smaller than a threshold $\frac{1}{\rho_{s_i}}$ which has not been exceeded so far.

We set the probability thresholds so that $\rho_1 = O(1)$, $\rho_j = 2\rho_{j-1}$ and $\frac{1}{\rho_n} \geq \frac{1}{T}$, so that $n \leq \log_2(T)$. In the beginning of each epoch, except for the first epoch, we check if $w_{t,i} < \frac{1}{\rho_1}$. If it is, we increase the step size as $\eta_{t+1,i} = \beta \eta_{t,i}$ and run the OMD step. The pseudocode for the algorithm is given in Algorithm 2. The routines OMD-STEP and PLAY-ROUND can be found in Algorithm 6 and Algorithm 7 (Appendix D) respectively. OMD-STEP essentially does the update described in Equation 3 and PLAY-ROUND samples and plays an algorithm, after which constructs an unbiased estimator of the losses and feeds these back to all of the sub-algorithms. We show the following regret bound for the corralling algorithm.

Theorem 5.2. *Let $\bar{R}_i(\cdot)$ be a function upper bounding the expected regret, $\mathbb{E}[R_i(\cdot)]$, of $\mathcal{A}_i$ for all $i \in [K]$. For $\beta = e^{1/\log(T)^2}$ and for η such that for all $i \in [K]$, $\eta_{1,i} \leq \min_{t \in [T]} \frac{(1 - \exp(-\frac{1}{\log(T)^2})) \sqrt{t}}{50 \bar{R}_i(t)}$, the expected*

Algorithm 3 NEG-REG-STEP(NRS)

Input: Prior iterate w_t , loss $\hat{\ell}_t$, step size η_t , set of rescaled step-sizes $\mathcal{R}_t$, cumulative loss $\hat{L}_{t-1}$
Output: Plays two rounds of the game and returns distribution w_{t+3} and cumulative loss $\hat{L}_{t+2}$

- 1: $(w_{t+1}, \hat{L}_t) = \text{OMD-STEP}(w_t, \hat{\ell}_t, \eta_t, \mathcal{R}_t, \hat{L}_{t-1})$
- 2: $\hat{\ell}_{t+1} = \text{PLAY-ROUND}(w_{t+1})$, $\hat{L}_{t+1} = \hat{L}_t + \hat{\ell}_{t+1}$
- 3: **for** all i such that $w_{t,i} \leq \frac{1}{\rho_1}$
- 4: $\eta_{t+2,i} = \beta \eta_{t,i}$, $\mathcal{R}_t = \mathcal{R}_t \cup \{i\}$ and restart $\mathcal{A}_i$ with updated environment $\theta_i = \frac{1}{2w_{t,i}}$
- 5: $w_{t+2} = \nabla \Phi_{t+2}(-\hat{L}_{t+1})$
- 6: $\hat{\ell}_{t+2} = \text{PLAY-ROUND}(w_{t+2})$
- 7: $\hat{L}_{t+2} = \hat{L}_{t+1} + \hat{\ell}_{t+2}$, $\eta_{t+3} = \eta_{t+2}$, $t = t + 2$
- 8: $w_{t+1} = \nabla \Phi_{t+1}(-\hat{L}_t)$, $t = t + 1$

regret of Algorithm 2 is bounded as follows: $\mathbb{E}[R(T)] \leq O\left(\sum_{i \neq i^*} \frac{\log(T)}{\eta_{1,i}^2 \Delta_i} + \mathbb{E}[R_{i^*}(T)]\right)$.

To parse the bound above, suppose $\{\mathcal{A}_i\}_{i \in [K]}$ are standard stochastic bandit algorithms such as UCB-I. In Theorem 5.4, we show that UCB-I is indeed $\frac{1}{2}$ -stable as long as we are allowed to rescale and introduce an additive factor to the confidence bounds. In this case, a worst-case upper bound on the regret of any $\mathcal{A}_i$ is $\mathbb{E}[R_i(t)] \leq c\sqrt{k_i \log(t)t}$ for all $t \in [T]$ and some universal constant c . We note that the min-max regret bound for the stochastic multi-armed bandit problem is $\Theta(\sqrt{KT})$ and most known any-time algorithms solving the problem achieve this bound up to poly-logarithmic factors. Further we note that $\left(1 - \exp\left(-\frac{1}{\log(T)^2}\right)\right) > \frac{1}{e \log(T)^2}$. This suggests that the bound in Theorem 5.2 on the regret of the corraling algorithm is at most $O\left(\sum_{i \neq i^*} \frac{k_i \log(T)^5}{\Delta_i} + \mathbb{E}[R_{i^*}(T)]\right)$. In particular, if we instantiate $\mathbb{E}[R_{i^*}(T)]$ to the instance-dependent bound of $O\left(\sum_{j \neq 1} \frac{\log(T)}{\Delta_{i^*,j}}\right)$, the regret of Algorithm 2 is bounded by $O\left(\sum_{i \neq i^*} \frac{k_i \log(T)^5}{\Delta_i} + \sum_{j \neq 1} \frac{\log(T)}{\Delta_{i^*,j}}\right)$. In general we cannot exactly compare the current bound with that of UCB-C (Algorithm 1), as the regret bound in Theorem 5.2 has worse scaling in the time horizon on the gap-dependent terms, compared to the regret bound in Theorem 4.2, but has no additional scaling in front of the $\mathbb{E}[R_{i^*}(T)]$ term. In practice we observe that Algorithm 2 outperforms Algorithm 1.

Since essentially all stochastic multi-armed bandit algorithms enjoy a regret bound, in time horizon, of the order $\tilde{O}(\sqrt{T})$, we are guaranteed that $1/\eta_{t,i}^2$ scales only poly-logarithmically with the time horizon. What happens, however, if algorithm $\mathcal{A}_i$ has a worst case regret bound of the order $\omega(\sqrt{T})$? For the next part of the

discussion, we only focus on time horizon dependence. As a simple example, suppose that $\mathcal{A}_i$ has worst case regret of $T^{2/3}$ and that $\mathcal{A}_{i^*}$ has a worst case regret of $\sqrt{T}$. In this case, Theorem 5.2 tells us that we should set $\eta_{1,i} = \tilde{O}(1/T^{1/6})$ and hence the regret bound scales at least as $\Omega(T^{1/3}/\Delta_i + \mathbb{E}[R_{i^*}(T)])$. In general, if the worst case regret bound of $\mathcal{A}_i$ is in the order of T^α we have a regret bound scaling at least as $T^{2\alpha-1}/\Delta_i$. This is not unique to Algorithm 2 and a similar scaling of the regret would occur in the bound for Algorithm 1 due to the scaling of confidence intervals.

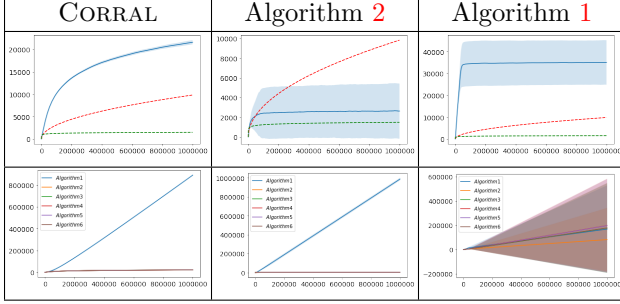
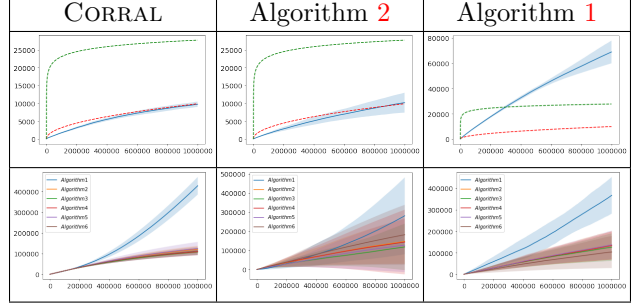
Corraling in an adversarial environment. Because Algorithm 2 is based on a best of both worlds algorithm, we can further handle the case when the losses/rewards are generated adversarially or whenever the best overall arm is shared across multiple algorithms, similarly to the settings studied by Agarwal et al. (2016); Pacchiano et al. (2020).

Theorem 5.3. Let $\bar{R}_{i^*}(\cdot)$ be a function upper bounding the expected regret of $\mathcal{A}_{i^*}$, $\mathbb{E}[R_{i^*}(\cdot)]$. For any $\eta_{1,i^*} \leq \min_{t \in [T]} \frac{(1 - \exp(-\frac{1}{\log(T)^2}))\sqrt{t}}{50\bar{R}_{i^*}(t)}$ and $\beta = e^{1/\log(T)^2}$ it holds that the expected regret of Algorithm 2 is bounded as follows: $\mathbb{E}[R(T)] \leq O\left(\max_{w \in \Delta^{K-1}} \sqrt{T} \sum_{i=1}^K \frac{\sqrt{w_i}}{\eta_{1,i}} + \mathbb{E}[R_{i^*}(T)]\right)$.

The bound in Theorem 5.3 essentially evaluates to $O(\max(\sqrt{TK}, \max_{i \in [K]} \bar{R}_i(T)) + \mathbb{E}[R_{i^*}(T)])$. Unfortunately, this is not quite enough to recover the results in (Agarwal et al., 2016; Pacchiano et al., 2020). This is attributed to the fact that we use the $\frac{1}{2}$ -Tsallis entropy as the regularizer instead of the log-barrier function. It is possible to improve the above bound for algorithms with stability $\gamma < 1/2$, however, because model selection is not the primary focus of this work, we will not present such results here.

Stability of UCB-I. We now briefly discuss how the regret bounds of UCB-I and similar algorithms change whenever the variance of the stochastic losses is rescaled by Algorithm 2. Let us focus on base learner $\mathcal{A}_i$ during epoch τ_j . During epoch τ_j , there is some largest threshold ρ_{s_i} which is never exceeded by the inverse probabilities, i.e., $\min_{t \in \tau_j} w_{t,i} \geq 1/\rho_{s_i}$. This implies that the rescaled losses are in $[0, \rho_{s_i}]$. Further, their variance is bounded by $\mathbb{E}[\hat{\ell}_t(i)^2] = \mathbb{E}[\ell_t(a_{i,j_t})^2/w_{t,i}] \leq \rho_{s_i}$. Using a version of Freedman's inequality (Freedman, 1975), we show the following.

Theorem 5.4 (Informal). Suppose that during epoch τ_j of size $\mathcal{T}_j$, UCB-I (Auer et al., 2002a) uses an upper confidence bound $\sqrt{\frac{4\rho_{s_i} \log(t)}{\mathcal{T}_{i,j}(t)}} + \frac{4\rho_{s_i} \log(t)}{3\mathcal{T}_{i,j}(t)}$ for arm j at time t . Then, the expected regret of $\mathcal{A}_i$ under the rescaled rewards is at most $\mathbb{E}[R_i(\mathcal{T}_j)] \leq \sqrt{8\rho_{s_i} k_i \mathcal{T}_j \log(\mathcal{T}_j)}$.


 Table 1: Regret for corraling when $\Delta_i = 0.2$

 Table 2: Regret for corraling when $\Delta_i = 0.02$

We expect that other UCB-type algorithms (Audibert et al., 2009; Garivier and Cappé, 2011; Bubeck et al., 2013; Garivier et al., 2018a) should also be $\frac{1}{2}$ -stable.

6 Empirical results

In this section, we further examine the empirical properties of our algorithms via experiments on synthetically generated datasets. We compare Algorithm 1 and Algorithm 2 to the Corral algorithm (Agarwal et al., 2016)[Algorithm 1], which is also used in (Pacchiano et al., 2020). We note that Pacchiano et al. (2020) also use Exp3.P as a corraling algorithm. Recent work (Lee et al., 2020) suggests that Corral exhibits similar high probability regret guarantees as Exp3.P and that Corral would completely outperform Exp3.P.

Experimental setup. The algorithms that we corral are UCB-I, Thompson sampling (TS), and FTRL with $\frac{1}{2}$ -Tsallis entropy regularizer (Tsallis-INF). When implementing Algorithm 2 and Corral, we make an important deviation from what theory prescribes: we *never* restart the corralled algorithms and run them with their default parameters. In all our experiments, we corral two instances of UCB-I, TS, and Tsallis-INF for a total of six algorithms. The algorithm containing the best arm plays over 10 arms. Every other algorithm plays over 5 arms. The rewards for each base algorithm are Bernoulli random variables with expectations set so that for all $i > 2$ and $j > 1$, $\Delta_{i,j} = 0.01$. We run two sets of experiments with Δ_i equal to either 0.2 or 0.02. This setting implies that Algorithm 1 always contains the best arm and that the best arm of each base algorithm is arm one. Even though $\Delta_{i,j} = 0.01$ implies large regret for all sub-optimal algorithms, it also reduces the variance of the total reward for these algorithms thereby making the corraling problem harder. Finally, the time horizon is set to $T = 10^6$. For a more extensive discussion, about our choice of algorithms and parameters for the experimental setup we refer the reader to Appendix A.

Large gap experiments. Table 1 reports the regret (top) and number of plays of each algorithm found in our experiments when $\Delta_i = 0.2$. The plots represent the average regret, in blue, and the average number of pulls of each algorithm (color according to the legend) over 75 runs of each experiment. The standard deviation is represented by the shaded blue region. The algorithm that contains the optimal arm is $\mathcal{A}_1$ and is an instance of UCB-I. The red dotted line in the top plots is given by $4\sqrt{KT} + \mathbb{E}[R_1(T)]$, and the green dotted line is given by $4 \sum_{i \neq 1} \frac{k_i \log(T)}{\Delta_i} + \mathbb{E}[R_1(T)]$. These lines serve as a reference across experiments and we believe they are more accurate upper bounds for the regret of the proposed and existing algorithms. As expected, we see that, in the large gap regime, the Corral algorithm exhibits $\Omega(\sqrt{T})$ regret, while the regret of Algorithm 2 remains bounded in $O(\log(T))$. Algorithm 1 admits two regret phases. In the initial phase, its regret is linear, while in the second phase it is logarithmic. This is typical of UCB strategies in the stochastic MAB problem (Garivier et al., 2018b).

Small gap experiments Table 2 reports the results of our experiments for $\Delta_i = 0.02$. The setting of the experiments is the same as in the large gap case. We observe that both Corral and Algorithm 2 behave according to the $O(\sqrt{T})$ bounds. This is expected since, when $\Delta_i = 0.02$, the optimistic bound dominates the $\sqrt{T}$ -bound. The result for Algorithm 1 might be somewhat surprising, as its regret exceeds both the green and red lines. We emphasize that this experiment does not contradict Theorem 4.2. Indeed, if we were to plot the green and red lines according to the bounds of Theorem 4.2, the regret would remain below both lines.

Our experiments suggest that Algorithm 2 is the best corraling algorithm. A tighter analysis would potentially yield optimistic regret bounds in the order of $O\left(\sum_{i \neq i^*} \frac{k_i \log(T)}{\Delta_i} + \mathbb{E}[R_{i^*}(T)]\right)$. Furthermore, we expect that the bounds of Theorem 4.2 are tight. For more detailed experiments, we refer the reader to Appendix A.

7 Model selection for linear bandits

While the main focus of the paper is corraling MAB base learners when there exists a best overall base algorithm, we now demonstrate that several known model selection results can be recovered using Algorithm 2.

We begin by recalling the model selection problem for linear bandits. The learner is given access to a set of loss functions $\mathcal{F}: \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}$ mapping from contexts $\mathcal{X}$ and actions $\mathcal{A}$ to losses. In the linear bandits setting, $\mathcal{F}$ is structured as a nested sequence of classes $\mathcal{F}_1 \subseteq \mathcal{F}_2 \subseteq \dots \subseteq \mathcal{F}_K = \mathcal{F}$, where each $\mathcal{F}_i$ is defined as

$$\mathcal{F}_i = \{(x, a) \rightarrow \langle \beta_i, \phi_i(x, a) \rangle : \beta_i \in \mathbb{R}^{d_i}\},$$

for some feature embedding $\phi_i: \mathcal{X} \times \mathcal{A} \rightarrow \mathbb{R}^{d_i}$. It is assumed that each feature embedding ϕ_i contains ϕ_{i-1} as its first d_{i-1} coordinates. It is further assumed that there exists a smallest $i^* \leq K$ to which the optimal parameter β^* belongs, that is the observed losses for each context-action pair (x, a) satisfy $\ell_t(x, a) = \mathbb{E}[\langle \beta^*, \phi_{i^*}(x, a) \rangle]$. The goal in the model selection problem is to identify i^* and compete against the smallest loss for the t -th context in $\mathbb{R}^{d_{i^*}}$ by minimizing the regret:

$$R_{i^*}(T) = \sum_{t=1}^T \left(\mathbb{E}[\langle \beta^*, \phi_{i^*}(x_t, a_t) \rangle] - \min_{a \in \mathcal{A}} \mathbb{E}[\langle \beta^*, \phi_{i^*}(x_t, a) \rangle] \right),$$

where the expectation is with respect to all randomness in the sampling of the contexts $x_t \sim \mathcal{D}$, actions and additional noise. We adopt the standard assumption that, given x_t , the observed loss for any a can be expressed as follows: $\langle \beta, \phi_i(x_t, a) \rangle + \xi_t$, where ξ_t is a zero-mean, sub-Gaussian random variable with variance proxy 1 and for each of the context-action pairs it holds that $\langle \beta, \phi_i(x_t, a) \rangle \in [0, 1]$.

7.1 Algorithm and main result

We assume that there are K base learners $\{\mathcal{A}_i\}_{i=1}^K$ such that the regret of $\mathcal{A}_i$, for $i \geq i^*$, is bounded by $\tilde{O}(d_i^\alpha \sqrt{T})$. That is, whenever the model is correctly specified, the i -th algorithm admits a meaningful regret guarantee. In the setting of Foster et al. (2019), $\mathcal{A}_i$ can be instantiated as LINUCB and in that case $\alpha = 1/2$. Further, in the setting of infinite arms, $\mathcal{A}_i$ can be instantiated as OFUL (Abbasi-Yadkori et al., 2011), in which case $\alpha = 1$. Both $\alpha = 1/2$ and $\alpha = 1$ govern the min-max optimal rates in the respective settings. Our algorithm is now a simple modification of Algorithm 2. At every time-step t , we update $\hat{L}_t = \hat{L}_{t-1} + \hat{\ell}_t + \mathbf{d}$, where $\mathbf{d}_i = \frac{d_i^{2\alpha}}{\sqrt{T}}$. Intuitively, our

modification creates a gap between the losses of $\mathcal{A}_{i^*}$ and any $\mathcal{A}_i$ for $i > i^*$ of the order $d_i^{2\alpha}$. On the other hand for any $i < i^*$, perturbing the loss can result in at most additional $d_{i^*}^{2\alpha} \sqrt{T}$ regret. With the above observations, the bound guaranteed by Theorem 5.2 implies that the modified algorithm should incur at most $\tilde{O}(d_{i^*}^{2\alpha} \sqrt{T})$ regret. In Appendix F, we show the following regret bound.

Theorem 7.1. *Assume that every base learner $\mathcal{A}_i$, $i \geq i^*$, admits a $\tilde{O}(d_i^\alpha \sqrt{T})$ regret. Then, there exists a corraling strategy with expected regret bounded by $\tilde{O}(d_{i^*}^{2\alpha} \sqrt{T} + K\sqrt{T})$. Moreover, under the additional assumption that the following holds for any $i < i^*$, for all $(x, a) \in \mathcal{X} \times \mathcal{A}$*

$$\mathbb{E}[\langle \beta_i, \phi_i(x, a) \rangle] - \min_{a \in \mathcal{A}} \mathbb{E}[\langle \beta^*, \phi_{i^*}(x, a) \rangle] \geq 2 \frac{d_{i^*}^{2\alpha} - d_i^{2\alpha}}{\sqrt{T}},$$

the expected regret of the same strategy is bounded as $\tilde{O}(d_{i^}^{2\alpha} \sqrt{T} + K\sqrt{T})$.*

Typically, we have $K = O(\log(T))$ and thus Theorem 7.1 guarantees a regret of at most $\tilde{O}(d_{i^*}^{2\alpha} \sqrt{T})$. Furthermore, under a gap-assumption, which implies that the value of the smallest loss for the optimal embedding i^* is sufficiently smaller compared to the value of any sub-optimal embedding $i < i^*$, we can actually achieve a corraling regret of the order $R_{i^*}(T)$. In particular, for the setting of Foster et al. (2019), our strategy yields the desired $\tilde{O}(\sqrt{d_{i^*} T})$ regret bound. Notice that the regret guarantees are only meaningful as long as $d_{i^*} = o(T^{1/(2\alpha)})$. In such a case, the second assumption on the gap is that the gap is lower bounded by $o(1)$. This is a completely problem-dependent assumption and in general we expect that it cannot be satisfied.

8 Conclusion

We presented an extensive analysis of the problem of corraling stochastic bandits. Our algorithms are applicable to a number of different contexts where this problem arises. There are also several natural extensions and related questions relevant to our study. One natural extension is the case where the set of arms accessible to the base algorithms admit some overlap and where the reward observed by one algorithm could serve as side-information to another algorithm. Another extension is the scenario of corraling online learning algorithms with feedback graphs. In addition to these and many other interesting extensions, our analysis may have some connection with the study of other problems such as model selection in contextual bandits (Foster et al., 2019) or active learning.

Acknowledgements

This research was supported in part by NSF BIG-DATA awards IIS-1546482, IIS-1838139, NSF CAREER award IIS-1943251, and by NSF CCF-1535987, NSF IIS-1618662, and a Google Research Award. RA would like to acknowledge support provided by Institute for Advanced Study and the Johns Hopkins Institute for Assured Autonomy. We warmly thank Julian Zimmert for insightful discussions regarding the Tsallis-INF approach.

References

- Yasin Abbasi-Yadkori, Dávid Pál, and Csaba Szepesvári. Improved algorithms for linear stochastic bandits. In *NIPS*, volume 11, pages 2312–2320, 2011.
- Alekh Agarwal, Haipeng Luo, Behnam Neyshabur, and Robert E Schapire. Corralling a band of bandit algorithms. *arXiv preprint arXiv:1612.06246*, 2016.
- Raman Arora, Ofer Dekel, and Ambuj Tewari. Deterministic MDPs with adversarial rewards and bandit feedback. In *Proceedings of the Twenty-Eighth Conference on Uncertainty in Artificial Intelligence*, pages 93–101, 2012.
- Jean-Yves Audibert and Sébastien Bubeck. Minimax policies for adversarial and stochastic bandits. In *COLT*, pages 217–226, 2009.
- Jean-Yves Audibert, Rémi Munos, and Csaba Szepesvári. Exploration–exploitation tradeoff using variance estimates in multi-armed bandits. *Theoretical Computer Science*, 410(19):1876–1902, 2009.
- Peter Auer and Chao-Kai Chiang. An algorithm with nearly optimal pseudo-regret for both stochastic and adversarial bandits. In *Conference on Learning Theory*, pages 116–120, 2016.
- Peter Auer, Nicolo Cesa-Bianchi, and Paul Fischer. Finite-time analysis of the multiarmed bandit problem. *Machine learning*, 47(2-3):235–256, 2002a.
- Peter Auer, Nicolo Cesa-Bianchi, Yoav Freund, and Robert E Schapire. The nonstochastic multiarmed bandit problem. *SIAM journal on computing*, 32(1):48–77, 2002b.
- Peter L Bartlett, Varsha Dani, Thomas Hayes, Sham Kakade, Alexander Rakhlin, and Ambuj Tewari. High-probability regret bounds for bandit online linear optimization. In *Conference on Learning Theory*, 2008.
- Chiradeep BasuMallick. What is display advertising? definition, targeting process, management, network, types, and examples. <https://marketing.toolbox.com/articles/what-is-display-advertising-definition-targeting-process-management-network-types-and-examples>, 2020.
- Haim Brezis. *Functional analysis, Sobolev spaces and partial differential equations*. Springer Science & Business Media, 2010.
- Sébastien Bubeck. *Bandits games and clustering foundations*. PhD thesis, INRIA Nord Europe (Lille, France), 2010.
- Sébastien Bubeck and Aleksandrs Slivkins. The best of both worlds: Stochastic and adversarial bandits. In *Conference on Learning Theory*, pages 42–1, 2012.
- Sébastien Bubeck, Nicolo Cesa-Bianchi, and Gábor Lugosi. Bandits with heavy tail. *IEEE Transactions on Information Theory*, 59(11):7711–7717, 2013.
- Sébastien Bubeck, Yin Tat Lee, and Ronen Eldan. Kernel-based methods for bandit convex optimization. In *Proceedings of the 49th Annual ACM SIGACT Symposium on Theory of Computing*, pages 72–85, 2017.
- Niladri S Chatterji, Vidya Muthukumar, and Peter L Bartlett. Osom: A simultaneously optimal algorithm for multi-armed and linear contextual bandits. *arXiv preprint arXiv:1905.10040*, 2019.
- Eyal Even-Dar, Shie Mannor, and Yishay Mansour. Pac bounds for multi-armed bandit and markov decision processes. In *International Conference on Computational Learning Theory*, pages 255–270. Springer, 2002.
- Dylan J. Foster, Zhiyuan Li, Thodoris Lykouris, Karthik Sridharan, and Éva Tardos. Learning in games: Robustness of fast convergence. In *Proceedings of NIPS*, pages 4727–4735, 2016.
- Dylan J Foster, Akshay Krishnamurthy, and Haipeng Luo. Model selection for contextual bandits. In *Advances in Neural Information Processing Systems*, pages 14714–14725, 2019.
- David A Freedman. On tail probabilities for martingales. *the Annals of Probability*, pages 100–118, 1975.
- Aurélien Garivier and Olivier Cappé. The kl-ucb algorithm for bounded stochastic bandits and beyond. In *Proceedings of the 24th annual conference on learning theory*, pages 359–376, 2011.
- Aurélien Garivier, Hédi Hadiji, Pierre Menard, and Gilles Stoltz. Kl-ucb-switch: optimal regret bounds for stochastic bandits from both a distribution-dependent and a distribution-free viewpoints. *arXiv preprint arXiv:1805.05071*, 2018a.
- Aurélien Garivier, Pierre Ménard, and Gilles Stoltz. Explore first, exploit next: The true shape of regret in bandit problems. *Mathematics of Operations Research*, 44(2):377–399, 2018b.

- Chung-Wei Lee, Haipeng Luo, Chen-Yu Wei, and Mengxiao Zhang. Bias no more: high-probability data-dependent regret bounds for adversarial bandits and mdps. *arXiv preprint arXiv:2006.08040*, 2020.
- Odalric-Ambrym Maillard and Rémi Munos. Adaptive bandits: Towards the best history-dependent strategy. In *Proceedings of AISTATS*, pages 570–578, 2011.
- Aldo Pacchiano, My Phan, Yasin Abbasi-Yadkori, Anup Rao, Julian Zimmert, Tor Lattimore, and Csaba Szepesvari. Model selection in contextual stochastic bandit problems. *arXiv preprint arXiv:2003.01704*, 2020.
- Antoine Salomon and Jean-Yves Audibert. Deviations of stochastic bandit regret. In *International Conference on Algorithmic Learning Theory*, pages 159–173. Springer, 2011.
- Yevgeny Seldin and Gábor Lugosi. An improved parametrization and analysis of the exp3++ algorithm for stochastic and adversarial bandits. In *The 30th Annual Conference on Learning Theory (COLT) Conference on Learning Theory*, pages 1743–1759. Proceedings of Machine Learning Research, 2017.
- Yevgeny Seldin and Aleksandrs Slivkins. One practical algorithm for both stochastic and adversarial bandits. In *ICML*, pages 1287–1295, 2014.
- Chen-Yu Wei and Haipeng Luo. More adaptive algorithms for adversarial bandits. In *Proceedings of COLT 2018*, pages 1263–1291, 2018.
- Julian Zimmert and Yevgeny Seldin. An optimal algorithm for stochastic and adversarial bandits. *arXiv preprint arXiv:1807.07623*, 2018.
- Julian Zimmert, Haipeng Luo, and Chen-Yu Wei. Beating stochastic and adversarial semi-bandits optimally and simultaneously. In *Proceedings of ICML*, pages 7683–7692, 2019.

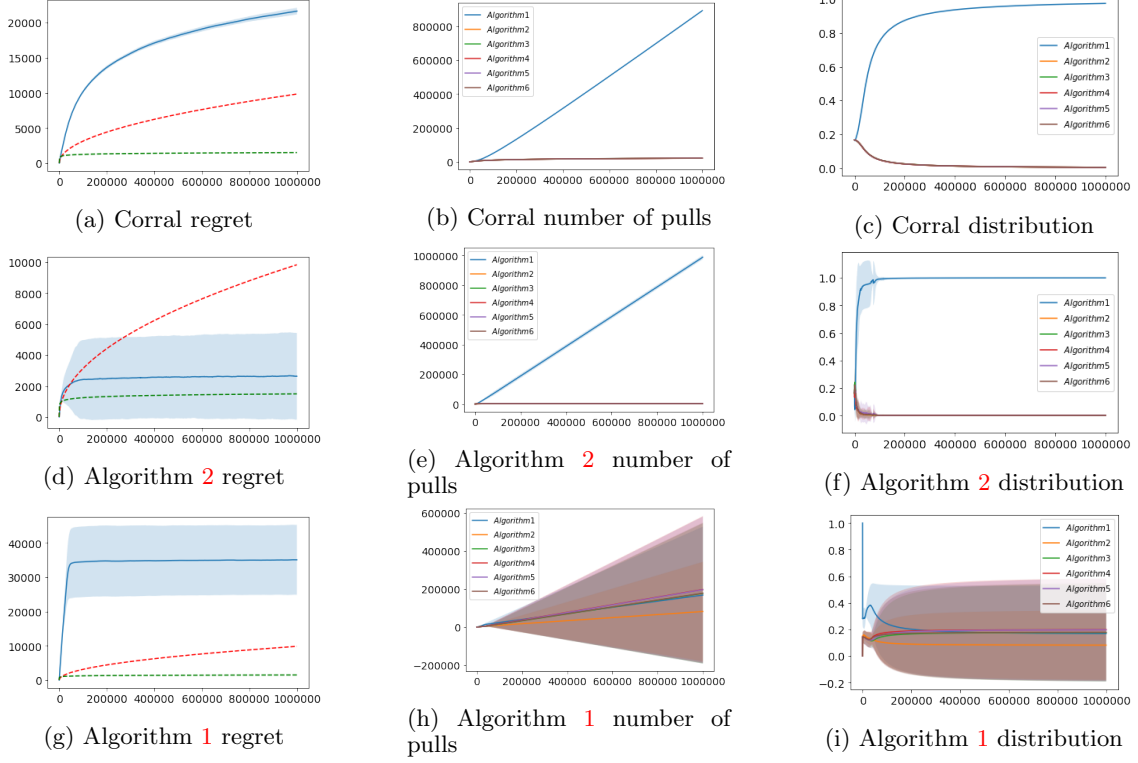
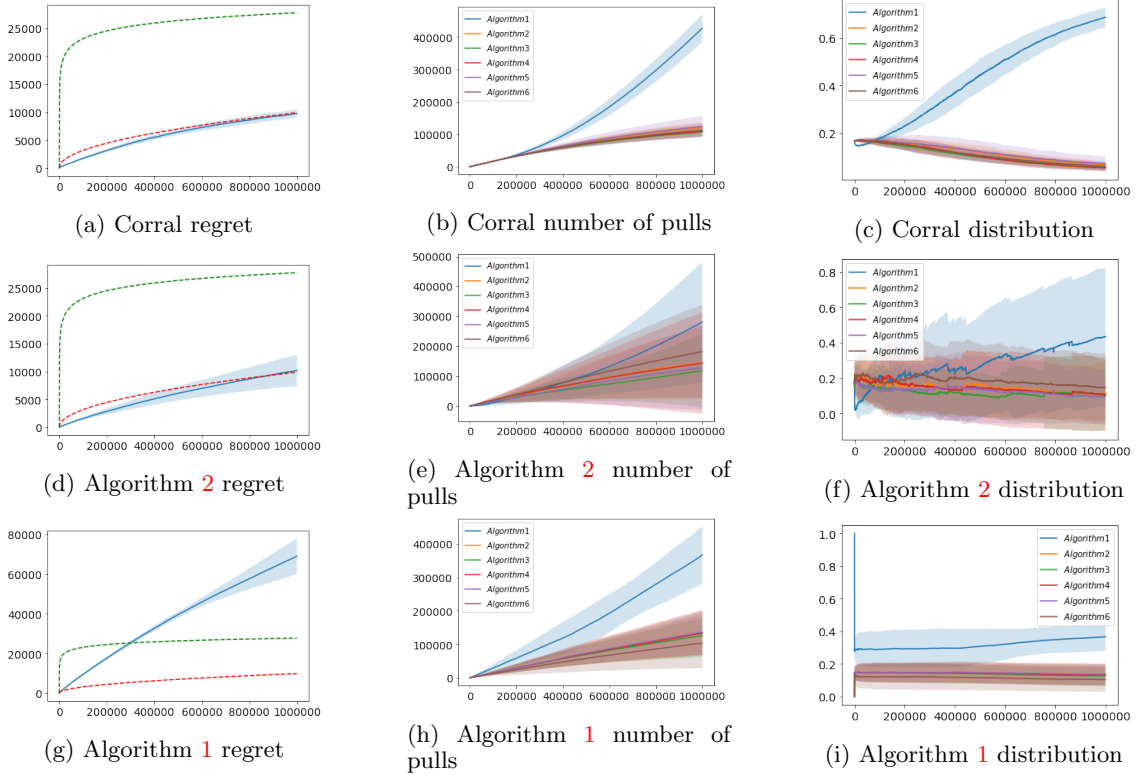
A Additional experiments

We now provide more detailed plots for our experiments, including number of times each corralled algorithm has been played and the distribution over corralled distribution each of the corraling algorithm keeps (in the case of Algorithm 1 this is just the empirical distribution of played algorithms). We additionally present experiments in which the corralled algorithm containing the best arm is FTRL with $\frac{1}{2}$ -Tsallis entropy regularization and Thompson sampling.

Detailed experimental setup. The algorithms which we corral are UCB-I, Thompson sampling (TS), and FTRL with $\frac{1}{2}$ -Tsallis entropy regularizer (Tsallis-INF). We chose these algorithms as they all come with regret guarantees for the stochastic multi-armed problem and they broadly represent three different classes of algorithms, i.e, algorithms based on the optimism in the face of uncertainty principle, algorithms based on posterior sampling, and algorithms based on online mirror descent. As already discussed in Section 6, when implementing Algorithm 2 and Corral, we *never* restart the corralled algorithms and run them with their default parameters. Even though, there are no theoretical guarantees for this modification of the corraling algorithms, we will see that the regret bounds remain meaningful in practice. In all of the experiments we corral two instances of UCB-I, TS, and FTRL for a total of six algorithms. The best algorithm plays over 10 arms. Every other algorithm plays over 5 arms. Intuitively, the higher the number of arms implies higher complexity of the best algorithm which would lead to higher regret and a harder corraling problem. The rewards for each algorithm are Bernoulli random variables setup according to the following parameters: BASE_REWARD, IN_GAP, OUT_GAP, and LOW_REWARD. The best overall arm has expected reward BASE_REWARD + IN_GAP + OUT_GAP. Every other arm of Algorithm 1 has expected reward equal to LOW_REWARD. For all other algorithms the best arm has reward BASE_REWARD+IN_GAP and other arms have reward BASE_REWARD. In all of the experiments we set BASE_REWARD = 0.5, IN_GAP = 0.01, LOW_REWARD = 0.2. While a small IN_GAP implies a large regret for the algorithms containing sub-optimal arms, it also reduces the likelihood that said algorithms would have small average reward. Combined with setting LOW_REWARD = 0.2, this will make the average reward of $\mathcal{A}_1$ look small in the initial number of rounds, compared to the average reward of $\mathcal{A}_i, i > 1$ and hence makes the corraling problem harder. We run two set of experiments, an easy set for which OUT_GAP = 0.19, which translates to gaps $\Delta_i = 0.2$ in our regret bounds, and a hard set for which OUT_GAP = 0.01 which implies $\Delta_i = 0.02$. Finally time horizon is set to $T = 1000000$.

A.1 UCB-I contains best arm

Experiments can be found in Figure 1 for $\Delta_i = 0.2$ and in Figure 2 for $\Delta_i = 0.02$.


 Figure 1: UCB-I contains best arm, $\Delta_i = 0.2$, $ALG_{1:2} = \text{UCB-I}$, $ALG_{3:4} = \text{Tsallis-INF}$, $ALG_{5:6} = \text{TS}$.

 Figure 2: UCB-I contains best arm, $\Delta_i = 0.02$, $ALG_{1:2} = \text{UCB-I}$, $ALG_{3:4} = \text{Tsallis-INF}$, $ALG_{5:6} = \text{TS}$.

A.2 Tsallis-INF contains best arm

Experiments can be found in Figure 3 for $\Delta_i = 0.2$ and in Figure 4 for $\Delta_i = 0.02$.

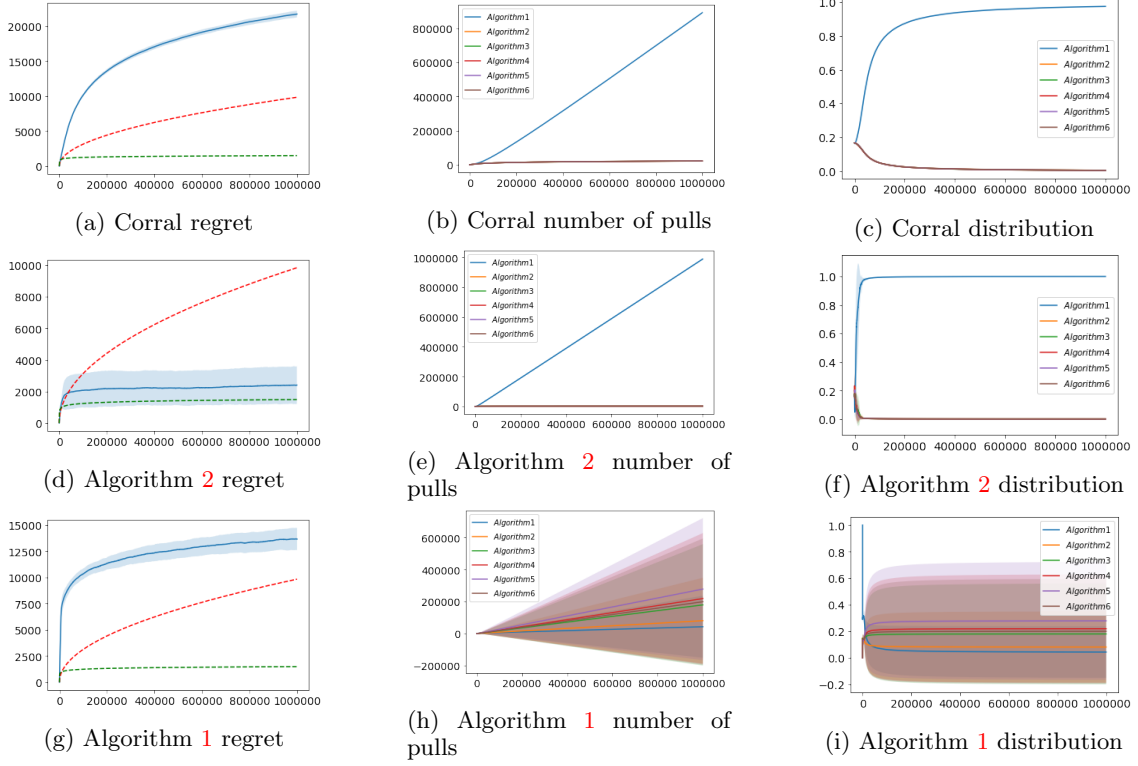


Figure 3: Tsallis-INF contains best arm, $\Delta_i = 0.2$, $ALG_{1:2} = \text{TSALLIS-INF}$, $ALG_{3:4} = \text{UCB-I}$, $ALG_{5:6} = \text{TS}$.

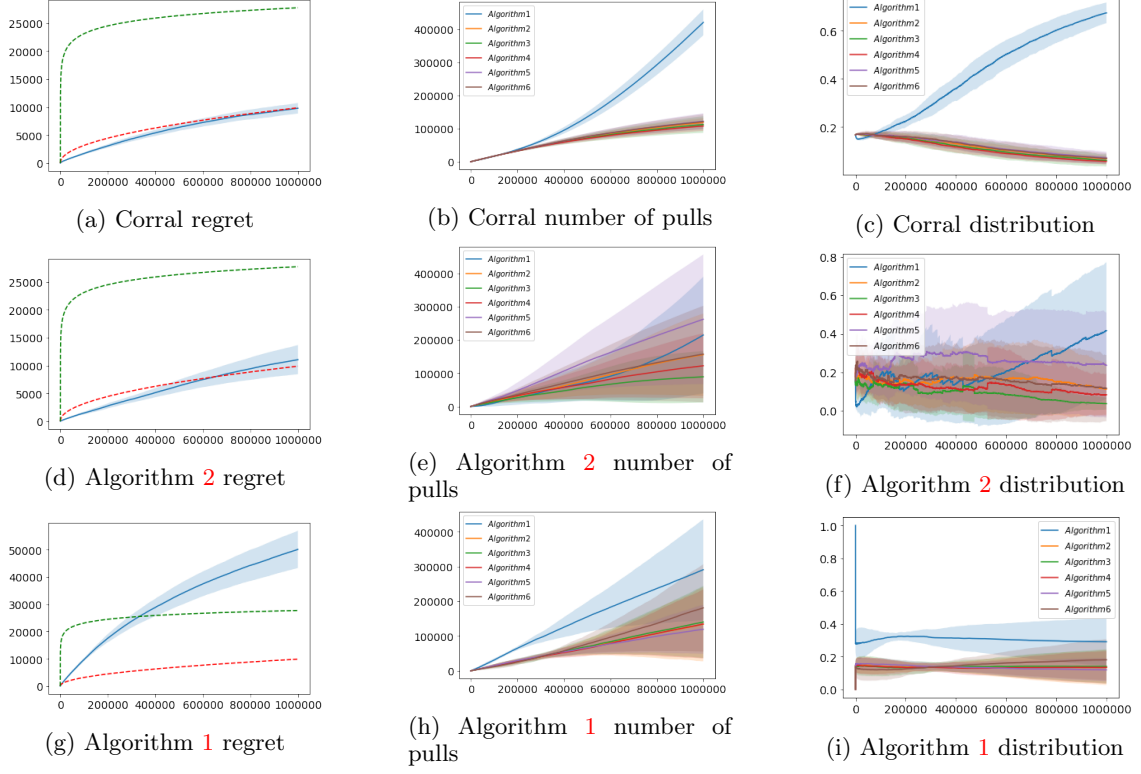


Figure 4: Tsallis-INF contains best arm, $\Delta_i = 0.02$, $\text{ALG}_{1:2} = \text{TSALLIS-INF}$, $\text{ALG}_{3:4} = \text{UCB-I}$, $\text{ALG}_{5:6} = \text{TS}$.

A.3 Thompson sampling contains best arm

Experiments can be found in Figure 5 for $\Delta_i = 0.2$ and in Figure 6 for $\Delta_i = 0.02$.

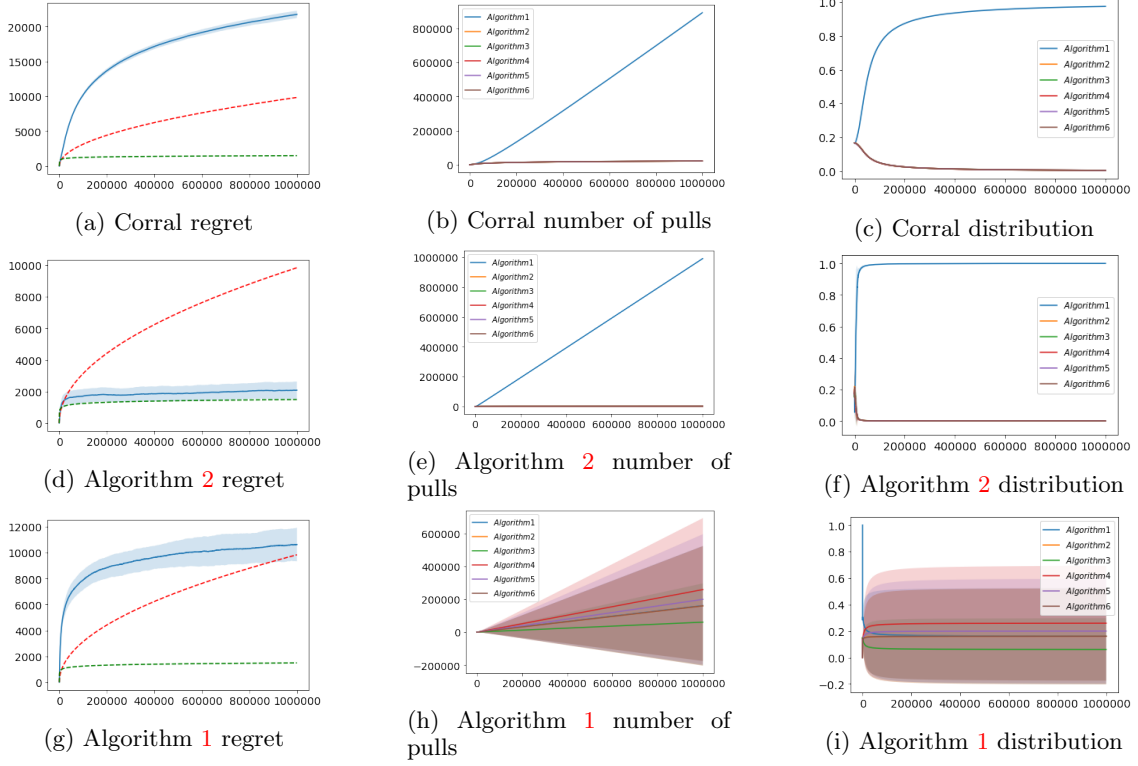


Figure 5: Thompson sampling (TS) contains best arm, $\Delta_i = 0.2$, $\text{ALG}_{1:2} = \text{TS}$, $\text{ALG}_{3:4} = \text{UCB-I}$, $\text{ALG}_{5:6} = \text{TSALLIS-INF}$.

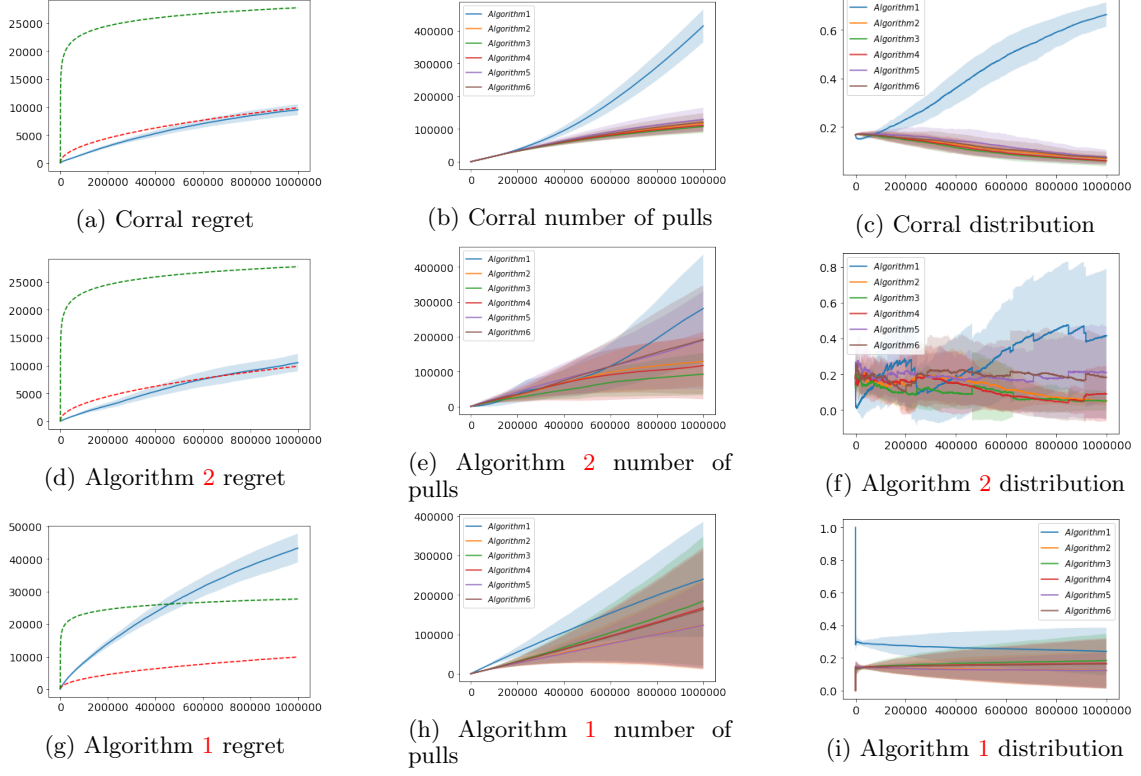


Figure 6: TS contains best arm, $\Delta_i = 0.02$, $\text{ALG}_{1:2} = \text{TS}$, $\text{ALG}_{3:4} = \text{UCB-I}$, $\text{ALG}_{5:6} = \text{TSALLIS-INF}$.

B Proofs from Section 3

We first introduce the formal construction briefly described in Section 3.

B.1 First lower bound

Assume that the corraling algorithm can play one of two algorithms, $\mathcal{A}_1$ or $\mathcal{A}_2$, with the rewards of each arm played by these algorithms distributed according to a Bernoulli random variable. Algorithm $\mathcal{A}_1$ plays a single arm with expected reward μ_1 and algorithm $\mathcal{A}_2$ is defined as follows.

Let β be drawn according to the Bernoulli distribution $\beta \sim \text{Ber}(\frac{1}{2})$ and let α be drawn uniformly over the unit interval, $\alpha \sim \text{Unif}[0, 1]$. If $\beta = 1$, $\mathcal{A}_2$ alternates between playing an arm with mean μ_2 and an arm with mean μ_3 every round, so that the algorithm incurs linear regret. We set μ_i s such that $\mu_2 > \mu_1 > \frac{\mu_2 + \mu_3}{2}$. If $\beta = 0$, then $\mathcal{A}_2$ behaves in the same way as if $\beta = 1$ for the first $T^{(1-\alpha)}$ rounds and for the remaining $T - T^{(1-\alpha)}$ rounds $\mathcal{A}_2$ only pulls the arm with mean μ_2 . Notice that, in this setting, $\mathcal{A}_2$ admits sublinear regret almost surely.

We denote by $\mathbb{P}(\cdot | r_1(a_{i_1, j_1}), \dots, r_t(a_{i_t, j_t}), \beta = i)$ the natural measure on the σ -algebra generated by the observed rewards under the environment $\beta = i$ and all the randomness of the player's algorithm. To simplify the notation, we denote by $r_{1:t}$ the sequence $\{r_s(a_{i_s, j_s})\}_{s=1}^t$. Let N denote the random variable counting the number of times the corraling strategy selected $\mathcal{A}_1$. Information-theoretically, the player can obtain a good approximation of μ_1 in time $O(\log(T))$ and, therefore, for simplicity, we assume that the player knows μ_1 exactly. Note that this can only make the problem easier for the player. Given this information, we can assume that the player begins by playing algorithm $\mathcal{A}_2$ for $T - N + 1$ rounds and then switches to $\mathcal{A}_1$ for the rest of the game. In particular, we assume that $T - N + 1$ is the time when the player can figure out that $\beta = 1$. We note that at time $T^{(1-\alpha)}$ we have $\mathbb{P}(\cdot | r_{1:T^{(1-\alpha)}}) = \mathbb{P}(\cdot | r_{1:T^{(1-\alpha)}})$, as the distribution of the rewards provided by $\mathcal{A}_2$ do not differ between $\beta = 1$ and $\beta = 0$. Furthermore, any random strategy would also need to select algorithm $\mathcal{A}_2$ at least $T^{(1-\alpha)} + 1$ rounds before it is able to distinguish between $\beta = 1$ or $\beta = 0$. It is also important to note that under the event that $\beta = 1$, the corraling algorithm does not receive any information about the value of α . This allows us to show that in the setting constructed above, with at least constant probability the best algorithm i.e., $\mathcal{A}_1$ when $\beta = 1$ and $\mathcal{A}_2$ when $\beta = 0$, has sublinear regret. Finally, a direct computation of the regret of this corraling strategy gives the following result.

Theorem B.1. *Let algorithms $\mathcal{A}_1$ and $\mathcal{A}_2$ follow the construction in Section 3. Then, with probability at least $1/2$ over the random choice of α , any corraling strategy incurs regret at least $\tilde{\Omega}(T)$, while the regret of the best algorithm is at most $O(\sqrt{T})$.*

Proof. Let $R(T)$ denote the regret of the corraling algorithm. Direct computation shows that if $\beta = 1$ the corraling regret is

$$\mathbb{E}[R(T) | \beta = 1, r_{1:T^{(1-\alpha)}}] = \mathbb{E} \left[\left(\mu_1 - \frac{\mu_2 + \mu_3}{2} \right) (T - N) | \beta = 1, r_{1:T^{(1-\alpha)}} \right]$$

Further if $\beta = 0$ and $\mathcal{A}_2$ is the best algorithm the regret of corraling is

$$\begin{aligned} \mathbb{E}[R(T) | \beta = 0, r_{1:T^{(1-\alpha)}}] &= \mathbb{E} \left[T^{(1-\alpha)} \mu_2 + (T - T^{(1-\alpha)}) \mu_2 | \beta = 0, r_{1:T^{(1-\alpha)}} \right] \\ &\geq \mathbb{E} \left[\frac{\mu_2 + \mu_3}{2} T^{(1-\alpha)} + \mu_2 (T - T^{(1-\alpha)}) \right. \\ &\quad \left. - \mu_1 N - \chi_{(N \leq T - T^{(1-\alpha)})} \left(\frac{\mu_2 + \mu_3}{2} T^{(1-\alpha)} + \mu_2 (T - T^{(1-\alpha)} - N) \right) \right. \\ &\quad \left. - \chi_{(N > T - T^{(1-\alpha)})} \frac{\mu_2 + \mu_3}{2} (T - N) | \beta = 0, r_{1:T^{(1-\alpha)}} \right], \end{aligned}$$

where the characteristic functions describe the event in which we pull $\mathcal{A}_1$ less times than is needed for $\mathcal{A}_2$ to switch to playing the best action. Notice that the total regret for corraling is at least the above as we also need to add the regret of the best algorithm to the above.

We first consider the case $\beta = 1$. Notice that in this case the corraling algorithm does not receive any information about α because $\mathcal{A}_2$ alternates between μ_2 and μ_3 at all rounds. This implies $\mathbb{E}[R(T) | \beta = 1, \alpha] = \mathbb{E}[R(T) | \beta = 1]$.

Condition on the event $N \leq T - T^{(1-\alpha)}$. We have

$$\begin{aligned} \mathbb{E}[R(T)|\beta = 1, N \leq T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha] &= \\ \mathbb{E}[R(T)|\beta = 1, N \leq T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}] &\geq \left(\mu_1 - \frac{\mu_2 + \mu_3}{2}\right) \mathbb{E}[T^{(1-\alpha)}|\beta = 1] \\ &= \left(\mu_1 - \frac{\mu_2 + \mu_3}{2}\right) \mathbb{E}[T^{(1-\alpha)}] \\ &= \left(\mu_1 - \frac{\mu_2 + \mu_3}{2}\right) \frac{T - 1}{\log(T)}, \end{aligned}$$

where in the first inequality we have replaced N by $T - T^{1-\alpha}$. Next consider the case $\beta = 0$. Condition on the event $N > T - T^{(1-\alpha)}$. We have

$$\begin{aligned} \mathbb{E}[R(T)|\beta = 0, N > T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha] &= \\ \mathbb{E}\left[\frac{\mu_2 - \mu_3}{2}(T - T^{(1-\alpha)}) - \left(\mu_1 - \frac{\mu_2 + \mu_3}{2}\right)N|\beta = 0, N > T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha\right] &= \\ \geq \mathbb{E}\left[(\mu_2 - \mu_1)T - \frac{\mu_2}{2}T^{(1-\alpha)}|\beta = 0, N > T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha\right] &= \\ = \mathbb{E}\left[(\mu_2 - \mu_1)T - \frac{\mu_2}{2}T^{1-\alpha}|\alpha\right], \end{aligned}$$

where in the inequality we have used the fact that $N > T - T^{1-\alpha}$ to bound $-\mu_1 N$ and $T \geq N$ to bound $\frac{\mu_1 + \mu_3}{2}N$. Let A denote the event $N \leq T - T^{(1-\alpha)}$. We are now ready to lower bound the regret of the player's strategy as follows.

$$\begin{aligned} \mathbb{E}[R(T)|\alpha] &= \frac{1}{2}\mathbb{E}[\mathbb{E}[R(T)|r_{1:T^{(1-\alpha)}}, \beta = 1, \alpha] + \mathbb{E}[R(T)|r_{1:T^{(1-\alpha)}}, \beta = 0, \alpha]|\alpha] \\ &\geq \frac{1}{2}\mathbb{E}[\mathbb{P}(A|r_{1:T^{(1-\alpha)}}, \beta = 1, \alpha)\mathbb{E}[R(T)|r_{1:T^{(1-\alpha)}}, \beta = 1, A, \alpha] \\ &\quad + \mathbb{P}(A^c|r_{1:T^{(1-\alpha)}}, \beta = 1, \alpha)\mathbb{E}[R(T)|r_{1:T^{(1-\alpha)}}, \beta = 0, A^c, \alpha]|\alpha] \\ &\geq \frac{1}{2}\mathbb{E}\left[\mathbb{P}(A|r_{1:T^{(1-\alpha)}}, \beta = 1, \alpha)\left(\mu_1 - \frac{\mu_2 + \mu_3}{2}\right)\frac{T - 1}{2\log(T)}\right. \\ &\quad \left.+ (1 - \mathbb{P}(A|r_{1:T^{(1-\alpha)}}, \beta = 1), \alpha)(\mu_2 - \mu_1)T - \frac{\mu_2}{2}T^{1-\alpha}|\alpha\right], \end{aligned}$$

where in the first inequality we have used the fact that the conditional measures induced by $\beta = 1$ and $\beta = 0$ are equal for the first $T^{1-\alpha}$ rounds. Because $\alpha \geq 1/2$ with probability at least $1/2$ it holds that the random variable $\mathbb{E}[R(T)|\alpha] > \tilde{\Omega}(T)$ with probability at least $1/2$ and that the regret of $\mathcal{A}_2$ when $\beta = 1$ is at most $O(\sqrt{T})$. $\square$

B.2 A realistic setting for Algorithm 2

The behavior of $\mathcal{A}_2$ for the setting given by $\beta = 0$, in the construction above, may seem somewhat artificial: a stochastic bandit algorithm may not be expected to behave in that manner when the gap between μ_2 and μ_3 is large enough. Here, we describe how to set μ_1 , μ_2 and μ_3 such that the successive elimination algorithm (Even-Dar et al., 2002) admits a similar behavior to $\mathcal{A}_2$ with $\beta = 0$. Recall that successive elimination needs at least $1/\Delta^2$ rounds to distinguish between the arm with mean μ_2 and the arm with mean μ_3 . In other words, for at least $1/\Delta^2$ rounds, it will alternate between the two arms. Therefore, we set $\frac{1}{\Delta^2} = T^{(1-\alpha)}$ or, equivalently, $\Delta = \frac{1}{T^{(1-\alpha)/2}}$, and $\mu_1 = \mu_2 - \frac{1}{4T^{(1-\alpha)/2}}$ to yield behavior similar to $\mathcal{A}_2$. For this construction, we show the following lower bound.

Theorem B.2 (Theorem 3.1 formal). *Let algorithms $\mathcal{A}_1$ and $\mathcal{A}_2$ follow the construction in Section B.2. With probability at least $1/4$ over the random choice of α any corraling strategy will incur regret at least $\tilde{\Omega}(\sqrt{T})$ while the gap between μ_2 and μ_3 is such that $\Delta > \omega(T^{-1/4})$ and hence the regret of the best algorithm is at most $o(T^{1/4})$.*

Proof. From the proof of Theorem B.1 we can compute, when $\beta = 1$, we can directly compute

$$\begin{aligned} & \mathbb{E} \left[R(T) | \beta = 1, N \leq T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha \right] \\ & \geq \mathbb{E} \left[\left(\mu_1 - \frac{\mu_2 + \mu_3}{2} \right) T^{(1-\alpha)} | \beta = 1, N \leq T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha \right] \\ & = \mathbb{E} \left[\frac{1}{4T^{(1-\alpha)/2}} T^{(1-\alpha)} | \beta = 1, N \leq T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}} \right] = \frac{\sqrt{T} - 1}{2 \log(T)}, \end{aligned}$$

Where in the equality we again used the fact that if $\beta = 1$, the corraling algorithm receives no information about α . Further when $\beta = 0$ we have

$$\begin{aligned} & \mathbb{E} \left[R(T) | \beta = 0, N > T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}(a_{T^{(1-\alpha)}}), \alpha \right] \\ & \geq \mathbb{E} \left[(\mu_2 - \mu_1)T - \frac{\mu_2}{2} T^{(1-\alpha)} | \beta = 0, N > T - T^{(1-\alpha)}, r_{1:T^{(1-\alpha)}}, \alpha \right] \\ & = \mathbb{E} \left[\frac{T^{(1+\alpha)/2}}{4} | \alpha \right] - \mathbb{E} \left[\frac{T^{(1-\alpha)}}{2} | \alpha \right]. \end{aligned}$$

Again we note that with probability $1/2$ we have $\alpha \geq 1/2$ and the above expression becomes asymptotically larger than $\sqrt{T}$. The same computation as in the proof of Theorem B.1 finishes the proof. $\square$

We note that, in our construction, if $\beta = 1$, then the inequality $\Delta \gg \frac{1}{\sqrt{T}}$ holds almost surely. In this setting, the instance-dependent regret bound for $\mathcal{A}_2$ and successive elimination is asymptotically smaller compared to the worst-case instance-independent regret bounds for stochastic bandit algorithms, which scale as $\tilde{O}(\sqrt{T})$ with the time horizon. This suggests that, even though $\mathcal{A}_2$ enjoys asymptotically better regret bounds than $\tilde{O}(\sqrt{T})$, the corraling algorithm will necessarily incur $\tilde{\Omega}(\sqrt{T})$ regret.

B.3 A lower bound when a worst case regret bound is known

Next, suppose that we know a worst case regret bound of $R_2(T)$ for algorithm $\mathcal{A}_2$. As before, we sample β according to a Bernoulli distribution. If $\beta = 1$, then algorithm $\mathcal{A}_2$ has a single arm with reward distributed as $\text{Ber}((\mu_2 + \mu_3)/2)$; in that case, $\mathcal{A}_2$ admits a regret equal to 0. If $\beta = 0$, then $\mathcal{A}_2$ has two arms distributed according to $\text{Ber}(\mu_2)$ and $\text{Ber}(\mu_3)$, respectively. We sample $\alpha \sim \text{Unif}[0, 1]$, and let $\mathcal{A}_2$ play an arm uniformly at random for the first $R_2(T)^{(1-\alpha)}$ rounds. In particular, during each of the first $R_2(T)^{(1-\alpha)}$ rounds, $\mathcal{A}_2$ plays with equal probability the arm with mean μ_2 and the arm with mean μ_3 . On round $R_2(T)^{(1-\alpha)}$, the algorithm switches to playing μ_1 until the rest of the game. Notice that the rewards up to time $R_2(T)^{(1-\alpha)}$, whether $\beta = 1$ or $\beta = 0$, have the same distribution. Hence, $\mathbb{P}(\cdot | r_{1:R_2(T)^{(1-\alpha)}}, \beta = 1) = \mathbb{P}(\cdot | r_{1:R_2(T)^{(1-\alpha)}}, \beta = 0)$. Then, following the arguments in the proof of Theorem B.1, we can prove the following lower bound.

Theorem B.3. *Let algorithms $\mathcal{A}_1$ and $\mathcal{A}_2$ follow the construction in Section B.3. Suppose that the worst case known regret bound for Algorithm is $R_2(T)$. With probability at least $1/2$ over the random choice of α any corraling strategy will incur regret at least $\Omega(R_2(T))$ while the regret of $\mathcal{A}_2$ is at most $O(\sqrt{R_2(T)})$.*

C Proofs from Section 4

Lemma C.1. *Suppose we run $2 \log(1/\delta)$ copies of algorithm $\mathcal{A}_i$ which satisfies Equation 2. Let $\mathcal{A}_{\text{med}_i}$ denote the algorithm with median reward at time t . Then,*

$$\mathbb{P} \left[t\mu_{\text{med}_i,1} - \sum_{s=1}^t r_s(a_{\text{med}_i,j_s}) \geq 2\bar{R}_{\text{med}_i}(t) \right] \leq \delta.$$

Proof of Lemma 4.1. First note that $\mu_{\text{med}_i,1} = \mu_{i,s,1}$ and $\bar{R}_{i_s}(t) = \bar{R}_{\text{med}_i}(t)$ for all s and t . The assumption in Equation 2 together with Markov's inequality implies that for every copy $\mathcal{A}_{i_s}$ of $\mathcal{A}_i$ at time t it holds that

$$\mathbb{P} \left[t\mu_{\text{med}_i,1} - \sum_{s=1}^t r_s(a_{i,j_s}) \geq 2\bar{R}_{\text{med}_i}(t) \right] \leq \frac{1}{2}.$$

Let $\mathcal{A}_{i_1}, \dots, \mathcal{A}_{i_n}$ be the algorithms which have reward smaller than $\mathcal{A}_{med_i}$ at time t . We have

$$\begin{aligned} \mathbb{P} \left[t\mu_{med_i,1} - \sum_{s=1}^t r_s(a_{med_i,j_s}) \geq 2\bar{R}_{med_i}(t) \right] &\leq \mathbb{P} \left[\bigcap_{l \in [n]} \left\{ t\mu_{l,1} - \sum_{s=1}^t r_s(a_{i_l,j_s}) \geq 2\bar{R}_{med_i}(t) \right\} \right] \\ &\leq \left(\frac{1}{2} \right)^{\log(1/\delta)} \leq \delta, \end{aligned}$$

where the first inequality follows from the definition of $\mathcal{A}_{med_i}$ and $\mathcal{A}_{i_l}$ for $l \in [n]$. $\square$

Theorem C.2. Suppose that algorithms $\mathcal{A}_1, \dots, \mathcal{A}_k$ satisfy the following regret bound $\mathbb{E}[R_i(t)] \leq \sqrt{\alpha k_i t \log(t)}$. Then after T rounds, Algorithm 1 produces a sequence of actions $a_1, \dots, a_T$, such that

$$\begin{aligned} T\mu_{1,1} - \mathbb{E} \left[\sum_{t=1}^T r_t(a_{i_t,j_t}) \right] &\leq O \left(\sum_{i \neq i^*} \frac{k_i \log(T)^2}{\Delta_i} + \log(T) \mathbb{E}[R_{i^*}(T)] \right), \\ T\mu_{1,1} - \mathbb{E} \left[\sum_{t=1}^T r_t(a_{i_t,j_t}) \right] &\leq O \left(\log(T) \sqrt{KT \log(T) \max_{i \in [K]}(k_i)} \right). \end{aligned}$$

Proof of Theorem 4.2. For simplicity we assume that $\lceil \log(T) \rceil = \log(T)$. For the rest of the proof we let $t_\ell = T_\ell(t)$ to simplify notation. Further, since $\bar{R}_{\ell_s} = \bar{R}_\ell, \forall s \in [\log(T)]$, we use $\bar{R}_\ell$ as the upper bound on the regret for all algorithms in $\mathbb{A}_\ell$. Let $\psi_\ell(t) = 2\sqrt{\frac{2 \log(t)}{t_\ell}} + \frac{\sqrt{2\bar{R}_\ell(t_\ell)}}{t_\ell}$. The proof follows the standard ideas behind analyses of UCB type algorithms. If at time t algorithm $\ell \neq 1$ is selected then one of the following must hold true:

$$\mu_{1,1} \geq \hat{\mu}_1(t_1) + \frac{\sqrt{2\bar{R}_1(t_1)} + \sqrt{2t_1 \log(t)}}{t_1}, \quad (5)$$

$$\hat{\mu}_{med_\ell}(t_\ell) > \mu_{1,1} + \sqrt{\frac{2 \log(t)}{t_\ell}}, \quad (6)$$

$$\Delta_\ell < 2\sqrt{\frac{2 \log(t)}{t_\ell}} + \frac{\sqrt{2\bar{R}_\ell(t_\ell)}}{t_\ell}. \quad (7)$$

The above conditions can be derived by considering the case when the UCB for $\mathcal{A}_1$ is smaller than the UCB for $\mathcal{A}_\ell$ and every algorithm has been selected a sufficient number of times. Suppose that the three conditions above are false at the same time. Then we have

$$\begin{aligned} \hat{\mu}_1(t_1) + \frac{\sqrt{2\bar{R}_1(t_1)} + \sqrt{2t_1 \log(t)}}{t_1} &> \mu_{1,1} = \mu_{1,\ell} + \Delta_\ell \\ &\geq \Delta_\ell + \hat{\mu}_{med_\ell}(t_\ell) - \sqrt{\frac{2 \log(t)}{t_\ell}} \\ &\geq \hat{\mu}_{med_\ell}(t_\ell) + \frac{\sqrt{2\bar{R}_\ell(t_\ell)} + \sqrt{2t_\ell \log(t)}}{t_\ell}, \end{aligned}$$

which contradicts the assumption that algorithm $\mathcal{A}_\ell$ was selected. With slight abuse of notation we use $[k_\ell]$ to denote the set of arms belonging to algorithm $\mathcal{A}_\ell$. Next we bound the expected number of times each sub-optimal algorithm is played up to time T . Let δ be an upper bound on the probability of the event that $\hat{\mu}_1(s)$ exceeds

the UCB for $\mathcal{A}_1$.

$$\begin{aligned}
 \mathbb{E}[T_\ell] &= \sum_{t=1}^T \mathbb{E}[\chi_{a_{i_t, j_t} \in [k_\ell]}] \leq \psi_\ell^{-1}(\Delta_\ell) + \sum_{t > \psi_\ell^{-1}(\Delta_\ell)} \mathbb{P}[\text{Equation 5 or Equation 6 hold}] \\
 &\leq \psi_\ell^{-1}(\Delta_\ell) + \sum_{t > \psi_\ell^{-1}(\Delta_\ell)} \mathbb{P}\left[\exists s \in [t] : \mu_{1,1} \geq \hat{\mu}_1(s) + \frac{\sqrt{2\bar{R}_1(s)} + \sqrt{2s \log(t)}}{s}\right] \\
 &\quad + \sum_{t > \psi_\ell^{-1}(\Delta_\ell)} \mathbb{P}\left[\exists s \in [t] : \hat{\mu}_{med_\ell}(s) > \mu_{1,1} + \sqrt{\frac{2 \log(t)}{s}}\right] \\
 &\leq \sum_{t > \psi_\ell^{-1}(\Delta_\ell)} t\delta + \sum_{t > \psi_\ell^{-1}(\Delta_\ell)} \frac{1}{t} + \psi^{-1}(\Delta_\ell),
 \end{aligned}$$

where the last inequality follows from the definition of δ and the fact that $\hat{\mu}_{med_\ell}(s) \leq \hat{\mu}_{med_\ell,1}(s)$ (empirical mean of arm 1 for algorithm $\mathcal{A}_{med_\ell}$ at time s) and the standard argument in the analysis of UCB-I. Setting $\delta = \frac{1}{T^2}$ finishes the bound on the number of suboptimal algorithm pulls. Next we consider bounding the regret incurred only by playing the median algorithms $\mathcal{A}_{med_\ell}$

$$\begin{aligned}
 t\mu_{1,1} - \mathbb{E}\left[\sum_{s=1}^t r_s(a_{i_s, j_s})\right] &= t\mu_{1,1} - \mathbb{E}\left[\sum_{\ell} t_\ell \mu_{\ell,1} + \sum_{\ell} \sum_i T_{med_\ell, i}(t_\ell) \mu_{\ell, i} - \sum_{\ell} t_\ell \mu_{\ell,1}\right] \\
 &= \mathbb{E}\left[\sum_{\ell \neq 1} t_\ell \Delta_\ell\right] + \sum_{\ell \neq 1} \mathbb{E}[R_{med_\ell}(t_\ell)] + \mathbb{E}[R_{med_1}(t)] \\
 &\leq \sum_{\ell \neq 1} \Delta_\ell \psi^{-1}(\Delta_\ell) + 2 \log(T) + \sum_{\ell \neq 1} \mathbb{E}[\sqrt{\alpha k_\ell t_\ell \log(t)}] + \mathbb{E}[R_{med_1}(t)] \\
 &\leq \sum_{\ell \neq 1} \Delta_\ell \psi^{-1}(\Delta_\ell) + 2 \log(T) + \sum_{\ell \neq 1} \sqrt{\alpha k_\ell \mathbb{E}[t_\ell] \log(t)} + \mathbb{E}[R_{med_1}(t)] \\
 &\leq \sum_{\ell \neq 1} \Delta_\ell \psi^{-1}(\Delta_\ell) + 2 \log(T) + \sum_{\ell \neq 1} \sqrt{\alpha k_\ell \psi^{-1}(\Delta_\ell) \log(t)} \\
 &\quad + \mathbb{E}[R_{med_1}(t)] + \sqrt{\alpha k_\ell \log(t)}.
 \end{aligned}$$

Now for the assumed regret bound on the algorithms, we have $\psi_\ell(t) = 2\sqrt{\frac{2 \log(t)}{t_\ell}} + \sqrt{2\frac{\alpha k_\ell \log(t_\ell)}{t_\ell}}$. This implies that $\psi_\ell^{-1}(\Delta_\ell) \leq \frac{\alpha' k_\ell \log(t)}{\Delta_\ell^2}$, for some other constant α' . To get the instance independent bound we first notice that by Jensen's inequality we have

$$\begin{aligned}
 \sum_{\ell} \sqrt{\alpha' k_\ell \mathbb{E}[t_\ell] \log(t)} &\leq K \sqrt{\frac{1}{K} \sum_{\ell} \alpha' \mathbb{E}[t_\ell] k_\ell \log(t)} \\
 &\leq \sqrt{\alpha' K t \log(t) \max_{\ell}(k_\ell)}.
 \end{aligned}$$

Next we can bound $\mathbb{E}\left[\sum_{\ell \neq 1} t_\ell \Delta_\ell\right]$ in the following way

$$\begin{aligned}
 \mathbb{E}\left[\sum_{\ell \neq 1} t_\ell \Delta_\ell\right] &\leq \sum_{\ell} \Delta_\ell \sqrt{\mathbb{E}[t_\ell]} \sqrt{\mathbb{E}[t_\ell]} = \sum_{\ell} \sqrt{\Delta_\ell^2 \mathbb{E}[t_\ell]} \sqrt{\mathbb{E}[t_\ell]} \\
 &= \sum_{\ell} \sqrt{\alpha' k_\ell \mathbb{E}[t_\ell] \log(t)} \leq \sqrt{\alpha' K t \log(t) \max_{\ell}(k_\ell)}
 \end{aligned}$$

The theorem now follows. $\square$

C.1 Proof of Theorem 4.3

Consider an instance of Algorithm 1, except that it runs a single copy of each base learner $\mathcal{A}_i$. Let $\mathcal{A}_1$ be a UCB algorithm with two arms with means $\mu_1 > \mu_2$, respectively. The arm with mean μ_1 is set according to a Bernoulli random variable, and the arm with mean μ_2 is deterministic. Let algorithm $\mathcal{A}_2$ have a single deterministic arm with mean μ_3 , such that $\mu_1 > \mu_3$ and $\mu_3 > \mu_2$. Let $\Delta = \mu_1 - \mu_3$. We now follow the lower bounding technique of Audibert et al. (2009).

Consider the event that in the first q pulls of arm $a_{1,1}^{A_1}$, we have $r_t(a_{1,1}) = 0$, i.e. $\mathcal{E} = \{r_1(a_{1,1}) = 0, r_2(a_{1,1}) = 0, \dots, r_q(a_{1,1}) = 0\}$. This event occurs with probability $(1 - \mu_1)^q$. Notice that on event $\mathcal{E}$, the upper confidence bound for μ_1 as per $\mathcal{A}_1$ is $\sqrt{\frac{\alpha \log(T_1(t))}{q}}$ during time t . This implies that for $a_{1,1}$ to be pulled again we need $\sqrt{\frac{\alpha \log(T_1(t))}{q}} > \mu_2$ and hence for the first $\exp(q\mu_2^2/\alpha)$ rounds in which $\mathcal{A}_1$ is selected by the corraling algorithm, $a_{1,1}$ is only pulled q times. Further, on $\mathcal{E}$, the upper confidence bound for $\mathcal{A}_1$ as per the corraling algorithm is of the form $\sqrt{\frac{2\beta \log(t)}{T_1(t)}}$. This implies that for $\mathcal{A}_1$ to be selected again we need $\mu_2 + \sqrt{\frac{2\beta \log(t)}{T_1(t)}} > \mu_3$. Let $\tilde{\Delta} = \mu_3 - \mu_2$. Then, the above implies that in the first $t \leq \exp(T_1(t)\tilde{\Delta}^2/(2\beta))$ rounds, $\mathcal{A}_1$ is pulled at most $T_1(t)$ times. Combining with the bound for the number of pulls of $a_{1,1}$ we arrive at the fact that on $\mathcal{E}$, $a_{1,1}$ can not be pulled more than q times in the first $\exp\left(\frac{\tilde{\Delta}^2 \exp(q\mu_2^2/\alpha)}{2\beta}\right)$ rounds. Let q be large enough so that $q \leq \frac{1}{2} \exp\left(\frac{\tilde{\Delta}^2 \exp(q\mu_2^2/\alpha)}{2\beta}\right)$. Then, for large enough T , we have that the pseudo-regret of the corraling algorithm is $\hat{R}(T) \geq \frac{1}{2} \Delta \exp\left(\frac{\tilde{\Delta}^2 \exp(q\mu_2^2/\alpha)}{2\beta}\right)$. Taking $q = \log\left(\frac{2\beta}{\tilde{\Delta}^2} \log\left(\tau \frac{\alpha}{\mu_2^2}\right)\right)$, we get

$$\mathbb{P}\left[\hat{R}(T) \geq \frac{1}{2} \Delta \tau\right] \geq \mathbb{P}[\mathcal{E}] = (1 - \mu_1)^q = \frac{1}{\exp(q)^{\log(1/(1-\mu_1))}} = \left(\frac{\tilde{\Delta}^2}{2\beta \log(\tau)}\right)^{\frac{\alpha}{\mu_2^2} \log(1/(1-\mu_1))}.$$

Let $\gamma = \frac{\alpha}{\mu_2^2} \log(1/(1 - \mu_1))$. We can now bound the expected pseudo-regret of the algorithm by integrating over $2 \leq \tau \leq T$, to get

$$\begin{aligned} \mathbb{E}[\hat{R}(T)] &\geq \frac{1}{2} \Delta \int_2^T \left(\frac{\tilde{\Delta}^2}{2\beta \log(\tau)}\right)^{\frac{\alpha}{\mu_2^2} \log(1/(1-\mu_1))} d\tau = \frac{1}{2} \Delta \left(\frac{\tilde{\Delta}^2}{2\beta}\right)^\gamma \int_2^T \left(\frac{1}{\log(\tau)}\right)^\gamma d\tau \\ &\geq \frac{1}{2} \Delta \left(\frac{\tilde{\Delta}^2}{2\beta}\right)^\gamma \frac{T-2}{(\log(\frac{T+2}{2}))^\gamma}, \end{aligned}$$

where the last inequality follows from the Hermite-Hadamart inequality.

It is important to note that the above reasoning will fail if γ is a function of T . This might occur if in the UCB for $\mathcal{A}_1$ we have $\alpha = \log(T)$. In such a case the lower bounds become meaningless as $\frac{1}{\log((T+2)/2)}^\gamma \leq o(1/T)$. Further, it should actually be possible to avoid boosting in this case as the tail bound of the regret will now be upper bounded as $\mathbb{P}[R_1(t) \geq \Delta \tau] \leq \frac{1}{T^c}$.

General Approach if Regret has a Polynomial Tail. Assume that, in general, the best algorithm has the following regret tail:

$$\mathbb{P}\left[R_1(t) \geq \frac{1}{2} \Delta_{1,1} \tau\right] \geq \frac{1}{\tau^c},$$

for some constant c . Results in Salomon and Audibert (2011) suggest that for stochastic bandit algorithms which enjoy anytime regret bounds we can not have a much tighter high probability regret bound. Let $\mathcal{E}_{T_1(t)} = \{R_1(T_1(t)) \geq T_1(t)(\mu_1 - \frac{1}{\sqrt{2}}\mu_3)\}$. After $T_1(t)$ pulls of $\mathcal{A}_1$ the reward plus the UCB for $\mathcal{A}_1$ is at most $\frac{\sum_{s=1}^{T_1(t)} r_s(a_{1,j_s})}{T_1(t)} + \sqrt{\frac{\alpha k_1 \log(t)}{T_1(t)}}$, and on $\mathcal{E}_{T_1(t)}$, we have $\frac{\sum_{s=1}^{T_1(t)} r_s(a_{1,j_s})}{T_1(t)} \leq \frac{1}{\sqrt{2}}\mu_3$. This implies that in the first t rounds, $\mathcal{A}_1$ could not have been pulled more than

$$\frac{\alpha k_1 \log(t)}{\left(\mu_3 - \frac{\sum_{s=1}^{T_1(t)} r_s(a_{1,j_s})}{T_1(t)}\right)^2} \leq \frac{2\alpha k_1 \log(t)}{\mu_3^2}.$$

Setting $T_1(t) = \frac{2\alpha k_1 \log(T)}{\mu_3^2}$, we have that $\mathcal{E}_{T_1(t)}$ occurs with probability at least $\left(\frac{\Delta_{1,2}\mu_3^2}{4\alpha k_1 \log(T)}\right)^c$ and hence the expected regret of the corralling algorithm is at least

$$\left(\frac{\Delta_{1,2}\mu_3^2}{4\alpha k_1 \log(T)}\right)^c \Delta \left(T - \frac{2\alpha k_1 \log(T)}{\mu_3^2}\right).$$

We have just showed the following.

Theorem C.3. *There exist instances $\mathcal{A}_1$ and $\mathcal{A}_2$ of UCB-I and a reward distribution, such that if Algorithm 1 runs a single copy of $\mathcal{A}_1$ and $\mathcal{A}_2$ the expected regret of the algorithm is at least*

$$\mathbb{E}[R(T)] \geq \tilde{\Omega}(\Delta T).$$

Further, for any algorithm $\mathcal{A}_1$ such that $\mathbb{P}[R_1(t) \geq \frac{1}{2}\Delta_{1,1}\tau] \geq \frac{1}{\tau^c}$, there exists a reward distribution such that if Algorithm 1 runs a single copy of $\mathcal{A}_1$ the expected regret of the algorithm is at least

$$\mathbb{E}[R(T)] \geq \tilde{\Omega}((\Delta_{1,2})^c \Delta T).$$

D Proof of Theorem 5.2

D.1 Potential function and auxiliary lemmas

First we recall the definition of conjugate of a convex function f , denoted as f^*

$$f^*(y) = \max_{x \in \mathbb{R}^d} \langle x, y \rangle - f(x).$$

In our algorithm, we are going to use the following potential at time t

$$\begin{aligned} \Psi_t(w) &= -4 \sum_{i=1}^K \frac{\sqrt{w_i} - \frac{1}{2}w_i}{\eta_{t,i}} \\ \nabla \Psi_t(w)_i &= -2 \frac{\frac{1}{\sqrt{w_i}} - 1}{\eta_{t,i}} \\ \nabla^2 \Psi_t(w)_{i,i} &= \frac{1}{w_i^{3/2} \eta_{t,i}}, \nabla^2 \Psi_t(w)_{i,j} = 0 \\ \nabla \Psi_t^*(Y)_i &= \frac{1}{(-\frac{\eta_{t,i}}{2} Y_i + 1)^2} \\ \Phi_t(Y) &= \max_{w \in \Delta^{K-1}} \langle Y, w \rangle - \Psi_t(w) = (\Psi_t + \mathbf{I}_{\Delta^{K-1}})^*(Y). \end{aligned} \tag{8}$$

Further for a function f we use $D_f(x, y)$ to denote the Bregman divergence between x and y induced by f equal to

$$D_f(x, y) = f(x) - f(y) - \langle \nabla f(y), x - y \rangle = f(x) + f^*(\nabla f(y)) - \langle \nabla f(y), x \rangle,$$

where the second inequality follows by the Fenchel duality equality $f^*(\nabla f(y)) + f(y) = \langle \nabla f(y), y \rangle$. We now

present a bandit algorithm is going to be the basis for the corralling algorithm. Let $\eta_t = \begin{pmatrix} \eta_{t,1} \\ \eta_{t,2} \\ \vdots \\ \eta_{t,K} \end{pmatrix}$ be the step

size schedule for time t . The algorithm proceeds in epochs. Each epoch is twice as large as the preceding and the step size schedule remains non-increasing throughout the epochs, except when an OMD step is taken. In each epoch the algorithm makes a choice to either take two mirror descent steps, while increasing the step size:

$$\begin{aligned} \hat{w}_{t+1} &= \operatorname{argmin}_{w \in \Delta^{K-1}} \langle \hat{\ell}_t, w \rangle + D_{\Psi_t}(w, w_t), \\ \eta_{t+1,i} &= \beta \eta_{t,i} \quad \text{for } i : w_{t,i} \leq 1/\rho_{s_i}, \\ \hat{w}_{t+2} &= \operatorname{argmin}_{w \in \Delta^{K-1}} \langle \hat{\ell}_{t+1}, w \rangle + D_{\Psi_{t+1}}(w, \hat{w}_{t+1}), \\ \rho_{s_i} &= 2\rho_{s_i}, \end{aligned} \tag{9}$$

or the algorithm takes a FTRL step

$$w_{t+1} = \operatorname{argmin}_{w \in \Delta^{K-1}} \langle \widehat{L}_t, w \rangle + \Psi_{t+1}(w), \quad (10)$$

where $\widehat{L}_t = \widehat{L}_{t-1} + \widehat{\ell}_t$ unless otherwise specified by the algorithm. We note that the algorithm can only increase the step size during the OMD step. For technical reasons we require a FTRL step after each OMD step. Further we require that the second step of each epoch be an OMD step, if there exists at least one $w_{t,i} \leq \frac{1}{\rho_1}$. The algorithm also can enter an OMD step during an epoch if at least one $w_{t,i} \leq \frac{1}{\rho_{n_i}}$. The intuition behind this behavior is as follows. Increasing the step size and doing an OMD step will give us negative regret during that round and we only require negative regret for a certain arm if the probability of pulling said arm becomes smaller than some threshold. The pseudo-code can be found in Algorithm 2. For the rest of the proofs and discussion we denote an iterate from the FTRL update as w_t and an iterate from the OMD update as $\widehat{w}_t$. Further, intermediate iterates of OMD are denoted as $\tilde{w}_t$. We now present a couple of auxiliary lemmas useful for analyzing the OMD and FTRL updates.

Lemma D.1. *For any $x, y \in \Delta^{K-1}$ it holds*

$$D_{\Psi_t}(x, y) = D_{\Phi_t}(\nabla \Phi_t^*(y), \nabla \Phi_t^*(x)).$$

Proof. Since $\Psi_t + \mathbf{I}_{\Delta^{K-1}}$ is a convex, closed function on Δ^{K-1} it holds that $\Psi_t + \mathbf{I}_{\Delta^{K-1}} = ((\Psi_t + \mathbf{I}_{\Delta^{K-1}})^*)^*$ (see for e.g. (Brezis, 2010) Theorem 1.11). Further, $\Phi_t^*(x) = ((\Psi_t + \mathbf{I}_{\Delta^{K-1}})^*)^*(x) = \Psi_t(x)$. The above implies

$$D_{\Psi_t}(x, y) = D_{\Phi_t^*}(x, y) = D_{\Phi_t}(\nabla \Phi_t^*(y), \nabla \Phi_t^*(x)).$$

□

Lemma D.2. *For any positive $\widehat{L}_t$ and w_{t+1} generated according to update 10 we have*

$$w_{t+1} = \nabla \Phi_{t+1}(-\widehat{L}_t) = \nabla \Psi_{t+1}^*(-\widehat{L}_t + \nu_{t+1} \mathbf{1}),$$

for some scalar ν_t . Further $(\widehat{L}_t - \nu_{t+1} \mathbf{1})_i > 0$ for all $i \in [K]$.

Proof. The proof is contained in Section 4.3 in Zimmert and Seldin (2018). □

Lemma D.3 (Lemma 16 Zimmert and Seldin (2018)). *Let $w \in \Delta^{K-1}$ and $\tilde{w} = \nabla \Psi_t^*(\nabla \Psi_t(w) - \ell)$. If $\eta_{t,i} \leq \frac{1}{4}$, then for all $\ell > -1$ it holds that $\tilde{w}_i^{3/2} \leq 2w_i^{3/2}$.*

D.2 Regret bound

We begin by studying the instantaneous regret of the FTRL update. The bound follows the one in Zimmert and Seldin (2018). Let $u = e_{i^*}$ be the unit vector corresponding to the optimal algorithm $\mathcal{A}_{i^*}$. First we decompose the regret into a stability term and a penalty term:

$$\begin{aligned} \langle \widehat{\ell}_t, w_t - u \rangle &= \langle \widehat{\ell}_t, w_t \rangle + \Phi_t(-\widehat{L}_t) - \Phi_t(-\widehat{L}_{t-1}) \quad (\text{Stability}) \\ &\quad - \Phi_t(-\widehat{L}_t) + \Phi_t(-\widehat{L}_{t-1}) - \langle \widehat{\ell}_t, u \rangle \quad (\text{Penalty}). \end{aligned}$$

The bound on the stability term follows from Lemma 11 in Zimmert and Seldin (2018), however, we will show this carefully, since parts of the proof will be needed to bound other terms. Recall the definition of $\Phi_t(Y) = \max_{w \in \Delta^{K-1}} \langle Y, w \rangle - \Psi_t(w)$. Since w is in the simplex we have $\Phi_t(Y + \alpha \mathbf{1}_k) = \max_{w \in \Delta^{K-1}} \langle Y, w \rangle + \langle \alpha \mathbf{1}, w \rangle - \Psi_t(w) = \Phi_t(Y) + \alpha$. We also note that from Lemma D.2 it follows that we can write $\nabla \Psi_t(w_t) = -\widehat{L}_{t-1} + \nu_t \mathbf{1}$. Combining

Algorithm 4 Corralling with Tsallis-INF

Input: Mult. constant β , thresholds $\{\rho_i\}_{i=1}^n$, initial step size η , epochs $\{\tau_i\}_{i=1}^m$, algorithms $\{\mathcal{A}_i\}_{i=1}^K$.

Output: Algorithm selection sequence $(i_t)_{t=1}^T$.

```

1: Initialize  $t = 1$ ,  $w_1 = \text{Unif}(\Delta^{K-1})$ ,  $\eta_1 = \eta$ 
2: Initialize current threshold list  $\theta \in [n]^K$  to  $\mathbf{1}$ 
3: while  $t \leq \log(T) + 1$ 
4:   for  $i \in [K]$ 
5:     Algorithm  $i$  plays action  $a_{i,j_t}$  and  $\hat{L}_{1,i} += \ell_t(a_{i,j_t})$ 
6:    $t += 1$ 
7:  $t = 2$ ,  $w_2 = \nabla \Phi_2(-\hat{L}_1)$ ,  $1/\eta_{t+1}^2 = 1/\eta_t^2 + 1$ 
8: while  $j \leq m$ 
9:   for  $t \in \tau_j$ 
10:     $\mathcal{R}_t = \emptyset$ ,  $\hat{\ell}_t = \text{PLAY-ROUND}(w_t)$ 
11:    if  $t$  is the first round of epoch  $\tau_j$  and  $\exists w_{t,i} \leq \frac{1}{\rho_1}$ 
12:      for  $i: w_{t,i} \leq \frac{1}{\rho_1}$ 
13:         $\theta_i = \min\{s \in [n]: w_{t,i} > \frac{1}{\rho_s}\}$ ,  $\mathcal{R}_t = \mathcal{R}_t \cup \{i\}$ .
14:       $(w_{t+3}, \hat{L}_{t+2}) = \text{NEG-REG-STEP}(w_t, \hat{\ell}_t, \eta_t, \mathcal{R}_t, \hat{L}_{t-1})$ ,  $t = t + 2$ ,  $\hat{\ell}_t = \text{PLAY-ROUND}(w_t)$ 
15:      if  $\exists i: w_{t,i} \leq \frac{1}{\rho_{\theta_i}}$  and prior step was not NEG-REG-STEP
16:        for  $i: w_{t,i} \leq \frac{1}{\rho_{\theta_i}}$ 
17:           $\theta_i += 1$ ,  $\mathcal{R}_t = \mathcal{R}_t \cup \{i\}$ .
18:         $(w_{t+3}, \hat{L}_{t+2}) = \text{NEG-REG-STEP}(w_t, \hat{\ell}_t, \eta_t, \mathcal{R}_t, \hat{L}_{t-1})$ ,  $t = t + 2$ ,  $\hat{\ell}_t = \text{PLAY-ROUND}(w_t)$ 
19:      else
20:         $1/\eta_{t+1}^2 = 1/\eta_t^2 + 1$ ,  $w_{t+1} = \nabla \Phi_{t+1}(-\hat{L}_t)$ 

```

the two facts we have

$$\begin{aligned}
 \langle \ell_t, w_t \rangle + \Phi_t(-\hat{L}_t) - \Phi_t(-\hat{L}_{t-1}) &= \langle \ell_t, w_t \rangle + \Phi_t(\nabla \Psi_t(w_t) - \hat{\ell}_t - \nu_t \mathbf{1}) - \Phi_t(\nabla \Psi_t(w_t) - \nu_t \mathbf{1}) \\
 &= \langle \ell_t - \alpha \mathbf{1}_k, w_t \rangle + \Phi_t(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k) - \Phi_t(\nabla \Psi_t(w_t)) \\
 &\leq \langle \ell_t - \alpha \mathbf{1}_k, w_t \rangle + \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k) - \Psi_t^*(\nabla \Psi_t(w_t)) \\
 &= D_{\Psi_t^*}(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k, \nabla \Psi_t(w_t)) \\
 &\leq \max_{z \in [\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k, \nabla \Psi_t(w_t)]} \frac{1}{2} \|\hat{\ell}_t - \alpha \mathbf{1}\|_{\nabla^2 \Psi_t^*(z)}^2 \\
 &= \max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k)]} \frac{1}{2} \|\hat{\ell}_t - \alpha \mathbf{1}\|_{\nabla^2 \Psi_t^{-1}(w)}^2,
 \end{aligned}$$

where the first inequality holds since $\Psi_t^* \geq \Phi_t$ and $\Psi_t^*(\nabla \Psi_t(w_t)) = \langle \nabla \Psi_t(w_t), w_t \rangle - \Psi_t(w_t) = \Phi_t(\nabla \Psi_t(w_t))$ and the second inequality follows since by Taylor's theorem there exists a z on the line segment between $\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k$ and $\nabla \Psi_t(w_t)$ such that $D_{\Psi_t^*}(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k, \nabla \Psi_t(w_t)) = \frac{1}{2} \|\hat{\ell}_t - \alpha \mathbf{1}\|_{\nabla^2 \Psi_t^*(z)}^2$.

Lemma D.4. Let $w_t \in \Delta^{K-1}$ and let $i_t \sim w_t$. Let $\hat{\ell}_{t,i_t} = \frac{\ell_{t,i_t}}{w_{t,i_t}}$ and $\hat{\ell}_{t,i} = 0$ for all $i \neq i_t$. It holds that

$$\begin{aligned}
 \mathbb{E} \left[\max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k)]} \|\hat{\ell}_t\|_{\nabla^2 \Psi_t^{-1}(w)}^2 \right] &\leq \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} \\
 \mathbb{E} \left[\max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k)]} \|\hat{\ell}_t - \chi_{(i_t=j)} \ell_{t,j} \mathbf{1}\|_{\nabla^2 \Psi_t^{-1}(w)}^2 \right] &\leq \sum_{i \neq j} \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} + \frac{\eta_{t,i} + \eta_{t,j}}{2} \mathbb{E}[w_{t,i}].
 \end{aligned}$$

Algorithm 5 NEG-REG-STEP

Input: Previous iterate w_t , current loss $\widehat{\ell}_t$, step size η_t , set of rescaled step-sizes $\mathcal{R}_t$, cumulative loss $\widehat{L}_{t-1}$

Output: Plays two rounds of the game and returns distribution w_{t+3} and cumulative loss $\widehat{L}_{t+2}$

- 1: $(w_{t+1}, \widehat{L}_t) = \text{OMD-STEP}(w_t, \widehat{\ell}_t, \eta_t, \mathcal{R}_t, \widehat{L}_{t-1})$
- 2: $\widehat{\ell}_{t+1} = \text{PLAY-ROUND}(w_{t+1}), \widehat{L}_{t+1} = \widehat{L}_t + \widehat{\ell}_{t+1}$
- 3: **for** $i \in \mathcal{R}_t$
- 4: $\eta_{t+2,i} = \beta \eta_{t,i}$ and restart $\mathcal{A}_i$ with updated environment $\theta_i = \frac{2}{w_{t,i}}$
- 5: $w_{t+2} = \nabla \Phi_{t+2}(-\widehat{L}_{t+1})$
- 6: $\widehat{\ell}_{t+2} = \text{PLAY-ROUND}(w_{t+2})$
- 7: $\widehat{L}_{t+2} = \widehat{L}_{t+1} + \widehat{\ell}_{t+2}, \eta_{t+3} = \eta_{t+2}, t = t + 2$
- 8: $w_{t+1} = \nabla \Phi_{t+1}(-\widehat{L}_t), t = t + 1$

Algorithm 6 OMD-STEP

Input: Previous iterate w_t , current loss $\widehat{\ell}_t$, step size η_t , set of rescaled step-sizes $\mathcal{R}_t$, cumulative loss $\widehat{L}_{t-1}$

Output: New iterate w_{t+1} , cumulative loss $\widehat{L}_t$

- 1: $\nabla \Psi_t(\tilde{w}_{t+1}) = \nabla \Psi_t(w_t) - \widehat{\ell}_t$
- 2: $w_{t+1} = \arg\min_{w \in \Delta^{K-1}} D_{\Phi_t}(w, \tilde{w}_{t+1})$.
- 3: $\mathbf{e} = \sum_{i \in \mathcal{R}_t} \mathbf{e}_i$
- 4: $\tilde{L}_{t-1} = (\mathbf{1}_k - \mathbf{e}) \odot (\widehat{L}_{t-1} - (\nu_{t-2} + \nu_{t-1})\mathbf{1}_k) + \frac{1}{\beta} \mathbf{e} \odot ((\widehat{L}_{t-1} - (\nu_{t-2} + \nu_{t-1})\mathbf{1}_k)) // \nu_{t-2}$ and ν_{t-1} are the Lagrange multipliers from the previous two FTRL steps.
- 5: $\widehat{L}_t = \tilde{L}_{t-1} + \widehat{\ell}_t$

Proof. First notice that:

$$\begin{aligned} & \mathbb{E} \left[\max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \widehat{\ell}_t + \alpha \mathbf{1}_k)]} \|\widehat{\ell}_t - \alpha \mathbf{1}_k\|_{\nabla^2 \Psi_t^{-1}(w)}^2 \right] \\ & \leq \mathbb{E} \left[\sum_{i=1}^K \max_{w_i \in [w_{t,i}, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \widehat{\ell}_t + \alpha \mathbf{1}_k)_i]} \frac{\eta_{t,i}}{2} w_i^{3/2} (\widehat{\ell}_{t,i} - \alpha)^2 \right] \end{aligned}$$

From the definition of $\nabla \Psi^*(Y)_i$ (Equation 8) we know that $\nabla \Psi^*(Y)_i$ is increasing on $(-\infty, 0]$ and hence for $\alpha = 0$ we have $w_{t,i} \geq \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \widehat{\ell}_t)_i$. This implies the maximum of each of the terms is attained at $w_i = w_{t,i}$. Thus

$$\begin{aligned} & \mathbb{E} \left[\sum_{i=1}^K \max_{w_i \in [w_{t,i}, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \widehat{\ell}_t + \alpha \mathbf{1}_k)_i]} \frac{\eta_{t,i}}{2} w_i^{3/2} (\widehat{\ell}_{t,i})^2 \right] \\ & = \mathbb{E} \left[\sum_{i=1}^K \frac{\eta_{t,i}}{2} w_{t,i}^{3/2} \chi_{(i_t=i)} \frac{\ell_{t,i}^2}{w_{t,i}^2} \right] = \mathbb{E} \left[\sum_{i=1}^K \frac{\eta_{t,i}}{2} w_{t,i}^{3/2} \frac{\ell_{t,i}^2}{w_{t,i}} \right] \leq \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]}. \end{aligned}$$

When $\alpha = \chi_{(i_t=j)} \ell_{t,j}$ we consider several cases. First if $i_t \neq j$ the same bound as above holds. Next if $i_t = j$ for all $i \neq j$ we have $\widehat{\ell}_{t,i} - \alpha = -\alpha = -\ell_{t,j} \geq -1$ and for $\nabla \Phi_t^*(\nabla \Phi_t(w_t) - \widehat{\ell}_t + \ell_{t,j}) = \nabla \Phi_t^*(\nabla \Phi_t(w_t) + \ell_{t,j}) \leq 2^{2/3} w_{t,i}$ by Lemma D.3. This implies that in this case the maximum in the terms is bounded by $2w_{t,i}^{3/2} \ell_{t,j}^2$. Finally if $i_t = j$ for the j -th term we again use the fact that $w_{t,j} \geq \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \widehat{\ell}_t + \ell_{t,j})_j$ since $-\widehat{\ell}_{t,j} + \ell_{t,j} \leq 0$.

Algorithm 7 PLAY-ROUND

Input: Sampling distribution w_t
Output: Loss vector $\widehat{\ell}_t$

- 1: Sample algorithm i_t according to $\bar{w}_t = (1 - \frac{1}{Tk})w_t + \frac{1}{Tk}\text{Unif}(\Delta^{K-1})$.
 - 2: Algorithm i_t plays action a_{i_t, j_t} . Observe loss $\ell_t(a_{i_t, j_t})$ and construct unbiased estimator $\widehat{\ell}_t = \frac{\ell_t(a_{i_t, j_t})}{\bar{w}_{t, i_t}}e_{i_t}$ of ℓ_t .
 - 3: Give feedback to i -th algorithm as $\widehat{\ell}_t(a_{i, j_t})$, where a_{i, j_t} was action provided by $\mathcal{A}_i$
-

Combining all of the above we have

$$\begin{aligned}
 & \mathbb{E} \left[\max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \widehat{\ell}_t + \alpha \mathbf{1}_K)]} \|\widehat{\ell}_t - \chi_{(i_t=j)} \ell_{t,j} \mathbf{1}\|_{\nabla^2 \Psi_t^{-1}(w)}^2 \right] \leq \sum_{i \neq j} \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} \\
 & + \mathbb{E} \left[\chi_{(i_t=j)} \left(\frac{\eta_{t,j}}{2} \left(\frac{\ell_{t,j}}{w_{t,j}} - \ell_{t,j} \right)^2 w_{t,j}^{3/2} + \sum_{i \neq j} \ell_{t,j}^2 \frac{\eta_{t,i}}{2} w_{t,i}^{3/2} \right) \right] \\
 & = \sum_{i \neq j} \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} + \mathbb{E} \left[\frac{\eta_{t,j}}{2} (\ell_{t,j}(1 - w_{t,j}))^2 w_{t,j}^{1/2} + \sum_{i \neq j} \ell_{t,j}^2 \frac{\eta_{t,i}}{2} w_{t,i}^{3/2} w_{t,j} \right] \\
 & \leq \sum_{i \neq j} \frac{\eta_{t,i} + \eta_{t,j}}{2} \left(\sqrt{\mathbb{E}[w_{t,i}]} + \mathbb{E}[w_{t,i}] \right).
 \end{aligned}$$

□

Now the stability term is bounded by Lemma D.4. Next we proceed to bound the penalty term in a slightly different way. Direct computation yields

$$\begin{aligned}
 D_{\Phi_t}(-\widehat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\widehat{L}_t, \nabla \Phi_t^*(u)) &= -\Phi_t(-\widehat{L}_t) + \Phi_t(-\widehat{L}_{t-1}) - \langle -\widehat{L}_{t-1} + \widehat{L}_t, u \rangle \\
 &\quad + \Phi_t(\nabla \Phi_t^*(u)) - \Phi_t(\nabla \Phi_t^*(u)) \\
 &= -\Phi_t(-\widehat{L}_t) + \Phi_t(-\widehat{L}_{t-1}) - \langle \widehat{\ell}_t, u \rangle.
 \end{aligned} \tag{11}$$

Using the next lemma and telescoping will result in a bound for the sum of the penalty terms

Lemma D.5. *Let $u = e_{i^*}$ be the optimal algorithm. For any w_{t+1} such that $w_{t+1} = \nabla \Phi_{t+1}(-\widehat{L}_t)$ and $\eta_{t+1} \leq \eta_t$ it holds that*

$$D_{\Phi_{t+1}}(-\widehat{L}_t, \nabla \Phi_{t+1}^*(u)) - D_{\Phi_t}(-\widehat{L}_t, \nabla \Phi_t^*(u)) \leq 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t+1,i}} - \frac{1}{\eta_{t,i}} \right) \left(\sqrt{w_{t+1,i}} - \frac{1}{2} w_{t+1,i} \right).$$

Proof.

$$\begin{aligned}
 D_{\Phi_{t+1}}(-\widehat{L}_t, \nabla \Phi_{t+1}^*(u)) - D_{\Phi_t}(-\widehat{L}_t, \nabla \Phi_t^*(u)) &= \Phi_{t+1}(-\widehat{L}_t) - \Phi_t(-\widehat{L}_t) + \Phi_{t+1}^*(u) - \Phi_t^*(u) \\
 &\quad - \langle u, \widehat{L}_t - \widehat{L}_t \rangle \\
 &= \Phi_{t+1}(-\widehat{L}_t) - \Phi_t(-\widehat{L}_t) + \Psi_{t+1}(u) - \Psi_t(u) \\
 &= \Phi_{t+1}(-\widehat{L}_t) - \Phi_t(-\widehat{L}_t) - 2 \left(\frac{1}{\eta_{t+1,i^*}} - \frac{1}{\eta_{t,i^*}} \right) \\
 &= \langle w_{t+1}, -\widehat{L}_t \rangle - \Psi_{t+1}(w_{t+1}) - \Phi_t(-\widehat{L}_t) \\
 &\quad - 2 \left(\frac{1}{\eta_{t+1,i^*}} - \frac{1}{\eta_{t,i^*}} \right) \\
 &\leq \langle w_{t+1}, -\widehat{L}_t \rangle - \Psi_{t+1}(w_{t+1}) - \langle w_{t+1}, -\widehat{L}_t \rangle + \Psi_t(w_{t+1}) \\
 &\quad - 2 \left(\frac{1}{\eta_{t+1,i^*}} - \frac{1}{\eta_{t,i^*}} \right) \\
 &\leq 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t+1,i}} - \frac{1}{\eta_{t,i}} \right) \left(\sqrt{w_{t+1,i}} - \frac{1}{2} w_{t+1,i} \right).
 \end{aligned}$$

The first equality holds by Fenchel duality and the definition of Bregman divergence. The second equality holds by the fact that on the simplex $\Phi_t^*(\cdot) = \Psi_t(\cdot)$. The third equality holds because $\Psi_t(u) = -4(\sqrt{1} - \frac{1}{2})$. The fourth equality holds because w_{t+1} is the maximizer of $\langle -\widehat{L}_t, w \rangle + \Psi_{t+1}(w)$ and this is exactly how $\Phi_{t+1}(-\widehat{L}_t)$ is defined. The first inequality holds because

$$\begin{aligned}
 -\Phi_t(-\widehat{L}_t) &= \max_{w \in \Delta^{K-1}} \langle -\widehat{L}_t, w \rangle + \Psi_t(w) \\
 &\leq \langle -\widehat{L}_t, w_{t+1} \rangle + \Psi_t(w_{t+1}).
 \end{aligned}$$

The final inequality holds because $\Psi_t(w_{t+1}) - \Psi_{t+1}(w_{t+1}) = 4 \sum_i (1/\eta_{t+1,i} - 1/\eta_{t,i})(\sqrt{w_{t+1,i}} - w_{t+1,i}/2)$ and the fact that $\sqrt{w_{t+1,i^*}} - \frac{1}{2} w_{t+1,i^*} \leq \frac{1}{2}$. $\square$

Next we focus on the OMD update. By the 3-point rule for Bregman divergence we can write

$$\begin{aligned}
 \langle \widehat{\ell}_t, w_t - u \rangle &= \langle \nabla \Psi_t(w_t) - \nabla \Psi_t(\tilde{w}_{t+1}), w_t - u \rangle = D_{\Psi_t}(u, w_t) - D_{\Psi_t}(u, \tilde{w}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}) \\
 &\leq D_{\Psi_t}(u, w_t) - D_{\Psi_t}(u, \widehat{w}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}), \\
 \langle \widehat{\ell}_{t+1}, \widehat{w}_{t+1} - u \rangle &\leq D_{\Psi_{t+1}}(u, \widehat{w}_{t+1}) - D_{\Psi_{t+1}}(u, \widehat{w}_{t+2}) + D_{\Psi_{t+1}}(\widehat{w}_{t+1}, \tilde{w}_{t+2}),
 \end{aligned}$$

where the first inequality follows from the fact that $D_{\Psi_t}(u, \tilde{w}_{t+1}) \leq D_{\Psi_t}(u, \widehat{w}_{t+1})$ as $\widehat{w}_{t+1}$ is the projection of $\tilde{w}_{t+1}$ with respect to the Bregman divergence onto Δ^{K-1} .

We now explain how to control each of the terms. First we begin by matching $D_{\Psi_{t+1}}(u, \hat{w}_{t+1})$ with $-D_{\Psi_t}(u, \hat{w}_{t+1})$.

$$\begin{aligned}
 D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1}) &= \Psi_{t+1}(u) - \Psi_t(u) + \Psi_t(\hat{w}_{t+1}) - \Psi_{t+1}(\hat{w}_{t+1}) \\
 &\quad + \langle \nabla \Psi_t(\hat{w}_{t+1}), u - \hat{w}_{t+1} \rangle - \langle \nabla \Psi_{t+1}(\hat{w}_{t+1}), u - \hat{w}_{t+1} \rangle \\
 &= -2 \left(\frac{1}{\eta_{t+1,i^*}} - \frac{1}{\eta_{t,i^*}} \right) \\
 &\quad - 4 \sum_i \left(\sqrt{\hat{w}_{t+1,i}} - \frac{1}{2} \hat{w}_{t+1,i} \right) \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t+1,i}} \right) \\
 &\quad - 2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 1 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \\
 &\quad + 2 \sum_i \hat{w}_{t+1,i} \left(\frac{1}{\sqrt{\hat{w}_{t+1,i}}} - 1 \right) \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t+1,i}} \right), \\
 &= 2 \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \\
 &\quad - 2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 1 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \\
 &\quad - 2 \sum_i \sqrt{\hat{w}_{t+1,i}} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t+1,i}} \right),
 \end{aligned}$$

where we have set $u = e_{i^*}$. Since the step size schedule is non-decreasing during OMD updates, we have that the above is bounded by

$$\begin{aligned}
 D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1}) &\leq 2 \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) - 2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 1 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \\
 &\leq -2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 2 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right).
 \end{aligned} \tag{12}$$

Next we explain how to control the terms $D_{\Psi_t}(w_t, \tilde{w}_{t+1})$ and $D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2})$. These can be thought of as the stability terms in the FTRL update.

Lemma D.6. *For iterates generated by the OMD step in Equation 9 and any j it holds that*

$$\begin{aligned}
 \mathbb{E}[D_{\Psi_t}(w_t, \tilde{w}_{t+1})] &\leq \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]}, \\
 \mathbb{E}[D_{\Psi_t}(w_t, \tilde{w}_{t+1})] &\leq \sum_{i \neq j} \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} + \frac{\eta_{t,i} + \eta_{t,j}}{2} \mathbb{E}[w_{t,i}], \\
 \mathbb{E}[D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2})] &\leq \sum_{i=1}^K \frac{\eta_{t+1,i}}{2} \sqrt{\mathbb{E}[\hat{w}_{t+1,i}]}, \\
 \mathbb{E}[D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2})] &\leq \sum_{i \neq j} \frac{\eta_{t+1,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} + \frac{\eta_{t+1,i} + \eta_{t+1,j}}{2} \mathbb{E}[w_{t,i}],
 \end{aligned}$$

where $\tilde{w}_{t+1}$ is any iterate such that $\hat{w}_{t+1} = \operatorname{argmin}_{w \in \Delta^{K-1}} D_{\Psi_t}(w, \tilde{w}_{t+1})$.

Proof. We show the first two inequalities. The second couple of inequalities follow similarly. First we notice that we have

$$\hat{w}_{t+1} = \operatorname{argmin}_{w \in \Delta^{K-1}} \langle w, \hat{\ell}_t \rangle + D_{\Psi_t}(w, w_{t+1}) = \operatorname{argmin}_{w \in \Delta^{K-1}} \langle w, \hat{\ell}_t - \alpha \mathbf{1}_k \rangle + D_{\Psi_t}(w, w_{t+1}),$$

for any α . This implies that $\hat{w}_{t+1} = \operatorname{argmin}_{w \in \Delta^{K-1}} D_{\Psi_t}(w, \tilde{w}_{t+1})$ for $\tilde{w}_{t+1} = \operatorname{argmin}_{w \in \mathbb{R}^K} \langle w, \hat{\ell}_t - \alpha \mathbf{1}_k \rangle + D_{\Psi_t}(w, w_{t+1})$. We can now write

$$\begin{aligned} D_{\Psi_t^*}(\nabla \Psi_t(\tilde{w}_{t+1}), \nabla \Psi_t(w_t)) &= D_{\Psi_t^*}(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k, \nabla \Psi_t(w_t)) \\ &\leq \max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_k)]} \|\hat{\ell}_t - \alpha \mathbf{1}_k\|_{\nabla^2 \Psi_t^{-1}(w)}^2. \end{aligned}$$

The proof is finished by Lemma D.4. $\square$

Finally we explain how to control $D_{\Psi_t}(u, w_t)$ and $D_{\Psi_{t+1}}(u, \hat{w}_{t+2})$. First by Lemma D.1 it holds that

$$D_{\Psi_t}(u, w_t) = D_{\Phi_t}(-L_{t-1}, \nabla \Phi_t^*(u)).$$

This term can now be combined with the term $-D_{\Phi_{t-1}}(-L_{t-1}, \nabla \Phi_{t-1}^*(u))$ coming from the prior FTRL update and both terms can be controlled through Lemma D.5. To control $-D_{\Psi_{t+1}}(u, \hat{w}_{t+2})$ we show that $-D_{\Psi_{t+1}}(u, \hat{w}_{t+2}) = -D_{\Phi_{t+1}}(-\hat{L}_{t+1}, \nabla \Phi_{t+1}^*(u))$. This is done by showing that if $\hat{w}_{t+1}$ and $\hat{w}_{t+2}$ are defined as in Equation 9 we can equivalently write $\hat{w}_{t+2}$ as an FTRL step coming from a slightly different loss.

Lemma D.7. *Let $\hat{w}_{t+2}$ be defined as in Equation 9. Let ν_{t+1} be the constant such that $\nabla \Phi_{t+1}(-\hat{L}_t) = \nabla \Psi_{t+1}^*(-\hat{L}_t + \nu_t \mathbf{1}_k)$. Let $\hat{L}_{t+1} = (\mathbf{1}_k - e) \odot (\hat{L}_t - (\nu_{t-1} + \nu_t) \mathbf{1}_k) + \frac{1}{\beta} e \odot ((\hat{L}_t - (\nu_{t-1} + \nu_t) \mathbf{1}_k)) + \hat{\ell}_{t+1}$ and $\eta_{t+2} = \eta_{t+1}$. Then $(\hat{L}_{t+1})_i \geq 0$ for all $i \in [K]$ and $\hat{w}_{t+2} = w_{t+2} = \nabla \Phi_{t+2}(-\hat{L}_{t+1})$.*

Proof. By the definition of the update we have

$$\begin{aligned} \hat{w}_{t+1} &= \nabla \Phi_t(\nabla \Psi_t(w_t) - \hat{\ell}_t) = \nabla \Phi_t(-\hat{L}_t + \nu_{t-1} \mathbf{1}_k) \\ &= \nabla \Psi_t^*(-\hat{L}_t + (\nu_{t-1} + \nu_t) \mathbf{1}_k), \\ \hat{w}_{t+2} &= \nabla \Phi_{t+1}(\nabla \Psi_{t+1}(\hat{w}_{t+1}) - \hat{\ell}_{t+1}), \end{aligned}$$

where in the first equality we have used the fact that $\nabla \Psi_t(w_t) = -L_{t-1} + \nu_{t-1} \mathbf{1}_k$. For any i such that the OMD update increased the step size, i.e. $\eta_{t+1,i} = \beta \eta_{t,i}$ it holds from the definition of $\nabla \Psi_{t+1}(\cdot)$ that $\nabla \Psi_{t+1}(w)_i = \frac{1}{\beta} \nabla \Psi_t(w)_i$. Since $\nabla \Psi_t^*$ inverts $\nabla \Psi_t$ coordinate wise, we can write

$$\nabla \Psi_{t+1}(\hat{w}_{t+1})_i = \frac{1}{\beta} \nabla \Psi_t(\hat{w}_{t+1})_i = \frac{1}{\beta} (-\hat{L}_t + (\nu_{t-1} + \nu_t) \mathbf{1}_k)_i.$$

If we let e be the the sum of all e_i 's such that $\eta_{t+1,i} = \beta \eta_{t,i}$ we can write

$$\hat{w}_{t+2} = \nabla \Phi_{t+1} \left((\mathbf{1}_k - e) \odot (-\hat{L}_t + (\nu_{t-1} + \nu_t) \mathbf{1}_k) + \frac{1}{\beta} e \odot ((-\hat{L}_t + (\nu_{t-1} + \nu_t) \mathbf{1}_k)) - \hat{\ell}_{t+1} \right).$$

The fact that $\hat{L}_{t+1,i} \geq 0$ for any i follows since any coordinate $\nabla \Psi_t(\hat{w}_{t+1})_i \leq 0$ which implies that any coordinate of $(-\hat{L}_t + (\nu_{t-1} + \nu_t) \mathbf{1}_k)_i \leq 0$. $\square$

We can finally couple $-D_{\Phi_{t+1}}(-\hat{L}_{t+1}, \nabla \Phi_{t+1}^*(u))$ with the term from the next FTRL step which is $D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u))$ and use Lemma D.5 to bound the sum of this two terms. Putting everything together we arrive at the following regret guarantee.

Theorem D.8. *The regret bound for Algorithm 2 for any step size schedule which is non-increasing on the FTRL steps and any T_0 satisfies*

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{3}{2} \eta_{t,i} \sqrt{w_{t,i}} + \frac{\eta_{t,i} + \eta_{t,i^*}}{2} w_{t,i} \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] \\ &\quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\ &\quad + \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right]. \end{aligned}$$

Proof. Let $\mathcal{T}_{FTRL}$ be the set of all rounds in which the FTRL step is taken except for all rounds immediately before the OMD step and immediately after the OMD step. Let $\mathcal{T}_{OMD}$ be the set of all round immediately before the OMD step. The regret is bounded as follows:

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] &= \sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t - u \rangle] + \sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t - u \rangle] \\ &= \sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t - u \rangle] + \sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t \rangle + \Phi_t(-\hat{L}_t) - \Phi_t(-\hat{L}_{t-1}) \\ &\quad + D_{\Phi_t}(-\hat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u))]. \end{aligned}$$

For any T_0 , by the stability bound in Lemma D.4 we have

$$\begin{aligned} \sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t \rangle + \Phi_t(-\hat{L}_t) - \Phi_t(-\hat{L}_{t-1})] &\leq \sum_{t \in \mathcal{T}_{FTRL} \cap \{[T_0]\}} \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]} \\ &\quad + \sum_{t \in \mathcal{T}_{FTRL} \setminus \{[T_0]\}} \sum_{i \neq i^*} \mathbb{E} \left[\frac{\eta_{t,i}}{2} (\sqrt{w_{t,i}} + w_{t,i}) \right]. \end{aligned}$$

Next we consider the penalty term

$$\begin{aligned} &\sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} [D_{\Phi_t}(-\hat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u))] = \mathbb{E} [D_{\Phi_1}(0, \nabla \Phi_1^*(u))] \\ &+ \sum_{t+1 \in \mathcal{T}_{FTRL}} \mathbb{E} [D_{\Phi_{t+1}}(-\hat{L}_t, \nabla \Phi_{t+1}^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u))] - \mathbb{E} \left[\sum_{t \in \mathcal{T}_{OMD}} D_{\Phi_{t-1}}(-\hat{L}_{t-1}, \nabla \Phi_{t-1}^*(u)) \right] \\ &+ \mathbb{E} \left[\sum_{t \in \mathcal{T}_{OMD}} D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) \right] - \mathbb{E} [D_{\Phi_T}(-\hat{L}_T, \nabla \Phi_T^*(u))]. \end{aligned}$$

We are now going to complete the penalty term by considering the extra terms which do not bring negative regret from $\sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t - u \rangle]$.

$$\begin{aligned} &\sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t - u \rangle] \leq \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_t}(u, w_t) - D_{\Psi_t}(u, \hat{w}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1})] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_{t+1}}(u, \hat{w}_{t+2}) + D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2})] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [\langle \hat{\ell}_{t+2}, w_{t+2} \rangle + \Phi_{t+2}(-\hat{L}_{t+2}) - \Phi_{t+2}(-\hat{L}_{t+1})] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) - D_{\Phi_{t+2}}(-\hat{L}_{t+2}, \nabla \Phi_{t+2}^*(u))] \\ &= \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [\langle \hat{\ell}_{t+2}, w_{t+2} \rangle + \Phi_{t+2}(-\hat{L}_{t+2}) - \Phi_{t+2}(-\hat{L}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}) + D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2})] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1})] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_t}(-\hat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_{t+2}}(-\hat{L}_{t+2}, \nabla \Phi_{t+2}^*(u))] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) - D_{\Psi_{t+1}}(u, \hat{w}_{t+2})], \end{aligned}$$

where in the first inequality we have used the 3-point rule for Bregman divergence and the definition of the set $\mathcal{T}_{FTRL}$. For any T_0 the term

$$\sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [\langle \hat{\ell}_{t+2}, w_{t+2} \rangle + \Phi_{t+2}(-\hat{L}_{t+2}) - \Phi_{t+2}(-\hat{L}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}) + D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2})]$$

is bounded by Lemma D.4 and Lemma D.6 as follows

$$\begin{aligned} & \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[\langle \hat{\ell}_{t+2}, w_{t+2} \rangle + \Phi_{t+2}(-\hat{L}_{t+2}) - \Phi_{t+2}(-\hat{L}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}) + D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2}) \right] \\ & \leq \sum_{t \in \mathcal{T}_{OMD} \setminus \{[T_0]\}} \sum_{i \neq i^*} \mathbb{E} \left[\eta_{t+2,i} \sqrt{w_{t+2,i}} + \frac{\eta_{t+2,i} + \eta_{t+2,i^*}}{2} w_{t+2,i} \right] + \sum_{t \in \mathcal{T}_{OMD} \cap \{[T_0]\}} \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t+2,i}]}, \end{aligned}$$

where we have used the $i \neq i^*$ bound from the above lemmas for all terms past T_0 and the bound which includes all $i \in [K]$ for the first T_0 terms. The term $\sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1})]$ is bounded from Equation 12 as follows

$$\sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1})] \leq \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 2 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right].$$

By Lemma D.7 and Lemma D.1

$$\begin{aligned} & \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) - D_{\Psi_{t+1}}(u, \hat{w}_{t+2}) \right] \\ & = \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) - D_{\Phi_{t+1}}(-\hat{L}_{t+1}, \nabla \Phi_{t+1}^*(u)) \right]. \end{aligned}$$

Combining all of the above we have

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] & \leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{3}{2} \eta_{t,i} \sqrt{w_{t,i}} + \frac{\eta_{t,i} + \eta_{t,i^*}}{2} w_{t,i} \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] \\ & \quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 2 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\ & \quad + \sum_{t \in [T] \setminus \mathcal{T}_{OMD}} \mathbb{E} \left[D_{\Phi_{t+1}}(-\hat{L}_t, \nabla \Phi_{t+1}^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u)) \right] \\ & \quad + \mathbb{E}[D_{\Phi_1}(0, \nabla \Phi_1^*(u))] - \mathbb{E}[D_{\Phi_T}(-\hat{L}_T, \nabla \Phi_T^*(u))]. \end{aligned} \tag{13}$$

Using Lemma D.5 we have that

$$\begin{aligned} & \sum_{t \in [T] \setminus \mathcal{T}_{OMD}} \mathbb{E} \left[D_{\Phi_{t+1}}(-\hat{L}_t, \nabla \Phi_{t+1}^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u)) \right] \leq \\ & \sum_{t \in [T] \setminus \mathcal{T}_{OMD}} \mathbb{E} \left[4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t+1,i}} - \frac{1}{\eta_{t,i}} \right) \left(\sqrt{w_{t+1,i}} - \frac{1}{2} w_{t+1,i} \right) \right]. \end{aligned}$$

By definition of w_1 we have $D_{\Phi_1}(0, \nabla \Phi_1^*(u)) = \Psi_1(u) - \Psi_1(w_1)$. Plugging back into Equation 13 we have

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] & \leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{3}{2} \eta_{t,i} \sqrt{w_{t,i}} + \frac{\eta_{t,i} + \eta_{t,i^*}}{2} w_{t,i} \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] \\ & \quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\ & \quad + \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right]. \end{aligned}$$

□

The algorithm begins by running each algorithm for $\log(T) + 1$ rounds. We set the probability thresholds so that $\rho_1 = 36$, $\rho_j = 2\rho_{j-1}$ and $\frac{1}{\rho_n} \geq \frac{1}{KT}$, because we mix each w_t with the uniform distribution weighted by $1/KT$. This implies $n \leq \log_2(T)$. The algorithm now proceeds in epochs. The sizes of the epochs are as follows. The first epoch was of size $K \log(T) + K$, each epoch after doubles the size of the preceding one so that the number of epochs is bounded by $\log(T)$. In the beginning of each epoch, except for the first epoch we check if $w_{t,i} < \frac{1}{\rho_1}$. If it is we increase the step size $\eta_{t+1,i} = \beta\eta_{t,i}$ and run the OMD step. Let the τ -th epoch have size s_τ . Let $\frac{1}{\rho_\tau}$ be the largest threshold which was not exceeded during epoch τ . We require that each of the algorithms have the following expected regret bound under the unbiased rescaling of the losses $\bar{R}_i(t)$: $\mathbb{E}[\bar{R}_i(\sum_{\tau=1}^S s_\tau)] \leq \sum_{\tau=1}^S \mathbb{E}[\sqrt{\rho_\tau} R_i(s_\tau)]$. This can be ensured by restarting the algorithms in the beginning of the epochs if at the beginning of epoch τ it happens that $w_{t,i} > \frac{1}{\rho_{\tau-1}}$. Let ℓ_t be the loss over all possible actions. Let i_t be the algorithm selected by the corraling algorithm at time t . Let a^* be the best overall action.

Lemma D.9. *Let $\bar{R}_{i^*}(\cdot)$ be a function upper bounding the expected regret of $\mathcal{A}_{i^*}$, $\mathbb{E}[R_{i^*}(\cdot)]$. For any η such that $\eta_{1,i} \leq \min_{t \in [T]} \frac{(1 - \exp(-\frac{1}{\log(T)^2}))\sqrt{t}}{50\bar{R}_{i^*}(t)}$, $\forall i \in [K]$ it holds that*

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \ell_t(a_{i_t, j_t}) - \ell_t(a^*) \right] &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{3}{2} (\eta_{t,i} + \eta_{t,i^*}) (\sqrt{w_{t,i}} + w_{t,i}) \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] \\ &+ \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right] + 1 + 36\mathbb{E}[R_{i^*}(T)]. \end{aligned}$$

Proof. First we note that $\mathbb{E}[\hat{\ell}_t(i^*)] = \mathbb{E} \left[w_{i^*,t} \frac{\ell_t(a_{i^*, j_t})}{w_{i^*,t}} \right] = \mathbb{E}[\ell_t(a_{i^*, j_t})]$. Using Theorem D.8 we have

$$\begin{aligned} \sum_{t=1}^T \mathbb{E} [\ell_t(a_{i_t, j_t}) - \ell_t(a^*)] &= \sum_{t=1}^T \mathbb{E} [\ell_t(a_{i^*, j_t}) - \ell_t(a^*)] + \sum_{t=1}^T \mathbb{E} [\hat{\ell}_t, \bar{w}_t - u] \leq \sum_{t=1}^T \mathbb{E} [\hat{\ell}_t(i^*) - \ell_t(a^*)] \\ &+ \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{3}{2} (\eta_{t,i} + \eta_{t,i^*}) (\sqrt{w_{t,i}} + w_{t,i}) \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] \\ &+ \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\ &+ \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right] + 1. \end{aligned}$$

Let us focus on $\sum_{t=1}^T \mathbb{E} [\hat{\ell}_t(i^*) - \ell_t(a^*)] - 2 \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[\left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right]$. By our assumption on $\mathcal{A}_{i^*}$ it holds that

$$\sum_{t=1}^T \mathbb{E} [\hat{\ell}_t(i^*) - \ell_t(a^*)] \leq \mathbb{E} \left[R_{i^*} \left(\sum_{\tau=1}^{\log(T)} s_\tau \right) \right] \leq \sum_{\tau=1}^{\log(T)} \mathbb{E} [\sqrt{\rho_\tau} R_{i^*}(s_\tau)].$$

We now claim that during epoch τ there is a t in that epoch such that also $t \in \mathcal{T}_{OMD}$ and for which $w_{t,i^*} \leq \frac{1}{\rho_{\tau-1}}$. We consider two cases, first if OMD was invoked because at least one of the probability thresholds ρ_s was passed by a w_{t_s, i^*} , we must have $\rho_s \leq \rho_\tau$. Also by definition of ρ_τ as the largest threshold not passed by any w_{t, i^*} there exists at least one $t' \geq t_s$ for which $\frac{1}{\rho_{\tau-1}} \geq w_{t', i^*} > \frac{1}{\rho_\tau}$. This implies that we have subtracted at least $2\mathbb{E} \left[\left(\frac{1}{\sqrt{\hat{w}_{t'+1, i^*}}} - 3 \right) \left(\frac{1}{\eta_{t', i^*}} - \frac{1}{\eta_{t'+1, i^*}} \right) \right] \geq 2\mathbb{E} \left[(\sqrt{\rho_{\tau-1}} - 3) \left(\frac{1}{\eta_{t', i^*}} - \frac{1}{\eta_{t'+1, i^*}} \right) \right]$. In the second case we have that for all t in epoch τ it holds that $\frac{1}{\rho_{\tau-1}} \geq w_{t, i^*} > \frac{1}{\rho_\tau}$ or $w_{t, i^*} > \frac{1}{\rho_1}$. In the second case we only incur regret $\mathbb{E}[R_1(t)]$ scaled by 36 and in the first case the OMD played in the beginning of the epoch has resulted in at least $-2\mathbb{E} \left[(\sqrt{\rho_{\tau-1}} - 3) \left(\frac{1}{\eta_{t, i^*}} - \frac{1}{\eta_{t+1, i^*}} \right) \right]$ negative contribution, where t indexes the beginning of the epoch. We set

$\beta = e^{1/\log(T)^2}$ and now evaluate the difference $\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \geq \left(1 - \frac{1}{\beta}\right) \frac{\sqrt{t}}{25\eta_{1,i^*}}$. Where we have used the fact that $\eta_{t,i^*} \leq \frac{\eta_{1,i^*} \beta^{\log_2(T)^2}}{\sqrt{t}} \leq \frac{25\eta_{1,i^*}}{\sqrt{t}}$. This follows by noting that there are $\log_2(T)$ epochs and during each epoch one can call the OMD step only $\log_2(T)$ times. Let $\beta' = \left(1 - \frac{1}{\beta}\right)$. Thus if t_τ is the beginning of epoch τ we subtract at least $\frac{\beta' \sqrt{t_\tau \rho_{\tau-1}}}{25\eta_{1,i^*}}$. Notice that the length of each epoch s_τ does not exceed $2t_\tau$, thus we have

$$\sum_{\tau=1}^{\log(T)} \mathbb{E}[\sqrt{\rho_\tau} R_{i^*}(s_\tau)] \leq \sum_{\tau=1}^{\log(T)} \mathbb{E}[\sqrt{\rho_\tau} R_{i^*}(2t_\tau)],$$

and so as long as we set $\eta_{1,i^*} \leq \frac{\beta' \sqrt{2t_\tau}}{50\bar{R}_{i^*}(2t_\tau)}$, where $\mathbb{E}[R_{i^*}(2t_\tau)] \leq \bar{R}_{i^*}(2t_\tau)$ we have

$$\begin{aligned} \sum_{\tau=1}^{\log(T)} \mathbb{E}[\sqrt{\rho_\tau} R_{i^*}(2t_\tau)] - 2 \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[\left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\ \leq \sum_{\tau=1}^{\log(T)} \mathbb{E}[\sqrt{\rho_\tau} R_{i^*}(2t_\tau) - \sqrt{\rho_\tau} R_{i^*}(2t_\tau)] \leq 0. \end{aligned}$$

□

We can now use the self-bounding trick of the regret as in [Zimmert and Seldin \(2018\)](#) to finish the proof. Let μ^* denote the reward of the best arm. First note that we can write

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \ell_t(a_{i_t, j_t}) - \ell_t(a^*) \right] &= \mathbb{E} \left[\sum_{t=1}^T \chi_{i_t \neq i^*} (\ell_t(a_{i_t, j_t}) - \mu^*) \right] + \mathbb{E}[R_{i^*}(T_{i^*}(T))] \\ &\geq \mathbb{E} \left[\sum_{t=1}^T \sum_{i=1}^K w_{t,i} \chi_{i_t \neq i^*} (\ell_t(a_{i_t, j_t}) - \mu^*) \right] \\ &\geq \mathbb{E} \left[\sum_{t=1}^T \sum_{i \neq i^*} w_{t,i} \Delta_i \right]. \end{aligned}$$

Theorem D.10. *Let $\bar{R}_{i^*}(\cdot)$ be a function upper bounding the expected regret of $\mathcal{A}_{i^*}$, $\mathbb{E}[R_{i^*}(\cdot)]$. For any η such that $\eta_{1,i} \leq \min_{t \in [T]} \frac{(1 - \exp(-\frac{1}{\log(T)^2})) \sqrt{t}}{50\bar{R}_{i^*}(t)}$, $\forall i \in [K]$ and $\beta = e^{1/\log(T)^2}$ it holds that the expected regret of Algorithm 2 is bounded as*

$$\begin{aligned} \mathbb{E}[R(T)] &\leq \sum_{i \neq i^*} \frac{1500(1/\eta_{1,i} + \eta_{1,i})^2}{\Delta_i} \left(\log \left(\frac{T\Delta_i - 15\eta_{1,i}}{T_0\Delta_i - 15\eta_{1,i}} \right) + \log(225\eta_{1,i}^2 \Delta_i / \Delta_1) \right) \\ &\quad + \sum_{i \in [K]} \frac{8}{\eta_{1,i} \sqrt{K}} + 2 + 72R_{i^*}(T), \end{aligned}$$

where $T_0 = \max_{i \neq i^*} \frac{225\eta_{1,i}^2}{\Delta_i}$.

Proof of Theorem 5.2. By Lemma D.9 we have that the overall regret is bounded by

$$\begin{aligned}
 \mathbb{E}[R(T)] &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{3}{2} (\eta_{t,i} + \eta_{t,i^*}) (\sqrt{w_{t,i}} + w_{t,i}) \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] \\
 &\quad + \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right] \\
 &\quad + 1 + 36\mathbb{E}[R_{i^*}(T)] \\
 &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{75\eta_{1,i}}{2\sqrt{t}} (\sqrt{w_{t,i}} + w_{t,i}) \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{25\eta_{1,i}}{2\sqrt{t}} \sqrt{w_{t,i}} \right] \\
 &\quad + \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t=1}^T \sum_{i \neq i^*} \frac{10}{\eta_{1,i}\sqrt{t}} (\sqrt{w_{t,i}}) \right] \\
 &\quad + 1 + 36\mathbb{E}[R_{i^*}(T)] \\
 &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{75\eta_{1,i}}{2\sqrt{t}} (\sqrt{w_{t,i}} + w_{t,i}) \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{25\eta_{1,i}}{2\sqrt{t}} \sqrt{w_{t,i}} \right] \\
 &\quad + \mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t=1}^T \sum_{i \neq i^*} \frac{10}{\eta_{1,i}\sqrt{t}} (\sqrt{w_{t,i}}) \right] \\
 &\quad + 1 + 36\mathbb{E}[R_{i^*}(T)] + \mathbb{E}[R(T)] - \mathbb{E} \left[\sum_{t=1}^T \sum_{i \neq i^*} w_{t,i} \Delta_i \right] \\
 &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \mathbb{E} \left[\frac{75\eta_{1,i}}{\sqrt{t}} (\sqrt{w_{t,i}} + w_{t,i}) \right] + \sum_{t=1}^{T_0} \sum_{i=1}^K \mathbb{E} \left[\frac{25\eta_{1,i}}{\sqrt{t}} \sqrt{w_{t,i}} \right] \\
 &\quad + 2\mathbb{E} [\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t=1}^T \sum_{i \neq i^*} \frac{20}{\eta_{1,i}\sqrt{t}} (\sqrt{w_{t,i}}) \right] \\
 &\quad + 2 + 3672\mathbb{E}[R_{i^*}(T)] - \mathbb{E} \left[\sum_{t=1}^T \sum_{i \neq i^*} w_{t,i} \Delta_i \right].
 \end{aligned}$$

In the first inequality we used the fact that for any i we have $\eta_{t,i} \leq 25\eta_{1,i}/\sqrt{t}$, in the second inequality we have used the self bounding property derived before the statement of the theorem and in the third inequality we again used the bound on the expected regret $\mathbb{E}[R(T)]$ from the first inequality. We are now going to use the fact that for any $w > 0$ it holds that $2\alpha\sqrt{w} - \beta w \leq \frac{\alpha^2}{\beta}$. For $t \leq T_0$ we have

$$\sum_{t=1}^{T_0} \sum_{i \neq i^*} \left(20 \frac{\sqrt{w_{t,i}}}{\sqrt{t}} \left(\frac{1}{\eta_{1,i}} + \eta_{1,i} \right) - \Delta_i w_{t,i} \right) \leq \sum_{t=1}^{T_0} \sum_{i \neq i^*} \frac{1500(1/\eta_{1,i} + \eta_{1,i})^2}{t\Delta_i}.$$

For $t > T_0$ we have

$$\begin{aligned}
 \sum_{t=T_0+1}^T \sum_{i \neq i^*} \left(\frac{\sqrt{w_{t,i}}}{\sqrt{t}} \left(\frac{20}{\eta_{1,i}} + 75\eta_{1,i} \right) - \left(\Delta_i - \frac{15\eta_{1,i}}{\sqrt{t}} \right) w_{t,i} \right) &\leq \sum_{t=T_0+1}^T \sum_{i \neq i^*} \frac{1500(1/\eta_{1,i} + \eta_{1,i})^2}{t\Delta_i - 15\eta_{1,i}\sqrt{t}} \\
 &\leq \sum_{i \neq i^*} \int_{T_0}^T \frac{1500(1/\eta_{1,i} + \eta_{1,i})^2}{t\Delta_i - 15\eta_{1,i}\sqrt{t}} dt \\
 &= \frac{1500(1/\eta_{1,i} + \eta_{1,i})^2}{\Delta_i} \log \left(\frac{15\eta_{1,i} - T\Delta_i}{15\eta_{1,i} - T_0\Delta_i} \right).
 \end{aligned}$$

We now choose $T_0 = \max_{i \neq i^*} \frac{225\eta_{1,i}^2}{\Delta_i}$. To bound $\mathbb{E}[\Psi_1(u) - \Psi_1(w_1)]$ we have set w_1 to be the uniform distribution over the K algorithms and recall that $\Psi_1(w) = -4 \sum_i \frac{\sqrt{w_i} - \frac{1}{2}w_i}{\eta_{1,i}}$. This implies $\Psi_1(u) - \Psi_1(w_1) \leq \sum_{i \in [K]} \frac{4}{\eta_{1,i}\sqrt{K}}$. Putting everything together we have

$$\begin{aligned} \mathbb{E}[R(T)] &\leq \sum_{i \neq i^*} \frac{1500(1/\eta_{1,i} + \eta_{1,i})^2}{\Delta_i} \left(\log \left(\frac{T\Delta_i - 15\eta_{1,i}}{T_0\Delta_i - 15\eta_{1,i}} \right) + \log(225\eta_{1,i}^2/\Delta_i) \right) \\ &\quad + \sum_{i \in [K]} \frac{8}{\eta_{1,i}\sqrt{K}} + 2 + 72R_{i^*}(T) \end{aligned}$$

□

To parse the above regret bound in the stochastic setting we note that the min-max regret bound for the k -armed problem is $\Theta(\sqrt{kT})$. Most popular algorithms like UCB, Thompson sampling and mirror descent have a regret bound which is (up to poly-logarithmic factors) $O(\sqrt{kT})$. If we were to corral only such algorithms, the condition of the theorem implies that $\frac{1}{\eta_{1,i}} \in \tilde{O}(\sqrt{k})$ as $\frac{\sqrt{i}}{R_i(t)} \leq O(1), \forall t \in [T]$. What happens, however, if algorithm $\mathcal{A}_i$ has a worst case regret bound of the order $\omega(\sqrt{T})$? For the next part of the discussion we only focus on time horizon dependence. As a simple example suppose that $\mathcal{A}_i$ has worst case regret of $T^{2/3}$ and that $\mathcal{A}_{i^*}$ has a worst case regret of $\sqrt{T}$. In this case Theorem D.8 tells us that we should set $\eta_{1,i} = \tilde{O}(1/T^{1/6})$ and hence the regret bound scales at least as $\Omega(T^{1/3}/\Delta_i + \mathbb{E}[R_{i^*}(T)])$. In general if the worst case regret bound of $\mathcal{A}_i$ is in the order of T^α we have a regret bound scaling at least as $T^{2\alpha-1}/\Delta_i$.

D.3 Stability of UCB and UCB-like algorithms under a change of environment

In this section we discuss how the regret bounds for UCB and similar algorithms change whenever the variance of the stochastic losses is rescaled by Algorithm 2. Assume that the UCB algorithm plays against stochastic rewards bounded in $[0, 1]$. We begin by noting that after every call to OMD-STEP (Algorithm 6) the UCB algorithm should be restarted with a change in the environment which reflects that the variance of the losses has now been rescaled. Let the UCB algorithm of interest be $\mathcal{A}_i$. If the OMD step occurred at time t' and it was the case that $\frac{1}{\rho_{s-1}} \geq w_{t',i} > \frac{1}{\rho_s}$, then we know that the rescaled rewards will be in $[0, \rho_s]$ until the next time the UCB algorithm is restarted. This suggests that the confidence bound for arm j at time t should become $\sqrt{\frac{\rho_s^2 \log(t)}{T_{i,j}(t)}}$.

However, we note that the second moment of the rescaled rewards is only $\frac{\ell_t(a_{i,j_t})^2}{w_{t,i}}$. A slightly more careful analysis using Bernstein's inequality for martingales (e.g. Lemma 10 Bartlett et al. (2008)) allows us to show the following.

Theorem D.11 (Theorem 5.4 formal). *Suppose that during epoch τ of size $\mathcal{T}$ UCB-I is restarted and its environment was changed by ρ_s so that the upper confidence bound is changed to $\sqrt{\frac{4\rho_s \log(t)}{T_{i,j}(t)}} + \frac{4\rho_s \log(t)}{3T_{i,j}(t)}$ for arm j at time t . Then the expected regret of the algorithm is bounded by*

$$\mathbb{E}[R_i(\mathcal{T})] \leq \sqrt{8\rho_s k_i \mathcal{T} \log(\mathcal{T})}$$

Proof of Theorem 5.4. Let the reward of arm j at time t be $r_{t,j}$ and the rescaled reward be $\hat{r}_{t,j}$. Without loss of generality assume that the arm with highest reward is $j = 1$. Denote the mean of arm j as μ_j and denote the mean of the best arm as μ^* . During this run of UCB we know that each $|\hat{r}_{t,j}| \leq \rho_s$. Further if we denote the probability with which the algorithm is sampled at time t as $w_{t,i}$ we have $\mathbb{E}[\hat{r}_{t,j} - \mu_j | w_{1:t-1,i}] = 0$ and hence $r_{t,j} - \mu_j$ is a martingale difference. Further notice that the conditional second moment of $r_{t,j}$ is $\mathbb{E}[\hat{r}_{t,j}^2 | w_{1:t-1,i}] = \mathbb{E}[w_{t,i} \frac{r_{t,j}^2}{w_{t,i}^2} + 0 | w_{1:t-1,i}] \leq \rho$. Let $Y_t = (\hat{r}_{\tau,j} - \mu_j)$. Bernstein's inequality for martingales (Bartlett et al. (2008)[Lemma 10]) now implies that $\mathbb{P}\left[\sum_{t=1}^{\mathcal{T}} Y_t > \sqrt{2\mathcal{T}\rho \log(1/\delta)} + \frac{2}{3}\rho \log(1/\delta)\right] \leq \delta$. This implies that the confidence bound should be changed to

$$\sqrt{\frac{4\rho_s \log(t)}{T_{i,j}(t)}} + \frac{4\rho_s \log(t)}{3T_{i,j}(t)}.$$

Following the standard proof of UCB we can now conclude that a suboptimal arm can be pulled at most $T_{i,j}(t)$ times up to time t where

$$2\Delta_j \geq \sqrt{\frac{4\rho_s \log(t)}{T_{i,j}(t)}} + \frac{4\rho_s \log(t)}{3T_{i,j}(t)}.$$

This implies that

$$\mathbb{E}[T_{i,j}(t)] \leq \frac{8\rho_s \log(t)}{\Delta_j^2}.$$

Next we bound the regret of the algorithm up to time t as follows:

$$\begin{aligned} \mathbb{E}[R_i(t)] &\leq \sum_{j \neq 1} \Delta_j \mathbb{E}[T_{i,j}(t)] = \sum_{j \neq j^*} \sqrt{\mathbb{E}[T_{i,j}(t)]} \sqrt{\Delta_j^2 \mathbb{E}[T_{i,j}(t)]} \\ &\leq \sum_{j \neq j^*} \sqrt{\mathbb{E}[T_{i,j}(t)]} \sqrt{8\rho_s \log(t)} \leq k_i \sqrt{\frac{1}{k_i} \sum_j \mathbb{E}[T_{i,j}(t)]} = \sqrt{8\rho_s k_i t \log(t)}. \end{aligned}$$

□

In general the argument can be repeated for other UCB-type algorithms (e.g. Successive Elimination) and hinges on the fact that the rescaled rewards $\hat{r}_{t,j}$ have second moment bounded by ρ since with probability $w_{t,i}$ we have $\hat{r}_{t,j}^2 = \frac{r_{t,j}^2}{w_{t,i}^2}$ and with probability $1 - w_{t,i}$ it equals $\hat{r}_{t,j}^2 = 0$. We are not sure if similar arguments can be carried out for more delicate versions of UCB, like KL-UCB and leave it as future work to check.

E Regret bound in the adversarial setting

We now consider the setting in which the best overall arm does not maintain a gap at every round. Following the proof of Theorem D.8 we are able to show the following.

Theorem E.1. *The regret bound for Algorithm 2 for any step size schedule which is non-increasing on the FTRL steps satisfies*

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] &\leq 4 \max_{w \in \Delta^{K-1}} \sqrt{T} \sum_{i=1}^K \left(\eta_{1,i} + \frac{1}{\eta_{1,i}} \right) \sqrt{w_i} \\ &\quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right]. \end{aligned}$$

Proof. From the proof of Theorem D.8 we have

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] &= \sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t - u \rangle] + \sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t \rangle + \Phi_t(-\hat{L}_t) - \Phi_t(-\hat{L}_{t-1}) \\ &\quad + D_{\Phi_t}(-\hat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u))]. \end{aligned}$$

Lemma D.4 implies

$$\sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} [\langle \hat{\ell}_t, w_t \rangle + \Phi_t(-\hat{L}_t) - \Phi_t(-\hat{L}_{t-1})] \leq \sum_{t \in \mathcal{T}_{FTRL}} \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t,i}]}.$$

As before the penalty term is decomposed as follows

$$\begin{aligned}
 & \sum_{t \in \mathcal{T}_{FTRL}} \mathbb{E} \left[D_{\Phi_t}(-\hat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u)) \right] = \mathbb{E} [D_{\Phi_1}(0, \nabla \Phi_1^*(u))] \\
 & + \sum_{t+1 \in \mathcal{T}_{FTRL}} \mathbb{E} \left[D_{\Phi_{t+1}}(-\hat{L}_t, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u)) \right] - \mathbb{E} \left[\sum_{t \in \mathcal{T}_{OMD}} D_{\Phi_{t-1}}(-\hat{L}_{t-1}, \nabla \Phi_{t-1}^*(u)) \right] \\
 & + \mathbb{E} \left[\sum_{t \in \mathcal{T}_{OMD}} D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) \right] - \mathbb{E} [D_{\Phi_T}(-\hat{L}_T, \nabla \Phi_T^*(u))].
 \end{aligned}$$

Next the term $\sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E}[\langle \hat{\ell}_t, w_t - u \rangle]$ is again decomposed as in the proof of Theorem D.8

$$\begin{aligned}
 & \sum_{t \in [T] \setminus \mathcal{T}_{FTRL}} \mathbb{E}[\langle \hat{\ell}_t, w_t - u \rangle] \\
 & \leq \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[\langle \hat{\ell}_{t+2}, w_{t+2} \rangle + \Phi_{t+2}(-\hat{L}_{t+2}) - \Phi_{t+2}(-\hat{L}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}) + D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2}) \right] \\
 & + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1})] \\
 & + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_t}(-\hat{L}_{t-1}, \nabla \Phi_t^*(u)) - D_{\Phi_{t+2}}(-\hat{L}_{t+2}, \nabla \Phi_{t+2}^*(u))] \\
 & + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*(u)) - D_{\Psi_{t+1}}(u, \hat{w}_{t+2})].
 \end{aligned}$$

Using Lemma D.4 and Lemma D.6 we bound the first term of the above inequality as

$$\begin{aligned}
 & \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[\langle \hat{\ell}_{t+2}, w_{t+2} \rangle + \Phi_{t+2}(-\hat{L}_{t+2}) - \Phi_{t+2}(-\hat{L}_{t+1}) + D_{\Psi_t}(w_t, \tilde{w}_{t+1}) + D_{\Psi_{t+1}}(\hat{w}_{t+1}, \tilde{w}_{t+2}) \right] \\
 & \leq \sum_{t \in \mathcal{T}_{OMD}} \sum_{i=1}^K \frac{\eta_{t,i}}{2} \sqrt{\mathbb{E}[w_{t+2,i}]}
 \end{aligned}$$

The term $\sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1})]$ is bounded from Equation 12 as follows

$$\sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Psi_{t+1}}(u, \hat{w}_{t+1}) - D_{\Psi_t}(u, \hat{w}_{t+1})] \leq \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 2 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right].$$

By Lemma D.7 and Lemma D.1

$$\begin{aligned}
 & \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*) - D_{\Psi_{t+1}}(u, \hat{w}_{t+2})] \\
 & = \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_{t+2}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*) - D_{\Phi_{t+1}}(-\hat{L}_{t+1}, \nabla \Phi_{t+2}^*)].
 \end{aligned}$$

Combining all of the above we have

$$\begin{aligned}
 \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] & \leq \sum_{t=1}^T \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 2 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\
 & + \sum_{t \in [T] \setminus \mathcal{T}_{OMD}} \mathbb{E} [D_{\Phi_{t+1}}(-\hat{L}_t, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u))] \\
 & + \mathbb{E}[D_{\Phi_1}(0, \nabla \Phi_1^*(u))] - \mathbb{E}[D_{\Phi_T}(-\hat{L}_T, \nabla \Phi_T^*(u))].
 \end{aligned} \tag{14}$$

The last two terms are bounded in the same way as in the proof of Theorem D.8

$$\begin{aligned}
 & \sum_{t \in [T] \setminus \mathcal{T}_{OMD}} \mathbb{E} \left[D_{\Phi_{t+1}}(-\hat{L}_t, \nabla \Phi_t^*(u)) - D_{\Phi_t}(-\hat{L}_t, \nabla \Phi_t^*(u)) \right] \\
 & \quad + \mathbb{E}[D_{\Phi_1}(0, \nabla \Phi_1^*(u))] - \mathbb{E}[D_{\Phi_T}(-\hat{L}_T, \nabla \Phi_T^*(u))] \\
 & \leq \mathbb{E}[\Psi_1(u) - \Psi_1(w_1)] + \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} 4 \sum_{i \neq i^*} \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right]
 \end{aligned}$$

Plugging back into Equation 14 we have

$$\begin{aligned}
 \mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right] & \leq \sum_{t=1}^T \sum_{i=1}^K \mathbb{E} \left[\frac{\eta_{t,i}}{2} \sqrt{w_{t,i}} \right] + \mathbb{E}[\Psi_1(u) - \Psi_1(w_1)] \\
 & \quad + 4 \mathbb{E} \left[\sum_{t \in [T] \setminus \mathcal{T}_{OMD}} \sum_{i=1}^K \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) (\sqrt{w_{t,i}}) \right] \\
 & \quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\
 & \leq \sum_{t=1}^T 4 \sum_{i=1}^K \left(\eta_{1,i} + \frac{1}{\eta_{1,i}} \right) \sqrt{\frac{w_{t,i}}{t}} \\
 & \quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] \\
 & \leq 4 \max_{w \in \Delta^{K-1}} \sqrt{T} \sum_{i=1}^K \left(\eta_{1,i} + \frac{1}{\eta_{1,i}} \right) \sqrt{w_i} \\
 & \quad + \sum_{t \in \mathcal{T}_{OMD}} \mathbb{E} \left[-2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right],
 \end{aligned}$$

where the last inequality follows from the fact that the maximizer of the function $\sum_{i=1}^K \sqrt{\frac{w_i}{t}} \alpha_i$ over the simplex, for $\alpha_i \geq 0$ is the same for all $t \in [T]$. $\square$

Following the proof of Lemma D.9 and replacing the bound on $\mathbb{E} \left[\sum_{t=1}^T \langle \hat{\ell}_t, w_t - u \rangle \right]$ from Theorem D.8 with the one from Theorem E.1 yields the next result.

Theorem E.2 (Theorem 5.3). *Let $\bar{R}_{i^*}(\cdot)$ be a function upper bounding the expected regret of $\mathcal{A}_{i^*}$, $\mathbb{E}[R_{i^*}(\cdot)]$. For any $\eta_{1,i^*} \leq \min_{t \in [T]} \frac{(1 - \exp(-\frac{1}{\log(T)^2})) \sqrt{t}}{50 R_{i^*}(t)}$ and $\beta = e^{1/\log(T)^2}$ it holds that the expected regret of Algorithm 2 is bounded as*

$$\mathbb{E} \left[\sum_{t=1}^T \ell_t(a_t) - \ell_t(a^*) \right] \leq 4 \max_{w \in \Delta^{K-1}} \sqrt{T} \sum_{i=1}^K \left(\eta_{1,i} + \frac{1}{\eta_{1,i}} \right) \sqrt{w_i} + 36 R_{i^*}(T).$$

A few remarks are in order. First, when the rewards obey the stochastically constrained adversarial setting i.e., there exists a gap Δ_i at every round between the best action and every other action during all rounds $t \in [T]$, then the regret for corraling bandit algorithms with worst case regret bounds of the order $\tilde{O}(\sqrt{T})$ in time horizon is at most $\tilde{O}(\sum_{i \neq i^*} \frac{\log(T)^5}{\Delta_i} + R_{i^*}(T))$. On the other hand, if there is no gap in the rewards then a worst case regret bound is still $\tilde{O}(\max\{\sqrt{KT}, \max_i \bar{R}_i(T)\} + R_{i^*}(T))$. This implies that Algorithm 2 can be used as a model selection tool when we are not sure what environment we are playing against. For example, if we are not sure if we should use a contextual bandit algorithm, a linear bandit algorithm or a stochastic multi-armed bandit algorithm, we can corral all of them and Algorithm 2 will perform almost as well as the algorithm for the best

environment. Further, if we are in a distributed setting where we have access to multiple algorithms of the same type but not the arms they are playing, we can do almost as well as an algorithm which plays on all the arms simultaneously. We believe that our algorithm will have numerous other applications outside of the scope of the above examples.

F Proof of Theorem 7.1

Recall the gap assumption made in Theorem 7.1:

Assumption F.1. For any $i < i^*$ it holds that for all $(x, a) \in \mathcal{X} \times \mathcal{A}$

$$\mathbb{E}[\langle \beta_i, \phi_i(x, a) \rangle] - \min_{a \in \mathcal{A}} \mathbb{E}[\langle \beta^*, \phi_{i^*}(x, a) \rangle] \geq 2 \frac{d_{i^*}^{2\alpha} - d_i^{2\alpha}}{\sqrt{T}}.$$

Since the losses might not be bounded in $[0, 1]$ as $d_K = \Theta(T)$ we need to slightly modify the bound for the Stability term in Lemma D.4 and the term $D_{\Psi_t}(w_t, \tilde{w}_{t+1})$ in Lemma D.6. Recall that we need to bound the term $\mathbb{E} \left[\max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_K)]} \|\hat{\ell}_t\|_{\nabla^2 \Psi_t^{-1}(w)}^2 \right]$. The argument is the same as in D.4 up to

$$\mathbb{E} \left[\max_{w \in [w_t, \nabla \Psi_t^*(\nabla \Psi_t(w_t) - \hat{\ell}_t + \alpha \mathbf{1}_K)]} \|\hat{\ell}_t\|_{\nabla^2 \Psi_t^{-1}(w)}^2 \right] \leq \mathbb{E} \left[\sum_{i=1}^K \frac{\eta_{t,i}}{2} w_{t,i}^{3/2} (\hat{\ell}_{t,i})^2 \right].$$

Let $\ell_{t,i} = \langle \beta_{i_t}, \phi_{i_t}(x_t, a_{i_t,j_t}) \rangle + \xi_t$, then we have

$$\mathbb{E} \left[\frac{\eta_{t,i}}{2} w_{t,i}^{3/2} (\hat{\ell}_{t,i})^2 \right] \leq \mathbb{E} \left[\eta_{t,i} \chi_{(i_t=i)} w_{t,i}^{3/2} \frac{\ell_{t,i}^2}{w_{t,i}^2} + \eta_{t,i} w_{t,i}^{3/2} \frac{d_i^{4\alpha}}{T} \right] \leq \mathbb{E} \left[\eta_{t,i} w_{t,i} \frac{d_i^{4\alpha}}{T} \right] + 2\mathbb{E} [\eta_{t,i} \sqrt{w_{t,i}}],$$

where in the last inequality we have used the fact that $w_{t,i} \geq w_{t,i}^{3/2}$ together with the our assumption that ξ_t is zero-mean with variance proxy 1. Following the proof of Lemma D.9 with the bound on the stability term we can bound

$$\begin{aligned} \mathbb{E} \left[\sum_{t=1}^T \ell_t(a_{i_t,j_t}) - \ell_t(a^*) \right] &= \sum_{t=1}^T \mathbb{E} [\ell_t(a_{i^*,j_t}) - \ell_t(a^*)] + \sum_{t=1}^T \mathbb{E} [\langle \hat{\ell}_t + \mathbf{d}, w_t - u \rangle] - \sum_{t=1}^T \mathbb{E} [\langle \mathbf{d}, w_t - u \rangle] + 1 \\ &\leq \sum_{t=1}^T \mathbb{E} [\hat{\ell}_t(i^*) - \ell_t(a^*)] + 2 \sum_{t=1}^T \sum_{i=1}^K \mathbb{E} \left[\eta_{t,i} \sqrt{w_{t,i}} + \eta_{t,i} w_{t,i} \frac{d_i^{4\alpha}}{T} \right] + 4 \sum_{t=1}^T \sum_{i=1}^K \left(\frac{1}{\eta_{t,i}} - \frac{1}{\eta_{t-1,i}} \right) \sqrt{w_{t,i}} \\ &\quad - \sum_{t=1}^T \mathbb{E} [\langle \mathbf{d}, w_t - u \rangle] - \sum_{t \in T_{OMD}} \mathbb{E} \left[2 \left(\frac{1}{\sqrt{\hat{w}_{t+1,i^*}}} - 3 \right) \left(\frac{1}{\eta_{t,i^*}} - \frac{1}{\eta_{t+1,i^*}} \right) \right] + \sqrt{K} + 1 \\ &\leq 4 \sum_{t=1}^T \sum_{i=1}^K \sqrt{\frac{w_{t,i}}{t}} \left(\eta_{1,i} + \frac{1}{\eta_{t,i}} \right) + 2 \sum_{t=1}^T \sum_{i=1}^K \mathbb{E} \left[\eta_{t,i} w_{t,i} \frac{d_i^{4\alpha}}{T} \right] - \sum_{t=1}^T \mathbb{E} [\langle \mathbf{d}, w_t - u \rangle] + 36\mathbb{E}[R_{i^*}(T)]. \end{aligned}$$

For a fixed t we have

$$-\langle \mathbf{d}, w_t - u \rangle = \frac{d_{i^*}^{2\alpha}}{\sqrt{T}} (1 - w_{t,i^*}) - \sum_{i \neq i^*} w_{t,i} \frac{d_i^{2\alpha}}{\sqrt{T}} = \sum_{i < i^*} w_{t,i} \frac{d_{i^*}^{2\alpha} - d_i^{2\alpha}}{\sqrt{T}} - \sum_{i > i^*} w_{t,i} \frac{d_i^{2\alpha} - d_{i^*}^{2\alpha}}{\sqrt{T}}.$$

First we consider the terms $i > i^*$. Assume WLOG that $d_K^{2\alpha} \leq T/4$, as otherwise the learning guarantees are trivial. For these terms we have

$$\sqrt{\frac{w_{t,i}}{t}} \frac{1}{\eta_{1,i}} + w_{t,i} \left(\frac{\eta_{1,i}}{\sqrt{t}} \frac{d_i^{4\alpha}}{T} - \frac{d_i^{2\alpha}}{\sqrt{T}} \right) \leq \sqrt{\frac{w_{t,i}}{t}} \frac{1}{\eta_{1,i}} - w_{t,i} \frac{d_i^{2\alpha}}{2\sqrt{T}} \leq \frac{\sqrt{T}}{td_i^{2\alpha} \eta_{1,i}^2}.$$

Since $\eta_{1,i} = \tilde{\Theta}(1/d_i^\alpha)$ we have that the above is further bounded by $\tilde{O}(\sqrt{T}/t)$.

Next we consider the terms for $i < i^*$ given by $w_{t,i} \frac{d_{i^*}^{2\alpha} - d_i^{2\alpha}}{\sqrt{T}}$. Here we use our assumption that the regret $\mathbb{E} \left[\sum_{t=1}^T \ell_t(a_{i_t, j_t}) - \ell_t(a^*) \right] \geq \mathbb{E} [w_{t,i} \Delta_i]$, where $\Delta_i = \mathbb{E}[\langle \beta_i, \phi_i(x, a) \rangle] - \min_{a \in \mathcal{A}} \mathbb{E}[\langle \beta^*, \phi_{i^*}(x, a) \rangle]$. Using the self-bounding trick we can cancel out the terms $w_{t,i} \frac{d_{i^*}^{2\alpha} - d_i^{2\alpha}}{\sqrt{T}}$ as soon as $\Delta_i \geq 2w_{t,i} \frac{d_{i^*}^{2\alpha} - d_i^{2\alpha}}{\sqrt{T}}$, which holds by Assumption F.1. All other terms in the regret bound are bounded by $\tilde{O}(d_{i^*}^\alpha \sqrt{T})$. Thus we have shown that the regret of the corralling algorithm is bounded as

$$\mathbb{E} \left[\sum_{t=1}^T \ell_t(a_{i_t, j_t}) - \ell_t(a^*) \right] \leq \tilde{O} \left(\mathbb{E}[R_{i^*}(T)] + K\sqrt{T} \right).$$