
Conic Blackwell Algorithm: Parameter-Free Convex-Concave Saddle-Point Solving

Julien Grand-Clément*

ISOM Department
HEC Paris
grand-clement@hec.fr

Christian Kroer

IEOR Department
Columbia University
christian.kroer@columbia.edu

Abstract

We develop new parameter-free and scale-free algorithms for solving convex-concave saddle-point problems. Our results are based on a new simple regret minimizer, the Conic Blackwell Algorithm⁺ (CBA⁺), which attains $O(1/\sqrt{T})$ average regret. Intuitively, our approach generalizes to other decision sets of interest ideas from the Counterfactual Regret minimization (CFR⁺) algorithm, which has very strong practical performance for solving sequential games on simplexes. We show how to implement CBA⁺ for the simplex, ℓ_p norm balls, and ellipsoidal confidence regions in the simplex, and we present numerical experiments for solving matrix games and distributionally robust optimization problems. Our empirical results show that CBA⁺ is a simple algorithm that outperforms state-of-the-art methods on synthetic data and real data instances, without the need for any choice of step sizes or other algorithmic parameters.

1 Introduction

We are interested in solving *saddle-point problems* (SPPs) of the form

$$\min_{\mathbf{x} \in \mathcal{X}} \max_{\mathbf{y} \in \mathcal{Y}} F(\mathbf{x}, \mathbf{y}), \quad (1)$$

where $\mathcal{X} \subset \mathbb{R}^n$, $\mathcal{Y} \subset \mathbb{R}^m$ are convex, compact sets, and $F : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ is a subdifferentiable convex-concave function. Convex-concave SPPs arise in a number of practical problems. For example, the problem of computing a Nash equilibrium of a zero-sum game can be formulated as a convex-concave SPP, and this is the foundation of most methods for solving sequential zero-sum games [vS96, ZJBP07, TBJB15, KWKS20]. They also arise in imaging [CP11], ℓ_∞ -regression [ST18], Markov Decision Processes [Iye05, WKR13, ST18], and in distributionally robust optimization, where the max term represents the distributional uncertainty [ND16, BTHKM15]. In this paper we propose efficient, *parameter-free* algorithms for solving (1) in many settings, i.e., algorithms that do not require any tuning or choices of step sizes.

Repeated game framework One way to solve convex-concave SPPs is by viewing the SPP as a repeated game between two players, where each step t consists of one player choosing $\mathbf{x}_t \in \mathcal{X}$, the other player choosing $\mathbf{y}_t \in \mathcal{Y}$, and then the players observe the payoff $F(\mathbf{x}_t, \mathbf{y}_t)$. If each player employs a regret-minimization algorithm, then a well-known folk theorem says that the uniform average strategy generated by two regret minimizers repeatedly playing an SPP against each other converges to a solution to the SPP. We will call this the “repeated game framework” (see Section 2). There are already well-known algorithms for instantiating the above repeated game framework

*Julien Grand-Clément acknowledges the financial support of Hi! Paris for this project.

for (1). For example, one can employ the *online mirror descent* (OMD) algorithm, which generates iterates as follows for the first player (and similarly for the second player):

$$\mathbf{x}_{t+1} = \arg \min_{\mathbf{x} \in \mathcal{X}} \langle \eta \mathbf{f}_t, \mathbf{x} \rangle + D(\mathbf{x} \| \mathbf{x}_t), \quad (2)$$

where $\mathbf{f}_t \in \partial_{\mathbf{x}} F(\mathbf{x}_t, \mathbf{y}_t)$ ($\partial_{\mathbf{x}}$ denotes here the set of subgradients as regards the variable $\mathbf{x}$), $D(\cdot \| \cdot)$ is a Bregman divergence which measures distance between pairs of points, and $\eta > 0$ is an appropriate step size. By choosing D appropriately for $\mathcal{X}$, the update step (2) becomes efficient, and one can achieve an overall regret on the order of $O(\sqrt{T})$ after T iterations. This regret can be achieved either by choosing a fixed step size $\eta = \alpha/L\sqrt{T}$, where L is an upper bound on the norms of the subgradients visited $(\mathbf{f}_t)_{t \geq 0}$, or by choosing adaptive step sizes $\eta_t = \alpha/\sqrt{t}$, for $\alpha > 0$. This is problematic, as 1) the upper bound L may be hard to obtain in many applications and may be too conservative in practice, and 2) adequately tuning the parameter α can be time- and resource-consuming, and even practically infeasible for very large instances, since we won't know if the step size will cause a slow convergence or a divergence until late in the optimization process. This is not just a theoretical issue, as we highlight in our numerical experiments (Section 4) and in the appendices (Appendices F). Similar results and challenges hold for the popular *follow the regularized leader* (FTRL) algorithm (see Appendix F).

The above issues can be circumvented by employing *adaptive* variants of OMD or FTRL, which lead to parameter- and scale-free algorithms that estimate the parameters through the observed subgradients, e.g., AdaHedge for the simplex setting [DRVEGK14] or AdaFTRL for general compact convex decisions sets [OP15]. Yet these adaptive variants have not seen practical adoption in large-scale game-solving, where regret-matching variants are preferred (see the next paragraph). As we show in our experiments, adaptive variants of OMD and FTRL perform much worse than our proposed algorithms. While these adaptive algorithms are referred to as parameter-free, this is only true in the sense that they are able to learn the necessary parameters. Our algorithm is parameter-free in the stronger sense that there are no parameters that even require learning. Formalizing this difference may be one interesting avenue for explaining the performance discrepancy on saddle-point problems.

Regret Matching In this paper, we introduce alternative regret-minimization schemes for instantiating the above framework. Our work is motivated by recent advances on solving large-scale zero-sum sequential games. In the zero-sum sequential game setting, $\mathcal{X}$ and $\mathcal{Y}$ are simplexes, the objective function becomes $F(\mathbf{x}, \mathbf{y}) = \langle \mathbf{x}, \mathbf{A}\mathbf{y} \rangle$, and thus (1) reduces to a *bilinear* SPP. Based on this bilinear SPP formulation, the best practical methods for solving large-scale sequential games use the repeated game framework, where each player minimizes regret via some variant of *counterfactual regret minimization* (CFR, [ZJBP07]). Variants of CFR were used in every recent poker AI challenge, where poker AIs beat human poker players [BBJT15, MSB⁺17, BS18, BS19]. The CFR framework itself is a decomposition of the overall regret of the bilinear SPP into local regrets at each decision point in a sequential game [FKS19a]. The key to the practical performance of CFR-based algorithms seems to be three ingredients (beyond the CFR decomposition itself): (1) a particular regret minimizer called *regret matching*⁺ (RM⁺) [TBJB15] which is employed at each decision point, (2) aggressive iterate averaging schemes that put greater weight on recent iterates (e.g. *linear averaging*, which weights iterate at period t by $2t/T(T+1)$), and (3) an alternation scheme where the updates of the repeated game framework are performed in an asymmetric fashion. The CFR framework itself is specific to sequential bilinear games on simplexes, but these last three ingredients could potentially be generalized to other problems of the form (1). That is the starting point of the present paper.

The most challenging aspect of generalizing the above ingredients is that RM⁺ is specifically designed for minimizing regret over a simplex. However, many problems of the form (1) have convex sets $\mathcal{X}, \mathcal{Y}$ that are not simplexes, e.g., box constraints or norm-balls for distributionally robust optimization [BTHKM15]. In principle, regret matching arises from a general theory called *Blackwell approachability* [Bla56, HMC00], and similar constructions can be envisioned for other convex sets. However, in practice the literature has only focused on developing concrete implementable instantiations of Blackwell approachability for simplexes. A notable deviation from this is the work of [ABH11], who showed a general reduction between regret minimization over general convex compact sets and Blackwell approachability. However, their general reduction still does not yield a practically implementable algorithm: among other things, their reduction relies on certain black-box projections that are not always efficient. We show how to implement these necessary projections for the setting where $\mathcal{X}$ and $\mathcal{Y}$ are simplexes, ℓ_p balls, and intersections of the ℓ_2 ball with a hyperplane

(with a focus on the case where an ℓ_2 ball is intersected with a simplex, which arises naturally as a confidence region). This yields an algorithm which we will refer to as the *conic Blackwell algorithm* (CBA), which is similar in spirit to the regret matching algorithm, but crucially generalizes to other decision sets. Motivated by the practical performance of RM^+ , we construct a variant of CBA which uses a thresholding operation similar to the one employed by RM^+ . We call this algorithm CBA^+ .

Our contributions We introduce CBA^+ , a parameter-free algorithm which achieves $O(\sqrt{T})$ regret in the worst-case and generalizes the strong performances of RM^+ for bilinear, simplex saddle-points solving to other more general settings. A major selling point for CBA^+ is that it does not require any step size choices. Instead, the algorithm implicitly adjusts to the structure of the domains and losses by being instantiations of Blackwell’s approachability algorithm. After developing the CBA^+ algorithm, we then develop analogues of another crucial components for large-scale game solving. In particular, we prove a generalization of the folk theorem for the repeated game framework for solving (1), which allows us to incorporate polynomial averaging schemes such as linear averaging. We then show that CBA^+ is compatible with linear averaging on the iterates. This mirrors the case of RM and RM^+ , where only RM^+ is compatible with linear averaging on the iterates. We also show that both CBA and CBA^+ are compatible with polynomial averaging when simultaneously performed on the regrets and the iterates. Combining all these ingredients, we arrive at a new class of algorithms for solving convex-concave SPPs. As long as efficient projection operations can be performed (which we show for several practical domains, including the simplex, ℓ_p balls and confidence regions in the simplex), one can apply the repeated game framework on (1), where one can use either CBA or CBA^+ as a regret minimizer for $\mathcal{X}$ and $\mathcal{Y}$ along with polynomial averaging on the generated iterates to solve (1) at a rate of $O\left(1/\sqrt{T}\right)$.

We highlight the practical efficacy of our algorithmic framework on several domains. First, we solve two-player zero-sum matrix games and extensive-form games, where RM^+ regret minimizer combined with linear averaging and alternation, and CFR^+ , lead to very strong practical algorithms [TBJB15]. We find that CBA^+ combined with linear averaging and alternation leads to a comparable performance in terms of the iteration complexity, and may even slightly outperform RM^+ and CFR^+ . On this simplex setting, we also find that CBA^+ outperforms both AdaHedge and AdaFTRL. Second, we apply our approach to a setting where RM^+ and CFR^+ do not apply: distributionally robust empirical risk minimization (DR-ERM) problems. Across two classes of synthetic problems and four real data sets, we find that our algorithm based on CBA^+ performs orders of magnitude better than online mirror descent and FTRL, as well as their optimistic variants, when using their theoretically-correct fixed step sizes. Even when considering adaptive step sizes, or fixed step sizes that are up to 10,000 larger than those predicted by theory, our CBA^+ algorithm performs better, with only a few cases of comparable performance (at step sizes that lead to divergence for some of the other non parameter-free methods). The fast practical performance of our algorithm, combined with its simplicity and the total lack of step sizes or parameters tuning, suggests that it should be seriously considered as a practical approach for solving convex-concave SPPs in various settings.

Finally, we make a brief note on accelerated methods. Our algorithms have a rate of convergence towards a saddle point of $O(1/\sqrt{T})$, similar to OMD and FTRL. In theory, it is possible to obtain a faster $O(1/T)$ rate of convergence when F is differentiable with Lipschitz gradients, for example via mirror prox [Nem04] or other primal-dual algorithms [CP16]. However, our experimental results show that CBA^+ is faster than optimistic variants of FTRL and OMD [SALS15], the latter being almost identical to the mirror prox algorithm, and both achieving $O(1/T)$ rate of convergence. A similar conclusion has been drawn in the context of sequential game solving, where the fastest $O(1/\sqrt{T})$ CFR-based algorithms have better practical performance than the theoretically-superior $O(1/T)$ -rate methods [KWKKS20, KFS18]. In a similar vein, using *error-bound conditions*, it is possible to achieve a linear rate, e.g., when solving bilinear saddle-point problems over polyhedral decision sets using the extragradient method [Tse95] or optimistic gradient descent-ascent [WLZL20]. However, these linear rates rely on unknown constants, and may not be indicative of practical performance.

2 Game setup and Blackwell Approachability

As stated in Section 1, we will solve (1) using a repeated game framework. The first player chooses strategies from $\mathcal{X}$ in order to minimize the sequence of payoffs in the repeated game, while the second player chooses strategies from $\mathcal{Y}$ in order to maximize payoffs. There are T iterations with indices $t = 1, \dots, T$. In this framework, each iteration t consists of the following steps:

1. Each player chooses strategies $\mathbf{x}_t \in \mathcal{X}, \mathbf{y}_t \in \mathcal{Y}$
2. First player observes $\mathbf{f}_t \in \partial_{\mathbf{x}} F(\mathbf{x}_t, \mathbf{y}_t)$ and uses $\mathbf{f}_t$ when computing the next strategy
3. Second player observes $\mathbf{g}_t \in \partial_{\mathbf{y}} F(\mathbf{x}_t, \mathbf{y}_t)$ and uses $\mathbf{g}_t$ when computing the next strategy

The goal of each player is to minimize their regret $R_{T,\mathbf{x}}, R_{T,\mathbf{y}}$ across the T iterations:

$$R_{T,\mathbf{x}} = \sum_{t=1}^T \langle \mathbf{f}_t, \mathbf{x}_t \rangle - \min_{\mathbf{x} \in \mathcal{X}} \sum_{t=1}^T \langle \mathbf{f}_t, \mathbf{x} \rangle, \quad R_{T,\mathbf{y}} = \max_{\mathbf{y} \in \mathcal{Y}} \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{y} \rangle - \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{y}_t \rangle.$$

The reason this repeated game framework leads to a solution to the SPP problem (1) is the following folk theorem. Relying on F being convex-concave and subdifferentiable, it connects the regret incurred by each player to the duality gap in (1).

Theorem 2.1 (Theorem 1, [Kro20]). *Let $(\bar{\mathbf{x}}_T, \bar{\mathbf{y}}_T) = \frac{1}{T} \sum_{t=1}^T (\mathbf{x}_t, \mathbf{y}_t)$ for any $(\mathbf{x}_t)_{t \geq 1}, (\mathbf{y}_t)_{t \geq 1}$. Then*

$$\max_{\mathbf{y} \in \mathcal{Y}} F(\bar{\mathbf{x}}_T, \mathbf{y}) - \min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x}, \bar{\mathbf{y}}_T) \leq (R_{T,\mathbf{x}} + R_{T,\mathbf{y}})/T.$$

Therefore, when each player runs a regret minimizer that guarantees regret on the order of $O(\sqrt{T})$, $(\bar{\mathbf{x}}_T, \bar{\mathbf{y}}_T)_{T \geq 0}$ converges to a solution to (1) at a rate of $O(1/\sqrt{T})$. Later we will show a generalization of Theorem 2.1 that will allow us to incorporate more aggressive averaging schemes that put additional weight on the later iterates. Given the repeated game framework, the next question becomes which algorithms to employ in order to minimize regret for each player. As mentioned in Section 1, for zero-sum games, variants of regret matching are used in practice.

Blackwell Approachability Regret matching arises from the *Blackwell approachability* framework [Bla56]. In Blackwell approachability, a decision maker repeatedly takes decisions $\mathbf{x}_t$ from some convex decision set $\mathcal{X}$ (this set plays the same role as $\mathcal{X}$ or $\mathcal{Y}$ in (1)). After taking decision $\mathbf{x}_t$ the player observes a vector-valued affine payoff function $\mathbf{u}_t(\mathbf{x}) \in \mathbb{R}^n$. The goal for the decision maker is to force the average payoff $\frac{1}{t} \sum_{\tau=1}^t \mathbf{u}_\tau(\mathbf{x}_\tau)$ to approach some convex target $\mathcal{S}$. Blackwell proved that a convex target set $\mathcal{S}$ can be approached if and only if for every halfspace $\mathcal{H} \supseteq \mathcal{S}$, there exists $\mathbf{x} \in \mathcal{X}$ such that for every possible payoff function $\mathbf{u}(\cdot)$, $\mathbf{u}(\mathbf{x})$ is guaranteed to lie in $\mathcal{H}$. The action $\mathbf{x}$ is said to *force* $\mathcal{H}$. Blackwell's proof is via an algorithm: at iteration t , his algorithm projects the average payoff $\bar{\mathbf{u}} = \frac{1}{t-1} \sum_{\tau=1}^{t-1} \mathbf{u}_\tau(\mathbf{x}_\tau)$ onto $\mathcal{S}$, and then the decision maker chooses an action $\mathbf{x}_t$ that forces the tangent halfspace to $\mathcal{S}$ generated by the normal $\bar{\mathbf{u}} - \pi_{\mathcal{S}}(\bar{\mathbf{u}})$, where $\pi_{\mathcal{S}}(\bar{\mathbf{u}})$ is the orthogonal projection of $\bar{\mathbf{u}}$ onto $\mathcal{S}$. We call this algorithm *Blackwell's algorithm*; it approaches $\mathcal{S}$ at a rate of $O(1/\sqrt{T})$. It is important to note here that Blackwell's algorithm is rather a meta-algorithm than a concrete algorithm. Even within the context of Blackwell's approachability problem, one needs to devise a way to compute the forcing actions needed at each iteration, i.e., to compute $\pi_{\mathcal{S}}(\bar{\mathbf{u}})$.

Details on Regret Matching Regret matching arises by instantiating Blackwell approachability with the decision space $\mathcal{X}$ equal to the simplex $\Delta(n)$, the target set $\mathcal{S}$ equal to the nonpositive orthant $\mathbb{R}_-^n$, and the vector-valued payoff function $\mathbf{u}_t(\mathbf{x}_t) = \mathbf{f}_t - \langle \mathbf{f}_t, \mathbf{x}_t \rangle \mathbf{e}$ equal to the regret associated to each of the n actions (which correspond to the corners of $\Delta(n)$). Here $\mathbf{e} \in \mathbb{R}^n$ has one on every component. [HMC00] showed that with this setup, playing each action with probability proportional to its positive regret up to time t satisfies the forcing condition needed in Blackwell's algorithm. Formally, regret matching (RM) keeps a running sum $\mathbf{r}_t = \sum_{\tau=1}^t (\mathbf{f}_\tau - \langle \mathbf{f}_\tau, \mathbf{x}_\tau \rangle \mathbf{e})$, and then action i is played with probability $x_{t+1,i} = [\mathbf{r}_{t,i}]^+ / \sum_{i=1}^n [\mathbf{r}_{t,i}]^+$, where $[\cdot]^+$ denotes thresholding at zero. By Blackwell's approachability theorem, this algorithm converges to zero average regret at a rate of $O(1/\sqrt{T})$. In zero-sum game-solving, it was discovered that a variant of

regret matching leads to extremely strong practical performance (but the same theoretical rate of convergence). In regret matching⁺ (RM⁺), the running sum is thresholded at zero at every iteration: $\mathbf{r}_t = [\mathbf{r}_{t-1} + \mathbf{f}_t - \langle \mathbf{f}_t, \mathbf{x}_t \rangle \mathbf{e}]^+$, and then actions are again played proportional to $\mathbf{r}_t$. In the next section, we describe a more general class of regret-minimization algorithms based on Blackwell's algorithm for general sets $\mathcal{X}$, introduced in [ABH11]. Note that a similar construction of a general class of algorithms can be achieved through the *Lagrangian Hedging* framework of [Gor07]. It would be interesting to construct a CBA⁺-like algorithm and efficient projection approaches for this framework as well.

3 Conic Blackwell Algorithm

We present the Conic Blackwell Algorithm Plus (CBA⁺), a no-regret algorithm which uses a variation of Blackwell's approachability procedure [Bla56] to perform regret minimization on general convex compact decision sets $\mathcal{X}$. We will assume that losses are coming from a bounded set; this occurs, for example, if there exists L_x, L_y (that we do not need to know), such that

$$\|\mathbf{f}\| \leq L_x, \|\mathbf{g}\| \leq L_y, \forall \mathbf{x} \in \mathcal{X}, \mathbf{y} \in \mathcal{Y}, \forall \mathbf{f} \in \partial_x F(\mathbf{x}, \mathbf{y}), \forall \mathbf{g} \in \partial_y F(\mathbf{x}, \mathbf{y}). \quad (3)$$

CBA⁺ is best understood as a combination of two steps. The first is the basic CBA algorithm, derived from Blackwell's algorithm, which we describe next. To convert Blackwell's algorithm to a regret minimizer on $\mathcal{X}$, we use the reduction from [ABH11], which considers the conic hull $\mathcal{C} = \text{cone}(\{\kappa\} \times \mathcal{X})$ where $\kappa = \max_{\mathbf{x} \in \mathcal{X}} \|\mathbf{x}\|_2$. The Blackwell approachability problem is then instantiated with $\mathcal{X}$ as the decision set, target set equal to the polar $\mathcal{C}^\circ = \{\mathbf{z} : \langle \mathbf{z}, \hat{\mathbf{z}} \rangle \leq 0, \forall \hat{\mathbf{z}} \in \mathcal{C}\}$ of $\mathcal{C}$, and payoff vectors $(\langle \mathbf{f}_t, \mathbf{x}_t \rangle, -\mathbf{f}_t)$. The conic Blackwell algorithm (CBA) is implemented by projecting the average payoff vector onto $\mathcal{C}$, calling this projection $\alpha(\kappa, \mathbf{x})$ with $\alpha \geq 0$ and $\mathbf{x} \in \mathcal{X}$, and playing the action $\mathbf{x}$.

The second step in CBA⁺ is to modify CBA to make it analogous to RM⁺ rather than to RM. To do this, the algorithm does not keep track of the average payoff vector. Instead, we keep a running aggregation of the payoffs, where we always add the newest payoff to the aggregate, and then project the aggregate onto $\mathcal{C}$. More concretely, pseudocode for CBA⁺ is given in Algorithm 1. This pseudocode relies on two functions: $\text{CHOOSEDECISION}_{\text{CBA}^+} : \mathbb{R}^{n+1} \rightarrow \mathbb{R}^n$, which maps the aggregate payoff vector $\mathbf{u}_t$ to a decision in $\mathcal{X}$, and $\text{UPDATEPAYOFF}_{\text{CBA}^+}$ which controls how we aggregate payoffs. Given an aggregate payoff vector $\mathbf{u} = (\tilde{u}, \hat{\mathbf{u}}) \in \mathbb{R} \times \mathbb{R}^n$, we have

$$\text{CHOOSEDECISION}_{\text{CBA}^+}(\mathbf{u}) = (\kappa/\tilde{u})\hat{\mathbf{u}}.$$

If $\tilde{u} = 0$, we just let $\text{CHOOSEDECISION}_{\text{CBA}^+}(\mathbf{u}) = \mathbf{x}_0$ for some chosen $\mathbf{x}_0 \in \mathcal{X}$.

The function $\text{UPDATEPAYOFF}_{\text{CBA}^+}$ is implemented by adding the most recent payoff to the aggregate payoffs, and then projecting onto $\mathcal{C}$. More formally, it is defined as

$$\text{UPDATEPAYOFF}_{\text{CBA}^+}(\mathbf{u}, \mathbf{x}, \mathbf{f}, \omega, S) = \pi_{\mathcal{C}} \left(\frac{S}{S + \omega} \mathbf{u} + \frac{\omega}{S + \omega} (\langle \mathbf{f}, \mathbf{x} \rangle / \kappa, -\mathbf{f}) \right),$$

where ω is the weight assigned to the most recent payoff and S the weight assigned to the previous aggregate payoff $\mathbf{u}$. Because of the projection step in $\text{UPDATEPAYOFF}_{\text{CBA}^+}$, we always have $\mathbf{u} \in \mathcal{C}$, which in turn guarantees that $\text{CHOOSEDECISION}_{\text{CBA}^+}(\mathbf{u}) \in \mathcal{X}$, since $\mathcal{C} = \text{cone}(\{\kappa\} \times \mathcal{X})$.

Let us give some intuition on the effect of projection onto $\mathcal{C}$. In a geometric sense, it is easier to visualize things in $\mathbb{R}^2$ with $\mathcal{C} = \mathbb{R}_+^2$ and $\mathcal{C}^\circ = \mathbb{R}_-^2$. The projection on $\mathcal{C}$ moves the vector along the edges of $\mathcal{C}^\circ$, maintaining the distance to $\mathcal{C}^\circ$ and moving toward the vector $\mathbf{0}$. This is illustrated in Figure 5 in Appendix B.1. From a game-theoretic standpoint, the projection on $\mathcal{C} = \mathbb{R}_+^2$ eliminates the components of the payoffs that are negative. It enables CBA⁺ to be less pessimistic than CBA, which may accumulate negative payoffs on actions for a long time and never resets the components of the aggregated payoff to 0, leading to some actions being chosen less frequently.

We will see in the next section that RM⁺ is related to CBA⁺ but replaces the exact projection step $\pi_{\mathcal{C}}(\mathbf{u})$ in $\text{UPDATEPAYOFF}_{\text{CBA}^+}$ by a suboptimal solution to the projection problem. Let us also note the difference between CBA⁺ and the algorithm introduced in [ABH11], which we have called

Algorithm 1 Conic Blackwell Algorithm Plus (CBA⁺)

- 1: **Input** A convex, compact set $\mathcal{X} \subset \mathbb{R}^n$, $\kappa = \max\{\|\mathbf{x}\|_2 \mid \mathbf{x} \in \mathcal{X}\}$.
 - 2: **Algorithm parameters** Weights $(\omega_\tau)_{\tau \geq 1} \in \mathbb{R}^{\mathbb{N}}$.
 - 3: **Initialization** $t = 1$, $\mathbf{x}_1 \in \mathcal{X}$.
 - 4: Observe $\mathbf{f}_1$ then set $\mathbf{u}_1 = (\langle \mathbf{f}_1, \mathbf{x}_1 \rangle / \kappa, -\mathbf{f}_1) \in \mathbb{R} \times \mathbb{R}^n$.
 - 5: **for** $t \geq 1$ **do**
 - 6: Choose $\mathbf{x}_{t+1} = \text{CHOOSEDECISION}_{\text{CBA}^+}(\mathbf{u}_t)$.
 - 7: Observe the loss $\mathbf{f}_{t+1} \in \mathbb{R}^n$.
 - 8: Update $\mathbf{u}_{t+1} = \text{UPDATEPAYOFF}_{\text{CBA}^+}(\mathbf{u}_t, \mathbf{x}_{t+1}, \mathbf{f}_{t+1}, \omega_{t+1}, \sum_{\tau=1}^t \omega_\tau)$.
 - 9: Increment $t \leftarrow t + 1$.
-

CBA. CBA uses different UPDATEPAYOFF and CHOOSEDECISION functions. In CBA the payoff update is defined as

$$\text{UPDATEPAYOFF}_{\text{CBA}}(\mathbf{u}, \mathbf{x}, \mathbf{f}, \omega, S) = \frac{S}{S + \omega} \mathbf{u} + \frac{\omega}{S + \omega} (\langle \mathbf{f}, \mathbf{x} \rangle / \kappa, -\mathbf{f}).$$

Note in particular the lack of projection as compared to CBA⁺, analogous to the difference between RM and RM⁺. The CHOOSEDECISION_{CBA} function then requires a projection onto $\mathcal{C}$:

$$\text{CHOOSEDECISION}_{\text{CBA}}(\mathbf{u}) = \text{CHOOSEDECISION}_{\text{CBA}^+}(\pi_{\mathcal{C}}(\mathbf{u})).$$

Based upon the analysis in [Bla56], [ABH11] show that CBA with uniform weights (both on payoffs and decisions) guarantees $O(1/\sqrt{T})$ average regret. The difference between CBA⁺ and CBA is similar to the difference between the RM and RM⁺ algorithms. In practice, RM⁺ performs significantly better than RM for solving matrix games, when combined with *linear averaging* on the decisions (as opposed to the uniform averaging used in Theorem 2.1). In the next theorem, we show that CBA⁺ is compatible with linear averaging on decisions only. We present a detailed proof in Appendix B.

Theorem 3.1. Consider $(\mathbf{x}_t)_{t \geq 0}$ generated by CBA⁺ with uniform weights: $\omega_\tau = 1, \forall \tau \geq 1$. Let $L = \max\{\|\mathbf{f}_t\|_2 \mid t \geq 1\}$ and $\kappa = \max\{\|\mathbf{x}\|_2 \mid \mathbf{x} \in \mathcal{X}\}$. Then

$$\frac{\sum_{t=1}^T t \langle \mathbf{f}_t, \mathbf{x}_t \rangle - \min_{\mathbf{x} \in \mathcal{X}} \sum_{t=1}^T t \langle \mathbf{f}_t, \mathbf{x} \rangle}{T(T+1)} = O\left(\kappa L / \sqrt{T}\right).$$

Note that in Theorem 3.1, we have *uniform* weights on the sequence of payoffs $(\mathbf{u}_t)_{t \geq 0}$, but *linearly increasing* weights on the sequence of decisions. The proof relies on properties specific to CBA⁺, and it does not extend to CBA. Numerically it also helps CBA⁺ but not CBA. In Appendix B, we show that both CBA and CBA⁺ achieve $O\left(\kappa L / \sqrt{T}\right)$ convergence rates when using a weighted average on *both* the decisions and the payoffs (Theorems B.2-B.3). In practice, using linear averaging only on the decisions, as in Theorem 3.1, performs vastly better than linear averaging on both decisions and payoffs. We present empirical evidence of this in Appendix B.

We can compare the $O\left(\kappa L / \sqrt{T}\right)$ average regret for CBA⁺ with the $O\left(\Omega L / \sqrt{T}\right)$ average regret for OMD [NY83, BTN01] and FTRL [AHR09, McM11], where $\Omega = \max\{\|\mathbf{x} - \mathbf{x}'\|_2 \mid \mathbf{x}, \mathbf{x}' \in \mathcal{X}\}$. We can always recenter $\mathcal{X}$ to contain $\mathbf{0}$, in which case the bounds for OMD/FTRL and CBA⁺ are equivalent since $\kappa \leq \Omega \leq 2\kappa$. Note that the bound on the average regret for Optimistic OMD (O-OMD, [CYL⁺12]) and Optimistic FTRL (O-FTRL, [RS13]) is $O\left(\Omega^2 L / T\right)$ in the game setup, a priori better than the bound for CBA⁺ as regards the number of iterations T . Nonetheless, we will see in Section 4 that the empirical performance of CBA⁺ is better than that of $O(1/T)$ methods. A similar situation occurs for RM⁺ compared to O-OMD and O-FTRL for solving poker games [FKS19b, KWKKS20].

The following theorem gives the convergence rate of CBA⁺ for solving saddle-points (1), based on our convergence rate on the regret of each player (Theorem 3.1). The proof is in Appendix C.

Theorem 3.2. Let $(\bar{\mathbf{x}}_T, \bar{\mathbf{y}}_T) = 2 \sum_{t=1}^T t (\mathbf{x}_t, \mathbf{y}_t) / (T(T+1))$, where $(\mathbf{x}_t)_{t \geq 0}, (\mathbf{y}_t)_{t \geq 0}$ are generated by the repeated game framework with CBA⁺ with uniform weights: $\omega_\tau = 1, \forall \tau \geq 1$. Let $L = \max\{L_x, L_y\}$ defined in (3) and $\kappa = \max\{\max\{\|\mathbf{x}\|_2, \|\mathbf{y}\|_2\} \mid \mathbf{x} \in \mathcal{X}, \mathbf{y} \in \mathcal{Y}\}$. Then

$$\max_{\mathbf{y} \in \mathcal{Y}} F(\bar{\mathbf{x}}_T, \mathbf{y}) - \min_{\mathbf{x} \in \mathcal{X}} F(\mathbf{x}, \bar{\mathbf{y}}_T) = O\left(\kappa L / \sqrt{T}\right).$$

3.1 Efficient implementations of CBA⁺

To obtain an implementation of CBA⁺ and CBA, we need to efficiently resolve the functions $\text{CHOOSEDECISION}_{\text{CBA}^+}$ and $\text{UPDATEPAYOFF}_{\text{CBA}^+}$. In particular, we need to compute $\pi_{\mathcal{C}}(\mathbf{u})$, the orthogonal projection of $\mathbf{u}$ onto the cone $\mathcal{C}$, where $\mathcal{C} = \text{cone}(\{\kappa\} \times \mathcal{X})$:

$$\pi_{\mathcal{C}}(\mathbf{u}) \in \arg \min_{\mathbf{y} \in \mathcal{C}} \|\mathbf{y} - \mathbf{u}\|_2^2. \quad (4)$$

Even for CBA this problem must be resolved, since [ABH11] did not study whether (4) can be efficiently solved. It turns out that (4) can be computed in closed-form or quasi closed-form for many decision sets $\mathcal{X}$ of interest. Interestingly, parts of the proofs rely on *Moreau's Decomposition Theorem* [CR13], which states that $\pi_{\mathcal{C}}(\mathbf{u})$ can be recovered from $\pi_{\mathcal{C}^\circ}(\mathbf{u})$ and vice-versa, because $\pi_{\mathcal{C}}(\mathbf{u}) + \pi_{\mathcal{C}^\circ}(\mathbf{u}) = \mathbf{u}$. We present the detailed complexity results and the proofs in Appendix D.

Simplex $\mathcal{X} = \Delta(n)$ is the classical setting used for matrix games. Also, for extensive-form games, CFR decomposes the decision sets (treeplexes) into a set of regret minimization problems over the simplex [FKS19a]. Here, n is the number of actions of a player and $\mathbf{x} \in \Delta(n)$ represents a randomized strategy. In this case, $\pi_{\mathcal{C}}(\mathbf{u})$ can be computed in $O(n \log(n))$. Note that RM and RM⁺ are obtained by choosing a suboptimal solution to (4), avoiding the $O(n \log(n))$ sorting operation, whereas CBA and CBA⁺ choose optimally (see Appendix D). Thus, RM and RM⁺ can be seen as approximate versions of CBA and CBA⁺, where (4) is solved approximatively at every iteration. In our numerical experiments, we will see that CBA⁺ slightly outperforms RM⁺ and CFR⁺ in terms of iteration count.

ℓ_p balls This is when $\mathcal{X} = \{\mathbf{x} \in \mathbb{R}^n \mid \|\mathbf{x}\|_p \leq 1\}$ with $p \geq 1$ or $p = \infty$. This is of interest for instance in distributionally robust optimization [BTHKM15, ND16], ℓ_∞ regression [ST18] and saddle-point reformulation of Markov Decision Process [JS20]. For $p = 2$, we can compute $\pi_{\mathcal{C}}(\mathbf{u})$ in closed-form, i.e., in $O(n)$ arithmetic operations. For $p \in \{1, \infty\}$, we can compute $\pi_{\mathcal{C}}(\mathbf{u})$ in $O(n \log(n))$ arithmetic operations using a sorting algorithm.

Ellipsoidal confidence region in the simplex Here, $\mathcal{X}$ is an *ellipsoidal subregion of the simplex*, defined as $\mathcal{X} = \{\mathbf{x} \in \Delta(n) \mid \|\mathbf{x} - \mathbf{x}_0\|_2 \leq \epsilon_x\}$. This type of decision set is widely used because it is associated with confidence regions when estimating a probability distribution from observed data [Iye05, BdHP19]. It can also be used in the Bellman update for robust Markov Decision Process [Iye05, WKR13, GGC18]. We also assume that the confidence region is “entirely contained in the simplex”: $\{\mathbf{x} \in \mathbb{R}^n \mid \mathbf{x}^\top \mathbf{e} = 1\} \cap \{\mathbf{x} \in \mathbb{R}^n \mid \|\mathbf{x} - \mathbf{x}_0\|_2 \leq \epsilon_x\} \subseteq \Delta(n)$, to avoid degenerate components. In this case, using a change of basis we show that it is possible to compute $\pi_{\mathcal{C}}(\mathbf{u})$ in closed-form, i.e., in $O(n)$ arithmetic operations.

Other potential sets of interests Other important decision sets include sets based on Kullback-Leibler divergence $\{\mathbf{x} \in \Delta(n) \mid KL(\mathbf{x}, \mathbf{x}_0) \leq \epsilon_x\}$, or, more generally, ϕ -divergence [BTDHDW⁺13]. For these sets, we did not find a closed-form solution to the projection problem (4). Still, as long as the domain $\mathcal{X}$ is a convex set, computing $\pi_{\mathcal{C}}(\mathbf{u})$ remains a convex problem, and it can be solved efficiently with solvers, although this results in a slower algorithm than with closed-form computations of $\pi_{\mathcal{C}}(\mathbf{u})$.

4 Numerical experiments

In this section we investigate the practical performances of our algorithms on several instances of saddle-point problems. We start by comparing CBA⁺ with RM⁺ in the matrix and extensive-form games setting. We then turn to comparing our algorithms on instances from the distributionally robust optimization literature. The code for all experiments is available in the supplemental material.

4.1 Matrix games on the simplex and Extensive-Form Games

Since the motivation for CBA⁺ is to obtain the strong empirical performances of RM⁺ and CFR⁺ on other decision sets than the simplex, we start by checking that CBA⁺ indeed provides comparable

performance on simplex settings. We compare these methods on matrix games

$$\min_{\mathbf{x} \in \Delta(n)} \max_{\mathbf{y} \in \Delta(m)} \langle \mathbf{x}, \mathbf{A}\mathbf{y} \rangle,$$

where $\mathbf{A}$ is the matrix of payoff, and on extensive-form games (EFGs). EFGs can also be written as SPPs with bilinear objective and $\mathcal{X}, \mathcal{Y}$ polytopes encoding the players' space of sequential strategies [vS96]. EFGs can be solved via simplex-based regret minimization by using the counterfactual regret minimization (CFR) framework to decompose regrets into local regrets at each simplex. Explaining CFR is beyond the scope of this work; we point the reader to [ZJBP07] or newer explanations [FKS19c, FKS19a]. For matrix games, we generate 70 synthetic 10-dimensional matrix games with $A_{ij} \sim U[0, 1]$ and compare the most efficient algorithms for matrix games with linear averaging: CBA^+ and RM^+ . We also compare with two other scale-free no-regret algorithms, AdaHedge [DRVEGK14] and AdaFTRL [OP15]. Figure 1a presents the duality gap of the current solutions vs. the number of steps. Here, both CBA^+ and RM^+ use *alternation*, which is a trick that is well-known to improve the performances of RM^+ [TBJB15], where the repeated game framework is changed such that players take turns updating their strategies, rather than performing these updates simultaneously, see Appendix E.1 for details.²

For EFGs, we compare CBA^+ and CFR^+ on many poker AI benchmark instances, including Leduc, Kuhn, search games and sheriff (see [FKS21] for game descriptions). We present our results in Figures 1b-1d. Additional details and experiments for EFGs are presented in Appendix E.4. Overall, we see in Figure 1 that CBA^+ may slightly outperform RM^+ and CFR^+ , two of the strongest algorithms for matrix games and EFGs, which were shown to achieve the best empirical performances compared to a wide range of algorithms, including Hedge and other first-order methods [Kro20, KFS18, FKS19b]. For matrix games, AdaHedge and AdaFTRL are both outperformed by RM^+ and CBA^+ ; we present more experiments to compare RM^+ and CBA^+ in Appendix E.2, and more experiments with matrix games in Appendix E.3. Recall that our goal is to generalize these strong performance to other settings: we present our numerical experiments for solving distributionally robust optimization problems in the next section.

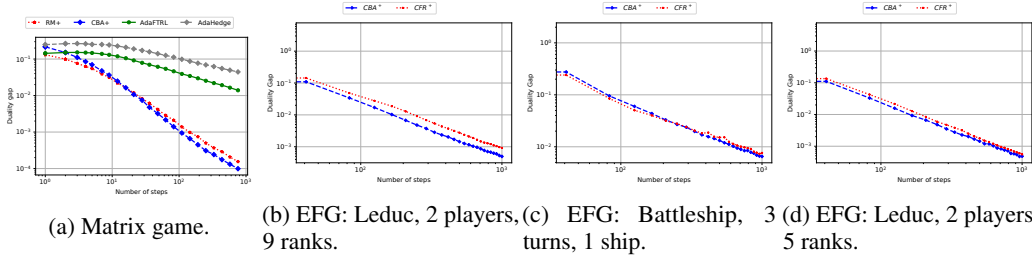


Figure 1: Comparison of CBA^+ with RM^+ and CFR^+ on matrix games and EFGs.

4.2 Distributionally Robust Optimization

Problem setup Broadly speaking, DRO attempts to exploit partial knowledge of the statistical properties of the model parameters to obtain risk-averse optimal solutions [RM19]. We focus on the following instance of distributionally robust classification with logistic losses [ND16, BTHKM15]. There are m observed feature-label pairs $(\mathbf{a}_i, b_i) \in \mathbb{R}^n \times \{-1, 1\}$, and we want to solve

$$\min_{\mathbf{x} \in \mathbb{R}^n, \|\mathbf{x} - \mathbf{x}_0\|_2 \leq R} \max_{\mathbf{y} \in \Delta(m), \|\mathbf{y} - \mathbf{y}_0\|_2^2 \leq \lambda} \sum_{i=1}^m y_i \ell_i(\mathbf{x}), \quad (5)$$

where $\ell_i(\mathbf{x}) = \log(1 + \exp(-b_i \mathbf{a}_i^\top \mathbf{x}))$. The formulation (5) takes a worst-case approach to put more weight on misclassified observations and provides some statistical guarantees, e.g., it can be seen as a convex regularization of standard empirical risk minimization instances [DGN21].

We compare CBA^+ (with linear averaging and alternation) with Online Mirror Descent (OMD), Optimistic OMD (O-OMD), Follow-The-Regularized-Leader (FTRL) and Optimistic FTRL (O-FTRL). We provide a detailed presentation of our implementations of these algorithms in Appendix

²We note that RM^+ is guaranteed to retain its convergence rate under alternation. In contrast, we leave resolving this property for CBA^+ to future work.

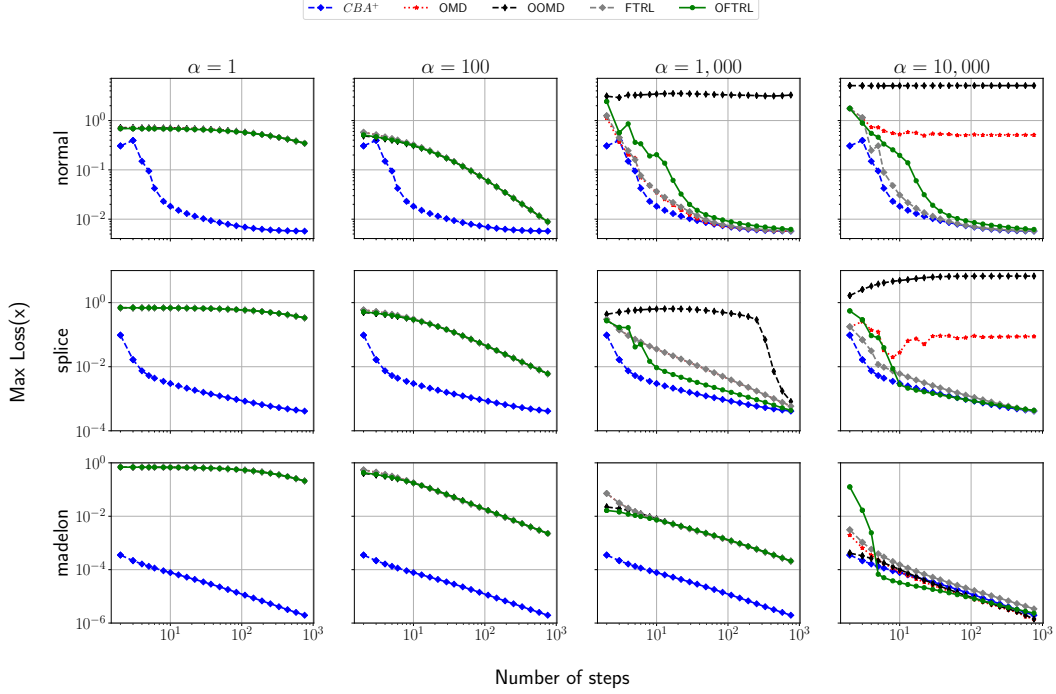


Figure 2: Comparisons of the performances of CBA^+ with OMD, FTRL, O-OMD and O-FTRL with fixed step sizes, on synthetic (with *normal* distribution) and real data sets (*splice* and *madelon*).

F. We compare the performances of these algorithms with CBA^+ on two synthetic data sets and four real data sets. We use linear averaging on decisions for all algorithms, and parameters $x_0 = \mathbf{0}$, $R = 10$, $y_0 = (1, \dots, 1)/m$, $\lambda = 1/2m$ in eq. (5).

Synthetic and real instances For the synthetic classification instances, we generate an optimal $x^* \in \mathbb{R}^n$, sample $a_i \sim N(0, I)$ for $i \in \{1, \dots, m\}$, set labels $b_i = \text{sign}(a_i^\top x^*)$, and then we flip 10% of them. For the real classification instances, we use the following data sets from the *libsvm* website³: *adult*, *australian*, *splice*, *madelon*. Details about the empirical setting, the data sets and additional numerical experiments are presented in Appendix G.

Choice of step sizes One of the main motivation for CBA^+ is to obtain a *parameter-free* algorithm. Choosing a *fixed* step size η for the other algorithms requires knowing a bound L on the norm of the instantaneous payoffs (see Appendix F.2 for our derivations of this upper bound). This is a major limitation in practice: these bounds may be very conservative, leading to small step sizes. We highlight this by showing the performance of all four algorithms, for various fixed step sizes $\eta = \alpha \times \eta_{\text{th}}$, where $\alpha \in \{1, 100, 1,000, 10,000\}$ is a multiplier and η_{th} is the theoretical step size which guarantees the convergence of the algorithms for each instance. We detail the computation of η_{th} in Appendix F.2. We present the results of our numerical experiments on synthetic and real data sets in Figure 2. Additional simulations with adaptive step sizes $\eta_t = 1/\sqrt{\sum_{\tau=1}^{t-1} \|f_\tau\|_2^2}$ [Ora19] are presented in Figure 3 and in Appendix G.

Results and discussion In Figure 2, we present the worst-case loss of the current solution $\bar{x}_T$ in terms of the number of steps T . We see that when the step sizes is chosen as the theoretical step sizes guaranteeing the convergence of the non-parameter free algorithms ($\alpha = 1$), CBA^+ vastly outperforms all of the algorithms. When we take more aggressive step sizes, the non-parameter-free algorithms become more competitive. For instance, when $\alpha = 1,000$, OMD, FTRL and O-FTRL are competitive with CBA^+ for the experiments on synthetic data sets. However, for this same instance and $\alpha = 1,000$, O-OMD diverges, because the step sizes are far greater than the theoretical step sizes guaranteeing convergence. At $\alpha = 10,000$, both OMD and O-OMD diverge. The same type of performances also hold for the *splice* data set. Finally, for the *madelon* data set, the non

³<https://www.csie.ntu.edu.tw/~cjlin/libsvmtools/datasets/>

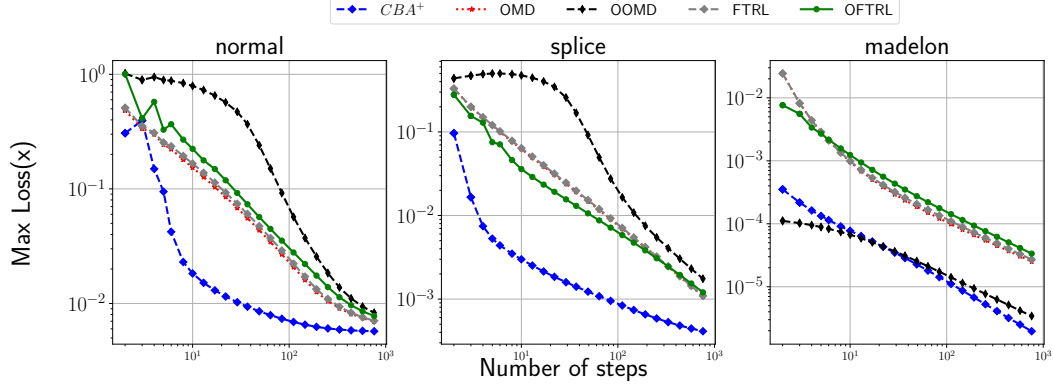


Figure 3: Comparisons of the performances of CBA^+ , OMD, FTRL, O-OMD and O-FTRL with adaptive step sizes, on synthetic (with *normal* distribution) and real data sets (*splice* and *madelon*).

parameter-free algorithms start to be competitive with CBA^+ only when $\alpha = 10,000$. Again, we note that this range of step sizes η is completely outside the values η_{th} that guarantee convergence of the algorithms, and fine-tuning the algorithms is time- and resource-consuming. In contrast, CBA^+ can be used without wasting time on exploring and finding the best convergence rates, and with confidence in the convergence of the algorithm. Similar observations hold for adaptive step sizes (see Figure 3 and Appendix G). The overall poor performances of the optimistic methods (compared to their $O(1/T)$ average regret guarantees) may reflect their sensibility to the choice of the step sizes. Additional experiments in Appendix G with other real and synthetic EFG and DRO instances show the robustness of the strong performances of CBA^+ across additional problem instances.

Running times compared to CBA^+ We would like to emphasize that all of our figures show the *number of steps* on the x -axis, and not the actual running times of the algorithms. Overall, CBA^+ converges to an optimal solution to the DRO instance (5) vastly faster than the other algorithms. In particular, empirically, CBA^+ is 2x-2.5x faster than OMD, FTRL and O-FTRL, and 3x-4x faster than O-OMD. This is because OMD, FTRL, O-OMD, and O-FTRL require binary searches at each step, see Appendix F. The functions used in the binary searches themselves require solving an optimization program (an orthogonal projection onto the simplex, see (33)) at each evaluation. Even though computing the orthogonal projection of a vector onto the simplex can be done in $O(n \log(n))$, this results in slower overall running time, compared to CBA^+ with (quasi) closed-form updates at each step. The situation is even worse for O-OMD, which requires two proximal updates at each iteration. We acknowledge that the same holds for CBA^+ compared to RM^+ . In particular, CBA^+ is slightly slower than RM^+ , because of the computation of $\pi_{\mathcal{C}}(u)$ in $O(n \log(n))$ operations at every iteration.

5 Conclusion

We have introduced CBA^+ , a new algorithm for convex-concave saddle-point solving, that is 1) simple to implement for many practical decision sets, 2) completely parameter-free and does not attempt to learn any step sizes, and 3) competitive with, or even better than, state-of-the-art approaches for the best choices of parameters, both for matrix games, extensive-form games, and distributionally robust instances. Our paper is based on Blackwell approachability, which has been used to achieved important breakthroughs in poker AI in recent years, and we hope to generalize the use and implementation of this framework to other important problem instances. Interesting future directions of research include developing a theoretical understanding of the improvements related to alternation in our setting, designing efficient implementations for other widespread decision sets (e.g., based on Kullback-Leibler divergence or ϕ -divergence), and novel accelerated versions based on strong convex-concavity or optimism.

Societal impact There is a priori no direct negative societal consequence to this work, since our methods simply return the same solutions as previous algorithms but in a simpler and more efficient way.

References

- [ABH11] Jacob Abernethy, Peter L Bartlett, and Elad Hazan. Blackwell approachability and no-regret learning are equivalent. In *Proceedings of the 24th Annual Conference on Learning Theory*, pages 27–46. JMLR Workshop and Conference Proceedings, 2011.
- [AHR09] Jacob D Abernethy, Elad Hazan, and Alexander Rakhlin. Competing in the dark: An efficient algorithm for bandit linear optimization. 2009.
- [BBJT15] Michael Bowling, Neil Burch, Michael Johanson, and Oskari Tammelin. Heads-up limit hold'em poker is solved. *Science*, 347(6218):145–149, 2015.
- [BdHP19] Dimitris Bertsimas, Dick den Hertog, and Jean Pauphilet. Probabilistic guarantees in robust optimization. 2019.
- [Bla56] David Blackwell. An analog of the minimax theorem for vector payoffs. *Pacific Journal of Mathematics*, 6(1):1–8, 1956.
- [BMS19] Neil Burch, Matej Moravcik, and Martin Schmid. Revisiting cfr+ and alternating updates. *Journal of Artificial Intelligence Research*, 64:429–443, 2019.
- [BS18] Noam Brown and Tuomas Sandholm. Superhuman AI for heads-up no-limit poker: Libratus beats top professionals. *Science*, 359(6374):418–424, 2018.
- [BS19] Noam Brown and Tuomas Sandholm. Superhuman AI for multiplayer poker. *Science*, 365(6456):885–890, 2019.
- [BT03] Amir Beck and Marc Teboulle. Mirror descent and nonlinear projected subgradient methods for convex optimization. *Operations Research Letters*, 31(3):167–175, 2003.
- [BTDHDW⁺13] Aharon Ben-Tal, Dick Den Hertog, Anja De Waegenaere, Bertrand Melenberg, and Gijs Rennen. Robust solutions of optimization problems affected by uncertain probabilities. *Management Science*, 59(2):341–357, 2013.
- [BTHKM15] Aharon Ben-Tal, Elad Hazan, Tomer Koren, and Shie Mannor. Oracle-based robust optimization via online learning. *Operations Research*, 63(3):628–638, 2015.
- [BTN01] Aharon Ben-Tal and Arkadi Nemirovski. *Lectures on modern convex optimization: analysis, algorithms, and engineering applications*, volume 2. Siam, 2001.
- [CP11] Antonin Chambolle and Thomas Pock. A first-order primal-dual algorithm for convex problems with applications to imaging. *Journal of mathematical imaging and vision*, 40(1):120–145, 2011.
- [CP16] Antonin Chambolle and Thomas Pock. On the ergodic convergence rates of a first-order primal-dual algorithm. *Mathematical Programming*, 159(1-2):253–287, 2016.
- [CR13] Patrick L Combettes and Noli N Reyes. Moreau’s decomposition in banach spaces. *Mathematical Programming*, 139(1):103–114, 2013.
- [CYL⁺12] Chao-Kai Chiang, Tianbao Yang, Chia-Jung Lee, Mehrdad Mahdavi, Chi-Jen Lu, Rong Jin, and Shenghuo Zhu. Online optimization with gradual variations. In *Conference on Learning Theory*, pages 6–1. JMLR Workshop and Conference Proceedings, 2012.
- [DGN21] John C Duchi, Peter W Glynn, and Hongseok Namkoong. Statistics of robust optimization: A generalized empirical likelihood approach. *Mathematics of Operations Research*, 2021.
- [DRVEGK14] Steven De Rooij, Tim Van Erven, Peter D Grünwald, and Wouter M Koolen. Follow the leader if you can, hedge if you must. *The Journal of Machine Learning Research*, 15(1):1281–1316, 2014.
- [DSSSC08] John Duchi, Shai Shalev-Shwartz, Yoram Singer, and Tushar Chandra. Efficient projections onto the L_1 ball for learning in high dimensions. In *Proceedings of the 25th international conference on Machine learning*, pages 272–279, 2008.
- [EPGMFBV03] Juan José Egozcue, Vera Pawlowsky-Glahn, Glòria Mateu-Figueras, and Carles Barcelo-Vidal. Isometric logratio transformations for compositional data analysis. *Mathematical Geology*, 35(3):279–300, 2003.

- [FKS19a] Gabriele Farina, Christian Kroer, and Tuomas Sandholm. Online convex optimization for sequential decision processes and extensive-form games. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 1917–1925, 2019.
- [FKS19b] Gabriele Farina, Christian Kroer, and Tuomas Sandholm. Optimistic regret minimization for extensive-form games via dilated distance-generating functions. In *Advances in Neural Information Processing Systems*, pages 5222–5232, 2019.
- [FKS19c] Gabriele Farina, Christian Kroer, and Tuomas Sandholm. Regret circuits: Composability of regret minimizers. In *International Conference on Machine Learning*, pages 1863–1872, 2019.
- [FKS21] Gabriele Farina, Christian Kroer, and Tuomas Sandholm. Faster game solving via predictive blackwell approachability: Connecting regret matching and mirror descent. In *Proceedings of the AAAI Conference on Artificial Intelligence*. AAAI, 2021.
- [GCK20a] Julien Grand-Clément and Christian Kroer. First-order methods for Wasserstein distributionally robust MDP. *arXiv preprint arXiv:2009.06790*, 2020.
- [GCK20b] Julien Grand-Clément and Christian Kroer. Scalable first-order methods for robust mdps. *arXiv preprint arXiv:2005.05434*, 2020.
- [GGC18] Vineet Goyal and Julien Grand-Clément. Robust Markov decision process: Beyond rectangularity. *arXiv preprint arXiv:1811.00215*, 2018.
- [GKG19] Yuan Gao, Christian Kroer, and Donald Goldfarb. Increasing iterate averaging for solving saddle-point problems. *arXiv preprint arXiv:1903.10646*, 2019.
- [Gor07] Geoffrey J Gordon. No-regret algorithms for online convex programs. In *Advances in Neural Information Processing Systems*, pages 489–496. Citeseer, 2007.
- [HMC00] Sergiu Hart and Andreu Mas-Colell. A simple adaptive procedure leading to correlated equilibrium. *Econometrica*, 68(5):1127–1150, 2000.
- [Iye05] Garud Iyengar. Robust dynamic programming. *Mathematics of Operations Research*, 30(2):257–280, 2005.
- [JS20] Yujia Jin and Aaron Sidford. Efficiently solving mdps with stochastic mirror descent. In *International Conference on Machine Learning*, pages 4890–4900. PMLR, 2020.
- [KFS18] Christian Kroer, Gabriele Farina, and Tuomas Sandholm. Solving large sequential games with the excessive gap technique. In *Advances in Neural Information Processing Systems*, pages 864–874, 2018.
- [Kro20] Christian Kroer. Ieor8100: Economics, ai, and optimization lecture note 5: Computing Nash equilibrium via regret minimization. 2020.
- [KWKKS20] Christian Kroer, Kevin Waugh, Fatma Kılınc-Karzan, and Tuomas Sandholm. Faster algorithms for extensive-form game solving via improved smoothing functions. *Mathematical Programming*, pages 1–33, 2020.
- [McM11] Brendan McMahan. Follow-the-regularized-leader and mirror descent: Equivalence theorems and ℓ_1 regularization. In *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*, pages 525–533. JMLR Workshop and Conference Proceedings, 2011.
- [MSB⁺17] Matej Moravčík, Martin Schmid, Neil Burch, Viliam Lisý, Dustin Morrill, Nolan Bard, Trevor Davis, Kevin Waugh, Michael Johanson, and Michael Bowling. Deepstack: Expert-level artificial intelligence in heads-up no-limit poker. *Science*, 356(6337):508–513, 2017.
- [ND16] Hongseok Namkoong and John C Duchi. Stochastic gradient methods for distributionally robust optimization with f -divergences. In *NIPS*, volume 29, pages 2208–2216, 2016.
- [Nem04] Arkadi Nemirovski. Prox-method with rate of convergence $O(1/t)$ for variational inequalities with lipschitz continuous monotone operators and smooth convex-concave saddle point problems. *SIAM Journal on Optimization*, 15(1):229–251, 2004.
- [NY83] Arkadi Nemirovski and David Yudin. *Problem complexity and method efficiency in optimization*. 1983.

- [OP15] Francesco Orabona and Dávid Pál. Scale-free algorithms for online linear optimization. In *International Conference on Algorithmic Learning Theory*, pages 287–301. Springer, 2015.
- [Ora19] Francesco Orabona. A modern introduction to online learning. *arXiv preprint arXiv:1912.13213*, 2019.
- [RM19] Hamed Rahimian and Sanjay Mehrotra. Distributionally robust optimization: A review. *arXiv preprint arXiv:1908.05659*, 2019.
- [RS13] Alexander Rakhlin and Karthik Sridharan. Online learning with predictable sequences. In *Conference on Learning Theory*, pages 993–1019. PMLR, 2013.
- [SALS15] Vasilis Syrgkanis, Alekh Agarwal, Haipeng Luo, and Robert E Schapire. Fast convergence of regularized learning in games. *arXiv preprint arXiv:1507.00407*, 2015.
- [Shi16] Nahum Shimkin. An online convex optimization approach to blackwell’s approachability. *The Journal of Machine Learning Research*, 17(1):4434–4456, 2016.
- [ST18] Aaron Sidford and Kevin Tian. Coordinate methods for accelerating l-infinity regression and faster approximate maximum flow. In *2018 IEEE 59th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 922–933. IEEE, 2018.
- [TBJB15] Oskari Tammelin, Neil Burch, Michael Johanson, and Michael Bowling. Solving heads-up limit Texas hold’em. In *Twenty-Fourth International Joint Conference on Artificial Intelligence*, 2015.
- [Tse95] Paul Tseng. On linear convergence of iterative methods for the variational inequality problem. *Journal of Computational and Applied Mathematics*, 60(1-2):237–252, 1995.
- [vS96] Bernhard von Stengel. Efficient computation of behavior strategies. *Games and Economic Behavior*, 14(2):220–246, 1996.
- [WKR13] W. Wiesemann, D. Kuhn, and B. Rustem. Robust Markov decision processes. *Operations Research*, 38(1):153–183, 2013.
- [WLZL20] Chen-Yu Wei, Chung-Wei Lee, Mengxiao Zhang, and Haipeng Luo. Linear last-iterate convergence in constrained saddle-point optimization. In *International Conference on Learning Representations*, 2020.
- [ZJP07] Martin Zinkevich, Michael Johanson, Michael Bowling, and Carmelo Piccione. Regret minimization in games with incomplete information. In *Advances in neural information processing systems*, pages 1729–1736, 2007.