
Shuffling Gradient-Based Methods for Nonconvex-Concave Minimax Optimization

Quoc Tran-Dinh

Department of Statistics and Operations Research
The University of North Carolina at Chapel Hill
quoctd@email.unc.edu

Trang H. Tran

School of OR and Information Engineering
Cornell University, Ithaca, NY
htt27@cornell.edu

Lam M. Nguyen

IBM Research, Thomas J. Watson Research Center
Yorktown Heights, NY
LamNguyen.MLTD@ibm.com

Abstract

This paper aims at developing novel shuffling gradient-based methods for tackling two classes of minimax problems: *nonconvex-linear* and *nonconvex-strongly concave* settings. The first algorithm addresses the nonconvex-linear minimax model and achieves the state-of-the-art oracle complexity typically observed in nonconvex optimization. It also employs a new shuffling estimator for the “hyper-gradient”, departing from standard shuffling techniques in optimization. The second method consists of two variants: *semi-shuffling* and *full-shuffling* schemes. These variants tackle the nonconvex-strongly concave minimax setting. We establish their oracle complexity bounds under standard assumptions, which, to our best knowledge, are the best-known for this specific setting. Numerical examples demonstrate the performance of our algorithms and compare them with two other methods. Our results show that the new methods achieve comparable performance with SGD, supporting the potential of incorporating shuffling strategies into minimax algorithms.

1 Introduction

Minimax problems arise in various applications across generative machine learning, game theory, robust optimization, online learning, and reinforcement learning (e.g., [1, 2, 3, 5, 12, 13, 17, 19, 21, 25, 35, 40]). These models often involve stochastic settings or large finite-sum objective functions. To tackle these problems, existing methods frequently adapt stochastic gradient descent (SGD) principles to develop algorithms for solving the underlying minimax problems [4, 13]. For instance, in generative adversarial networks (GANs), early algorithms employed stochastic gradient descent-ascent methods where two routines, each using an SGD loop, ran iteratively [13]. However, practical implementations of SGD often incorporate shuffling strategies, as seen in popular deep learning libraries like TensorFlow and PyTorch. This has motivated recent research on developing shuffling techniques specifically for optimization algorithms [4, 5, 8, 16, 26, 32, 38]. Our work builds upon this trend by developing shuffling methods for two specific classes of minimax problems.

Problem statement. In this paper, we study the following minimax optimization problem:

$$\min_{w \in \mathbb{R}^p} \max_{u \in \mathbb{R}^q} \left\{ \mathcal{L}(w, u) := f(w) + \mathcal{H}(w, u) - h(u) \equiv f(w) + \frac{1}{n} \sum_{i=1}^n \mathcal{H}_i(w, u) - h(u) \right\}, \quad (1)$$

where $f : \mathbb{R}^p \rightarrow \mathbb{R} \cup \{+\infty\}$ is a proper, closed, and convex function, $\mathcal{H}_i : \mathbb{R}^p \times \mathbb{R}^q \rightarrow \mathbb{R}$ are smooth for all $i \in [n] := \{1, 2, \dots, n\}$, and $h : \mathbb{R}^q \rightarrow \mathbb{R} \cup \{+\infty\}$ is also a proper, closed, and convex function. In this paper, we will focus on two classes of problems in (1), overlapped to each other.

- (NL) $\mathcal{H}_i$ is nonconvex in w and linear in u as $\mathcal{H}_i(w, u) := \langle F_i(w), Ku \rangle$ for a given function $F_i : \mathbb{R}^p \rightarrow \mathbb{R}^m$ and a matrix $K \in \mathbb{R}^{q \times m}$ for all $i \in [n]$ and $(w, u) \in \text{dom}(\mathcal{L})$.
- (NC) $\mathcal{H}_i$ is nonconvex in w and $\mathcal{H}_i(w, \cdot) - h(\cdot)$ is strongly concave in u for all $(w, u) \in \text{dom}(\mathcal{L})$.

Although (NC) looks more general than (NL), both cases can be overlapped, but one is not a special case of the other. Under these two settings, our approach will rely on a *bilevel optimization* approach, where the lower-level problem is to solve $\max_u \mathcal{L}(w, u)$, while the upper-level one is $\min_w \mathcal{L}(w, u)$.

Challenges. The setting (NL) is a special case of stochastic nonconvex-concave minimax problems because the objective term $\mathcal{H}(w, u) := \langle F(w), Ku \rangle$ is linear in u . It is equivalent to the compositional model (CO) described below. However, if h is only merely convex and not strongly convex (e.g., the indicator of a standard simplex), then Φ_0 in (CO) becomes nonsmooth regardless of F 's properties. This presents our first challenge. A natural approach to address this issue, as discussed in Section 2, is to smooth Φ_0 . The second challenge arises from the composition between the outer function h^* and the finite sum $F(\cdot)$ in (CO). Unlike standard finite-sum optimization, this composition prevents any direct use of existing techniques, requiring a novel approach for algorithmic development and analysis. The third challenge involves unbiased estimators for gradients or “hyper-gradients” in minimax problems. Most existing methods rely on unbiased estimators for objective gradients, with limited work exploring biased estimators. While biased estimators can be used, they require variance reduction properties (see, e.g., [10]). The setting (NC) faces the same second and third challenges as the setting (NL). Additionally, when reformulating it as a minimization problem using a bilevel optimization approach (3), constructing a shuffling estimator for the “hyper-gradient” $\nabla \Phi_0$ becomes unclear. This requires solving the lower-level maximization problem (2). Therefore, it remains an open question whether shuffling gradient-type methods can be extended to this bilevel optimization approach to address (1). In this paper, we address the following research question:

Can we efficiently develop shuffling gradient methods to solve (1) for both (NL) and (NC) settings?

Our attempt to tackle this question leads to a novel way of constructing shuffling estimators for the hyper-gradient $\nabla \Phi_0$ or its smoothed counterpart. This allows us to develop two shuffling gradient-based algorithms with rigorous theoretical guarantees on oracle complexity, matching state-of-the-art complexity results in shuffling-type algorithms for nonconvex optimization.

Related work. Shuffling optimization algorithms have gained significant attention in optimization and machine communities, demonstrating advantages over standard SGDs, see, e.g., [4, 5, 8, 16, 26, 32, 38]. Nevertheless, applying these techniques to minimax problems like (1) remains challenging, with limited existing literature (e.g., [3, 8, 11]). Das *et al.* in [8] explored a specific case of (1) without nonsmooth terms f and h , assuming strong monotonicity and L -Lipschitz continuity of the gradient $\nabla \mathcal{H} := [\nabla_w \mathcal{H}, -\nabla_u \mathcal{H}]$ of the joint objective $\mathcal{H}$. Their algorithm simplifies to a shuffling variant of fixed-point iteration or a gradient descent-ascent scheme, not applicable to our settings. Cho and Yun in [3] built upon [8] by relaxing the strong monotonicity to Polyak-Łojasiewicz (PL) conditions. This work is perhaps the most closely related one to our algorithm, Algorithm 2, for the (NC) setting. Note that the method in [3] exploits Nash’s equilibrium perspective with a simultaneous update, which is different from our alternative update. Moreover, [3] only considers the noncomposite case with $f = 0$ and $h = 0$. Though we only focus on a nonconvex-strongly-concave setting (NC), our results here can be extended to the PL condition as in [3]. Very recently, Konstantinos *et al.* in [11] introduced shuffling extragradient methods for variational inequalities, which encompass convex-concave minimax problems as a special case. However, this also falls outside the scope of our work due to the nonconvexity of (1) in w . Again, all the existing works in [3, 8, 11] utilize a Nash’s equilibrium perspective, while ours leverages a bilevel optimization technique. Besides, in contrast to our sampling-without-replacement approach, stochastic and randomized methods (i.e. using i.i.d. sampling strategies) have been extensively studied for minimax problems, see, e.g., [9, 14, 15, 18, 22, 23, 31, 37, 42]. A comprehensive comparison can be found, e.g., in [3].

Contribution. Our main contribution can be summarized as follows.

- (a) For setting (NL), we suggest to reformulate (1) into a compositional minimization and exploit a smoothing technique to treat this reformulation. We propose a new way of constructing shuffling estimators for the “hyper-gradient” $\nabla \Phi_\gamma$ (cf. (10)) and establish their properties.

- (b) We propose a novel shuffling gradient-based algorithm (*cf.* Algorithm 1) to approximate an ϵ -KKT point of (1) for the setting (NL). Our method requires $\mathcal{O}(n\epsilon^{-3})$ evaluations of F_i and ∇F_i under the strong convexity of h , and $\mathcal{O}(n\epsilon^{-7/2})$ evaluations of F_i and ∇F_i without the strong convexity of h , for a desired accuracy $\epsilon > 0$.
- (c) For setting (NC), we develop two variants of the shuffling gradient method: *semi-shuffling* and *full-shuffling* schemes (*cf.* Algorithm 2). The semi-shuffling variant combines both gradient ascent and shuffling gradient methods to construct a new algorithm, which requires $\mathcal{O}(n\epsilon^{-3})$ evaluations of both $\nabla_w \mathcal{H}_i$ and $\nabla_u \mathcal{H}_i$. The full-shuffling scheme allows to perform both shuffling schemes on the maximization and the minimization alternatively, requiring either $\mathcal{O}(n\epsilon^{-3})$ or $\mathcal{O}(n\epsilon^{-4})$ evaluations of $\nabla_u \mathcal{H}_i$ depending on our assumptions, while maintaining $\mathcal{O}(n\epsilon^{-3})$ evaluations of $\nabla_w \mathcal{H}_i$ for a given desired accuracy $\epsilon > 0$.

If a random shuffling strategy is used in our algorithms, then the oracle complexity in all the cases presented above is improved by a factor of $\sqrt{n}$. Our settings (NL) and (NC) of (1) are different from existing works [3, 8, 11], as we work with general nonconvexity in w , and linearity or [strong] concavity in u , and both f and h are possibly nonsmooth. Our algorithms are not reduced or similar to existing shuffling methods for optimization, but we use shuffling strategies to form estimators for the hyper-gradient $\nabla \Phi_0$ in (5). The oracle complexity in both settings (NL) and (NC) is similar to the ones in nonconvex optimization and in a special case of (1) from [3] (up to a constant factor).

Paper outline. The rest of this paper is organized as follows. Section 2 presents our bilevel optimization approach to (1) and recalls necessary preliminary results. Section 3 develops our shuffling algorithm to solve the setting (NL) of (1) and establishes its convergence. Section 4 proposes new shuffling methods to solve the setting (NC) and investigates their convergence. Section 5 presents numerical experiments, while technical proofs and supporting results are deferred to Supp. Docs.

Notations. For a function f , we use $\text{dom}(f)$ to denote its effective domain, and ∇f for its gradient or Jacobian. If f is convex, then ∇f denotes a subgradient, ∂f is its subdifferential, and prox_f is its proximal operator. We use $\mathcal{F}_t$ to denote $\sigma(w_0, w_1, \dots, w_t)$, a σ -algebra generated by random vectors $w_0, w_1, \dots, w_t$, $\mathbb{E}_t[\cdot] = \mathbb{E}[\cdot | \mathcal{F}_t]$ is a conditional expectation, and $\mathbb{E}[\cdot]$ is the full expectation. As usual, $\mathcal{O}(\cdot)$ denotes Big-O notation in the theory of algorithm complexity.

2 Bilevel Optimization Approach and Preliminary Results

Our approach relies on a bilevel optimization technique [9] in contrast to Nash's game viewpoint [24], which treats the maximization as a lower level and the minimization as an upper level problem.

2.1 Bilevel optimization approach

The minimax model (1) is split into a *lower-level (i.e. a follower) maximization problem* of the form:

$$\begin{aligned}\Phi_0(w) &:= \max_{u \in \mathbb{R}^q} \{ \mathcal{H}(w, u) - h(u) \equiv \frac{1}{n} \sum_{i=1}^n \mathcal{H}_i(w, u) - h(u) \}, \\ u_0^*(w) &:= \arg\max_{u \in \mathbb{R}^q} \{ \mathcal{H}(w, u) - h(u) \equiv \frac{1}{n} \sum_{i=1}^n \mathcal{H}_i(w, u) - h(u) \}.\end{aligned}\quad (2)$$

For Φ_0 defined by (2), then the *upper-level (i.e. the leader) minimization problem* can be written as

$$\Psi_0^* := \min_{w \in \mathbb{R}^p} \{ \Psi_0(w) := \Phi_0(w) + f(w) \}.\quad (3)$$

Clearly, this approach is sequential, and only works if Φ_0 is well-defined, i.e. (2) is globally solvable. Hence, the concavity of $\mathcal{H}(w, \cdot) - h(\cdot)$ w.r.t. to u is crucial for this approach as stated below. However, this assumption can be relaxed to a global solvability of (2) combined with a PL condition as in [3].

Assumption 1 (Basic). *Problems (1) and (3) satisfy the following assumptions for all $i \in [n]$:*

- (a) $\Psi_0^* := \inf_w \Psi_0(w) > -\infty$.
- (b) $\mathcal{H}_i$ is differentiable w.r.t. $(w, u) \in \text{dom}(\mathcal{L})$ and $\mathcal{H}_i(w, \cdot)$ is concave in u for any w .
- (c) Both $f : \mathbb{R}^p \rightarrow \mathbb{R} \cup \{+\infty\}$ and $h : \mathbb{R}^q \rightarrow \mathbb{R} \cup \{+\infty\}$ are proper, closed, and convex.

This assumption remains preliminary. To develop our algorithms, we will need more conditions on $\mathcal{H}_i$ and possibly on f and h , which will be stated later. In addition, we can work with a sublevel set

$$\mathcal{L}_{\Psi_0}(w_0) := \{w \in \text{dom}(\Psi_0) : \Psi_0(w) \leq \Psi_0(w_0)\} \quad (4)$$

of Ψ_0 for a given initial point w_0 from our methods. If $u_0^*(w)$ is uniquely well-defined for given $w \in \mathcal{L}_{\Psi_0}(w_0)$, then by the well-known Danskin's theorem, Φ_0 is differential at w and its gradient is

$$\nabla \Phi_0(w) = \nabla_w \mathcal{H}(w, u_0^*(w)) = \frac{1}{n} \sum_{i=1}^n \nabla_w \mathcal{H}_i(w, u_0^*(w)). \quad (5)$$

We adopt the term “hyper-gradient” from bilevel optimization to name $\nabla \Phi_0$ in this paper.

2.2 Technical assumptions and properties of Φ_0 for nonconvex-linear setting (NL)

(a) **Compositional minimization formulation.** If $\mathcal{H}_i(w, u) := \langle F_i(w), Ku \rangle$ as in setting (NL), then (1) is equivalently reformulated into the following *nonconvex compositional minimization* problem:

$$\min_{w \in \mathbb{R}^p} \left\{ \Psi_0(w) := f(w) + \Phi_0(w) = f(w) + h^* \left(\frac{1}{n} \sum_{i=1}^n K^T F_i(w) \right) \right\}, \quad (\text{CO})$$

where $h^*(v) := \sup_u \{ \langle v, u \rangle - h(u) \}$, the Fenchel conjugate of h , and $\Phi_0(w) = h^*(K^T F(w))$. If h is not strongly convex, then h^* is convex but possibly nonsmooth.

(b) **Technical assumptions.** To develop our algorithms, we also need the following assumptions.

Assumption 2. h is μ_h -strongly convex with $\mu_h \geq 0$, and $\text{dom}(h)$ is bounded by $M_h < +\infty$.

Assumption 3 (For F_i). For setting (NL) with $\mathcal{H}_i(w, u) := \langle F_i(w), Ku \rangle$ ($i \in [n]$), assume that

- (a) F_i is continuously differentiable, and its Jacobian ∇F_i is L_{F_i} -Lipschitz continuous.
- (b) F_i is also M_{F_i} -Lipschitz continuous or equivalently, its Jacobian ∇F_i is M_{F_i} -bounded.
- (c) There exists a positive constant $\sigma_J \in (0, +\infty)$ such that

$$\frac{1}{n} \sum_{i=1}^n \|\nabla F_i(w) - \nabla F(w)\|^2 \leq \sigma_J^2, \quad \forall w \in \text{dom}(F). \quad (6)$$

Assumption 2 allows $\mu_h = 0$ that also covers the non-strong convexity of h . Assumption 3 is rather standard to develop gradient-based methods for solving (1). Under Assumption 3, the finite-sum F is also M_F -Lipschitz continuous and the Jacobian ∇F of F is also L_F -Lipschitz continuous with

$$M_F := \max\{M_{F_i} : i \in [n]\} \quad \text{and} \quad L_F := \max\{L_{F_i} : i \in [n]\}. \quad (7)$$

Condition (6) can be relaxed to the form $\frac{1}{n} \sum_{i=1}^n \|\nabla F_i(w) - \nabla F(w)\|^2 \leq \sigma_J^2 + \Theta_J \|\nabla \Phi_0(w)\|^2$ for some $\Theta_J \geq 0$, where $\nabla \Phi_0$ is a [sub]gradient of Φ_0 or Φ_γ (its smoothed approximation). Moreover, under Assumption 3, if $\mu_h > 0$, then ∇h^* is L_{h^*} -Lipschitz continuous with $L_{h^*} := \frac{1}{\mu_h}$. Thus it is possible (see [9]) to prove that Φ_0 is differentiable, and $\nabla \Phi_0$ is also L_{Φ_0} -Lipschitz continuous with $L_{\Phi_0} := M_h \|K\| L_F + \frac{M_F^2 \|K\|^2}{\mu_h}$ as a consequence of Lemma 4 when $\gamma \downarrow 0^+$ in Supp. Doc. A.

(c) **Smoothing technique for lower-level maximization problem (2).** If h is only merely convex (i.e. $\mu_h = 0$), then (2) may not be uniquely solvable, leading to the possible non-differentiability of Φ_0 . Let us define the following convex function:

$$\phi_0(v) := \max_{u \in \mathbb{R}^q} \{ \langle v, Ku \rangle - h(u) \} = h^*(K^T v). \quad (8)$$

Then, Φ_0 in (2) or (CO) can be written as $\Phi_0(w) = \phi_0(F(w)) = \phi_0 \left(\frac{1}{n} \sum_{i=1}^n F_i(w) \right)$. Our goal is to smooth ϕ_0 if h is not strongly convex, leading to

$$\begin{cases} \phi_\gamma(v) := \max_u \{ \langle v, Ku \rangle - h(u) - \gamma b(u) \}, \\ u_\gamma^*(v) := \operatorname{argmax}_u \{ \langle v, Ku \rangle - h(u) - \gamma b(u) \}, \end{cases} \quad (9)$$

where $\gamma > 0$ is a given smoothness parameter and $b : \mathbb{R}^q \rightarrow \mathbb{R}$ is a proper, closed, and 1-strongly convex function such that $\text{dom}(h) \subseteq \text{dom}(b)$. We also denote $D_b := \sup\{\|\nabla b(u)\| : u \in \text{dom}(h)\}$. In particular, if we choose $b(u) := \frac{1}{2} \|u - \bar{u}\|^2$ for a fixed $\bar{u}$, then $u_\gamma^*(v) = \text{prox}_{h/\gamma}(\bar{u} - K^T v)$.

Using ϕ_γ , problem (CO) can be approximated by its smoothed formulation:

$$\min_{w \in \mathbb{R}^p} \left\{ \Psi_\gamma(w) := f(w) + \Phi_\gamma(w) = f(w) + \phi_\gamma(F(w)) \equiv f(w) + \phi_\gamma \left(\frac{1}{n} \sum_{i=1}^n F_i(w) \right) \right\}. \quad (10)$$

To develop our method, one key step is to approximate the hyper-gradient of Φ_γ in (10), where

$$\nabla \Phi_\gamma(w) = \nabla F(w)^T \nabla \phi_\gamma(F(w)) = \frac{1}{n} \sum_{i=1}^n \nabla F_i(w)^T \nabla \phi_\gamma(F(w)). \quad (11)$$

Then, $\nabla \Phi_\gamma$ is L_{Φ_γ} -Lipschitz continuous with $L_{\Phi_\gamma} := M_h \|K\| L_F + \frac{M_F^2 \|K\|^2}{\mu_h + \gamma}$ (see Lemma 4).

2.3 Technical assumptions and properties of Φ_0 for the nonconvex-strongly-concave setting

To develop our shuffling gradient-based algorithms for solving (1) under the nonconvex-strongly-concave setting (NC), we impose the following assumptions.

Assumption 4 (For $\mathcal{H}_i$). $\mathcal{H}_i$ for all $i \in [n]$ in (1) satisfies the following conditions:

- (a) For any given w such that $(w, u) \in \text{dom}(\mathcal{H})$, $\mathcal{H}_i(w, \cdot)$ is μ_H -strongly concave w.r.t. u .
- (b) $\nabla \mathcal{H}_i$ is (L_w, L_u) -Lipschitz continuous, i.e. for all $(w, u), (\hat{w}, \hat{u}) \in \text{dom}(\mathcal{H})$:

$$\|\nabla \mathcal{H}_i(w, u) - \nabla \mathcal{H}_i(\hat{w}, \hat{u})\|^2 \leq L_w^2 \|w - \hat{w}\|^2 + L_u^2 \|u - \hat{u}\|^2. \quad (12)$$

- (c) There exist two constants $\Theta_w \geq 0$ and $\sigma_w \geq 0$ such that for $(w, u) \in \text{dom}(\mathcal{H})$, we have

$$\frac{1}{n} \sum_{i=1}^n \|\nabla_w \mathcal{H}_i(w, u) - \nabla_w \mathcal{H}(w, u)\|^2 \leq \Theta_w \|\nabla_w \mathcal{H}(w, u)\|^2 + \sigma_w^2. \quad (13)$$

There exist two constants $\Theta_u \geq 0$ and $\sigma_u \geq 0$ such that for all $(w, u) \in \text{dom}(\mathcal{H})$, we have

$$\frac{1}{n} \sum_{i=1}^n \|\nabla_u \mathcal{H}_i(w, u) - \nabla_u \mathcal{H}(w, u)\|^2 \leq \Theta_u \|\nabla_u \mathcal{H}(w, u)\|^2 + \sigma_u^2. \quad (14)$$

Assumption 4(a) makes sure that our lower-level maximization of (1) is well-defined. Assumption 4(b) and (c) are standard in shuffling gradient-type methods as often seen in nonconvex optimization [9].

Lemma 1 (Smoothness of Φ_0). Under Assumptions 2 and 4, $u_0^*(\cdot)$ in (2) is κ -Lipschitz continuous with $\kappa := \frac{L_u}{\mu_H + \mu_h}$. Moreover, $\nabla \Phi_0$ in (5) is L_{Φ_0} -Lipschitz continuous with $L_{\Phi_0} := (1 + \kappa)L_w$.

2.4 Approximate KKT points and approximate stationary points

(a) **Exact and approximate KKT points and stationary points.** A pair $(w^*, u^*) \in \text{dom}(\mathcal{L})$ is called a KKT (Karush-Kuhn-Tucker) point of (1) if

$$0 \in \nabla_w \mathcal{H}(w^*, u^*) + \partial f(w^*) \quad \text{and} \quad 0 \in -\nabla_u \mathcal{H}(w^*, u^*) + \partial h(u^*). \quad (15)$$

Given a tolerance $\epsilon > 0$, our goal is to find an ϵ -approximate KKT point $(\hat{w}, \hat{u})$ of (1) defined as

$$r_w \in \nabla_w \mathcal{H}(\hat{w}, \hat{u}) + \partial f(\hat{w}), \quad r_u \in -\nabla_u \mathcal{H}(\hat{w}, \hat{u}) + \partial h(\hat{u}), \quad \text{and} \quad \mathbb{E}[\|r_w, r_u\|^2] \leq \epsilon^2. \quad (16)$$

A vector $w^* \in \text{dom}(\Psi_0)$ is said to be a stationary point of (3) if

$$0 \in \nabla \Phi_0(w^*) + \partial f(w^*). \quad (17)$$

Since f is possibly nonsmooth, we can define a stationary point of (3) via a gradient mapping as:

$$\mathcal{G}_\eta(w) := \eta^{-1}(w - \text{prox}_{\eta f}(w - \eta \nabla \Phi_0(w))), \quad (18)$$

where $\eta > 0$ is given. It is well-known that $\mathcal{G}_\eta(w^*) = 0$ iff w^* is a stationary point of (3). Again, since we cannot exactly compute w^* , we expect to find an ϵ -stationary point $\hat{w}_T$ of (3) such that $\mathbb{E}[\|\mathcal{G}_\eta(\hat{w}_T)\|^2] \leq \epsilon^2$ for a given tolerance $\epsilon > 0$.

(b) **Constructing an approximate stationary point and KKT point from algorithms.** Our algorithms below generate a sequence $\{\tilde{w}_t\}_{t=0}^T$ such that $\frac{1}{T+1} \sum_{t=0}^T \mathbb{E}[\|\mathcal{G}_\eta(\tilde{w}_t)\|^2] \leq \epsilon^2$. Hence, we construct an ϵ -stationary point $\hat{w}_T$ using one of the following two options:

$$\hat{w}_T := \tilde{w}_{t_*}, \quad \text{where} \quad \begin{cases} t_* := \arg\min\{\|\mathcal{G}_\eta(\tilde{w}_t)\| : 0 \leq t \leq T\}, & \text{(Option 1) or} \\ t_* \text{ is uniformly randomly chosen from } \{0, 1, \dots, T\} & \text{(Option 2).} \end{cases} \quad (19)$$

Clearly, we have $\mathbb{E}[\|\mathcal{G}_\eta(\hat{w}_T)\|^2] \leq \frac{1}{T+1} \sum_{t=0}^T \mathbb{E}[\|\mathcal{G}_\eta(\tilde{w}_t)\|^2] \leq \epsilon^2$. We need the following result.

Lemma 2. (a) If (w^*, u^*) is a KKT point of (1), then w^* is a stationary point of (3). Conversely, if w^* is a stationary point of (3), then $(w^*, u_0^*(w^*))$ is a KKT point of (1).

(b) If $\hat{w}_T$ is an ϵ -stationary point of (3) and $\nabla \Phi_0$ is L_{Φ_0} -Lipschitz continuous, then $(\bar{w}_T, \bar{u}_T)$ is an $\hat{\epsilon}$ -KKT point of (1), where $\bar{w}_T := \text{prox}_{\eta f}(\hat{w}_T - \eta \nabla \Phi_0(\hat{w}_T))$, $\bar{u}_T := u_0^*(\bar{w}_T)$, and $\hat{\epsilon} := (1 + L_{\Phi_0}\eta)\epsilon$.

(c) If $\hat{w}_T$ is an ϵ -stationary point of (10), then $(\bar{w}_T, \bar{u}_T)$ is an $\hat{\epsilon}$ -KKT point of (1), where $\bar{w}_T := \text{prox}_{\eta f}(\hat{w}_T - \eta \nabla \Phi_\gamma(\hat{w}_T))$, $\bar{u}_T := u_\gamma^*(F(\bar{w}_T))$, and $\hat{\epsilon} := \max\{(1 + L_{\Phi_\gamma}\eta)\epsilon, \gamma D_b\}$.

Lemma 2 allows us to construct an $\hat{\epsilon}$ -approximate KKT point $(\bar{w}_T, \bar{u}_T)$ of (1) from an ϵ -stationary point $\hat{w}_T$ of either (3) or its smoothed problem (10), where $\hat{\epsilon} = \mathcal{O}(\max\{\epsilon, \gamma\})$.

2.5 Technical condition to handle the possible nonsmooth term f

To handle the nonsmooth term f of (1) in our algorithms we require one more condition as in [5].

Assumption 5. Let Φ_γ be defined by (10), which reduces to Φ_0 given by (2) as $\gamma \downarrow 0^+$, and $\mathcal{G}_\eta$ be defined by (18). Assume that there exist two constants $\Lambda_0 \geq 1$ and $\Lambda_1 \geq 0$ such that:

$$\|\nabla\Phi_\gamma(w)\|^2 \leq \Lambda_0\|\mathcal{G}_\eta(w)\|^2 + \Lambda_1, \quad \forall w \in \text{dom}(\Phi_0). \quad (20)$$

If $f = 0$, then $\mathcal{G}_\eta(w) \equiv \nabla\Phi_\gamma(w)$, and Assumption 5 automatically holds with $\Lambda_0 = 1$ and $\Lambda_1 = 0$. If $f \neq 0$, then it is crucial to have $\Lambda_0 \geq 1$ in (20). Let us consider two examples to see why?

- (i) If f is M_f -Lipschitz continuous (e.g., ℓ_1 -norm), then (20) also holds with $\Lambda_0 := 1 + \nu > 1$ and $\Lambda_1 := \frac{1+\nu}{\nu} M_f$ for a given $\nu > 0$.
- (ii) If $f = \delta_{\mathcal{W}}$, the indicator of a nonempty, closed, convex, and bounded set $\mathcal{W}$, then Assumption 5 also holds by the same reason as in Example (i) (see Supp. Doc. A).

3 Shuffling Gradient Method for Nonconvex-Linear Minimax Problems

We first propose a new construction using shuffling techniques to approximate the true gradient $\nabla\Phi_\gamma$ in (11) for any $\gamma \geq 0$. Next, we propose our algorithm and analyze its convergence.

3.1 The shuffling gradient estimators for $\nabla\Phi_\gamma$

Challenges. To evaluate $\nabla\Phi_\gamma(w)$ in (11), we need to evaluate both $\nabla F(w)$ and $F(w)$ at each w . However, in SGD or shuffling gradient methods, we want to approximate both quantities at each iteration. Note that this gradient can be written in a finite-sum $\frac{1}{n} \sum_{i=1}^n \nabla F_i(w)^T \nabla \phi_\gamma(F(w))$ (see (11)), but every summand requires $\nabla \phi_\gamma(F(w))$, which involves the full evaluation of F .

Our estimators. Let $F_{\pi^{(t)}(i)}(w_{i-1}^{(t)})$ and $\nabla F_{\hat{\pi}^{(t)}(i)}(w_{i-1}^{(t)})$ be the function value and the Jacobian component evaluated at $w_{i-1}^{(t)}$ respectively for $i \in [n]$, where $\pi^{(t)} = (\pi^{(t)}(1), \pi^{(t)}(2), \dots, \pi^{(t)}(n))$ and $\hat{\pi}^{(t)} = (\hat{\pi}^{(t)}(1), \hat{\pi}^{(t)}(2), \dots, \hat{\pi}^{(t)}(n))$ are two permutations of $[n] := \{1, 2, \dots, n\}$. We want to use these quantities to approximate the function value $F(w_0^{(t)})$ and its Jacobian $\nabla F(w_0^{(t)})$ of F at $w_0^{(t)}$, respectively, where $w_0^{(t)}$ the iterate vector at the beginning of each epoch t .

For function value $F(w_0^{(t)})$, we suggest the following approximation at each *inner iteration* $i \in [n]$:

$$\textbf{Option 1:} \quad F_i^{(t)} := \frac{1}{n} \left[\sum_{j=1}^i F_{\pi^{(t)}(j)}(w_{j-1}^{(t)}) + \sum_{j=i+1}^n F_{\pi^{(t)}(j)}(w_0^{(t)}) \right]. \quad (21)$$

Alternative to (21), for all $i \in [n]$, we can simply choose another option:

$$\textbf{Option 2:} \quad F_i^{(t)} := \frac{1}{n} \sum_{j=1}^n F_j(w_0^{(t)}) = \frac{1}{n} \sum_{j=1}^n F_{\pi^{(t)}(j)}(w_0^{(t)}). \quad (22)$$

For Jacobian $\nabla F(w_0^{(t)})$, we suggest to use the following standard shuffling estimator for all $i \in [n]$:

$$\nabla F_i^{(t)} := \nabla F_{\hat{\pi}^{(t)}(i)}(w_{i-1}^{(t)}). \quad (23)$$

For $F_i^{(t)}$ from (21) (or (22)) and for $\nabla F_i^{(t)}$ from (23), we form an approximation of $\nabla\Phi_\gamma(w_0^{(t)})$ as

$$\tilde{\nabla}\Phi_\gamma(w_{i-1}^{(t)}) := (\nabla F_i^{(t)})^T \nabla \phi_\gamma(F_i^{(t)}) \equiv (\nabla F_i^{(t)})^T K u_\gamma^*(F_i^{(t)}). \quad (24)$$

Discussion. The estimator $F_i^{(t)}$ for F requires $n - i$ more function evaluations $F_{\pi^{(t)}(j)}(w_0^{(t)})$ at each epoch t . The first option (21) for F uses $2n$ function evaluations F_i , while the second one in (22) only needs n function evaluations at each epoch $t \geq 0$. However, (21) uses the most updated information up to the *inner iteration* i compared to (22), which is expected to perform better. The Jacobian estimator $\nabla F_i^{(t)}$ is standard and only uses one sample or a mini-batch at each iteration i .

3.2 The shuffling gradient-type algorithm for nonconvex-linear setting (NL)

We propose Algorithm 1, a shuffling gradient-type method, to approximate a stationary point of (10).

Discussion. First, the cost per epoch of Algorithm 1 consists of either $2n$ or n function evaluations F_i , and n Jacobian evaluations ∇F_i . Compare to standard shuffling gradient-type methods, e.g., in [8], Algorithm 1 has either n more evaluations of F_i or the same cost. Second, when implementing

Algorithm 1 (Shuffling Proximal Gradient-Based Algorithm for Solving (10))

```
1: Initialization: Choose an initial point  $\tilde{w}_0 \in \text{dom}(\Phi_0)$  and a smoothness parameter  $\gamma > 0$ .
2: for  $t = 1, 2, \dots, T$  do
3:   Set  $w_0^{(t)} := \tilde{w}_{t-1}$ ;
4:   Generate two permutations  $\pi^{(t)}$  and  $\hat{\pi}^{(t)}$  of  $[n]$  (identically or randomly and independently)
5:   for  $i = 1, \dots, n$  do
6:     Evaluate  $F_i^{(t)}$  by either (21) or (22) using  $\pi^{(t)}$ , and  $\nabla F_i^{(t)}$  by (23) using  $\hat{\pi}^{(t)}$ .
7:     Solve (9) to get  $u_\gamma^*(F_i^{(t)})$  and form  $\tilde{\nabla}\Phi_\gamma(w_{i-1}^{(t)}) := (\nabla F_i^{(t)})^T K u_\gamma^*(F_i^{(t)})$ .
8:     Update  $w_i^{(t)} := w_{i-1}^{(t)} - \frac{\eta_t}{n} \tilde{\nabla}\Phi_\gamma(w_{i-1}^{(t)})$ ;
9:   end for
10:  Compute  $\tilde{w}_t := \text{prox}_{\eta_t f}(w_n^{(t)})$ ;
11: end for
```

Algorithm 1, we do not need to evaluate the full Jacobian $\nabla F_i^{(t)}$, but rather the product of matrix $(\nabla F_i^{(t)})^T$ and vector $\nabla\Phi_\gamma(F_i^{(t)})$ as $\tilde{\nabla}\Phi_\gamma(w_{i-1}^{(t)}) := (\nabla F_i^{(t)})^T \nabla\Phi_\gamma(F_i^{(t)})$. Evaluating this matrix-vector multiplication is much more efficient than evaluating the full Jacobian $\nabla F_i^{(t)}$ and $\nabla\Phi_\gamma(F_i^{(t)})$ individually. Third, thanks to Assumption 5, the proximal step $\tilde{w}_t := \text{prox}_{\eta_t f}(w_n^{(t)})$ is only required at the end of each epoch t . This significantly reduces the computational cost if $\text{prox}_{\eta_t f}$ is expensive.

3.3 Convergence Analysis of Algorithm 1 for Nonconvex-Linear Setting (NL)

Now, we are ready to state the convergence result of Algorithm 1 in a short version: Theorem 1. The full version of this theorem is Theorem 6, which can be found in Supp. Doc. B.

Theorem 1. Suppose that Assumptions 1, 2, 3, and 5 holds for the setting (NL) of (1) and $\epsilon > 0$ is a sufficiently small tolerance. Let $\{\tilde{w}_t\}$ be generated by Algorithm 1 after $T = \mathcal{O}(\epsilon^{-3})$ epochs using arbitrarily permutations $\pi^{(t)}$ and $\hat{\pi}^{(t)}$ and a learning rate $\eta_t = \eta := \mathcal{O}(\epsilon)$ (see Theorem 6 in Supp. Doc. B for the exact formulas of T and η). Then, we have $\frac{1}{T+1} \sum_{t=0}^T \|\mathcal{G}_{\eta_t}(\tilde{w}_t)\|^2 \leq \epsilon^2$.

Alternatively, if $\{\tilde{w}_t\}$ is generated by Algorithm 1 after $T := \mathcal{O}(n^{-1/2}\epsilon^{-3})$ epochs using two random and independent permutations $\pi^{(t)}$ and $\hat{\pi}^{(t)}$ and a learning rate $\eta_t = \eta := \mathcal{O}(n^{1/2}\epsilon)$ (see Theorem 6 in Supp. Doc. B for the exact formulas). Then, we have $\frac{1}{T+1} \sum_{t=0}^T \mathbb{E}[\|\mathcal{G}_{\eta_t}(\tilde{w}_t)\|^2] \leq \epsilon^2$.

Our first goal is to approximate a stationary point w^* of (CO) as $\mathbb{E}[\|\mathcal{G}_\eta(\hat{w})\|^2] \leq \epsilon^2$, while Algorithm 1 only provides an ϵ -stationary of (10). For a proper choice of γ , it is also an ϵ -stationary point of (3).

Corollary 1. Let $\hat{w}_T$ defined by (19) be generated from $\{\tilde{w}_t\}$ of Algorithm 1. Under the conditions of Theorem 1 and any permutations $\pi^{(t)}$ and $\hat{\pi}^{(t)}$, the following statements hold.

- (a) If h is μ_h -strongly convex with $\mu_h > 0$, then we can set $\gamma = 0$, and Algorithm 1 requires $\mathcal{O}(n\epsilon^{-3})$ evaluations of F_i and ∇F_i to achieve an ϵ -stationary $\hat{w}_T$ of (3).
- (b) If h is only convex (i.e. $\mu_h = 0$), then we can set $\gamma := \mathcal{O}(\epsilon)$, and Algorithm 1 needs $\mathcal{O}(n\epsilon^{-7/2})$ evaluations of F_i and ∇F_i to achieve an ϵ -stationary $\hat{w}_T$ of (3).

If, in addition, $\pi^{(t)}$ and $\hat{\pi}^{(t)}$ are sampled uniformly at random without replacement and independently, and $\Lambda_1 = \mathcal{O}(n^{-1})$, then the numbers of evaluations of F_i and ∇F_i are reduced by a factor of $\sqrt{n}$.

4 Shuffling Method for Nonconvex-Strongly Concave Minimax Problems

In this section, we develop shuffling gradient-based methods to solve (1) under the nonconvex-strongly concave setting (NC). Since this setting does not cover the nonconvex-linear setting (NL) in Section 3 as a special case, we need to treat it separately using different ideas and proof techniques.

4.1 The construction of algorithm

Unlike the linear case with $\mathcal{H}_i(w, u) = \langle F_i(w), Ku \rangle$ in Section 3, we cannot generally compute the solution $u_0^*(\tilde{w}_{t-1})$ in (2) exactly for a given $\tilde{w}_{t-1}$. We can only approximate $u_0^*(\tilde{w}_{t-1})$ by some $\tilde{u}_t$. This leads to another level of inexactness in an approximate “hyper-gradient” $\tilde{\nabla}\Phi_0(w_{i-1}^{(t)})$ defined by

$$\tilde{\nabla}\Phi_0(w_{i-1}^{(t)}) := \nabla_w \mathcal{H}_{\hat{\pi}^{(t)}(i)}(w_{i-1}^{(t)}, \tilde{u}_t). \quad (25)$$

There are different options to approximate $u_0^*(\tilde{w}_{t-1})$. We propose two options below, but other choices are possible, including accelerated gradient ascent methods and stochastic algorithms [6, 20].

(a₁) **Gradient ascent scheme for the lower-level problem.** We apply a standard gradient ascent scheme to update $\tilde{u}_t$: Starting from $s = 0$ with $u_0^{(t)} := \tilde{u}_{t-1}$, at each epoch $s = 1, \dots, S$, we update

$$\hat{u}_s^{(t)} := \text{prox}_{\hat{\eta}_t h}(\hat{u}_{s-1}^{(t)} + \frac{\hat{\eta}_t}{n} \sum_{i=1}^n \nabla_u \mathcal{H}_i(\tilde{w}_{t-1}, \hat{u}_{s-1}^{(t)})), \quad (26)$$

for a given learning rate $\hat{\eta}_t > 0$. Then, we finally output $\tilde{u}_t := \hat{u}_S^{(t)}$ to approximate $u_0^*(\tilde{w}_{t-1})$.

To make our method more flexible, we allow to perform either only *one iteration* (i.e. $S = 1$) or *multiple iterations* (i.e. $S > 1$) of (26). Each iteration s requires n evaluations of $\nabla_u \mathcal{H}_i$.

(a₂) **Shuffling gradient ascent scheme for the lower-level problem.** We can also construct $\tilde{u}_t$ by a *shuffling gradient ascent scheme*. Again, we allow to run either only *one epoch* (i.e. $S = 1$) or *multiple epochs* (i.e. $S > 1$) of the shuffling algorithm to update $\tilde{u}_t$, leading to the following scheme: Starting from $s := 1$ with $\hat{u}_0^{(t)} := \tilde{u}_{t-1}$, at each epoch $s = 1, 2, \dots, S$, having $\hat{u}_{s-1}^{(t)}$, we generate a permutation $\pi^{(s,t)}$ of $[n]$ and run a *shuffling gradient ascent scheme* as

$$\begin{cases} u_0^{(s,t)} := \hat{u}_{s-1}^{(t)}, \\ \text{For } i = 1, 2, \dots, n, \text{ update} \\ \quad u_i^{(s,t)} := u_{i-1}^{(s,t)} + \frac{\hat{\eta}_t}{n} \nabla_u \mathcal{H}_{\pi^{(s,t)}(i)}(\tilde{w}_{t-1}, u_{i-1}^{(s,t)}), \\ \hat{u}_s^{(t)} := \text{prox}_{\hat{\eta}_t h}(u_n^{(s,t)}). \end{cases} \quad (27)$$

At the end of the S -th epoch, we output $\tilde{u}_t := \hat{u}_S^{(t)}$ as an approximation to $u_0^*(\tilde{w}_{t-1})$. Here, we use the same learning rate $\hat{\eta}_t$ for all epochs $s \in [S]$. Each epoch s requires n evaluations of $\nabla_u \mathcal{H}_i$.

(b) **Shuffling gradient descent scheme for the upper-level minimization problem.** Having $\tilde{u}_t$ from either (26) or (27), we run a *shuffling gradient descent epoch* to update $\tilde{w}_t$ from $\tilde{w}_{t-1}$ as

$$\begin{cases} w_0^{(t)} := \tilde{w}_{t-1}, \\ \text{For } i = 1, 2, \dots, n, \text{ update} \\ \quad w_i^{(t)} := w_{i-1}^{(t)} - \frac{\eta_t}{n} \tilde{\nabla} \Phi_0(w_{i-1}^{(t)}) \equiv w_{i-1}^{(t)} - \frac{\eta_t}{n} \nabla_w \mathcal{H}_{\hat{\pi}^{(t)}(i)}(w_{i-1}^{(t)}, \tilde{u}_t), \\ \tilde{w}_t := \text{prox}_{\eta_t f}(w_n^{(t)}). \end{cases} \quad (28)$$

These two steps (26) (or (27)) in u and (28) in w are implemented alternatively for $t = 1, \dots, T$.

(c) **The full algorithm.** Combining both steps (26) (or (27)) and (28), we can present an *alternating shuffling proximal gradient algorithm* for solving (1) as in Algorithm 2.

Algorithm 2 (Alternating Shuffling Proximal Gradient Algorithm for Solving (1) under setting (NC))

- 1: **Initialization:** Choose an initial point $(\tilde{w}_0, \tilde{u}_0) \in \text{dom}(\mathcal{L})$.
 - 2: **for** $t = 1, 2, \dots, T$ **do**
 - 3: Compute $\tilde{u}_t$ using either (26) or (27).
 - 4: Set $w_0^{(t)} := \tilde{w}_{t-1}$ and generate a permutation $\hat{\pi}^{(t)}$ of $[n]$.
 - 5: **for** $i = 1, \dots, n$ **do**
 - 6: Evaluate $\tilde{\nabla} \Phi_0(w_{i-1}^{(t)}) := \nabla_w \mathcal{H}_{\hat{\pi}^{(t)}(i)}(w_{i-1}^{(t)}, \tilde{u}_t)$.
 - 7: Update $w_i^{(t)} := w_{i-1}^{(t)} - \frac{\eta_t}{n} \tilde{\nabla} \Phi_0(w_{i-1}^{(t)})$.
 - 8: **end for**
 - 9: Compute $\tilde{w}_t := \text{prox}_{\eta_t f}(w_n^{(t)})$.
 - 10: **end for**
-

Discussion. Algorithm 2 has a similar form as Algorithm 1, where $u_0^*(\tilde{w}_{t-1})$ is approximated by $\tilde{u}_t$. In Algorithm 1, $u_0^*(\tilde{w}_{t-1})$ is approximated by $u_\gamma^*(F_i^{(t)})$. Moreover, Algorithm 1 solves the smoothed problem (10) of (3), while Algorithm 2 directly solves (3). Depending on the choice of method to approximate $u_0^*(\tilde{w}_{t-1})$, we obtain different variants of Algorithm 2. We have proposed two variants:

- **Semi-shuffling variant:** We use (26) for computing $\tilde{u}_t$ to approximate $u_0^*(\tilde{w}_{t-1})$.
- **Full-shuffling variant:** We use (27) for computing $\tilde{u}_t$ to approximate $u_0^*(\tilde{w}_{t-1})$.

Note that Algorithm 2 works in an alternative manner, where it approximates $u_0^*(\tilde{w}_{t-1})$ up to a certain accuracy before updating $\tilde{w}_t$. This alternating update is very natural and has been widely applied to solve minimax optimization as well as bilevel optimization problems, see, e.g., [1, 9, 13].

4.2 Convergence analysis

Now, we state the convergence of both variants of Algorithm 2: *semi-shuffling* and *full-shuffling* variants. The full proof of the following theorems can be found in Supp. Doc. C.

(a) **Convergence of the semi-shuffling variant.** Our first result is as follows.

Theorem 2. Suppose that Assumptions 1, 2, 4, and 5 hold for (1), and $\mathcal{G}_\eta$ is defined by (18).

Let $\{(\tilde{w}_t, \tilde{u}_t)\}$ be generated by Algorithm 2 using the **gradient ascent scheme** (26) with $\eta := \mathcal{O}(\epsilon)$ explicitly given in Theorem 8 of Supp. Doc. C, $\hat{\eta} \in (0, \frac{2}{L_u + \mu_h}]$, $S := \mathcal{O}(\frac{1}{\hat{\eta}}(\mu_h + \frac{4L_u\mu_H}{L_u + \mu_H})^{-1}) = \mathcal{O}(1)$, and $T := \mathcal{O}(\epsilon^{-3})$ explicitly given in Theorem 8. Then, we have $\frac{1}{T+1} \sum_{t=0}^T \|\mathcal{G}_\eta(\tilde{w}_t)\|^2 \leq \epsilon^2$.

Consequently, Algorithm 2 requires $\mathcal{O}(n\epsilon^{-3})$ evaluations of both $\nabla_w \mathcal{H}_i$ and $\nabla_u \mathcal{H}_i$ to achieve an ϵ -stationary point $\hat{w}_T$ of (3) computed by (19).

Note that Theorem 2 holds for both $S > 1$ and $S = 1$ (i.e. we perform only one iteration of (26)).

(b) **Convergence of the full-shuffling variant – The case $S > 1$ with multiple epochs.** We state our results for two separated cases: only $\mathcal{H}_i$ is μ_H -strongly convex, and only h is μ_h -strongly convex.

Theorem 3 (Strong convexity of $\mathcal{H}_i$). Suppose that Assumptions 1, 2, 4, and 5 hold, and $\mathcal{H}_i$ is μ_H -strongly concave with $\mu_H > 0$ for $i \in [n]$, but h is only merely convex.

Let $\{(\tilde{w}_t, \tilde{u}_t)\}$ be generated by Algorithm 2 using S epochs of the **shuffling routine** (27) and fixed learning rates $\eta_t = \eta := \mathcal{O}(\epsilon)$ as given in Theorem 8 of Supp. Doc. C for a given $\epsilon > 0$, $\hat{\eta}_t := \hat{\eta} = \mathcal{O}(\epsilon)$, $S := \lfloor \frac{\ln(7/2)}{\mu_H \hat{\eta}} \rfloor$, and $T := \mathcal{O}(\epsilon^{-3})$. Then, we have $\frac{1}{T+1} \sum_{t=0}^T \|\mathcal{G}_\eta(\tilde{w}_t)\|^2 \leq \epsilon^2$.

Consequently, Algorithm 2 requires $\mathcal{O}(n\epsilon^{-3})$ evaluations of $\nabla_w \mathcal{H}_i$ and $\mathcal{O}(n\epsilon^{-4})$ evaluations of $\nabla_u \mathcal{H}_i$ to achieve an ϵ -stationary point $\hat{w}_T$ of (3) computed by (19).

Theorem 4 (Strong convexity of h). Suppose that Assumptions 1, 2, 4, and 5 hold for (1), and h is μ_h -strongly convex with $\mu_h > 0$, but $\mathcal{H}_i$ is only merely concave for all $i \in [n]$. Then, under the same settings as in Theorem 3, but with $S := \lfloor \frac{\ln(7/2)}{\mu_h \hat{\eta}} \rfloor$, the conclusions of Theorem 3 still hold.

(c) **Convergence of the full-shuffling variant – The case $S = 1$ with one epoch.** Both Theorems 3 and 4 require $\mathcal{O}(n\epsilon^{-4})$ evaluations of $\nabla_u \mathcal{H}_i$. To improve this complexity, we need two additional assumptions but can perform only one epoch of (27), i.e. $S = 1$.

Assumption 6. Let $\hat{\mathcal{G}}_\eta(u) := \eta^{-1}(u - \text{prox}_{\eta h}(u + \eta \nabla_u \mathcal{H}(w, u)))$ be the gradient mapping of $\psi(w, \cdot) := -\mathcal{H}(w, \cdot) + h(\cdot)$. Assume that there exist $\hat{\Lambda}_0 \geq 1$ and $\hat{\Lambda}_1 \geq 0$ such that

$$\|\nabla_u \mathcal{H}(w, u)\|^2 \leq \hat{\Lambda}_0 \|\hat{\mathcal{G}}_\eta(u)\|^2 + \hat{\Lambda}_1, \quad \forall (w, u) \in \text{dom}(\mathcal{L}). \quad (29)$$

Clearly, if $h = 0$, then $\hat{\mathcal{G}}_\eta(u) = -\nabla_u \mathcal{H}(w, u)$ and (20) automatically holds for $\hat{\Lambda}_0 = 1$ and $\hat{\Lambda}_1 = 0$. Assumption 6 is similar to Assumption 5, and it is required to handle the prox operator of h in (27).

Assumption 7. For f in (1), there exists $L_f \geq 0$ such that

$$f(y) \leq f(x) + \langle f'(x), y - x \rangle + \frac{L_f}{2} \|y - x\|^2, \quad \forall x, y \in \text{dom}(f), \quad f'(x) \in \partial f(x). \quad (30)$$

Clearly, if f is L_f -smooth, then (30) holds. If f is also convex, then (30) implies that f is L_f -smooth.

Under these additional assumptions, we have the following result.

Theorem 5. Suppose that Assumptions 1, 2, 4, 5, 6, and 7 hold and $\mathcal{G}_\eta$ is defined by (18).

Let $\{(\tilde{w}_t, \tilde{u}_t)\}$ be generated by Algorithm 2 using **one epoch** ($S = 1$) of the **shuffling routine** (27), and fixed learning rates $\eta_t = \eta := \mathcal{O}(\epsilon)$ as in Theorem 9 of Supp. Doc. C for a given $\epsilon > 0$, $\hat{\eta}_t := \hat{\eta} = 30\kappa^2\eta$, and $T := \mathcal{O}(\epsilon^{-3})$, where $\kappa := \frac{L_u}{\mu_H + \mu_h}$. Then, we have $\frac{1}{T+1} \sum_{t=0}^T \|\mathcal{G}_\eta(\tilde{w}_t)\|^2 \leq \epsilon^2$.

Consequently, Algorithm 2 requires $\mathcal{O}(n\epsilon^{-3})$ evaluations of both $\nabla_w \mathcal{H}_i$ and of $\nabla_u \mathcal{H}_i$ to achieve an ϵ -stationary point $\hat{w}_T$ of (3) computed by (19).

Similar to Algorithm 1, if $\pi^{(s,t)}$ and $\hat{\pi}^{(t)}$ are generated randomly and independently, $\Lambda_1 = \mathcal{O}(1/n)$, and $\hat{\Lambda}_1 = \mathcal{O}(1/n)$, then our complexity stated above can be improved by a factor of $\sqrt{n}$. Nevertheless, we omit this analysis. Finally, we can combine each Theorem 2, 3, 4 or 5 and Lemma 2 to construct an $\hat{\epsilon}$ -KKT point of (1). Theorem 5 has a better complexity than Theorems 3 and 4, but requires stronger assumptions. Algorithm 2 is also different from the one in [3] both in terms of algorithmic form and the underlying problem to be solved, while achieving the same oracle complexity.

5 Numerical Experiments

We perform some experiments to illustrate Algorithm 1 and compare it with two existing and related algorithms. Further details and additional experiments can be found in Supp. Doc. D.

We consider the following regularized stochastic minimax problem studied, e.g., in [9, 33]:

$$\min_{w \in \mathbb{R}^p} \left\{ \max_{1 \leq j \leq m} \left\{ \frac{1}{n} \sum_{i=1}^n F_{i,j}(w) \right\} + \frac{\lambda}{2} \|w\|^2 \right\}, \quad (31)$$

where $F_{i,j} : \mathbb{R}^p \times \Omega \rightarrow \mathbb{R}_+$ can be viewed as the loss of the j -th model for data point $i \in [n]$. If we define $\phi_0(v) := \max_{1 \leq j \leq m} \{v_j\}$ and $f(w) := \frac{\lambda}{2} \|w\|^2$, then (31) can be reformulated into (3). Since $v_j \geq 0$, we have $\phi_0(v) := \max_{1 \leq j \leq m} \{v_j\} = \|v\|_\infty = \max_{\|u\|_1 \leq 1} \langle v, u \rangle$, which is nonsmooth. Thus we can smooth ϕ_0 as $\phi_\gamma(v) := \max_{\|u\|_1 \leq 1} \{\langle v, u \rangle - (\gamma/2) \|u\|^2\}$ using $b(u) := \frac{1}{2} \|u\|^2$.

Here, we apply our problem (31) to solve a model selection problem in binary classification with nonnegative nonconvex losses, see, e.g., [41]. Each function $F_{i,j}$ belongs to 4 different nonconvex losses ($m = 4$): $F_{i,1}(w, \xi) := 1 - \tanh(b_i \langle a_i, w \rangle)$, $F_{i,2}(w, \xi) := \log(1 + \exp(-b_i \langle a_i, w \rangle)) - \log(1 + \exp(-b_i \langle a_i, w \rangle - 1))$, $F_{i,3}(w, \xi) := (1 - 1/(\exp(-b_i \langle a_i, w \rangle) + 1))^2$, and $F_{i,4}(w, \xi) := \log(1 + \exp(-b_i \langle a_i, w \rangle))$ (see [41] for more details), where (a_i, b_i) represents data samples.

We implement 4 algorithms: our SGM with 2 options, SGD from [10], and Prox-Linear from [11]. We test these algorithms on two datasets from LIBSVM [6]. We set $\lambda := 10^{-4}$ and update the smoothing parameter γ_t as $\gamma_t := \frac{1}{2(t+1)^{1/3}}$. The learning rate η for all algorithms is finely tuned from $\{100, 50, 10, 5, 1, 0.5, 0.1, 0.05, 0.01, 0.001, 0.0001\}$, and the results are shown in Figure 1 for **w8a** and **rcv1** datasets using $k_b = 32$ blocks. The details of this experiment is given in Supp. Doc. D.

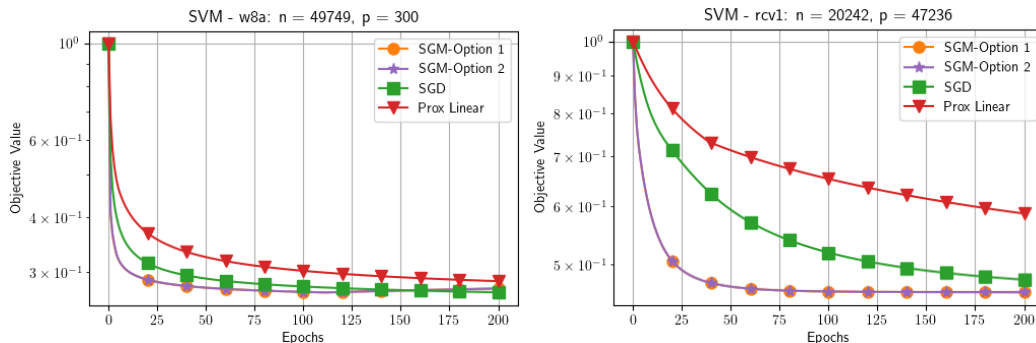


Figure 1: The performance of 4 algorithms for solving (31) on two datasets after 200 epochs.

As shown in Figure 1, the two variants of our SGM have a comparable performance with SGD and Prox-Linear, providing supportive evidence for using shuffling strategies in minimax algorithms.

6 Conclusions

This work explores a bilevel optimization approach to address two prevalent classes of nonconvex-concave minimax problems. These problems find numerous applications in practice, including robust learning and generative AIs. Motivated by the widespread use of shuffling strategies in implementing gradient-based methods within the machine learning community, we develop novel shuffling-based algorithms for solving these problems under standard assumptions. The first algorithm uses a non-standard shuffling strategy and achieves the state-of-the-art oracle complexity typically observed in nonconvex optimization. The second algorithm is also new, flexible, and offers a promising possibility for further exploration. Our results are expected to provide theoretical justification for incorporating shuffling strategies into minimax optimization algorithms, especially in nonconvex settings.

Acknowledgments and Disclosure of Funding

This work was partly supported by the National Science Foundation (NSF): NSF-RTG grant No. NSF DMS-2134107 and the Office of Naval Research (ONR), grant No. N00014-23-1-2588.

References

- [1] M. Arjovsky, S. Chintala, and L. Bottou. Wasserstein generative adversarial networks. In *International Conference on Machine Learning*, pages 214–223, 2017.
- [2] M. G. Azar, I. Osband, and R. Munos. Minimax regret bounds for reinforcement learning. In *International Conference on Machine Learning*, pages 263–272. PMLR, 2017.
- [3] A. Ben-Tal, L. El Ghaoui, and A. Nemirovski. *Robust optimization*. Princeton University Press, 2009.
- [4] A. Beznosikov, E. Gorbunov, H. Berard, and N. Loizou. Stochastic gradient descent-ascent: Unified theory and new efficient methods. In *International Conference on Artificial Intelligence and Statistics*, pages 172–235. PMLR, 2023.
- [5] K. Bhatia and K. Sridharan. Online learning with dynamics: A minimax perspective. *Advances in Neural Information Processing Systems*, 33:15020–15030, 2020.
- [6] C.-C. Chang and C.-J. Lin. LIBSVM: A library for Support Vector Machines. *ACM Transactions on Intelligent Systems and Technology*, 2:27:1–27:27, 2011.
- [7] H. Cho and C. Yun. SGDA with shuffling: faster convergence for nonconvex-PL minimax optimization. *The 11th International Conference on Learning Representations*, pp. 1–10, 2022.
- [8] A. Das, B. Schölkopf, and M. Muehlebach. Sampling without replacement leads to faster rates in finite-sum minimax optimization. *Advances in Neural Information Processing Systems*, 35:6749–6762, 2022.
- [9] S. Dempe. *Foundations of Bilevel Programming*. Springer Science & Business Media, 2002.
- [10] D. Driggs, J. Liang, and C.-B. Schönlieb. On biased stochastic gradient estimation. *Journal of Machine Learning Research*, vol. 23, no. 24, pp. 1–43, 2022.
- [11] K. Emmanouilidis, R. Vidal, and N. Loizou. Stochastic extragradient with random reshuffling: Improved convergence for variational inequalities. In *International Conference on Artificial Intelligence and Statistics*, pages 3682–3690. PMLR, 2024.
- [12] G. Gidel, H. Berard, G. Vignoud, P. Vincent, and S. Lacoste-Julien. A variational inequality perspective on generative adversarial networks. *International Conference on Learning Representations*, pp. 1–10, 2019.
- [13] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. In *Advances in neural information processing systems*, pages 2672–2680, 2014.
- [14] E. Gorbunov, H. Berard, G. Gidel, and N. Loizou. Stochastic extragradient: General analysis and improved rates. In *International Conference on Artificial Intelligence and Statistics*, pages 7865–7901. PMLR, 2022.
- [15] E. Y. Hamedani, A. Jalilzadeh, N. S. Aybat, and U. V. Shanbhag. Iteration complexity of randomized primal-dual methods for convex-concave saddle point problems. *arXiv preprint arXiv:1806.04118*, 2018.
- [16] J. Z. HaoChen and S. Sra. Random shuffling beats SGD after finite epochs. *International Conference on Machine Learning*, pp. 2624–2633, 2019.
- [17] E. Ho, A. Rajagopalan, A. Skvortsov, S. Arulampalam, and M. Piraveenan. Game theory in defence applications: A review. *Sensors*, 22(3):1032, 2022.

- [18] Y. Hsieh, F. Iutzeler, J. Malick, and P. Mertikopoulos. Explore aggressively, update conservatively: Stochastic extragradient methods with variable stepsize scaling. *Advances in Neural Information Processing Systems*, 33:16223–16234, 2020.
- [19] Abdul Jabbar, Xi Li, and Bourahla Omar. A survey on generative adversarial networks: Variants, applications, and training. *ACM Computing Surveys (CSUR)*, 54(8):1–49, 2021.
- [20] G. Lan. *First-order and Stochastic Optimization Methods for Machine Learning*. Springer, 2020.
- [21] F. Lin, X. Fang, and Z. Gao. Distributionally robust optimization: A review on theory and applications. *Numerical Algebra, Control & Optimization*, 12(1):159, 2022.
- [22] N. Loizou, H. Berard, G. Gidel, I. Mitliagkas, and S. Lacoste-Julien. Stochastic gradient descent-ascent and consensus optimization for smooth games: Convergence analysis under expected co-coercivity. *Advances in Neural Information Processing Systems*, 34:19095–19108, 2021.
- [23] L. Luo, H. Ye, and T. Zhang. Stochastic recursive gradient descent ascent for stochastic nonconvex-strongly-concave minimax problems. *Advances in Neural Information Processing Systems*, vol. 33, pp. 20566–20577, 2020.
- [24] Z. Luo, J. Pang, and D. Ralph. *Mathematical Programs with Equilibrium Constraints*. Cambridge University Press, Cambridge, 1996.
- [25] A. Madry, A. Makelov, L. Schmidt, D. Tsipras, and A. Vladu. Towards deep learning models resistant to adversarial attacks. In *International Conference on Learning Representations*, 2018.
- [26] Q. Meng, W. Chen, Y. Wang, Z.-M. Ma, and T.-Y. Liu. Convergence analysis of distributed stochastic gradient descent with shuffling. *Neurocomputing*, 337:46–57, 2019.
- [27] K. Mishchenko, A. Khaled, and P. Richtárik. Random reshuffling: Simple analysis with vast improvements. *Advances in Neural Information Processing Systems*, 33:17309–17320, 2020.
- [28] K. Mishchenko, A. Khaled, and P. Richtárik. Proximal and federated random reshuffling. In *International Conference on Machine Learning*, pages 15718–15749. PMLR, 2022.
- [29] Y. Nesterov. *Introductory lectures on convex optimization: A basic course*, volume 87 of *Applied Optimization*. Kluwer Academic Publishers, 2004.
- [30] L. M. Nguyen, Q. Tran-Dinh, D. T. Phan, P. H. Nguyen, and M. van Dijk. A unified convergence analysis for shuffling-type gradient methods. *Journal of Machine Learning Research*, 22(207):1–44, 2021.
- [31] B. Palaniappan and F. Bach. Stochastic variance reduction methods for saddle-point problems. In *Advances in Neural Information Processing Systems*, pages 1416–1424, 2016.
- [32] I. Safran and O. Shamir. How good is SGD with random shuffling? *Conference on Learning Theory*, pp. 3250–3284, 2020.
- [33] A. Shapiro and A. Kleywegt. Minimax analysis of stochastic problems. *Optim. Methods Softw.*, 17(3):523–542, 2002.
- [34] Q. Tran-Dinh, D. Liu, and L. M. Nguyen. Hybrid variance-reduced SGD algorithms for nonconvex-concave minimax problems. *The 34th Conference on Neural Information Processing Systems (NeurIPS 2020)*, 2020.
- [35] J. Wang, T. Zhang, S. Liu, P.-Y. Chen, J. Xu, M. Fardad, and B. Li. Adversarial attack generation empowered by min-max optimization. *Advances in Neural Information Processing Systems*, 34:16020–16033, 2021.
- [36] M. Wang, E. Fang, and L. Liu. Stochastic compositional gradient descent: algorithms for minimizing compositions of expected-value functions. *Math. Program.*, 161(1-2):419–449, 2017.

- [37] J. Yang, N. Kiyavash, and N. He. Global convergence and variance-reduced optimization for a class of nonconvex-nonconcave minimax problems. *arXiv preprint arXiv:2002.09621*, 2020.
- [38] B. Ying, K. Yuan, and A. H. Sayed. Convergence of variance-reduced stochastic learning under random reshuffling. *arXiv preprint arXiv:1708.01383*, 2(3):6, 2017.
- [39] J. Zhang and L. Xiao. Stochastic variance-reduced prox-linear algorithms for nonconvex composite optimization. *Mathematical Programming*, pp. 1–43, 2022.
- [40] K. Zhang, Z. Yang, and T. Başar. Multi-agent reinforcement learning: A selective overview of theories and algorithms. *Handbook of reinforcement learning and control*, pages 321–384, 2021.
- [41] L. Zhao, M. Mammadov, and J. Yearwood. From convex to nonconvex: a loss function analysis for binary classification. In *IEEE International Conference on Data Mining Workshops (ICDMW)*, pages 1281–1288. IEEE, 2010.
- [42] R. Zhao. Optimal stochastic algorithms for convex-concave saddle-point problems. *arXiv preprint arXiv:1903.01687*, 2019.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: Our claims made in the abstract reflects our contribution stated in the introduction, see the "Contribution" paragraph in the introduction section. Our contribution consists of two algorithms, Algorithm 1 and Algorithm 2, and their theoretical convergence guarantees stated in the subsequent theorems.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: This paper has limitation as it only focuses on two classes of minimax problems defined in (1). Yes, we only consider two classes of minimax problems: nonconvex-linear (NL) and convex-strongly concave (NC), covered by our assumption, Assumptions 1 to 5.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory Assumptions and Proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: We state all required assumptions in Assumptions 1 to 5. Our theoretical results stated in each theorem also refers to these assumptions when required. Our full proofs are given in Supp. Docs. due to space limit, and we believe that our technical proofs are correct.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental Result Reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: We provide the details of our experiments, including mathematical models, the detailed implementation of algorithms, the choice of parameters, and datasets. We also upload the code with examples to run and verify.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: Our data is available online from LIBSVM. The code is implemented in Python. The code for all experiments is also provided with instruction.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental Setting/Details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Supp. Doc. D provides all the details of our experiments, including how to select parameters, and how to report our results.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment Statistical Significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [NA]

Justification: The paper does not have such a result to report.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments Compute Resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: Our experiments were run on a MacBook Pro. 2.8GHz Quad-Core Intel Core I7, 16Gb Memory specified at the beginning of Supp. Doc. D.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code Of Ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Our data is publicly available online from LIBSVM.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader Impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: We do not yet know if our paper has an immediate broader impact. However, since our problems and our algorithms are sufficiently general, we hope they will create broader impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: We do not have our own real data or specific model that has a high risk for misuse.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: Our code is open-source and will be made available online under a standard public license.

Guidelines:

- The answer NA means that the paper does not use existing assets.

- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New Assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: It does not have new asset.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and Research with Human Subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: It does not relate to crowdsourcing experiments and research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional Review Board (IRB) Approvals or Equivalent for Research with Human Subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: It does not require any approval.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.