

Dueling Optimization with a Monotone Adversary

Avrim Blum

Toyota Technological Institute Chicago

AVRIM@TTIC.EDU

Meghal Gupta

University of California Berkeley

MEGHAL@BERKELEY.EDU

Gene Li

Toyota Technological Institute Chicago

GENE@TTIC.EDU

Naren Sarayu Manoj

Toyota Technological Institute Chicago

NSM@TTIC.EDU

Aadirupa Saha

Toyota Technological Institute Chicago (currently at Apple Research)

AADIRUPA@TTIC.EDU

Yuanyuan Yang

University of Washington

YYANGH@CS.WASHINGTON.EDU

Editors: Claire Vernade and Daniel Hsu

Keywords: Online learning, Monotone adversary, Convex optimization

Abstract

We introduce and study the problem of *dueling optimization with a monotone adversary*, which is a generalization of (noiseless) dueling convex optimization. The goal is to design an online algorithm to find a minimizer $\mathbf{x}^*$ for a function $f: \mathcal{X} \rightarrow \mathbb{R}$, where $\mathcal{X} \subseteq \mathbb{R}^d$. In each round, the algorithm submits a pair of guesses, i.e., $\mathbf{x}^{(1)}$ and $\mathbf{x}^{(2)}$, and the adversary responds with *any* point in the space that is at least as good as both guesses. The cost of each query is the suboptimality of the worse of the two guesses; i.e., $\max(f(\mathbf{x}^{(1)}), f(\mathbf{x}^{(2)})) - f(\mathbf{x}^*)$. The goal is to minimize the number of iterations required to find an ε -optimal point and to minimize the total cost (regret) of the guesses over many rounds. Our main result is an efficient randomized algorithm for several natural choices of the function f and set $\mathcal{X}$ that incurs cost $O(d)$ and iteration complexity $O(d \log(1/\varepsilon)^2)$. Moreover, our dependence on d is asymptotically optimal, as we show examples in which any randomized algorithm for this problem must incur $\Omega(d)$ cost and iteration complexity.

1. Introduction

A growing body of literature studies learning with preference-based feedback [BV06; SJ11], with tremendous empirical success in recommendation systems, search engine optimization, information retrieval, and robotics. More recently, preference-based feedback has received a lot of attention as a mechanism to train large language models [OWJ⁺22]. Moreover, in recommender systems [BOHG13], a natural approach is to learn from users’ preferences relations on a set of recommended items and update the system’s belief for better future recommendations [JRTZ16] (e.g., given these items, which one do you prefer the most?).

Such preference-based feedback is not readily addressed by classical formulations for online decision making, such as bandits and reinforcement learning. In particular, algorithms for these problems rely on ordinal feedback per item (e.g., on a scale of 1 to 10, how much did the user like a particular item?). To address this, a long line of work studies the *dueling bandit framework* for online decision making under pairwise/preference-based feedback. There exist efficient algorithms with provable guarantees for the standard multi-armed bandit setup [YBKJ12; AKJ14; KHK15], contextual bandits [DHSSZ15; SK22], as well as dueling convex optimization [JNR12; SKM21; SKM22], to name a few. The dueling bandit framework is especially applicable in settings where real-valued feedback is scarce or impossible to obtain, but preference-based feedback is readily available.

However, a key limitation of the dueling bandit framework is that the feedback that the learner receives is essentially “in-list”. That is, the users are restricted to selecting items exclusively from the list of recommended items. This feedback model fails to capture the real-world scenarios where the users might select an out-of-list item they prefer. To illustrate, music streaming services like Spotify create personalized playlists for users. Concretely, each song can be encoded as a feature vector $\mathbf{x} \in \mathbb{R}^d$, and the goal is to recommend the songs with the highest utility for a hidden, well-structured utility function of $\mathbf{x}$. However, the users can also search for and play the songs they have a stronger preference (i.e., higher utility) than all recommendations.

This out-of-list feedback model falls into a monotone adversarial framework (see the chapter by Feige [Fei21]). In such models, an adversary is only allowed to make “helpful” changes. For example, in a graph clustering problem, the adversary is only allowed to add edges within communities and delete edges that cross communities (see, e.g., the chapter by Moitra [Moi21]). In our setting, the adversary is only allowed to respond with an item that is at least as good as any recommended item. A clear adaptation of the dueling bandit framework to this new feedback type is not evident.

1.1. Problem statement

As our main conceptual contribution, we introduce a theoretical formulation for this setting that we call *dueling optimization with a monotone adversary*. As we will see, our formulation supports “out-of-list” feedback.

Problem 1 (Dueling optimization with a monotone adversary). *Let $\mathcal{X} \subseteq \mathbb{R}^d$ be a decision space, and let $f: \mathcal{X} \rightarrow \mathbb{R}$ be a cost function with an unknown global minimum $\mathbf{x}^*$. A learner interacts with an adversary over rounds $t = 1, 2, \dots$, where each round is of the following form.*

1. *The learner proposes two points $\mathbf{x}_t^{(1)}, \mathbf{x}_t^{(2)} \in \mathcal{X}$.*

2. The adversary responds with a point $\mathbf{r}_t$ that satisfies $f(\mathbf{r}_t) \leq \min \left\{ f(\mathbf{x}_t^{(1)}), f(\mathbf{x}_t^{(2)}) \right\}$.

The goal is to design algorithms that:

1. for some prespecified $\varepsilon > 0$, minimize the number of iterations to find a point $\mathbf{x}$ for which $f(\mathbf{x}) - f(\mathbf{x}^*) \leq \varepsilon$;
2. minimize the total cost $\sum_{t=1}^{\infty} \left(\max \left\{ f(\mathbf{x}_t^{(1)}), f(\mathbf{x}_t^{(2)}) \right\} - f(\mathbf{x}^*) \right)$.

Note that in Problem 1, we are interested in both the iteration complexity and the total cost. The first objective is a standard metric for measuring the performance of an iterative optimization algorithm. The second objective is motivated by online settings in which a practitioner may wish to minimize the total regret (cost) of its recommendations over an indefinitely long interaction with a user. In fact, the algorithms we propose in this paper simultaneously achieve both small iteration complexity (for any choice of ε) as well as total cost — see our technical overview in Section 1.3 for more details.

Problem 1 is a natural extension of (noiseless) dueling optimization [JNR12; SKM21; SKM22] to handle “out-of-list” responses, as in the Spotify recommendation example. The vanilla (noiseless) dueling optimization setup corresponds to the requirement that the user’s response satisfies $\mathbf{r}_t \in \{\mathbf{x}_t^{(1)}, \mathbf{x}_t^{(2)}\}$. We allow the user to be potentially adversarial by allowing it to respond with any improvement to the learner’s suggestions (in the sequel, we exclusively refer to the user as the adversary).

Even though the monotone adversary is only improving upon the learner’s suggestions, existing algorithms for dueling optimization cannot be freely extended to handle the monotone feedback. At a high level, the difficulty arises from the fact that existing algorithms carefully select the queries $\mathbf{x}_t^{(1)}, \mathbf{x}_t^{(2)}$ so that learning whether $f(\mathbf{x}_t^{(1)}) > f(\mathbf{x}_t^{(2)})$ reveals information about the underlying f . However, a monotone adversary can return a point $\mathbf{r}_t$ that reveals no information about the relationship between $\mathbf{x}_t^{(1)}$ and $\mathbf{x}_t^{(2)}$.

To illustrate this point, consider a natural coordinate-wise binary search algorithm for the dueling optimization problem when $f(\mathbf{x}) = \|\mathbf{x} - \mathbf{x}^*\|_2^2$ for some $\mathbf{x}^* \in \mathcal{B}_2^d := \{\mathbf{x} : \|\mathbf{x}\|_2 \leq 1\}$. For coordinates $i = 1, \dots, d$, query points of the form $\mathbf{x}_t^{(1)} = c_1 \cdot \mathbf{e}_i$, $\mathbf{x}_t^{(2)} = c_2 \cdot \mathbf{e}_i$ and progressively refine the values $c_1, c_2 \in \mathbb{R}$ to search for the value of $\mathbf{x}^*[i]$ (i.e., the i -th entry of $\mathbf{x}^*$). It is easy to show that this approach has a query complexity of $O(d \log(1/\varepsilon))$ in the vanilla dueling optimization setting. However, a monotone adversary can return orthogonal responses of the form $\mathbf{r}_t = C\mathbf{e}_j$ (where $j \neq i$ and C is a constant) that do not allow the learner to search along the intended coordinate i . Furthermore, Jamieson, Nowak, and Recht [JNR12] and Saha, Koren, and Mansour [SKM21] give more sophisticated algorithms for the dueling optimization problem that inherently depend upon the “in-list” feedback, which clearly cannot apply to our setting. We therefore need novel insights to solve Problem 1.

1.2. Our results

We study Problem 1 for various natural classes of functions f and provide tight upper and lower bounds on the number of queries required to find an ε -optimal point.

Upper bound for linear functions. First, we study dueling optimization with a monotone adversary when the function f is linear. This is a natural class to consider. In particular, an algorithm that solves Problem 1 can be adapted to achieve constant regret for (noiseless) linear contextual bandits [CLRS11], where the reward function is $r(\mathbf{x}) := \langle \mathbf{x}, \mathbf{x}^* \rangle$. Note that the key difference in the setup is that the learner does not get to observe the actual linear costs but instead only an improvement to the actions (points) that the learner selects.

Theorem 1. *Let $\mathcal{X} = \mathbb{S}_2^d$, let $\mathbf{x}^*$ be such that $\|\mathbf{x}^*\|_2 = 1$, and let $f: \mathcal{X} \rightarrow \mathbb{R}$ be $f(\mathbf{x}) = -\langle \mathbf{x}, \mathbf{x}^* \rangle$. Fix any $\varepsilon > 0$. There exists an algorithm that, in the setting of Problem 1, with probability at least $1 - \exp(-O(d))$:*

- *outputs a point $\mathbf{x}$ satisfying $\langle \mathbf{x}^* - \mathbf{x}, \mathbf{x}^* \rangle \leq \varepsilon$ within $O(d \log(1/\varepsilon)^2)$ iterations;*
- *incurs total cost $O(d)$.*

Each pair of guesses at time t can be computed in $O(d)$ time.

We prove Theorem 1 in Section 2.2, and the cost is near-optimal with respect to d .

Gollapudi, Guruganesh, Kollias, Manurangsi, Leme, and Schneider [GGKMLS21] study a closely related setup that they call *local contextual recommendation*. Their result (see their Theorem 6.4) can be interpreted as showing that if the action set $\mathcal{X}$ is a discrete set (namely a packing over the unit sphere), there exists a $2^{\Omega(d)}$ lower bound on the iteration complexity to find a point with constant suboptimality. In contrast, our Theorem 1 shows a much smaller upper bound when the domain is the entire unit sphere.

Upper bound for smooth and PL functions. Next, we study whether we can show guarantees for a large class of functions. We show a positive result for functions that are both β -smooth and α -Polyak-Łojasiewicz (abbreviated as PL). These assumptions are standard in optimization.

Definition 1 (β -smooth function [Bub15, Lemma 3.4]). *We say f is β -smooth if it satisfies (1.1).*

$$\text{For all } \mathbf{x}, \mathbf{y} \in \mathbb{R}^d : \quad |f(\mathbf{x}) - f(\mathbf{y}) - \langle \nabla f(\mathbf{y}), \mathbf{x} - \mathbf{y} \rangle| \leq \frac{\beta}{2} \cdot \|\mathbf{x} - \mathbf{y}\|_2^2 \quad (1.1)$$

Definition 2 (α -PL function). *We say f is α -PL if it satisfies (1.2).*

$$\text{For all } \mathbf{x} \in \mathbb{R}^d \text{ and minimizers } \mathbf{x}^* : \quad f(\mathbf{x}) - f(\mathbf{x}^*) \leq \frac{1}{2\alpha} \|\nabla f(\mathbf{x})\|_2^2 \quad (1.2)$$

Our main result for this setting is Theorem 2.

Theorem 2. *Let $\mathcal{X} = \mathbb{R}^d$, and suppose f is β -smooth (Definition 1) and α -PL (Definition 2). Fix any $\varepsilon > 0$, as well as a known point $\mathbf{x}_1$ and a value B satisfying $B \geq f(\mathbf{x}_1) - f(\mathbf{x}^*)$. There exists an algorithm that, in the setting of Problem 1, with probability at least $1 - \exp(-O(d))$:*

- *outputs a point $\mathbf{x}$ satisfying $f(\mathbf{x}) - f(\mathbf{x}^*) \leq \varepsilon$ within $O\left(\frac{\beta}{\alpha} \cdot d \cdot \log(B/\varepsilon)^2\right)$ iterations;*
- *incurs total cost $O\left(\frac{\beta}{\alpha} \cdot B \cdot d\right)$.*

Each pair of guesses at time t can be computed in $O(d)$ time.

We prove Theorem 2 in Section 2.3.

As an application, we show a positive result when the loss function is the Euclidean distance, and the decision space $\mathcal{X} = \mathcal{B}_2^d$ is a unit ball:

Theorem 3. *Let $\mathcal{X} = \mathcal{B}_2^d$, let $\mathbf{x}^*$ be such that $\|\mathbf{x}^*\|_2 \leq 1$, and let $f: \mathcal{X} \rightarrow \mathbb{R}$ be $f(\mathbf{x}) = \|\mathbf{x} - \mathbf{x}^*\|_2$. Fix any $\varepsilon > 0$. There exists an algorithm that, in the setting of Problem 1, with probability at least $1 - \exp(-O(d))$:*

- *outputs a point $\mathbf{x}$ satisfying $\|\mathbf{x} - \mathbf{x}^*\|_2 \leq \varepsilon$ within $O\left(d \cdot \log(B/\varepsilon)^2\right)$ iterations;*
- *incurs total cost $O(d)$.*

Each pair of guesses at time t can be computed in $O(d)$ time.

We prove Theorem 3 in Section 2.4.

Note that unlike in Theorem 2, Theorem 3 applies to the setting where the algorithm must guess points belonging to a given constraint set $\mathcal{X}$. Hence, in the proof of Theorem 3, we have to be careful to ensure that the convergence argument still holds when we apply the algorithm for Theorem 2 along with a projection step. It is not clear that this argument holds by default for all f satisfying the conditions requested by Theorem 2. Furthermore, as will become evident, we really only require that $\mathcal{X}$ be any convex body (though we state the result with $\mathcal{X} = \mathcal{B}_2^d$ to emphasize the consistency with our following lower bounds).

Lower bounds. We also prove that the dependence on d in our results is tight. In particular, when f is either a linear function or the distance to the target (as in Theorem 3), then $\Omega(d)$ queries are necessary to identify $\mathbf{x}^*$. This will translate to a $\Omega(d)$ cost over an infinite number of rounds. In fact, our lower bound is valid when the adversary must return one of the two queried points, as in vanilla dueling optimization framework.

Our lower bound also covers a more general setting than that stated in Problem 1. Thus far, we have only discussed the setting where the algorithm can query only two points and is told the better of the two. In many practical instances, the algorithm can query m points and learn the point with the best objective value (we call this m -ary dueling optimization). Then, one may ask why we study only the $m = 2$ in this paper. In our construction, we prove that unless m is polynomial in d , we cannot decrease the total cost substantially below $\Omega(d)$. Thus, the most interesting case for constant m is when $m = 2$, and in this case, our result is optimal in d .

See Theorem 4 for a formal statement of our lower bound.

Theorem 4 (Lower bound, ℓ_2 distance). *Let $\mathcal{X} = \mathcal{B}_2^d$. For any randomized algorithm for m -ary dueling optimization, there exists a choice of minimizer $\mathbf{x}^* \in \mathcal{B}_2^d$ and function $f(\mathbf{x}) := \|\mathbf{x} - \mathbf{x}^*\|_2$ such that the algorithm must:*

- *perform $\Omega(d/\log m)$ iterations in expectation to find a point $\mathbf{x}$ for which $f(\mathbf{x}) - f(\mathbf{x}^*) \leq \varepsilon$.*
- *incur cost $\Omega(d/\log m)$ in expectation.*

Here, $\varepsilon > 0$ is an absolute numerical constant.

We prove Theorem 4 in Section 3. Using the same construction, we can also demonstrate that Theorem 1 is tight when $\mathcal{X}$ is the unit sphere.

Corollary 5 (Lower bound, linear f). *Let $\mathcal{X} = \mathbb{S}_2^d$. For any randomized algorithm for m -ary dueling optimization there exists a choice of minimizer $\mathbf{x}^* \in \mathbb{S}_2^d$ and function $f(\mathbf{x}) := -\langle \mathbf{x}, \mathbf{x}^* \rangle$ such that the same conclusions as in Theorem 4 hold.*

1.3. Technical overview

At a high level, our algorithms maintain a guess $\mathbf{x}_t$ for the optimal solution $\mathbf{x}^*$. They will update this guess over many interactions with the adversary.

A general recipe. We first describe the primitives that our methods depend on. Our first technical innovation is the notion of *progress distributions*. Loosely speaking, these are distributions from which a learner is likely to sample a new guess $\mathbf{x}_{t+1}$ that decreases its suboptimality. See Definition 3.

Definition 3 (Progress Distribution). *Let $f: \mathcal{X} \rightarrow \mathbb{R}$ for $\mathcal{X} \subseteq \mathbb{R}^d$. For $\mathbf{x} \in \mathcal{X}$ and $1 \leq p < 2$, we say a distribution $\mathcal{D}(\mathbf{x})$ over vectors in $\mathbb{R}^d$ is a (p, γ, ρ) -progress distribution for $\mathbf{x}$ if we have the below.*

$$\Pr_{\mathbf{x}^+ \sim \mathcal{D}(\mathbf{x})} \left[\frac{f(\mathbf{x}) - f(\mathbf{x}^+)}{(f(\mathbf{x}) - f(\mathbf{x}^*))^p} \geq \frac{\rho}{d} \right] \geq \gamma.$$

So, if for every $\mathbf{x}_t$ the learner had sample access to some progress distribution $\mathcal{D}(\mathbf{x}_t)$, the learner can significantly improve its solution (e.g. when $p = 1$, roughly $\sim d/\rho$ steps are sufficient for the learner to decrease its suboptimality by a constant factor). It is therefore natural that repeating such a sample-then-guess approach ad infinitum will yield an approximately optimal solution. In Theorem 8, we prove this whenever there exist families of progress distributions for every range of possible suboptimality. Thus, assuming the learner can maintain a (possibly quite pessimistic) estimate of its suboptimality over all the rounds, we obtain a template for proving the iteration complexities of Theorems 1, 2, and 3. Note that ρ can be an arbitrarily small positive constant; even if there is a slim chance of decreasing the suboptimality, this is still sufficient because the monotone adversary ensures that the algorithm can never make negative progress.

Specifying progress distributions. We now discuss how we instantiate the above template for the β -smooth (Definition 1) and α -PL (Definition 2) case (Theorem 2). We focus on Theorem 2 for the sake of brevity; the proofs of Theorems 1 and 3 require some additional care but at a high level follow a similar structure. At step t , the algorithm maintains a guess $\mathbf{x}_t$ for the target $\mathbf{x}^*$. It chooses some step size ε_t and a random vector $\mathbf{g}_t$ from $\varepsilon_t \cdot \text{Unif}(\mathbb{S}_2^{d-1})$, where $\mathbb{S}_2^{d-1} := \{\mathbf{x} \in \mathbb{R}^d : \|\mathbf{x}\|_2 = 1\}$. We then query $\mathbf{x}_t$ and $\mathbf{x}_t - \mathbf{g}_t$. The key observation is that with a constant probability, the angle between $\mathbf{g}_t$ and the gradient $\nabla f(\mathbf{x})$ is small. We will use this to show that the distribution $\mathbf{x}_t - \varepsilon_t \cdot \text{Unif}(\mathbb{S}_2^{d-1})$ is a $(1, C_1, C_2)$ -progress distribution (Definition 3) for constants C_1, C_2 . Intuitively, this means that $\mathbf{x}_t - \mathbf{g}_t$ almost behaves like a step of gradient descent. To turn this observation into an algorithm, we need two main insights.

Step size schedule. The principal difficulty of this approach is to choose the step size ε_t . It is not immediately obvious how to do so since the algorithm does not observe any actual gradients

or function values. Hence, if our step sizes are too large, the algorithm may overshoot the optimal solution $\mathbf{x}^*$ and therefore not actually improve the quality of its current solution $\mathbf{x}_t$. On the other hand, if our step sizes are too small, the algorithm may not make enough progress in each step, which undesirably increases both the iteration complexity and the total cost.

To address this, we carefully construct a step size schedule that relies on a pessimistic upper bound on the suboptimality of the algorithm’s current solution. With this schedule, we show that in every step, one of two things happens – either the step size ε_t is small enough such that there is the possibility of the algorithm decreasing the cost, or it is too large. For the first case, we use β -smoothness (Definition 1) to prove that there is a constant probability that the algorithm finds a descent direction, which decreases the cost of its current solution substantially. For the second case, we use the α -PL condition (Definition 2) to prove that the cost the algorithm incurs in such steps is low. After enough steps, we can show that either the second case always holds (i.e. that the suboptimality is already desirably small) or the maximum cost that the algorithm can pay per round is small. We then decrease the step size ε_t by a constant factor, update the suboptimality estimate accordingly, and infinitely recurse.

Bounding the failure probability over infinite rounds. It now remains to show that the probability that the algorithm fails to make enough progress over *infinitely many rounds* is small. This is where the distinction between the two goals of Problem 1 becomes apparent. Specifically, even if we have a subroutine that, with high probability, outputs an ε -approximate solution, this does not immediately convert to an algorithm that can achieve bounded cost over an infinite number of rounds – note that the failure probabilities may accumulate in a divergent manner. Hence, we will require a more careful probabilistic analysis.

To overcome this challenge, we design the algorithm to run in phases $i = 1, 2, \dots$. In phase i , we use a step size ε_t proportional to $2^{-i/2}$ and run phase i for $\sim id$ steps. Using the fact that the family of distributions we are using for sampling next steps are $(1, C_1, C_2)$ -progress distributions, it will be enough to prove that $\sim d \cdot \beta/\alpha$ steps yield enough improving steps to decrease the suboptimality by a constant factor. We can therefore apply a Chernoff bound to conclude that the probability that the algorithm fails to make enough progress in phase i is at most $\exp(-id \cdot \beta/\alpha)$. Finally, we apply a union bound that the total probability of failure by $\exp(-d \cdot \beta/\alpha) \leq \exp(-d)$.

To bound the total cost over all phases $i \in \mathbb{N}_{\geq 1}$, we note that the sum of the suboptimality in each round is of the form $d \sum_{i \geq 1} i 2^{-i} = O(d)$. The guarantee on the iteration complexity follows by noting that to achieve a suboptimality of 2^{-i} , the algorithm runs $d \sum_{j \leq i} j = O(i^2 \cdot d)$ iterations.

1.4. Related works

Dueling convex optimization. As already mentioned, our formulation in Problem 1 is an extension of dueling convex optimization in the noiseless setting [JNR12; SKM21; SKM22]. Jamieson, Nowak, and Recht [JNR12] employ a coordinate-descent algorithm to show for α -smooth and β -strongly convex f , $\tilde{O}(d\beta/\alpha \log(1/\varepsilon))$ queries suffice to learn an ε -optimal point. As mentioned earlier, it is not clear how to adapt their algorithm to handle monotone feedback. In addition, the papers [SKM21; SKM22] show results for more general classes of f and in the presence of noise (where the adversary can return invalid response with nonzero probability); see Section 1.5 for a discussion on extending our results to noisy settings. However, their algorithms explicitly rely on

sign feedback $f(\mathbf{x}_t^{(1)}) \stackrel{?}{>} f(\mathbf{x}_t^{(2)})$ to construct gradient estimators, which are not possible in the monotone adversary setting.

Monotone adversaries. Our setting is an example of learning with a monotone adversary, where an adversary can choose to improve the feedback or information the algorithm gets. A common characteristic is that the improved information may paradoxically break or harm the performance of a given algorithm that works with non-improved information. Monotone adversaries are often studied in the semi-random model literature [BS95; Fei21; Moi21] for statistical estimation problems [CG18; Moi21; KLLST22] as well as learning problems, i.e., linear classification with Massart noise [MN06; DGT19].

Preference-based feedback. Our formulation in this paper falls within the growing body of literature that tackles learning with preference-based feedback, where the algorithm does not learn *how good* its options were in an absolute sense, just which one(s) were better than others. Other natural problems with preference-based feedback are contextual search [LS18; LLV18; LLS20], contextual recommendation (also called contextual inverse optimization) [BFL21; GGKMLS21], and 1-bit matrix completion [DPVW14].

1.5. Future work

We now discuss two interesting directions our work leaves open.

Dueling optimization with a monotone adversary under noise models. The most pressing next step is to determine noise models under which we can either obtain algorithmic results or hardness for solving Problem 1. There are several natural noise models that one could study. As a first step, one can consider the most analogous extension of the noise model studied by Jamieson, Nowak, and Recht [JNR12] and Saha, Koren, and Mansour [SKM21]. In the simplest form, they study a noise model where with probability $1/2 + \nu$ for some parameter $0 < \nu \leq 1/2$, the adversary returns $\arg\min_{i \in \{1,2\}} f(\mathbf{x}_t^{(i)})$, and with probability $1/2 - \nu$, the adversary returns $\arg\max_{i \in \{1,2\}} f(\mathbf{x}_t^{(i)})$. Note that this is a straightforward noise model to handle if the adversary must return one of $\mathbf{x}_t^{(1)}$ or $\mathbf{x}_t^{(2)}$. The algorithm simply queries the same pair of points roughly ν^{-2} times, which by Hoeffding’s inequality is enough to determine the index i that corresponded to the better of the two guesses. Additionally, this strategy can be implemented even without knowledge of the noise parameter ν ; see Section 6 of [SKM21] for more details.

The natural extension of this noise model to our monotone feedback setting is as follows. With probability $1/2 + \nu$, the adversary returns a response $\mathbf{r}_t$ that satisfies $f(\mathbf{r}_t) \leq \min\{f(\mathbf{x}_t^{(1)}), f(\mathbf{x}_t^{(2)})\}$. On the other hand, with probability $1/2 - \nu$, the adversary returns an arbitrary point in $\mathcal{X}$. It is immediate that there is no analogue of the “majority vote” strategy in this setting. We therefore believe that entirely new algorithmic ideas will be needed to address this noise model. On the other hand, it would be very interesting to show an impossibility result for this noise model. An impossibility result would imply that the monotone feedback makes the problem provably harder than the vanilla dueling optimization setting.

Another noise model of interest is one where the adversary’s valid monotone feedback is perturbed by a random variable Δ where $\mathbb{E}[\Delta] = 0$ and $\text{Cov}[\Delta] \preceq \sigma^2 \mathbf{I}_d$. To our knowledge, this noise model has not been considered in past works dealing with optimization with preference-based feedback,

even in the vanilla dueling setting (particularly the works of Jamieson, Nowak, and Recht [JNR12] and Saha, Koren, and Mansour [SKM21]).

Dueling optimization for other function classes. An orthogonal thread would be to identify other function classes and feasible regions $\mathcal{X}$ for which we can build algorithms for Problem 1. For instance, Saha, Koren, and Mansour [SKM21] obtain results for functions that are just β -smooth (and not necessarily α -PL or α -strongly convex). Our Theorem 2 uses the α -PL condition to prove that steps that do not make much progress also do not incur much cost. Hence, to remove the α -PL assumption, one would need to either avoid this line of reasoning or find another way to argue that the costs in low-progress rounds are not too high.

2. Proofs of upper bound results

In this section, we prove Theorem 8 (in which we construct and analyze a meta-algorithm for Problem 1 when the algorithm can sample next steps from progress distributions (Definition 3)). We then show how to use this framework to prove Theorem 1 (results for $f(\mathbf{x}) = \langle -\mathbf{x}, \mathbf{x}^* \rangle$), Theorem 2 (results for $f(\mathbf{x})$ being β -smooth and α -PL), and Theorem 3 (results for $f(\mathbf{x}) = \|\mathbf{x} - \mathbf{x}^*\|_2$), in that order. It will be helpful to recall the overview from Section 1.3 throughout this section.

We prove Theorem 8 in Section 2.1, Theorem 1 in Section 2.2, Theorem 2 in Section 2.3, and Theorem 3 in Section 2.4.

Before we jump into the main proofs, we will need a couple straightforward numerical inequalities.

Lemma 6. *For $r \in (0, 1)$ and $1 \leq p < 2$, we have $\sum_{i \geq 0} i \cdot r^{(1-p/2)i} \leq \frac{r^{p/2+1}}{(r - r^{p/2})^2}$.*

Proof of Lemma 6. Recall that

$$\sum_{i \geq 0} r^{(1-p/2)i} = \frac{1}{1 - r^{1-p/2}}.$$

Taking the derivative of both sides with respect to r yields

$$\sum_{i \geq 0} (1 - p/2) i \cdot r^{(1-p/2)i-1} = \frac{(2-p)r^{p/2}}{2(r - r^{p/2})^2}.$$

We multiply both sides by r and divide both sides by $1 - p/2$; we conclude that

$$\sum_{i \geq 0} i \cdot r^{(1-p/2)i} = \frac{r^{p/2+1}}{(r - r^{p/2})^2}$$

which recovers the statement of Lemma 6. □

Lemma 7 (Inner product with a random vector). *Let $\mathbf{g} \sim \text{Unif}(\mathbb{S}_2^{d-1})$ and let $\mathbf{y} \in \mathbb{S}_2^{d-1}$ be fixed. Then*

$$\Pr_{\mathbf{g}} \left[\langle \mathbf{g}, \mathbf{y} \rangle \geq \frac{1}{2\sqrt{d}} \right] \geq \frac{1}{8}.$$

Proof. By rotational invariance, without loss of generality, we can let $\mathbf{y} = \mathbf{e}_1$. We apply Lemma 2.2 (a) due to Dasgupta and Gupta [DG03] with $\beta = 1/4$ to conclude that

$$\Pr_{\mathbf{g}} \left[\mathbf{g}_1^2 \leq \frac{1}{4d} \right] \leq \exp \left(\frac{1}{2} \left(1 - \frac{1}{4} + \ln \left(\frac{1}{4} \right) \right) \right) < \frac{3}{4}$$

which means that $\Pr_{\mathbf{g}} \left[|\langle \mathbf{g}, \mathbf{y} \rangle| \geq \frac{1}{2\sqrt{d}} \right] \geq \frac{1}{4}$. The result of Lemma 7 follows by symmetry. $\square$

2.1. A general algorithm for Problem 1 with progress distributions

The goal of this subsection is to develop the general tools we need to prove our main results.

The key primitive of our analysis is a general algorithm (Algorithm 1) that solves Problem 1 when we are given certain convenient distributions from which we sample new guesses. We call these *progress distributions*; recall Definition 3.

Let us describe Algorithm 1. In each step, Algorithm 1 maintains a current guess $\mathbf{x}_t$ and chooses a slight perturbation of that guess $\mathbf{x}_t^+ \sim \mathcal{D}(\mathbf{x}_t)$, where $\mathcal{D}(\mathbf{x}_t)$ is a (p, γ, ρ) -progress distribution (Definition 3). Algorithm 1 then submits the pair of guesses $\{\mathbf{x}_t, \mathbf{x}_t^+\}$. To analyze Algorithm 1, the main observation is that with probability $\geq \gamma$, the point $\mathbf{x}_t^+$ substantially improves over the cost of $\mathbf{x}_t$ – this follows directly from Definition 3. We exploit this intuition to give our most general result (Theorem 8) and to prove the correctness of Algorithm 1.

Theorem 8. *Let $f: \mathcal{X} \rightarrow \mathbb{R}$. Let B and $\mathbf{x}_1 \in \mathcal{X}$ be such that $f(\mathbf{x}_1) - f(\mathbf{x}^*) \leq B$. For $C > 0$, constant $r \in (0, 0.99)$, and for all $i \in \mathbb{N}_{\geq 1}$, suppose there exists intervals of the form $C \cdot [r^{i+1}, r^i]$ such that their union covers the interval $[0, B]$.*

If there exists a (p, γ, ρ) -progress distribution $\mathcal{D}_i(\mathbf{x})$ whenever $f(\mathbf{x}) - f(\mathbf{x}^) \in C \cdot [r^{i+1}, r^i]$ for all $i \geq 1$ and where p, γ, ρ do not depend on $\mathbf{x}$ and i , then there is an algorithm (Algorithm 1) for Problem 1 that, with probability at least $1 - \exp \left(-O \left(\frac{d}{\rho B^{p-1}} \right) \right)$, incurs total cost*

$$O \left(\frac{B \log(1/r)}{B^{p-1} \gamma \rho \min \left\{ r^{p(p-1/2-p)}, (r - r^{p/2})^2 \right\}} \cdot d \right).$$

Additionally, Algorithm 1 finds a point $\mathbf{x}$ satisfying $f(\mathbf{x}) - f(\mathbf{x}^) \leq \varepsilon$ in*

$$O \left(\frac{1}{B^{p-1} \gamma \rho} \cdot d \cdot \log \left(\frac{B}{\varepsilon} \right)^2 \right)$$

iterations with at least the aforementioned probability.

The proof of Theorem 8 has two main parts. In the first part, we will prove that for each value of i (call the set of timesteps belonging to a particular value of i “phase i ”), the number of steps $T(i)$ is sufficient to ensure that the cost of the algorithm’s solution decays gracefully with sufficiently large probability. In the second part, we will prove that the total cost the algorithm pays over all phases $i \geq 1$ is $\sim B \cdot d/\gamma\rho$ as promised. Theorem 8 will easily follow by combining these facts.

We start with stating and proving Lemma 9.

Algorithm 1 General recipe algorithm for dueling convex optimization

- 1: **Input:** Interaction with a monotone adversary $\mathcal{M}$ as defined in Problem 1; initial point $\mathbf{x}_1$ and bound B satisfying $f(\mathbf{x}_1) - f(\mathbf{x}^*) \leq B$; values C and r for which there exist corresponding intervals and (p, γ, ρ) -progress distribution families $\mathcal{D}_i$ (see the statement of Theorem 8).
 - 2: Initialize $\mathbf{x}_1 = 0, t = 1$.
 - 3: **for** $i = 1, \dots$ **do**
 - 4: **for** $T(i) := 2i/(\gamma \min(1, \rho)) \cdot (Cr^{i+1})^{-(p-1)} \cdot \log(1/r) \cdot d$ iterations **do**
 - 5: Sample $\mathbf{x}_t^+$ from $\mathcal{D}_i(\mathbf{x}_t)$.
 - 6: **Submit guesses** $\{\mathbf{x}_t, \mathbf{x}_t^+\}$ and **receive response** $\mathbf{r}_t$.
 - 7: Let $\mathbf{x}_{t+1} = \mathbf{r}_t$.
 - 8: Update $t \leftarrow t + 1$.
 - 9: **end for**
 - 10: **end for**
-

Lemma 9. Let $i \geq 1$ and let $T(i)$ be defined below (or see Line 4 of Algorithm 1).

$$T(i) := \frac{2i}{\gamma \min(1, \rho) \cdot (Cr^{i+1})^{p-1}} \cdot \log(1/r) \cdot d.$$

Let t_i be the first iteration of phase i . If $f(\mathbf{x}_{t_i}) - f(\mathbf{x}^*) \leq Cr^i$, then, with probability $\geq 1 - r^{\frac{dB^{-(p-1)}}{4\rho} \cdot i}$, we have $f(\mathbf{x}_{t_i+T(i)+1}) - f(\mathbf{x}^*) \leq Cr^{i+1}$.

Proof of Lemma 9. Assume that we have $f(\mathbf{x}_{t_i}) - f(\mathbf{x}^*) \in C \cdot [r^{i+1}, r^i]$ (otherwise, we are done immediately).

Define the indicator random variable Y_t as follows.

$$Y_t := \mathbb{1} \left\{ \frac{f(\mathbf{x}_t) - f(\mathbf{x}_t^+)}{(f(\mathbf{x}_t) - f(\mathbf{x}^*))^p} \geq \frac{\rho}{d} \right\}.$$

Consider the distribution of guesses $\mathcal{D}_i$ (let us omit the argument $\mathbf{x}_t$ for the sake of brevity). Since $\mathcal{D}_i$ is a (p, γ, ρ) -progress distribution, we have

$$\Pr_{\mathbf{x}^+ \sim \mathcal{D}_i} \left[\frac{f(\mathbf{x}_t) - f(\mathbf{x}_t^+)}{f(\mathbf{x}_t) - f(\mathbf{x}^*)} \geq \frac{\rho}{d} \cdot (Cr^{i+1})^{p-1} \right] \geq \Pr_{\mathbf{x}^+ \sim \mathcal{D}_i} [Y_t = 1] \geq \gamma.$$

Call every step t for which $Y_t = 1$ a “successful step.” Let us give a high-probability count on the number of successful steps. Recall that a form of the Chernoff bound states that, for $\delta \in [0, 1]$ and independent indicator random variables Y_j ,

$$\Pr \left[\sum_{j=t_i}^{t_i+T(i)} Y_j \leq (1 - \delta) \mathbb{E} \left[\sum_{j=t_i}^{t_i+T(i)} Y_j \right] \right] \leq \exp \left(- \frac{\delta^2 \cdot \mathbb{E} \left[\sum_{j=t_i}^{t_i+T(i)} Y_j \right]}{2} \right).$$

Applying the Chernoff bound with $\delta = 1/2$ yields

$$\Pr \left[\sum_{j=t_i}^{t_i+T(i)} Y_j \leq \frac{T(i)\gamma}{2} \right] \leq \exp \left(- \frac{i \cdot \frac{d}{\rho(Cr^{i+1})^{p-1}} \cdot \log(1/r)}{4} \right) \leq r^{\frac{dB^{-(p-1)}}{4\rho} \cdot i}.$$

where we use $Cr^{i+1} \leq Cr^2 \leq B$.

It remains to show that after at least $T(i)\gamma/2$ successful steps, we have $f(\mathbf{x}_{t_i+T(i)+1}) - f(\mathbf{x}^*) \leq Cr^{i+1}$. Recall that we assume that $f(\mathbf{x}_{t_i}) - f(\mathbf{x}^*) \geq Cr^{i+1}$ and note that for every successful step, we have

$$\frac{f(\mathbf{x}_t) - f(\mathbf{x}_t^+)}{f(\mathbf{x}_t) - f(\mathbf{x}^*)} \geq \frac{\rho}{d} \cdot (Cr^{i+1})^{p-1}$$

which implies

$$\frac{f(\mathbf{x}_t^+) - f(\mathbf{x}^*)}{f(\mathbf{x}_t) - f(\mathbf{x}^*)} \leq 1 - \frac{\rho}{d} \cdot (Cr^{i+1})^{p-1}.$$

We multiply over all steps in phase i , giving

$$\begin{aligned} \frac{f(\mathbf{x}_{t_i+T(i)+1}) - f(\mathbf{x}^*)}{f(\mathbf{x}_{t_i}) - f(\mathbf{x}^*)} &= \prod_{t=t_i}^{t_i+T(i)} \frac{f(\mathbf{x}_{t+1}) - f(\mathbf{x}^*)}{f(\mathbf{x}_t) - f(\mathbf{x}^*)} \leq \left(1 - \frac{\rho}{d} \cdot (Cr^{i+1})^{p-1}\right)^{T(i)\gamma/2} \\ &\leq \left(1 - \frac{2i}{\gamma} \cdot \frac{\log(1/r)}{T(i)}\right)^{T(i)\gamma/2} \leq \exp\left(\frac{2i}{\gamma} \cdot \frac{\log(1/r)}{T(i)} \cdot \frac{T(i)\gamma}{2}\right) = r^i \leq r. \end{aligned}$$

Finally, recall that $f(\mathbf{x}_{t_i}) - f(\mathbf{x}^*) \leq Cr^i$. Combining this with the above gives $f(\mathbf{x}_{t_i+T(i)+1}) - f(\mathbf{x}^*) \leq Cr^{i+1}$, concluding the proof of Lemma 9. $\square$

Next, we have Lemma 10, which controls the total cost that Algorithm 1 incurs assuming that the cost is sufficiently low in each phase.

Lemma 10. *For a timestep t , let $i(t)$ be the phase that t belongs to.*

If for all t we have $f(\mathbf{x}_t) - f(\mathbf{x}^) \leq Cr^{i(t)}$, then Algorithm 1 incurs total cost*

$$O\left(\frac{B \log(1/r)}{C^{p-1}\gamma\rho \min\left\{r^{p(p-1/2-p)}, (r - r^{p/2})^2\right\}} \cdot d\right).$$

Proof of Lemma 10. Recall throughout this proof that $r \leq 0.99$ and p is a constant such that $p < 2$.

Observe that in phase i , the algorithm incurs cost at most

$$T(i) \cdot Cr^i = \frac{2i}{\gamma} \cdot \frac{dCr^i}{\rho(Cr^{i+1})^{p-1}} \cdot \log(1/r) = \frac{2i}{\gamma} \cdot \frac{d \log(1/r)}{\rho C^{p-2}} \cdot r^{i-(i+1)(p-1)}.$$

We will find a threshold i_p for which for all $i \geq i_p$, the above cost is exponentially decaying. This will allow us to control the sum of the costs over infinitely many rounds. We choose $i_p = 2 \cdot \lceil ((p-1)/(2-p)) \rceil$. Notice that for all $i \geq i_p$, the exponent on r can be bounded as

$$i - (i+1)(p-1) = i(2-p) - (p-1) \geq \left(1 - \frac{p}{2}\right)i.$$

Note that this also implies that $(i_p+1)(p-1) \leq p/2 \cdot i_p = p \lceil (p-1)/(2-p) \rceil$.

To total the cost, we consider two cases. First, suppose $1 \leq i \leq i_p - 1$. Observe that in each of these phases, we pay cost at most B , so we have

$$\begin{aligned} \sum_{i=1}^{i_p-1} B \cdot T(i) &\leq B (T_{i_p} \cdot i_p) = 2B \left(2 \cdot \frac{p-1}{2-p}\right)^2 \cdot \frac{\log(1/r)}{\gamma} \cdot \frac{d}{\rho C^{p-1} r^{(i_p+1)(p-1)}} \\ &\leq 2B \left(2 \cdot \frac{p-1}{2-p}\right)^2 \cdot \frac{\log(1/r)}{\gamma} \cdot \frac{d}{\rho C^{p-1} r^{p(p-1/2-p)}}. \end{aligned} \quad (2.1)$$

Next, we sum over all phases $i \geq i_p$. We obtain a cost that is at most

$$\sum_{i \geq i_p} \frac{2i}{\gamma} \cdot \frac{d \log(1/r)}{\rho C^{p-2}} \cdot r^{i-(i+1)(p-1)} \leq \frac{2d \log(1/r)}{\gamma \rho \cdot C^{p-2}} \sum_{i \geq i_p} i \cdot r^{(1-p/2)i} \leq \frac{2d \log(1/r)}{\gamma \rho \cdot C^{p-2}} \cdot \frac{r^{p/2+1}}{(r - r^{p/2})^2} \quad (2.2)$$

where the last inequality follows from Lemma 6. Combining (2.1) and (2.2) yields

$$\begin{aligned} &2B \left(2 \cdot \frac{p-1}{2-p}\right)^2 \cdot \frac{\log(1/r)}{\gamma} \cdot \frac{d}{\rho C^{p-1} r^{p(p-1/2-p)}} + \frac{2d \log(1/r)}{\gamma \rho \cdot C^{p-2}} \cdot \frac{r^{p/2+1}}{(r - r^{p/2})^2} \\ &\leq 2B \left(2 \cdot \frac{p-1}{2-p}\right)^2 \cdot \frac{\log(1/r)}{\gamma} \cdot \frac{d}{\rho C^{p-1} r^{p(p-1/2-p)}} + \frac{2d \log(1/r)}{\gamma \rho \cdot C^{p-1}} \cdot \frac{B}{(r - r^{p/2})^2} \\ &= O \left(\frac{B \log(1/r)}{C^{p-1} \gamma \rho \min \left\{ r^{p(p-1/2-p)}, (r - r^{p/2})^2 \right\}} \cdot d \right) \end{aligned}$$

This concludes the proof of Lemma 10. $\square$

We are now ready to prove Theorem 8.

Proof of Theorem 8. It is sufficient to prove that with probability $\geq 1 - \exp\left(-O\left(\frac{d}{\rho B^{p-1}}\right)\right)$, at the end of phase i , we have $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq C r^{i+1}$. Recall the conclusion of Lemma 9 and that $f(\mathbf{x}_1) - f(\mathbf{x}^*) \leq B \leq C r$; by a union bound, we have for all phases i that $f(\mathbf{x}_t) - f(\mathbf{x}^*) \leq C r^{i+1}$ with probability

$$1 - \sum_{i \geq 1} r^{\frac{dB^{-(p-1)}}{4\rho} \cdot i} \geq 1 - \exp\left(-O\left(\frac{d}{\rho B^{p-1}}\right)\right)$$

where we use $0 < r < 0.99$. The first part of Theorem 8 now follows directly from applying Lemma 10. The rest of the statement of Theorem 8 follows by noting that

$$\sum_{j \leq i} T(j) = \frac{2C^{-(p-1)} \log(1/r)}{\gamma \min(1, \rho)} \cdot d \cdot \sum_{j \leq i} j \left(r^{-(j+1)(p-1)} \right) \lesssim \frac{2C^{-(p-1)}}{\gamma \min(1, \rho)} \cdot d \cdot i^2.$$

where we again use $r < 0.99$. We set $\varepsilon = C r^i$ and conclude. $\square$

2.2. Proof of Theorem 1

The goal of this subsection is to prove Theorem 1.

Our plan will be to use the general guarantee of Theorem 8. Thus, the main task is to prove that there is an appropriate interval cover and corresponding sequence $\mathcal{D}_i(\mathbf{x})$ of progress distributions for all $\mathbf{x}$ belonging to phase i that satisfy the conditions of Theorem 8.

We prove this fact in Lemma 11. We remark that we made no effort to optimize the numerical constants; we choose the constants that appear in the Lemma statement to simplify calculations, as these will not impact our asymptotic results.

Lemma 11. *Let $f: \mathbb{R}^d \rightarrow \mathbb{R}$ be the negative inner product function defined on $\mathbb{S}_2^{d-1}$ with respect to some unknown target $\mathbf{x}^*$. Then for any $\mathbf{x}$ for which $f(\mathbf{x}) - f(\mathbf{x}^*) \in [10^{-(i+1)}, 10^{-i}]$ and for which $\langle \mathbf{x}^*, \mathbf{x} \rangle > 0$, there is a $(1.5, 10^{-1}, 10^{-4})$ -progress distribution (Definition 3) that can be computed in time $O(d)$.*

Proof of Lemma 11. We explain the construction of the distribution $\mathcal{D}(\mathbf{x})$.

Impose the following coordinates on $\mathbb{R}^d$. Let the first coordinate x_1 be the direction of $\mathbf{x}$, and the remaining $d-1$ coordinates be an arbitrary coordinate system for the perpendicular directions. Then, $\mathbf{x}$ has coordinates $(1, 0 \dots)$. Next, let $z := 10^{-(i+1)}$ and $s := \frac{z}{10\sqrt{d-1}}$. Let $\mathbf{s}$ be a point randomly drawn from a $d-1$ dimensional sphere of radius s whose coordinates are denoted $s_1 \dots s_{d-1}$. Then, the distribution $\mathcal{D}(\mathbf{x})$ is the distribution of $(\sqrt{1-s^2}, s_1 \dots s_{d-1})$. It is easy to verify that these points lie on $\mathbb{S}_2^{d-1}$.

This distribution can be computed in time $O(d)$. We will now show that it is a $(1.5, 10^{-1}, 10^{-4})$ progress distribution.

Let $z_0 := f(\mathbf{x}) - f(\mathbf{x}^*)$; recall that $z_0 \in [0, 1]$ and $z \leq z_0 \leq 10z$. We write the target vector $\mathbf{x}^* = (1 - z_0)\mathbf{x} + \sqrt{1 - (1 - z_0)^2}\mathbf{y} = (1 - z_0)\mathbf{x} + \sqrt{2z_0 - z_0^2}\mathbf{y}$ where $\mathbf{y}$ is a unit vector and $\langle \mathbf{y}, \mathbf{x} \rangle = 0$. Note that this expression holds because $\langle \mathbf{x}^*, \mathbf{x} \rangle = 1 - z_0$.

Let $\mathbf{x}^+$ be a random point chosen from $\mathcal{D}(\mathbf{x})$. Let $y_1 \dots y_{d-1}$ be the coordinates of $\mathbf{y}$ in the $d-1$ dimensional coordinate plane perpendicular to $\mathbf{x}$ defined above. We compute

$$\langle \mathbf{x}^*, \mathbf{x}^+ \rangle = \left(1 - z_0, \sqrt{2z_0 - z_0^2}y_1 \dots \sqrt{2z_0 - z_0^2}y_{d-1}\right) \cdot \left(\sqrt{1 - s^2}, s_1 \dots s_{d-1}\right) \quad (2.3)$$

$$= (1 - z_0)\sqrt{1 - s^2} + \left(\sqrt{2z_0 - z_0^2}y_1 \dots \sqrt{2z_0 - z_0^2}y_{d-1}\right) \cdot (s_1 \dots s_{d-1}). \quad (2.4)$$

By Lemma 7, we have (note that we weaken the constants from Lemma 7 for numerical convenience later in the proof)

$$\Pr_s \left[\left(\sqrt{2z_0 - z_0^2}y_1 \dots \sqrt{2z_0 - z_0^2}y_{d-1} \right) \cdot (s_1 \dots s_{d-1}) \geq \frac{0.1}{\sqrt{d-1}} \cdot \sqrt{2z_0 - z_0^2} \cdot s \right] \geq 0.1.$$

Because $2z_0 - z_0^2 \geq z_0 \geq z$, we have

$$\frac{0.1}{\sqrt{d-1}} \cdot \sqrt{2z_0 - z_0^2} \cdot s \geq \frac{s\sqrt{z}}{10\sqrt{d-1}}.$$

In turn, this shows

$$\Pr_s \left[\left(\sqrt{2z_0 - z_0^2} y_1 \dots \sqrt{2z_0 - z_0^2} y_{d-1} \right) \cdot (s_1 \dots s_{d-1}) \geq \frac{s\sqrt{z}}{10\sqrt{d-1}} \right] \geq 0.1.$$

Combining this with Equation 2.3, we obtain

$$\Pr_{\mathbf{x}^+ \sim \mathcal{D}_{\mathbf{x}}} \left[\langle \mathbf{x}^*, \mathbf{x}^+ \rangle \geq (1 - z_0)\sqrt{1 - s^2} + \frac{s\sqrt{z}}{10\sqrt{d-1}} \right] \geq 0.1. \quad (2.5)$$

Now, we will find a lower bound for $(1 - z_0)\sqrt{1 - s^2} + \frac{s\sqrt{z}}{10\sqrt{d-1}}$. We have

$$\begin{aligned} (1 - z_0)\sqrt{1 - s^2} + \frac{s\sqrt{z}}{10\sqrt{d-1}} &\geq (1 - z_0)(1 - s^2) + \frac{s\sqrt{z}}{10\sqrt{d-1}} \\ &= (1 - z_0)(-s^2) + \frac{s\sqrt{z}}{10\sqrt{d-1}} + (1 - z_0) \\ &\geq (1 - z)(-s^2) + \frac{s\sqrt{z}}{10\sqrt{d-1}} + (1 - z_0). \end{aligned}$$

Now, using that $s = \frac{z}{10\sqrt{d-1}}$, we get

$$\begin{aligned} (1 - z)(-s^2) + \frac{s\sqrt{z}}{10\sqrt{d-1}} + (1 - z_0) &= (1 - 10s\sqrt{d-1})(-s^2) + \frac{s^{3/2}}{\sqrt{10}(d-1)^{1/4}} + (1 - z_0) \\ &= -s^2 + 10s^3\sqrt{d-1} + \frac{s^{3/2}}{\sqrt{10}(d-1)^{1/4}} + (1 - z_0) \\ &\geq \left(10s^3\sqrt{d-1} + \frac{s^{3/2}}{8(d-1)^{1/4}} - s^2 \right) \\ &\quad + \frac{s^{3/2}}{6(d-1)^{1/4}} + (1 - z_0) \end{aligned}$$

where the last line follows from $1/\sqrt{10} > 1/8 + 1/6$. Finally, applying weighted AM-GM lets us see

$$10s^3\sqrt{d-1} + \frac{s^{3/2}}{8(d-1)^{1/4}} \geq \frac{3}{2^{2/3}} \left(10s^3\sqrt{d-1} \right)^{1/3} \left(\frac{s^{3/2}}{8(d-1)^{1/4}} \right)^{2/3} = \frac{3 \cdot 10^{1/3}}{2^{2/3} 2^2} s^2 > s^2$$

where we use a weight of $1/3$ on the first term and a weight of $2/3$ on the second term. We now write

$$(1 - z_0)\sqrt{1 - s^2} + \frac{s\sqrt{z}}{10\sqrt{d-1}} \geq \frac{s^{3/2}}{6(d-1)^{1/4}} + (1 - z_0).$$

Substituting s once again and recalling that $\langle \mathbf{x}^*, \mathbf{x} \rangle = (1 - z_0)$ and $z \geq \frac{z_0}{10} = \frac{\langle \mathbf{x}^*, \mathbf{x}^* - \mathbf{x} \rangle}{10}$, we get

$$(1 - z_0)\sqrt{1 - s^2} + \frac{s\sqrt{z}}{10\sqrt{d-1}} \geq \frac{z^{3/2}}{6 \cdot 10^{3/2} \cdot d} + \langle \mathbf{x}^*, \mathbf{x} \rangle > \frac{10^{-4} \langle \mathbf{x}^*, \mathbf{x}^* - \mathbf{x} \rangle^{3/2}}{d} + \langle \mathbf{x}^*, \mathbf{x} \rangle.$$

Combining this with (2.5), we now have

$$\Pr_{\mathbf{x}^+ \sim \mathcal{D}(\mathbf{x})} \left[\langle \mathbf{x}^*, \mathbf{x}^+ \rangle \geq \frac{10^{-4} \langle \mathbf{x}^*, \mathbf{x}^* - \mathbf{x} \rangle^{3/2}}{d} + \langle \mathbf{x}^*, \mathbf{x} \rangle \right] \geq 0.1$$

which means that

$$\Pr_{\mathbf{x}^+ \sim \mathcal{D}(\mathbf{x})} \left[\langle \mathbf{x}^*, \mathbf{x}^+ - \mathbf{x} \rangle \geq \frac{10^{-4} \langle \mathbf{x}^*, \mathbf{x}^* - \mathbf{x} \rangle^{3/2}}{d} \right] \geq 0.1.$$

This exactly aligns with the definition of a $(1.5, 10^{-1}, 10^{-4})$ progress distribution, completing the proof of Lemma 11. $\square$

We will now conclude Theorem 1 using Theorem 8.

Proof of Theorem 1. To apply Theorem 8, we need to present C , r , and a sequence of parameterizations $\mathcal{D}_i$ that satisfy the premises.

Set $r = 0.1$ and $C = 10$. By Lemma 11, we can find progress distributions for each interval $[Cr^i, Cr^{i+1}]$ of suboptimality of the current function value, since we can find such progress distributions as long as the suboptimality of the function is at most 1.

Note that the algorithm can begin with a point $\mathbf{x}$ where $f(\mathbf{x}) - f(\mathbf{x}^*) < 1$ by first querying two opposite points on a sphere; one can easily see that at least one of the two points queried satisfies $f(\mathbf{x}) - f(\mathbf{x}^*) < 1$.

We therefore conclude the proof of Theorem 1. $\square$

2.3. Proof of Theorem 2

The goal of this subsection is to prove Theorem 2.

As before, we use the general guarantee of Theorem 8 via proving that there is an appropriate interval cover and corresponding sequence $\mathcal{D}_i(\mathbf{x})$ of progress distributions for all $\mathbf{x}$ for which $f(\mathbf{x}) - f(\mathbf{x}^*) \in [Cr^{i+1}, Cr^i]$.

We prove this fact in Lemma 12.

Lemma 12. Fix $i \in \mathbb{N}_{\geq 1}$. Let $\varepsilon_i = \sqrt{2B\alpha/\beta^2} \cdot 1/\sqrt{d} \cdot 2^{-i/2-1}$. If f is β -smooth and α -PL, and if we have $f(\mathbf{x}) - f(\mathbf{x}^*) \in B \cdot [2^{-i}, 2^{-i+1}]$, then the distribution $\mathcal{D}_i(\mathbf{x}) = \mathbf{x} + \varepsilon_i \cdot \text{Unif}(\mathbb{S}_2^{d-1})$ is a $(1, \gamma, \rho)$ -progress distribution for $(\gamma, \rho) = (1/8, \alpha/8\beta)$.

Proof of Lemma 12. Let $\mathbf{g} := \mathbf{x} - \mathbf{x}^*$.

It is sufficient to consider the case where we have $\|\mathbf{g}\|_2 \leq \frac{1}{2\beta} \cdot \frac{\|\nabla f(\mathbf{x})\|_2}{\sqrt{d}}$. To see this, suppose this is not the case. We apply the PL inequality and write

$$f(\mathbf{x}) - f(\mathbf{x}^*) \leq \frac{1}{2\alpha} \|\nabla f(\mathbf{x})\|_2^2 \leq d \cdot \frac{2\beta^2}{\alpha} \|\mathbf{g}\|_2^2 = d \cdot \frac{2\beta^2}{\alpha} \left(\sqrt{\frac{2B\alpha}{\beta^2 d}} \cdot \frac{1}{2^{i/2+1}} \right)^2 = \frac{B}{2^i}$$

which implies that the suboptimality $f(\mathbf{x}) - f(\mathbf{x}^*)$ does not belong to the range we are considering.

Next, we use Lemma 7 to write the below.

$$\Pr_{\mathbf{g}} \left[\left\langle \frac{\nabla f(\mathbf{x})}{\|\nabla f(\mathbf{x})\|_2}, \frac{\mathbf{g}}{\|\mathbf{g}\|_2} \right\rangle \geq \frac{1}{2\sqrt{d}} \right] \geq \frac{1}{8}.$$

By Definition 1, we have for a β -smooth function and for any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$ that

$$|f(\mathbf{x}) - f(\mathbf{y}) - \langle \nabla f(\mathbf{y}), \mathbf{x} - \mathbf{y} \rangle| \leq \frac{\beta}{2} \cdot \|\mathbf{x} - \mathbf{y}\|_2^2,$$

from which it easily follows that

$$|f(\mathbf{x} - \mathbf{g}) - f(\mathbf{x}) + \langle \nabla f(\mathbf{x}), \mathbf{g} \rangle| \leq \frac{\beta}{2} \cdot \|\mathbf{g}\|_2^2.$$

The above rearranges to

$$\begin{aligned} f(\mathbf{x}) - f(\mathbf{x} - \mathbf{g}) &\geq \langle \nabla f(\mathbf{x}), \mathbf{g} \rangle - \frac{\beta}{2} \cdot \|\mathbf{g}\|_2^2 \\ &= \|\nabla f(\mathbf{x})\|_2 \cdot \|\mathbf{g}\|_2 \cdot \left(\left\langle \frac{\nabla f(\mathbf{x})}{\|\nabla f(\mathbf{x})\|_2}, \frac{\mathbf{g}}{\|\mathbf{g}\|_2} \right\rangle - \frac{\beta/2 \cdot \|\mathbf{g}\|_2}{\|\nabla f(\mathbf{x})\|_2} \right) \\ &\geq \frac{\beta}{2} \cdot \|\mathbf{g}\|_2^2 \quad \text{with probability } > 1/8 \end{aligned}$$

We therefore conclude that with probability $> \gamma := 1/8$,

$$f(\mathbf{x}) - f(\mathbf{x} - \mathbf{g}) \geq \frac{\beta}{2} \cdot \|\mathbf{g}\|_2^2 = \frac{\beta}{2} \cdot \left(\sqrt{\frac{2B\alpha}{\beta^2 d}} \cdot \frac{1}{2^{i/2+1}} \right)^2 = \frac{\alpha}{\beta} \cdot \frac{B}{d} \cdot \frac{1}{2^{i+2}}.$$

This means that

$$\frac{f(\mathbf{x}) - f(\mathbf{x} - \mathbf{g})}{f(\mathbf{x}) - f(\mathbf{x}^*)} \geq \frac{\alpha/\beta \cdot B/d \cdot 1/2^{i+2}}{B/2^{i-1}} = \frac{\alpha}{\beta} \cdot \frac{1}{8d}$$

which means we can take $\rho = \alpha/8\beta$. This concludes the proof of Lemma 12. $\square$

The proof of Theorem 2 follows very easily from Lemma 12.

Proof of Theorem 2. Our plan is to apply Theorem 8. To do so, we need to present C , r , and a sequence of $\mathcal{D}_i$ that satisfy the premise of Theorem 8. We will use the settings of these objects guaranteed by Lemma 12.

Let $C = 2B$ and $r = 1/2$. It is clear that the intervals given by Lemma 12 cover $[0, B]$, and so for every $i \geq 1$, there exists a corresponding $(1, 1/8, \alpha/8\beta)$ -progress distribution family $\mathcal{D}_i$. We now apply Lemma 12 along with Theorem 8 to conclude the proof of Theorem 2. $\square$

2.4. Proof of Theorem 3

In this subsection, we prove Theorem 3.

Again, we present an appropriate interval cover and corresponding sequence of progress distributions $\mathcal{D}_i(\mathbf{x})$ that satisfy the conditions of Theorem 8. See Lemma 13.

Lemma 13. *Fix $i \in \mathbb{N}_{\geq 1}$. Let $\varepsilon = 1/\sqrt{d} \cdot 2^{-i/2}$. If $f: \mathcal{B}_2^d \rightarrow \mathbb{R}$ is $f(\mathbf{x}) = \|\mathbf{x} - \mathbf{x}^*\|_2$ for $\mathbf{x}^* \in \mathcal{B}_2^d$ and if $\|\mathbf{x} - \mathbf{x}^*\|_2 \leq \sqrt{2} \cdot [2^{-(i+1)/2}, 2^{-i/2}]$, then there exists a distribution $\mathcal{D}(\mathbf{x})$ that can be efficiently sampled from and is a $(1, \gamma, \rho)$ -progress distribution for $(\gamma, \rho) = (1/8, 1/8)$.*

Proof of Lemma 13. Let $\mathbf{x}^+$ have distribution

$$\frac{\mathbf{x} - \mathbf{g}}{\max\{1, \|\mathbf{x} - \mathbf{g}\|_2\}}, \quad \text{where } \mathbf{g} \sim \varepsilon_t \cdot \text{Unif}(\mathbb{S}_2^{d-1}). \quad (2.6)$$

Note that this distribution can be described as, “add a uniformly random direction of length ε_t to $\mathbf{x}$ and project the result back onto $\mathcal{X} = \mathcal{B}_2^d$.”

It is easy to see that $\|\mathbf{x}^+\|_2 \leq 1$, so the iterates of Algorithm 1 will always remain inside $\mathcal{B}_2^d$. We now prove that $\mathcal{D}$ as described above in fact is a $(1, \gamma, \rho)$ -progress distribution for the promised parameters.

First, use the fact that $\|\mathbf{x} - \mathbf{x}^*\|_2^2$ is 2-smooth and 2-PL along with Lemma 12 to conclude that

$$\Pr_{\mathbf{g}} \left[\frac{\|\mathbf{x} - \mathbf{x}^*\|_2^2 - \|\mathbf{x} - \mathbf{g} - \mathbf{x}^*\|_2^2}{\|\mathbf{x} - \mathbf{x}^*\|_2^2} \geq \frac{2\rho}{d} \right] \geq \gamma.$$

Condition on this event. A basic property of the Euclidean projection onto a convex set implies that

$$\|\mathbf{x}^+ - \mathbf{x}^*\|_2^2 \leq \|(\mathbf{x} - \mathbf{g}) - \mathbf{x}^*\|_2^2$$

which yields

$$\Pr_{\mathbf{g}} \left[\frac{\|\mathbf{x} - \mathbf{x}^*\|_2^2 - \|\mathbf{x}^+ - \mathbf{x}^*\|_2^2}{\|\mathbf{x} - \mathbf{x}^*\|_2^2} \geq \frac{2\rho}{d} \right] \geq \gamma.$$

Finally, observe that the above event implies

$$\left(\frac{\|\mathbf{x}^+ - \mathbf{x}^*\|_2}{\|\mathbf{x} - \mathbf{x}^*\|_2} \right)^2 \leq \left(\sqrt{1 - \frac{2\rho}{d}} \right)^2 \leq \left(1 - \frac{\rho}{d} \right)^2.$$

Taking the square root of both sides and rearranging concludes the proof of Lemma 13. $\square$

We remark that the above proof goes through if $\mathcal{X}$ is an arbitrary convex set; we simply replace (2.6) with $\Pi_{\mathcal{X}}(\mathbf{x} - \mathbf{g})$, where $\Pi_{\mathcal{X}}(\mathbf{z})$ is the Euclidean projection of $\mathbf{z}$ onto $\mathcal{X}$.

Now, the proof of Theorem 3 will follow in a very similar manner to that of Theorem 2.

Proof of Theorem 3. To apply Theorem 8, we need to present C , r , and a sequence of distribution parameterizations $\mathcal{D}_i$ that satisfy the premises.

Let $\mathbf{x}_1 = 0$. It is clear that $\|\mathbf{x}^*\|_2 \leq 1 = B$, which means that the intervals of the form $\sqrt{2} \cdot [2^{-(i+1)/2}, 2^{-i/2}]$ for $i \geq 1$ cover the interval $[0, 1]$. Hence, for every $i \geq 1$, there exists a corresponding $(1, 1/8, 1/8)$ -progress distribution. Theorem 3 follows immediately. $\square$

3. Proofs of lower bound results

In this section, we will prove Theorem 4 and Corollary 5. We first state the following well-known fact (see, e.g., [Ver18]) that there exist $2^{\Omega(d)}$ points inside the unit ℓ_2 ball which are sufficiently far apart from one another.

Fact 14. *There exists a subset $S \subset \mathcal{B}_2^d$ such that $|S| = 2^{\Omega(d)}$, and for all $\mathbf{x}, \mathbf{y} \in S$ such that $\mathbf{x} \neq \mathbf{y}$, we have $\|\mathbf{x} - \mathbf{y}\|_2 \geq 0.1$.*

We are now ready to prove Theorem 4.

Proof of Theorem 4. We actually prove the lower bound even when the adversary must return the item *in the list* with smallest function value (breaking ties consistently, e.g., according to lexicographic order). Since the adversary is only weaker in this case, this implies the lower bound for the monotone adversary.

By Yao's Lemma [Yao77], it suffices to give a distribution over instances such that every deterministic algorithm satisfies the conclusions of the theorem. Hence, choose S from Fact 14 and let $\mathbf{x}^*$ be sampled uniformly from S .

Fix any deterministic algorithm. The deterministic algorithm branches into at most m states every round, depending on the response the adversary gives. Therefore after $r := \lfloor \log_m |S| \rfloor - 1$ rounds, the algorithm has at most $m^r < \frac{1}{2}|S|$ distinct states. Each of these states Q can be represented as a tuple of the form $\{(\mathbf{x}_t^{(1)}, \dots, \mathbf{x}_t^{(m)}, i_t)\}_{t \in [r]}$, where the $\mathbf{x}_t^{(i)} \in \mathcal{X}$ and the $i_t \in [m]$, which represents a set of the algorithm's guesses as well as the closest-point responses for the first r rounds.

Cost lower bound. Let us denote $c_r(\mathbf{x}^*, Q)$ to be the total cost incurred for the state Q if the target is $\mathbf{x}^* \in S$. We claim that all but at most one $\mathbf{x}^* \in S$ have $c_r(\mathbf{x}^*, Q) > 0.05r$. Suppose there were two points $\mathbf{x}^*$ and $\mathbf{x}^{*'}$ which had $c_r(\mathbf{x}^*, Q) < 0.05r$. Then $c_r(\mathbf{x}^*, Q) + c_r(\mathbf{x}^{*'}, Q) < 0.1r$, so there exists some round $t \in [r]$ for which

$$\max \left\{ \|\mathbf{x}_t^{(1)} - \mathbf{x}^*\|, \dots, \|\mathbf{x}_t^{(m)} - \mathbf{x}^*\| \right\} + \max \left\{ \|\mathbf{x}_t^{(1)} - \mathbf{x}^{*'}\|, \dots, \|\mathbf{x}_t^{(m)} - \mathbf{x}^{*'}\| \right\} < 0.1.$$

However, this cannot hold by triangle inequality since $\mathbf{x}^*$ and $\mathbf{x}^{*'}$ are well-separated.

For any target $\mathbf{x}^*$, the cost paid in the first r steps is at least $c_r(\mathbf{x}^*, Q(\mathbf{x}^*))$, where $Q(\mathbf{x}^*)$ is the state of the algorithm after r rounds when the target is $\mathbf{x}^*$. In particular, it is of the form $c_r(\mathbf{x}^*, Q)$ for some Q . Since there are only $\frac{1}{2}|S|$ possible algorithm states, at most $\frac{1}{2}|S|$ values of $\mathbf{x}^*$ can have total cost less than $0.05r$. Therefore, the average cost over instances uniformly drawn from S must be at least

$$\frac{1}{|S|} \left(|S| - \frac{1}{2}|S| \right) \cdot 0.05r \geq 0.025 \left(\frac{\log S}{\log m} - 2 \right) = \Omega(d/\log m).$$

Iteration lower bound. We will use the cost lower bound to prove the iteration lower bound. Recall that we proved that for any algorithm, there existed an instance $\mathbf{x}^*$ for which the algorithm incurs $\Omega(d/\log m)$ cost over the first r rounds.

Now suppose we had an algorithm $\mathcal{A}$ which achieved an expected iteration complexity of finding an ε -optimal point of $C \cdot d/\log m$ for any $\mathbf{x}^* \in S$, where $\varepsilon, C > 0$ are sufficiently small numerical constants. We can convert this into a low-cost algorithm $\mathcal{A}'$ for the first r rounds that (1) runs $\mathcal{A}$ to find an ε -optimal point $\mathbf{x}$; then (2) until round r repeatedly suggests $\mathbf{x}_t^{(1)} = \dots = \mathbf{x}_t^{(m)} = \mathbf{x}$. The expected cost of algorithm $\mathcal{A}'$ for the first r rounds is at most

$$2 \cdot \frac{Cd}{\log m} + \varepsilon \cdot r \leq (2C + \varepsilon) \frac{d}{\log m}.$$

For sufficiently small ε and C , we have a contradiction with the previous cost lower bound; thus we can conclude that any algorithm must perform $\Omega(d/\log m)$ iterations in expectation to find an ε -optimal point $\mathbf{x}$.

This concludes the proof of Theorem 4. □

Proof of Corollary 5. The argument for linear f is a reprise of the lower bound for ℓ_2 distance. Observe that Fact 14 implies that the points in S also satisfy $\langle \mathbf{x}, \mathbf{y} \rangle \leq 0.995$ for any $\mathbf{x} \neq \mathbf{y}$.

Therefore, we again use Yao's Lemma and consider deterministic algorithms that branch into m states in every round. Letting $c_r(\mathbf{x}^*, Q)$ denote the total cost incurred for state Q if the target is $\mathbf{x}^*$, we again have the claim that all but at most one $\mathbf{x}^* \in S$ have $c_r(\mathbf{x}^*, Q) > C \cdot r$ for some constant $C > 0$, from which it follows that at most $\frac{1}{2}|S|$ values of $\mathbf{x}^*$ can have total cost less than $C \cdot r$. We conclude that the average cost over instances drawn uniformly from S must be $\Omega(d/\log m)$.

The argument for the iteration lower bound also proceeds similarly, so we omit the details.

This concludes the proof of Corollary 5. □

Acknowledgements

AB is supported by NSF Awards CCF-2212968 and ECCS-2216899 and by the Defense Advanced Research Projects Agency under cooperative agreement HR00112020003. MG is supported by NSF Graduate Research Fellowship. GL is supported by the Institute for Data, Econometrics, Algorithms, and Learning (IDEAL). NSM is supported by NSF Graduate Research Fellowship and NSF Award ECCS-2216899. YY is supported by NSF Award CCF-2045402 and NSF Award CCF-2019844. YY thanks her advisor Jamie Morgenstern for her continued support and encouragement. We thank Aditya Bhaskara for suggesting that our algorithmic guarantees hold for smooth functions.

References

- [AKJ14] Nir Ailon, Zohar Karnin, and Thorsten Joachims. Reducing dueling bandits to cardinal bandits. In *International Conference on Machine Learning*, pages 856–864. PMLR, 2014 (cited on page 2).
- [BV06] Eyal Beigman and Rakesh Vohra. Learning from revealed preference. In *Proceedings of the 7th ACM Conference on Electronic Commerce*, pages 36–42, 2006 (cited on page 2).

- [BFL21] Omar Besbes, Yuri Fonseca, and Ilan Lobel. Contextual inverse optimization: of-line and online learning, 2021. DOI: [10.48550/ARXIV.2106.14015](https://arxiv.org/abs/2106.14015). URL: <https://arxiv.org/abs/2106.14015> (cited on page 8).
- [BS95] Avrim Blum and Joel Spencer. Coloring random and semi-random k-colorable graphs. *Journal of Algorithms*, 19(2):204–234, 1995 (cited on page 8).
- [BOHG13] Jesús Bobadilla, Fernando Ortega, Antonio Hernando, and Abraham Gutiérrez. Recommender systems survey. *Knowledge-based systems*, 46:109–132, 2013 (cited on page 2).
- [Bub15] Sébastien Bubeck. Convex optimization: algorithms and complexity, 2015. arXiv: [1405.4980](https://arxiv.org/abs/1405.4980) [math.OA] (cited on page 4).
- [CG18] Yu Cheng and Rong Ge. Non-convex matrix completion against a semi-random adversary. In *Conference On Learning Theory*, pages 1362–1394. PMLR, 2018 (cited on page 8).
- [CLRS11] Wei Chu, Lihong Li, Lev Reyzin, and Robert Schapire. Contextual bandits with linear payoff functions. In *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*, pages 208–214. JMLR Workshop and Conference Proceedings, 2011 (cited on page 4).
- [DG03] Sanjoy Dasgupta and Anupam Gupta. An elementary proof of a theorem of johnson and lindenstrauss. *Random Structures & Algorithms*, 22(1):60–65, 2003. DOI: <https://doi.org/10.1002/rsa.10073> (cited on page 10).
- [DPVW14] Mark A Davenport, Yaniv Plan, Ewout Van Den Berg, and Mary Wootters. 1-bit matrix completion. *Information and Inference: A Journal of the IMA*, 3(3):189–223, 2014 (cited on page 8).
- [DGT19] Ilias Diakonikolas, Themis Gouleakis, and Christos Tzamos. Distribution-independent pac learning of halfspaces with massart noise. *Advances in Neural Information Processing Systems*, 32, 2019 (cited on page 8).
- [DHSSZ15] Miroslav Dudík, Katja Hofmann, Robert E Schapire, Aleksandrs Slivkins, and Masrour Zoghi. Contextual dueling bandits. In *Conference on Learning Theory*, pages 563–587. PMLR, 2015 (cited on page 2).
- [Fei21] Uriel Feige. *Introduction to semirandom models*. In *Beyond the Worst-Case Analysis of Algorithms*. TimEditor Roughgarden, editor. Cambridge University Press, 2021, pages 189–211. DOI: [10.1017/9781108637435.013](https://doi.org/10.1017/9781108637435.013) (cited on pages 2, 8).
- [GGKMLS21] Sreenivas Gollapudi, Guru Guruganesh, Kostas Kollias, Pasin Manurangsi, Renato Paes Leme, and Jon Schneider. Contextual recommendations and low-regret cutting-plane algorithms, 2021. DOI: [10.48550/ARXIV.2106.04819](https://arxiv.org/abs/2106.04819). URL: <https://arxiv.org/abs/2106.04819> (cited on pages 4, 8).
- [JNR12] Kevin G Jamieson, Robert Nowak, and Ben Recht. Query complexity of derivative-free optimization. *Advances in Neural Information Processing Systems*, 25, 2012 (cited on pages 2, 3, 7–9).
- [JRTZ16] Dietmar Jannach, Paul Resnick, Alexander Tuzhilin, and Markus Zanker. Recommender systems—beyond matrix completion. *Communications of the ACM*, 59(11):94–102, 2016 (cited on page 2).

- [KLLST22] Jonathan A Kelner, Jerry Li, Allen Liu, Aaron Sidford, and Kevin Tian. Semi-random sparse recovery in nearly-linear time. *arXiv preprint arXiv:2203.04002*, 2022 (cited on page 8).
- [KHKN15] Junpei Komiyama, Junya Honda, Hisashi Kashima, and Hiroshi Nakagawa. Regret lower bound and optimal algorithm in dueling bandit problem. In *Conference on learning theory*, pages 1141–1154. PMLR, 2015 (cited on page 2).
- [LS18] Renato Paes Leme and Jon Schneider. Contextual search via intrinsic volumes, 2018. DOI: [10.48550/ARXIV.1804.03195](https://arxiv.org/abs/1804.03195). URL: <https://arxiv.org/abs/1804.03195> (cited on page 8).
- [LLS20] Allen Liu, Renato Paes Leme, and Jon Schneider. Optimal contextual pricing and extensions, 2020. DOI: [10.48550/ARXIV.2003.01703](https://arxiv.org/abs/2003.01703). URL: <https://arxiv.org/abs/2003.01703> (cited on page 8).
- [LLV18] Ilan Lobel, Renato Paes Leme, and Adrian Vladu. Multidimensional binary search for contextual decision-making. *Operations Research*, 66(5):1346–1361, 2018 (cited on page 8).
- [MN06] Pascal Massart and Élodie Nédélec. Risk bounds for statistical learning, 2006 (cited on page 8).
- [Moi21] Ankur Moitra. *Semirandom stochastic block models*. In *Beyond the Worst-Case Analysis of Algorithms*. TimEditor Roughgarden, editor. Cambridge University Press, 2021, pages 212–233. DOI: [10.1017/9781108637435.014](https://doi.org/10.1017/9781108637435.014) (cited on pages 2, 8).
- [OWJ⁺22] Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, et al. Training language models to follow instructions with human feedback. *Advances in Neural Information Processing Systems*, 35:27730–27744, 2022 (cited on page 2).
- [SKM21] Aadirupa Saha, Tomer Koren, and Yishay Mansour. Dueling convex optimization. In *International Conference on Machine Learning*, pages 9245–9254. PMLR, 2021 (cited on pages 2, 3, 7–9).
- [SKM22] Aadirupa Saha, Tomer Koren, and Yishay Mansour. Dueling convex optimization with general preferences. *arXiv preprint arXiv:2210.02562*, 2022 (cited on pages 2, 3, 7).
- [SK22] Aadirupa Saha and Akshay Krishnamurthy. Efficient and optimal algorithms for contextual dueling bandits under realizability. In *International Conference on Algorithmic Learning Theory*, pages 968–994. PMLR, 2022 (cited on page 2).
- [SJ11] Pannagadatta K Shivaswamy and Thorsten Joachims. Online learning with preference feedback. *arXiv preprint arXiv:1111.0712*, 2011 (cited on page 2).
- [Ver18] R. Vershynin. *High-Dimensional Probability: An Introduction with Applications in Data Science*. Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press, 2018. ISBN: 9781108415194. URL: <https://books.google.com/books?id=NDdqDwAAQBAJ> (cited on page 19).
- [Yao77] Andrew Chi-Chin Yao. Probabilistic computations: toward a unified measure of complexity. In *18th Annual Symposium on Foundations of Computer Science (sfcs 1977)*, pages 222–227. IEEE Computer Society, 1977 (cited on page 19).

- [YBKJ12] Yisong Yue, Josef Broder, Robert Kleinberg, and Thorsten Joachims. The k-armed dueling bandits problem. *Journal of Computer and System Sciences*, 78(5):1538–1556, 2012 (cited on page [2](#)).