

An Algorithm with Optimal Dimension-Dependence for Zero-Order Nonsmooth Nonconvex Stochastic Optimization

Guy Kornowski

GUY.KORNOWSKI@WEIZMANN.AC.IL

*Department of Computer Science and Applied Mathematics
Weizmann Institute of Science
Rehovot, Israel*

Ohad Shamir

OHAD.SHAMIR@WEIZMANN.AC.IL

*Department of Computer Science and Applied Mathematics
Weizmann Institute of Science
Rehovot, Israel*

Editor: Francesco Orabona

Abstract

We study the complexity of producing (δ, ϵ) -stationary points of Lipschitz objectives which are possibly neither smooth nor convex, using only noisy function evaluations. Recent works proposed several stochastic zero-order algorithms that solve this task, all of which suffer from a dimension-dependence of $\Omega(d^{3/2})$ where d is the dimension of the problem, which was conjectured to be optimal. We refute this conjecture by providing a faster algorithm that has complexity $O(d\delta^{-1}\epsilon^{-3})$, which is optimal (up to numerical constants) with respect to d and also optimal with respect to the accuracy parameters δ, ϵ , thus solving an open question due to Lin et al. (2022). Moreover, the convergence rate achieved by our algorithm is also optimal for smooth objectives, proving that in the nonconvex stochastic zero-order setting, *nonsmooth optimization is as easy as smooth optimization*. We provide algorithms that achieve the aforementioned convergence rate in expectation as well as with high probability. Our analysis is based on a simple yet powerful lemma regarding the Goldstein-subdifferential set, which allows utilizing recent advancements in first-order nonsmooth nonconvex optimization.

Keywords: nonsmooth nonconvex optimization, stochastic optimization, zero-order, gradient-free, Goldstein subdifferential

1. Introduction

We consider the problem of optimizing a stochastic objective of the form

$$f(\mathbf{x}) = \mathbb{E}_{\xi \sim \Xi}[F(\mathbf{x}; \xi)]$$

where the stochastic components $F(\cdot; \xi) : \mathbb{R}^d \rightarrow \mathbb{R}$ are Lipschitz continuous, yet possibly not smooth nor convex. We consider stochastic zero-order (also known as *gradient-free* or *derivative-free*) algorithms that have access only to noisy function evaluations. At each time step, the algorithm draws $\xi \sim \Xi$ and can observe $F(\mathbf{x}, \xi)$ for points $\mathbf{x} \in \mathbb{R}^d$ of its choice. Problems of this type arise throughout machine learning, control theory and finance, in applications in which gradients are expensive (or even impossible) to evaluate, see for example the book by Spall (2005) for an overview. Although in the convex setting the

complexity of such algorithms is relatively well understood (Agarwal and Dekel, 2010; Duchi et al., 2015; Nesterov and Spokoiny, 2017; Shamir, 2017), much less is known about the nonsmooth nonconvex setting. This challenging setting is of major interest in modern deep learning applications, where objective functions of interest are typically neither smooth or convex. For example, stochastic zero-order optimization methods were applied to fine-tune large language models (Malladi et al., 2023).

Recently, Lin et al. (2022) proposed a gradient-free algorithm that produces a (δ, ϵ) -stationary point using $O(d^{3/2}\delta^{-1}\epsilon^{-4})$ function evaluations. Following Zhang et al. (2020), recall that a point $\mathbf{x}$ is called a (δ, ϵ) -stationary point if there exists a convex combination of gradients in a δ -neighborhood of $\mathbf{x}$ whose norm is less than ϵ (see Section 2 for a reminder of relevant definitions). Lin et al. (2022) posed the question as to whether this super-linear dimension dependence is inevitable or not. The aforementioned complexity was very recently improved to $O(d^{3/2}\delta^{-1}\epsilon^{-3})$ by Chen et al. (2023), yet notably, this result still suffers from the same super-linear dimension dependence.

In particular, as pointed out by Lin et al., this dimension dependence is $\Omega(\sqrt{d})$ worse than that of stochastic zero-order *smooth* nonconvex optimization, a setting in which it is possible to find an ϵ -stationary point (i.e. $\mathbf{x}$ such that $\|\nabla f(\mathbf{x})\| \leq \epsilon$) using $O(d\epsilon^{-4})$ noisy function evaluations (Ghadimi and Lan, 2013). This led the authors to conjecture that stochastic zero-order nonsmooth nonconvex optimization is “likely to be intrinsically harder” than its smooth counterpart.

Our main contribution resolves this open question, showing that this is actually *not* the case. We propose a faster zero-order algorithm for nonsmooth nonconvex optimization, which requires only $O(d\delta^{-1}\epsilon^{-3})$ noisy function evaluations. As we will soon argue, this complexity has an optimal linear-dependence on the dimension d , while also obtaining the optimal dependence with respect to the accuracy parameters δ and ϵ . All of these dependencies are known to be optimal even if $f(\cdot)$ is smooth, implying that in the stochastic zero-order setting, *nonsmooth nonconvex optimization is as easy as smooth nonconvex optimization*. Moreover, when the objective is smooth, our proposed algorithm automatically recovers the $O(d\epsilon^{-4})$ complexity of stochastic gradient-free smooth nonconvex optimization (Ghadimi and Lan, 2013). Whether this adaptivity property is possible was originally raised as an open question by Zhang et al. (2020) in the context of first-order algorithms (that have access to gradient information), and was recently confirmed by Cutkosky et al. (2023). Our result extends the resolution of this question to the case of zero-order algorithms.

As previously mentioned, the dependencies on d, δ and ϵ we obtain are all optimal. Indeed, the linear dimension dependence is well-known to be inevitable for gradient-free algorithms even in the strictly-easier cases of smooth or convex optimization (Duchi et al., 2015), while the implied ϵ^{-4} factor is known to be inevitable even in the strictly-easier case of stochastic first-order smooth optimization with exact function evaluations (Arjevani et al., 2023).¹ Interestingly, in terms of the dependence on δ and ϵ , the convergence rate we obtain is as fast as the currently best-known *deterministic first-order* algorithms for nonsmooth nonconvex optimization (Zhang et al., 2020; Tian et al., 2022; Davis et al., 2022). This is in stark contrast to smooth nonconvex optimization, in which optimal stochastic and deterministic methods have disparate complexities on the order of ϵ^{-4} and ϵ^{-2} , respectively

1. Technically, the lower bound construction in Arjevani et al. (2023) is not globally Lipschitz, yet a slight modification of it which appears in Cutkosky et al. (2023, Appendix F) is.

(Arjevani et al., 2023; Carmon et al., 2020). We also note that our algorithm is a factor of $\Omega(\sqrt{d})$ faster even than the previously best-known rate for deterministic (i.e. noiseless, when $\Xi = \{\xi\}$) zero-order nonsmooth nonconvex optimization, and that it also obtains an improved dependence on the Lipschitz parameter when compared to the previously mentioned works in this setting. Finally, we remark that while the dependence on each parameter by itself (i.e. d, δ and ϵ) is already known to be optimal, currently there is no result in the literature formally proving a lower bound *jointly* in these parameters, namely of the form $\Omega(d\delta^{-1}\epsilon^{-3})$, which will automatically follow from a smooth (joint) lower bound of the form $\Omega(d\epsilon^{-4})$. We conjecture these results should be obtainable by modifying the analysis of Arjevani et al. (2023) to incorporate zero-order oracles.

2. Preliminaries.

Notation. We use bold-faced font to denote vectors, e.g. $\mathbf{x} \in \mathbb{R}^d$, and denote by $\|\mathbf{x}\|$ the Euclidean norm. We denote by $[n] := \{1, \dots, n\}$, $\mathbb{B}(\mathbf{x}, \delta) := \{\mathbf{y} \in \mathbb{R}^d : \|\mathbf{y} - \mathbf{x}\| \leq \delta\}$, and by $\mathbb{S}^{d-1} \subset \mathbb{R}^d$ the unit sphere. We denote by $\text{conv}(\cdot)$ the convex hull operator, and by $\text{Unif}(A)$ the uniform measure over a set A . We use the standard big-O notation, with $O(\cdot)$, $\Theta(\cdot)$ and $\Omega(\cdot)$ hiding absolute constants that do not depend on problem parameters, $\tilde{O}(\cdot)$ and $\tilde{\Omega}(\cdot)$ hiding absolute constants and additional logarithmic factors.

Nonsmooth analysis. We call a function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ L -Lipschitz if for any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$: $|f(\mathbf{x}) - f(\mathbf{y})| \leq L \|\mathbf{x} - \mathbf{y}\|$, and H -smooth if it is differentiable and $\nabla f : \mathbb{R}^d \rightarrow \mathbb{R}^d$ is H -Lipschitz, namely for any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$: $\|\nabla f(\mathbf{x}) - \nabla f(\mathbf{y})\| \leq H \|\mathbf{x} - \mathbf{y}\|$. By Rademacher's theorem, Lipschitz functions are differentiable almost everywhere (in the sense of Lebesgue). Hence, for any Lipschitz function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ and point $\mathbf{x} \in \mathbb{R}^d$ the Clarke subgradient set (Clarke, 1990) can be defined as

$$\partial f(\mathbf{x}) := \text{conv}\{\mathbf{g} : \mathbf{g} = \lim_{n \rightarrow \infty} \nabla f(\mathbf{x}_n), \mathbf{x}_n \rightarrow \mathbf{x}\} ,$$

namely, the convex hull of all limit points of $\nabla f(\mathbf{x}_n)$ over all sequences of differentiable points which converge to $\mathbf{x}$. Note that if the function is continuously differentiable at a point or convex, the Clarke subdifferential reduces to the gradient or subgradient in the convex analytic sense, respectively. We say that a point $\mathbf{x}$ is an ϵ -stationary point of $f(\cdot)$ if $\min\{\|\mathbf{g}\| : \mathbf{g} \in \partial f(\mathbf{x})\} \leq \epsilon$. Furthermore, given $\delta \geq 0$ the Goldstein δ -subdifferential (Goldstein, 1977) of f at $\mathbf{x}$ is the set

$$\partial_\delta f(\mathbf{x}) := \text{conv}\left(\bigcup_{\mathbf{y} \in \mathbb{B}(\mathbf{x}, \delta)} \partial f(\mathbf{y})\right) ,$$

namely all convex combinations of gradients at points in a δ -neighborhood of $\mathbf{x}$. We denote the minimum-norm element of the Goldstein δ -subdifferential by

$$\bar{\partial}_\delta f(\mathbf{x}) := \arg \min_{\mathbf{g} \in \partial_\delta f(\mathbf{x})} \|\mathbf{g}\| .$$

Definition 1 A point $\mathbf{x}$ is called a (δ, ϵ) -stationary point of $f(\cdot)$ if $\|\bar{\partial}_\delta f(\mathbf{x})\| \leq \epsilon$.

Note that a point is ϵ -stationary if and only if it is (δ, ϵ) -stationary for all $\delta \geq 0$ (Zhang et al., 2020, Lemma 7). Moreover, if f is H -smooth and $\mathbf{x}$ is a $(\frac{\epsilon}{3H}, \frac{\epsilon}{3})$ -stationary point of f , then it is also ϵ -stationary (Zhang et al., 2020, Proposition 6).

Randomized smoothing. Given a Lipschitz function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, we define its uniform smoothing

$$f_\rho(\mathbf{x}) := \mathbb{E}_{\mathbf{z} \sim \text{Unif}(\mathbb{B}(\mathbf{0}, 1))} [f(\mathbf{x} + \rho \mathbf{z})] .$$

It is well known (cf. Yousefian et al., 2012) that if f is L_0 -Lipschitz, then

- f_ρ is L_0 -Lipschitz;
- f_ρ is $O(\sqrt{d}L_0\rho^{-1})$ -smooth;
- $|f(\mathbf{x}) - f_\delta(\mathbf{x})| \leq \rho L_0$ for all $\mathbf{x} \in \mathbb{R}^d$.

Setting. We consider optimization objectives of the form $f(\mathbf{x}) = \mathbb{E}_{\xi \sim \Xi} [F(\mathbf{x}; \xi)]$, where $\xi \sim \Xi$ is a random variable. We impose the assumption that the stochastic components $F(\cdot; \xi) : \mathbb{R}^d \rightarrow \mathbb{R}$ are Lipschitz continuous, possibly with a varying Lipschitz constant:

Assumption 2 *For any ξ , the function $F(\cdot; \xi)$ is $L(\xi)$ -Lipschitz. Moreover, we assume $L(\xi)$ has a bounded second moment: Namely, there exists $L_0 > 0$ such that*

$$\mathbb{E}_{\xi \sim \Xi} [L(\xi)^2] \leq L_0^2 .$$

We note that Assumption 2 is weaker than assuming $F(\cdot; \xi)$ is L_0 -Lipschitz for all ξ . We also remark that in case Ξ is supported on a single point then the optimization problem reduces to that of a L_0 -Lipschitz objective using exact evaluations.

3. Algorithms and Main Results

Before formally presenting our main result, we find it insightful to stress out the key idea, and in particular how our algorithm differs than those of Lin et al. (2022); Chen et al. (2023). The main strategy employed by both of these papers is based on the following result.

Proposition 3 (Lin et al., 2022, Theorem 3.1) *For any $\rho \geq 0 : \nabla f_\rho(\mathbf{x}) \in \partial_\rho f(\mathbf{x})$. Hence, for $\rho = \delta$, if $\mathbf{x}$ is an ϵ -stationary point of f_δ , then it is a (δ, ϵ) -stationary point of f .*

Following this observation, both papers set out to design algorithms that produce an ϵ -stationary point of f_δ . A well known technique (which we formally recall later on) allows to use two possibly noisy evaluations of f in order to produce a stochastic *first-order* oracle of f_δ whose second moment is bounded by $\sigma^2 = O(d)$. Noting that f_δ is $L_1 = O(\sqrt{d}/\delta)$ -smooth, the standard analysis of stochastic gradient descent (SGD) for smooth nonconvex optimization shows that it obtains an ϵ -stationary point of f_δ within $O(\sigma^2 L_1 \epsilon^{-4}) = O(d^{3/2} \delta^{-1} \epsilon^{-4})$ oracle calls, recovering the main result of Lin et al. (2022). The improved ϵ -dependence due to Chen et al. (2023) was achieved by employing a variance-reduction method instead of plain SGD, though other than that, their main algorithmic strategy and analysis are the same.

Moreover, the algorithmic strategy we have described seems to reveal a barrier, (mistakenly) suggesting the $d^{3/2}$ dependence is unavoidable. Indeed, it is relatively straightforward to see that any gradient estimator which is based on a constant number of function evaluations must have variance of at least $\sigma^2 = \Omega(d)$, while it is also known that any efficient

smoothing technique must suffer from a smoothness parameter of at least $L_1 = \Omega(\sqrt{d})$ (Kornowski and Shamir, 2022). Since the complexity of any stochastic first-order method for smooth nonconvex optimization must scale at least as $\Omega(\sigma^2 L_1)$ (Arjevani et al., 2023) which in this case is unavoidably $\Omega(d^{3/2})$, we are stuck with this factor.

The main technical ingredient that allows us to reduce this factor is the following result (to be proved in Section 4), which examines the Goldstein δ -subdifferential set under randomized smoothing.

Lemma 4 *For any $\rho, \nu \geq 0$: $\partial_\nu f_\rho(\mathbf{x}) \subseteq \partial_{\rho+\nu} f(\mathbf{x})$. Hence, if $\mathbf{x}$ is an (ν, ϵ) -stationary point of f_ρ , then it is a $(\rho + \nu, \epsilon)$ -stationary point of f . In particular, as long as $\rho + \nu \leq \delta$ it is a (δ, ϵ) -stationary point of f .*

Note that the lemma above strictly generalizes Proposition 3 (Lin et al., 2022, Theorem 3.1) which is readily recovered by plugging $\nu = 0$. The utility of this result is that it allows to replace the task of finding an ϵ -stationary point of f_δ to that of finding a (δ, ϵ) -stationary point of it (disregarding a constant factor multiplying δ). To see why this is beneficial, recall that while f_δ is $O(\sqrt{d}L_0\delta^{-1})$ -smooth, it is merely L_0 -Lipschitz! Thus using a stochastic first-order *nonsmooth* nonconvex algorithm which scales with the Lipschitz parameter (instead of the smoothness parameter), we save a whole $\Omega(\sqrt{d})$ factor, yielding the optimal dimension dependence. In particular, using the optimal stochastic first-order algorithm of Cutkosky et al. (2023) that has complexity $O(\sigma^2 \delta^{-1} \epsilon^{-3})$, as described in Algorithm 1, results in the following convergence guarantee:²

Theorem 5 *Let $\delta, \epsilon \in (0, 1)$, and suppose $f(\mathbf{x}_0) - \inf_{\mathbf{x}} f(\mathbf{x}) \leq \Delta$. Under Assumption 2, there exists*

$$T = O\left(\frac{dL_0^2\Delta}{\delta\epsilon^3}\right)$$

such that setting

$$\rho = \min\left\{\frac{\delta}{2}, \frac{\Delta}{L_0}\right\}, \quad \nu = \max\left\{\frac{\delta}{2}, \delta - \frac{\Delta}{L_0}\right\}, \quad D = \left(\frac{(\Delta + \rho L_0)\sqrt{\nu}}{\sqrt{d}L_0T}\right)^{2/3}, \quad \eta = \frac{\Delta + \rho L_0}{dL_0^2T},$$

and running Algorithm 1 with Algorithm 2 as a subroutine, outputs a point $\mathbf{x}^{\text{out}}$ satisfying $\mathbb{E}[\|\bar{\partial}_\delta f(\mathbf{x}^{\text{out}})\|] \leq \epsilon$ using $2T$ noisy function evaluations.

Parallel complexity. At each iteration, Algorithm 1 determines $\mathbf{g}_t$ by calling Algorithm 2 (GRADESTIMATOR), which requires 2 evaluations of $F(\cdot; \xi_t)$. More generally, $\mathbf{g}_t$ can be set as the average of k independent, possibly parallel calls to this subroutine, which would require $2k$ function evaluations. By an easy generalization of the second-order moment bound which we present in Lemma 7 (as part of the proof of Theorem 5), this averaging would decrease the second-moment on the order of

$$\mathbb{E}[\|\mathbf{g}_t\|^2 | \mathbf{x}_t, s_t, \Delta_t] \lesssim \frac{dL_0^2}{k} + \|\mathbb{E}[\mathbf{g}_t]\|^2 \leq L_0^2 \left(\frac{d}{k} + 1\right).$$

2. It is interesting to note that using the stochastic algorithm of Zhang et al. (2020) (instead of Algorithm 1), when paired with our analysis, yields the desired linear dimension dependence as well – albeit with a worse convergence rate with respect to ϵ , on the order of $d\delta^{-1}\epsilon^{-4}$.

Algorithm 1 Optimal Stochastic Nonsmooth Nonconvex Optimization Algorithm

```

1: Input: Initialization  $\mathbf{x}_0 \in \mathbb{R}^d$ , smoothing parameter  $\rho > 0$ , accuracy parameter  $\nu > 0$ ,
   clipping parameter  $D > 0$ , step size  $\eta > 0$ , iteration budget  $T \in \mathbb{N}$ .
2: Initialize:  $\Delta_1 = \mathbf{0}$ 
3: for  $t = 1, \dots, T$  do
4:   Sample  $\xi_t \sim \Xi$ 
5:   Sample  $s_t \sim \text{Unif}[0, 1]$ 
6:    $\mathbf{x}_t = \mathbf{x}_{t-1} + \Delta_t$ 
7:    $\mathbf{z}_t = \mathbf{x}_{t-1} + s_t \Delta_t$ 
8:    $\mathbf{g}_t = \text{GRADESTIMATOR}(\mathbf{z}_t, \rho, \xi_t)$  ▷ Uses two noisy function evaluations
9:    $\Delta_{t+1} = \min\left(1, \frac{D}{\|\Delta_t - \eta \mathbf{g}_t\|}\right) \cdot (\Delta_t - \eta \mathbf{g}_t)$ 
10: end for
11:  $M = \lfloor \frac{\nu}{D} \rfloor$ ,  $K = \lfloor \frac{T}{M} \rfloor$ 
12: for  $k = 1, \dots, K$  do
13:    $\bar{\mathbf{x}}_k = \frac{1}{M} \sum_{m=1}^M \mathbf{z}_{(k-1)M+m}$ 
14: end for
15:  $k_{\text{out}} \sim \text{Unif}\{1, \dots, K\}$ 
16:  $\mathbf{x}^{\text{out}} = \bar{\mathbf{x}}_{k_{\text{out}}}$ 
17: Output:  $\mathbf{x}^{\text{out}}$ ,  $(\mathbf{z}_{(k_{\text{out}}-1)M+m})_{m=1}^M$ .

```

Algorithm 2 GRADESTIMATOR($\mathbf{x}, \rho, \xi$)

```

1: Input: Point  $\mathbf{x} \in \mathbb{R}^d$ , smoothing parameter  $\rho > 0$ , random seed  $\xi$ .
2: Sample  $\mathbf{w} \sim \text{Unif}(\mathbb{S}^{d-1})$ 
3: Evaluate  $F(\mathbf{x} + \rho \mathbf{w}; \xi)$  and  $F(\mathbf{x} - \rho \mathbf{w}; \xi)$ 
4:  $\mathbf{g} = \frac{d}{2\rho} (F(\mathbf{x} + \rho \mathbf{w}; \xi) - F(\mathbf{x} - \rho \mathbf{w}; \xi)) \mathbf{w}$  ▷ Unbiased estimator of  $\nabla f_\rho(\mathbf{x})$ 
5: Output:  $\mathbf{g}$ .

```

With the rest of the proof of Theorem 5 as is, this yields an expected number of rounds of

$$T = O\left(\left(\frac{d}{k} + 1\right) \cdot \frac{\Delta L_0^2}{\delta \epsilon^3}\right),$$

though the total number of queries would be k times larger than above, and equal to

$$O\left((d + k) \cdot \frac{\Delta L_0^2}{\delta \epsilon^3}\right).$$

In particular, letting $k = \Theta(d)$ removes the dimension dependence in the parallel complexity altogether, while maintaining the same complexity overall (up to a constant). Notably, this even matches the currently best-known complexity for deterministic first-order algorithms under the discussed setting (Zhang et al., 2020; Tian et al., 2022; Davis et al., 2022).

The trick of averaging gradient estimators in order to reduce the dimension-factor in the parallel complexity is applicable to smooth or convex optimization as well, though under those settings the resulting overall complexity is still significantly worse (in terms of the ϵ -dependence) than optimal deterministic first-order methods, as mentioned in the introduction (cf. Duchi et al., 2018; Bubeck et al., 2019).

High probability guarantee. While Theorem 1 shows that Algorithm 1 yields the desired expected complexity, many practical applications require high probability bounds, namely producing a point $\mathbf{x}$ such that

$$\Pr [\|\bar{\partial}_\delta f(\mathbf{x})\| \leq \epsilon] \geq 1 - \gamma$$

for some small $\gamma > 0$. A naive application of Markov's inequality to the expected complexity shows that Algorithm 1 produces such a point within

$$T = O\left(\frac{dL_0^2\Delta}{\delta\epsilon^3\gamma^3}\right) \quad (1)$$

noisy function evaluations, which is rather crude with respect to the probability parameter γ . Adapting a technique due to Ghadimi and Lan (2013) to our setting, we can design an algorithm with a significantly tighter high-probability bound. The original idea of Ghadimi and Lan (2013) for the case of smooth stochastic optimization, which was also used by Lin et al. (2022), consists of several independent calls to the main algorithm, yielding a list of candidate points. Subsequently, a post-optimization phase estimates the gradient norm of any such point, returning the minimal — which is proved likely to succeed due to a concentration argument. We note that adapting this technique to our setting is not trivial, since the post-optimization phase should attempt at estimating $\|\bar{\partial}_\delta f(\cdot)\|$ rather than $\|\nabla f(\cdot)\|$, which is hard in general. Luckily, using Lemma 4, the former can be bounded by $\|\bar{\partial}_\nu f_\rho(\cdot)\|$, which in turn can be bounded (with high probability) using a sequence of evaluations at nearby points. This procedure is described in Algorithm 3, whose convergence rate is presented in the following theorem.

Theorem 6 *Let $\gamma, \delta, \epsilon \in (0, 1)$, and suppose $f(\mathbf{x}_0) - \inf_{\mathbf{x}} f(\mathbf{x}) \leq \Delta$. Under Assumption 2, there exist $T = O\left(\frac{d\Delta L_0^2}{\delta\epsilon^3}\right)$, $R = O(\log(1/\gamma))$, $S = O\left(\frac{\log(1/\gamma)}{\gamma}\right)$ such that setting $\rho = \min\left\{\frac{\delta}{2}, \frac{\Delta}{L_0}\right\}$, $\nu = \max\left\{\frac{\delta}{2}, \delta - \frac{\Delta}{L_0}\right\}$, $D = \left(\frac{(\Delta + \rho L_0)\sqrt{\nu}}{\sqrt{d}L_0 T}\right)^{2/3}$, $\eta = \frac{\Delta + \rho L_0}{dL_0^2 T}$, and running Algorithm 3 outputs a point $\mathbf{x}^{\text{out}}$ satisfying*

$$\Pr [\|\bar{\partial}_\delta f(\mathbf{x}^{\text{out}})\| \leq \epsilon] \geq 1 - \gamma$$

using

$$O\left(\frac{dL_0^2\Delta \log(1/\gamma)}{\delta\epsilon^3} + \frac{dL_0^2 \log^2(1/\gamma)}{\gamma\epsilon^2}\right)$$

noisy function evaluations.

Notably, the number of function evaluations guaranteed by the theorem above is significantly smaller than in Eq. (1). We also remark that even in the easier case in which the function evaluations are noiseless, when relying on a local information (e.g. zero-order or even first-order evaluations), a lack of smoothness or convexity provably necessitates the use of randomization in the optimization algorithm when applied to Lipschitz objectives, thus resorting to high probability guarantees is inevitable (Jordan et al., 2023).

Algorithm 3 Algorithm with Post-Optimization Validation

```

1: Input: Initialization  $\mathbf{x}_0 \in \mathbb{R}^d$ , smoothing parameter  $\rho > 0$ , accuracy parameter  $\nu > 0$ ,
   clipping parameter  $D > 0$ , step size  $\eta > 0$ , iteration budget per round  $T \in \mathbb{N}$ , number
   of rounds  $R \in \mathbb{N}$ , validation sample size  $S \in \mathbb{N}$ .
2: Initialize:  $M = \lfloor \frac{\nu}{D} \rfloor$ 
3: for  $r = 1, \dots, R$  do
4:    $\mathbf{x}_{\text{out}}^r, (\mathbf{z}_m^r)_{m=1}^M = \text{Algorithm 1}(\mathbf{x}_0, \rho, \nu, D, \eta, T)$ 
5:   for  $s = 1, \dots, S$  do
6:     for  $m = 1, \dots, M$  do
7:       Sample  $\xi_{m,s} \sim \Xi$ 
8:        $\mathbf{g}_{m,s}^r = \text{GRADESTIMATOR}(\mathbf{z}_m^r, \rho, \xi_{m,s}) \triangleright \text{Unbiased estimator of } \nabla f_\rho(\mathbf{z}_m^r)$ 
9:     end for
10:     $\hat{\mathbf{g}}_s^r = \frac{1}{M} \sum_{m=1}^M \mathbf{g}_{m,s}^r$ 
11:  end for
12:   $\hat{\mathbf{g}}^r = \frac{1}{S} \sum_{s=1}^S \hat{\mathbf{g}}_s^r$ 
13: end for
14:  $r^* = \arg \min_{r \in [R]} \|\hat{\mathbf{g}}^r\|$ 
15: Output:  $\mathbf{x}_{\text{out}}^{r^*}$ .

```

4. Proof of Theorem 5

As previously discussed, the key to obtaining the improved rate is Lemma 4. We start by proving it, followed by two additional propositions, after which we combine the ingredients in order to conclude the proof.

Proof [Proof of Lemma 4] Let $\mathbf{g} \in \partial_\nu f_\rho(\mathbf{x})$. Then, by definition, there exist $\mathbf{y}_1, \dots, \mathbf{y}_k \in \mathbb{B}(\mathbf{x}, \nu)$ (for some $k \in \mathbb{N}$) such that $\mathbf{g} = \sum_{i \in [k]} \lambda_i \nabla f_\rho(\mathbf{y}_i)$, where $\lambda_1, \dots, \lambda_k \geq 0$ with $\sum_{i \in [k]} \lambda_i = 1$. By Proposition 3 we have for all $i \in [k]$:

$$\nabla f_\rho(\mathbf{y}_i) \in \partial_\rho f(\mathbf{y}_i) . \quad (2)$$

Further note that since $\|\mathbf{y}_i - \mathbf{x}\| \leq \nu$, then by definition

$$\partial_\rho f(\mathbf{y}_i) \subseteq \partial_{\rho+\nu} f(\mathbf{x}) . \quad (3)$$

By combining Eq. (2) and Eq. (3) we get that for all $i \in [k]$: $\nabla f_\rho(\mathbf{y}_i) \in \partial_{\rho+\nu} f(\mathbf{x})$. Since $\partial_{\rho+\nu} f(\mathbf{x})$ is a convex set, we get that

$$\mathbf{g} = \sum_{i \in [k]} \lambda_i \nabla f_\rho(\mathbf{y}_i) \in \partial_{\rho+\nu} f(\mathbf{x}) .$$

■

The following lemma is essentially due to Shamir (2017), showing that it is possible to construct a gradient estimator whose second moment scales linearly with respect to the dimension d .

Lemma 7 *Let*

$$\mathbf{g}_t = \frac{d}{2\rho} (F(\mathbf{x}_t + s_t \mathbf{\Delta}_t + \rho \mathbf{w}_t; \xi_t) - F(\mathbf{x}_t + s_t \mathbf{\Delta}_t - \rho \mathbf{w}_t; \xi_t)) \mathbf{w}_t ,$$

as generated by `GRADESTIMATOR` (Algorithm 2) when called at iteration t of Algorithm 1. Then

$$\mathbb{E}_{\xi_t, \mathbf{w}_t} [\mathbf{g}_t | \mathbf{x}_{t-1}, s_t, \mathbf{\Delta}_t] = \nabla f_\rho(\mathbf{x}_{t-1} + s_t \mathbf{\Delta}_t) = \nabla f_\rho(\mathbf{z}_t)$$

and

$$\mathbb{E}_{\xi_t, \mathbf{w}_t} [\|\mathbf{g}_t\|^2 | \mathbf{x}_{t-1}, s_t, \mathbf{\Delta}_t] \leq 16\sqrt{2\pi}dL_0^2 .$$

Proof For the sake of notational simplicity, we omit the subscript t throughout the proof. For the first claim, since $-\mathbf{w} \sim \mathbf{w}$ are identically distributed, we have

$$\begin{aligned} \mathbb{E}_{\xi, \mathbf{w}} [\mathbf{g} | \mathbf{x}, s, \mathbf{\Delta}] &= \mathbb{E}_{\xi, \mathbf{w}} \left[\frac{d}{2\rho} (F(\mathbf{x} + s\mathbf{\Delta} + \rho\mathbf{w}; \xi) - F(\mathbf{x} + s\mathbf{\Delta} - \rho\mathbf{w}; \xi)) \mathbf{w} | \mathbf{x}, s, \mathbf{\Delta} \right] \\ &= \frac{1}{2} \left(\mathbb{E}_{\xi, \mathbf{w}} \left[\frac{d}{\rho} F(\mathbf{x} + s\mathbf{\Delta} + \rho\mathbf{w}; \xi) \mathbf{w} | \mathbf{x}, s, \mathbf{\Delta} \right] \right. \\ &\quad \left. + \mathbb{E}_{\xi, \mathbf{w}} \left[\frac{d}{\rho} F(\mathbf{x} + s\mathbf{\Delta} + \rho(-\mathbf{w}); \xi) (-\mathbf{w}) | \mathbf{x}, s, \mathbf{\Delta} \right] \right) \\ &= \mathbb{E}_{\xi, \mathbf{w}} \left[\frac{d}{\rho} F(\mathbf{x} + s\mathbf{\Delta} + \rho\mathbf{w}; \xi) \mathbf{w} | \mathbf{x}, s, \mathbf{\Delta} \right] . \end{aligned}$$

Using the law of total expectation, we get

$$\begin{aligned} \mathbb{E}_{\xi, \mathbf{w}} [\mathbf{g} | \mathbf{x}, s, \mathbf{\Delta}] &= \mathbb{E}_{\mathbf{w}} \left[\frac{d}{\rho} \mathbb{E}_{\xi} [F(\mathbf{x} + s\mathbf{\Delta} + \rho\mathbf{w}; \xi) \mathbf{w} | \mathbf{w}, \mathbf{x}, s, \mathbf{\Delta}] | \mathbf{x}, s, \mathbf{\Delta} \right] \\ &= \mathbb{E}_{\mathbf{w}} \left[\frac{d}{\rho} f(\mathbf{x} + s\mathbf{\Delta} + \rho\mathbf{w}) \mathbf{w} | \mathbf{x}, s, \mathbf{\Delta} \right] \\ &= \nabla f_\rho(\mathbf{x} + s\mathbf{\Delta}) , \end{aligned}$$

where the last equality is due to Flaxman et al. (2005, Lemma 2.1). The second moment bound follows from Shamir (2017, Lemma 10), with the explicit constant pointed out by Lin et al. (2022, Lemma E.1). ■

The following result of Cutkosky et al. (2023) provides a stochastic first-order nonsmooth nonconvex optimization method, whose convergence scales linearly with the second-moment of the gradient estimator.

Theorem 8 (Cutkosky et al., 2023) *Let $\nu, \epsilon \in (0, 1)$, and let $h : \mathbb{R}^d \rightarrow \mathbb{R}$ be an L -Lipschitz function such that $h(\mathbf{x}_0) - \inf_{\mathbf{x}} h(\mathbf{x}) \leq \Delta_h$. Suppose `GRADESTIMATOR`($\mathbf{x}, \xi$) returns an unbiased gradient estimator of $\nabla h(\mathbf{x})$ whose second moment is bounded by σ^2 . Then there exists $T = O\left(\frac{\sigma^2 \Delta_h}{\nu \epsilon^3}\right)$ such that setting $\eta = \frac{\Delta_h}{\sigma^2 T}$, $D = \left(\frac{(\nu)^{1/2} \Delta_h}{\sigma T}\right)^{2/3}$ and running Algorithm 1 uses T calls to `GRADESTIMATOR` and satisfies*

- $\mathbf{z}_{(k-1)M+m} \in \mathbb{B}(\bar{\mathbf{x}}_k, \nu)$ for all $m \in [M], k \in [K]$ (where $M, K, (\mathbf{z}_t)_{t=1}^T$ are defined in the algorithm).

$$\bullet \mathbb{E}_{\mathbf{z}_1, \dots, \mathbf{z}_T} \left[\frac{1}{K} \sum_{k=1}^K \left\| \frac{1}{M} \sum_{m=1}^M \nabla h(\mathbf{z}_{(k-1)M+m}) \right\| \right] \leq \epsilon .$$

In particular, its output $\mathbf{x}^{\text{out}} \sim \text{Unif}\{\bar{\mathbf{x}}_1, \dots, \bar{\mathbf{x}}_K\}$ satisfies $\mathbb{E} [\|\bar{\partial}_\nu f(\mathbf{x}^{\text{out}})\|] \leq \epsilon$.

We are now ready to complete the proof of Theorem 5. By Lemma 7, `GRADESTIMATOR` (Algorithm 2) returns an unbiased estimator of ∇f_ρ whose second moment is bounded by $\sigma^2 = O(dL_0^2)$, using two evaluations of $F(\cdot; \xi)$. Thus applying Theorem 8 to $h = f_\rho$ ensures that Algorithm 1 returns a (ν, ϵ) -stationary point of f_ρ , which by Lemma 4 is a $(\nu + \rho, \epsilon)$ -stationary point of f . Recall that $\|f - f_\rho\|_\infty \leq \rho L_0$ and $f(\mathbf{x}_0) - \inf_{\mathbf{x}} f(\mathbf{x}) \leq \Delta$, thus $f_\rho(\mathbf{x}_0) - \inf_{\mathbf{x}} f_\rho(\mathbf{x}) \leq \Delta + \rho L_0 =: \Delta_h$. Overall we have obtained a $(\nu + \rho)$ -stationary point of f using $2T$ function evaluations, where

$$T = O\left(\frac{\sigma^2 \Delta_h}{\nu \epsilon^3}\right) = O\left(\frac{dL_0^2(\Delta + \rho L_0)}{\nu \epsilon^3}\right) .$$

Setting $\rho = \min\left\{\frac{\delta}{2}, \frac{\Delta}{L_0}\right\}$ and $\nu = \delta - \rho = \max\left\{\frac{\delta}{2}, \delta - \frac{\Delta}{L_0}\right\}$ completes the proof, since $\nu + \rho = \delta$ and

$$T = O\left(\frac{dL_0^2(\Delta + \min\left\{\frac{\delta}{2}, \frac{\Delta}{L_0}\right\} L_0)}{\max\left\{\frac{\delta}{2}, \delta - \frac{\Delta}{L_0}\right\} \epsilon^3}\right) = O\left(\frac{dL_0^2(\Delta + \frac{\Delta}{L_0} \cdot L_0)}{\frac{\delta}{2} \cdot \epsilon^3}\right) = O\left(\frac{dL_0^2 \Delta}{\delta \epsilon^3}\right) .$$

5. Proof of Theorem 6

Recall that we saw in the proof of Theorem 5 that $\partial_\nu f_\rho(\cdot) \subseteq \partial_{\nu+\rho} f(\cdot)$ according to Lemma 4, and that for all $m \in [M]$: $\mathbf{z}_m^{r*} \in \mathbb{B}(\mathbf{x}_{\text{out}}^{r*}, \nu)$ by according to Theorem 8. Thus

$$\left\| \bar{\partial}_\delta f(\mathbf{x}_{\text{out}}^{r*}) \right\| = \left\| \bar{\partial}_{\nu+\rho} f(\mathbf{x}_{\text{out}}^{r*}) \right\| \leq \left\| \bar{\partial}_\nu f_\rho(\mathbf{x}_{\text{out}}^{r*}) \right\| \leq \left\| \frac{1}{M} \sum_{m \in [M]} \nabla f_\rho(\mathbf{z}_m^{r*}) \right\| .$$

By denoting $\mathbf{g}^r := \frac{1}{M} \sum_{m \in [M]} \nabla f_\rho(\mathbf{z}_m^r)$, $r \in [R]$ we get that it suffices to show that

$$\Pr \left[\left\| \mathbf{g}^{r*} \right\|^2 \leq \epsilon^2 \right] = \Pr \left[\left\| \mathbf{g}^{r*} \right\| \leq \epsilon \right] \geq 1 - \gamma . \quad (4)$$

By definition of r^* we have

$$\left\| \hat{\mathbf{g}}^{r*} \right\|^2 = \min_{r \in [R]} \left\| \hat{\mathbf{g}}^r \right\|^2 \leq \min_{r \in [R]} \left(2 \left\| \mathbf{g}^r \right\|^2 + 2 \left\| \hat{\mathbf{g}}^r - \mathbf{g}^r \right\|^2 \right) \leq 2 \left(\min_{r \in [R]} \left\| \mathbf{g}^r \right\|^2 + \max_{r \in [R]} \left\| \hat{\mathbf{g}}^r - \mathbf{g}^r \right\|^2 \right) ,$$

thus

$$\begin{aligned} \left\| \mathbf{g}^{r*} \right\|^2 &\leq 2 \left\| \hat{\mathbf{g}}^{r*} \right\|^2 + 2 \left\| \hat{\mathbf{g}}^{r*} - \mathbf{g}^{r*} \right\|^2 \\ &\leq 4 \left(\min_{r \in [R]} \left\| \mathbf{g}^r \right\|^2 + \max_{r \in [R]} \left\| \hat{\mathbf{g}}^r - \mathbf{g}^r \right\|^2 \right) + 2 \left\| \hat{\mathbf{g}}^{r*} - \mathbf{g}^{r*} \right\|^2 \\ &\leq 4 \cdot \min_{r \in [R]} \left\| \mathbf{g}^r \right\|^2 + 6 \cdot \max_{r \in [R]} \left\| \hat{\mathbf{g}}^r - \mathbf{g}^r \right\|^2 . \end{aligned} \quad (5)$$

We now turn to bound each of the summand above with high probability. First, by Theorem 5 we can set $T = O\left(\frac{d\Delta L_0^2}{\delta\epsilon^3}\right)$ so that $\mathbb{E}[\|\mathbf{g}^r\|] \leq \frac{\epsilon}{8}$ for any $r \in [R]$, hence by Markov's inequality

$$\Pr\left[4 \cdot \min_{r \in R} \|\mathbf{g}^r\|^2 > \frac{\epsilon^2}{4}\right] = \Pr\left[\min_{r \in R} \|\mathbf{g}^r\| > \frac{\epsilon}{4}\right] \leq \prod_{r \in [R]} \Pr\left[\|\mathbf{g}^r\| > \frac{\epsilon}{4}\right] \leq 2^{-R}.$$

By setting $R \geq \lceil \log_2(2/\gamma) \rceil$ so that $2^{-R} \leq \frac{\gamma}{2}$, we conclude that

$$\Pr\left[4 \cdot \min_{r \in R} \|\mathbf{g}^r\|^2 > \frac{\epsilon^2}{4}\right] \leq \frac{\gamma}{2}. \quad (6)$$

For the second summand in Eq. (5), note that for all $r \in [R]$:

$$\begin{aligned} \mathbb{E}[\hat{\mathbf{g}}^r] &= \mathbb{E}\left[\frac{1}{S} \sum_{s=1}^S \mathbf{g}_s^r\right] = \frac{1}{S} \sum_{s=1}^S \left(\frac{1}{M} \sum_{m=1}^M \mathbb{E}[\mathbf{g}_{m,s}^r]\right) \\ &= \frac{1}{M} \sum_{m=1}^M \left(\frac{1}{S} \sum_{s=1}^S \mathbb{E}[\mathbf{g}_{m,s}^r]\right) \stackrel{\text{Lemma 7}}{=} \frac{1}{M} \sum_{m=1}^M \nabla f_\rho(\mathbf{z}_m^r) = \mathbf{g}^r, \end{aligned}$$

thus $\mathbb{E}[\hat{\mathbf{g}}^r - \mathbf{g}^r] = 0$, and that it follows from the second claim in Lemma 7 that for any $r \in [R], s \in [S]$:

$$\mathbb{E}[\|\hat{\mathbf{g}}_s^r - \mathbf{g}^r\|^2] \leq \frac{16\sqrt{2\pi}dL_0^2}{M}.$$

Noting that $(\mathbf{g}_1^r - \mathbf{g}^r), \dots, (\mathbf{g}_S^r - \mathbf{g}^r)$ are independent as they are functions of the independent samples $\xi_{m,1}, \dots, \xi_{m,S}$, $m \in [M]$, we apply a simple tail bound for the norm of a sum of independent vectors (Lemma 9 in the appendix) to get for any $\lambda > 0$:

$$\begin{aligned} \Pr\left[\|\hat{\mathbf{g}}^r - \mathbf{g}^r\|^2 \geq \lambda \frac{16\sqrt{2\pi}dL_0^2}{MS}\right] &= \Pr\left[\left\|\frac{1}{S} \sum_{s=1}^S (\mathbf{g}_s^r - \mathbf{g}^r)\right\|^2 \geq \lambda \frac{16\sqrt{2\pi}dL_0^2}{MS}\right] \\ &= \Pr\left[\left\|\sum_{s=1}^S (\mathbf{g}_s^r - \mathbf{g}^r)\right\|^2 \geq \lambda S \cdot \frac{16\sqrt{2\pi}dL_0^2}{M}\right] \leq \frac{1}{\lambda}, \end{aligned}$$

hence by the union bound

$$\Pr\left[\max_{r \in [R]} \|\hat{\mathbf{g}}^r - \mathbf{g}^r\|^2 \geq \lambda \frac{16\sqrt{2\pi}dL_0^2}{MS}\right] \leq \frac{R}{\lambda}.$$

Setting $\lambda := \lceil \frac{2R}{\gamma} \rceil$ so that $\frac{R}{\lambda} \leq \frac{\gamma}{2}$, we see that $S \gtrsim \frac{dL_0^2 \log(1/\gamma)}{M\epsilon^2\gamma}$ suffices for having $\lambda \frac{16\sqrt{2\pi}dL_0^2}{MS} \leq \frac{\epsilon^2}{4}$, under which the inequality above shows that

$$\Pr\left[\max_{r \in [R]} \|\hat{\mathbf{g}}^r - \mathbf{g}^r\|^2 \geq \frac{\epsilon^2}{4}\right] \leq \frac{\gamma}{2}. \quad (7)$$

By combining Eq. (6) and Eq. (7) and applying the union bound to get Eq. (5), we have proved Eq. (4) as required. Finally, recalling that `GRADESTIMATOR` (Algorithm 2) requires 2 noisy function evaluations, it is clear that the total number of evaluations performed by Algorithm 3 is bounded by

$$2R \cdot (T + MS) = O\left(\frac{dL_0^2 \Delta \log(1/\gamma)}{\delta \epsilon^3} + \frac{dL_0^2 \log^2(1/\gamma)}{\gamma \epsilon^2}\right).$$

Acknowledgments

The authors would like to thank the anonymous reviewers for their helpful suggestions, and in particular for pointing out a modified assignment of hyper-parameters which led to an improved result with respect to the Lipschitz constant. This research is supported in part by European Research Council (ERC) grant 754705. GK is supported by an Azrieli Foundation graduate fellowship.

Appendix A.

Lemma 9 *Let $X_1, \dots, X_N \in \mathbb{R}^d$ be independent random vectors such that for all $i \in [N]$: $\mathbb{E}[X_i] = \mathbf{0}$, $\mathbb{E}[\|X_i\|^2] \leq \sigma_i^2$. Then $\mathbb{E}\left[\left\|\sum_{i=1}^N X_i\right\|^2\right] \leq \sum_{i=1}^N \sigma_i^2$. In particular, for any $\lambda > 0$:*

$$\Pr\left[\left\|\sum_{i=1}^N X_i\right\|^2 \geq \lambda \cdot \sum_{i=1}^N \sigma_i^2\right] \leq \lambda^{-1}.$$

Proof By linearity of expectation we have

$$\mathbb{E}\left[\left\|\sum_{i \in [N]} X_i\right\|^2\right] = \sum_{i \in [N]} \mathbb{E}[\|X_i\|^2] + \sum_{i \neq j \in [N]} \mathbb{E}[\langle X_i, X_j \rangle] = \sum_{i \in [N]} \mathbb{E}[\|X_i\|^2] \leq \sum_{i \in [N]} \sigma_i^2,$$

where we used the assumption that for any $i \neq j$: X_i, X_j are independent, thus $\mathbb{E}[\langle X_i, X_j \rangle] = 0$. The second claim follows from Markov's inequality. $\blacksquare$

References

- Alekh Agarwal and Ofer Dekel. Optimal algorithms for online convex optimization with multi-point bandit feedback. In *Conference on Learning Theory*, pages 28–40. Citeseer, 2010.
- Yossi Arjevani, Yair Carmon, John C Duchi, Dylan J Foster, Nathan Srebro, and Blake Woodworth. Lower bounds for non-convex stochastic optimization. *Mathematical Programming*, 199(1-2):165–214, 2023.

- Sébastien Bubeck, Qijia Jiang, Yin-Tat Lee, Yuanzhi Li, and Aaron Sidford. Complexity of highly parallel non-smooth convex optimization. *Advances in neural information processing systems*, 32, 2019.
- Yair Carmon, John C Duchi, Oliver Hinder, and Aaron Sidford. Lower bounds for finding stationary points i. *Mathematical Programming*, 184(1-2):71–120, 2020.
- Lesi Chen, Jing Xu, and Luo Luo. Faster gradient-free algorithms for nonsmooth nonconvex stochastic optimization. In *International Conference on Machine Learning*, pages 5219–5233. PMLR, 2023.
- F. H. Clarke. *Optimization and Nonsmooth Analysis*. SIAM, 1990.
- Ashok Cutkosky, Harsh Mehta, and Francesco Orabona. Optimal stochastic non-smooth non-convex optimization through online-to-non-convex conversion. In *International Conference on Machine Learning*, pages 6643–6670. PMLR, 2023.
- Damek Davis, Dmitriy Drusvyatskiy, Yin Tat Lee, Swati Padmanabhan, and Guanhao Ye. A gradient sampling method with complexity guarantees for lipschitz functions in high and low dimensions. *Advances in Neural Information Processing Systems*, 35:6692–6703, 2022.
- John Duchi, Feng Ruan, and Chulhee Yun. Minimax bounds on stochastic batched convex optimization. In *Conference On Learning Theory*, pages 3065–3162. PMLR, 2018.
- John C Duchi, Michael I Jordan, Martin J Wainwright, and Andre Wibisono. Optimal rates for zero-order convex optimization: The power of two function evaluations. *IEEE Transactions on Information Theory*, 61(5):2788–2806, 2015.
- Abraham D Flaxman, Adam Tauman Kalai, and H Brendan McMahan. Online convex optimization in the bandit setting: gradient descent without a gradient. In *Proceedings of the sixteenth annual ACM-SIAM symposium on Discrete algorithms*, pages 385–394, 2005.
- Saeed Ghadimi and Guanhui Lan. Stochastic first-and zeroth-order methods for nonconvex stochastic programming. *SIAM Journal on Optimization*, 23(4):2341–2368, 2013.
- A. Goldstein. Optimization of Lipschitz continuous functions. *Mathematical Programming*, 13(1):14–22, 1977.
- Michael Jordan, Guy Kornowski, Tianyi Lin, Ohad Shamir, and Manolis Zampetakis. Deterministic nonsmooth nonconvex optimization. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 4570–4597. PMLR, 2023.
- Guy Kornowski and Ohad Shamir. Oracle complexity in nonsmooth nonconvex optimization. *Journal of Machine Learning Research*, 23(314):1–44, 2022.
- Tianyi Lin, Zeyu Zheng, and Michael Jordan. Gradient-free methods for deterministic and stochastic nonsmooth nonconvex optimization. *Advances in Neural Information Processing Systems*, 35:26160–26175, 2022.

- Sadhika Malladi, Tianyu Gao, Eshaan Nichani, Alex Damian, Jason D Lee, Danqi Chen, and Sanjeev Arora. Fine-tuning language models with just forward passes. *Advances in Neural Information Processing Systems*, 2023.
- Yurii Nesterov and Vladimir Spokoiny. Random gradient-free minimization of convex functions. *Foundations of Computational Mathematics*, 17:527–566, 2017.
- Ohad Shamir. An optimal algorithm for bandit and zero-order convex optimization with two-point feedback. *The Journal of Machine Learning Research*, 18(1):1703–1713, 2017.
- James C Spall. *Introduction to stochastic search and optimization: estimation, simulation, and control*. John Wiley & Sons, 2005.
- L. Tian, K. Zhou, and A. M-C. So. On the finite-time complexity and practical computation of approximate stationarity concepts of Lipschitz functions. In *ICML*, pages 21360–21379. PMLR, 2022.
- Farzad Yousefian, Angelia Nedić, and Uday V Shanbhag. On stochastic gradient and sub-gradient methods with adaptive steplength sequences. *Automatica*, 48(1):56–67, 2012.
- J. Zhang, H. Lin, S. Jegelka, S. Sra, and A. Jadbabaie. Complexity of finding stationary points of nonconvex nonsmooth functions. In *ICML*, pages 11173–11182. PMLR, 2020.