
Efficiently Learning Adversarially Robust Halfspaces with Noise

Omar Montasser¹ Surbhi Goel² Ilias Diakonikolas³ Nathan Srebro¹

Abstract

We study the problem of learning adversarially robust halfspaces in the distribution-independent setting. In the realizable setting, we provide necessary and sufficient conditions on the adversarial perturbation sets under which halfspaces are efficiently robustly learnable. In the presence of random label noise, we give a simple computationally efficient algorithm for this problem with respect to any ℓ_p -perturbation.

1. Introduction

Learning predictors that are robust to adversarial examples remains a major challenge in machine learning. A line of work has shown that predictors learned by deep neural networks are *not* robust to adversarial examples (Szegedy et al., 2014; Biggio et al., 2013; Goodfellow et al., 2015). This has led to a long line of research studying different aspects of robustness to adversarial examples.

In this paper, we consider the problem of distribution-independent learning of halfspaces that are robust to adversarial examples at test time, also referred to as robust PAC learning of halfspaces. Halfspaces are binary predictors of the form $h_w(x) = \text{sign}(\langle w, x \rangle)$, where $w \in \mathbb{R}^d$.

In adversarially robust PAC learning, given an instance space $\mathcal{X}$ and label space $\mathcal{Y} = \{\pm 1\}$, we formalize an adversary – that we would like to be robust against – as a map $\mathcal{U} : \mathcal{X} \mapsto 2^{\mathcal{X}}$, where $\mathcal{U}(x) \subseteq \mathcal{X}$ represents the set of perturbations (adversarial examples) that can be chosen by the adversary at test time (i.e., we require that $x \in \mathcal{U}(x)$). For an unknown distribution $\mathcal{D}$ over $\mathcal{X} \times \mathcal{Y}$, we observe m i.i.d. samples $S \sim \mathcal{D}^m$, and our goal is to learn a predictor $\hat{h} : \mathcal{X} \mapsto \mathcal{Y}$ that achieves small robust risk,

$$R_{\mathcal{U}}(\hat{h}; \mathcal{D}) \triangleq \mathbb{E}_{(x,y) \sim \mathcal{D}} \left[\sup_{z \in \mathcal{U}(x)} \mathbb{1}[\hat{h}(z) \neq y] \right]. \quad (1)$$

¹Toyota Technological Institute at Chicago ²University of Texas at Austin ³University of Wisconsin-Madison. Correspondence to: Omar Montasser <omar@ttic.edu>.

The information-theoretic aspects of adversarially robust learning have been studied in recent work, see, e.g., (Schmidt et al., 2018; Cullina et al., 2018; Khim & Loh, 2018; Bubeck et al., 2019; Yin et al., 2019; Montasser et al., 2019). This includes studying what learning rules should be used for robust learning and how much training data is needed to guarantee high robust accuracy. It is now known that any hypothesis class $\mathcal{H}$ with finite VC dimension is robustly learnable, though sometimes improper learning is necessary and the sample complexity may be exponential in the VC dimension (Montasser et al., 2019).

On the other hand, the computational aspects of adversarially robust PAC learning are less understood. In this paper, we take a first step towards studying this broad algorithmic question with a focus on the fundamental problem of learning adversarially robust halfspaces.

A first question to ask is whether efficient PAC learning implies efficient *robust* PAC learning, i.e., whether there is a general reduction that solves the adversarially robust learning problem. Recent work has provided strong evidence that this is not the case. Specifically, (Bubeck et al., 2019) showed that there exists a learning problem that can be learned efficiently non-robustly, but is computationally intractable to learn robustly (under plausible complexity-theoretic assumptions). There is also more recent evidence that suggests that this is also the case in the PAC model. (Awasthi et al., 2019) showed that it is computationally intractable to even weakly robustly learn degree-2 polynomial threshold functions (PTFs) with ℓ_∞ perturbations in the realizable setting, while PTFs of any constant degree are known to be efficiently PAC learnable non-robustly in the realizable setting. (Gourdeau et al., 2019) showed that there are hypothesis classes that are hard to robustly PAC learn, under the assumption that it is hard to non-robustly PAC learn.

The aforementioned discussion suggests that when studying robust PAC learning, we need to characterize which types of perturbation sets $\mathcal{U}$ admit *computationally efficient* robust PAC learners and under which noise assumptions. In the agnostic PAC setting, it is known that even weak (non-robust) learning of halfspaces is computationally intractable (Feldman et al., 2006; Guruswami & Raghavendra, 2009; Diakonikolas et al., 2011; Daniely, 2016). For ℓ_2 -

perturbations, where $\mathcal{U}(x) = \{z : \|x - z\|_2 \leq \gamma\}$, it was recently shown that the complexity of proper learning is exponential in $1/\gamma$ (Diakonikolas et al., 2019b). In this paper, we focus on the realizable case and the (more challenging) case of random label noise.

We can be more optimistic in the realizable setting. Halfspaces are efficiently PAC learnable non-robustly via Linear Programming (Maass & Turán, 1994), and under the margin assumption via the Perceptron algorithm (Rosenblatt, 1958). But what can we say about robustly PAC learning halfspaces? Given a perturbation set $\mathcal{U}$ and under the assumption that there is a halfspace h_w that robustly separates the data, can we efficiently learn a predictor with small robust risk?

Just as empirical risk minimization (ERM) is central for non-robust PAC learning, a core component of adversarially robust learning is minimizing the *robust* empirical risk on a dataset S ,

$$\hat{h} \in \text{RERM}_{\mathcal{U}}(S) \triangleq \underset{h \in \mathcal{H}}{\text{argmin}} \frac{1}{m} \sum_{i=1}^m \sup_{z \in \mathcal{U}(x_i)} \mathbb{1}[h(z) \neq y_i].$$

In this paper, we provide necessary and sufficient conditions on perturbation sets $\mathcal{U}$, under which the robust empirical risk minimization (RERM) problem is *efficiently* solvable in the realizable setting. We show that an efficient separation oracle for $\mathcal{U}$ yields an efficient solver for $\text{RERM}_{\mathcal{U}}$, while an efficient *approximate* separation oracle for $\mathcal{U}$ is necessary for even computing the robust loss $\sup_{z \in \mathcal{U}(x)} \mathbb{1}[h_w(z) \neq y]$ of a halfspace h_w . In addition, we relax our realizability assumption and show that under random classification noise (Angluin & Laird, 1987), we can efficiently robustly PAC learn halfspaces with respect to any ℓ_p perturbation.

Main Contributions Our main contributions can be summarized as follows:

1. In the realizable setting, the class of halfspaces is efficiently robustly PAC learnable with respect to $\mathcal{U}$, given an efficient separation oracle for $\mathcal{U}$.
2. To even compute the robust risk with respect to $\mathcal{U}$ efficiently, an efficient *approximate* separation oracle for $\mathcal{U}$ is *necessary*.
3. In the random classification noise setting, the class of halfspaces is efficiently robustly PAC learnable with respect to any ℓ_p perturbation.

1.1. Related Work

Here we focus on the recent work that is most closely related to the results of this paper. (Awasthi et al., 2019) studied the tractability of RERM with respect to ℓ_∞ perturbations, obtaining efficient algorithms for halfspaces in the realizable setting, but showing that RERM for degree-2 polynomial

threshold functions is computationally intractable (assuming $\text{NP} \neq \text{RP}$). (Gourdeau et al., 2019) studied robust learnability of hypothesis classes defined over $\{0, 1\}^n$ with respect to hamming distance, and showed that monotone conjunctions are robustly learnable when the adversary can perturb only $\mathcal{O}(\log n)$ bits, but are *not* robustly learnable even under the uniform distribution when the adversary can flip $\omega(\log n)$ bits.

In this work, we take a more general approach, and instead of considering specific perturbation sets, we provide methods in terms of oracle access to a separation oracle for the perturbation set $\mathcal{U}$, and aim to characterize which perturbation sets $\mathcal{U}$ admit tractable RERM.

In the non-realizable setting, the only prior work we are aware of is by (Diakonikolas et al., 2019b) who studied the complexity of robustly learning halfspaces in the agnostic setting under ℓ_2 perturbations.

2. Problem Setup

Let $\mathcal{X} = \mathbb{R}^d$ be the instance space and $\mathcal{Y} = \{\pm 1\}$ be the label space. We consider halfspaces $\mathcal{H} = \{x \mapsto \text{sign}(\langle w, x \rangle) : w \in \mathbb{R}^d\}$.

The following definitions formalize the notion of adversarially robust PAC learning in the realizable and random classification noise settings:

Definition 2.1 (Realizable Robust PAC Learning). *We say $\mathcal{H}$ is robustly PAC learnable with respect to an adversary $\mathcal{U}$ in the realizable setting, if there exists a learning algorithm $\mathcal{A} : (\mathcal{X} \times \mathcal{Y})^* \mapsto \mathcal{Y}^{\mathcal{X}}$ with sample complexity $m : (0, 1) \rightarrow \mathbb{N}$ such that: for any $\epsilon, \delta \in (0, 1)$, for every data distribution $\mathcal{D}$ over $\mathcal{X} \times \mathcal{Y}$ where there exists a predictor $h^* \in \mathcal{H}$ with zero robust risk, $R_{\mathcal{U}}(h^*; \mathcal{D}) = 0$, with probability at least $1 - \delta$ over $S \sim \mathcal{D}^m$,*

$$R_{\mathcal{U}}(\mathcal{A}(S); \mathcal{D}) \leq \epsilon.$$

Definition 2.2 (Robust PAC Learning with Random Classification Noise). *Let $h^* \in \mathcal{H}$ be an unknown halfspace. Let $\mathcal{D}_x$ be an arbitrary distribution over $\mathcal{X}$ such that $R_{\mathcal{U}}(h^*; \mathcal{D}_{h^*}) = 0$, and $\eta \leq 0 < 1/2$. A noisy example oracle, $\text{EX}(h^*, \mathcal{D}_x, \eta)$ works as follows: Each time $\text{EX}(h^*, \mathcal{D}_x, \eta)$ is invoked, it returns a labeled example (x, y) , where $x \sim \mathcal{D}_x$, $y = h^*(x)$ with probability $1 - \eta$ and $y = -h^*(x)$ with probability η . Let $\mathcal{D}$ be the joint distribution on (x, y) generated by the above oracle.*

We say $\mathcal{H}$ is robustly PAC learnable with respect to an adversary $\mathcal{U}$ in the random classification noise model, if $\exists m(\epsilon, \delta, \eta) \in \mathbb{N} \cup \{0\}$ and a learning algorithm $\mathcal{A} : (\mathcal{X} \times \mathcal{Y})^ \mapsto \mathcal{Y}^{\mathcal{X}}$, such that for every distribution $\mathcal{D}$ over $\mathcal{X} \times \mathcal{Y}$ (generated as above by a noisy oracle), with probability at*

least $1 - \delta$ over $S \sim \mathcal{D}^m$,

$$R_{\mathcal{U}}(\mathcal{A}(S); \mathcal{D}) \leq \eta + \epsilon.$$

Sample Complexity of Robust Learning Denote by $\mathcal{L}_{\mathcal{H}}^{\mathcal{U}}$ the robust loss class of $\mathcal{H}$,

$$\mathcal{L}_{\mathcal{H}}^{\mathcal{U}} = \left\{ (x, y) \mapsto \sup_{z \in \mathcal{U}(x)} \mathbb{1}[h(z) \neq y] : h \in \mathcal{H} \right\}.$$

It was shown by (Cullina et al., 2018) that for any set $\mathcal{B} \subseteq \mathcal{X}$ that is nonempty, closed, convex, and origin-symmetric, and an adversary $\mathcal{U}$ that is defined as $\mathcal{U}(x) = x + \mathcal{B}$ (e.g., ℓ_p -balls), the VC dimension of the robust loss of halfspaces $\text{vc}(\mathcal{L}_{\mathcal{H}}^{\mathcal{U}})$ is at most the standard VC dimension $\text{vc}(\mathcal{H}) = d + 1$. Based on Vapnik’s “General Learning” (Vapnik, 1982), this implies that we have uniform convergence of robust risk with $m = \mathcal{O}(\frac{d + \log(1/\delta)}{\epsilon^2})$ samples. Formally, for any $\epsilon, \delta \in (0, 1)$ and any distribution $\mathcal{D}$ over $\mathcal{X} \times \mathcal{Y}$, with probability at least $1 - \delta$ over $S \sim \mathcal{D}^m$,

$$\forall \mathbf{w} \in \mathbb{R}^d, |R_{\mathcal{U}}(h_{\mathbf{w}}; \mathcal{D}) - R_{\mathcal{U}}(h_{\mathbf{w}}; S)| \leq \epsilon. \quad (2)$$

In particular, this implies that for any adversary $\mathcal{U}$ that satisfies the conditions above, $\mathcal{H}$ is robustly PAC learnable w.r.t. $\mathcal{U}$ by minimizing the *robust* empirical risk on S ,

$$\text{RERM}_{\mathcal{U}}(S) = \underset{\mathbf{w} \in \mathbb{R}^d}{\text{argmin}} \frac{1}{m} \sum_{i=1}^m \sup_{\mathbf{z}_i \in \mathcal{U}(x_i)} \mathbb{1}[y_i \langle \mathbf{w}, \mathbf{z}_i \rangle \leq 0]. \quad (3)$$

Thus, it remains to efficiently solve the $\text{RERM}_{\mathcal{U}}$ problem. We discuss necessary and sufficient conditions for solving $\text{RERM}_{\mathcal{U}}$ in the following section.

3. The Realizable Setting

In this section, we show necessary and sufficient conditions for minimizing the robust empirical risk $\text{RERM}_{\mathcal{U}}$ on a dataset $S = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_m, y_m)\} \in (\mathcal{X} \times \mathcal{Y})^m$ in the realizable setting, i.e. when the dataset S is *robustly* separable with a halfspace $h_{\mathbf{w}^*}$ where $\mathbf{w}^* \in \mathbb{R}^d$. In Theorem 3.5, we show that an efficient separation oracle for $\mathcal{U}$ yields an efficient solver for $\text{RERM}_{\mathcal{U}}$. While in Theorem 3.10, we show that an efficient *approximate* separation oracle for $\mathcal{U}$ is necessary for even computing the robust loss $\sup_{\mathbf{z} \in \mathcal{U}(x)} \mathbb{1}[h_{\mathbf{w}}(\mathbf{z}) \neq y]$ of a halfspace $h_{\mathbf{w}}$.

Note that the set of allowed perturbations $\mathcal{U}$ can be non-convex, and so it might seem difficult to imagine being able to solve the $\text{RERM}_{\mathcal{U}}$ problem in full generality. But, it turns out that for halfspaces it suffices to consider only convex perturbation sets due to the following observation:

Observation 3.1. Given a halfspace $\mathbf{w} \in \mathbb{R}^d$ and an example $(x, y) \in \mathcal{X} \times \mathcal{Y}$. If $\forall \mathbf{z} \in \mathcal{U}(x), y \langle \mathbf{w}, \mathbf{z} \rangle > 0$,

then $\forall \mathbf{z} \in \text{conv}(\mathcal{U}(x)), y \langle \mathbf{w}, \mathbf{z} \rangle > 0$. And if $\exists \mathbf{z} \in \mathcal{U}(x), y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$, then $\exists \mathbf{z} \in \text{conv}(\mathcal{U}(x)), y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$, where $\text{conv}(\mathcal{U}(x))$ denotes the convex-hull of $\mathcal{U}(x)$.

Observation 3.1 shows that for any dataset S that is *robustly* separable w.r.t. $\mathcal{U}$ with a halfspace $\mathbf{w}^*$, S is also robustly separable w.r.t. the convex hull $\text{conv}(\mathcal{U})$ using the same halfspace $\mathbf{w}^*$, where $\text{conv}(\mathcal{U})(x) = \text{conv}(\mathcal{U}(x))$. Thus, in the remainder of this section we only consider perturbation sets $\mathcal{U}$ that are convex, i.e., for each x , $\mathcal{U}(x)$ is convex.

Definition 3.2. Denote by $\text{SEP}_{\mathcal{U}}$ a separation oracle for $\mathcal{U}$. $\text{SEP}_{\mathcal{U}}(x, z)$ takes as input $x, z \in \mathcal{X}$ and either:

- asserts that $z \in \mathcal{U}(x)$, or
- returns a separating hyperplane $\mathbf{w} \in \mathbb{R}^d$ such that $\langle \mathbf{w}, \mathbf{z}' \rangle \leq \langle \mathbf{w}, \mathbf{z} \rangle$ for all $\mathbf{z}' \in \mathcal{U}(x)$.

Definition 3.3. For any $\eta > 0$, denote by $\text{SEP}_{\mathcal{U}}^{\eta}$ an approximate separation oracle for $\mathcal{U}$. $\text{SEP}_{\mathcal{U}}^{\eta}$ takes as input $x, z \in \mathcal{X}$ and either:

- asserts that $z \in \mathcal{U}(x)^{+\eta} \stackrel{\text{def}}{=} \{z : \exists \mathbf{z}' \in \mathcal{U}(x) \text{ s.t. } \|z - \mathbf{z}'\|_2 \leq \eta\}$, or
- returns a separating hyperplane $\mathbf{w} \in \mathbb{R}^d$ such that $\langle \mathbf{w}, \mathbf{z}' \rangle \leq \langle \mathbf{w}, \mathbf{z} \rangle + \eta$ for all $\mathbf{z}' \in \mathcal{U}(x)^{-\eta} \stackrel{\text{def}}{=} \{z' : B(z', \eta) \subseteq \mathcal{U}(x)\}$.

Definition 3.4. Denote by $\text{MEM}_{\mathcal{U}}$ a membership oracle for $\mathcal{U}$. $\text{MEM}_{\mathcal{U}}(x, z)$ takes as input $x, z \in \mathcal{X}$ and either:

- asserts that $z \in \mathcal{U}(x)$, or
- asserts that $z \notin \mathcal{U}(x)$.

When discussing a separation or membership oracle for a fixed convex set K , we overload notation and write SEP_K , SEP_K^{η} , and MEM_K (in this case only one argument is required).

3.1. An efficient separation oracle for $\mathcal{U}$ is sufficient to solve $\text{RERM}_{\mathcal{U}}$ efficiently

Let $\text{Soln}_S^{\mathcal{U}} = \{\mathbf{w} \in \mathbb{R}^d : \forall (x, y) \in S, \forall \mathbf{z} \in \mathcal{U}(x), y \langle \mathbf{w}, \mathbf{z} \rangle > 0\}$ denote the set of valid solutions for $\text{RERM}_{\mathcal{U}}(S)$ (see Equation 3). Note that $\text{Soln}_S^{\mathcal{U}}$ is not empty since we are considering the realizable setting. Although the treatment we present here is for homogeneous halfspaces (where a bias term is not needed), the results extend trivially to the non-homogeneous case.

Below, we show that we can efficiently find a solution $\mathbf{w} \in \text{Soln}_S^{\mathcal{U}}$ given access to a separation oracle for $\mathcal{U}$, $\text{SEP}_{\mathcal{U}}$.

Theorem 3.5. Let $\mathcal{U}$ be an arbitrary convex adversary. Given access to a separation oracle for $\mathcal{U}$, $\text{SEP}_{\mathcal{U}}$ that runs in time $\text{poly}(d, b)$. There is an algorithm that finds $\mathbf{w} \in \text{Soln}_S^{\mathcal{U}}$

in $\text{poly}(m, d, b)$ time where b is an upper bound on the bit complexity of the valid solutions in $\text{Soln}_S^{\mathcal{U}}$ and the examples and perturbations in S .

Note that the polynomial dependence on b in the runtime is unavoidable even in standard non-robust ERM for halfspaces, unless we can solve linear programs in strongly polynomial time, which is currently an open problem.

Theorem 3.5 implies that for a broad family of perturbation sets $\mathcal{U}$, halfspaces $\mathcal{H}$ are efficiently robustly PAC learnable with respect to $\mathcal{U}$ in the realizable setting, as we show in the following corollary:

Corollary 3.6. *Let $\mathcal{U} : \mathcal{X} \mapsto 2^{\mathcal{X}}$ be an adversary such that $\mathcal{U}(\mathbf{x}) = \mathbf{x} + \mathcal{B}$ where $\mathcal{B}$ is nonempty, closed, convex, and origin-symmetric. Then, given access to an efficient separation oracle $\text{SEP}_{\mathcal{U}}$ that runs in time $\text{poly}(d, b)$, $\mathcal{H}$ is robustly PAC learnable w.r.t. $\mathcal{U}$ in the realizable setting in time $\text{poly}(d, b, 1/\epsilon, \log(1/\delta))$.*

Proof. This follows from the uniform convergence guarantee for the robust risk of halfspaces (see Equation (2)) and Theorem 3.5. $\square$

This covers many types of perturbation sets that are considered in practice. For example, $\mathcal{U}$ could be perturbations of distance at most γ w.r.t. some norm $\|\cdot\|$, such as the ℓ_{∞} norm considered in many applications: $\mathcal{U}(\mathbf{x}) = \{\mathbf{z} \in \mathcal{X} : \|\mathbf{x} - \mathbf{z}\|_{\infty} \leq \gamma\}$. In addition, Theorem 3.5 also implies that we can solve the RERM problem for other natural perturbation sets such as translations and rotations in images (see, e.g., (Engstrom et al., 2019)), and perhaps mixtures of perturbations of different types (see, e.g., (Kang et al., 2019)), as long as we have access to efficient separation oracles for these sets.

Benefits of handling general perturbation sets $\mathcal{U}$: One important implication of Theorem 3.5 that highlights the importance of having a treatment that considers general perturbation sets (and not just ℓ_p perturbations for example) is the following: for any efficiently computable feature map $\varphi : \mathbb{R}^r \rightarrow \mathbb{R}^d$, we can efficiently solve the robust empirical risk problem over the induced halfspaces $\mathcal{H}_{\varphi} = \{\mathbf{x} \mapsto \text{sign}(\langle \mathbf{w}, \varphi(\mathbf{x}) \rangle) : \mathbf{w} \in \mathbb{R}^d\}$, as long as we have access to an efficient separation oracle for the image of the perturbations $\varphi(\mathcal{U}(\mathbf{x}))$. Observe that in general $\varphi(\mathcal{U}(\mathbf{x}))$ maybe non-convex and complicated even if $\mathcal{U}(\mathbf{x})$ is convex, however Observation 3.1 combined with the realizability assumption imply that it suffices to have an efficient separation oracle for the convex-hull $\text{conv}(\varphi(\mathcal{U}(\mathbf{x})))$.

Before we proceed with the proof of Theorem 3.5, we state the following requirements and guarantees for the Ellipsoid method which will be useful for us in the remainder of the section:

Lemma 3.7 (see, e.g., Theorem 2.4 in (Bubeck et al., 2015)). *Let $K \subseteq \mathbb{R}^d$ be a convex set, and SEP_K a separation oracle for K . Then, the Ellipsoid method using $\mathcal{O}(d^2b)$ oracle queries to SEP_K , will find a $\mathbf{w} \in K$, or assert that K is empty. Furthermore, the total runtime is $\mathcal{O}(d^4b)$.*

The proof of Theorem 3.5 relies on two key lemmas. First, we show that efficient robust certification yields an efficient solver for the RERM problem. Given a halfspace $\mathbf{w} \in \mathbb{R}^d$ and an example $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathcal{Y}$, efficient robust certification means that there is an algorithm that can *efficiently* either: (a) assert that $\mathbf{w}$ is robust on $\mathcal{U}(\mathbf{x})$, i.e. $\forall \mathbf{z} \in \mathcal{U}(\mathbf{x}), y \langle \mathbf{w}, \mathbf{z} \rangle > 0$, or (b) return a perturbation $\mathbf{z} \in \mathcal{U}(\mathbf{x})$ such that $y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$.

Lemma 3.8. *Let $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$ be a procedure that either: (a) Asserts that $\mathbf{w}$ is robust on $\mathcal{U}(\mathbf{x})$, i.e., $\forall \mathbf{z} \in \mathcal{U}(\mathbf{x}), y \langle \mathbf{w}, \mathbf{z} \rangle > 0$, or (b) Finds a perturbation $\mathbf{z} \in \mathcal{U}(\mathbf{x})$ such that $y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$. If $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$ can be solved in $\text{poly}(d, b)$ time, then there is an algorithm that finds $\mathbf{w} \in \text{Soln}_S^{\mathcal{U}}$ in $\text{poly}(m, d, b)$ time.*

Proof. Observe that $\text{Soln}_S^{\mathcal{U}}$ is a convex set since

$$\begin{aligned} \mathbf{w}_1, \mathbf{w}_2 \in \text{Soln}_S^{\mathcal{U}} &\Rightarrow \forall (\mathbf{x}, y) \in S, \forall \mathbf{z} \in \mathcal{U}(\mathbf{x}), y \langle \mathbf{w}_1, \mathbf{z} \rangle > 0 \\ &\quad \text{and } y \langle \mathbf{w}_2, \mathbf{z} \rangle > 0 \\ &\Rightarrow \forall \alpha \in [0, 1], \forall (\mathbf{x}, y) \in S, \forall \mathbf{z} \in \mathcal{U}(\mathbf{x}), \\ &\quad y \langle \alpha \mathbf{w}_1 + (1 - \alpha) \mathbf{w}_2, \mathbf{z} \rangle > 0 \\ &\Rightarrow \forall \alpha \in [0, 1], \alpha \mathbf{w}_1 + (1 - \alpha) \mathbf{w}_2 \in \text{Soln}_S^{\mathcal{U}}. \end{aligned}$$

Our goal is to find a $\mathbf{w} \in \text{Soln}_S^{\mathcal{U}}$. Let $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$ be an efficient robust certifier that runs in $\text{poly}(d, b)$ time. We will use $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$ to implement a separation oracle for $\text{Soln}_S^{\mathcal{U}}$ denoted $\text{SEP}_{\text{Soln}_S^{\mathcal{U}}}$. Given a halfspace $\mathbf{w} \in \mathbb{R}^d$, we simply check if $\mathbf{w}$ is robustly correct on all datapoints by running $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}_i, y_i))$ on each $(\mathbf{x}_i, y_i) \in S$. If there is a point $(\mathbf{x}_i, y_i) \in S$ where $\mathbf{w}$ is not robustly correct, then we get a perturbation $\mathbf{z}_i \in \mathcal{U}(\mathbf{x}_i)$ where $y_i \langle \mathbf{w}, \mathbf{z}_i \rangle \leq 0$, and we return $-\mathbf{y}_i \mathbf{z}_i$ as a separating hyperplane. Otherwise, we know that $\mathbf{w}$ is robustly correct on all datapoints, and we just assert that $\mathbf{w} \in \text{Soln}_S^{\mathcal{U}}$.

Once we have a separation oracle $\text{SEP}_{\text{Soln}_S^{\mathcal{U}}}$, we can use the Ellipsoid method (see Lemma 3.7) to solve the $\text{RERM}_{\mathcal{U}}(S)$ problem. More specifically, with a query complexity of $\mathcal{O}(d^2b)$ to $\text{SEP}_{\text{Soln}_S^{\mathcal{U}}}$ and overall runtime of $\text{poly}(m, d, b)$ (this depends on runtime of $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$), the Ellipsoid method will return a $\mathbf{w} \in \text{Soln}_S^{\mathcal{U}}$. $\square$

Next, we show that we can do efficient robust certification when given access to an efficient separation oracle for $\mathcal{U}$, $\text{SEP}_{\mathcal{U}}$.

Lemma 3.9. *If we have an efficient separation oracle $\text{SEP}_{\mathcal{U}}$ that runs in $\text{poly}(d, b)$ time. Then, we can efficiently solve $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$ in $\text{poly}(d, b)$ time.*

Proof. Given a halfspace $\mathbf{w} \in \mathbb{R}^d$ and $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathcal{Y}$, we want to either: (a) assert that $\mathbf{w}$ is robust on $\mathcal{U}(\mathbf{x})$, i.e. $\forall \mathbf{z} \in \mathcal{U}(\mathbf{x}), y \langle \mathbf{w}, \mathbf{z} \rangle > 0$, or (b) find a perturbation $\mathbf{z} \in \mathcal{U}(\mathbf{x})$ such that $y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$. Let $M(\mathbf{w}, y) = \{\mathbf{z}' \in \mathcal{X} : y \langle \mathbf{w}, \mathbf{z}' \rangle \leq 0\}$ be the set of all points that $\mathbf{w}$ mislabels. Observe that by definition $M(\mathbf{w}, y)$ is convex, and therefore $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$ is also convex. We argue that having an efficient separation oracle for $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$ suffices to solve our robust certification problem. Because if $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$ is not empty, then by definition, we can find a perturbation $\mathbf{z} \in \mathcal{U}(\mathbf{x})$ such that $y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$ with a separation oracle $\text{SEP}_{\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)}$ and the Ellipsoid method (see Lemma 3.7). If $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$ is empty, then by definition, $\mathbf{w}$ is robustly correct on $\mathcal{U}(\mathbf{x})$, and the Ellipsoid method will terminate and assert that $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$ is empty.

Thus, it remains to implement $\text{SEP}_{\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)}$. Given a point $\mathbf{z} \in \mathbb{R}^d$, we simply ask the separation oracle for $\mathcal{U}(\mathbf{x})$ by calling $\text{SEP}_{\mathcal{U}}(\mathbf{x}, \mathbf{z})$ and the separation oracle for $M(\mathbf{w}, y)$ by checking if $y \langle \mathbf{w}, \mathbf{z} \rangle \leq 0$. If $\mathbf{z} \notin \mathcal{U}(\mathbf{x})$ then we get a separating hyperplane $\mathbf{c}$ from $\text{SEP}_{\mathcal{U}}$ and we can use it to separate $\mathbf{z}$ from $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$. Similarly, if $\mathbf{z} \notin M(\mathbf{w}, y)$, by definition, $y \langle \mathbf{w}, \mathbf{z} \rangle > 0$ and so we can use $y\mathbf{w}$ as a separating hyperplane to separate $\mathbf{z}$ from $\mathcal{U}(\mathbf{x}) \cap M(\mathbf{w}, y)$. The overall runtime of this separation oracle is $\text{poly}(d, b)$, and so we can efficiently solve $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$ in $\text{poly}(d, b)$ time using the Ellipsoid method (Lemma 3.7). $\square$

We are now ready to proceed with the proof of Theorem 3.5.

Proof of Theorem 3.5. We want to efficiently solve $\text{RERM}_{\mathcal{U}}(S)$. Given that we have a separation oracle for $\mathcal{U}$, $\text{SEP}_{\mathcal{U}}$ that runs in $\text{poly}(d, b)$. By Lemma 3.9, we get an efficient robust certification procedure $\text{CERT}_{\mathcal{U}}(\mathbf{w}, (\mathbf{x}, y))$. Then, by Lemma 3.8, we get an efficient solver for $\text{RERM}_{\mathcal{U}}$. In particular, the runtime complexity is $\text{poly}(m, d, b)$. $\square$

3.2. An efficient approximate separation oracle for $\mathcal{U}$ is necessary for computing the robust loss

Our efficient algorithm for $\text{RERM}_{\mathcal{U}}$ requires a separation oracle for $\mathcal{U}$. We now show that even efficiently computing the robust loss of a halfspace $(\mathbf{w}, b_0) \in \mathbb{R}^d \times \mathbb{R}$ on an example $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathcal{Y}$ requires an efficient *approximate* separation oracle for $\mathcal{U}$.

Theorem 3.10. *Given a halfspace $\mathbf{w} \in \mathbb{R}^d$ and an example $(\mathbf{x}, y) \in \mathbb{R}^d \times \mathcal{Y}$, let $\text{EVAL}_{\mathcal{U}}((\mathbf{w}, b_0), (\mathbf{x}, y))$ be a procedure that computes the robust loss $\sup_{\mathbf{z} \in \mathcal{U}(\mathbf{x})} \mathbb{1}[y(\langle \mathbf{w}, \mathbf{z} \rangle + b_0) \leq$*

$0]$ in $\text{poly}(d, b)$ time, then for any $\gamma > 0$, we can implement an efficient γ -approximate separation oracle $\text{SEP}_{\mathcal{U}}^{\gamma}(\mathbf{x}, \mathbf{z})$ in $\text{poly}(d, b, \log(1/\gamma), \log(R))$ time, where $\mathcal{U}(\mathbf{x}) \subseteq B(0, R)$.

Proof. Let $\gamma > 0$. We will describe how to implement a γ -approximate separation oracle for $\mathcal{U}$ denoted $\text{SEP}_{\mathcal{U}}^{\gamma}(\mathbf{x}, \mathbf{z})$. Fix the first argument to an arbitrary $\mathbf{x} \in \mathcal{X}$. Upon receiving a point $\mathbf{z} \in \mathcal{X}$ as input, the main strategy is to search for a halfspace $\mathbf{w} \in \mathbb{R}^d$ that can label all of $\mathcal{U}(\mathbf{x})$ with +1, and label the point $\mathbf{z}$ with -1. If $\mathbf{z} \notin \mathcal{U}(\mathbf{x})$ then there is a halfspace $\mathbf{w}$ that separates $\mathbf{z}$ from $\mathcal{U}(\mathbf{x})$ because $\mathcal{U}(\mathbf{x})$ is convex, but this is impossible if $\mathbf{z} \in \mathcal{U}(\mathbf{x})$. Since we are only concerned with implementing an *approximate* separation oracle, we will settle for a slight relaxation which is to either:

- assert that $\mathbf{z}$ is γ -close to $\mathcal{U}(\mathbf{x})$, i.e., $\mathbf{z} \in B(\mathcal{U}(\mathbf{x}), \gamma)$, or
- return a separating hyperplane $\mathbf{w}$ such that $\langle \mathbf{w}, \mathbf{z}' \rangle \leq \langle \mathbf{w}, \mathbf{z} \rangle$ for all $\mathbf{z}' \in \mathcal{U}(\mathbf{x})$.

Let $K = \{(\mathbf{w}, b_0) : \forall \mathbf{z}' \in \mathcal{U}(\mathbf{x}), \langle \mathbf{w}, \mathbf{z}' \rangle + b_0 > 0\}$ denote the set of halfspaces that label all of $\mathcal{U}(\mathbf{x})$ with +1. Since $\mathcal{U}(\mathbf{x})$ is nonempty, it follows by definition that K is nonempty. To evaluate membership in K , given a query $\mathbf{w}_q, b_q$, we just make a call to $\text{EVAL}_{\mathcal{U}}((\mathbf{w}_q, b_q), (\mathbf{x}, +))$. Let $\text{MEM}_K(\mathbf{w}_q, b_q) = 1 - \text{EVAL}_{\mathcal{U}}((\mathbf{w}_q, b_q), (\mathbf{x}, +))$. This can be efficiently computed in $\text{poly}(d, b)$ time. Next, for any $\eta \in (0, 0.5)$, we can get an η -approximate separation oracle for K denoted SEP_K^{η} (see Definition 3.3) using $\mathcal{O}(db \log(d/\eta))$ queries to the membership oracle MEM_K described above (Lee et al., 2018). When queried with a halfspace $\tilde{\mathbf{w}} = (\mathbf{w}, b_0)$, SEP_K^{η} either:

- asserts that $\tilde{\mathbf{w}} \in K^{+\eta}$, or
- returns a separating hyperplane $\mathbf{c}$ such that $\langle \mathbf{c}, \tilde{\mathbf{w}}' \rangle \leq \langle \mathbf{c}, \tilde{\mathbf{w}} \rangle + \eta$ for all halfspaces $\tilde{\mathbf{w}}' \in K^{-\eta}$.

Observe that by definition, $K^{-\eta} \subseteq K \subseteq K^{+\eta}$. Furthermore, for any $\mathbf{w} \in K^{+\eta}$, by definition, $\exists \tilde{\mathbf{w}}' \in K$ such that $\|\tilde{\mathbf{w}} - \tilde{\mathbf{w}}'\|_2 \leq \eta$. Since, for each $\mathbf{z}' \in \mathcal{U}(\mathbf{x})$, by definition of K , we have $\langle \tilde{\mathbf{w}}', (\mathbf{z}', 1) \rangle = \langle \mathbf{w}', \mathbf{z}' \rangle + b_0 > 0$, it follows by Cauchy-Schwarz inequality that $(\forall \tilde{\mathbf{w}} \in K^{+\eta}) (\forall \mathbf{z}' \in \mathcal{U}(\mathbf{x}))$:

$$\langle \tilde{\mathbf{w}}, (\mathbf{z}', 1) \rangle = \langle \tilde{\mathbf{w}} - \tilde{\mathbf{w}}', \mathbf{z}' \rangle + \langle \tilde{\mathbf{w}}', (\mathbf{z}', 1) \rangle > -\eta 2R. \quad (4)$$

Let $\text{SEP}_K^{\gamma/4R}$ be a $\frac{\gamma}{4R}$ -approximate separation oracle for K . Observe that if the distance between $\mathbf{z}$ and $\mathcal{U}(\mathbf{x})$ is greater than γ , it follows that there is $(\mathbf{w}, b_0)$ such that:

$$\langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2 \text{ and } \langle \mathbf{w}, \mathbf{z}' \rangle + b_0 > 0 (\forall \mathbf{z}' \in \mathcal{U}(\mathbf{x})).$$

By definition of K , this implies that $K \cap \{(\mathbf{w}, b_0) : \langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2\}$ is not empty,

which implies that the intersection $K^{+\frac{\gamma}{4R}} \cap \{(\mathbf{w}, b_0) : \langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2\}$ is nonempty. We also have the contrapositive, which is, if the intersection $K^{+\frac{\gamma}{4R}} \cap \{(\mathbf{w}, b_0) : \langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2\}$ is empty, then we know that $\mathbf{z} \in B(\mathcal{U}(\mathbf{x}), \gamma)$. To conclude the proof, we run the Ellipsoid method with the approximate separation oracle $\text{SEP}_K^{\gamma/4R}$ to search over the restricted space $\{(\mathbf{w}, b_0) : \langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2\}$. Restricting the space is easily done because we will use the query point $\mathbf{z}$ as the separating hyperplane. Either the Ellipsoid method will find $(\mathbf{w}, b_0) \in K^{+\frac{\gamma}{4R}} \cap \{(\mathbf{w}, b_0) : \langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2\}$, in which case by Equation 4, $(\mathbf{w}, b_0)$ has the property that:

$$\langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\frac{\gamma}{2} \text{ and } \langle \mathbf{w}, \mathbf{z}' \rangle + b_0 > -\frac{\gamma}{2} \ (\forall \mathbf{z}' \in \mathcal{U}(\mathbf{x})),$$

and so we return $\mathbf{w}$ as a separating hyperplane between $\mathbf{z}$ and $\mathcal{U}(\mathbf{x})$. If the Ellipsoid terminates without finding any such $(\mathbf{w}, b_0)$, this implies that the intersection $K^{+\frac{\gamma}{4R}} \cap \{(\mathbf{w}, b_0) : \langle \mathbf{w}, \mathbf{z} \rangle + b_0 \leq -\gamma/2\}$ is empty, and therefore, by the contrapositive above, we assert that $\mathbf{z} \in B(\mathcal{U}(\mathbf{x}), \gamma)$. $\square$

4. Random Classification Noise

In this section, we relax the realizability assumption to random classification noise (Angluin & Laird, 1987). We show that for any adversary $\mathcal{U}$ that represents perturbations of bounded norm (i.e., $\mathcal{U}(\mathbf{x}) = \mathbf{x} + \mathcal{B}$, where $\mathcal{B} = \{\delta \in \mathbb{R}^d : \|\delta\|_p \leq \gamma\}$, $p \in [1, \infty]$), the class of halfspaces $\mathcal{H}$ is efficiently robustly PAC learnable with respect to $\mathcal{U}$ in the random classification noise model.

Theorem 4.1. *Let $\mathcal{U} : \mathcal{X} \mapsto 2^{\mathcal{X}}$ be an adversary such that $\mathcal{U}(\mathbf{x}) = \mathbf{x} + \mathcal{B}$ where $\mathcal{B} = \{\delta \in \mathbb{R}^d : \|\delta\|_p \leq \gamma\}$ and $p \in [1, \infty]$. Then, $\mathcal{H}$ is robustly PAC learnable w.r.t $\mathcal{U}$ under random classification noise in time $\text{poly}(d, 1/\varepsilon, 1/\gamma, 1/(1-2\eta), \log(1/\delta))$.*

The proof of Theorem 4.1 relies on the following key lemma. We show that the structure of the perturbations $\mathcal{B}$ allows us to relate the robust loss of a halfspace $\mathbf{w} \in \mathbb{R}^d$ with the γ -margin loss of $\mathbf{w}$. Before we state the lemma, recall that the dual norm of $\mathbf{w}$ denoted $\|\mathbf{w}\|_*$ is defined as $\sup\{\langle \mathbf{u}, \mathbf{w} \rangle : \|\mathbf{u}\| \leq 1\}$.

Lemma 4.2. *For any $\mathbf{w}, \mathbf{x} \in \mathbb{R}^d$ and any $y \in \mathcal{Y}$,*

$$\sup_{\delta \in \mathcal{B}} \mathbb{1}\{h_{\mathbf{w}}(\mathbf{x} + \delta) \neq y\} = \mathbb{1}\left\{y \left\langle \frac{\mathbf{w}}{\|\mathbf{w}\|_*}, \mathbf{x} \right\rangle \leq \gamma\right\}.$$

Proof. First observe that

$$\begin{aligned} \sup_{\delta \in \mathcal{B}} \mathbb{1}\{h_{\mathbf{w}}(\mathbf{x} + \delta) \neq y\} &= \sup_{\delta \in \mathcal{B}} \mathbb{1}\{y \langle \mathbf{w}, \mathbf{x} + \delta \rangle \leq 0\} \\ &= \mathbb{1}\left\{\inf_{\delta \in \mathcal{B}} y \langle \mathbf{w}, \mathbf{x} + \delta \rangle \leq 0\right\}. \end{aligned}$$

This holds because when $\inf_{\delta \in \mathcal{B}} y \langle \mathbf{w}, \mathbf{x} + \delta \rangle > 0$, by definition $\forall \delta \in \mathcal{B}, y \langle \mathbf{w}, \mathbf{x} + \delta \rangle > 0$, which implies that $\sup_{\delta \in \mathcal{B}} \mathbb{1}\{h_{\mathbf{w}}(\mathbf{x} + \delta) \neq y\} = 0$. For the other direction, when $\sup_{\delta \in \mathcal{B}} \mathbb{1}\{h_{\mathbf{w}}(\mathbf{x} + \delta) \neq y\} = 1$, by definition $\exists \delta \in \mathcal{B}$ such that $y \langle \mathbf{w}, \mathbf{x} + \delta \rangle \leq 0$, which implies that $\inf_{\delta \in \mathcal{B}} y \langle \mathbf{w}, \mathbf{x} + \delta \rangle \leq 0$. To conclude the proof, by definition of the set $\mathcal{B}$ and the dual norm $\|\cdot\|_*$, we have

$$\begin{aligned} \inf_{\delta \in \mathcal{B}} y \langle \mathbf{w}, \mathbf{x} + \delta \rangle &= y \langle \mathbf{w}, \mathbf{x} \rangle - \sup_{\delta \in \mathcal{B}} \langle -y\mathbf{w}, \delta \rangle \\ &= y \langle \mathbf{w}, \mathbf{x} \rangle - \|\mathbf{w}\|_* \gamma. \end{aligned}$$

$\square$

Lemma 4.2 implies that for any distribution $\mathcal{D}$ over $\mathcal{X} \times \mathcal{Y}$, to solve the γ -robust learning problem

$$\argmin_{\mathbf{w} \in \mathbb{R}^d} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}} \left[\sup_{\delta \in \mathcal{B}} \mathbb{1}\{h_{\mathbf{w}}(\mathbf{x} + \delta) \neq y\} \right], \quad (5)$$

it suffices to solve the γ -margin learning problem

$$\argmin_{\|\mathbf{w}\|_* = 1} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}} [\mathbb{1}\{y \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma\}]. \quad (6)$$

We will solve the γ -margin learning problem in Equation (6) in the random classification noise setting using an appropriately chosen convex surrogate loss. Our convex surrogate loss and its analysis build on a convex surrogate that appears in the appendix of (Diakonikolas et al., 2019a) for learning large ℓ_2 -margin halfspaces under random classification noise w.r.t. the 0-1 loss. We note that the idea of using a convex surrogate to (non-robustly) learn large margin halfspaces in the presence of random classification noise is implicit in a number of prior works, starting with (Bylander, 1994).

Our robust setting is more challenging for the following reasons. First, we are not interested in only ensuring small 0-1 loss, but rather ensuring small γ -margin loss. Second, we want to be able to handle all ℓ_p norms, as opposed to just the ℓ_2 norm. As a result, our analysis is somewhat delicate.

Let

$$\phi(s) = \begin{cases} \lambda(1 - \frac{s}{\gamma}), & s > \gamma \\ (1 - \lambda)(1 - \frac{s}{\gamma}), & s \leq \gamma \end{cases}.$$

We will show that solving the following convex optimization problem:

$$\argmin_{\|\mathbf{w}\|_* \leq 1} G_{\lambda}^{\gamma}(\mathbf{w}) \stackrel{\text{def}}{=} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}} [\phi(y \langle \mathbf{w}, \mathbf{x} \rangle)], \quad (7)$$

where $\lambda = \frac{\varepsilon\gamma/2+\eta}{1+\varepsilon\gamma}$, suffices to solve the γ -margin learning problem in Equation (6). Intuitively, the idea here is that for $\lambda = 0$, the ϕ objective is exactly a scaled hinge loss, which gives a learning guarantee w.r.t to the γ -margin loss when

there is no noise ($\eta = 0$). When the noise $\eta > 0$, we slightly adjust the slopes, such that even correct prediction encounters a loss. The choice of the slope is based on λ which will depend on the noise rate η and the ε -suboptimality that is required for Equation (6).

We can solve Equation (7) with a standard first-order method through samples using Stochastic Mirror Descent, when the dual norm $\|\cdot\|_*$ is an ℓ_q -norm. We state the following properties of Mirror Descent we will require:

Lemma 4.3 (see, e.g., Theorem 6.1 in (Bubeck et al., 2015)).

Let $G(\mathbf{w}) \stackrel{\text{def}}{=} \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}} [\ell(\mathbf{w}, (\mathbf{x}, y))]$ be a convex function that is L -Lipschitz w.r.t. $\|\mathbf{w}\|_q$ where $q > 1$. Then, using the potential function $\psi(\mathbf{w}) = \frac{1}{2}\|\mathbf{w}\|_q^2$, a suitable step-size η , and a sequence of iterates $\mathbf{w}^k$ computed by the following update:

$$\mathbf{w}^{k+1} = \Pi_{B_q}^{\psi} (\nabla \psi^{-1} (\nabla \psi(\mathbf{w}^k) - \eta \mathbf{g}^k)) ,$$

Stochastic Mirror Descent with $\mathcal{O}(L^2/(q-1)\varepsilon^2)$ stochastic gradients $\mathbf{g}$ of G , will find an ε -suboptimal point $\hat{\mathbf{w}}$ such that $\|\hat{\mathbf{w}}\|_q \leq 1$ and $G(\hat{\mathbf{w}}) \leq \inf_{\mathbf{w}} G(\mathbf{w}) + \varepsilon$.

Remark 4.4. When $q = 1$, we will use the entropy potential function $\psi(\mathbf{w}) = \sum_{i=1}^d w_i \log w_i$. In this case, Stochastic Mirror Descent will require $\mathcal{O}(\frac{L^2 \log d}{\varepsilon^2})$ stochastic gradients.

We are now ready to state our main result for this section:

Theorem 4.5. Let $\mathcal{X} = \{\mathbf{x} \in \mathbb{R}^d : \|\mathbf{x}\|_p \leq 1\}$. Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \mathcal{Y}$ such that there exists a halfspace $\mathbf{w}^* \in \mathbb{R}^d$ with $\Pr_{\mathbf{x} \sim \mathcal{D}_x} [\langle \mathbf{w}^*, \mathbf{x} \rangle > \gamma] = 1$ and y is generated by $h_{\mathbf{w}^*}(\mathbf{x}) := \text{sign}(\langle \mathbf{w}^*, \mathbf{x} \rangle)$ corrupted by RCN with noise rate $\eta < 1/2$. An application of Stochastic Mirror Descent on $G_\lambda^\gamma(\mathbf{w})$, returns, with high probability, a halfspace $\mathbf{w}$ where $\|\mathbf{w}\|_q \leq 1$ with $\gamma/2$ -robust misclassification error $\mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}} [\mathbb{1}\{y \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \leq \eta + \varepsilon$ in $\text{poly}(d, 1/\varepsilon, 1/\gamma, 1/(1-2\eta))$ time.

With Theorem 4.5, the proof of Theorem 4.1 immediately follows.

Proof of Theorem 4.1. This follows from Lemma 4.2 and Theorem 4.5. $\square$

Remark 4.6. In Theorem 4.5, we get a $\gamma/2$ -robustness guarantee assuming γ -robust halfspace $\mathbf{w}^*$ that is corrupted with random classification noise. This can be strengthened to get a guarantee of $(1-c)\gamma$ -robustness for any constant $c > 0$.

The rest of this section is devoted to the proof of Theorem 4.5. The high-level strategy is to show that an ε' -suboptimal solution to Equation (7) gives us an ε -suboptimal solution to Equation (6) (for a suitably chosen

ε'). In Lemma 4.8, we bound from above the $\gamma/2$ -margin loss in terms of our convex surrogate objective G_λ^γ , and in Lemma 4.9 we show that there are minimizers of our convex surrogate G_λ^γ such that it is sufficiently small. These are the two key lemmas that we will use to piece everything together.

For any $\mathbf{w}, \mathbf{x} \in \mathbb{R}^d$, consider the contribution of the objective G_λ^γ of $\mathbf{x}$, denoted by $G_\lambda^\gamma(\mathbf{w}, \mathbf{x})$. This is defined as $G_\lambda^\gamma(\mathbf{w}, \mathbf{x}) = \mathbb{E}_{y \sim \mathcal{D}_y(\mathbf{x})} [\phi(y \langle \mathbf{w}, \mathbf{x} \rangle)] = \eta \phi(-z) + (1-\eta) \phi(z)$ where $z = h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle$. In the following lemma, we provide a decomposition of $G_\lambda^\gamma(\mathbf{w}, \mathbf{x})$ that will help us in proving Lemmas 4.8 and 4.9.

Lemma 4.7. For any $\mathbf{w}, \mathbf{x} \in \mathbb{R}^d$, let $z = h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle$. Then, we have that:

$$\begin{aligned} G_\lambda^\gamma(\mathbf{w}, \mathbf{x}) &= (\eta - \lambda) \left(\frac{z}{\gamma} \right) + \lambda + \eta - 2\lambda\eta \\ &\quad + \mathbb{1}\{-\gamma \leq z \leq \gamma\} (1-\eta)(1-2\lambda) \left(1 - \frac{z}{\gamma} \right) \\ &\quad + \mathbb{1}\{z < -\gamma\} (1-2\lambda) \left(1 - 2\eta - \frac{z}{\gamma} \right). \end{aligned}$$

Proof. Based on the definition of the surrogate loss, it suffices to consider three cases:

Case $z > \gamma$:

$$\begin{aligned} l_1(z) &\stackrel{\text{def}}{=} \eta \phi(-z) + (1-\eta) \phi(z) \\ &= \eta(1-\lambda)(1+z/\gamma) + (1-\eta)\lambda(1-z/\gamma) \\ &= (\eta(1-\lambda) - (1-\eta)\lambda) \left(\frac{z}{\gamma} \right) + (1-\eta)\lambda \\ &\quad + \eta(1-\lambda) \\ &= (\eta - \lambda) \left(\frac{z}{\gamma} \right) + \lambda + \eta - 2\lambda\eta. \end{aligned}$$

Case $-\gamma \leq z \leq \gamma$:

$$\begin{aligned} l_2(z) &\stackrel{\text{def}}{=} \eta \phi(-z) + (1-\eta) \phi(z) \\ &= \eta(1-\lambda)(1+z/\gamma) \\ &\quad + (1-\eta)(1-\lambda)(1-z/\gamma) \\ &= -(1-2\eta)(1-\lambda) \left(\frac{z}{\gamma} \right) + 1 - \lambda. \end{aligned}$$

Case $z < -\gamma$:

$$\begin{aligned} l_3(z) &\stackrel{\text{def}}{=} \eta \lambda(1+z/\gamma) + (1-\eta)(1-\lambda)(1-z/\gamma) \\ &= (\eta\lambda - (1-\eta)(1-\lambda)) \left(\frac{z}{\gamma} \right) + \eta\lambda \\ &\quad + (1-\eta)(1-\lambda) \\ &= (\eta + \lambda - 1) \left(\frac{z}{\gamma} \right) + 1 - \eta - \lambda + 2\eta\lambda. \end{aligned}$$

Considering the three cases above, we can write

$$G(z) = l_1(z) + \mathbb{1}\{-\gamma \leq z \leq \gamma\} (l_2(z) - l_1(z)) \\ + \mathbb{1}\{z < -\gamma\} (l_3(z) - l_1(z)) .$$

Then, we calculate $l_2(z) - l_1(z)$ and $l_3(z) - l_1(z)$ as follows:

$$\begin{aligned} l_2(z) - l_1(z) &= -(1 - 2\eta)(1 - \lambda) \left(\frac{z}{\gamma} \right) + 1 - \lambda \\ &\quad - (\eta - \lambda) \left(\frac{z}{\gamma} \right) - \lambda - \eta + 2\lambda\eta \\ &= -((1 - 2\eta)(1 - \lambda) + \eta - \lambda) \left(\frac{z}{\gamma} \right) \\ &\quad + (1 - \eta)(1 - 2\lambda) \\ &= (1 - \eta)(1 - 2\lambda) \left(1 - \frac{z}{\gamma} \right) , \text{ and} \\ l_3(z) - l_1(z) &= (\eta + \lambda - 1) \left(\frac{z}{\gamma} \right) + 1 - \eta - \lambda + 2\eta\lambda \\ &\quad - (\eta - \lambda) \left(\frac{z}{\gamma} \right) - \lambda - \eta + 2\lambda\eta \\ &= (2\lambda - 1) \left(\frac{z}{\gamma} \right) + (1 - 2\eta)(1 - 2\lambda) \\ &= (2\lambda - 1) \left(\frac{z}{\gamma} + 2\eta - 1 \right) . \end{aligned}$$

Using the above, we have that

$$G(z) = l_1(z) + \mathbb{1}\{-\gamma \leq z \leq \gamma\} (1 - \eta)(1 - 2\lambda) \left(1 - \frac{z}{\gamma} \right) \\ + \mathbb{1}\{z < -\gamma\} (1 - 2\lambda) \left(1 - 2\eta - \frac{z}{\gamma} \right) .$$

□

The following lemma allows us to bound from below our convex surrogate $\mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}, \mathbf{x})]$ in terms of the $\gamma/2$ -margin loss of $\mathbf{w}$.

Lemma 4.8. Assume that λ is chosen such that $\lambda < 1/2$ and $\eta < \lambda$. Then, for any $\mathbf{w} \in \mathbb{R}^d$, $\mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}, \mathbf{x})] \geq \frac{\eta - \lambda}{\gamma} + \frac{1}{2}(1 - 2\lambda)(1 - \eta) \mathbb{E}_{\mathbf{x}} [\mathbb{1}\{z \leq \frac{\gamma}{2}\}] + \lambda + \eta - 2\lambda\eta$.

Proof. By Lemma 4.7, and linearity of expectation, we have that

$$\begin{aligned} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}, \mathbf{x})] &= (\eta - \lambda) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} \left[\frac{z}{\gamma} \right] + \lambda + \eta - 2\lambda\eta \\ &\quad + (1 - \eta)(1 - 2\lambda) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} \left[\mathbb{1}\{-\gamma \leq z \leq \gamma\} \left(1 - \frac{z}{\gamma} \right) \right] \\ &\quad + (1 - 2\lambda) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} \left[\mathbb{1}\{z < -\gamma\} \left(1 - 2\eta - \frac{z}{\gamma} \right) \right] . \end{aligned}$$

First, observe that for any $\mathbf{x}$, $-1 \leq z = h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq 1$ and since $\eta \leq \lambda$, we have

$$(\eta - \lambda) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} \left[\frac{z}{\gamma} \right] \geq \frac{\eta - \lambda}{\gamma} .$$

Then we observe that whenever $z < -\gamma$, $1 - \frac{z}{\gamma} - 2\eta > 2(1 - \eta) > (1 - \eta)/2$ and $\lambda \leq 1/2$, thus we can bound from below the third term

$$(1 - 2\lambda) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} \left[\mathbb{1}\{z < -\gamma\} \left(1 - 2\eta - \frac{z}{\gamma} \right) \right] \geq \frac{1}{2}(1 - 2\lambda)(1 - \eta) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{z < -\gamma\}] .$$

Next we note that whenever $-\gamma \leq z \leq \gamma$, $1 - \frac{z}{\gamma} \geq 0$. This implies that instead of considering $\mathbb{1}\{-\gamma \leq z \leq \gamma\}$, we can relax this and consider the subset $\mathbb{1}\{-\gamma \leq z \leq \frac{\gamma}{2}\}$, and on this subset $1 - \frac{z}{\gamma} \geq 1/2$. Thus, we can bound the second term from below as follows:

$$(1 - 2\lambda)(1 - \eta) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} \left[\mathbb{1}\{-\gamma \leq z \leq \gamma\} \left(1 - \frac{z}{\gamma} \right) \right] \geq \frac{1}{2}(1 - 2\lambda)(1 - \eta) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{-\gamma \leq z \leq \frac{\gamma}{2}\}] .$$

Combining the above, we obtain

$$\begin{aligned} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}, \mathbf{x})] &\geq \frac{\eta - \lambda}{\gamma} \\ &\quad + \frac{1}{2}(1 - 2\lambda)(1 - \eta) \mathbb{E}_{\mathbf{x}} [\mathbb{1}\{-\gamma \leq z \leq \frac{\gamma}{2}\}] + \mathbb{1}\{z < -\gamma\}] \\ &\quad + \lambda + \eta - 2\lambda\eta \\ &\geq \frac{\eta - \lambda}{\gamma} + \frac{1}{2}(1 - 2\lambda)(1 - \eta) \mathbb{E}_{\mathbf{x}} [\mathbb{1}\{z \leq \frac{\gamma}{2}\}] \\ &\quad + \lambda + \eta - 2\lambda\eta , \end{aligned}$$

as desired. □

We now show that there exist minimizers of the convex surrogate G_{λ}^{γ} such that it is sufficiently small, which will be useful later in choosing the suboptimality parameter ϵ' .

Lemma 4.9. Assume that λ is chosen such that $\lambda < 1/2$ and $\eta < \lambda$. Then we have that

$$\inf_{\mathbf{w} \in \mathbb{R}^d} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}, \mathbf{x})] \leq 2\eta(1 - \lambda) .$$

Proof. By definition, we have that $\inf_{\mathbf{w} \in \mathbb{R}^d} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}, \mathbf{x})] \leq \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [G_{\lambda}^{\gamma}(\mathbf{w}^*, \mathbf{x})]$. By assumption, with probability 1 over $\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}$, we have $h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}^*, \mathbf{x} \rangle > \gamma$. Thus, by Lemma 4.7, we have

$$\begin{aligned} \mathbb{E}_{\mathbf{x}} [G_{\lambda}^{\gamma}(\mathbf{w}^*, \mathbf{x})] &= (\eta - \lambda) \left(\frac{z}{\gamma} \right) + \lambda + \eta - 2\lambda\eta \\ &\leq \eta - \lambda + \lambda + \eta - 2\lambda\eta \\ &= 2\eta(1 - \lambda) , \end{aligned}$$

where the last inequality follows from the fact that $\eta < \lambda$ and the fact that $\mathbb{E}_x \left[\frac{z}{\gamma} \right] > 1$. $\square$

Using the above lemmas, we are now able to bound from above the $\gamma/2$ -margin loss of a halfspace $\mathbf{w}$ that is ε' -suboptimal for our convex optimization problem (see Equation (7)).

Lemma 4.10. *For any $\varepsilon' \in (0, 1)$ and any $\mathbf{w} \in \mathbb{R}^d$ such that $\mathbb{E}_{x \sim \mathcal{D}_x} [G_\lambda^\gamma(\mathbf{w}, \mathbf{x})] \leq \mathbb{E}_{x \sim \mathcal{D}_x} [G_\lambda^\gamma(\mathbf{w}^*, \mathbf{x})] + \varepsilon'$, the $\gamma/2$ -misclassification error of $\mathbf{w}$ satisfies*

$$\mathbb{E}_{(x,y) \sim \mathcal{D}} [\mathbb{1}\{y \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \leq \eta + \frac{2}{(1-2\lambda)} \left(\varepsilon' + (\lambda - \eta) \left(\frac{1}{\gamma} - 1 \right) \right).$$

Proof. By Lemma 4.8 and Lemma 4.9, we have

$$\frac{\eta - \lambda}{\gamma} + \frac{1}{2}(1-2\lambda)(1-\eta) \mathbb{E}_x \left[\mathbb{1}\left\{z \leq \frac{\gamma}{2}\right\} \right] + \lambda + \eta - 2\lambda\eta \leq 2\eta(1-\lambda) + \varepsilon'.$$

This implies

$$(1-\eta) \mathbb{E}_x \left[\mathbb{1}\left\{z \leq \frac{\gamma}{2}\right\} \right] \leq \frac{2}{(1-2\lambda)} \left(\varepsilon' + (\lambda - \eta) \left(\frac{1}{\gamma} - 1 \right) \right)$$

Since $\mathbb{E}_{(x,y) \sim \mathcal{D}} [\mathbb{1}\{y \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \leq \eta + (1-\eta) \mathbb{E}_x [\mathbb{1}\{z \leq \frac{\gamma}{2}\}]$, we get the desired result. $\square$

We are now ready to prove Theorem 4.5.

Proof of Theorem 4.5. Based on Lemma 4.10, we will choose λ, ε' such that

$$\frac{2}{(1-2\lambda)} \left(\varepsilon' + (\lambda - \eta) \left(\frac{1}{\gamma} - 1 \right) \right) \leq \varepsilon.$$

By setting $\varepsilon' = \lambda - \eta$, this condition reduces to

$$\frac{2(\lambda - \eta)}{(1-2\lambda)} \leq \gamma\varepsilon.$$

This implies that we need $\lambda \leq \frac{\varepsilon\gamma/2 + \eta}{1 + \varepsilon\gamma}$. We will choose $\lambda = \frac{\varepsilon\gamma/2 + \eta}{1 + \varepsilon\gamma}$. Note that our analysis relied on having $\lambda \leq 1/2$ and $\eta \leq \lambda$. These conditions combined imply that we should choose λ such that $\eta \leq \lambda \leq 1/2$. Our choice of $\lambda = \frac{\varepsilon\gamma/2 + \eta}{1 + \varepsilon\gamma}$ satisfies these conditions, since

$$\frac{\varepsilon\gamma/2 + \eta}{1 + \varepsilon\gamma} - \eta = \frac{\varepsilon\gamma/2 + \eta - \eta(1 + \varepsilon\gamma)}{1 + \varepsilon\gamma} = \frac{\varepsilon\gamma(1/2 - \eta)}{1 + \varepsilon\gamma} \geq 0,$$

and

$$\frac{\varepsilon\gamma/2 + \eta}{1 + \varepsilon\gamma} - \frac{1}{2} = \frac{\varepsilon\gamma/2 + \eta - 1/2(1 + \varepsilon\gamma)}{1 + \varepsilon\gamma} = \frac{\eta - 1/2}{1 + \varepsilon\gamma} \leq 0.$$

By our choice of λ , we have that $\varepsilon' = \lambda - \eta = \frac{\varepsilon\gamma(1/2 - \eta)}{1 + \varepsilon\gamma} = \frac{\varepsilon\gamma(1-2\eta)}{2(1 + \varepsilon\gamma)}$. By the guarantees of Stochastic Mirror Descent (see Lemma 4.3), our theorem follows with $\mathcal{O}\left(\frac{1}{\varepsilon^2\gamma^2(1-2\eta)^2(q-1)}\right)$ samples for $q > 1$ and $\mathcal{O}\left(\frac{\log d}{\varepsilon^2\gamma^2(1-2\eta)^2}\right)$ samples for $q = 1$. $\square$

5. Conclusion

In this paper, we provide necessary and sufficient conditions for perturbation sets $\mathcal{U}$, under which we can efficiently solve the robust empirical risk minimization (RERM) problem. We give a polynomial time algorithm to solve RERM given access to a polynomial time separation oracle for $\mathcal{U}$. In addition, we show that an efficient *approximate* separation oracle for $\mathcal{U}$ is necessary for even computing the robust loss of a halfspace. As a corollary, we show that halfspaces are efficiently robustly PAC learnable for a broad range of perturbation sets. By relaxing the realizability assumption, we show that under random classification noise, we can efficiently robustly PAC learn halfspaces with respect to any ℓ_p perturbations. An interesting direction for future work is to understand the computational complexity of robustly PAC learning halfspaces under stronger noise models, including Massart noise and agnostic noise.

Acknowledgments

We thank Sepideh Mahabadi for many insightful and helpful discussions, and Brian Bullins for helpful discussions on Mirror Descent. We also thank the anonymous reviewers for their helpful feedback. This work was initiated when the authors were visiting the Simons Institute for the Theory of Computing as part of the Summer 2019 program on *Foundations of Deep Learning*. Work by N.S. and O.M. was partially supported by NSF award IIS-1546500 and by DARPA¹ cooperative agreement HR00112020003. I.D. was supported by NSF Award CCF-1652862 (CAREER), a Sloan Research Fellowship, and a DARPA Learning with Less Labels (LwLL) grant. S.G. was supported by the JP Morgan AI Research PhD Fellowship.

Appendix: Another Approach to Large Margin Learning under Random Classification Noise

We remark that γ -margin learning of halfspaces has been studied in earlier work, and we have algorithms such as Margin Perceptron (Balcan et al., 2008) and SVM. The Margin Perceptron (for ℓ_2 margin) and other ℓ_p margin algorithms have been also implemented in the SQ model

¹This paper does not reflect the position or the policy of the Government, and no endorsement should be inferred.

(Feldman et al., 2017). But no explicit connection has been made to adversarial robustness.

We present here a simple approach to learn γ -margin halfspaces under random classification noise using only a convex surrogate loss and Stochastic Mirror Descent. The construction of the convex surrogate is based on learning generalized linear models with a suitable link function $u : \mathbb{R} \rightarrow \mathbb{R}$. To the best of our knowledge, the result of this section is not explicit in prior work.

Theorem 5.1. *Let $\mathcal{X} = \{x \in \mathbb{R}^d : \|x\|_p \leq 1\}$. Let $\mathcal{D}$ be a distribution over $\mathcal{X} \times \mathcal{Y}$ such that there exists a halfspace $w^* \in \mathbb{R}^d$, $\|w^*\|_q = 1$ with $\Pr_{x \sim \mathcal{D}_x} [|\langle w^*, x \rangle| > \gamma] = 1$ and y is generated by $h_{w^*}(x) := \text{sign}(\langle w^*, x \rangle)$ corrupted with random classification noise rate $\eta < 1/2$. Then, running Stochastic Mirror Descent on the following convex optimization problem:*

$$\min_{w \in \mathbb{R}^d, \|w\|_q \leq 1} \mathbb{E}_{(x,y) \sim \mathcal{D}} [\ell(w, (x, y))]$$

where the convex loss function ℓ is defined in Equation 8, returns with high probability, a halfspace w with $\gamma/2$ -robust misclassification error $\mathbb{E}_{(x,y) \sim \mathcal{D}} [\mathbb{1}\{y \langle w, x \rangle \leq \gamma/2\}] \leq \eta + \varepsilon$.

We prove Theorem 4.5 in the remainder of this section. We will connect our problem to that of solving generalized linear models. We define the link function as follows,

$$u(s) = \begin{cases} \eta & s < -\gamma \\ \frac{1-2\eta}{2\gamma} s + \frac{1}{2} & -\gamma \leq s \leq \gamma \\ 1-\eta & s > \gamma \end{cases}.$$

Observe that u is monotone and $\frac{1-2\eta}{2\gamma}$ -Lipschitz.

First, we will relate our loss of interest, which is the $\gamma/2$ -margin loss with the squared loss defined in terms of the link function u ,

Lemma 5.2. *For any $w \in \mathbb{R}^d$,*

$$\mathbb{E}_{x \sim \mathcal{D}_x} [\mathbb{1}\{h_{w^*}(x) \langle w, x \rangle \leq \gamma/2\}] \leq \frac{16}{(1-2\eta)^2} \mathbb{E}_{x \sim \mathcal{D}_x} [(u(\langle w, x \rangle) - u(\langle w^*, x \rangle))^2].$$

Proof. Let $E^+ = \{h_{w^*}(x) = +\}$ and $E^- = \{h_{w^*}(x) = -\}$. Let $U(x) = (u(\langle w, x \rangle) - u(\langle w^*, x \rangle))^2$. By law of total expectation and definition of the link function u , we have

$$\begin{aligned} \mathbb{E}_{x \sim \mathcal{D}_x} [U(x)] &= \mathbb{E}_{x \sim \mathcal{D}_x} [U(x) \mathbb{1}\{E^+\}] + \mathbb{E}_{x \sim \mathcal{D}_x} [U(x) \mathbb{1}\{E^-\}] \\ &= \mathbb{E}_{x \sim \mathcal{D}_x} [(u(\langle w, x \rangle) - (1-\eta))^2 \mathbb{1}\{E^+\}] \\ &+ \mathbb{E}_{x \sim \mathcal{D}_x} [(u(\langle w, x \rangle) - \eta)^2 \mathbb{1}\{E^-\}]. \end{aligned}$$

We will lower bound both terms:

$$\begin{aligned} &\mathbb{E}_{x \sim \mathcal{D}_x} [(u(\langle w, x \rangle) - (1-\eta))^2 \mathbb{1}\{E^+\}] \\ &\geq \mathbb{E}_{x \sim \mathcal{D}_x} [(a - (1-\eta))^2 \mathbb{1}\{E^+\} \mathbb{1}\{u(\langle w, x \rangle) \leq a\}] , \\ &\mathbb{E}_{x \sim \mathcal{D}_x} [(u(\langle w, x \rangle) - \eta)^2 \mathbb{1}\{E^-\}] \\ &\geq \mathbb{E}_{x \sim \mathcal{D}_x} [(b - \eta)^2 \mathbb{1}\{E^-\} \mathbb{1}\{u(\langle w, x \rangle) \geq b\}] . \end{aligned}$$

Then, observe that the event $\{\langle w, x \rangle \leq \gamma/2\}$ implies the event $\{u(\langle w, x \rangle) \leq \frac{3-2\eta}{4}\}$, and similarly the event $\{\langle w, x \rangle \geq -\gamma/2\}$ implies the event $\{u(\langle w, x \rangle) \geq \frac{2\eta+1}{4}\}$. This means that

$$\begin{aligned} &\mathbb{E}_{x \sim \mathcal{D}_x} \left[\left(\frac{3-2\eta}{4} - (1-\eta) \right)^2 \mathbb{1}\{E^+\} \mathbb{1}\left\{u(\langle w, x \rangle) \leq \frac{3-2\eta}{4}\right\} \right] \\ &\geq \left(\frac{3-2\eta}{4} - (1-\eta) \right)^2 \mathbb{E}_{x \sim \mathcal{D}_x} [\mathbb{1}\{E^+\} \mathbb{1}\{\langle w, x \rangle \leq \gamma/2\}] , \text{ and} \\ &\mathbb{E}_{x \sim \mathcal{D}_x} \left[\left(\frac{2\eta+1}{4} - \eta \right)^2 \mathbb{1}\{E^-\} \mathbb{1}\left\{u(\langle w, x \rangle) \geq \frac{2\eta+1}{4}\right\} \right] \\ &\geq \left(\frac{2\eta+1}{4} - \eta \right)^2 \mathbb{E}_{x \sim \mathcal{D}_x} [\mathbb{1}\{E^-\} \mathbb{1}\{\langle w, x \rangle \geq -\gamma/2\}] . \end{aligned}$$

We combine these observations to conclude the proof,

$$\begin{aligned} &\mathbb{E}_{x \sim \mathcal{D}_x} [(u(\langle w, x \rangle) - u(\langle w^*, x \rangle))^2] \geq \frac{(1-2\eta)^2}{16} \times \\ &\mathbb{E}_{x \sim \mathcal{D}_x} [\mathbb{1}(E^+) \mathbb{1}\{\langle w, x \rangle \leq \gamma/2\} + \mathbb{1}(E^-) \mathbb{1}\{\langle w, x \rangle \geq -\gamma/2\}] \\ &\geq \frac{(1-2\eta)^2}{16} \mathbb{E}_{x \sim \mathcal{D}_x} [\mathbb{1}\{h_{w^*}(x) \langle w, x \rangle \leq \gamma/2\}] . \end{aligned}$$

□

But note that the squared loss is non-convex and so it may not be easy to optimize. Luckily, we can get a tight upper-bound with the following surrogate loss (see (Kanade, 2018)):

$$\ell(w, (x, y)) = \int_0^{\langle w, x \rangle} (u(s) - y) ds. \quad (8)$$

Note that $\ell(w, (x, y))$ is convex w.r.t w since the Hessian $\nabla_w^2 \ell(w, (x, y)) = u'(\langle w, x \rangle) x x^T$ is positive semi-definite.

Assuming our labels y have been transformed to $\{0, 1\}$ from $\{\pm 1\}$, observe that $\mathbb{E}[y|x] = u(\langle w^*, x \rangle)$. We now have the following guarantee (see e.g., (Cohen, 2014; Kanade, 2018)) for $U(x) = (u(\langle w, x \rangle) - u(\langle w^*, x \rangle))^2$:

$$\mathbb{E}_{x \sim \mathcal{D}_x} [U(x)] \leq \frac{1-2\eta}{\gamma} \mathbb{E}_{(x,y) \sim \mathcal{D}} [\ell(w, (x, y)) - \ell(w^*, (x, y))]. \quad (9)$$

Proof of Theorem 5.1. Combining Lemma 5.2 and Equation 9, we get the following guarantee for any $w \in \mathbb{R}^d$,

$$\begin{aligned} &(1-2\eta) \mathbb{E}_{x \sim \mathcal{D}_x} [\mathbb{1}\{h_{w^*}(x) \langle w, x \rangle \leq \gamma/2\}] \leq \\ &\frac{16}{\gamma} \mathbb{E}_{(x,y) \sim \mathcal{D}} [\ell(w, (x, y)) - \ell(w^*, (x, y))]. \end{aligned}$$

Thus, running Stochastic Mirror Descent with $\varepsilon' = (\varepsilon\gamma(1 - 2\eta))/16$ and $O(1/\varepsilon'^2)$ samples (labels transformed to $\{0, 1\}$), returns with high probability, a halfspace $\mathbf{w}$ such that

$$\mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \leq \varepsilon. \quad (10)$$

Then, to conclude the proof, observe that

$$\begin{aligned} & \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}} [\mathbb{1}\{y \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] = \eta \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{-h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \\ & + (1 - \eta) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \\ & = \eta(1 - \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq -\gamma/2\}]) \\ & + (1 - \eta) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \\ & \stackrel{(i)}{\leq} \eta + (1 - \eta) \mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq \gamma/2\}] \\ & \stackrel{(ii)}{\leq} \eta + \varepsilon, \end{aligned}$$

where (i) follows from that fact that $\mathbb{E}_{\mathbf{x} \sim \mathcal{D}_{\mathbf{x}}} [\mathbb{1}\{h_{\mathbf{w}^*}(\mathbf{x}) \langle \mathbf{w}, \mathbf{x} \rangle \leq -\gamma/2\}] \geq 0$, and (ii) follows from Equation 10. $\square$

References

- Angluin, D. and Laird, P. D. Learning from noisy examples. *Mach. Learn.*, 2(4):343–370, 1987. doi: 10.1007/BF00116829. URL <https://doi.org/10.1007/BF00116829>.
- Awasthi, P., Dutta, A., and Vijayaraghavan, A. On robustness to adversarial examples and polynomial optimization. In *Advances in Neural Information Processing Systems*, pp. 13737–13747, 2019.
- Balcan, M., Blum, A., and Srebro, N. A theory of learning with similarity functions. *Machine Learning*, 72(1-2):89–112, 2008. doi: 10.1007/s10994-008-5059-5. URL <https://doi.org/10.1007/s10994-008-5059-5>.
- Biggio, B., Corona, I., Maiorca, D., Nelson, B., Šrđić, N., Laskov, P., Giacinto, G., and Roli, F. Evasion attacks against machine learning at test time. In *Joint European conference on machine learning and knowledge discovery in databases*, pp. 387–402. Springer, 2013.
- Bubeck, S., Lee, Y. T., Price, E., and Razenshteyn, I. Adversarial examples from computational constraints. In *International Conference on Machine Learning*, pp. 831–840, 2019.
- Bubeck, S. et al. Convex optimization: Algorithms and complexity. *Foundations and Trends® in Machine Learning*, 8(3-4):231–357, 2015.
- Bylander, T. Learning linear threshold functions in the presence of classification noise. In Warmuth, M. K. (ed.), *Proceedings of the Seventh Annual ACM Conference on Computational Learning Theory, COLT 1994, New Brunswick, NJ, USA, July 12-15, 1994*, pp. 340–347. ACM, 1994. doi: 10.1145/180139.181176. URL <https://doi.org/10.1145/180139.181176>.
- Cohen, A. *Surrogate Loss Minimization*. PhD thesis, Hebrew University of Jerusalem, 2014.
- Cullina, D., Bhagoji, A. N., and Mittal, P. Pac-learning in the presence of adversaries. In Bengio, S., Wallach, H., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 31*, pp. 230–241. Curran Associates, Inc., 2018.
- Daniely, A. Complexity theoretic limitations on learning halfspaces. In *Proceedings of the forty-eighth annual ACM symposium on Theory of Computing*, pp. 105–117, 2016.
- Diakonikolas, I., O’Donnell, R., Servedio, R. A., and Wu, Y. Hardness results for agnostically learning low-degree polynomial threshold functions. In *Proceedings of the Twenty-Second Annual ACM-SIAM Symposium on Discrete Algorithms, SODA 2011*, pp. 1590–1606. SIAM, 2011.
- Diakonikolas, I., Gouleakis, T., and Tzamos, C. Distribution-independent pac learning of halfspaces with massart noise. In *Advances in Neural Information Processing Systems*, pp. 4751–4762, 2019a.
- Diakonikolas, I., Kane, D., and Manurangsi, P. Nearly tight bounds for robust proper learning of halfspaces with a margin. In *Advances in Neural Information Processing Systems*, pp. 10473–10484, 2019b.
- Engstrom, L., Tran, B., Tsipras, D., Schmidt, L., and Madry, A. Exploring the landscape of spatial robustness. In Chaudhuri, K. and Salakhutdinov, R. (eds.), *Proceedings of the 36th International Conference on Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pp. 1802–1811, Long Beach, California, USA, 09–15 Jun 2019. PMLR. URL <http://proceedings.mlr.press/v97/engstrom19a.html>.
- Feldman, V., Gopalan, P., Khot, S., and Ponnuswami, A. K. New results for learning noisy parities and halfspaces. In *2006 47th Annual IEEE Symposium on Foundations of Computer Science (FOCS’06)*, pp. 563–574. IEEE, 2006.
- Feldman, V., Guzman, C., and Vempala, S. S. Statistical query algorithms for mean vector estimation and stochastic convex optimization. In Klein, P. N. (ed.), *Proceedings*

- of the Twenty-Eighth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA 2017, Barcelona, Spain, Hotel Porta Fira, January 16-19, pp. 1265–1277. SIAM, 2017. doi: 10.1137/1.9781611974782.82. URL <https://doi.org/10.1137/1.9781611974782.82>.
- Goodfellow, I. J., Shlens, J., and Szegedy, C. Explaining and harnessing adversarial examples. In Bengio, Y. and LeCun, Y. (eds.), *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, 2015. URL <http://arxiv.org/abs/1412.6572>.
- Gourdeau, P., Kanade, V., Kwiatkowska, M., and Worrell, J. On the hardness of robust classification. In *Advances in Neural Information Processing Systems*, pp. 7444–7453, 2019.
- Guruswami, V. and Raghavendra, P. Hardness of learning halfspaces with noise. *SIAM Journal on Computing*, 39(2):742–765, 2009.
- Kanade, V. Computational learning theory notes - 8 : Learning real-valued functions, 2018.
- Kang, D., Sun, Y., Hendrycks, D., Brown, T., and Steinhardt, J. Testing robustness against unforeseen adversaries. *CoRR*, abs/1908.08016, 2019. URL <http://arxiv.org/abs/1908.08016>.
- Khim, J. and Loh, P.-L. Adversarial risk bounds for binary classification via function transformation. *arXiv preprint arXiv:1810.09519*, 2018.
- Lee, Y. T., Sidford, A., and Vempala, S. S. Efficient convex optimization with membership oracles. In *Conference On Learning Theory*, pp. 1292–1294, 2018.
- Maass, W. and Turán, G. How fast can a threshold gate learn? In *Proceedings of a workshop on Computational learning theory and natural learning systems (vol. 1): constraints and prospects: constraints and prospects*, pp. 381–414, 1994.
- Montasser, O., Hanneke, S., and Srebro, N. V_c classes are adversarially robustly learnable, but only improperly. In Beygelzimer, A. and Hsu, D. (eds.), *Proceedings of the Thirty-Second Conference on Learning Theory*, volume 99 of *Proceedings of Machine Learning Research*, pp. 2512–2530, Phoenix, USA, 25–28 Jun 2019. PMLR.
- Rosenblatt, F. The perceptron: a probabilistic model for information storage and organization in the brain. *Psychological review*, 65(6):386, 1958.
- Schmidt, L., Santurkar, S., Tsipras, D., Talwar, K., and Madry, A. Adversarially robust generalization requires more data. In Bengio, S., Wallach, H. M., Larochelle, H., Grauman, K., Cesa-Bianchi, N., and Garnett, R. (eds.), *Advances in Neural Information Processing Systems 31: Annual Conference on Neural Information Processing Systems 2018, NeurIPS 2018, 3-8 December 2018, Montréal, Canada*, pp. 5019–5031, 2018.
- Szegedy, C., Zaremba, W., Sutskever, I., Bruna, J., Erhan, D., Goodfellow, I. J., and Fergus, R. Intriguing properties of neural networks. In Bengio, Y. and LeCun, Y. (eds.), *2nd International Conference on Learning Representations, ICLR 2014, Banff, AB, Canada, April 14-16, 2014, Conference Track Proceedings*, 2014. URL <http://arxiv.org/abs/1312.6199>.
- Vapnik, V. *Estimation of Dependencies Based on Empirical Data*. Springer-Verlag, New York, 1982.
- Yin, D., Ramchandran, K., and Bartlett, P. L. Rademacher complexity for adversarially robust generalization. In Chaudhuri, K. and Salakhutdinov, R. (eds.), *Proceedings of the 36th International Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA*, volume 97 of *Proceedings of Machine Learning Research*, pp. 7085–7094. PMLR, 2019. URL <http://proceedings.mlr.press/v97/yin19b.html>.