

Online Minimax Multiobjective Optimization: Multicalibeating and Other Applications

Daniel Lee¹, Georgy Noarov¹, Mallesh Pai², Aaron Roth¹

¹ University of Pennsylvania, ² Rice University

daniellee@alumni.upenn.edu, gnoarov@seas.upenn.edu,

mallesh.pai@rice.edu, aaroth@cis.upenn.edu

October 14, 2022

Abstract

We introduce a simple but general online learning framework in which a learner plays against an adversary in a vector-valued game that changes every round. Even though the learner’s objective is not convex-concave (and so the minimax theorem does not apply), we give a simple algorithm that can compete with the setting in which the adversary must announce their action first, with optimally diminishing regret. We demonstrate the power of our framework by using it to (re)derive optimal bounds and efficient algorithms across a variety of domains, ranging from multicalibration to a large set of no regret algorithms, to a variant of Blackwell’s approachability theorem for polytopes with fast convergence rates. As a new application, we show how to “(multi)calibeat” an arbitrary collection of forecasters — achieving an exponentially improved dependence on the number of models we are competing against, compared to prior work.

1 Introduction

We introduce and study a simple but powerful framework for online adversarial multiobjective minimax optimization. At each round t , an adaptive adversary chooses an environment for the learner to play in, defined by a convex compact action set $\mathcal{X}^t$ for the learner, a convex compact action set $\mathcal{Y}^t$ for the adversary, and a d -dimensional continuous loss function $\ell^t : \mathcal{X}^t \times \mathcal{Y}^t \rightarrow [-1, 1]^d$ that, in each coordinate, is convex in the learner’s action and concave in the adversary’s action. The learner then chooses an action, or distribution over actions, x^t , and the adversary responds with an action y^t . This results in a loss vector $\ell^t(x^t, y^t)$, which accumulates over time. The learner’s goal is to minimize the maximum accumulated loss over each of the d dimensions: $\max_{j \in [d]} \left(\sum_{t=1}^T \ell_j^t(x^t, y^t) \right)$.

One may view the environment chosen at each round t as defining a zero-sum game in which the learner wishes to minimize the *maximum* coordinate of the resulting loss vector. The objective of the learner in the stage game in isolation can be written as:¹

$$w_L^t = \inf_{x^t \in \mathcal{X}^t} \max_{y^t \in \mathcal{Y}^t} \left(\max_{j \in [d]} \ell_j^t(x^t, y^t) \right).$$

Unfortunately, although ℓ_j^t is convex-concave in each coordinate, the maximum over coordinates does not preserve concavity for the adversary. Thus the minimax theorem does not hold, and the value of the game in which the learner moves first (defined above) is larger than the value of the game in which the adversary moves first—that is, $w_L^t > w_A^t$, where w_A^t is defined as:

¹ A brief aside about the “inf max max” structure of w_L^t : since each ℓ_j^t is continuous, so is $\max_j \ell_j^t$, and hence $\max_y (\max_j \ell_j^t)$ is attained on the compact set $\mathcal{Y}^t$. However, $\max_y (\max_j \ell_j^t)$ may not be a continuous function of x and therefore the infimum over $\mathcal{X}^t$ need not be attained.

$$w_A^t = \sup_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} \left(\max_{j \in [d]} \ell_j^t(x^t, y^t) \right).$$

Nevertheless, fixing a series of T environments chosen by the adversary, this defines in hindsight an aspirational quantity $W_A^T = \sum_{t=1}^T w_A^t$, summing the adversary-moves-first value of the constituent zero sum games. Despite the fact that these values are not individually obtainable in the stage games, we show that they are approachable on average over a sequence of rounds, i.e., there is an algorithm for the learner that guarantees that against any adversary,

$$\max_{j \in [d]} \left(\frac{1}{T} \sum_{t=1}^T \ell_j^t(x^t, y^t) \right) \leq \frac{1}{T} W_A^T + 4\sqrt{\frac{2 \ln d}{T}}.$$

Our derivation is elementary and based on a minimax argument, and is a development of a game-theoretic argument from the calibration literature due to Hart [2020] and Fudenberg and Levine [1999].² The generic algorithm plays actions at every round t according to a minimax equilibrium strategy in a surrogate game that is derived both from the environment chosen by the adversary at round t , as well as from the history of play so far on previous rounds $t' < t$. The loss in the surrogate game is convex-concave (and so we may apply minimax arguments), and can be used to upper bound the loss in the original games.

We then show that this simple framework can be instantiated to derive a wide array of optimal bounds, and that the corresponding algorithms can be derived in closed form by solving for the minimax equilibrium of the corresponding surrogate game. Despite its simplicity, our framework has a number of applications to online learning—we sketch these below.

“Multi-Calibeating”: Foster and Hart [2021] recently introduced the notion of “calibeating” an arbitrary online forecaster: making online calibrated predictions about an adversarially chosen sequence of inputs that are guaranteed to have lower squared error than an arbitrary predictor f , where the improvement in error approaches f ’s calibration error in hindsight. Foster and Hart give two methods for calibeating an arbitrary collection of predictors $\mathcal{F}$ simultaneously, but these methods have an exponential and polynomial dependence in their convergence bounds on $|\mathcal{F}|$, respectively.

Using our framework, we can derive optimal online bounds for online *multicalibration* [Hébert-Johnson et al., 2018, Gupta et al., 2022], and as an application, obtain bounds for calibeating arbitrary collection of models with only a *logarithmic* dependence on $|\mathcal{F}|$. Our algorithm naturally extends to the more general problem of online “multi-calibeating” — i.e. combining the goals of online multicalibration and calibeating. Namely, we give an algorithm for making real-valued predictions given contexts from some space Θ . The algorithm is parameterized by (i) a collection $\mathcal{G} \subseteq 2^\Theta$ of (arbitrary, potentially intersecting) subsets of Θ that we might envision to represent e.g. different demographic groups in a setting in which we are making predictions about people; and (ii) an arbitrary collection of predictors $\mathcal{F}$. We promise that our predictions are calibrated not just overall, but simultaneously within each group $g \in \mathcal{G}$ — and moreover, that we *calibeat* each predictor $f \in \mathcal{F}$ not just overall, but simultaneously within each group $g \in \mathcal{G}$. We do this by proving an online analogue of what Hébert-Johnson et al. [2018] call a “do no harm” property in the batch setting using a similar technique: multicalibrating with respect to the level sets of the predictors.

Fast Polytope Blackwell Approachability: We give a variant of Blackwell’s Approachability Theorem [Blackwell, 1956] for approaching a polytope. Standard methods approach a set in Euclidean distance, at a rate polynomial in the payoff dimension. In contrast, we give a dimension-independent approachability guarantee: we approximately satisfy all halfspace constraints defining the polytope, after *logarithmically* many rounds in the number of constraints, a significant improvement over a polynomial dimensional dependence in many settings. It is equivalent to the results of Perchet [2015], which show that the negative orthant $\mathbb{R}_{\leq 0}^d$ is approachable in the ℓ_∞ metric with a $\log(d)$ dependence in the convergence rate. This result follows immediately from a specialization of our framework that does not require changing the environment at each round, highlighting the connection between our framework and approachability. We remark that approachability has been extended in a number of ways in recent years [Mannor et al., 2014a,b, Perchet and Mannor,

²This argument was extended in Gupta et al. [2022] to obtain fast rates and explicit algorithms for multicalibration and multivalidity.

2013]. However most of our other applications take advantage of the flexibility of our framework to play a different game at each round (which can be defined by context) with potentially different action sets, and so do not directly follow from Blackwell approachability. Therefore, while many of our regret bounds could be derived from approachability to the negative orthant by enlarging the action space exponentially to simulate aspects of our framework, this approach would not easily lead to *efficient* algorithms.

Recovering Expert Learning Bounds: Algorithms and optimal bounds for various expert learning problems fall naturally out of our framework as corollaries. This includes external regret [Vovk, 1990, Littlestone and Warmuth, 1994], internal and swap regret [Foster and Vohra, 1998, Hart and Mas-Colell, 2000, Blum and Mansour, 2007], adaptive regret [Littlestone and Warmuth, 1994, Hazan and Seshadhri, 2009, Adamskiy et al., 2012], sleeping experts [Freund et al., 1997, Blum, 1997, Blum and Mansour, 2007, Kleinberg et al., 2010], and the recently introduced multi-group regret [Blum and Lykouris, 2020, Rothblum and Yona, 2021]. Multi-group regret refers to a contextual prediction problem in which the learner gets contexts from Θ before each round. It is parameterized by a collection of groups $\mathcal{G} \subseteq 2^\Theta$: e.g., if the predictions concern people, $\mathcal{G}$ may represent an arbitrary, intersecting set of demographic groups. Here the “experts” are different models that make predictions on each instance; the goal is to attain no-regret not just overall, but also on the subset of rounds corresponding to contexts from each $g \in \mathcal{G}$. Multi-group regret, like multicalibration, is one of the few solution concepts in the algorithmic fairness literature known not to involve tradeoffs with overall accuracy [Globus-Harris et al., 2022]. Blum and Lykouris [2020] derived their algorithm for online multigroup regret via a reduction to sleeping experts, and Gupta et al. [2022] derived their algorithm for online multicalibration via a direct argument. Here we derive online algorithms for both multicalibration and multigroup regret as corollaries of the same fundamental framework.

1.1 Additional Related Work

Papers by Azar et al. [2014] and Kesselheim and Singla [2020] study a related problem: an online setting with vector-valued losses, where the goal is to minimize the ℓ_∞ norm of the accumulated loss vector (they also consider other ℓ_p -norms). However, they study an incomparable benchmark that in our notation would be written as $\min_{x^* \in \mathcal{X}} \max_{j \in [d]} \frac{1}{T} \sum_{t=1}^T \ell_j(x^*, y^t)$ (which is well-defined in their setting, where loss functions $\ell^t = \ell$ and action sets $\mathcal{X}^t = \mathcal{X}, \mathcal{Y}^t = \mathcal{Y}$ are fixed throughout the interaction). On the one hand, this benchmark is stronger than ours in the sense that the maximum over coordinates is taken outside the sum over time, whereas our benchmark considers a “greedy” per-round maximum. On the other hand, in our setting the game can be different at every round, so our benchmark allows a comparison to a different action at each round rather than a single fixed action. In the setting of Kesselheim and Singla [2020], it is impossible to give any regret bound to their benchmark, so they derive an algorithm obtaining a $\log(d)$ *competitive ratio* to this benchmark. In contrast, our benchmark admits a regret bound. Hence, our results are quite different in kind despite the outward similarity of the settings: none of our applications follow from their theorems (since in all of our applications, we derive regret bounds).

A different line of work [Rakhlin et al., 2010, 2011] takes a very general minimax approach towards deriving bounds in online learning, including regret minimization, calibration, and approachability. Their approach is substantially more powerful than the framework we introduce here (e.g. it can be used to derive bounds for infinite dimensional problems, and characterizes online learnability in the sense that it can also be used to prove lower bounds). However, it is also correspondingly more complex, and requires analyzing the continuation value of a T round dynamic program. Such analyses are generally technically challenging; as an example, a recent line of work by Drenska and Kohn [2020] and Kobzar et al. [2020] considers a Rakhlin et al.-style minimax formulation of the standard experts problem, and shows how to find nonlinear PDE-based minimax solutions for the Learner and the Adversary that can be optimal not just asymptotically in the number of experts (dimensions) d , but also nonasymptotically for small d such as 2 or 3; their PDE approach is also conducive to bounding not just the maximum regret across dimensions, but also more general functions of the individual dimensions’ losses.

Overall, results derived from the Rakhlin et al. framework (with some notable exceptions, including

Rakhlin et al. [2012]) are generically nonconstructive, whereas our framework is simple and inherently constructive, in that the algorithm derives from repeatedly solving a one-round stage zero-sum game. Relative to this literature, we view our framework as a “user-friendly” power tool that can be used to derive a wide variety of algorithms and bounds without much additional work — at the cost of not being universally expressive.

2 General Framework

2.1 The Setting

A Learner (she) plays against an Adversary (he) over rounds $t \in [T] := \{1, \dots, T\}$. Over these rounds, she accumulates a d -dimensional loss vector ($d \geq 1$), where each round’s loss vector lies in $[-C, C]^d$ for some $C > 0$. At each round t , the Learner and the Adversary interact as follows:

1. Before round t , the Adversary selects and reveals to the Learner an *environment* comprising:
 - (a) The Learner’s and Adversary’s respective convex compact action sets $\mathcal{X}^t, \mathcal{Y}^t$ embedded into a finite-dimensional Euclidean space;
 - (b) A continuous vector loss function $\ell^t(\cdot, \cdot) : \mathcal{X}^t \times \mathcal{Y}^t \rightarrow [-C, C]^d$, with each $\ell_j^t(\cdot, \cdot) : \mathcal{X}^t \times \mathcal{Y}^t \rightarrow [-C, C]$ (for $j \in [d]$) convex in the 1st and concave in the 2nd argument.
2. The Learner selects some $x^t \in \mathcal{X}^t$.
3. The Adversary observes the Learner’s selection x^t , and responds with some $y^t \in \mathcal{Y}^t$.
4. The Learner suffers (and observes) the loss vector $\ell^t(x^t, y^t)$.

The Learner’s objective is to minimize the value of the maximum dimension of the accumulated loss vector after T rounds—in other words, to minimize: $\max_{j \in [d]} \sum_{t \in [T]} \ell_j^t(x^t, y^t)$.

To benchmark the Learner’s performance, we consider the following quantity at each round t :

Definition 2.1 (The Adversary-Moves-First (AMF) Value at Round t). *The Adversary-Moves-First value of the game defined by the environment $(\mathcal{X}^t, \mathcal{Y}^t, \ell^t)$ at round t is:*

$$w_A^t := \sup_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} \left(\max_{j \in [d]} \ell_j^t(x^t, y^t) \right).$$

If the Adversary had to reveal y^t first and the Learner could best respond, w_A^t would be the smallest value of the maximum coordinate of ℓ^t she could guarantee. However, the function $\max_{j \in [d]} \ell_j^t(x^t, y^t)$ is not convex-concave (as the max does not preserve concavity); hence the minimax theorem does not apply, making this value unobtainable for the Learner, who is in fact obligated to reveal x^t first. However, we can define regret to a benchmark given by the cumulative AMF values of the games:

Definition 2.2 (Adversary-Moves-First (AMF) Regret). *On transcript $\pi^t = \{(\mathcal{X}^s, \mathcal{Y}^s, \ell^s), x^s, y^s\}_{s=1}^t$, we define the Learner’s Adversary Moves First (AMF) Regret for the j^{th} dimension at time t to be:*

$$R_j^t(\pi^t) := \sum_{s=1}^t \ell_j^s(x^s, y^s) - \sum_{s=1}^t w_A^s.$$

*The overall AMF Regret is then defined as follows: $R^t(\pi^t) = \max_{j \in [d]} R_j^t$.*³

Again, the game played at each round is *not* convex-concave, so we cannot get $R^T \leq 0$. Instead, we will aim to obtain *sublinear* AMF regret, worst-case over adaptive adversaries: $R^T = o(T)$.

³We will generally elide the dependence on the transcript and simply write R_j^t and R^t .

2.2 General Algorithm

Our algorithmic framework will be based on a natural idea: instead of directly grappling with the maximum coordinate of the cumulative vector valued loss, we upper bound the AMF regret with a one-dimensional “soft-max” surrogate loss function, which the algorithm will then aim to minimize.

Definition 2.3 (Surrogate loss). *Fixing a parameter $\eta \in (0, 1)$, we define our surrogate loss function (that implicitly depends on the transcript π^t through the respective round t) as:*

$$L^t := \sum_{j \in [d]} \exp(\eta R_j^t) \quad \text{for } t \in [T], \quad \text{and } L^0 := d.$$

This surrogate loss tightly bounds the AMF regret $R^T = \max_{j \in [d]} R_j^T$:

Lemma 2.1. *The Learner’s AMF Regret is upper bounded using the surrogate loss as: $R^T \leq \frac{\ln L^T}{\eta}$.*

Next we observe a simple but important bound on the per-round increase in the surrogate loss.

Lemma 2.2. *For any t , any transcript through round t , and any $\eta \leq \frac{1}{2C}$, it holds that:*

$$L^t \leq (4\eta^2 C^2 + 1) L^{t-1} + \eta \sum_{j \in [d]} \exp(\eta R_j^{t-1}) \cdot (\ell_j^t(x^t, y^t) - w_A^t).$$

The proof is very simple (see Appendix A.1): we write out the quantity $L^t - L^{t-1}$, use the definition of AMF regret R^t , and then bound $L^t - L^{t-1}$ via the inequality $e^x \leq 1 + x + x^2$ for $|x| \leq 1$.

We now exploit Lemma 2.2 to bound the final surrogate loss L^T and obtain a game-theoretic algorithm for the Learner that attains this bound. While the above steps should remind the reader of a standard derivation of the celebrated Exponential Weights algorithm via bounding a log-sum-exp potential function, the next lemma is the novel ingredient that makes our framework significantly more general by relying on Sion’s powerful generalization of the Minimax Theorem to convex-concave games.

Lemma 2.3. *For any $\eta \leq \frac{1}{2C}$, the Learner can ensure that the final surrogate loss is bounded as:*

$$L^T \leq d (4\eta^2 C^2 + 1)^T.$$

Proof sketch; see Appendix A.1. Define, for $t \in [T]$, continuous convex-concave functions $u^t : \mathcal{X}^t \times \mathcal{Y}^t \rightarrow \mathbb{R}$ by: $u^t(x, y) := \sum_{j \in [d]} \exp(\eta R_j^{t-1}) (\ell_j^t(x, y) - w_A^t)$. If the Learner can ensure $u^t(x^t, y^t) \leq 0$ on all rounds $t \in [T]$ regardless of the Adversary’s play, then Lemma 2.2 implies $L^t \leq (4\eta^2 C^2 + 1) L^{t-1}$ for all $t \in [T]$, leading to the desired bound on L^T . Due to the continuous convex-concave nature of each u^t (inherited from the loss coordinates ℓ_j^t), we can apply Sion’s Minimax Theorem to conclude that: $\min_{x^t \in \mathcal{X}^t} \max_{y^t \in \mathcal{Y}^t} u^t(x^t, y^t) = \max_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} u^t(x^t, y^t)$.

In words, the Learner has a so-called *minimax-optimal* strategy x^t , that achieves (worst-case over all $y^t \in \mathcal{Y}^t$) value $u^t(x^t, y^t)$ as low as if the Adversary moved first and the Learner could best-respond. But in the latter counterfactual scenario, using the definitions of u^t and the Adversary-moves-first value w_A^t , we can easily see that by best-responding to the Adversary, the Learner would always guarantee herself value ≤ 0 : that is, $\max_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} u^t(x^t, y^t) \leq 0$. Thus, $\min_{x^t \in \mathcal{X}^t} \max_{y^t \in \mathcal{Y}^t} u^t(x^t, y^t) \leq 0$, and so by playing minimax-optimally at every round $t \in [T]$, the Learner will guarantee $u^t(x^t, y^t) \leq 0$ for all t , leading to the desired regret bound. $\square$

In fact, via a simple algebraic transformation (see Appendix A.1) taking advantage of the values w_A^t being independent of the actions x^t, y^t , we can explicitly express the Learner’s minimax optimal strategies at all rounds as: $\operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} u^t(x, y) = \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \frac{\exp(\eta \sum_{s=1}^{t-1} \ell_j^s(x^s, y^s))}{\sum_{i \in [d]} \exp(\eta \sum_{s=1}^{t-1} \ell_i^s(x^s, y^s))} \ell_j^t(x, y)$. Together with the

proof of Lemma 2.3, this immediately gives the following algorithm for the Learner that achieves the desired bound on L^T (and thus, as we will show, on the AMF regret R^T).

Algorithm 1: General Algorithm for the Learner that Achieves Sublinear AMF Regret

for rounds $t = 1, \dots, T$ **do**

Learn adversarially chosen $\mathcal{X}^t, \mathcal{Y}^t$, and loss function $\ell^t(\cdot, \cdot)$.

Let
$$\chi_j^t := \frac{\exp\left(\eta \sum_{s=1}^{t-1} \ell_j^s(x^s, y^s)\right)}{\sum_{i \in [d]} \exp\left(\eta \sum_{s=1}^{t-1} \ell_i^s(x^s, y^s)\right)} \text{ for } j \in [d].$$

Play
$$x^t \in \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \chi_j^t \cdot \ell_j^t(x, y).$$

Observe the Adversary's selection of $y^t \in \mathcal{Y}^t$.

Theorem 2.1 (AMF Regret guarantee of Algorithm 1). *For any $T \geq \ln d$, Algorithm 1 with learning rate $\eta = \sqrt{\frac{\ln d}{4TC^2}}$ obtains, against any Adversary, AMF regret bounded by: $R^T \leq 4C\sqrt{T \ln d}$.*

Indeed, using Lemma 2.1, then Lemma 2.3, then $1 + x \leq e^x$, and finally setting $\eta = \sqrt{\frac{\ln d}{4TC^2}}$, we get:

$$R^T \leq \frac{\ln L^T}{\eta} \leq \frac{\ln(d(4\eta^2 C^2 + 1)^T)}{\eta} \leq \frac{\ln(d \exp(4T\eta^2 C^2))}{\eta} = \frac{\ln d}{\eta} + 4TC^2\eta = 4C\sqrt{T \ln d}.$$

Remark 2.1. *Our framework is easy to adapt to the setting where the Learner randomizes, at each round, amongst a finite set of actions $\mathcal{A}^t$ (i.e. $\mathcal{X}^t = \Delta \mathcal{A}^t$), and wishes to obtain in-expectation and high-probability AMF regret bounds. This is useful in all our applications below. Additionally, our AMF regret bounds are robust to the Learner playing only an approximate (rather than exact) minimax strategy at each round: we use this to derive our simple multicalibration algorithm below. See Appendix A.2 for both these extensions.*

3 Deriving No-X-Regret Algorithms from Our Framework

The core of our framework — the Adversary-Moves-First regret — is strictly more general than a very large variety of known regret notions including: *external*, *internal*, *swap*, *adaptive*, *sleeping-experts*, *multigroup*, and *wide-range* (Φ) *regret*. Specifically, in Appendix E, we use our framework to derive simple $O(\sqrt{T})$ -regret algorithms for what we call *subsequence regret*, which encapsulates all these regret forms. In each of these cases, our generic algorithm is efficient, and often specializes (by computing a minimax equilibrium strategy in closed form) to simple combinatorial algorithms that had been derived from first principles in prior work. We note that in any problem that involves context or changing action spaces (as the sleeping experts problem does), we are taking advantage of the flexibility of our framework to present a different environment at every round, which distinguishes our framework from more standard Blackwell approachability arguments. In fact, as we will see in Section 5 below, our framework recovers fast Blackwell approachability as a special case.

For our general subsequence regret algorithms, please see Appendix E. Now, as a warm-up application of our framework, we directly instantiate it for the simplest case of obtaining $O(\sqrt{T})$ *external* regret.

Simple Learning From Expert Advice: External Regret In the classical experts learning setting Littlestone and Warmuth [1994], the Learner has a set of pure actions (“experts”) $\mathcal{A}$. At the outset of each round $t \in [T]$, the Learner chooses a distribution over experts $x^t \in \Delta \mathcal{A}$. The Adversary then comes up with a vector of losses $r^t = (r_a^t)_{a \in \mathcal{A}} \in [0, 1]^{\mathcal{A}}$ corresponding to each expert. Next, the Learner samples $a^t \sim x^t$, and experiences loss corresponding to the expert she chose: $r_{a^t}^t$. The Learner also gets to observe the entire vector of losses r^t for that round. The goal of the Learner is to achieve sublinear *external regret* — that is, to ensure that the difference between her cumulative loss and the loss of the best fixed expert in hindsight grows sublinearly with T : $R_{\text{ext}}^T(\pi^T) := \sum_{t \in [T]} r_{a^t}^t - \min_{j \in \mathcal{A}} \sum_{t \in [T]} r_j^t = o(T)$.

Theorem 3.1. Fix a finite pure action set $\mathcal{A}$ for the Learner and a time horizon $T \geq \ln |\mathcal{A}|$. Then, an instantiation of our framework's Algorithm 2 lets the Learner achieve the following regret bounds:

$$\mathbb{E}_{\pi^T} [R_{\text{ext}}^T(\pi^T)] \leq 4\sqrt{T \ln |\mathcal{A}|}, \quad \text{and } R_{\text{ext}}^T(\pi^T) \leq 8\sqrt{T \ln \frac{|\mathcal{A}|}{\delta}} \text{ with prob. } 1 - \delta.$$

Proof. We instantiate (the probabilistic version of) our framework (see Section A.2.1).

At all rounds, the Learner's pure action set is $\mathcal{A}$, and the Adversary's strategy space is the convex and compact set $[0, 1]^{|\mathcal{A}|}$, from which each round's collection $(r_a^t)_{a \in \mathcal{A}}$ of all actions' losses is selected. Next, we define a $|\mathcal{A}|$ -dimensional loss function $\ell^t = (\ell_j^t)_{j \in \mathcal{A}}$, where each coordinate loss ℓ_j^t expresses the regret of the Learner's chosen action a relative to action $j \in \mathcal{A}$:

$$\ell_j^t(a, r^t) = r_a^t - r_j^t, \quad \text{for } a \in \mathcal{A}, r^t \in [0, 1]^{|\mathcal{A}|}.$$

By Theorem A.1, $\mathbb{E} \left[\max_{j \in \mathcal{A}} \sum_{t \in [T]} \ell_j^t(a^t, r^t) - \sum_{t \in [T]} w_A^t \right] \leq 4\sqrt{T \ln |\mathcal{A}|}$, where w_A^t is the AMF value at round t . Using this AMF regret bound, we can bound the Learner's external regret as:

$$\mathbb{E} [R_{\text{ext}}^T] = \mathbb{E} \left[\max_{j \in \mathcal{A}} \sum_{t \in [T]} r_{a^t}^t - r_j^t \right] = \mathbb{E} \left[\max_{j \in \mathcal{A}} \sum_{t \in [T]} \ell_j^t(a^t, r^t) \right] \leq 4\sqrt{T \ln |\mathcal{A}|} + \sum_{t \in [T]} w_A^t.$$

It thus remains to show that the AMF value $w_A^t \leq 0$ for all t . This holds, since if the Learner knew the Adversary's choice of losses $(r_a^t)_{a \in \mathcal{A}}$ before round t , then picking the action $a \in \mathcal{A}$ with the smallest loss r_a^t would get her 0 regret in that round.⁴ This gives the in-expectation regret bound; the high-probability bound follows in the same way from Theorem A.2. $\square$

A bound of $\sqrt{T \ln |\mathcal{A}|}$ is optimal for external regret in the experts learning setting, and so serves to witness the optimality of our framework's general AMF regret bound in Theorem 2.1.

In fact, the above instantiation of Algorithm 2 yields the classical Exponential Weights algorithm Littlestone and Warmuth [1994]: at each round t , the action a^t is sampled with $\Pr[a^t = j] \sim \exp \left(-\eta \sum_{s=1}^{t-1} r_j^s \right)$, for $j \in \mathcal{A}$. We denote this distribution by $\text{EW}_\eta(\pi^{t-1}) \in \Delta(\mathcal{A})$.

Indeed, given the above defined loss ℓ^t , the Learner solves the following problem at each round:

$$x^t \in \operatorname{argmin}_{x \in \Delta \mathcal{A}} \max_{r^t \in [0, 1]^{|\mathcal{A}|}} \sum_{j \in \mathcal{A}} \chi_j^t \mathbb{E}_{a \sim x} [r_a^t - r_j^t],$$

where $\chi_j^t = \frac{\exp(\eta \sum_{s=1}^{t-1} (r_{a^s}^s - r_j^s))}{\sum_{i \in \mathcal{A}} \exp(\eta \sum_{s=1}^{t-1} (r_{a^s}^s - r_i^s))} = \frac{\exp(-\eta \sum_{s=1}^{t-1} r_j^s)}{\sum_{i \in \mathcal{A}} \exp(-\eta \sum_{s=1}^{t-1} r_i^s)}$. That is, the per-coordinate weights $(\chi_j^t)_{j \in \mathcal{A}}$ themselves form the Exponential Weights distribution with rate η .

For any choice of r^t by the Adversary, the quantity inside the expectation, $\ell_j^t(a, r^t) = r_a^t - r_j^t$, is *antisymmetric* in a and j : that is, $\ell_j^t(a, r^t) = -\ell_a^t(j, r^t)$. Due to this antisymmetry, no matter which r^t gets selected by the Adversary, by playing $a \sim \text{EW}_\eta(\pi^{t-1})$ the Learner obtains $\mathbb{E}_{a, j \sim \text{EW}_\eta(\pi^{t-1})} [r_a^t - r_j^t] = 0$, thus achieving the value of the game. It is also easy to see that $x^t = \text{EW}_\eta(\pi^{t-1})$ is the unique choice of x^t that guarantees nonnegative value, hence Algorithm 2, when specialized to the external regret setting, is *equivalent* to the Exponential Weights Algorithm 5.

⁴Formally, for any vector of actions' losses r^t , define $a_{r^t}^* := \operatorname{argmin}_{a \in \mathcal{A}} r_a^t$, and notice that

$$\min_{a \in \mathcal{A}} \max_{j \in \mathcal{A}} \ell_j^t(a, r^t) \leq \max_{j \in \mathcal{A}} \ell_j^t(a_{r^t}^*, r^t) = \max_{j \in \mathcal{A}} \left(r_{a_{r^t}^*}^t - r_j^t \right) = \min_{a \in \mathcal{A}} r_a^t - \min_{j \in \mathcal{A}} r_j^t = 0.$$

Hence, the AMF value is indeed nonpositive at each round: $w_A^t = \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \min_{a \in \mathcal{A}} \max_{j \in \mathcal{A}} \ell_j^t(a, r^t) \leq 0$.

4 Multicalibration and Multicalibeating

We now apply our framework to derive an online contextual prediction algorithm which simultaneously satisfies a (potentially very large) family of strong adversarial accuracy and calibration conditions. Namely, given an arbitrarily complex family $\mathcal{G}$ of subsets of the context space (we call them “groups”, a term from the fairness literature), the predictor will be both calibrated and accurate on each group $g \in \mathcal{G}$ (that is, over those online rounds when the context belongs to g).

The accuracy benchmark that we aim to satisfy was recently proposed by Foster and Hart [2021], who called it *calibeating*: given any collection $\mathcal{F}$ of online forecasters, the goal is (intuitively) to “beat” the (squared) error of each $f \in \mathcal{F}$ by at least the *calibration score* of f .

In Section 4.1, we use our framework to rederive the online multigroup calibration (known as *multicalibration*) algorithm of Gupta et al. [2022]. In Section 4.2, we show that by appropriately augmenting the original collection of groups $\mathcal{G}$, this algorithm will, *in addition to* multicalibration, calibrate any family of predictors $f \in \mathcal{F}$ on every group $g \in \mathcal{G}$, which we call *multicalibeating*.

4.1 Multicalibration

Setting There is a *feature* (or *context*) space Θ encoding the set of possible feature vectors representing *individuals* $\theta \in \Theta$. There is also a label space $[0, 1]$. At every round $t \in [T]$:

1. The Adversary announces a particular individual $\theta^t \in \Theta$, whose label is to be predicted;
2. The Learner predicts a label distribution x^t over $[0, 1]$;
3. The Adversary observes x^t , and fixes the true label distribution y^t over $[0, 1]$;
4. The (pure) guessed label $a^t \sim x^t$ and the (pure) true label $b^t \sim y^t$ are sampled.

Objective: Multicalibration The Learner is initially given an arbitrary collection $\mathcal{G} \subseteq 2^\Theta$ of protected population *groups*. Her goal, *multicalibration*, is empirical calibration not just marginally over the whole population, but also conditionally on individual membership in each $g \in \mathcal{G}$. Formally, for any $n \geq 1$ we let the n -*bucketing* of the label interval $[0, 1]$ be its partition into subintervals $[0, 1/n), \dots, [1 - 2/n, 1 - 1/n), [1 - 1/n, 1]$. The i^{th} of these intervals (buckets) is denoted B_n^i .

Definition 4.1 ((α, n) -Multicalibration with respect to $\mathcal{G}$). *Fix a real $\alpha > 0$ and an integer $n \geq 1$. Given the transcript of the interaction $\{(a^t, b^t)\}_{t \in [T]}$, the Learner’s sequence of guessed labels $\{a^t\}_{t \in [T]}$ is (α, n) -multicalibrated with respect to the collection of groups $\mathcal{G}$ if:*

$$\frac{1}{T} \left| \sum_{t=1}^T 1_{\theta^t \in g} \cdot 1_{a^t \in B_n^i} \cdot (b^t - a^t) \right| \leq \alpha, \text{ for every group } g \in \mathcal{G} \text{ and every bucket } B_n^i \text{ (for } i \in [n]).$$

Using our framework, we now derive the guarantee on α that matches that of Gupta et al. [2022].

Theorem 4.1 (Multicalibration). *Fix a family of groups $\mathcal{G}$, a time horizon $T \geq \ln(2|\mathcal{G}|n)$, and any natural $n, r \geq 1$. Then, our framework’s Algorithm 2 can be instantiated as Algorithm 3 to produce (α, n) -multicalibrated predictions w.r.t. $\mathcal{G}$, where α satisfies (over transcript randomness):*

$$\mathbb{E}[\alpha] \leq \frac{1}{rn} + 4\sqrt{\frac{\ln(2|\mathcal{G}|n)}{T}} \quad \text{and} \quad \Pr \left[\alpha \leq \frac{1}{rn} + 8\sqrt{\frac{1}{T} \ln \left(\frac{2|\mathcal{G}|n}{\delta} \right)} \right] \geq 1 - \delta \quad \forall \delta \in (0, 1).$$

Sketch. Setting up the game: The adversary’s strategy space is $\mathcal{Y} = [0, 1]$. The learner will randomize over $\mathcal{A}_r = \{0, 1/(rn), 2/(rn), \dots, 1\}$, for any choice of integer $r \geq 1$ (this will ensure continuity of the loss functions that we are about to define), i.e., her strategy space is $\mathcal{X} = \Delta \mathcal{A}_r$.

Loss functions: The definition of multicalibration consists of $2|\mathcal{G}|n$ constraints (one for each $\pm$ sign, group g , and bucket i) of the following form: $\pm \frac{1}{T} \sum_{t=1}^T 1_{\theta^t \in g} \cdot 1_{a^t \in B_n^i} \cdot (b^t - a^t) \leq \alpha$. Thus, we define (for each $t \in [T]$, $\sigma = \pm 1$, g , and i) a loss function over $(a^t, b^t) \in \mathcal{A}_r \times \mathcal{Y}$ as: $\ell_{i,g,\sigma}^t(a^t, b^t) := \sigma \cdot 1_{\theta^t \in g} \cdot 1_{a^t \in B_n^i} \cdot (b^t - a^t)$. Now, defining a $2|\mathcal{G}|n$ dimensional loss vector $\ell^t := (\ell_{i,g,\sigma}^t)_{i \in [n], g \in \mathcal{G}, \sigma \in \{-1,1\}}$ for each $t \in [T]$ recasts multicalibration in our framework as requiring that $\max_{i \in [n], g \in \mathcal{G}, \sigma \in \{-1,1\}} \sum_{t=1}^T \ell_{i,g,\sigma}^t(a^t, b^t) \leq \alpha T$.

Bounding the AMF regret: To bound the Adversary-Moves-First value with these loss functions, suppose the Adversary announces $b^t \in [0, 1]$. Then, we easily see that by (deterministically) responding with $a^t = \arg\min_{a \in \mathcal{A}_r} |b^t - a|$, for all σ, g, i , $\ell_{i,g,\sigma}^t(a^t, b^t) \leq \frac{1}{2rn}$. Hence,

$$w_A^t = \sup_{b^t \in [0,1]} \min_{x^t \in \Delta \mathcal{A}_r} \max_{i \in [n], g \in \mathcal{G}, \sigma \in \{-1,1\}} \mathbb{E}_{a^t \sim x^t} [\ell_{i,g,\sigma}^t(a^t, b^t)] \leq \frac{1}{2rn} \quad \text{for every } t \in [T].$$

Now, for $T \geq \ln(2|\mathcal{G}|n)$, the AMF regret $R^T = \max_{i \in [n], g \in \mathcal{G}, \sigma \in \{-1,1\}} \sum_{t=1}^T \ell_{i,g,\sigma}^t(a^t, b^t) - \sum_{t=1}^T w_A^t$, by our framework's guarantees, satisfies $\mathbb{E}[R^T] \leq 4\sqrt{T \ln(2|\mathcal{G}|n)}$ over the Learner's randomness. Since $\sum_{t=1}^T w_A^t \leq \frac{T}{2rn}$, we get $\mathbb{E}[\max_{i \in [n], g \in \mathcal{G}, \sigma \in \{-1,1\}} \sum_{t=1}^T \ell_{i,g,\sigma}^t(a^t, b^t)] \leq \frac{T}{2rn} + 4\sqrt{T \ln(2|\mathcal{G}|n)}$.

This gives (α, n) -multicalibration with $\mathbb{E}[\alpha] \leq \frac{1}{T} \left(\frac{T}{2rn} + 4\sqrt{T \ln(2|\mathcal{G}|n)} \right) = \frac{1}{2rn} + 4\sqrt{\frac{\ln(2|\mathcal{G}|n)}{T}}$. The high-probability bound on α is obtained similarly.

Simplifying Learner's algorithm: To attain the AMF value $w_A^t = \frac{1}{2rn}$ at each round, our framework has the Learner solve a linear program (that encodes her minimax strategy). However, she can obtain the *almost* optimal value $\frac{1}{rn}$ *without* solving an LP: this observation gives Algorithm 3 (see Appendix B). The guarantees on α only differ from optimal ones by replacing $\frac{1}{2rn} \rightarrow \frac{1}{rn}$. $\square$

4.2 Multicalibeating

We now give an approach to “beating” arbitrary collections of online forecasters via online multicalibration. The goal, called *calibeating* by Foster and Hart [2021] who introduce the problem, is to make calibrated forecasts that are more accurate than each of an arbitrary set of forecasters, by exactly the calibration error in hindsight of that forecaster. They achieve optimal calibeating bounds for a single forecaster, but their extension to calibeating multiple forecasters incurs at least a polynomial dependence on the number of forecasters. We achieve a logarithmic dependence on the number of forecasters. Additionally, we are able to simultaneously calibeat forecasters on *all* (big enough) subgroups in some set $\mathcal{G}$, with still only a logarithmic dependence on $|\mathcal{G}|$ and the number of forecasters in the group-wise convergence bound. We call this multicalibeating. We now give an overview of our setting, results, and techniques. For full details, see Appendix C.

Setting The Learner (predictor $a = \{a^t\}_{t \in [T]}$) and the Adversary (true labels $b = \{b^t\}_{t \in [T]}$) interact in the same way as in Section 4.1, but the Adversary additionally reveals to the Learner a finite set of forecasters $\mathcal{F}$, where each $f \in \mathcal{F}$ is a function $f : \Theta \rightarrow D_f$. Here $D_f \subset [0, 1]$ is assumed to be a finite set of all possible forecasts that f makes: it will characterize the *level sets* of f . We often suppress the dependence on the transcript, denoting $f^t \in D_f$ the forecast at time t .

The Learner's goal is to “improve on” the forecasts of all $f \in \mathcal{F}$, for some suitable scoring of the predictions. We measure the Learner's and the forecasters' accuracy via the squared error, alternatively known as the Brier score.

Definition 4.2 (Brier Score). *The Brier score of a forecaster f over all rounds $t \in [T]$ is defined as: $\mathcal{B}^f(\pi^T) := \frac{1}{T} \sum_{t \in [T]} (f^t - b^t)^2$.*

The Brier score can be decomposed into so-called *calibration* and *refinement* parts. The former quantifies the extent to which the predictor is calibrated, while the latter expresses the average amount of variance in predictions within every calibration bucket.

To define this decomposition, we need some extra notation. We denote by S_i the subsequence of days on which the Learner’s prediction is in bucket i .⁵ Similarly, $S^d(f)$ (eliding (f) when clear from context) denotes days on which forecaster f predicts d . We let $S_i^d(f) = S_i \cap S^d(f)$. Finally, we use bars to indicate average predictions over given subsequences. For instance, $\bar{a}(S)$ is the Learner’s average prediction over a given subsequence S .

Definition 4.3 (Calibration and Refinement). *The calibration score $\mathcal{K}$ and refinement score $\mathcal{R}$ of a forecaster f over the full transcript π^T are defined as:*

$$\mathcal{K}^f(\pi^T) := \frac{1}{T} \sum_{d \in D_f} |S^d| (d - \bar{b}(S^d))^2, \quad \mathcal{R}^f(\pi^T) := \frac{1}{T} \sum_{d \in D_f} \sum_{t \in S^d} (b^t - \bar{b}(S^d))^2.$$

Fact 1 (Calibration-Refinement Decomposition of Brier Score [DeGroot and Fienberg, 1983]). $\mathcal{B}^f(\pi^T) = \mathcal{K}^f(\pi^T) + \mathcal{R}^f(\pi^T)$.

The goal of calibrating is to beat the forecaster’s Brier score by an amount equal to its calibration score. Or equivalently, to attain a Brier score (almost) equal to the refinement score of the forecaster.

Definition 4.4 (Calibrating). *The Learner’s predictor a is said to τ -calibeat a forecaster f if: $\mathcal{B}^a(\pi^T) \leq \mathcal{R}^f(\pi^T) + \tau$.*

We will now extend the definition of calibrating simultaneously along two natural directions. First, we will want to calibeat multiple forecasters at once. The second extension is that we will want to calibeat the forecasters not just overall, but also on each of the subsequences corresponding to each “population group” $g \in \mathcal{G}$ in a given family of subpopulations $\mathcal{G} \subseteq 2^\Theta$.

Definition 4.5 (Multicalibrating). *Given a family of forecasters $\mathcal{F}$, groups $\mathcal{G} \subseteq 2^\Theta$, and a mapping $\beta : \mathcal{F} \times \mathcal{G} \rightarrow \mathbb{R}_{\geq 0}$, the Learner’s predictor a is an $(\mathcal{F}, \mathcal{G}, \beta)$ -multicalibeater if for every $g \in \mathcal{G}$: $\mathcal{B}^a(\pi^T|_{\{t: \theta^t \in g\}}) \leq \min_{f \in \mathcal{F}} \{ \mathcal{R}^f(\pi^T|_{\{t: \theta^t \in g\}}) + \beta(f, g) \}$*

Note that $(\{f\}, \{\Theta\}, \beta(f, \Theta) := \tau)$ -multicalibrating is equivalent to τ -calibrating a forecaster f .

We first show how to calibeat a single forecaster (Definition 4.4). The modularity of multicalibration will then let us easily extend this result to multiple forecasters and population subgroups.

The idea is to show that if our predictor is multicalibrated with respect to the *level sets* of f , then we achieve calibrating. Hébert-Johnson et al. [2018] give a similar bound in the batch setting. We denote the collection of level sets of f as: $\mathcal{S}(f) := \{\theta \in \Theta : f(\theta) = d\}_{d \in D_f}$.

Theorem 4.2 (Calibrating One Forecaster). *Suppose that the Learner’s predictions a are (α, n) -multicalibrated on the collection of groups $\mathcal{S}(f) \cup \{\Theta\}$. Then the Learner is (α, n) -calibrated on Θ , and she $(\alpha n(|D_f| + 2) + \frac{2}{n})$ -calibeats forecaster f .*

Proof sketch. We show that a has small calibration score, and refinement score close to that of f .

Step 1: Replace $\mathcal{B}^a$ with a surrogate Brier score $\mathcal{B}_n^a$. Consider a (pseudo-)predictor $\tilde{a}$ given by $\tilde{a}^t = \bar{a}(S_{i_{a^t}})$ for $t \in [T]$ (where i_{a^t} is the bucket of a^t). That is, whenever $a^t \in B_n^i$, $\tilde{a}^t$ predicts the average of a over all such rounds $s \in [T]$ that $a^s \in B_n^i$. This is a pseudo-predictor, as the bucket averages of a are unknown until after round T . Thus, $\tilde{a}$ has precisely n level sets, unlike a . Now, we define $\mathcal{B}_n^a, \mathcal{K}_n^a, \mathcal{R}_n^a$ to be the Brier, calibration, and refinement scores of $\tilde{a}$. We can show $\mathcal{B}^a \leq \mathcal{B}_n^a + 1/n$, allowing us to switch to bounding the more manageable Brier loss $\mathcal{B}_n^a = \mathcal{K}_n^a + \mathcal{R}_n^a$.

Step 2: Bound the surrogate calibration score $\mathcal{K}_n^a$. Since the Learner is (α, n) -calibrated on the domain Θ , the calibration error per level set is at most α . There are n level sets, so $\mathcal{K}_n^a \leq \alpha n$.

Step 3: Bound the surrogate refinement score $\mathcal{R}_n^a$. We connect $\mathcal{R}^f$ and $\mathcal{R}_n^a$ via a *joint* refinement score: $\mathcal{R}^{f \times a}$, which measures the average variance of the partition generated by all intersections of the level sets of a and f . The finer the partition, the smaller the refinement score, so $\mathcal{R}^f \geq \mathcal{R}^{f \times a}$. Next, informally, multicalibration ensures that a has already “captured” most of the variance explained by f . Therefore, refining a ’s level sets by f does little to reduce variance. More precisely, we show that $\mathcal{R}_n^a \leq \mathcal{R}^{f \times a} + \alpha n(|D_f| + 1) + \frac{1}{n}$. Combining with our previous inequality, we have: $\mathcal{R}_n^a \leq \mathcal{R}^f + \alpha n(|D_f| + 1) + \frac{1}{n}$.

Combining the above, we get: $\mathcal{B}^a \leq \mathcal{R}_n^a + \mathcal{K}_n^a + \frac{1}{n} \leq (\mathcal{R}^f + \alpha n(|D_f| + 1) + \frac{1}{n}) + \alpha n + \frac{1}{n}$. $\square$

⁵Note that S_i depends implicitly on the bucketing parameter n and the transcript π^T .

Calibeating many forecasters Generalizing the above construction, we can easily calibeat any collection of forecasters $\mathcal{F}$ on the entire context space Θ : it suffices to ask for multicalibration with respect to the level sets of *all* forecasters, i.e. $\left(\bigcup_{f \in \mathcal{F}} \mathcal{S}(f)\right) \cup \{\Theta\}$. Theorem 4.2 applies separately to each f ; the only degradation in the guarantees will come in the form of a larger α , since we are asking for multicalibration with respect to more groups than before. But this effect will be small, since α depends on the number of required groups $|\mathcal{G}'|$ as $O(\sqrt{\ln |\mathcal{G}'|})$. See Corollary C.2.

However, to fully satisfy Definition 4.5 of multicalibeating, we need to calibeat all $f \in \mathcal{F}$ on *all* groups $g \in \mathcal{G}$ in a given collection $\mathcal{G} \subseteq 2^\Theta$. For that, we simply extend the above construction by requiring multicalibration with respect to all pairwise *intersections* of the forecasters' level sets with the groups $g \in \mathcal{G}$. By further augmenting this collection with the protected groups $\mathcal{G}$ themselves, we finally achieve our ultimate goal: *simultaneous* multicalibeating and multicalibration.

Theorem 4.3 (Multicalibeating + Multicalibration). *Let $\mathcal{G} \subseteq 2^\Theta$, and $\mathcal{F}$ some set of forecasters $f : \Theta \rightarrow D_f$. The multicalibration algorithm on $\mathcal{G}' := \left(\bigcup_{f \in \mathcal{F}} \{g \cap S : (g, S) \in \mathcal{G} \times \mathcal{S}(f)\}\right) \cup \mathcal{G}$ with parameters $r, n \geq 1$, after T rounds, attains expected $(\mathcal{F}, \mathcal{G}, \beta)$ -multicalibeating, where: ⁶ $\mathbb{E}[\beta(f, g)] \leq \frac{2}{n} + \frac{|D_f|+2}{r \cdot |S(g)|/T} + 4n(|D_f| + 2)\sqrt{\frac{1}{|S(g)|^2/T} \ln \left(2n|\mathcal{G}|(1 + \sum_f |D_f|)\right)} \forall f \in \mathcal{F}, g \in \mathcal{G}$, while maintaining (α, n) -multicalibration on $\mathcal{G}$, with: $\mathbb{E}[\alpha] \leq \frac{1}{rn} + 4\sqrt{\frac{1}{T} \ln \left(2n|\mathcal{G}|(1 + \sum_f |D_f|)\right)}$.*

In particular, for any group g occurring more than $T^{-1/2}$ of the time, we asymptotically converge to $\frac{1}{n}$ -calibeating as $T \rightarrow \infty$, thus combining the goals of online multicalibration and multigroup regret.

5 Polytope Blackwell Approachability

Consider a setting where the Learner and the Adversary are playing a repeated game with *vector-valued* payoffs, in which the Learner always goes first and aims to force the average payoff over the entire interaction to approach a given convex set. Blackwell's Theorem [1956] states that a convex set is approachable if and only if it is *response-satisfiable* (roughly, for any choice of the Adversary, the Learner has a response forcing the one-round payoff inside the convex set). The rate of approachability typically depends on the dimension of the payoff vectors.

This is a specialization of our framework to a case in which the environment is fixed at every round. Thus our framework can be used to obtain a *dimension-independent* rate bound in the fundamental case where the approachable set is a convex *polytope*. Our bound is only logarithmic in the polytope's number of facets, and is achievable via an efficient convex-programming based algorithm.

Let us formalize our setting. In rounds $t = 1, 2, \dots$, the Learner and the Adversary play a repeated game. Their respective pure strategy sets are $\mathcal{A}$ and $\mathcal{Y}$, where $\mathcal{A}$ is a finite set and $\mathcal{Y} \subseteq \mathbb{R}^m$ (for some integer $m \geq 1$) is convex and compact. The game's utility function is λ -dimensional (for some integer $\lambda \geq 1$), continuous, concave in the second argument, and is denoted by $u : \mathcal{A} \times \mathcal{Y} \rightarrow \mathbb{R}^\lambda$. At each round t , the Learner plays a mixed strategy $x^t \in \Delta \mathcal{A}$, the Adversary responds with some $y^t \in \mathcal{Y}$, and the Learner then samples a pure action $a^t \sim x^t$. This gives rise to the utility vector $u(a^t, y^t)$. The *average play* up to any round $t \geq 1$ is then defined as $\bar{u}^t = \frac{1}{t} \sum_{s=1}^t u(a^s, y^s)$.

The target convex set that the Learner wants to approach is a polytope $\mathcal{P}(\mathcal{H}) \subseteq \mathbb{R}^\lambda$, defined as the intersection of a finite collection of halfspaces $\mathcal{H} = (h_{\alpha, \beta})$, where for any given $\alpha \in \mathbb{R}^\lambda, \beta \in \mathbb{R}$ we denote $h_{\alpha, \beta} = \{x \in \mathbb{R}^\lambda : \langle \alpha, x \rangle - \beta \leq 0\}$. Finally, by way of normalization, consider any two dual norms $\|\cdot\|_p$ and $\|\cdot\|_q$. We require, first, that $\|\alpha\|_p \leq 1$ and $|\beta| \leq 1$ for each halfspace $h_{\alpha, \beta} \in \mathcal{H}$; and second, that the payoffs be in the $\|\cdot\|_q$ -unit ball: $\|u(a, y)\|_q \leq 1$ for $a \in \mathcal{A}, y \in \mathcal{Y}$.

⁶ $S(g)$ denotes the subsequence of days on which a group g occurs, suppressing dependence on transcript.

Theorem 5.1 (Polytope Blackwell Approachability). *Suppose the target convex polytope $\mathcal{P}(\mathcal{H})$ is response-satisfiable, in the sense that for any Adversary’s action $y \in \mathcal{Y}$, the Learner has a mixed response $x \in \Delta\mathcal{A}$ that places the expected payoff inside $\mathcal{P}(\mathcal{H})$: that is, $\mathbb{E}_{a \sim x}[u(a, y)] \in \mathcal{P}(\mathcal{H})$.*

Then, $\mathcal{P}(\mathcal{H})$ is approachable, both in expectation and with high probability with respect to the transcript of the interaction. Namely, the Learner has an efficient convex programming based algorithm which guarantees both following conditions simultaneously:

1. *For any margin $\epsilon > 0$, the average play $\bar{u}^t$ up to any round $t \geq \frac{64 \ln |\mathcal{H}|}{\epsilon^2}$ will satisfy $\mathbb{E} [\max_{h_{\alpha, \beta} \in \mathcal{H}} (\langle \alpha, \bar{u}^t \rangle - \beta)] \leq \epsilon$.*
2. *For any $\delta \in (0, 1)$, the average play $\bar{u}^t$ up to any round $t \geq \ln |\mathcal{H}|$ will satisfy $\max_{h_{\alpha, \beta} \in \mathcal{H}} (\langle \alpha, \bar{u}^t \rangle - \beta) \leq 16 \sqrt{\frac{1}{t} \ln \left(\frac{|\mathcal{H}|}{\delta} \right)}$ with probability at least $1 - \delta$.*

Acknowledgments

We thank Ira Globus-Harris, Chris Jung, and Kunal Talwar for helpful conversations at an early stage of this work. Supported in part by NSF grants AF-1763307, FAI-2147212, CCF-2217062, and CCF-1934876 and the Simons collaboration on algorithmic fairness.

References

- Dmitry Adamskiy, Wouter M Koolen, Alexey Chernov, and Vladimir Vovk. A closer look at adaptive regret. In *International Conference on Algorithmic Learning Theory*, pages 290–304. Springer, 2012.
- Yossi Azar, Uriel Felge, Michal Feldman, and Moshe Tennenholtz. Sequential decision making with vector outcomes. In *Proceedings of the 5th conference on Innovations in Theoretical Computer Science*, pages 195–206, 2014.
- David Blackwell. An analog of the minimax theorem for vector payoffs. *Pacific Journal of Mathematics*, 6(1):1–8, 1956.
- Avrim Blum. Empirical support for winnow and weighted-majority algorithms: Results on a calendar scheduling domain. *Machine Learning*, 26(1):5–23, 1997.
- Avrim Blum and Thodoris Lykouris. Advancing subgroup fairness via sleeping experts. In *Innovations in Theoretical Computer Science Conference (ITCS)*, volume 11, 2020.
- Avrim Blum and Yishay Mansour. From external to internal regret. *Journal of Machine Learning Research*, 8(6), 2007.
- Morris H. DeGroot and Stephen E. Fienberg. The comparison and evaluation of forecasters. *The Statistician*, 32:12–22, 1983.
- Nadejda Drenska and Robert V Kohn. Prediction with expert advice: A pde perspective. *Journal of Nonlinear Science*, 30(1):137–173, 2020.
- Devdatt P Dubhashi and Alessandro Panconesi. *Concentration of measure for the analysis of randomized algorithms*. Cambridge University Press, 2009.
- Dean P Foster and Sergiu Hart. “calibeating”: Beating forecasters at their own game. <https://www.ma.huji.ac.il/~hart/papers/calib-beat.pdf>, 2021.
- Dean P Foster and Rakesh V Vohra. Asymptotic calibration. *Biometrika*, 85(2):379–390, 1998.

- Yoav Freund, Robert E Schapire, Yoram Singer, and Manfred K Warmuth. Using and combining predictors that specialize. In *Proceedings of the twenty-ninth annual ACM symposium on Theory of computing*, pages 334–343, 1997.
- Drew Fudenberg and David K Levine. An easier way to calibrate. *Games and Economic Behavior*, 29(1-2): 131–137, 1999.
- Ira Globus-Harris, Michael Kearns, and Aaron Roth. Beyond the frontier: Fairness without accuracy loss. *arXiv preprint arXiv:2201.10408*, 2022.
- Amy Greenwald and Amir Jafari. A general class of no-regret learning algorithms and game-theoretic equilibria. In *Learning theory and kernel machines*, pages 2–12. Springer, 2003.
- Varun Gupta, Christopher Jung, Georgy Noarov, Mallesh M. Pai, and Aaron Roth. Online Multivalid Learning: Means, Moments, and Prediction Intervals. In *13th Innovations in Theoretical Computer Science Conference (ITCS 2022)*, pages 82:1–82:24, 2022.
- Sergiu Hart. Calibrated forecasts: The minimax proof, 2020. URL <http://www.ma.huji.ac.il/~hart/papers/calib-minmax.pdf>.
- Sergiu Hart and Andreu Mas-Colell. A simple adaptive procedure leading to correlated equilibrium. *Econometrica*, 68(5):1127–1150, 2000.
- Elad Hazan and Comandur Seshadhri. Efficient learning algorithms for changing environments. In *Proceedings of the 26th Annual International Conference on Machine Learning*, pages 393–400, 2009.
- Ursula Hébert-Johnson, Michael Kim, Omer Reingold, and Guy Rothblum. Multicalibration: Calibration for the (computationally-identifiable) masses. In *International Conference on Machine Learning*, pages 1939–1948. PMLR, 2018.
- Thomas Kesselheim and Sahil Singla. Online learning with vector costs and bandits with knapsacks. In *Conference on Learning Theory*, pages 2286–2305. PMLR, 2020.
- Robert Kleinberg, Alexandru Niculescu-Mizil, and Yogeshwer Sharma. Regret bounds for sleeping experts and bandits. *Machine Learning*, 80(2):245–272, 2010.
- Vladimir A Kobzar, Robert V Kohn, and Zhilei Wang. New potential-based bounds for prediction with expert advice. In *Conference on Learning Theory*, pages 2370–2405. PMLR, 2020.
- Ehud Lehrer. A wide range no-regret theorem. *Games and Economic Behavior*, 42(1):101–115, 2003.
- Nick Littlestone and Manfred K Warmuth. The weighted majority algorithm. *Information and computation*, 108(2):212–261, 1994.
- Shie Mannor, Vianney Perchet, and Gilles Stoltz. Approachability in unknown games: Online learning meets multi-objective optimization. In *Conference on Learning Theory*, pages 339–355. PMLR, 2014a.
- Shie Mannor, Vianney Perchet, and Gilles Stoltz. Set-valued approachability and online learning with partial monitoring. *The Journal of Machine Learning Research*, 15(1):3247–3295, 2014b.
- Vianney Perchet. Exponential weight approachability, applications to calibration and regret minimization. *Dynamic Games and Applications*, 5(1):136–153, 2015.
- Vianney Perchet and Shie Mannor. Approachability, fast and slow. In *Conference on Learning Theory*, pages 474–488. PMLR, 2013.
- T.E.S. Raghavan. Zero-sum two-person games. In R.J. Aumann and S. Hart, editors, *Handbook of Game Theory with Economic Applications*, volume 2 of *Handbook of Game Theory with Economic Applications*, chapter 20, pages 735–768. Elsevier, 1994. URL <https://ideas.repec.org/h/eee/gamchp/2-20.html>.

Alexander Rakhlin, Karthik Sridharan, and Ambuj Tewari. Online learning: Random averages, combinatorial parameters, and learnability. *Advances in Neural Information Processing Systems*, 23:1984–1992, 2010.

Alexander Rakhlin, Karthik Sridharan, and Ambuj Tewari. Online learning: Beyond regret. In *Proceedings of the 24th Annual Conference on Learning Theory*, pages 559–594. JMLR Workshop and Conference Proceedings, 2011.

Alexander Rakhlin, Ohad Shamir, and Karthik Sridharan. Relax and randomize: from value to algorithms. In *Proceedings of the 25th International Conference on Neural Information Processing Systems-Volume 2*, pages 2141–2149, 2012.

Guy N Rothblum and Gal Yona. Multi-group agnostic pac learnability. *arXiv preprint arXiv:2105.09989*, 2021.

Volodimir G Vovk. Aggregating strategies. *Proc. of Computational Learning Theory, 1990*, 1990.

A The General Framework with Extensions to Probabilistic and Approximate Learners: Full Proofs and Algorithms

A.1 Omitted Proofs from Section 2

Proof of Lemma 2.1. After taking the log and dividing by η , this lemma follows from the following chain:

$$\exp(\eta R^T) = \exp\left(\eta \max_{j \in [d]} R_j^T\right) = \exp\left(\max_{j \in [d]} \eta R_j^T\right) = \max_{j \in [d]} \exp(\eta R_j^T) \leq \sum_{j \in [d]} \exp(\eta R_j^T) = L^T.$$

□

Proof of Lemma 2.2. By definition of the surrogate loss, we have:

$$\begin{aligned} L^t - L^{t-1} &= \sum_{j \in [d]} \exp(\eta R_j^t) - \sum_{j \in [d]} \exp(\eta R_j^{t-1}), \\ &= \sum_{j \in [d]} \exp(\eta R_j^{t-1} + \eta(\ell_j^t(x^t, y^t) - w_A^t)) - \sum_{j \in [d]} \exp(\eta R_j^{t-1}), \\ &= \sum_{j \in [d]} \exp(\eta R_j^{t-1}) (\exp(\eta(\ell_j^t(x^t, y^t) - w_A^t)) - 1). \end{aligned}$$

Using the fact that $\exp(x) - 1 \leq x + x^2$ for $|x| \leq 1$, we have, for $\eta \cdot 2C \leq 1$,

$$\begin{aligned} &\leq \sum_{j \in [d]} \exp(\eta R_j^{t-1}) \left(\eta(\ell_j^t(x^t, y^t) - w_A^t) + \eta^2(\ell_j^t(x^t, y^t) - w_A^t)^2 \right), \\ &\leq \eta \sum_{j \in [d]} \exp(\eta R_j^{t-1}) (\ell_j^t(x^t, y^t) - w_A^t) + \eta^2(2C)^2 L^{t-1}. \end{aligned}$$

□

Proof of Lemma 2.3. We begin by recalling that $L^0 = d$. Thus, the desired bound on L^T follows via Lemma 2.2 and a telescoping argument, if only we can show that for every $t \in [T]$ the Learner has an action $x^t \in \mathcal{X}^t$ which guarantees that for any $y^t \in \mathcal{Y}^t$,

$$\eta \sum_{j \in [d]} \exp(\eta R_j^{t-1}) (\ell_j^t(x^t, y^t) - w_A^t) \leq 0.$$

To this end, we define a zero-sum game between the Learner and the Adversary, with action space $\mathcal{X}^t$ for the Learner and $\mathcal{Y}^t$ for the Adversary, and with the objective function (which the Adversary wants to maximize and the Learner wants to minimize):

$$u^t(x, y) := \sum_{j \in [d]} \exp(\eta R_j^{t-1}) (\ell_j^t(x, y) - w_A^t), \text{ for all } x \in \mathcal{X}^t, y \in \mathcal{Y}^t.$$

Recall from the definition of our framework that $\mathcal{X}^t, \mathcal{Y}^t$ are convex, compact and finite-dimensional, as well as that each ℓ_j^t is continuous, convex in the first argument, and concave in the second argument. Since u^t is defined as an affine function of the individual coordinate functions ℓ_j^t , u^t is also convex-concave and continuous. This means that we may invoke Sion's Minimax Theorem:

Fact 2 (Sion's Minimax Theorem). *Given finite-dimensional convex compact sets $\mathcal{X}, \mathcal{Y}$, and a continuous function $f : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ which is convex in the first argument and concave in the second argument, it holds that*

$$\min_{x \in \mathcal{X}} \max_{y \in \mathcal{Y}} f(x, y) = \max_{y \in \mathcal{Y}} \min_{x \in \mathcal{X}} f(x, y).$$

Using Sion's Theorem to switch the order of play (so that the Adversary is compelled to move first), and then recalling the definition of w_A^t (the value of the maximum coordinate value of ℓ^t that the Learner can obtain when the Adversary is compelled to move first), we obtain:⁷

$$\begin{aligned} \min_{x^t \in \mathcal{X}^t} \max_{y^t \in \mathcal{Y}^t} u^t(x^t, y^t) &= \max_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} u^t(x^t, y^t) \\ &= \max_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} \sum_{j' \in [d]} \exp(\eta R_{j'}^{t-1}) \cdot (\ell_{j'}^t(x^t, y^t) - w_A^t), \\ &\leq \sup_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} \sum_{j' \in [d]} \exp(\eta R_{j'}^{t-1}) \cdot \max_{j \in [d]} (\ell_j^t(x^t, y^t) - w_A^t), \\ &= \sum_{j' \in [d]} \exp(\eta R_{j'}^{t-1}) \cdot \sup_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} \max_{j \in [d]} (\ell_j^t(x^t, y^t) - w_A^t), \\ &= \sum_{j' \in [d]} \exp(\eta R_{j'}^{t-1}) \cdot (w_A^t - w_A^t), \\ &= 0. \end{aligned}$$

Thus, the Learner can ensure that $L^t \leq (4\eta^2 C^2 + 1) L^{t-1}$ by playing at every round t :

$$x^t \in \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} u^t(x, y).$$

This concludes the proof. □

An equivalent description of Learner's space of minimax optimal strategies at each round t
We observe that the Learner's optimal action at each round, derived in the proof, can be expressed without

⁷Note that in the third step, $\max_{y^t \in \mathcal{Y}^t}$ turns into $\sup_{y^t \in \mathcal{Y}^t}$. This is because after each $(\ell_{j'}^t(x^t, y^t) - w_A^t)$ is replaced with $\max_j (\ell_j^t(x^t, y^t) - w_A^t)$, the maximum over y generally becomes unachievable (recall Footnote 1).

any reference to the quantities w_A^t :

$$\begin{aligned}
x^t &\in \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \exp(\eta R_j^{t-1}) (\ell_j^t(x, y) - w_A^t), \\
&= \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \exp(\eta R_j^{t-1}) \ell_j^t(x, y), \\
&= \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \frac{\exp\left(\eta \sum_{s=1}^{t-1} \ell_j^s(x^s, y^s)\right) \ell_j^t(x, y)}{\exp\left(\eta \sum_{s=1}^{t-1} w_A^s\right)}, \\
&= \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \exp\left(\eta \sum_{s=1}^{t-1} \ell_j^s(x^s, y^s)\right) \ell_j^t(x, y), \\
&= \operatorname{argmin}_{x \in \mathcal{X}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \frac{\exp\left(\eta \sum_{s=1}^{t-1} \ell_j^s(x^s, y^s)\right)}{\sum_{i \in [d]} \exp\left(\eta \sum_{s=1}^{t-1} \ell_i^s(x^s, y^s)\right)} \ell_j^t(x, y).
\end{aligned}$$

The weights placed on the loss coordinates $\ell_j^s(x^t, y^t)$ in the final expression form a probability distribution which should remind the reader of the well known Exponential Weights distribution.

A.2 Extensions

Before presenting applications of our framework, we pause to discuss two natural extensions that are called for in some of our applications. Both extensions only require very minimal changes to the notation in Section 2.1 and to the general algorithmic framework in Section 2.2.

We begin by discussing, in Section A.2.1, how to adapt our framework to the setting where the Learner is allowed to randomize at each round amongst a finite set of actions, and wishes to obtain probabilistic guarantees for her AMF regret with respect to her randomness. This will be useful in all three of our applications.

We then proceed to show, in Section A.2.2, that our AMF regret bounds are robust to the case in which at each round, the Learner, who is playing according to the general Algorithm 1 given above, computes and plays according to an approximate (rather than exact) minimax strategy. This is useful for settings where it may be desirable (for computational or other reasons) to implement our algorithmic framework approximately, rather than exactly. In particular, in one of our applications — mean multicalibration, which is discussed in Section 4.1 — we will illustrate this point by deriving a multicalibration algorithm that has the Learner play only extremely (computationally and structurally) simple strategies, at the cost of adding an arbitrarily small term to the multicalibration bounds, compared to the Learner that plays the exact minimax equilibrium.

A.2.1 Performance Bounds for a Probabilistic Learner

So far, we have described the interaction between the Learner and the Adversary as deterministic. In many applications, however, the convex action space for the Learner is the simplex over some finite set of base actions, representing *probability distributions* over actions. In this case, the Adversary chooses his action in response to the *probability distribution* over base actions chosen by the Learner, at which point the Learner samples a single base action from her chosen distribution.

We will use the following notation. The Learner's pure action set at time t is denoted by $\mathcal{A}^t$. Before each round t , the Adversary reveals a vector valued loss function $\ell^t : \mathcal{A}^t \times \mathcal{Y}^t \rightarrow [-C, C]^d$. At the beginning of round t , the Learner chooses a probabilistic mixture over her action set $\mathcal{A}^t$, which we will usually denote as $x^t \in \Delta \mathcal{A}^t$; after the Adversary has made his move, the Learner samples her pure action a^t for the round, which is recorded into the transcript of the interaction.

The redefined vector valued losses ℓ^t now take as their first argument a *pure* action $a \in \mathcal{A}^t$. We extend this to $\mathcal{X}^t := \Delta\mathcal{A}^t$ as $\ell^t(x^t, y^t) := \mathbb{E}_{a^t \sim x^t}[\ell^t(a^t, y^t)]$ for any $x^t \in \Delta\mathcal{A}^t$. In this notation, holding the second argument fixed, the loss function is linear (hence convex and continuous) and has a convex, compact domain (the simplex $\Delta\mathcal{A}^t$). Using this extended notation, it is now easy to see how to define the probabilistic analog of the AMF value.

Definition A.1 (Probabilistic AMF Value).

$$w_A^t := \sup_{y^t \in \mathcal{Y}^t} \min_{x^t \in \mathcal{X}^t} \max_{j \in [d]} \ell_j^t(x^t, y^t) = \sup_{y^t \in \mathcal{Y}^t} \min_{x^t \in \Delta\mathcal{A}^t} \max_{j \in [d]} \mathbb{E}_{a^t \sim x^t} [\ell_j^t(a^t, y^t)].$$

For a more detailed discussion of the probabilistic setting, please refer to Appendix A.3.

Adapting the algorithm to the probabilistic Learner setting Above, Algorithm 1 was given for the deterministic case of our framework. In the probabilistic setting, when computing the probability distribution for the current round, the Learner should take into account the *realized* losses from the past rounds. We present the modified algorithm below.

Algorithm 2: General Algorithm for the *Probabilistic* Learner

for rounds $t = 1, \dots, T$ **do**

 Learn adversarially chosen $\mathcal{A}^t, \mathcal{Y}^t$, and vector loss function $\ell^t(\cdot, \cdot) : \mathcal{A}^t \times \mathcal{Y}^t \rightarrow [-C, C]^d$.

 Let

$$\chi_j^t := \frac{\exp\left(\eta \sum_{s=1}^{t-1} \ell_j^s(a^s, y^s)\right)}{\sum_{i \in [d]} \exp\left(\eta \sum_{s=1}^{t-1} \ell_i^s(a^s, y^s)\right)} \text{ for } j \in [d].$$

 Select a mixed action $x^t \in \Delta\mathcal{A}^t$, where

$$x^t \in \operatorname{argmin}_{x \in \Delta\mathcal{A}^t} \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \chi_j^t \cdot \ell_j^t(x, y).$$

 Observe the Adversary's selection of $y^t \in \mathcal{Y}^t$.

 Sample pure action $a^t \sim x^t$.

Probabilistic performance guarantees Algorithm 2 provides two crucial blackbox guarantees to the probabilistic Learner. First, the guarantees on Algorithm 1 from Theorem 2.1 almost immediately translate into a bound on the *expected* AMF regret of the Learner who uses Algorithm 2, over the randomness in her actions. Second, a *high-probability* AMF regret bound, also over the Learner's randomness, can be derived in a straightforward way.

Theorem A.1 (In-Expectation Bound). *Given $T \geq \ln d$, Algorithm 2 with learning rate $\eta = \sqrt{\frac{\ln d}{4TC^2}}$ guarantees that ex-ante, with respect to the randomness in the Learner's realized outcomes, the expected AMF regret is bounded as:*

$$\mathbb{E}[R^T] \leq 4C\sqrt{T \ln d}.$$

Proof Sketch. Using Jensen's inequality to switch expectations and exponentials, it is easy to modify the proof of Lemma 2.1 to obtain the following in-expectation bound:

$$\mathbb{E}[R^T] \leq \frac{\ln \mathbb{E}[L^T]}{\eta}.$$

The rest of the proof is similar to the proofs of Lemma 2.2 and Lemma 2.3. □

Theorem A.2 (High-Probability Bound). *Fix any $\delta \in (0, 1)$. Given $T \geq \ln d$, Algorithm 2 with learning rate $\eta = \sqrt{\frac{\ln d}{4TC^2}}$ guarantees that the AMF regret will satisfy, with ex-ante probability $1 - \delta$ over the randomness in the Learner’s realized outcomes,*

$$R^T \leq 8C \sqrt{T \ln \left(\frac{d}{\delta} \right)}.$$

Proof Sketch. The proof proceeds by constructing a martingale with bounded increments that tracks the increase in the surrogate loss L^T , and then using Azuma’s inequality to conclude that the final surrogate loss (and hence the AMF regret) is bounded above with high probability. For a detailed proof, see Appendix A.3. $\square$

A.2.2 Performance Bounds for a Suboptimal Learner

Our general Algorithms 1 and 2 involve the Learner solving a convex program at each round in order to identify her minimax optimal strategy. However, in some applications of our framework it may be necessary or desirable for the Learner to restrict herself to playing *approximately* minimax optimal strategies instead of exactly optimal ones. This can happen for a variety of reasons:

1. *Computational efficiency.* While the convex program that the Learner must solve at each round is polynomial-sized in the description of the environment, one may wish for a better running time dependence — e.g. in settings in which the action space for the Learner is exponential in some other relevant parameter of the problem. In such cases, we will want to trade off run-time for approximation error in the minimax equilibrium computation at each round.
2. *Structural simplicity of strategies.* One may wish to restrict the Learner to only playing “simple” strategies (for example, distributions over actions with small support), or more generally, strategies belonging to a certain predefined strict subset of the Learner’s strategy space. This subset may only contain approximately optimal minimax strategies.
3. *Numerical precision.* As the convex programs solved by the Learner at each round generally have irrational coefficients (due to the exponents), using finite-precision arithmetic to solve these programs will lead to a corresponding precision error in the solution, making the computed strategy only approximately minimax optimal for the Learner. This kind of approximation error can generally be driven to be arbitrarily small, but still necessitates being able to reason about approximate solutions.

Given a suboptimal instantiation of Algorithm 1 or 2, we thus want to know: how much worse will its achieved regret bound be, compared to the existential guarantee? We will now address this question for both the deterministic setting of Sections 2.1 and 2.2, and the probabilistic setting of Section A.2.1.

Recall that at each round $t \in [T]$, both Algorithm 1 and Algorithm 2 (with the weights χ_j^t defined accordingly) have the Learner solve for the minimizer x of the function $\psi^t : \mathcal{X}^t \rightarrow [-C, C]$ defined as:

$$\psi^t(x) := \max_{y \in \mathcal{Y}^t} \sum_{j \in [d]} \chi_j^t \cdot \ell_j^t(x, y).$$

The range of ψ^t is $[-C, C]$ as indicated, since it is a linear combination of loss coordinates $\ell_j^t(x, y) \in [-C, C]$, where the weights $(\chi_1^t, \dots, \chi_d^t)$ form a probability distribution over $[d]$.

Now suppose the Learner ends up playing actions $x^1, \dots, x^T$ which do not necessarily minimize the respective objectives $\psi^t(\cdot)$. The following definition helps capture the degree of suboptimality in the Learner’s play at each round.

Definition A.2 (Achieved AMF Value Bound). *Consider any round $t \in [T]$, and suppose the Learner plays action $x^t \in \mathcal{X}^t$ at round t . Then, any number*

$$w_{\text{bd}}^t \in [\psi^t(x^t), C]$$

is called an achieved AMF value bound for round t .

This definition has two aspects. Most importantly, w_{bd}^t upper bounds the Learner's *achieved* objective function value at round t . Furthermore, we restrict w_{bd}^t to be $\leq C$ — otherwise it would be a meaningless bound as the Learner gets objective value $\leq C$ no matter what x^t she plays.

We now formulate the desired bounds on the performance of a suboptimal Learner. The upshot is that for a suboptimal Learner, the bounds of Theorems 2.1, A.1, A.2 hold with each w_A^t replaced with the corresponding achieved AMF bound w_{bd}^t .

Theorem A.3 (Bounds for a Suboptimal Learner). *Consider a Learner who does not necessarily play optimally at all rounds, and a sequence $w_{\text{bd}}^1, \dots, w_{\text{bd}}^T$ of achieved AMF value bounds.*

In the deterministic setting, the Learner achieves the following regret bound analogous to Theorem 2.1:

$$\max_{j \in [d]} \sum_{t=1}^T \ell_j^t(x^t, y^t) \leq \sum_{t=1}^T w_{\text{bd}}^t + 4C\sqrt{T \ln d}.$$

In the probabilistic setting, the Learner achieves the following in-expectation regret bound analogous to Theorem A.1:

$$\mathbb{E} \left[\max_{j \in [d]} \sum_{t=1}^T \ell_j^t(a^t, y^t) \right] \leq \sum_{t=1}^T w_{\text{bd}}^t + 4C\sqrt{T \ln d},$$

and the following high-probability bound analogous to Theorem A.2:

$$\max_{j \in [d]} \sum_{t=1}^T \ell_j^t(a^t, y^t) \leq \sum_{t=1}^T w_{\text{bd}}^t + 8C\sqrt{T \ln \left(\frac{d}{\delta} \right)} \text{ with probability } \geq 1 - \delta, \text{ for any } \delta \in (0, 1).$$

Proof Sketch. We use the deterministic case for illustration. The main idea is to redefine the Learner's regret to be relative to her achieved AMF value bounds $(w_{\text{bd}}^t)_{t \in [T]}$ rather than the AMF values $(w_A^t)_{t \in [T]}$. Namely, we let $R_{\text{bd}}^t := \max_{j \in [d]} (R_{\text{bd}}^t)_j$, where $(R_{\text{bd}}^t)_j := \sum_{s=1}^t \ell_j^s(x^s, y^s) - \sum_{s=1}^t w_{\text{bd}}^s$. The surrogate loss is defined in the same way as before, namely $L_{\text{bd}}^t := \sum_{j \in [d]} \exp(\eta \cdot (R_{\text{bd}}^t)_j)$.

First, Lemma 2.1 still holds: $R_{\text{bd}}^T \leq (\ln L_{\text{bd}}^T) / \eta$, with the same proof. Lemma 2.2 also holds after replacing each w_A^t with w_{bd}^t : namely, $L_{\text{bd}}^t \leq (4\eta^2 C^2 + 1) L_{\text{bd}}^{t-1} + \eta \sum_{j \in [d]} \exp(\eta (R_{\text{bd}}^{t-1})_j) \cdot (\ell_j^t(x^t, y^t) - w_{\text{bd}}^t)$. The proof is almost the same: we formerly used $w_A^t \leq C$, and now use that $w_{\text{bd}}^t \leq C$ by Definition A.2.

Now, following the proofs of Lemma 2.3 and Theorem 2.1, to obtain the declared regret bound it suffices to show for $t \in [T]$ that the Learner's action x^t guarantees $\sum_{j \in [d]} \exp(\eta (R_{\text{bd}}^{t-1})_j) \cdot (\ell_j^t(x^t, y^t) - w_{\text{bd}}^t) \leq 0$, no matter what y^t is played by the Adversary. For any $y^t \in \mathcal{Y}^t$, we can rewrite this objective as:

$$\sum_{j \in [d]} \exp(\eta (R_{\text{bd}}^t)_j) \cdot (\ell_j^t(x^t, y^t) - w_{\text{bd}}^t) = \frac{\sum_{i \in [d]} \exp(\eta \sum_{s=1}^{t-1} \ell_i^s(x^s, y^s))}{\exp(\sum_{s=1}^{t-1} w_{\text{bd}}^s)} \sum_{j \in [d]} \chi_j^t \cdot (\ell_j^t(x^t, y^t) - w_{\text{bd}}^t).$$

It now follows that action x^t achieves $\sum_{j \in [d]} \exp(\eta (R_{\text{bd}}^{t-1})_j) \cdot (\ell_j^t(x^t, y^t) - w_{\text{bd}}^t) \leq 0$, from observing that:

$$\sum_{j \in [d]} \chi_j^t \cdot (\ell_j^t(x^t, y^t) - w_{\text{bd}}^t) = \sum_{j \in [d]} \chi_j^t \cdot \ell_j^t(x^t, y^t) - w_{\text{bd}}^t \leq \psi^t(x^t) - w_{\text{bd}}^t \leq 0,$$

where the final inequality holds since the Learner achieves AMF value bound w_{bd}^t at round t . $\square$

A.3 Omitted Proofs and Details from Section A.2.1: Bounds for the Probabilistic Learner

First, we define our probabilistic setting, emphasizing the differences to the deterministic protocol. At each round $t \in [T]$, the interaction between the Learner and the Adversary proceeds as follows:

1. At the beginning of each round t , the Adversary selects an environment consisting of the following, and reveals it to the Learner:
 - (a) The Learner's *simplex action set* $\mathcal{X}^t = \Delta\mathcal{A}^t$, where $\mathcal{A}^t$ is a finite set of pure actions;
 - (b) The Adversary's convex compact action set $\mathcal{Y}^t$, embedded in a finite-dimensional Euclidean space;
 - (c) A vector valued loss function $\ell^t(\cdot, \cdot) : \mathcal{A}^t \times \mathcal{Y}^t \rightarrow [-C, C]^d$. Every dimension $\ell_j^t(\cdot, \cdot) : \mathcal{A}^t \times \mathcal{Y}^t \rightarrow [-C, C]$ (where $j \in [d]$) of the loss function is continuous and concave in the second argument.
2. The Learner selects some $x^t \in \mathcal{X}^t$;
3. The Adversary observes the Learner's selection x^t , and chooses some action $y^t \in \mathcal{Y}^t$ in response;
4. The Learner's action $x^t \in \Delta\mathcal{A}^t$ is interpreted as a mixture over the pure actions in $\mathcal{A}^t$, and an *outcome* $a^t \in \mathcal{A}^t$ is sampled from it; that is, $a^t \sim x^t$.
5. The Learner suffers (and observes) $\ell^t(a^t, y^t)$, the loss vector with respect to the outcome a^t .

Thus, the probabilistic setting is simply a specialization of our framework to the case of the Learner's action set being a simplex at each round.

Unlike in the above deterministic setting, where the transcript through any round t was defined as $\{(x^s, y^s)\}_{s=1}^t$, in the present case we define the transcript through round t as

$$\pi^t := \{(a^1, y^1), \dots, (a^t, y^t)\},$$

that is, the transcript now records the Learner's *realized outcomes* rather than her chosen mixtures at all rounds. Furthermore, we will denote by Π^t the set of transcripts through round t , for $t \in [T]$.

Now, let us fix any Adversary Adv (that is, all of the Adversary's decisions through round T). With respect to this fixed Adversary, any *algorithm* for the Learner (defined as the collection of the Learner's decision mappings $\{\pi^{t-1} \rightarrow \Delta\mathcal{A}^t\}_{t \in [T]}$ for all rounds) induces an ex-ante distribution $\mathcal{P}_{\text{Adv}}$ over the set of transcripts Π^T .

Now, we give two types of probabilistic guarantees on the performance of Algorithm 2, namely, an in-expectation bound and a high-probability bound. Both bounds hold for any choice of Adversary Adv, and are *ex-ante with respect to the algorithm-induced distribution $\mathcal{P}_{\text{Adv}}$ over the final transcripts*.

Theorem A.1 (In-Expectation Bound). *Given $T \geq \ln d$, Algorithm 2 with learning rate $\eta = \sqrt{\frac{\ln d}{4TC^2}}$ guarantees that ex-ante, with respect to the randomness in the Learner's realized outcomes, the expected AMF regret is bounded as:*

$$\mathbb{E}[R^T] \leq 4C\sqrt{T \ln d}.$$

As mentioned in Section A.2.1, the proof of Theorem A.1 is much the same as the proofs of Theorem 2.1 and the helper Lemmas 2.1, 2.2, 2.3, with the exception of using Jensen's inequality to switch the order of taking expectations when necessary. We omit further details.

Theorem A.2 (High-Probability Bound). *Fix any $\delta \in (0, 1)$. Given $T \geq \ln d$, Algorithm 2 with learning rate $\eta = \sqrt{\frac{\ln d}{4TC^2}}$ guarantees that the AMF regret will satisfy, with ex-ante probability $1 - \delta$ over the randomness in the Learner's realized outcomes,*

$$R^T \leq 8C\sqrt{T \ln \left(\frac{d}{\delta}\right)}.$$

Proof. Throughout this proof, we put tildes over random variables to distinguish them from their realized values. For instance, $\tilde{\pi}^t$ is the random transcript through round t , while π^t is a realization of $\tilde{\pi}^t$. Also, we explicitly specify the dependence of the surrogate loss L^t on the (random or realized) transcript.

Consider the following random process $\{\tilde{Z}^t\}$, defined recursively for $t = 0, 1, \dots, T$ and adapted to the sequence of random variables $\tilde{\pi}^1, \dots, \tilde{\pi}^T$. We let $\tilde{Z}^0 := 0$ deterministically, and for $t \in [T]$ we let

$$\tilde{Z}^t := \tilde{Z}^{t-1} + \ln L^t(\tilde{\pi}^t) - \mathbb{E}_{\tilde{\pi}^t} [\ln L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}].$$

It is easy to see that for all $t \in [T]$, we have $\mathbb{E}_{\tilde{\pi}^t} [\tilde{Z}^t | \tilde{\pi}^{t-1}] = \tilde{Z}^{t-1}$, and thus $\{\tilde{Z}^t\}$ is a martingale.

We next show that this martingale has bounded increments. In brief, this follows from $\{\tilde{Z}^t\}$ being defined in terms of the *logarithm* of the surrogate loss.

Lemma A.1. *The martingale $\{\tilde{Z}^t\}$ has bounded increments: $|\tilde{Z}^t - \tilde{Z}^{t-1}| \leq 4\eta C$ for all $t \in [T]$.*

Proof. It suffices to establish the bounded increments property for an arbitrary realization of the process. Towards this, fix the full transcript π^T of the interaction, and consider any round $t \in [T]$.

Recall from the definition of the surrogate loss that

$$L^t(\pi^t) = \sum_{j \in [d]} \exp(\eta R_j^{t-1}(\pi^{t-1})) \cdot \exp(\eta(\ell_j^t(a^t, y^t) - w_A^t)).$$

Thus, noting that $|\ell_j^t(a^t, y^t) - w_A^t| \leq 2C$ for all $j \in [d]$, we have

$$\frac{L^t(\pi^t)}{L^{t-1}(\pi^{t-1})} = \frac{L^t(\pi^t)}{\sum_{j \in [d]} \exp(\eta R_j^{t-1}(\pi^{t-1}))} \in [\exp(-\eta \cdot 2C), \exp(\eta \cdot 2C)].$$

Taking the logarithm yields

$$|\ln L^t(\pi^t) - \ln L^{t-1}(\pi^{t-1})| \leq 2\eta C.$$

In fact, this argument shows that $|\ln L^t(\pi^t) - \ln L^{t-1}(\pi^{t-1})| \leq 2\eta C$ for *any* transcript π^t that equals π^{t-1} on the first $t-1$ rounds. Hence, taking the expectation over $\tilde{\pi}^t$ conditioned on π^{t-1} , we obtain:

$$|\mathbb{E} [\ln L^t(\tilde{\pi}^t) | \pi^{t-1}] - \ln L^{t-1}(\pi^{t-1})| \leq 2\eta C.$$

To conclude the proof, it now suffices to observe that:

$$\begin{aligned} |\tilde{Z}^t - \tilde{Z}^{t-1}| &= |\ln L^t(\tilde{\pi}^t) - \mathbb{E} [\ln L^t(\tilde{\pi}^t) | \pi^{t-1}]| \\ &\leq |\ln L^t(\tilde{\pi}^t) - \ln L^{t-1}(\tilde{\pi}^{t-1})| + |\ln L^{t-1}(\tilde{\pi}^{t-1}) - \mathbb{E} [\ln L^t(\tilde{\pi}^t) | \pi^{t-1}]| \\ &\leq 2\eta C + 2\eta C = 4\eta C. \end{aligned}$$

□

Having established that $\{\tilde{Z}^t\}$ is a martingale with bounded increments, we can now apply the following concentration bound (see e.g. Dubhashi and Panconesi [2009]).

Fact 3 (Azuma's Inequality). *Fix $\epsilon > 0$. For any martingale $\{\tilde{Z}^t\}_{t=0}^T$ with $|\tilde{Z}^t - \tilde{Z}^{t-1}| \leq \xi$ for $t \in [T]$,*

$$\Pr \left[\tilde{Z}^T - \tilde{Z}^0 \geq \epsilon \right] \leq \exp \left(-\frac{\epsilon^2}{2\xi^2 T} \right).$$

We instantiate this bound for our martingale with $\tilde{Z}^0 = 0$, $\xi = 4\eta C$, and $\epsilon = \xi \sqrt{2T \ln \frac{1}{\delta}} = 4\eta C \sqrt{2T \ln \frac{1}{\delta}}$, and obtain that for any $\delta \in (0, 1)$,

$$\tilde{Z}_T \leq 4\eta C \sqrt{2T \ln \frac{1}{\delta}} \quad \text{with prob. } 1 - \delta. \tag{1}$$

At this point, let us express $\tilde{Z}^T$ as follows:

$$\tilde{Z}^T = \sum_{t=1}^T \left(\ln L^t(\tilde{\pi}^t) - \mathbb{E}_{\tilde{\pi}^t} [\ln L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] \right) = \ln L^T(\tilde{\pi}^T) - \ln L^0 - \sum_{t=1}^T \left(\mathbb{E}_{\tilde{\pi}^t} [\ln L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] - \ln L^{t-1}(\tilde{\pi}^{t-1}) \right).$$

Now, with an eye toward bounding the latter sum, observe that for $t \in [T]$,

$$\begin{aligned} \mathbb{E}_{\tilde{\pi}^t} [\ln L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] - \ln L^{t-1}(\tilde{\pi}^{t-1}) &\leq \ln \mathbb{E}_{\tilde{\pi}^t} [L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] - \ln L^{t-1}(\tilde{\pi}^{t-1}) \\ &\leq \ln \left((4\eta^2 C^2 + 1) L^{t-1}(\tilde{\pi}^{t-1}) \right) - \ln L^{t-1}(\tilde{\pi}^{t-1}) \\ &= \ln(4\eta^2 C^2 + 1) \\ &\leq 4\eta^2 C^2. \end{aligned}$$

Here, the first step is via Jensen's inequality and the last step is via $\ln(1+x) \leq x$ for $x > -1$. The second step holds since we can show (via reasoning similar to Lemma 2.3) that for any $T \geq \ln d$, at each round $t \in [T]$ Algorithm 2 with learning rate $\eta = \sqrt{\frac{\ln d}{4TC^2}}$ achieves:

$$\mathbb{E}_{\tilde{\pi}^t} [L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] \leq (4\eta^2 C^2 + 1) L^{t-1}(\tilde{\pi}^{t-1}).$$

Combining the above observations with Bound 1 and recalling $L^0 = d$ yields, with probability $\geq 1 - \delta$,

$$\begin{aligned} \tilde{Z}_T &\leq 4\eta C \sqrt{2T \ln \frac{1}{\delta}} \iff \ln L^T(\tilde{\pi}^T) - \ln d - \sum_{t=1}^T \left(\mathbb{E}_{\tilde{\pi}^t} [\ln L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] - \ln L^{t-1}(\tilde{\pi}^{t-1}) \right) \leq 4\eta C \sqrt{2T \ln \frac{1}{\delta}} \\ &\iff \ln L^T(\tilde{\pi}^T) \leq \ln d + \sum_{t=1}^T \left(\mathbb{E}_{\tilde{\pi}^t} [\ln L^t(\tilde{\pi}^t) | \tilde{\pi}^{t-1}] - \ln L^{t-1}(\tilde{\pi}^{t-1}) \right) + 4\eta C \sqrt{2T \ln \frac{1}{\delta}} \\ &\implies \ln L^T(\tilde{\pi}^T) \leq \ln d + 4\eta^2 C^2 T + 4\eta C \sqrt{2T \ln \frac{1}{\delta}}. \end{aligned}$$

Using the last inequality, with $\eta = \sqrt{\frac{\ln d}{4TC^2}}$, and the fact that $R^T(\tilde{\pi}^T) \leq \frac{L^T(\tilde{\pi}^T)}{\eta}$ (which is easy to deduce via Lemma 2.1), we thus obtain the desired high-probability AMF regret bound. Specifically, with probability $1 - \delta$ we have:

$$\begin{aligned} R^T(\tilde{\pi}^T) &\leq \frac{L^T(\tilde{\pi}^T)}{\eta} \leq \frac{\ln d}{\eta} + 4\eta C^2 T + 4C \sqrt{2T \ln \frac{1}{\delta}} = 2\sqrt{4C^2 T \ln d} + 4C \sqrt{2T \ln \frac{1}{\delta}} \\ &= 4C\sqrt{T} \left(\sqrt{\ln d} + \sqrt{2 \ln \frac{1}{\delta}} \right) \leq 4C\sqrt{T} \cdot \sqrt{2} \cdot \sqrt{\ln d + 2 \ln \frac{1}{\delta}} \leq 8C \sqrt{T \ln \frac{d}{\delta}}. \end{aligned}$$

In the last line, we used that $\sqrt{x} + \sqrt{y} \leq \sqrt{2}\sqrt{x+y}$ for $x, y \geq 0$. □

B Multicalibration: The Algorithm and Full Proofs

A simple and efficient algorithm for the Learner As mentioned in the proof sketch of Theorem 4.1, in the setting of multicalibration, our framework's general Algorithm 2 has a particularly simple *approximate* version (originally derived in Gupta et al. [2022]) that lets the Learner (almost) match the above bounds on the multicalibration constant α . This approximate algorithm is very efficient and has “low” randomization:

namely, at each round the Learner plays an explicitly given distribution which randomizes over at most two points in $\mathcal{A}_r$.

Algorithm 3: Simple Multicalibrated Learner

for $t = 1, \dots, T$ **do**

 Observe θ^t .

 For each $i \in [n]$, compute:

$$C_{t-1}^i := \sum_{g \in \mathcal{G}: \theta^t \in g} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{i,g,+1}^s(a^s, b^s) \right) - \exp \left(-\eta \sum_{s=1}^{t-1} \ell_{i,g,+1}^s(a^s, b^s) \right).$$

if $C_{t-1}^i > 0$ for all $i \in [n]$ **then**

 Predict $a^t = 1$.

else if $C_{t-1}^i < 0$ for all $i \in [n]$ **then**

 Predict $a^t = 0$.

else

 Find $j \in [n-1]$ such that $C_{t-1}^j \cdot C_{t-1}^{j+1} \leq 0$.

 Define $q^t \in [0, 1]$ as follows (using the convention that $0/0 = 1$):

$$q^t := \left| C_{t-1}^{j+1} \right| / \left(\left| C_{t-1}^{j+1} \right| + \left| C_{t-1}^j \right| \right).$$

 Sample $a^t = \frac{j}{n} - \frac{1}{rn}$ with probability q^t and $a^t = \frac{j}{n}$ with probability $1 - q^t$.

Theorem B.1. *Algorithm 3 achieves the multicalibration guarantees of Theorem 4.1.*

Proof. Let us instantiate the generic probabilistic Algorithm 2 with our current set of loss functions. In parallel with the notation of Algorithm 2, for any bucket i , group g and $\sigma \in \{-1, +1\}$, we define

$$\chi_{i,g,\sigma}^t := \frac{1}{Z^t} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{i,g,\sigma}^s(a^s, b^s) \right),$$

where

$$Z^t := \sum_{i' \in [n], g' \in \mathcal{G}, \sigma' = \pm 1} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{i',g',\sigma'}^s(a^s, b^s) \right).$$

In this notation, at each round $t \in [T]$, the Learner has to solve the following zero-sum game:

$$x^t \in \operatorname{argmin}_{x \in \Delta_{\mathcal{A}_r}} \max_{b \in [0,1]} \mathbb{E}_{a \sim x} [\xi^t(a, b)],$$

where we define

$$\xi^t(a, b) := \sum_{i \in [n], g \in \mathcal{G}, \sigma \in \{-1, 1\}} \chi_{i,g,\sigma}^t \cdot \ell_{i,g,\sigma}^t(a, b) \quad \text{for } a \in \mathcal{A}_r, b \in [0, 1].$$

For any a , let i_a denote the unique bucket index $i \in [n]$ such that $a \in B_n^i$. Substituting

$$\ell_{i,g,\sigma}^t(a, b) = \sigma \cdot 1_{\theta^t \in g} \cdot 1_{a \in B_n^i} \cdot (b - a),$$

we see that most terms in the summation disappear, and what remains is precisely

$$\xi^t(a, b) = \sum_{g \in \mathcal{G}: \theta^t \in g} \sum_{\sigma \in \{-1, 1\}} \chi_{i_a, g, \sigma}^t \cdot \sigma (b - a) = (b - a) \cdot \frac{C_{t-1}^{i_a}}{Z^t},$$

where $C_{t-1}^{i_a} = Z^t \sum_{g \in \mathcal{G}: \theta^t \in g} \chi_{i_a, g, +1}^t - \chi_{i_a, g, -1}^t$ is as defined in the pseudocode for Algorithm 3.

Crucially, for any distribution x chosen by the Learner, her attained utility after the Adversary best-responds has a simple closed form. Namely, given any x played by the Learner, we have

$$\begin{aligned} \max_{b \in [0,1]} \mathbb{E}_{a \sim x} [\xi^t(a, b)] &= \frac{1}{Z^t} \left(\max_{b \in [0,1]} \left(b \cdot \mathbb{E}_{a \sim x} [C_{t-1}^{i_a}] \right) - \mathbb{E}_{a \sim x} [a \cdot C_{t-1}^{i_a}] \right), \\ &= \frac{1}{Z^t} \left(\max_{a \sim x} \left(\mathbb{E}_{a \sim x} [C_{t-1}^{i_a}], 0 \right) - \mathbb{E}_{a \sim x} [a \cdot C_{t-1}^{i_a}] \right). \end{aligned}$$

With this in mind, the Learner can easily achieve value 0 in the following two cases. When $C_{t-1}^{i_a} > 0$ for all $i \in [n]$, playing $a = 1$ deterministically gives: $\max_{a \sim x} \left(\mathbb{E}_{a \sim x} [C_{t-1}^{i_a}], 0 \right) - \mathbb{E}_{a \sim x} [a \cdot C_{t-1}^{i_a}] = \mathbb{E}_{a \sim x} [C_{t-1}^{i_a}] - \mathbb{E}_{a \sim x} [C_{t-1}^{i_a}] = 0$. When $C_{t-1}^{i_a} < 0$ for all $i \in [n]$, she can play $a = 0$ deterministically, ensuring that $\max_{a \sim x} \left(\mathbb{E}_{a \sim x} [C_{t-1}^{i_a}], 0 \right) - \mathbb{E}_{a \sim x} [a \cdot C_{t-1}^{i_a}] = 0 - 0 = 0$.

In the final case, when there are nonpositive and nonnegative quantities among $\{C_{t-1}^{i_a}\}_{i \in [n]}$, note that there exists an intermediate index $j \in [n-1]$ such that $C_{t-1}^j \cdot C_{t-1}^{j+1} \leq 0$. Then, it is easy to check that q^t , as defined in Algorithm 3, satisfies

$$q^t C_{t-1}^j + (1 - q^t) C_{t-1}^{j+1} = 0.$$

Using this relation, we obtain that when the Learner plays $a^t = \frac{j}{n} - \frac{1}{rn}$ with probability q^t and $a^t = \frac{j}{n}$ with probability $1 - q^t$, she accomplishes value

$$\begin{aligned} \max_{b \in [0,1]} \mathbb{E}_{a^t} [\xi^t(a^t, b)] &= \frac{1}{Z^t} \left(\max_{b \in [0,1]} \left(\mathbb{E}_{a^t} [C_{t-1}^{i_{a^t}}], 0 \right) - \mathbb{E}_{a^t} [a^t \cdot C_{t-1}^{i_{a^t}}] \right) \\ &= \frac{1}{Z^t} \left(\max_{b \in [0,1]} \left(q^t \cdot C_{t-1}^j + (1 - q^t) C_{t-1}^{j+1}, 0 \right) - \left(q^t \left(\frac{j}{n} - \frac{1}{rn} \right) C_{t-1}^j + (1 - q^t) \frac{j}{n} C_{t-1}^{j+1} \right) \right) \\ &= \frac{1}{Z^t} \cdot \frac{1}{rn} C_{t-1}^j, \end{aligned}$$

and thus, recalling that $C_{t-1}^{i_a} = Z^t \sum_{g \in \mathcal{G}: \theta^t \in g} \chi_{i_a, g, +1}^t - \chi_{i_a, g, -1}^t$, we obtain

$$\max_{b \in [0,1]} \mathbb{E}_{a^t} [\xi^t(a^t, b)] = \frac{1}{rn} \sum_{g \in \mathcal{G}: \theta^t \in g} \chi_{j, g, +1}^t - \chi_{j, g, -1}^t \leq \frac{1}{rn} \sum_{i \in [n], g \in \mathcal{G}, \sigma = \pm 1} \chi_{i, g, \sigma}^t = \frac{1}{rn},$$

where the last line is due to the quantities $\chi_{i, g, \sigma}$ forming a probability distribution.

Therefore, in the language of Section A.2.2, the Learner who uses Algorithm 3 guarantees herself *achieved AMF value bounds*

$$w_{\text{bd}}^t = \frac{1}{rn} \text{ for } t \in [T].$$

Hence, by Theorem A.3, our (suboptimal) Learner achieves the claimed multicalibration bounds. $\square$

C Multicalibeating: Full Statements and Proofs

C.1 Calibeating a Single Forecaster: Proof of Theorem 4.2

Proof of Theorem 4.2. For the exposition of this full proof, we will employ some probabilistic notation that we have not seen in the main Section 4.2. We briefly define it here.

For any subsequence $S \subseteq [T]$ of rounds, $t \sim S$ denotes a uniformly random round in S . We denote the empirical distributions of the values of f , a , (f, a) on $S \subseteq [T]$ by $\mathcal{D}^f(S)$, $\mathcal{D}^a(S)$, $\mathcal{D}^{f \times a}(S)$ (or simply $\mathcal{D}^f$, $\mathcal{D}^a$, $\mathcal{D}^{f \times a}$ when $S = [T]$). In this notation, we e.g. have $\mathcal{R}^f(\pi^T) = \mathbb{E}_{d \sim \mathcal{D}^f} [\text{Var}_{t \sim S^d} [b^t]]$.

Our quantity of interest, the Brier score $\mathcal{B}^a$ of the Learner's predictions a , is inconvenient to handle: indeed, the calibration-refinement decomposition of $\mathcal{B}^a$ is of little utility since the Learner's predictions can take arbitrary real values (in particular, they might all be distinct, in which case the refinement score would be 0, and all of the Brier score would be contained in the calibration error). Instead, we define a convenient surrogate notion of bucketed Brier/calibration/refinement score.

$$\begin{aligned}\mathcal{K}_n^a(\pi^T) &:= \frac{1}{T} \sum_{i \in [n]} |S_i| (\bar{a}(S_i) - \bar{b}(S_i))^2. \\ \mathcal{R}_n^a(\pi^T) &:= \frac{1}{T} \sum_{i \in [n]} \sum_{t \in S_i} (b^t - \bar{b}(S_i))^2 = \frac{1}{T} \sum_{i \in [n]} |S_i| \text{Var}_{t \in S_i}[b^t] = \mathbb{E}_{i \sim \mathcal{D}^i} [\text{Var}_{t \sim S_i}[b^t]]. \\ \mathcal{B}_n^a(\pi^T) &:= \mathcal{K}_n^a(\pi^T) + \mathcal{R}_n^a(\pi^T).\end{aligned}$$

The following lemma shows that as long as n is large enough, the surrogate Brier score is a good estimate of the true Brier score of our predictions (i.e. our squared error).

Lemma C.1. $\mathcal{B}^a \leq \mathcal{B}_n^a + \frac{1}{n}$.

Proof. We first compute that the original Brier score $\mathcal{B}^a$ equals

$$\mathcal{B}^a := \frac{1}{T} \sum_{t=1}^T (a^t - b^t)^2 = \frac{1}{T} \sum_{i=1}^n \sum_{t \in S_i} (a^t - b^t)^2 = \frac{1}{T} \sum_{i=1}^n |S_i| \sum_{t \in S_i} \frac{1}{|S_i|} (a^t - b^t)^2.$$

The inner sum is the expectation, over the transcript, of $(a^t - b^t)^2$ conditioned on $a^t \in B_n^i$, so we can write:

$$\mathcal{B}^a = \frac{1}{T} \sum_{i=1}^n |S_i| \mathbb{E}_{t \sim S_i} [(a^t - b^t)^2].$$

We can decompose the expected value as:

$$\mathbb{E}_{t \sim S_i} [(a^t - b^t)^2] = (\mathbb{E}_{t \sim S_i} [a^t - b^t])^2 + \text{Var}_{t \sim S_i} [a^t - b^t].$$

By linearity of expectation, the expectation-squared term satisfies:

$$(\mathbb{E}_{t \sim S_i} [a^t - b^t])^2 = (\bar{a}(S_i) - \bar{b}(S_i))^2.$$

Meanwhile, the variance term can be upper bounded using the following fact:

Fact 4. For any random variables X, Y :

$$\text{Var}[X + Y] = \text{Var}[X] + \text{Var}[Y] + 2\text{Cov}(X, Y) \leq \text{Var}[X] + \text{Var}[Y] + 2\sqrt{\text{Var}[X] \text{Var}[Y]}.$$

where the inequality follows from an application of Cauchy-Schwartz.

Instantiating $X = a^t$ and $Y = -b^t$, and upper bounding $\sqrt{\text{Var}[X]} \leq \frac{1}{2n}$, $\sqrt{\text{Var}[Y]} \leq \frac{1}{2}$, we get:

$$\begin{aligned}\text{Var}_{t \sim S_i} [a^t - b^t] &\leq \text{Var}_{t \sim S_i} [a^t] + \text{Var}_{t \sim S_i} [b^t] + 2\sqrt{\text{Var}_{t \sim S_i} [a^t] \text{Var}_{t \sim S_i} [b^t]}, \\ &\leq \frac{1}{(2n)^2} + \text{Var}_{t \sim S_i} [b^t] + \frac{1}{2n}, \\ &\leq \text{Var}_{t \sim S_i} [b^t] + \frac{1}{n}.\end{aligned}$$

Putting the above back together gives the desired bound on the difference of $\mathcal{B}^a$ and $\mathcal{B}_n^a$:

$$\begin{aligned}
\mathcal{B}^a &= \frac{1}{T} \sum_{i=1}^n |S_i| \mathbb{E}_{t \sim S_i} [(a^t - b^t)^2], \\
&\leq \frac{1}{T} \sum_{i=1}^n |S_i| \left((\bar{a}(S_i) - \bar{b}(S_i))^2 + \text{Var}_{t \sim S_i}[b^t] + \frac{1}{n} \right), \\
&= \frac{1}{T} \sum_{i=1}^n |S_i| (\bar{a}(S_i) - \bar{b}(S_i))^2 + \frac{1}{T} \sum_{i=1}^n |S_i| \text{Var}_{t \sim S_i}[b^t] + \frac{1}{n}, \\
&= \mathcal{K}_n^a + \mathcal{R}_n^a + \frac{1}{n}.
\end{aligned}$$

□

Having shown that the surrogate Brier score $\mathcal{B}_n^a$ closely approximates the Learner's original score $\mathcal{B}^a$, we can now focus on bounding the calibration and refinement scores associated with $\mathcal{B}_n^a$.

Calibration: Our multicalibration condition on Θ implies that $\frac{|S_i|}{T} |\bar{b}(S_i) - \bar{a}(S_i)| \leq \alpha$ for $i \in [n]$. The calibration score bound then follows directly.

$$\mathcal{K}_n^a = \frac{1}{T} \sum_{i \in [n]} |S_i| (\bar{b}(S_i) - \bar{a}(S_i))^2 \leq \frac{1}{T} \sum_{i \in [n]} |S_i| |\bar{b}(S_i) - \bar{a}(S_i)| \leq \sum_{i \in [n]} \alpha = \alpha n.$$

Refinement: We claim that the Learner's surrogate refinement score relates to the refinement score of the forecaster f as follows:

$$\mathcal{R}_n^a \leq \mathcal{R}^f + \alpha n (|D_f| + 1) + \frac{1}{n}.$$

The proof proceeds in two steps, connecting $\mathcal{R}^f$ and $\mathcal{R}^a$ via a quantity we call $\mathcal{R}^{f \times a}$.

Definition C.1 (Joint Refinement Score).

$$\mathcal{R}^{f \times a} := \mathbb{E}_{d, i \sim \mathcal{D}^{f \times a}} \mathbb{E}_{t \sim S_i^d} [\text{Var}[b^t]] = \frac{1}{T} \sum_{d \in D_f, i \in [n]} |S_i^d| \text{Var}_{t \sim S_i^d}[b^t].$$

Recall that refinement score, although we defined it for a forecaster, is really a property of a *partition* of the days. It's equally well defined if, instead of partitioning by days on which a forecaster makes a certain forecast, we partition on say, even and odd days, or sunny vs cloudy vs rainy vs snowy days. Or, in the case of Definition C.1, the partition $\{S_i^d\}_{i \in [n], d \in D}$.

First, note that the joint refinement score of a and f is no worse than the refinement score of f .

Observation 1. $\mathcal{R}^f \geq \mathcal{R}^{f \times a}$.

Intuitively this should make sense, since $\{S_i^d\}$ is a *refinement* of f 's level sets by a 's level sets. If a is "useful", then this inequality would be strict, as combining with a would explain away more of the variance. Refining by a cannot decrease the amount of variance captured by the partition.

Reversing our perspective, we can think of $\{S_i^d\}$ as a refinement of a 's level sets by f 's level sets. The key idea is to use multicalibration to show that refining by f is not "useful." Multicalibration ensures us that almost all of f 's explanatory power is captured by a .

Observation 2. $\mathcal{R}_n^a = \mathcal{R}^{f \times a} + \mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)]]$.

Observation 3. *The extra error term is small:* $\mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)]] \leq \alpha n (|D| + 1) + \frac{1}{n}$.

Combining these three observations will give us our desired refinement score bound:

$$\mathcal{R}_n^a(b) \leq \mathcal{R}^f + \alpha n(|D| + 1) + \frac{1}{n}.$$

We therefore now prove these observations one by one.

Proof of Observation 1. We recall the following fact from probability:

Fact 5 (Law of Total Variance). *For any random variables $W, Z : \Omega \rightarrow \mathbb{R}$ in a probability space,*

$$\text{Var}[Z] = \mathbb{E}[\text{Var}[Z|W]] + \text{Var}[\mathbb{E}[Z|W]].$$

In particular, since variance is always non-negative:

$$\text{Var}[Z] \geq \mathbb{E}[\text{Var}[Z|W]].$$

For each fixed d , we instantiate this fact with $\Omega = S^d$ (equipped with the discrete σ -algebra and uniform distribution). $Z(t) := b^t$ and $W(t) := i_{a^t}$, the unique i s.t. $a^t \in B_n^i$. This gives us:

$$\text{Var}[b^t] \geq \mathbb{E}_{t \sim S^d} [\text{Var}[b^t | a^t \in B_n^i]] = \mathbb{E}_{i \sim \mathcal{D}^a(S^d)} [\text{Var}[b^t]].$$

Since this is true for all d , the inequality continues to hold in expectation over the d 's:

$$\mathcal{R}^f = \mathbb{E}_{d \sim \mathcal{D}^f} [\text{Var}[b^t]] \geq \mathbb{E}_{d, i \sim \mathcal{D}^{f \times a}} [\text{Var}[b^t]] = \mathcal{R}^{f \times a}.$$

□

Proof of Observation 2. Recall the definition of bucketed refinement:

$$\mathcal{R}_n^a = \mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}[b^t]].$$

To relate this back to $\mathcal{R}^{f \times a}$, we instantiate Fact 5 again, but flipping the roles of f and a : we take the underlying spaces to be the sequences S_i defined by calibrated buckets, and let W , the variable we condition on, be the level sets of f .

For any fixed i representing a level set of a , Fact 5 tells us:

$$\text{Var}_{t \sim S_i}[b^t] = \mathbb{E}_{d \sim \mathcal{D}^f(S_i)} [\text{Var}_{t \sim S_i}[b^t | f^t = d]] + \text{Var}_{d \sim \mathcal{D}^f(S_i)} [\mathbb{E}_{t \sim S_i}[b^t | f^t = d]] = \mathbb{E}_{d \sim \mathcal{D}^f(S_i)} [\text{Var}_{t \sim S_i^d}[b^t]] + \text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)].$$

Like before, we take the expectation over all $i \in [n]$, giving us the desired result:

$$\mathcal{R}_n^a = \mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{t \sim S_i}[b^t]] = \mathbb{E}_{d, i \sim \mathcal{D}^{f \times a}} [\text{Var}_{t \sim S_i^d}[b^t]] + \mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)]] = \mathcal{R}^{f \times a} + \mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)]].$$

□

Proof of Observation 3. We have to bound the extra error term:

$$\mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)]].$$

In words, this is the expected variance of the true averages on S_i^d , conditioned on the buckets i . Intuitively, if these true averages vary a lot, then the calibration error on the S_i^d s must be large since the prediction on each

of the S_i^d s is (close to) i/n ; in particular, they are almost constant across d . Conversely, if multicalibration error is low, then the variance must be low as well. Formally,

$$\begin{aligned}
\mathbb{E}_{i \sim \mathcal{D}^a} [\text{Var}_{d \sim \mathcal{D}^f(S_i)} [\bar{b}(S_i^d)]] &= \sum_{i \in [n]} \frac{|S_i|}{T} (\text{Var}_d [\mathbb{E}_{t \sim S_i^d} [b^t]]), \\
&= \sum_{i \in [n]} \frac{|S_i|}{T} \left(\sum_{d \in D} \frac{|S_i^d|}{|S_i|} (\bar{b}(S_i^d) - \bar{b}(S_i))^2 \right), \\
&= \sum_{i,d} \frac{|S_i^d|}{T} (\bar{b}(S_i^d) - \bar{b}(S_i))^2, \\
&\leq \sum_{i,d} \frac{|S_i^d|}{T} |\bar{b}(S_i^d) - \bar{b}(S_i)|, \\
&\leq \sum_{i,d} \frac{|S_i^d|}{T} (|\bar{b}(S_i^d) - \bar{a}(S_i^d)| + |\bar{a}(S_i^d) - \bar{a}(S_i)| + |\bar{a}(S_i) - \bar{b}(S_i)|), \\
&\leq \sum_{i,d} \frac{|S_i^d|}{T} (T\alpha/|S_i^d| + \frac{1}{n} + T\alpha/|S_i|), \\
&\leq \frac{1}{n} + \sum_i \alpha + \sum_{i,d} \alpha, \\
&= \frac{1}{n} + \alpha n(|D| + 1).
\end{aligned}$$

The first line is just expanding out the definition. In the third line, we upperbound square with absolute value, since all values are at most 1. In the forth line, we break apart the error term into the difference between our average prediction on S_i^d and the true average (upperbounded by $T\alpha/|S_i^d|$, by calibration guarantees w.r.t $\mathcal{S}(f)$), the difference between our prediction on S_i^d and our average prediction on S_i (which is upper bounded by $1/n$, the size of our bucketing), and the difference between our average prediction on S_i and the true average (upperbounded by $T\alpha/|S_i|$). $\square$

We have shown that $\mathcal{K}_n^a \leq \alpha n$, and our three observations have given us that $\mathcal{R}_n^a(b) \leq \mathcal{R}^f + \alpha n(|D| + 1) + \frac{1}{n}$. Combining these results and Lemma C.1, we obtain the desired bound: $\mathcal{B}^a \leq \mathcal{R}^f + \alpha n(|D| + 2) + \frac{2}{n}$. This concludes the proof of Theorem 4.2. $\square$

C.2 Applying Theorem 4.2: Explicit Rates and Multiple Forecasters

First, we show how to instantiate Theorem 4.2 with our efficiently achievable multicalibration guarantees on α of Theorem 4.1.

Corollary C.1. *When run with parameters $r, n \geq 1$ on the collection $\mathcal{G}' := \mathcal{S}(f) \cup \{\Theta\}$, the multicalibration algorithm (Algorithm 3) τ -calibrates f , where*

$$\mathbb{E}[\tau] \leq \frac{2}{n} + n(|D_f| + 2) \left(\frac{1}{rn} + 4\sqrt{\frac{\ln(2(|D_f| + 1)n)}{T}} \right),$$

and for any $\delta \in (0, 1)$, with probability $1 - \delta$,

$$\tau \leq \frac{2}{n} + n(|D_f| + 2) \left(\frac{1}{rn} + 8\sqrt{\frac{1}{T} \ln \left(\frac{2(|D_f| + 1)n}{\delta} \right)} \right).$$

The calibration error overall of the algorithm is bounded, for any $\delta \in (0, 1)$, as:

$$\mathbb{E}[\mathcal{K}_n^a] \leq \frac{1}{r} + 4n\sqrt{\frac{\ln(2(|D_f| + 1)n)}{T}} \quad \text{and} \quad \mathcal{K}_n^a \leq \frac{1}{r} + 8n\sqrt{\frac{1}{T} \ln\left(\frac{2(|D_f| + 1)n}{\delta}\right)} \quad \text{w. prob. } 1 - \delta.$$

Proof. Using our online multicalibration guarantees, we get (by Theorem B.1):

$$\mathbb{E}[\alpha] \leq \frac{1}{rn} + 4\sqrt{\frac{\ln(2(|D_f| + 1)n)}{T}},$$

and, for any $\delta \in (0, 1)$, with probability $1 - \delta$:

$$\alpha \leq \frac{1}{rn} + 8\sqrt{\frac{1}{T} \ln\left(\frac{2(|D_f| + 1)n}{\delta}\right)},$$

Plugging this into the result from Theorem 4.2:

$$\mathcal{B}^a - \mathcal{R}^f \leq \alpha n(|D| + 2) + \frac{2}{n},$$

we obtain the desired in-expectation bound on τ :

$$\mathbb{E}[\tau] \leq \frac{2}{n} + n(|D_f| + 2) \mathbb{E}[\alpha] \leq \frac{2}{n} + n(|D_f| + 2) \left(\frac{1}{rn} + 4\sqrt{\frac{\ln(2(|D_f| + 1)n)}{T}} \right).$$

We can do so similarly for the high probability bound, so that with probability $1 - \delta$:

$$\tau \leq \frac{2}{n} + n(|D_f| + 2) \left(\frac{1}{rn} + 8\sqrt{\frac{1}{T} \ln\left(\frac{2(|D_f| + 1)n}{\delta}\right)} \right).$$

Finally, the overall calibration error follows directly by plugging in for α . □

The main utility in our approach to calibrating is that it easily extends to multicalibrating. As a warm up, we start by deriving calibration with respect to an ensemble of forecasters. The main result then combines this with calibrating on groups to attain the multicalibrating from Definition 4.5.

Calibrating an ensemble of forecasters Since our result above is based on bounds on multicalibration, we can easily extend it to calibrating an ensemble of forecasters $\mathcal{F}$ by asking for multicalibration with respect to the level sets of all forecasters. More formally, define the groups as:

$$\left(\bigcup_{f \in \mathcal{F}} \mathcal{S}(f) \right) \cup \{\Theta\}.$$

Theorem 4.2 applies separately to each f . The only degradation comes in the α , since we're asking for multicalibration with respect to more groups. But this effect is small, since the dependence on the number of groups is only $O(\sqrt{\ln |\mathcal{G}|})$.

Corollary C.2 (Ensemble Calibrating). *On groups $\mathcal{G}' := \left(\bigcup_{f \in \mathcal{F}} \mathcal{S}(f) \right) \cup \{\Theta\}$, the multicalibration algorithm with parameters $r, n \geq 1$, after T rounds attains $(\mathcal{F}, \{\Theta\}, \beta)$ -multicalibrating with*

$$\mathbb{E}[\beta(f, \Theta)] \leq \frac{2}{n} + n(|D_f| + 2) \left(\frac{1}{rn} + 4\sqrt{\frac{\ln(2(1 + \sum_{f' \in \mathcal{F}} |D_{f'}|)n)}{T}} \right).$$

Proof. We instantiate Theorem B.1 with group collection size $|\mathcal{G}'| = 1 + \sum_{f' \in \mathcal{F}} |D_{f'}|$ to conclude that the multicalibrated algorithm achieves (α, n) -multicalibration, with

$$\mathbb{E}[\alpha] \leq \frac{1}{rn} + 4\sqrt{\frac{\ln(2(1 + \sum_{f' \in \mathcal{F}} |D_{f'}|)n)}{T}}.$$

Now, $\forall f \in \mathcal{F} : \mathcal{S}(f) \cup \{\Theta\} \subseteq \mathcal{G}'$ for every $f \in \mathcal{F}$, so we can instantiate Theorem 4.2 for every forecaster $f \in \mathcal{F}$ to give us:

$$\mathcal{B}^a - \mathcal{R}^f \leq \alpha n(|D_f| + 2) + \frac{2}{n} \quad \forall f \in \mathcal{F}.$$

Plugging in the in-expectation bound on α , we conclude:

$$\mathbb{E}[\beta(f, \Theta)] \leq \mathbb{E}[\alpha] \cdot n(|D_f| + 2) + \frac{2}{n} \leq \frac{2}{n} + n(|D_f| + 2) \left(\frac{1}{rn} + 4\sqrt{\frac{\ln(2(1 + \sum_{f' \in \mathcal{F}} |D_{f'}|)n)}{T}} \right).$$

□

C.3 Multicalbeating + Multicalibration Theorem 4.3: Full Statement and Proof

Recall that for every group $g \in \mathcal{G}$, we let $S(g)$ denote the subsequence of days on which g occurs, where the transcript is left implicit.

Theorem C.1 (Multicalbeating + Multicalibration: Full version with high-probability bounds). *Let $\mathcal{G} \subseteq 2^\Theta$, and $\mathcal{F}$ some set of forecasters $f : \Theta \rightarrow D_f$. The multicalibration algorithm on $\mathcal{G}' := \left(\bigcup_{f \in \mathcal{F}} \{g \cap S : (g, S) \in \mathcal{G} \times \mathcal{S}(f)\} \right) \cup \mathcal{G}$ with parameters $r, n \geq 1$, after T rounds, attains expected $(\mathcal{F}, \mathcal{G}, \beta)$ -multicalbeating, where:⁸*

$$\mathbb{E}[\beta(f, g)] \leq \frac{2}{n} + \frac{|D_f| + 2}{r \cdot |S(g)|/T} + 4n(|D_f| + 2) \sqrt{\frac{1}{|S(g)|^2/T} \ln \left(2n|\mathcal{G}|(1 + \sum_f |D_f|) \right)} \quad \forall f \in \mathcal{F}, g \in \mathcal{G},$$

while maintaining (α, n) -multicalibration on the original collection $\mathcal{G}$, with:

$$\mathbb{E}[\alpha] \leq \frac{1}{rn} + 4\sqrt{\frac{1}{T} \ln \left(2n|\mathcal{G}|(1 + \sum_f |D_f|) \right)}.$$

We also have the corresponding high probability bounds. For any $\delta \in (0, 1)$, with probability $1 - \delta$:

$$\beta(f, g) \leq \frac{2}{n} + \frac{|D_f| + 2}{r \cdot |S(g)|/T} + 8n(|D_f| + 2) \sqrt{\frac{1}{|S(g)|^2/T} \ln \left(\frac{2n|\mathcal{G}|(1 + \sum_f |D_f|)}{\delta} \right)} \quad \forall f \in \mathcal{F}, g \in \mathcal{G},$$

and on the original collection $\mathcal{G}$, the multicalibration constant α satisfies, with probability $1 - \delta$,

$$\alpha \leq \frac{1}{rn} + 8\sqrt{\frac{1}{T} \ln \left(\frac{2n|\mathcal{G}|(1 + \sum_f |D_f|)}{\delta} \right)}.$$

Proof. We begin with a preliminary observation that translates our overall multicalibration assumptions into guarantees over the individual sequences $S(g)$, for $g \in \mathcal{G}$.

Observation 4. *Let a be (α, n) -multicalibrated on groups $\mathcal{G}'$ over the entire time sequence $[T]$. Then, for any g , on the subsequence of days $S(g)$ the predictor a is $\left(\alpha \frac{T}{|S(g)|}, n \right)$ -multicalibrated with respect to groups $\left(\bigcup_{f \in \mathcal{F}} \mathcal{S}(f) \right) \cup \{\Theta\}$.*

⁸ $S(g)$ denotes the subsequence of days on which a group g occurs, suppressing dependence on transcript.

Proof. Let $g \in \mathcal{G}$ be some particular group. Also, fix any $f \in \mathcal{F}$ and $S \in \mathcal{S}(f) \cup \{\Theta\}$. Using multicalibration guarantees (Definition 4.1), we have that for every $i \in [n]$:

$$\left| \sum_{t \in S(g): \theta^t \in S \text{ and } a^t \in B_n^i} b^t - a^t \right| = \left| \sum_{t \in [T]: \theta^t \in g \cap S \text{ and } a^t \in B_n^i} b^t - a^t \right| \leq \alpha T = \left(\alpha \frac{T}{|S(g)|} \right) |S(g)|.$$

The first equality is by definition of $S(g)$; in particular, $\theta^t \in g \cap S \iff t \in S(g) \wedge \theta^t \in S$. This concludes the proof of our observation. $\square$

With this observation in hand, the proof is again a direct application of Theorem 4.2.

We can instantiate Theorem B.1 with groups $\mathcal{G}'$ to conclude that the multicalibrated algorithm achieves (α, n) -multicalibration, with (choosing any $\delta \in (0, 1)$):

$$\mathbb{E}[\alpha] \leq \frac{1}{rn} + 4\sqrt{\frac{\ln(2|\mathcal{G}'|n)}{T}}, \quad \text{and} \quad \alpha \leq \frac{1}{rn} + 8\sqrt{\frac{1}{T} \ln \left(\frac{2|\mathcal{G}'|n}{\delta} \right)} \text{ w. prob. } 1 - \delta.$$

where $|\mathcal{G}'| = |\mathcal{G}| + |\mathcal{G}|(\sum_f D_f) = |\mathcal{G}|(1 + \sum_f D_f)$.

Now, fix any $g \in \mathcal{G}$ and $f \in \mathcal{F}$. By our observation above, we are $\alpha \frac{T}{|S(g)|}$ multicalibrated w.r.t. $S(f) \cup \{\Theta\}$ on the sequence of days on which g occurs. Therefore, we can instantiate Theorem 4.2:

$$\mathcal{B}^a(\pi^T|_{\{t: \theta^t \in g\}}) - \mathcal{R}^f(\pi^T|_{\{t: \theta^t \in g\}}) \leq \frac{2}{n} + n(|D_f| + 2) \alpha \frac{T}{|S(g)|}.$$

Inserting the above bounds on α yields our in-expectation and high-probability bounds on $\beta(\cdot, \cdot)$.

Additionally, the theorem posits that the predictor is also (α, n) -multicalibrated on the base collection of subgroups $\mathcal{G}$. Indeed, we have included the family $\mathcal{G}$ into the collection $\mathcal{G}'$, hence the predictor will be (α, n) -multicalibrated on $\mathcal{G}$. $\square$

D Blackwell Approachability: The Algorithm and Full Proofs

of Theorem 5.1. We instantiate our probabilistic framework of Section A.2.1. The Learner's and Adversary's action sets are inherited from the underlying Polytope Blackwell game.

Defining the loss functions. For all $t = 1, 2, \dots$, we consider the following losses:

$$\ell_{h_{\alpha, \beta}}^t(x, y) := \langle \alpha, u(x, y) \rangle - \beta, \quad \text{for } h_{\alpha, \beta} \in \mathcal{H}, x \in \mathcal{X}, y \in \mathcal{Y},$$

where here and below the notational convention is that for $x \in \mathcal{X}, y \in \mathcal{Y}$, $u(x, y) := \mathbb{E}_{a \sim x}[u(a, y)]$. The coordinates of the resulting vector loss $\ell_{\mathcal{H}}^t(x, y) := \left(\ell_{h_{\alpha, \beta}}^t(x, y) \right)_{h_{\alpha, \beta} \in \mathcal{H}}$ correspond to the collection $\mathcal{H}$ of the halfspaces that define the polytope. By Holder's inequality, each vector loss function $\ell_{\mathcal{H}}^t \in [-2, 2]^d$ — this follows because we required that for some p, q with $\frac{1}{p} + \frac{1}{q} = 1$, the family $\mathcal{H}$ is p -normalized, and the range of u is contained in B_q^d . In addition, each $\ell_{h_{\alpha, \beta}}^t$ is continuous and convex-concave by virtue of being a linear function of the continuous and affine-concave function u .

Bounding the Adversary-Moves-First value. We observe that for $t \in [T]$, the AMF value $w_A^t \leq 0$. Indeed, if the Adversary moves first and selects any $y^t \in \mathcal{Y}$, then by the assumption of response satisfiability, the Learner has some $x^t \in \mathcal{X}$ guaranteeing that $u(x^t, y^t) \in P(\mathcal{H})$. The latter is equivalent to $\ell_{h_{\alpha, \beta}}^t(x^t, y^t) = \langle \alpha, u(x^t, y^t) \rangle - \beta \leq 0$ for all $h_{\alpha, \beta} \in \mathcal{H}$, letting us conclude that for any round t ,

$$w_A^t = \sup_{y^t \in \mathcal{Y}} \min_{x^t \in \mathcal{X}} \left(\max_{h_{\alpha, \beta} \in \mathcal{H}} \ell_{h_{\alpha, \beta}}^t(x^t, y^t) \right) \leq 0.$$

Applying AMF regret bounds. Given this instantiation of our framework, Theorem A.1 implies that for any response satisfiable Polytope Blackwell game, the Learner can use Algorithm 2 (instantiated with the above loss functions) to ensure that after any round $T \geq \ln |\mathcal{H}|$,

$$\mathbb{E} \left[\max_{h_{\alpha, \beta} \in \mathcal{H}} \sum_{t \in [T]} (\langle \alpha, u(a^t, y^t) \rangle - \beta) \right] \leq \mathbb{E} \left[\max_{h_{\alpha, \beta} \in \mathcal{H}} \sum_{t \in [T]} \ell_{h_{\alpha, \beta}}^t(a^t, y^t) - \sum_{t=1}^T w_A^t \right] \leq 8\sqrt{T \ln |\mathcal{H}|},$$

where the expectation is with respect to the Learner's randomness. Given this guarantee, we obtain, using the definition of $\bar{u}^T$, that

$$\max_{h_{\alpha, \beta} \in \mathcal{H}} \mathbb{E} [\langle \alpha, \bar{u}^T \rangle - \beta] \leq 8\sqrt{\frac{\ln |\mathcal{H}|}{T}}.$$

Using $T = T(\epsilon) \geq \ln |\mathcal{H}|$, we have that for every $h_{\alpha, \beta} \in \mathcal{H}$,

$$\mathbb{E} [\langle \alpha, \bar{u}^{T(\epsilon)} \rangle - \beta] \leq 8\sqrt{\frac{\ln |\mathcal{H}|}{T(\epsilon)}} = 8\sqrt{\frac{\ln |\mathcal{H}|}{64 \ln |\mathcal{H}| / \epsilon^2}} = \epsilon.$$

This concludes the proof of our in-expectation guarantee for Polytope Blackwell games.

The high-probability statement follows directly from Theorem A.2, using $C = 2$. $\square$

An LP based algorithm when the Adversary has a finite pure strategy space. Algorithm 2, which achieves the guarantees of Theorem 5.1, generally involves solving a convex program at each round. It is worth pointing out that only a *linear program* will need to be solved at each round in the commonly studied special case of Blackwell approachability where *both* the Learner and the Adversary randomize between actions in their respective finite action sets $\mathcal{A}$ and $\mathcal{B}$.

Formally, in the setting above, suppose additionally that the Adversary's action space is $\mathcal{Y} = \Delta \mathcal{B}$, where $\mathcal{B}$ is a finite set of pure actions for the Adversary. At each round t , *both* the Learner and the Adversary randomize over their respective action sets. First, the Learner selects a mixture $x^t \in \Delta \mathcal{A}$, and then the Adversary selects a mixture $y^t \in \Delta \mathcal{B}$ in response. Next, pure actions $a^t \sim x^t$ and $b^t \sim y^t$ are sampled from the chosen mixtures, and the vector valued utility in that round is set to $u(a^t, b^t)$.

In this fully probabilistic setting, at each round t Algorithm 2 has the Learner solve a normal-form zero-sum game with pure action sets $\mathcal{A}, \mathcal{B}$, where the utility to the Adversary (the max player) is

$$\xi^t(a, b) := \sum_{h_{\alpha, \beta} \in \mathcal{H}} \exp \left(\eta \sum_{s=1}^{t-1} (\langle \alpha, u(a^s, b^s) \rangle - \beta) \right) \cdot (\langle \alpha, u(a, b) \rangle - \beta) \text{ for } a \in \mathcal{A}, b \in \mathcal{B}. \quad (2)$$

A standard LP-based approach to solving this zero-sum game (see e.g. Raghavan [1994]) is for the Learner to select among distributions $x^t \in \Delta \mathcal{A}$ with the goal of minimizing the maximum payoff to the Adversary over all pure responses $b \in \mathcal{B}$. Writing this down as a linear program, we obtain the following algorithm:

Algorithm 4: Linear Programming Based Learner for Polytope Blackwell Approachability

for $t = 1, \dots, T$ **do**

 Choose a mixture $x^t = (x_a^t)_{a \in \mathcal{A}} \in \Delta \mathcal{A}$ that solves the following linear program (where $\xi^t(\cdot, \cdot)$ is defined in (2), and z is an unconstrained variable):

$$\begin{aligned} & \text{Minimize } z \\ & \text{s.t. } \forall b \in \mathcal{B}: \quad z \geq \sum_{a \in \mathcal{A}} x_a^t \xi^t(a, b). \end{aligned}$$

 Sample $a^t \sim x^t$.

E No-X-Regret: Definitions, Examples, Algorithms, and Proofs

As a warmup, we begin this subsection by carefully demonstrating how to use our framework to derive bounds and algorithms for the very fundamental *external regret* setting. Then, we derive the same types of existential guarantees in the much more general *subsequence regret* setting. We then specialize these subsequence regret bounds into tight bounds for various existing regret notions (such as internal, adaptive, sleeping experts, and multigroup regret). We conclude this subsection by deriving a general no-subsequence-regret algorithm which in turn specializes to an efficient algorithm in all of our applications.

E.1 Simple Learning From Expert Advice: External Regret

In the classical experts learning setting Littlestone and Warmuth [1994], the Learner has a set of pure actions (“experts”) $\mathcal{A}$. At the outset of each round $t \in [T]$, the Learner chooses a distribution over experts $x^t \in \Delta\mathcal{A}$. The Adversary then comes up with a vector of losses $r^t = (r_a^t)_{a \in \mathcal{A}} \in [0, 1]^{\mathcal{A}}$ corresponding to each expert. Next, the Learner samples $a^t \sim x^t$, and experiences loss corresponding to the expert she chose: $r_{a^t}^t$. The Learner also gets to observe the entire vector of losses r^t for that round. The goal of the Learner is to achieve sublinear *external regret* — that is, to ensure that the difference between her cumulative loss and the loss of the best fixed expert in hindsight grows sublinearly with T :

$$R_{\text{ext}}^T(\pi^T) := \sum_{t \in [T]} r_{a^t}^t - \min_{j \in \mathcal{A}} \sum_{t \in [T]} r_j^t = o(T).$$

Theorem E.1. *Fix a finite pure action set $\mathcal{A}$ for the Learner and a time horizon $T \geq \ln |\mathcal{A}|$. Then, Algorithm 2 can be instantiated to guarantee that the Learner’s expected external regret is bounded as*

$$\mathbb{E}_{\pi^T} [R_{\text{ext}}^T(\pi^T)] \leq 4\sqrt{T \ln |\mathcal{A}|},$$

and furthermore that for any $\delta \in (0, 1)$, with ex-ante probability $1 - \delta$ over the Learner’s randomness,

$$R_{\text{ext}}^T(\pi^T) \leq 8\sqrt{T \ln \frac{|\mathcal{A}|}{\delta}}.$$

Proof. We instantiate our probabilistic framework (see Section A.2.1).

Defining the strategy spaces. We define the Learner’s pure action set at each round to be the set $\mathcal{A}$, and the Adversary’s strategy space to be the convex and compact set $[0, 1]^{|\mathcal{A}|}$, from which the Adversary chooses each round’s collection $(r_a^t)_{a \in \mathcal{A}}$ of all actions’ losses.

Defining the loss functions. For $d = |\mathcal{A}|$, we define a d -dimensional vector valued loss function $\ell^t = (\ell_j^t)_{j \in \mathcal{A}}$, where for every action $j \in \mathcal{A}$, the corresponding coordinate $\ell_j^t : \mathcal{A} \times [0, 1]^{|\mathcal{A}|} \rightarrow [-1, 1]$ is given by

$$\ell_j^t(a, r^t) = r_a^t - r_j^t, \quad \text{for } a \in \mathcal{A}, r^t \in [0, 1]^{|\mathcal{A}|}.$$

It is easy to see that $\ell_j^t(a, \cdot)$ is continuous and concave — in fact, linear — in the second argument for all $j, a \in \mathcal{A}$ and $t \in [T]$. Furthermore, its range is $[-C, C]$, for $C = 1$. This verifies the technical conditions imposed by our framework on the loss functions.

Applying AMF regret bounds. We may now invoke Theorem A.1, which implies the following in-expectation AMF regret bound after round T for the instantiation of Algorithm 2 with the just defined vector losses $(\ell^t)_{t \in [T]}$:

$$\mathbb{E} \left[\max_{j \in \mathcal{A}} \sum_{t \in [T]} \ell_j^t(a^t, r^t) - \sum_{t \in [T]} w_A^t \right] \leq 4C\sqrt{T \ln d} = 4\sqrt{T \ln |\mathcal{A}|},$$

where recall that w_A^t is the Adversary-Moves-First (AMF) value at round t . Connecting the instantiated AMF regret to the Learner's external regret, we get:

$$\mathbb{E}[R_{\text{ext}}^T] = \mathbb{E}\left[\max_{j \in \mathcal{A}} \sum_{t \in [T]} r_{a^t}^t - r_j^t\right] = \mathbb{E}\left[\max_{j \in \mathcal{A}} \sum_{t \in [T]} \ell_j^t(a^t, r^t)\right] \leq 4\sqrt{T \ln |\mathcal{A}|} + \sum_{t \in [T]} w_A^t.$$

Bounding the Adversary-Moves-First value. To obtain the claimed in-expectation external regret bound, it suffices to show that the AMF value at each round $t \in [T]$ satisfies $w_A^t \leq 0$. Intuitively, this holds because if at some round the Learner knew the Adversary's choice of losses $(r_a^t)_{a \in \mathcal{A}}$ in advance, then she could guarantee herself no added loss in that round by picking the action $a \in \mathcal{A}$ with the smallest loss r_a^t .

Formally, for any vector of actions' losses r^t , define $a_{r^t}^* := \operatorname{argmin}_{a \in \mathcal{A}} r_a^t$, and notice that

$$\min_{a \in \mathcal{A}} \max_{j \in \mathcal{A}} \ell_j^t(a, r^t) \leq \max_{j \in \mathcal{A}} \ell_j^t(a_{r^t}^*, r^t) = \max_{j \in \mathcal{A}} (r_{a_{r^t}^*}^t - r_j^t) = \min_{a \in \mathcal{A}} r_a^t - \min_{j \in \mathcal{A}} r_j^t = 0.$$

The third step follows by definition of $a_{r^t}^*$. Hence, the AMF value is indeed nonpositive at each round:

$$w_A^t = \sup_{r^t \in [0,1]^{|\mathcal{A}|}} \min_{a \in \mathcal{A}} \max_{j \in \mathcal{A}} \ell_j^t(a, r^t) \leq 0.$$

This completes the proof of the in-expectation external regret bound. The high-probability external regret bound follows in the same way from Theorem A.2 of Section A.2.1. $\square$

A bound of $\sqrt{T \ln |\mathcal{A}|}$ is optimal for external regret in the experts learning setting, and so serves to witness the optimality of Theorem 2.1.

In fact, it is easy to demonstrate that in the external regret setting, the generic probabilistic Algorithm 2 amounts to the well known Exponential Weights algorithm (Algorithm 5 below) Littlestone and Warmuth [1994]. To see this, note that Algorithm 2, when instantiated with the above defined loss functions, has the Learner solve the following problem at each round:

$$\begin{aligned} x^t &\in \operatorname{argmin}_{x \in \Delta \mathcal{A}} \max_{r^t \in [0,1]^{|\mathcal{A}|}} \sum_{j \in \mathcal{A}} \frac{\exp\left(\eta \sum_{s=1}^{t-1} (r_{a^s}^s - r_j^s)\right)}{\sum_{i \in \mathcal{A}} \exp\left(\eta \sum_{s=1}^{t-1} (r_{a^s}^s - r_i^s)\right)} \mathbb{E}_{a \sim x} [r_a^t - r_j^t], \\ &= \operatorname{argmin}_{x \in \Delta \mathcal{A}} \max_{r^t \in [0,1]^{|\mathcal{A}|}} \sum_{j \in \mathcal{A}} \frac{\exp\left(-\eta \sum_{s=1}^{t-1} r_j^s\right)}{\sum_{i \in \mathcal{A}} \exp\left(-\eta \sum_{s=1}^{t-1} r_i^s\right)} \mathbb{E}_{a \sim x} [r_a^t - r_j^t], \\ &= \operatorname{argmin}_{x \in \Delta \mathcal{A}} \max_{r^t \in [0,1]^{|\mathcal{A}|}} \mathbb{E}_{a \sim x, j \sim \text{EW}_\eta(\pi^{t-1})} [r_a^t - r_j^t], \end{aligned}$$

where we denoted the exponential weights distribution as

$$\text{EW}_\eta(\pi^{t-1}) := \left(\frac{\exp\left(-\eta \sum_{s=1}^{t-1} r_j^s\right)}{\sum_{i \in \mathcal{A}} \exp\left(-\eta \sum_{s=1}^{t-1} r_i^s\right)} \right)_{j \in \mathcal{A}} \in \Delta \mathcal{A}.$$

For any choice of r^t by the Adversary, the quantity inside the expectation, $\ell_j^t(a, r^t) = r_a^t - r_j^t$, is *antisymmetric* in a and j : that is, $\ell_j^t(a, r^t) = -\ell_a^t(j, r^t)$. Due to this antisymmetry, no matter which r^t gets selected by the Adversary, by playing $a \sim \text{EW}_\eta(\pi^{t-1})$ the Learner obtains

$$\mathbb{E}_{a, j \sim \text{EW}_\eta(\pi^{t-1})} [r_a^t - r_j^t] = 0,$$

thus achieving the value of the game. It is also easy to see that $x^t = \text{EW}_\eta(\pi^{t-1})$ is the unique choice of x^t that guarantees nonnegative value, hence Algorithm 2, when specialized to the external regret setting, is *equivalent* to the Exponential Weights Algorithm 5.

Algorithm 5: The Exponential Weights Algorithm with Learning Rate η

for $t = 1, \dots, T$ **do**

Sample a^t such that $a^t = j$ with probability proportional to $\exp\left(-\eta \sum_{s=1}^{t-1} r_j^s\right)$, for $j \in \mathcal{A}$.

E.2 Generalization to Subsequence Regret

Here, we present a generalization of the experts learning framework from which we will be able to derive our other applications to no-regret learning problems. There is again a Learner and an Adversary playing over the course of rounds $t \in [T]$. Initially, the Learner is endowed with a finite set of pure actions $\mathcal{A}$. At each round t , the Adversary restricts the Learner's set of available actions for that round to some subset $\mathcal{A}^t \subseteq \mathcal{A}$. The Learner plays a mixture $x^t \in \Delta \mathcal{A}^t$ over the available actions. The Adversary responds by selecting a vector of losses $(r_a^t)_{a \in \mathcal{A}} \in [0, 1]^{|\mathcal{A}|}$ associated with the Learner's pure actions. Next, the Learner samples a pure action $a^t \sim x^t$.

Unlike in the standard setting, the Learner's regret will now be measured not just on the entire sequence of rounds $1, 2, \dots, T$, but more generally on an arbitrary collection $\mathcal{F}$ of *weighted subsequences* $f : [T] \times \mathcal{A} \rightarrow [0, 1]$. The understanding is that for any $f \in \mathcal{F}, t \in [T], a \in \mathcal{A}^t$, the quantity $f(t, a)$ is the "weight" with which round t will be included in the subsequence if the Learner's sampled action is a at that round. The Learner does *not* need to know the subsequences ahead of time; instead the Adversary may announce the values $\{f(t, a)\}_{a \in \mathcal{A}^t, f \in \mathcal{F}}$ to the Learner before the corresponding round $t \in [T]$.

Definition E.1 (Subsequence Regret). *Given a family of functions $\mathcal{F}$, where each $f \in \mathcal{F}$ is a mapping $f : [T] \times \mathcal{A} \rightarrow [0, 1]$, chosen adaptively by the Adversary, and a set of finitely many pure actions $\mathcal{A}$ for the Learner, consider a collection of action-subsequence pairs $\mathcal{H} \subseteq \mathcal{A} \times \mathcal{F}$.*

The Learner's subsequence regret after round T with respect to the collection $\mathcal{H}$ is defined by

$$R_{\mathcal{H}}^T(\pi^T) := \max_{(j, f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) (r_{a^t}^t - r_j^t),$$

where $\pi^T = \{(a^t, r^t)\}_{t \in [T]}$ is the transcript of the interaction.

For intuition, suppose $\mathcal{F} = \{\mathbf{1}\}$, where $\mathbf{1} : [T] \times \mathcal{A} \rightarrow [0, 1]$ satisfies $\mathbf{1}(t, a) = 1$ for all t, a . That is, the only relevant subsequence is the entire sequence of rounds $1, 2, \dots, T$. If we then set $\mathcal{H} = \mathcal{A} \times \mathcal{F}$, subsequence regret specializes to the classical notion of (external) regret which was discussed above.

Moreover, we shall require the following condition on $\mathcal{H}$ and the action sets $\{\mathcal{A}^t\}_{t \in [T]}$, which simply asks that at each round, the Learner be responsible for regret only to currently available actions.

Definition E.2 (No regret to unavailable actions). *A collection of action-subsequence pairs $\mathcal{H}$, paired with action sets $\{\mathcal{A}^t\}_{t \in [T]}$, satisfy the no-regret-to-unavailable-actions property if at each round $t \in [T]$, for every $f \in \mathcal{F}$ such that $(j, f) \in \mathcal{H}$ for some $j \notin \mathcal{A}^t$, it holds that $f(t, a) = 0$ for all $a \in \mathcal{A}^t$.*

It is worth noting that this condition is trivially satisfied whenever the Learner's action set is invariant across rounds ($\mathcal{A}^t = \mathcal{A}$ for all t).

Theorem E.2. *Consider a sequence of action sets $\{\mathcal{A}^t\}_{t \in [T]}$ for the Learner, a collection $\mathcal{H}$ of action-subsequence pairs, and a time horizon $T \geq \ln |\mathcal{H}|$. If $\mathcal{H}$ and $\{\mathcal{A}^t\}_{t \in [T]}$ satisfy no-regret-to-unavailable-actions, then an appropriate instantiation of Algorithm 2 guarantees that the Learner's expected subsequence regret is bounded as*

$$\mathbb{E}_{\pi^T} [R_{\mathcal{H}}^T(\pi^T)] \leq 4\sqrt{T \ln |\mathcal{H}|},$$

and furthermore, for any $\delta \in (0, 1)$, that with ex-ante probability $1 - \delta$ over the Learner's randomness,

$$R_{\mathcal{H}}^T(\pi^T) \leq 8\sqrt{T \ln \frac{|\mathcal{H}|}{\delta}}.$$

Proof. We instantiate our probabilistic framework of Section A.2.1.

Defining the strategy spaces. At each round t , the Learner's pure strategy set will be $\mathcal{A}^t$, and the Adversary's strategy space will be the convex and compact set $[0, 1]^{|\mathcal{A}|}$.

Defining the loss functions. For all action-subsequence pairs $(j, f) \in \mathcal{H}$, we define the corresponding loss $\ell_{(j,f)}^t : \mathcal{A}^t \times [0, 1]^{|\mathcal{A}|} \rightarrow [-1, 1]$ as

$$\ell_{(j,f)}^t(a, r^t) = f(t, a)(r_a^t - r_j^t), \quad \text{for } a \in \mathcal{A}^t, r^t \in [0, 1]^{|\mathcal{A}|}.$$

It is easy to see that for all $(j, f) \in \mathcal{H}$ and each $a \in \mathcal{A}^t$, the function $\ell_{(j,f)}^t(a, \cdot)$ is continuous and concave — in fact, linear — in the second argument, as well as bounded within $[-C, C]$ for $C = 1$. Therefore, the technical conditions imposed by our framework on the loss functions are met.

Bounding the Adversary-Moves-First value. At each round t , the AMF value $w_A^t = 0$. Trivially, $w_A^t \geq 0$, as the Adversary can always set $r_a^t = 0$ for all a . Conversely, $w_A^t \leq 0$ as an easy consequence of the no-regret-to-unavailable-actions property. To see this, for any vector of actions' losses r^t , define

$$a_{r^t}^* := \operatorname{argmin}_{a \in \mathcal{A}^t} r_a^t,$$

and notice that

$$\begin{aligned} w_A^t &= \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \min_{a \in \mathcal{A}^t} \left(\max_{(j,f) \in \mathcal{H}} \ell_{(j,f)}^t(a, r^t) \right), \\ &= \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \min_{a \in \mathcal{A}^t} \max \left(\max_{(j,f) \in \mathcal{H}: j \in \mathcal{A}^t} \ell_{(j,f)}^t(a, r^t), 0 \right), \quad (\text{no regret to unavailable actions}) \\ &\leq \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \max \left(\max_{(j,f) \in \mathcal{H}: j \in \mathcal{A}^t} \ell_{(j,f)}^t(a_{r^t}^*, r^t), 0 \right), \\ &= \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \max \left(\max_{(j,f) \in \mathcal{H}: j \in \mathcal{A}^t} f(t, a_{r^t}^*)(r_{a_{r^t}^*}^t - r_j^t), 0 \right), \\ &\leq \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \max \left(\max_{(j,f) \in \mathcal{H}: j \in \mathcal{A}^t} f(t, a_{r^t}^*)(r_j^t - r_j^t), 0 \right), \quad (\text{by definition of } a_{r^t}^*) \\ &= \sup_{r^t \in [0, 1]^{|\mathcal{A}|}} \max(0, 0), \\ &= 0. \end{aligned}$$

We thus conclude that Theorems A.1 and A.2 apply (with $C = 1$ and all $w_A^t = 0$) to the subsequence regret setting, yielding the claimed in-expectation and high-probability regret bounds. $\square$

We now instantiate subsequence regret with various choices of subsequence families, in order to get bounds and efficient algorithms for several standard notions of regret from the literature. For brevity, for each notion of regret considered below we only exhibit the existential in-expectation guarantee for that type of regret, and omit the corresponding high-probability bounds (which are all easily derivable from Theorem A.2). We also point out that all in-expectation bounds cited below are efficiently achievable by instantiating, with appropriate loss functions, the no-subsequence regret Algorithms 6 and 7 derived in the following Section E.3.

In all no-regret settings discussed below, except for Sleeping Experts, the Learner has a pure and finite action set $\mathcal{A}$ at every round $t \in [T]$; furthermore — as usual — the Adversary’s role at each round consists in selecting the vector of per-action losses $(r_a^t)_{a \in \mathcal{A}} \in [0, 1]^{|\mathcal{A}|}$.

Internal and Swap Regret To introduce the notion of *internal regret* [Foster and Vohra, 1998], consider the following collection $\mathcal{M}_{\text{int}} \subset \mathcal{A}^{\mathcal{A}}$ of mappings from the action set $\mathcal{A}$ to itself. $\mathcal{M}_{\text{int}}$ consists of the identity map μ_{id} (such that $\mu_{\text{id}}(a) = a$ for all $a \in \mathcal{A}$), together with all $|\mathcal{A}|(|\mathcal{A}| - 1)$ maps $\mu_{i \rightarrow j}$ that pair two particular actions: i.e., $\mu_{i \rightarrow j}(i) = j$, and $\mu_{i \rightarrow j}(a) = a$ for $a \neq i$. The Learner’s internal regret is then defined as

$$R_{\text{int}}^T := \max_{\mu \in \mathcal{M}_{\text{int}}} \sum_{t \in [T]} r_{a^t}^t - r_{\mu(a^t)}^t.$$

In other words, the Learner’s total loss is being compared to all possible counterfactual worlds, for $i, j \in \mathcal{A}$, in which whenever the Learner played some action i , it got replaced with action j (and other actions remain fixed).

We can reduce the problem of obtaining no-internal-regret to the problem of obtaining no subsequence regret for a simple choice of subsequences. Let us define the following set of subsequences: $\mathcal{F} := \{f_i : i \in \mathcal{A}\}$, where each f_i is defined to be the indicator of the subsequence where the Learner played action i — that is, for all $t \in [T]$, we let $f_i(t, a) = 1_{a=i}$. Then, we let $\mathcal{H} := \mathcal{A} \times \mathcal{F}$. By the in-expectation no-subsequence-regret guarantee, we then have

$$\mathbb{E} \left[\max_{(j, f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) (r_{a^t}^t - r_j^t) \right] \leq 4\sqrt{T \ln |\mathcal{H}|} = 4\sqrt{2T \ln |\mathcal{A}|},$$

since $|\mathcal{H}| = |\mathcal{A}| \cdot |\mathcal{F}| = |\mathcal{A}|^2$.

But observe that the Learner’s internal regret precisely coincides with the just defined instance of subsequence regret:

$$\begin{aligned} R_{\text{int}}^T &= \max_{\mu \in \mathcal{M}_{\text{int}}} \sum_{t \in [T]} r_{a^t}^t - r_{\mu(a^t)}^t = \max_{i, j \in \mathcal{A}} \sum_{t \in [T]: a^t = i} r_i^t - r_j^t = \max_{j \in \mathcal{A}} \max_{f_i: i \in \mathcal{A}} \sum_{t \in [T]} f_i(t, a^t) (r_{a^t}^t - r_j^t) \\ &= \max_{(j, f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) (r_{a^t}^t - r_j^t). \end{aligned}$$

Therefore, we have established the following existential in-expectation internal regret bound:

$$\mathbb{E} [R_{\text{int}}^T] \leq 4\sqrt{2T \ln |\mathcal{A}|},$$

which is optimal.

The notion of *swap regret*, introduced in Blum and Mansour [2007], is strictly more demanding than internal regret in that it considers strategy modification rules μ that can perform more than one action swap at a time. Consider the set $\mathcal{M}_{\text{swap}}$ of all $|\mathcal{A}|^{|\mathcal{A}|}$ *swapping rules* $\mu : \mathcal{A} \rightarrow \mathcal{A}$. The Learner’s swap regret is defined to be the maximum of her regret to all swapping rules:

$$R_{\text{swap}}^T := \max_{\mu \in \mathcal{M}_{\text{swap}}} \sum_{t \in [T]} r_{a^t}^t - r_{\mu(a^t)}^t.$$

The interpretation is that the Learner’s total loss is being compared to the total loss of any remapping of her action sequence.

An easy reduction shows that the swap regret is upper-bounded by $|\mathcal{A}|$ times the internal regret. For completeness, we provide the details of this reduction in Appendix E.4. The reduction implies an in-expectation bound of $4|\mathcal{A}|\sqrt{2T \ln |\mathcal{A}|}$ on swap regret, which, compared to the optimal bound of $O(\sqrt{T|\mathcal{A}| \ln |\mathcal{A}|})$ (see Blum and Mansour [2007]), has suboptimal dependence on $|\mathcal{A}|$.

Adaptive Regret In this setting, consider all contiguous time intervals within rounds $1, \dots, T$, namely, all intervals $[t_1, t_2]$, where t_1, t_2 are integers such that $1 \leq t_1 \leq t_2 \leq T$. The Learner’s regret on each interval $[t_1, t_2]$ is defined as her total loss over the rounds $t \in [t_1, t_2]$, minus the loss of the best action for that interval in hindsight. The Learner’s adaptive regret is then defined to be her maximum regret over all contiguous time intervals:

$$R_{\text{adaptive}}^T := \max_{[t_1, t_2]: 1 \leq t_1 \leq t_2 \leq T} \max_{j \in \mathcal{A}} \sum_{t=t_1}^{t_2} r_{a^t}^t - r_j^t.$$

We observe that adaptive regret corresponds to subsequence regret with respect to $\mathcal{H} := \mathcal{A} \times \mathcal{F}$, where $\mathcal{F} := \{f_{[t_1, t_2]} : 1 \leq t_1 \leq t_2 \leq T\}$ is the collection of subinterval indicator subsequences — that is, $f_{[t_1, t_2]}(t, a) := 1_{t_1 \leq t \leq t_2}$ for all $t \in [T]$ and $a \in \mathcal{A}$. Observe that $|\mathcal{F}| \leq T^2$, and therefore, the expected regret upper bound for subsequence regret specializes to the following expected adaptive regret bound:

$$\mathbb{E} [R_{\text{adaptive}}^T] \leq 4\sqrt{T \ln(|\mathcal{A}||\mathcal{F}|)} \leq 4\sqrt{T(\ln|\mathcal{A}| + 2 \ln T)}.$$

Sleeping Experts Following Blum and Mansour [2007], we define the sleeping experts setting as follows. Suppose that the Learner is initially given a set of pure actions $\mathcal{A}$, and before each round t , the Adversary chooses a subset of pure actions $\mathcal{A}^t \subseteq \mathcal{A}$ available to the Learner at that round — these are known as the “awake experts”, and the rest of the experts are the “sleeping experts” at that round.

The Learner’s regret to each action $j \in \mathcal{A}$ is defined to be the excess total loss of the Learner during rounds where j was “awake”, compared to the total loss of j over those rounds. Formally, the Learner’s sleeping experts regret after round T is defined to be

$$R_{\text{sleeping}}^T := \max_{j \in \mathcal{A}} \sum_{t \in [T]: j \in \mathcal{A}^t} r_{a^t}^t - r_j^t.$$

This is clearly an instance of subsequence regret — indeed, we may consider the family of subsequences $\mathcal{F} := \{f_j : j \in \mathcal{A}\}$, where $f_j(t, a) := 1_{j \in \mathcal{A}^t}$ for all j, a, t , and let $\mathcal{H} := \{(j, f_j)\}_{j \in \mathcal{A}}$. It is easy to verify that the no-regret-to-unavailable-actions property holds, and thus the guarantees of the subsequence regret setting carry over to this sleeping experts setting. In particular, the following existential in-expectation sleeping experts regret bound holds:

$$\mathbb{E} [R_{\text{sleeping}}^T] \leq 4\sqrt{T \ln |\mathcal{A}|},$$

which is also optimal in this setting.

Multi-Group Regret We imagine that before each round, the Adversary selects and reveals to the Learner some *context* θ^t from an underlying feature space Θ . The interpretation is that the Learner’s decision at round t will pertain to an individual with features θ^t . Additionally, there is a fixed collection $\mathcal{G} \subset 2^\Theta$, where each $g \in \mathcal{G}$ is interpreted as a (demographic) group of individuals within the population Θ . Here $\mathcal{G}$ may be large and may consist of overlapping groups. The Learner’s goal is to minimize regret to each action $a \in \mathcal{A}$ not just over the entire population, but also separately for each population group $g \in \mathcal{G}$. Explicitly, the Learner’s multi-group regret after round T is defined to be

$$R_{\text{multi}}^T := \max_{g \in \mathcal{G}} \max_{j \in \mathcal{A}} \sum_{t \in [T]: \theta^t \in g} r_{a^t}^t - r_j^t.$$

It is easy to see that multi-group regret corresponds to subsequence regret with $\mathcal{H} := \mathcal{A} \times \mathcal{F}$, where $\mathcal{F} := \{f_g : g \in \mathcal{G}\}$ is the collection of group indicator subsequences — that is, $f_g(t, a) := 1_{\theta^t \in g}$ for all t, a . Here we are taking advantage of the fact that the functions f on which subsequences are defined need not be known to the algorithm ahead of time, and can be revealed sequentially by the Adversary, allowing us to model adversarially chosen contexts. Therefore, multi-group regret inherits subsequence regret guarantees, and in particular, we obtain the following existential in-expectation multi-group regret bound:

$$\mathbb{E} [R_{\text{multi}}^T] \leq 4\sqrt{T \ln(|\mathcal{A}||\mathcal{G}|)}.$$

Observe that this bound scales only as $\sqrt{\ln |\mathcal{G}|}$ with respect to the number of population groups, which we can therefore take to be exponentially large in the parameters of the problem.

E.3 Deriving No-Subsequence-Regret Algorithms

We now present a way to specialize Algorithm 2 to the setting of subsequence regret with no-regret-to-unavailable-actions. At each round, instead of solving a convex-concave problem, the specialized algorithm will only need to solve a polynomial-sized linear program.

Algorithm 6: Efficient No Subsequence Regret Algorithm for the Learner

for $t = 1, \dots, T$ **do**

 Learn the current set of feasible actions $\mathcal{A}^t$ (potentially selected by an Adversary).

 Learn the values $f(t, a)$ for every $a \in \mathcal{A}^t$ and $f \in \mathcal{F}$ (potentially selected by an Adversary).

 Solve for $x^t = (x_a^t)_{a \in \mathcal{A}^t} \in \Delta \mathcal{A}^t$ defined by the following linear inequalities for all $a \in \mathcal{A}^t$:

$$x_a^t \sum_{(j,f) \in \mathcal{H}} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s) \right) f(t, a) - \sum_{j \in \mathcal{A}^t} x_j^t \sum_{f: (a,f) \in \mathcal{H}} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(a,f)}^s(a^s, r^s) \right) f(t, j) \leq 0$$

 Sample $a^t \sim x^t$.

Theorem E.3. *Algorithm 6 implements Algorithm 2 in the subsequence regret setting, and achieves the same guarantees.*

Proof. In parallel to the notation of Algorithm 2, we define the following set of weights at round $t \in [T]$:

$$\chi_{(j,f)}^t := \frac{1}{Z^t} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s) \right),$$

where

$$Z^t := \sum_{(j,f) \in \mathcal{H}} \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s) \right).$$

When instantiated with our current set of loss functions, Algorithm 2 solves the following zero-sum game at round $t \in [T]$, where we denote $\ell_{(j,f)}^t(x, r^t) := \mathbb{E}_{a \sim x}[\ell_{(j,f)}^t(a, r^t)]$:

$$x^t \in \operatorname{argmin}_{x \in \Delta \mathcal{A}^t} \max_{r^t \in [0,1]^{|\mathcal{A}|}} \sum_{(j,f) \in \mathcal{H}} \chi_{(j,f)}^t \cdot \ell_{(j,f)}^t(x, r^t).$$

By definition of the loss functions in the subsequence regret setting, the objective function is linear in the Adversary's choice of r^t . Thus, let us rewrite the objective as a linear combination of $(r_a^t)_{a \in \mathcal{A}^t}$:

$$\begin{aligned} & \sum_{(j,f) \in \mathcal{H}} \chi_{(j,f)}^t \cdot \ell_{(j,f)}^t(x, r^t), \\ &= \sum_{(j,f) \in \mathcal{H}} \chi_{(j,f)}^t \sum_{a \in \mathcal{A}^t} x_a \cdot f(t, a) \cdot (r_a^t - r_j^t), \\ &= \sum_{(j,f) \in \mathcal{H}} \sum_{a \in \mathcal{A}^t} r_a^t \cdot x_a \cdot f(t, a) \cdot \chi_{(j,f)}^t - \sum_{(j,f) \in \mathcal{H}} \sum_{a \in \mathcal{A}^t} r_j^t \cdot x_a \cdot f(t, a) \cdot \chi_{(j,f)}^t, \end{aligned}$$

which, by the no-regret-to-unavailable actions property,

$$= \sum_{a \in \mathcal{A}^t} r_a^t \cdot x_a \sum_{(j,f) \in \mathcal{H}} f(t, a) \cdot \chi_{(j,f)}^t - \sum_{j \in \mathcal{A}^t} r_j^t \sum_{a \in \mathcal{A}^t} x_a \sum_{f: (j,f) \in \mathcal{H}} f(t, a) \cdot \chi_{(j,f)}^t,$$

and now, swapping j and a in the second summation,

$$\begin{aligned}
&= \sum_{a \in \mathcal{A}^t} r_a^t \cdot x_a \sum_{(j,f) \in \mathcal{H}} f(t,a) \cdot \chi_{(j,f)}^t - \sum_{a \in \mathcal{A}^t} r_a^t \sum_{j \in \mathcal{A}^t} x_j \sum_{f:(a,f) \in \mathcal{H}} f(t,j) \cdot \chi_{(a,f)}^t, \\
&= \sum_{a \in \mathcal{A}^t} r_a^t \underbrace{\left(x_a \sum_{(j,f) \in \mathcal{H}} f(t,a) \cdot \chi_{(j,f)}^t - \sum_{j \in \mathcal{A}^t} x_j \sum_{f:(a,f) \in \mathcal{H}} f(t,j) \cdot \chi_{(a,f)}^t \right)}_{:= c_a(x)}.
\end{aligned}$$

Thus, the zero-sum game played at round t has objective function $\sum_{a \in \mathcal{A}^t} c_a(x^t) \cdot r_a^t$, where the coefficients $c_a(x^t)$ do not depend on the Adversary's action r^t . Recall that this game has value at most $w_A^t = 0$. Hence, $\max_{a \in \mathcal{A}^t} c_a(x^t) \leq 0$ for any minimax optimal strategy x^t for the Learner — since otherwise, if some $c_{a'}(x^t) > 0$, the Adversary would get value $c_{a'}(x^t) > 0$ by setting $r_{a'}^t = 1$ and $r_a^t = 0$ for $a \neq a'$. Conversely, by playing x^t such that $\max_{a \in \mathcal{A}^t} c_a(x^t) \leq 0$, the Learner gets value ≤ 0 , as $r_a^t \geq 0$ for all a .

Therefore, the Learner's choice of x^t is minimax optimal if and only if for all $a \in \mathcal{A}^t$,

$$\begin{aligned}
c_a(x^t) \leq 0 &\iff Z^t \cdot c_a(x^t) \leq 0 \iff \\
x_a^t \sum_{(j,f) \in \mathcal{H}} f(t,a) \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s) \right) - \sum_{j \in \mathcal{A}^t} x_j^t \sum_{f:(a,f) \in \mathcal{H}} f(t,j) \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(a,f)}^s(a^s, r^s) \right) &\leq 0.
\end{aligned}$$

This recovers Algorithm 6, concluding the proof. $\square$

Simplification for Action Independent Subsequences The above Algorithm 6 requires solving a linear feasibility problem. This mirrors how existing algorithms for the special case of minimizing internal regret operate (Blum and Mansour [2007]); recall that internal regret corresponds to subsequence regret for a certain collection of $|\mathcal{A}|$ subsequences that depend on the Learner's action in the current round t .

By contrast, if all of our subsequence indicators $f \in \mathcal{F}$ are *action independent*, that is, satisfy $f(t,a) = f(t,a')$ for all $a, a' \in \mathcal{A}$ and $t \in [T]$, then it turns out that we can avoid solving a system of linear inequalities: our equilibrium has a closed form. In what follows, we abuse notation and simply write $f(t)$ for the value of the subsequence f at round t .

Observe that if each $f \in \mathcal{F}$ is action independent, then we can rewrite our equilibrium characterization in Algorithm 6 as the requirement that the Learner's chosen distribution $x^t \in \Delta \mathcal{A}^t$ must satisfy, for each $a \in \mathcal{A}^t$ (provided that $f(t) \neq 0$ for at least some $f \in \mathcal{F}$), the following inequality:

$$\begin{aligned}
x_a^t &\leq \frac{\sum_{j \in \mathcal{A}^t} x_j^t \sum_{f:(a,f) \in \mathcal{H}} f(t) \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(a,f)}^s(a^s, r^s) \right)}{\sum_{(j,f) \in \mathcal{H}} f(t) \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s) \right)}, \\
&= \frac{\sum_{f:(a,f) \in \mathcal{H}} f(t) \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(a,f)}^s(a^s, r^s) \right)}{\sum_{(j,f) \in \mathcal{H}} f(t) \exp \left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s) \right)}.
\end{aligned}$$

Here the equality follows because $x^t \in \Delta \mathcal{A}^t$ is a probability distribution.

We now observe that setting each x_a^t to be its upper bound, for $a \in \mathcal{A}^t$, yields a probability distribution over $\mathcal{A}^t$, which is consequently the unique feasible solution to the above system. Hence, for action independent subsequences, we have a closed-form implementation of Algorithm 6 that does not require solving a linear

feasibility problem:

Algorithm 7: An Efficient Learner for Action Independent Subsequences

for $t = 1, \dots, T$ **do**

 Learn the current set of feasible actions $\mathcal{A}^t$ and the values $f(t)$ for every $f \in \mathcal{F}$ (potentially selected by an Adversary).

 Sample $a^t \sim x^t$, where for all $a \in \mathcal{A}^t$,

$$x_a^t = \frac{\sum_{f:(a,f) \in \mathcal{H}} f(t) \exp\left(\eta \sum_{s=1}^{t-1} \ell_{(a,f)}^s(a^s, r^s)\right)}{\sum_{(j,f) \in \mathcal{H}} f(t) \exp\left(\eta \sum_{s=1}^{t-1} \ell_{(j,f)}^s(a^s, r^s)\right)}.$$

E.4 Omitted Reductions between Different Notions of Regret

Reducing swap regret to internal regret We can upper bound the swap regret by reusing the instance of subsequence regret that we defined to capture internal regret. Recall that it was defined as follows. We let $\mathcal{F} := \{f_i : i \in \mathcal{A}\}$, where each f_i is the indicator of the subsequence of rounds where the Learner played action i — that is, for all $t \in [T]$, we let $f(t, a) = 1_{a=i}$. Then, we let $\mathcal{H} := \mathcal{A} \times \mathcal{F}$. We then obtained the in-expectation regret guarantee

$$\mathbb{E} \left[\max_{(j,f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) (r_{a^t}^t - r_j^t) \right] \leq 4\sqrt{2T \ln |\mathcal{A}|}.$$

Returning to swap regret, note that for any fixed swapping rule $\mu : \mathcal{A} \rightarrow \mathcal{A}$, we have

$$\begin{aligned} \sum_{t \in [T]} r_{a^t}^t - r_{\mu(a^t)}^t &= \sum_{i \in \mathcal{A}} \sum_{t \in [T]: a^t = i} r_{a^t}^t - r_{\mu(i)}^t \\ &\leq \sum_{i \in \mathcal{A}} \max_{j \in \mathcal{A}} \sum_{t \in [T]: a^t = i} r_{a^t}^t - r_j^t \\ &\leq |\mathcal{A}| \max_{i \in \mathcal{A}} \max_{j \in \mathcal{A}} \sum_{t \in [T]: a^t = i} r_{a^t}^t - r_j^t \\ &= |\mathcal{A}| \max_{(j,f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) (r_{a^t}^t - r_j^t), \end{aligned}$$

where in the last line we simply reparametrized the maximum over $i \in \mathcal{A}$ as the maximum over all $f \in \mathcal{F}$. Since the above holds for any $\mu \in \mathcal{M}_{\text{swap}}$, we have

$$R_{\text{swap}}^t = \max_{\mu \in \mathcal{M}_{\text{swap}}} \sum_{t \in [T]} r_{a^t}^t - r_{\mu(a^t)}^t \leq |\mathcal{A}| \max_{(j,f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) (r_{a^t}^t - r_j^t),$$

and therefore, we conclude that there exists an efficient algorithm that achieves expected swap regret

$$\mathbb{E} [R_{\text{swap}}^T] \leq 4|\mathcal{A}| \sqrt{2T \ln |\mathcal{A}|}.$$

Wide-range regret and its connection to subsequence regret The wide-range regret setting was first introduced in Lehrer [2003] and then studied, in particular, in Blum and Mansour [2007] and Greenwald and Jafari [2003]. It is quite general, and is in fact equivalent to the subsequence regret setting, up to a reparametrization.

Just as in the subsequence regret setting, imagine there is a finite family of subsequences $\mathcal{F}$, where each $f \in \mathcal{F}$ has the form $f : [T] \times \mathcal{A} \rightarrow [0, 1]$. Moreover, suppose there is a finite family $\mathcal{M}$ of *modification rules*.

Each modification rule $\mu \in \mathcal{M}$ is defined as a mapping $\mu : [T] \times \mathcal{A} \rightarrow \mathcal{A}$, which has the interpretation that if at time t , the Learner plays action a^t , then the modification rule modifies this action into another action $\mu(t, a^t) \in \mathcal{A}$. Now, consider a collection of modification rule-subsequence pairs $\mathcal{H} \subseteq \mathcal{M} \times \mathcal{F}$. The Learner's wide-range regret with respect to $\mathcal{H}$ is defined as

$$R_{\text{wide}}^T := \max_{(\mu, f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) \left(r_{a^t}^t - r_{\mu(t, a^t)}^t \right).$$

It is evident that wide-range regret has subsequence regret (when the Learner's action set $\mathcal{A}^t = \mathcal{A}$ for all $t \in [T]$) as a special case, where each modification rule $\mu \in \mathcal{M}$ always outputs the same action: that is, for all t, a^t , we have $\mu(t, a^t) = j$ for some $j \in \mathcal{A}$.

It is also not hard to establish the converse. Indeed, suppose we have an instance of no-wide-range-regret learning with $\mathcal{H} \subseteq \mathcal{M} \times \mathcal{F}$, where $\mathcal{M}$ is a family of modification rules and $\mathcal{F}$ is a family of subsequences. Fix any pair $(\mu, f) \in \mathcal{H}$. Then, let us define, for all $j \in \mathcal{A}$, the subsequence

$$\phi_j^{(\mu, f)} : [T] \times \mathcal{A} \rightarrow [0, 1] \text{ such that } \phi_j^{(\mu, f)}(t, a) := f(t, a) \cdot 1_{\mu(t, a)=j} \text{ for all } t \in [T], a \in \mathcal{A}.$$

Now, let us instantiate our subsequence regret setting with

$$\mathcal{H}_{\text{wide}} := \bigcup_{(\mu, f) \in \mathcal{H}} \bigcup_{j \in \mathcal{A}} (j, \phi_j^{(\mu, f)}).$$

Observe in particular that $|\mathcal{H}_{\text{wide}}| = |\mathcal{A}| \cdot |\mathcal{H}|$.

Computing the subsequence regret of this family $\mathcal{H}_{\text{wide}}$, we have

$$R_{\mathcal{H}_{\text{wide}}}^T = \max_{(\mu, f) \in \mathcal{H}} \max_{j \in \mathcal{A}} \sum_{t \in [T]: \mu(t, a^t)=j} f(t, a^t) (r_{a^t}^t - r_j^t).$$

Now, we have the following upper bound on the wide-range regret:

$$\begin{aligned} R_{\text{wide}}^T &= \max_{(\mu, f) \in \mathcal{H}} \sum_{t \in [T]} f(t, a^t) \left(r_{a^t}^t - r_{\mu(t, a^t)}^t \right) \\ &= \max_{(\mu, f) \in \mathcal{H}} \sum_{j \in \mathcal{A}} \sum_{t \in [T]: \mu(t, a^t)=j} f(t, a^t) (r_{a^t}^t - r_j^t) \\ &\leq \max_{(\mu, f) \in \mathcal{H}} |\mathcal{A}| \max_{j \in \mathcal{A}} \sum_{t \in [T]: \mu(t, a^t)=j} f(t, a^t) (r_{a^t}^t - r_j^t) \\ &= |\mathcal{A}| R_{\mathcal{H}_{\text{wide}}}^T. \end{aligned}$$

Since our subsequence regret results imply the existence of an algorithm such that $\mathbb{E} [R_{\mathcal{H}_{\text{wide}}}^T] \leq 4\sqrt{T \ln |\mathcal{H}'|} = 4\sqrt{T(\ln |\mathcal{A}| + \ln |\mathcal{H}|)}$, we have the following expected wide-range regret bound:

$$\mathbb{E} [R_{\text{wide}}^T] \leq 4|\mathcal{A}| \sqrt{T(\ln |\mathcal{A}| + \ln |\mathcal{H}|)}.$$

This figure "frog.jpg" is available in "jpg" format from:

<http://arxiv.org/ps/2108.03837v3>