

Online Convex Optimization with Heavy Tails: Old Algorithms, New Regrets, and Applications

Zijian Liu

ZL3067@STERN.NYU.EDU

Stern School of Business, New York University

Editors: Matus Telgarsky and Jonathan Ullman

Abstract

In Online Convex Optimization (OCO), when the stochastic gradient has a finite variance, many algorithms provably work and guarantee a sublinear regret. However, limited results are known if the gradient estimate has a heavy tail, i.e., the stochastic gradient only admits a finite p -th central moment for some $p \in (1, 2]$. Motivated by it, this work examines different old algorithms for OCO (e.g., Online Gradient Descent) in the more challenging heavy-tailed setting. Under the standard bounded domain assumption, we establish new regrets for these classical methods without any algorithmic modification. Remarkably, these regret bounds are fully optimal in all parameters (can be achieved even without knowing p), suggesting that OCO with heavy tails can be solved effectively without any extra operation (e.g., gradient clipping). Our new results have several applications. A particularly interesting one is the first provable and optimal convergence result for nonsmooth nonconvex optimization under heavy-tailed noise without gradient clipping.

Keywords: Online Learning, Online Convex Optimization, Heavy Tails

1. Introduction

This paper studies the online learning problem with convex losses, also known as Online Convex Optimization (OCO), a widely applicable framework that learns under streaming data (Cesa-Bianchi and Lugosi, 2006; Shalev-Shwartz, 2012; Hazan, 2016; Orabona, 2019). OCO has tons of implications for both designing and analyzing algorithms in different areas, for example, stochastic optimization (McMahan and Streeter, 2010; Duchi et al., 2011; Kingma and Ba, 2014), PAC learning (Cesa-Bianchi et al., 2004), control theory (Agarwal et al., 2019; Hazan and Singh, 2025), etc.

In an OCO problem, a learning algorithm A would interact with the environment in T rounds, where $T \in \mathbb{N}$ can be either known or unknown. Formally, in each round t , the learner A first decides an output $\mathbf{x}_t \in \mathcal{X}$ from a convex feasible set $\mathcal{X} \subseteq \mathbb{R}^d$, then the environment reveals a convex loss function $\ell_t : \mathcal{X} \rightarrow \mathbb{R}$, and A incurs a loss of $\ell_t(\mathbf{x}_t)$. After T many rounds, the quantity measuring the algorithm’s performance is called regret, defined relative to any fixed competitor $\mathbf{x} \in \mathcal{X}$ as follows:

$$R_T^A(\mathbf{x}) \triangleq \sum_{t=1}^T \ell_t(\mathbf{x}_t) - \ell_t(\mathbf{x}).$$

In the classical setting, instead of observing full information about ℓ_t , the learner A is only guaranteed to receive a subgradient $\nabla \ell_t(\mathbf{x}_t) \in \partial \ell_t(\mathbf{x}_t)$ at its decision, where $\partial \ell_t(\mathbf{x}_t)$ denotes the subdifferential set of ℓ_t at $\mathbf{x}_t$ (Rockafellar, 1997). This turns out to be enough for our purpose of minimizing the regret, since any OCO problem can be reduced to an Online Linear Optimization

0. This is a short, self-contained version of the full paper (Liu, 2025). Compared to the full paper, this version focuses only on nonsmooth problems. For broader settings (e.g., smooth OCO) and applications, please refer to Liu (2025).

(OLO) instance via the inequality $\ell_t(\mathbf{x}_t) - \ell_t(\mathbf{x}) \leq \langle \nabla \ell_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x} \rangle$, which holds due to convexity. Under the standard bounded domain assumption, i.e., $\mathcal{X}$ has a finite diameter D , many classical algorithms, e.g., Online Gradient Descent (OGD) (Zinkevich, 2003), guarantee an optimal sublinear regret $GD\sqrt{T}$ for G -Lipschitz ℓ_t . Even better, in the case that computing an exact subgradient is intractable, and one could only query a stochastic estimate $\mathbf{g}_t$ satisfying $\mathbb{E}[\mathbf{g}_t | \mathbf{x}_t] \in \partial \ell_t(\mathbf{x}_t)$, the OGD algorithm can still solve OCO effectively with a provable $(G + \sigma)D\sqrt{T}$ regret bound in expectation if the stochastic noise $\mathbf{g}_t - \nabla \ell_t(\mathbf{x}_t)$ has a bounded second moment σ^2 for some $\sigma \geq 0$, which is called the finite variance condition.

However, many works have pointed out that even for the easier stochastic optimization (i.e., $\ell_t = F$ for a common F), the typical finite variance assumption is too optimistic and can be violated in different tasks (Simsekli et al., 2019; Zhang et al., 2020a; Hodgkinson and Mahoney, 2021), and their observations suggest that the stochastic gradient only admits a finite p -th central moment upper bounded by σ^p for some $p \in (1, 2]$, which is named heavy-tailed noise. This new assumption generalizes the classical finite variance condition ($p = 2$) and becomes challenging when $p < 2$. A particular evidence is that the famous Stochastic Gradient Descent (SGD) algorithm (Robbins and Monro, 1951) (which is exactly OGD for stochastic optimization) provably diverges (Zhang et al., 2020a).

Though heavy-tailed stochastic optimization has been extensively studied (Sadiev et al., 2023; Liu and Zhou, 2023; Nguyen et al., 2023), limited results are known for OCO with heavy tails. The only work under this topic that we are aware of is Zhang and Cutkosky (2022), which established a parameter-free regret bound in high probability (more discussions provided later). However, their algorithm includes many nontrivial modifications like gradient clipping and significantly deviates from the existing simple OCO algorithms used in practice. Especially, consider OGD as an example. Though the heavy-tailed issue is known, OGD (or just think of it as SGD) still works (sometimes very well) in practice even without gradient clipping and is arguably one of the most popular optimizers, which seemingly contradicts the theory of nonconvergence mentioned before. This indicates that, for classical OCO algorithms under heavy-tailed noise, a huge gap exists between the empirical convergence (or even the effective practical performance) and theoretical guarantees. Therefore, we are naturally led to the following question:

In what context can old OCO algorithms work under heavy tails, in what sense, and to what extent?

1.1. Contributions

Motivated by the above question, we examine three classical algorithms for OCO: Online Gradient Descent (OGD) (Zinkevich, 2003), Dual Averaging (DA) (Nesterov, 2009; Xiao, 2009), and AdaGrad (McMahan and Streeter, 2010; Duchi et al., 2011), and answer it as follows:

Under the standard bounded domain assumption, the in-expectation regret $\mathbb{E}[\mathbf{R}_T^A(\mathbf{x})]$ is finite and optimal for any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$, without any algorithmic modification.

In detail, our new results for heavy-tailed OCO are summarized here:

- We prove the only and the first optimal regret bound $\mathbb{E}[\mathbf{R}_T^A(\mathbf{x})] \lesssim GD\sqrt{T} + \sigma DT^{1/p}, \forall \mathbf{x} \in \mathcal{X}$ for any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$. Remarkably, AdaGrad can achieve this result without knowing any of the Lipschitz parameter G , noise level σ , and tail index p .

- We extend the analysis of OGD to Online Strongly Convex Optimization with heavy tails and establish the first provable result $\mathbb{E} [\mathbf{R}_T^{\text{OGD}}(\mathbf{x})] \lesssim \frac{G^2 \log T}{\mu} + \frac{\sigma^p D^{2-p}}{\mu^{p-1}} T^{2-p}, \forall \mathbf{x} \in \mathcal{X}$, where $\mu > 0$ is the modulus of strong convexity and T^0 should be read as $\log T$.

Based on the new regret bounds for OCO with heavy tails, we provide the following applications:

- For nonsmooth convex optimization with heavy tails, we show the first optimal in-expectation rate $GD/\sqrt{T} + \sigma D/T^{1-1/p}$ achieved without gradient clipping, which applies to both the average iterate and last iterate, demonstrating that SGD does converge once the domain is bounded. Moreover, we also give the rate when strong convexity additionally holds.
- For nonsmooth nonconvex optimization with heavy tails, we show the first provable sample complexity of $G^2 \delta^{-1} \epsilon^{-3} + \sigma^{\frac{p}{p-1}} \delta^{-1} \epsilon^{-\frac{2p-1}{p-1}}$ for finding a (δ, ϵ) -stationary point without gradient clipping. In addition, we also establish the first lower bound for nonsmooth nonconvex optimization under heavy tails, matching our sample complexity in the high accuracy and noisy regime (i.e., ϵ is small enough with $\sigma > 0$). These two results together provide a nearly complete characterization of the complexity of finding (δ, ϵ) -stationary points in the heavy-tailed setting. Moreover, we give the first convergence result when the problem-dependent parameters (like G , σ , and p) are unknown in advance, resolving a question asked by [Liu et al. \(2024\)](#).

1.2. Discussion on [Zhang and Cutkosky \(2022\)](#)

As noted, [Zhang and Cutkosky \(2022\)](#) is the only work for OCO with heavy tails, as far as we know. There are two major discrepancies between them and us. First, they consider the case where the feasible set $\mathcal{X}$ is unbounded and aim to establish a parameter-free regret bound, i.e., the regret bound has a linear dependency on $\|\mathbf{x}\|$ (up to an extra polylog $\|\mathbf{x}\|$) for any competitor $\mathbf{x} \in \mathcal{X}$. Second, they focus on high-probability rather than in-expectation analysis. As such, their regret is in the form of $\mathbf{R}_T^A(\mathbf{x}) \lesssim (G + \sigma) \|\mathbf{x}\| T^{1/p}, \forall \mathbf{x} \in \mathcal{X}$ (up to extra polylogarithmic factors) with high probability. Without a doubt, their setting is harder than ours implying their bound is stronger as it can convert to an in-expectation regret $\mathbb{E} [\mathbf{R}_T^A(\mathbf{x})] \lesssim (G + \sigma) DT^{1/p}$ for any bounded domain $\mathcal{X}$ with a diameter D .

We emphasize that the motivation behind [Zhang and Cutkosky \(2022\)](#) differs heavily from ours. They aim to solve heavy-tailed OCO with a new proposed method that contains many nontrivial technical tricks, including gradient clipping, artificially added regularization, and solving the additional fixed-point equation. However, their result cannot reflect why the existing simple OCO algorithms like OGD work in practice under heavy-tailed noise. In contrast, our goal is to examine whether, when, and how the classical OCO algorithms work under heavy tails, thereby filling the missing piece in the literature.

Moreover, we would like to mention two limitations of [Zhang and Cutkosky \(2022\)](#). First, though the $T^{1/p}$ regret seems tight as it matches the lower bound ([Nemirovski and Yudin, 1983](#); [Raginsky and Rakhlin, 2009](#); [Vural et al., 2022](#)), this may not be the best, since an optimal bound should recover the standard $\sqrt{T}$ regret in the deterministic case (i.e., $\sigma = 0$), as one can imagine. This suggests that their bound is not entirely optimal. Second, we remark that they require knowing both problem-dependent parameters G , σ , p and time horizon T in the algorithm, which may be hard to satisfy in the online setting. In comparison, our regret bound $GD\sqrt{T} + \sigma DT^{1/p}$ is fully

optimal in all parameters¹. Importantly, AdaGrad can achieve it while oblivious to the problem information.

1.3. Discussion on High-Probability Bounds

One may notice that all of our new results are measured in expectation and naturally wonder whether high-probability bounds (i.e., polylogarithmic dependency on the failure probability) are achievable in the same setting of this work. Though we cannot fully address this question at this time, some preliminary discussions are provided in Appendix A, including a negative result for OGD/SGD.

2. Preliminary

Notation. $\mathbb{N}$ denotes the set of natural numbers (excluding 0). $[T] \triangleq \{1, \dots, T\}$, $\forall T \in \mathbb{N}$. $a \wedge b \triangleq \min\{a, b\}$ and $a \vee b \triangleq \max\{a, b\}$. We write $a \lesssim b$ (resp., $a \gtrsim b$) if $a \leq Cb$ (resp., $a \geq Cb$) for a universal constant $C > 0$. $\lfloor \cdot \rfloor$ and $\lceil \cdot \rceil$ respectively represent the floor and ceiling functions. $\langle \cdot, \cdot \rangle$ denotes the Euclidean inner product and $\|\cdot\| \triangleq \sqrt{\langle \cdot, \cdot \rangle}$ is the standard 2-norm. Given $\mathbf{x} \in \mathbb{R}^d$ and $D > 0$, $\mathcal{B}^d(\mathbf{x}, D)$ is the Euclidean ball in $\mathbb{R}^d$ centered at $\mathbf{x}$ with a radius D . In the case $\mathbf{x} = \mathbf{0}$, we use the shorthand $\mathcal{B}^d(D)$. Given $A \subseteq \mathbb{R}^d$, $\text{int}A$ stands for the interior points of A . For nonempty closed convex $A \subseteq \mathbb{R}^d$, Π_A is the Euclidean projection operator onto A . For a convex function f , $\partial f(\mathbf{x})$ denotes its subgradient set at $\mathbf{x}$.

Remark 1 We choose the Euclidean norm only for simplicity. Extending the results of this work to any general norm by Online Mirror Descent (OMD) (Nemirovski and Yudin, 1983; Beck and Teboulle, 2003) via the Bregman divergence is straightforward.

This work studies OCO in the context of Assumption 1.

Assumption 1 We consider the following series of assumptions:

- $\mathcal{X} \subset \mathbb{R}^d$ is a nonempty closed convex set bounded by D , i.e., $\sup_{\mathbf{x}, \mathbf{y} \in \mathcal{X}} \|\mathbf{x} - \mathbf{y}\| \leq D$.
- $\ell_t : \mathcal{X} \rightarrow \mathbb{R}$ is closed convex for all $t \in [T]$.
- ℓ_t is G -Lipschitz on $\mathcal{X}$, i.e., $\|\nabla \ell_t(\mathbf{x})\| \leq G, \forall \mathbf{x} \in \mathcal{X}, \nabla \ell_t(\mathbf{x}) \in \partial \ell_t(\mathbf{x})$, for all $t \in [T]$.
- Given a point $\mathbf{x}_t \in \mathcal{X}$ at the t -th iteration, one can query $\mathbf{g}_t \in \mathbb{R}^d$ satisfying $\nabla \ell_t(\mathbf{x}_t) \triangleq \mathbb{E}[\mathbf{g}_t \mid \mathcal{F}_{t-1}] \in \partial \ell_t(\mathbf{x}_t)$ and $\mathbb{E}[\|\epsilon_t\|^p] \leq \sigma^p$ for some $p \in (1, 2]$ and $\sigma \geq 0$, where $\mathcal{F}_t \triangleq \sigma(\mathbf{g}_1, \dots, \mathbf{g}_t)$ denotes the natural filtration and $\epsilon_t \triangleq \mathbf{g}_t - \nabla \ell_t(\mathbf{x}_t)$ is the stochastic noise.

Remark 2 D is recognized as known, like ubiquitously assumed in the OCO literature. Moreover, $\mathbf{x}_t$ denotes the decision/output of the online learning algorithm by default.

Remark 3 The idea presented in this work can be extended to the composite setting (i.e., $\ell_t + \psi_t$, where ψ_t is convex and assumed to be known at time $t - 1$) via the proximal update, provably.

In Assumption 1, the first three points are standard, and the fourth is the heavy-tailed noise assumption. In particular, $p = 2$ recovers the standard finite variance condition.

1. The optimality is justified by a matching lower bound $GD\sqrt{T} + \sigma DT^{1/p}$ that can be obtained by combining Theorem 5.1 of Orabona (2019) and Section 5.3.1 of Nemirovski and Yudin (1983).

3. Old Algorithms under Heavy Tails

In this section, we revisit three classical algorithms for OCO: OGD, DA, and AdaGrad, whose regret bounds are well-studied in the finite variance case but remain unknown under heavy-tailed noise.

The basic idea of proving these algorithms work under heavy tails is to leverage the boundedness property of $\mathcal{X}$. We will describe it in more detail using OGD as an illustrated example. The analysis of DA follows a similar way at a high level, but differs in some details. However, though AdaGrad can be viewed as OGD with an adaptive stepsize, the way to utilize the boundedness property is entirely different. All formal proofs are deferred to the appendix due to space limitations.

3.1. New Regret for Online Gradient Descent

Algorithm 1 Online Gradient Descent (OGD) (Zinkevich, 2003)

Input: initial point $\mathbf{x}_1 \in \mathcal{X}$, stepsize $\eta_t > 0$

for $t = 1$ **to** T **do**

$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}}(\mathbf{x}_t - \eta_t \mathbf{g}_t)$

end for

We begin from arguably the most basic algorithm for OCO, Online Gradient Descent (OGD).

A well known analysis. The regret bound of OGD has been extensively studied (Shalev-Shwartz, 2012; Hazan, 2016; Orabona, 2019). The most well known analysis is perhaps the following one: for any $\mathbf{x} \in \mathcal{X}$, there is

$$\|\mathbf{x}_{t+1} - \mathbf{x}\|^2 = \|\Pi_{\mathcal{X}}(\mathbf{x}_t - \eta_t \mathbf{g}_t) - \Pi_{\mathcal{X}}(\mathbf{x})\|^2 \leq \|\mathbf{x}_t - \eta_t \mathbf{g}_t - \mathbf{x}\|^2,$$

where the inequality holds by the nonexpansive property of $\Pi_{\mathcal{X}}$. Expanding both sides and rearranging terms yield that

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x}_t - \mathbf{x}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2\eta_t} + \frac{\eta_t \|\mathbf{g}_t\|^2}{2}. \quad (1)$$

If $\mathbf{g}_t$ admits a finite variance, i.e., $p = 2$ in Assumption 1, taking expectations on both sides, then following a standard analysis for $\eta_t = \frac{D}{(G+\sigma)\sqrt{t}}$ (or $\eta_t = \frac{D}{(G+\sigma)\sqrt{T}}$ if T is known) gives the regret

$$\mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] \lesssim (G + \sigma) D \sqrt{T}, \forall \mathbf{x} \in \mathcal{X}.$$

However, the step of taking expectations on the R.H.S. of (1) crucially relies on the finite variance condition of $\mathbf{g}_t$. Therefore, one may naturally think OGD would not guarantee a finite regret if $p < 2$.

A less well known analysis². As discussed, the failure of the above proof under heavy-tailed noise is due to (1). Therefore, if a tighter inequality than (1) exists, then it might be possible to show that OGD still works for $p < 2$. However, does it exist?

Actually, there is another less well known analysis to produce a better inequality than (1). That is, first showing for any $\mathbf{x} \in \mathcal{X}$, by the optimality condition of the update rule,

$$\langle \mathbf{g}_t, \mathbf{x}_{t+1} - \mathbf{x} \rangle \leq \frac{\langle \mathbf{x}_t - \mathbf{x}_{t+1}, \mathbf{x}_{t+1} - \mathbf{x} \rangle}{\eta_t} = \frac{\|\mathbf{x}_t - \mathbf{x}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{x}\|^2 - \|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_t},$$

2. To clarify, the phrase “less well known” is compared to the first one. This analysis itself is also famous in the literature. For example, see Lemma 6.10 of Orabona (2019) and Lemma 3.1 of Lan (2020).

and then obtaining

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x}_t - \mathbf{x}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2\eta_t} + \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_t}. \quad (2)$$

Note that (2) is tighter than (1) as $\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle \leq \|\mathbf{g}_t\| \|\mathbf{x}_t - \mathbf{x}_{t+1}\| \leq \frac{\eta_t \|\mathbf{g}_t\|^2}{2} + \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_t}$, where the first step is due to Cauchy-Schwarz inequality and the second one is by AM-GM inequality.

Handle $p < 2$ in a simple way. Though we have tightened (1) into (2), can inequality (2) help to overcome heavy tails? The answer is surprisingly positive, and our solution is fairly simple. Instead of directly applying AM-GM inequality in the second step, we recall $\mathbf{g}_t = \nabla \ell_t(\mathbf{x}_t) + \epsilon_t$ and use triangle inequality to obtain

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle \leq \|\mathbf{g}_t\| \|\mathbf{x}_t - \mathbf{x}_{t+1}\| \leq (\|\nabla \ell_t(\mathbf{x}_t)\| + \|\epsilon_t\|) \|\mathbf{x}_t - \mathbf{x}_{t+1}\|. \quad (3)$$

On the one hand, by $\|\nabla \ell_t(\mathbf{x}_t)\| \leq G$ and AM-GM inequality, there is

$$\|\nabla \ell_t(\mathbf{x}_t)\| \|\mathbf{x}_t - \mathbf{x}_{t+1}\| \leq G \|\mathbf{x}_t - \mathbf{x}_{t+1}\| \leq \eta_t G^2 + \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{4\eta_t}. \quad (4)$$

On the other hand, let $p_\star \triangleq \frac{p}{p-1}$ and $C(p) \triangleq \frac{(4p-4)^{p-1}}{p^p}$, we have

$$\begin{aligned} \|\epsilon_t\| \|\mathbf{x}_t - \mathbf{x}_{t+1}\| &= \left(\frac{4\eta_t}{p_\star} \right)^{\frac{1}{p_\star}} \|\epsilon_t\| \|\mathbf{x}_t - \mathbf{x}_{t+1}\|^{1-\frac{2}{p_\star}} \cdot \left(\frac{p_\star \|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{4\eta_t} \right)^{\frac{1}{p_\star}} \\ &\stackrel{(a)}{\leq} \frac{\left(\frac{4\eta_t}{p_\star} \right)^{\frac{p}{p_\star}} \|\epsilon_t\|^p \|\mathbf{x}_t - \mathbf{x}_{t+1}\|^{p-\frac{2p}{p_\star}}}{p} + \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{4\eta_t} \\ &\stackrel{(b)}{\leq} C(p) \eta_t^{p-1} \|\epsilon_t\|^p D^{2-p} + \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{4\eta_t}, \end{aligned} \quad (5)$$

where (a) is by Young's inequality and (b) is due to $\|\mathbf{x}_t - \mathbf{x}_{t+1}\| \leq D$, $p_\star = \frac{p}{p-1}$, and $C(p) = \frac{(4p-4)^{p-1}}{p^p}$. Next, we plug (4) and (5) back into (3), then combine with (2) to know

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x}_t - \mathbf{x}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2\eta_t} + \eta_t G^2 + C(p) \eta_t^{p-1} \|\epsilon_t\|^p D^{2-p}. \quad (6)$$

Notably, the term $\|\epsilon_t\|^p$ has a correct exponent p . Thus, we can safely take expectations on both sides. Finally, a standard analysis yields the following Theorem 1 (see Appendix B for a formal proof).

Theorem 1 Under Assumption 1, taking $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/p}}$ in OGD (Algorithm 1), we have

$$\mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}, \forall \mathbf{x} \in \mathcal{X}.$$

As far as we know, Theorem 1 is the first and only provable result for OGD under heavy tails. Remarkably, it is not only tight in T (Nemirovski and Yudin, 1983; Raginsky and Rakhlin, 2009; Vural et al., 2022) but also fully optimal in all parameters, in contrast to the bound $(G + \sigma)DT^{1/p}$

of [Zhang and Cutkosky \(2022\)](#). This reveals that OCO with heavy tails can be optimally solved as effectively as the finite variance case once the domain is bounded, a classical condition adopted in many existing works.

Strongly convex functions. We highlight that the above idea can also be applied to Online Strongly Convex Optimization and leads to a sublinear regret T^{2-p} better than $T^{1/p}$. This extension can be found in [Appendix B](#).

3.2. New Regret for Dual Averaging

Algorithm 2 Dual Averaging (DA) ([Nesterov, 2009](#); [Xiao, 2009](#))

Input: initial point $\mathbf{x}_1 \in \mathcal{X}$, stepsize $\eta_t > 0$

for $t = 1$ **to** T **do**

$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}}(\mathbf{x}_1 - \eta_t \sum_{s=1}^t \mathbf{g}_s)$$

end for

Remark 4 *It is known that DA is a special realization of the more general Follow-the-Regularized-Leader (FTRL) framework ([McMahan, 2011](#)). To keep the work concise, we focus only on DA. The key idea for proving [Theorem 2](#) can be directly extended to show new regret for FTRL under heavy-tailed noise.*

We turn our attention to the second candidate, the Dual Averaging (DA) algorithm, which is given in [Algorithm 2](#). Though DA coincides with OGD when $\mathcal{X} = \mathbb{R}^d$ and $\eta_t = \eta$, these two methods in general are not equivalent and can have significant performance differences in practice. Therefore, it is also important to understand DA under heavy tails.

Despite the proof strategies for OGD and DA are in different flavors (even for $p = 2$), the basic idea presented before for OGD still works here, i.e., apply the boundedness property of $\mathcal{X}$ to make the term $\|\epsilon_t\|$ have a correct exponent. Armed with this thought, we can prove the following new regret bound for DA under heavy-tailed noise. We refer the reader to [Appendix C](#) for its proof.

Theorem 2 *Under [Assumption 1](#), taking $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/p}}$ in DA ([Algorithm 2](#)), we have*

$$\mathbb{E} \left[R_T^{\text{DA}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}, \forall \mathbf{x} \in \mathcal{X}.$$

As far as we know, [Theorem 2](#) is the first provable and optimal regret for DA under heavy tails. It guarantees the same tight bound as in [Theorem 1](#) up to different constants.

3.3. New Regret for AdaGrad

Algorithm 3 AdaGrad ([McMahan and Streeter, 2010](#); [Duchi et al., 2011](#))

Input: initial point $\mathbf{x}_1 \in \mathcal{X}$, stepsize $\eta > 0$

for $t = 1$ **to** T **do**

$$\eta_t = \eta V_t^{-1/2} \text{ where } V_t = \sum_{s=1}^t \|\mathbf{g}_s\|^2$$

$$\mathbf{x}_{t+1} = \Pi_{\mathcal{X}}(\mathbf{x}_t - \eta_t \mathbf{g}_t)$$

end for

Remark 5 Algorithm 3 is also named AdaGrad-Norm (e.g., Ward et al. (2019)). We simply call it AdaGrad. It is straightforward to generalize Theorem 3 below to the per-coordinate update version.

Although Theorems 1 and 2 are optimal, they both suffer from an undesired point. That is, the stepsize $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/p}}$ requires knowing all problem-dependent parameters. However, it may not be easy to obtain them in an online setting. Especially, it heavily depends on the prior information about the tail index p , which is hard to know (even approximately) in advance. In other words, they both lack the adaptive property to an unknown environment.

To handle this issue, we consider AdaGrad, a classical adaptive algorithm for OCO. As can be seen, AdaGrad is just OGD with an adaptive stepsize. However, it is this adaptive stepsize that can help us to overcome the above undesired point.

Theorem 3 Under Assumption 1, taking $\eta = D/\sqrt{2}$ in AdaGrad (Algorithm 3), we have

$$\mathbb{E} \left[R_T^{\text{AdaGrad}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}, \forall \mathbf{x} \in \mathcal{X}.$$

Remark 6 We also establish a similar result for DA with an adaptive stepsize. See Theorem 8 in Appendix C for details.

Theorem 3 provides the first regret bound for AdaGrad under heavy tails. Impressively, it is optimal even without knowing any of G , σ , and p . This surprising result once again demonstrates the power of the adaptive method, indicating it is robust to an unknown environment and even heavy-tailed noise, which may partially explain the favorable performance of many adaptive optimizers designed based on AdaGrad like RMSProp (Tieleman et al., 2012) and Adam (Kingma and Ba, 2014).

We point out that the key to establishing Theorem 3 differs from the idea used before for OGD and DA. Actually, Theorem 3 can be obtained in an embarrassingly simple way. It is known that AdaGrad with $\eta = D/\sqrt{2}$ on a bounded domain guarantees the following path-wise regret

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \lesssim D \sqrt{\sum_{t=1}^T \|\mathbf{g}_t\|^2}. \quad (7)$$

Observe that $\sqrt{\sum_{t=1}^T \|\mathbf{g}_t\|^2} \lesssim \sqrt{\sum_{t=1}^T \|\nabla \ell_t(\mathbf{x}_t)\|^2} + \sqrt{\sum_{t=1}^T \|\epsilon_t\|^2} \leq G\sqrt{T} + \left(\sum_{t=1}^T \|\epsilon_t\|^p \right)^{\frac{1}{p}}$, where the last step is due to $\|\cdot\|_2 \leq \|\cdot\|_p$ for any $p \in [1, 2]$. After taking expectations on both sides of (7) and applying Hölder's inequality to obtain $\mathbb{E} \left[\left(\sum_{t=1}^T \|\epsilon_t\|^p \right)^{\frac{1}{p}} \right] \leq \left(\sum_{t=1}^T \mathbb{E} [\|\epsilon_t\|^p] \right)^{\frac{1}{p}} \leq \sigma T^{\frac{1}{p}}$, we conclude Theorem 3. To make the work self-consistent, we produce the formal proof of Theorem 3 in Appendix D.

4. Applications

We provide some applications based on the new regret bounds established in Section 3. The basic problem we study is optimizing a single objective F , which could be either convex or nonconvex.

4.1. Nonsmooth Convex Optimization

In this section, we consider nonsmooth convex optimization with heavy tails.

Convergence of the average iterate. First, we focus on the average-iterate convergence. By the classical online-to-batch conversion (Cesa-Bianchi et al., 2004), the following corollary holds.

Corollary 1 *Under Assumption 1 for $\ell_t(\mathbf{x}) = \langle \nabla F(\mathbf{x}_t), \mathbf{x} \rangle$ and let $\bar{\mathbf{x}}_T \triangleq \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$, for any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$, we have*

$$\mathbb{E}[F(\bar{\mathbf{x}}_T) - F(\mathbf{x})] \leq \frac{\mathbb{E}[\mathbf{R}_T^A(\mathbf{x})]}{T} \lesssim \frac{GD}{\sqrt{T}} + \frac{\sigma D}{T^{1-\frac{1}{p}}}, \forall \mathbf{x} \in \mathcal{X}.$$

Proof By convexity, $F(\bar{\mathbf{x}}_T) - F(\mathbf{x}) \leq \frac{\sum_{t=1}^T F(\mathbf{x}_t) - F(\mathbf{x})}{T} \leq \frac{\mathbf{R}_T^A(\mathbf{x})}{T}$ is valid for any OCO algorithm A. We conclude from invoking Theorems 1, 2 and 3. $\blacksquare$

To the best of our knowledge, Corollary 1 gives the first and optimal convergence rate for these three algorithms in stochastic optimization with heavy tails. Especially, it implies that once the domain is bounded, the widely implemented SGD algorithm provably converges under heavy-tailed noise without any algorithmic change considered in many prior works, e.g., gradient clipping (Liu and Zhou, 2023; Nguyen et al., 2023).

We are only aware of two works (Vural et al., 2022; Liu and Zhou, 2024) based on Stochastic Mirror Descent (SMD) (Nemirovski and Yudin, 1983) that gave convergence results without clipping. However, they share a common shortcoming, i.e., their bounds are both in the form of $(G + \sigma)D/T^{1-1/p}$, which cannot recover the optimal rate $GD/\sqrt{T}$ when $\sigma = 0$.

Remark 7 *As mentioned in Remark 1, our analysis also provably extends to OMD, which degenerates to SMD for stochastic optimization. However, we highlight a major difference between the existing works mentioned above and ours: the condition on the mirror map. Concretely, the mirror map in Vural et al. (2022); Liu and Zhou (2024) is required to be $\frac{p}{p-1}$ -uniformly convex w.r.t. the studied norm. In contrast, as one can check, our proof technique only needs the mirror map to be 1-strongly convex w.r.t. the studied norm regardless of the value of p . We emphasize that this difference is already significant for the standard 2-norm considered in the work. As explicitly discussed in Section 3.1 of Vural et al. (2022), their framework cannot recover standard SGD when $p \neq 2$, since their mirror map is chosen to be proportional to $\|\cdot\|^{\frac{p}{p-1}}$ to satisfy the $\frac{p}{p-1}$ -uniformly convex requirement. In comparison, for any $p \in (1, 2]$, OMD/SMD with the 1-strongly convex mirror map $\frac{1}{2} \|\cdot\|^2$ exactly corresponds to the OGD/SGD algorithm.*

Lastly, we highlight that for $A = \text{AdaGrad}$, Corollary 1 is not only optimal but also adaptive to the tail index p . As far as we know, no result has achieved this property before. This once again evidences the benefit of adaptive gradient methods.

Convergence of the last iterate. Next, we consider the more challenging last-iterate convergence, which has a long history in stochastic optimization and fruitful results in the case of $p = 2$ (see, e.g., Zhang (2004); Shamir and Zhang (2013); Harvey et al. (2019); Orabona (2020); Jain et al. (2021)). However, less is known about heavy-tailed problems. So far, only two works (Liu and Zhou, 2024; Parletta et al., 2025) have established the last-iterate convergence. The former is based on SMD, and the latter employs gradient clipping in SGD. Unfortunately, their rates are both in the suboptimal order $(G + \sigma)D/T^{1-1/p}$.

We will provide an optimal last-iterate rate based on the following lemma, which reduces the last-iterate convergence to an online learning problem.

Lemma 1 (Theorem 1 of Defazio et al. (2023)) Suppose $\mathbf{x}_1, \dots, \mathbf{x}_T$ and $\mathbf{y}_1, \dots, \mathbf{y}_T$ are two sequences of vectors satisfying $\mathbf{x}_t \in \mathcal{X}$, $\mathbf{x}_1 = \mathbf{y}_1$ and

$$\mathbf{y}_{t+1} = \mathbf{y}_t + \frac{T-t}{T} (\mathbf{x}_{t+1} - \mathbf{x}_t). \quad (8)$$

Given a convex function $F(\mathbf{x})$, let $\ell_t(\mathbf{x}) = \langle \nabla F(\mathbf{y}_t), \mathbf{x} \rangle$. Then for any online learner A , we have

$$F(\mathbf{y}_T) - F(\mathbf{x}) \leq \frac{R_T^A(\mathbf{x})}{T}, \forall \mathbf{x} \in \mathcal{X}.$$

We emphasize that the stochastic gradient $\mathbf{g}_t$ received by A is an estimate of $\nabla F(\mathbf{y}_t)$ instead of $\nabla F(\mathbf{x}_t)$. This flexibility is due to the generality of the OCO framework. Moreover, for OGD, suppose there is no projection step, then (8) is equivalent to $\mathbf{y}_{t+1} = \mathbf{y}_t - \frac{T-t}{T} \eta_t \mathbf{g}_t$, which can be viewed as SGD with a stepsize $\frac{T-t}{T} \eta_t$. For proof of Lemma 1, we refer the interested reader to Defazio et al. (2023).

Corollary 2 Under Assumption 1 for $\ell_t(\mathbf{x}) = \langle \nabla F(\mathbf{y}_t), \mathbf{x} \rangle$, where $\mathbf{y}_t$ satisfies (8), for any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$, we have

$$\mathbb{E}[F(\mathbf{y}_T) - F(\mathbf{x})] \leq \frac{\mathbb{E}[R_T^A(\mathbf{x})]}{T} \lesssim \frac{GD}{\sqrt{T}} + \frac{\sigma D}{T^{1-\frac{1}{p}}}, \forall \mathbf{x} \in \mathcal{X}.$$

Proof Combine Lemma 1 and Theorems 1, 2 and 3 to conclude. ■

As far as we know, Corollary 2 is the first optimal last-iterate convergence rate for stochastic convex optimization with heavy tails, closing the gap in existing works.

One may notice that $\mathbf{y}_t$ itself is not the decision made by the online learner and naturally may ask whether $\mathbf{x}_t$ ensures the last-iterate convergence if we simply pick $\ell_t = F$. The answer turns out to be positive at least for OGD (which is equivalent to SGD now). However, to prove this result, we rely on a technique specialized to stochastic optimization recently developed by Zamani and Glineur (2025); Liu and Zhou (2024). To not diverge from the topic of OCO, we defer the last-iterate convergence of OGD (i.e., $\mathbb{E}[F(\mathbf{x}_T)]$) to Appendix E, in which Theorem 9 gives a general result for any stepsize η_t and Corollary 4 provides a last-iterate rate similar to Corollary 2 (up to an extra logarithmic factor) under the same stepsize $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/p}}$ as in Theorem 1. If T is assumed to be known, Corollary 4 also shows how to set the stepsize in OGD to converge in the optimal rate $\frac{GD}{\sqrt{T}} + \frac{\sigma D}{T^{1-\frac{1}{p}}}$.

Remark 8 If strong convexity further holds, Corollary 5 gives a last-iterate rate in the order $\frac{\log T}{T^{p-1}}$.

4.2. Nonsmooth Nonconvex Optimization

This section contains another application, nonsmooth nonconvex optimization with heavy tails. Due to limited space, we will provide only the necessary background. For more details, we refer the reader to Davis et al. (2022); Kornowski and Shamir (2022a,b); Tian et al. (2022); Jordan et al. (2023); Tian and So (2024) for recent progress. We start with a new set of conditions.

Assumption 2 We consider the following series of assumptions:

- The objective F is lower bounded by $F_\star \triangleq \inf_{\mathbf{x} \in \mathbb{R}^d} F(\mathbf{x}) \in \mathbb{R}$.
- F is differentiable and well-behaved, i.e., $F(\mathbf{x}) - F(\mathbf{y}) = \int_0^1 \langle \nabla F(\mathbf{y} + t(\mathbf{x} - \mathbf{y})), \mathbf{x} - \mathbf{y} \rangle dt$.
- F is G -Lipschitz on $\mathbb{R}^d$, i.e., $\|\nabla F(\mathbf{x})\| \leq G, \forall \mathbf{x} \in \mathbb{R}^d$.
- Given $\mathbf{z}_t \in \mathbb{R}^d$ at the t -th iteration, one can query $\mathbf{g}_t \in \mathbb{R}^d$ satisfying $\mathbb{E}[\mathbf{g}_t \mid \mathcal{F}_{t-1}] = \nabla F(\mathbf{z}_t)$ and $\mathbb{E}[\|\epsilon_t\|^p] \leq \sigma^p$ for some $p \in (1, 2]$ and $\sigma \geq 0$, where $\mathcal{F}_t$ denotes the natural filtration and $\epsilon_t \triangleq \mathbf{g}_t - \nabla F(\mathbf{z}_t)$ is the stochastic noise.

Remark 9 The second point is a mild regularity condition introduced by [Cutkosky et al. \(2023\)](#) and becomes standard in the literature ([Liu et al., 2024](#); [Zhang and Cutkosky, 2024](#); [Ahn and Cutkosky, 2024](#)). See Definition 1 and Proposition 2 of [Cutkosky et al. \(2023\)](#) for more details. In the fourth point, we use the same notation $\mathbf{z}_t$ as in the algorithm being studied later. In fact, it can be arbitrary.

In nonsmooth nonconvex optimization, we aim to find a (δ, ϵ) -stationary point ([Zhang et al., 2020b](#)) (see the formal Definition 2 in Appendix F). This goal can be reduced to finding a point $\mathbf{x} \in \mathbb{R}^d$ such that $\|\nabla F(\mathbf{x})\|_\delta \leq \epsilon$, where $\|\nabla F(\mathbf{x})\|_\delta$ is a quantity introduced by [Cutkosky et al. \(2023\)](#) as follows.

Definition 1 (Definition 5 of [Cutkosky et al. \(2023\)](#)) Given a point $\mathbf{x} \in \mathbb{R}^d$, a number $\delta > 0$ and a differentiable function F , define $\|\nabla F(\mathbf{x})\|_\delta \triangleq \inf_{S \subset B(\mathbf{x}, \delta), \frac{1}{|S|} \sum_{\mathbf{y} \in S} \mathbf{y} = \mathbf{x}} \left\| \frac{1}{|S|} \sum_{\mathbf{y} \in S} \nabla F(\mathbf{y}) \right\|$.

The only existing sample complexity under Assumption 2 is $(G + \sigma)^{\frac{p}{p-1}} \delta^{-1} \epsilon^{-\frac{2p-1}{p-1}}$ in high probability ([Liu et al., 2024](#)), where we only report the dominant term and hide the dependency on the failure probability.

However, on the theoretical side, their result cannot recover the optimal bound $G^2 \delta^{-1} \epsilon^{-3}$ ([Cutkosky et al., 2023](#)) in the deterministic case. On the practical side, their method also employs the gradient clipping step, which introduces a new clipping parameter to tune. In fact, as stated in their Section 5, they observed in experiments that their algorithm without the clipping operation (exactly the algorithm we study next) still works under heavy tails. In addition, in their Section 6, they also explicitly ask whether the requirement to know G and A can be removed.

As will be seen later, we can address these points with the new regret bounds presented before.

4.2.1. ONLINE-TO-NONCONVEX CONVERSION UNDER HEAVY TAILS

Algorithm 4 Online-to-Nonconvex Conversion (O2NC) ([Cutkosky et al., 2023](#))

Input: initial point $\mathbf{y}_0 \in \mathbb{R}^d$, $K \in \mathbb{N}$, $T \in \mathbb{N}$, online learning algorithm A .

for $n = 1$ **to** KT **do**

 Receive $\mathbf{x}_n$ from A

$\mathbf{y}_n = \mathbf{y}_{n-1} + \mathbf{x}_n$

$\mathbf{z}_n = \mathbf{y}_{n-1} + s_n \mathbf{x}_n$ where $s_n \sim \text{Uniform}[0, 1]$ i.i.d.

 Query a stochastic gradient $\mathbf{g}_n$ at $\mathbf{z}_n$

 Send $\mathbf{g}_n$ to A

end for

Remark 10 Note that O2NC is a randomized algorithm. Therefore, the definition of the natural filtration is adjusted to $\mathcal{F}_n \triangleq \sigma(s_1, \mathbf{g}_1, \dots, s_n, \mathbf{g}_n, s_{n+1})$ accordingly.

We provide the Online-to-Nonconvex Conversion (O2NC) framework in Algorithm 4, which serves as a meta algorithm. Roughly speaking, Algorithm 4 reduces a nonconvex optimization problem to an OCO (in fact, OLO) problem, for which the K -shifting regret (see (9)) of the online learner A crucially affects the final convergence rate. However, the existing Theorem 8 of Cutkosky et al. (2023), a general convergence result for the above reduction, cannot directly apply to heavy-tailed noise, since its proof relies on the finite variance condition on $\mathbf{g}_n$ (see Appendix F for more details).

Theorem 4 Under Assumption 2 and let $\mathbf{v}_k \triangleq -D \frac{\sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n)}{\|\sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n)\|}$, $\forall k \in [K]$ for arbitrary $D > 0$, then for any online learning algorithm A in O2NC (Algorithm 4), we have

$$\mathbb{E} \left[\sum_{k=1}^K \frac{1}{K} \left\| \frac{1}{T} \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\| \right] \lesssim \frac{F(\mathbf{y}_0) - F_\star}{DKT} + \frac{\mathbb{E} [\mathbf{R}_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)]}{DKT} + \frac{\sigma}{T^{1-\frac{1}{p}}}.$$

$\mathbf{R}_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)$ in Theorem 4 is called K -shifting regret (Cutkosky et al., 2023), defined as follows:

$$\mathbf{R}_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K) \triangleq \sum_{k=1}^K \sum_{n=(k-1)T+1}^{kT} \ell_n(\mathbf{x}_n) - \ell_n(\mathbf{v}_k) \quad \text{where} \quad \ell_n(\mathbf{x}) \triangleq \langle \mathbf{g}_n, \mathbf{x} \rangle. \quad (9)$$

Theorem 4 here provides a new and the first theoretical guarantee for O2NC under heavy tails. Especially, it recovers Theorem 8 of Cutkosky et al. (2023) when $p = 2$. A remarkable point is that the O2NC algorithm itself does not need any information about p . The proof of Theorem 4 can be found in Appendix F.

4.2.2. CONVERGENCE RATES

Theorem 4 enables us to apply the results presented in Section 3. Concretely, for $\mathcal{X} = \mathcal{B}^d(D)$ and any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$, if we reset the stepsize in A after every T iterations, there will be $\mathbb{E} [\mathbf{R}_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)] \lesssim GDK\sqrt{T} + \sigma DK T^{1/p}$ by our new regret bounds, since $\mathbf{v}_k \in \mathcal{X}$. With a carefully picked D , we obtain the following Theorem 5. Its proof is deferred to Appendix F.

Theorem 5 Under Assumption 2 and let $\Delta \triangleq F(\mathbf{y}_0) - F_\star$ and $\bar{\mathbf{z}}_k \triangleq \frac{1}{T} \sum_{n=(k-1)T+1}^{kT} \mathbf{z}_n$, $\forall k \in [K]$, setting any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$ in O2NC (Algorithm 4) with a domain $\mathcal{X} = \mathcal{B}^d(D)$ for $D = \delta/T$ and resetting the stepsize in A after every T iterations, we have

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{\mathbf{z}}_k)\|_\delta \right] \lesssim \frac{\Delta}{\delta K} + \frac{G}{\sqrt{T}} + \frac{\sigma}{T^{1-\frac{1}{p}}}.$$

Notably, this is the first time confirming that gradient clipping is indeed unnecessary for the O2NC framework, matching the experimental observation of Liu et al. (2024).

Corollary 3 Under the same setting of Theorem 5, suppose we have $N \geq 2$ stochastic gradient budgets, taking $K = \lfloor N/T \rfloor$ and $T = \lceil N/2 \rceil \wedge \left(\left\lceil (\delta GN/\Delta)^{\frac{2}{3}} \right\rceil \vee \left\lceil (\delta \sigma N/\Delta)^{\frac{2}{2p-1}} \right\rceil \right)$, we have

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{\mathbf{z}}_k)\|_\delta \right] \lesssim \frac{G}{\sqrt{N}} + \frac{\sigma}{N^{1-\frac{1}{p}}} + \frac{\Delta}{\delta N} + \frac{G^{\frac{2}{3}} \Delta^{\frac{1}{3}}}{(\delta N)^{\frac{1}{3}}} + \frac{\sigma^{\frac{p}{2p-1}} \Delta^{\frac{p-1}{2p-1}}}{(\delta N)^{\frac{p-1}{2p-1}}}.$$

Corollary 3 is obtained by optimizing K and T in Theorem 5. It implies a sample complexity of $G^2\delta^{-1}\epsilon^{-3} + \sigma^{\frac{p}{p-1}}\delta^{-1}\epsilon^{-\frac{2p-1}{p-1}}$ for finding a (δ, ϵ) -stationary point, improved over the previous bound $(G + \sigma)^{\frac{p}{p-1}}\delta^{-1}\epsilon^{-\frac{2p-1}{p-1}}$ (Liu et al., 2024). Furthermore, leveraging the adaptive feature of AdaGrad, Corollary 6 in Appendix F shows how to set K and T without G , σ , and p , resulting in the first provable rate for O2NC when no problem information is known in advance, which solves the problem asked by Liu et al. (2024).

4.2.3. LOWER BOUNDS

In this part, we provide the first lower bound for nonsmooth nonconvex optimization under heavy tails in the following Theorem 6, the proof of which follows the framework first established in Arjevani et al. (2023) and later developed by Cutkosky et al. (2023) but with some necessary (though minor) variation to make it compatible with heavy-tailed noise.

Theorem 6 *For any given $\Delta > 0$, $G > 0$, $p \in (1, 2]$, $\sigma \geq 0$, $\delta > 0$ and $0 < \epsilon \lesssim \frac{\Delta}{\delta} \wedge \frac{G^2\delta}{\Delta}$, there exists a dimension $d > 0$ depending on the previous parameters such that, for any randomized first-order algorithm (see Definition 4 for a formal description), there exists a G -Lipschitz differentiable function $F : \mathbb{R}^d \rightarrow \mathbb{R}$ satisfying $F(\mathbf{0}) - F_\star \leq \Delta$ and a function $\mathbf{g} : \mathbb{R}^d \times \{0, 1\} \rightarrow \mathbb{R}$ satisfying $\mathbb{E}_r[\mathbf{g}(\mathbf{x}, r)] = \nabla F(\mathbf{x})$ and $\mathbb{E}_r[\|\mathbf{g}(\mathbf{x}, r) - \nabla F(\mathbf{x})\|^p] \leq \sigma^p$ where r follows a certain probability distribution over $\{0, 1\}$ such that the algorithm requires $\gtrsim \Delta\delta^{-1}\epsilon^{-1} + \Delta\sigma^{\frac{p}{p-1}}\delta^{-1}\epsilon^{-\frac{2p-1}{p-1}}$ queries of $\mathbf{g}$ to find a point $\mathbf{z}$ such that $\mathbb{E}[\|\nabla F(\mathbf{z})\|_\delta] \leq \epsilon$.*

For small enough ϵ and $\sigma > 0$, Theorem 6 can be further simplified into a lower bound of $\Delta\sigma^{\frac{p}{p-1}}\delta^{-1}\epsilon^{-\frac{2p-1}{p-1}}$, matching the leading term in the sample complexity derived from the previous Corollary 3. Therefore, Theorem 6 suggests that our Corollary 3 is tight in the high accuracy and noisy regime. As such, Corollary 3 and Theorem 6 together provide a nearly complete characterization of the complexity of finding (δ, ϵ) -stationary points in the heavy-tailed setting.

However, for any general $\epsilon > 0$ or the case $\sigma = 0$, there is still a gap between the upper and lower bounds. Closing this gap could be an interesting direction for the future.

5. Conclusion and Future Work

This paper shows that three classical OCO algorithms, OGD, DA, and AdaGrad, can achieve the optimal in-expectation regret under heavy tails without any algorithmic modification if the feasible set is bounded, and provides some applications in stochastic optimization. The main limitation of our work is that all the proof crucially relies on the bounded domain assumption, which may not always be suitable in practice. Finding a weaker sufficient condition, under which the classical OCO algorithms work with heavy tails provably, is a direction worth studying in the future.

Acknowledgments

The author thanks anonymous reviewers for their helpful comments and valuable feedback.

References

Naman Agarwal, Brian Bullins, Elad Hazan, Sham Kakade, and Karan Singh. Online control with adversarial disturbances. In Kamalika Chaudhuri and Ruslan Salakhutdinov, editors, *Proceedings*

- of the 36th International Conference on Machine Learning, volume 97 of *Proceedings of Machine Learning Research*, pages 111–119. PMLR, 09–15 Jun 2019. URL <https://proceedings.mlr.press/v97/agarwal19c.html>.
- Kwangjun Ahn and Ashok Cutkosky. Adam with model exponential moving average is effective for nonconvex optimization. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, editors, *Advances in Neural Information Processing Systems*, volume 37, pages 94909–94933. Curran Associates, Inc., 2024. URL https://proceedings.neurips.cc/paper_files/paper/2024/file/ac8ec9b4d94c03f0af8c4fe3d5fad4fd-Paper-Conference.pdf.
- 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.
- Amir Beck and Marc Teboulle. Mirror descent and nonlinear projected subgradient methods for convex optimization. *Operations Research Letters*, 31(3):167–175, 2003. ISSN 0167-6377. doi: [https://doi.org/10.1016/S0167-6377\(02\)00231-6](https://doi.org/10.1016/S0167-6377(02)00231-6). URL <https://www.sciencedirect.com/science/article/pii/S0167637702002316>.
- N. Cesa-Bianchi, A. Conconi, and C. Gentile. On the generalization ability of on-line learning algorithms. *IEEE Transactions on Information Theory*, 50(9):2050–2057, 2004. doi: 10.1109/TIT.2004.833339.
- Nicolo Cesa-Bianchi and Gabor Lugosi. *Prediction, Learning, and Games*. Cambridge University Press, 2006.
- Ashok Cutkosky, Harsh Mehta, and Francesco Orabona. Optimal stochastic non-smooth non-convex optimization through online-to-non-convex conversion. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett, editors, *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 6643–6670. PMLR, 23–29 Jul 2023. URL <https://proceedings.mlr.press/v202/cutkosky23a.html>.
- 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. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 6692–6703. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/2c8d9636f74d0207ff4f65956010f450-Paper-Conference.pdf.
- Aaron Defazio, Ashok Cutkosky, Harsh Mehta, and Konstantin Mishchenko. Optimal linear decay learning rate schedules and further refinements. *arXiv preprint arXiv:2310.07831*, 2023.
- John Duchi, Elad Hazan, and Yoram Singer. Adaptive subgradient methods for online learning and stochastic optimization. *Journal of Machine Learning Research*, 12(61):2121–2159, 2011. URL <http://jmlr.org/papers/v12/duchilla.html>.

- Nicholas J. A. Harvey, Christopher Liaw, Yaniv Plan, and Sikander Randhawa. Tight analyses for non-smooth stochastic gradient descent. In Alina Beygelzimer and Daniel Hsu, editors, *Proceedings of the Thirty-Second Conference on Learning Theory*, volume 99 of *Proceedings of Machine Learning Research*, pages 1579–1613. PMLR, 25–28 Jun 2019. URL <https://proceedings.mlr.press/v99/harvey19a.html>.
- Elad Hazan. Introduction to online convex optimization. *Foundations and Trends in Optimization*, 2(3-4):157–325, 2016. ISSN 2167-3888. doi: 10.1561/24000000013. URL <http://dx.doi.org/10.1561/24000000013>.
- Elad Hazan and Karan Singh. Introduction to online control, 2025. URL <https://arxiv.org/abs/2211.09619>.
- Liam Hodgkinson and Michael Mahoney. Multiplicative noise and heavy tails in stochastic optimization. In Marina Meila and Tong Zhang, editors, *Proceedings of the 38th International Conference on Machine Learning*, volume 139 of *Proceedings of Machine Learning Research*, pages 4262–4274. PMLR, 18–24 Jul 2021. URL <https://proceedings.mlr.press/v139/hodgkinson21a.html>.
- Prateek Jain, Dheeraj M. Nagaraj, and Praneeth Netrapalli. Making the last iterate of sgd information theoretically optimal. *SIAM Journal on Optimization*, 31(2):1108–1130, 2021. doi: 10.1137/19M128908X. URL <https://doi.org/10.1137/19M128908X>.
- Michael Jordan, Guy Kornowski, Tianyi Lin, Ohad Shamir, and Manolis Zampetakis. Deterministic nonsmooth nonconvex optimization. In Gergely Neu and Lorenzo Rosasco, editors, *Proceedings of Thirty Sixth Conference on Learning Theory*, volume 195 of *Proceedings of Machine Learning Research*, pages 4570–4597. PMLR, 12–15 Jul 2023. URL <https://proceedings.mlr.press/v195/jordan23a.html>.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Guy Kornowski and Ohad Shamir. Oracle complexity in nonsmooth nonconvex optimization. *Journal of Machine Learning Research*, 23(314):1–44, 2022a. URL <http://jmlr.org/papers/v23/21-1507.html>.
- Guy Kornowski and Ohad Shamir. On the complexity of finding small subgradients in nonsmooth optimization. In *OPT 2022: Optimization for Machine Learning (NeurIPS 2022 Workshop)*, 2022b. URL <https://openreview.net/forum?id=SaRQ4oTqWbP>.
- Guanghui Lan. *First-order and stochastic optimization methods for machine learning*. Springer, 2020.
- Langqi Liu, Yibo Wang, and Lijun Zhang. High-probability bound for non-smooth non-convex stochastic optimization with heavy tails. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 32122–32138. PMLR, 21–27 Jul 2024. URL <https://proceedings.mlr.press/v235/liu24bo.html>.

- Zijian Liu. Online convex optimization with heavy tails: Old algorithms, new regrets, and applications. *arXiv preprint arXiv:2508.07473*, 2025.
- Zijian Liu and Zhengyuan Zhou. Stochastic nonsmooth convex optimization with heavy-tailed noises: High-probability bound, in-expectation rate and initial distance adaptation. *arXiv preprint arXiv:2303.12277*, 2023.
- Zijian Liu and Zhengyuan Zhou. Revisiting the last-iterate convergence of stochastic gradient methods. In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=xxaEhwC1I4>.
- Zijian Liu and Zhengyuan Zhou. Nonconvex stochastic optimization under heavy-tailed noises: Optimal convergence without gradient clipping. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=NKotdPUc3L>.
- Brendan McMahan. Follow-the-regularized-leader and mirror descent: Equivalence theorems and ℓ_1 regularization. In Geoffrey Gordon, David Dunson, and Miroslav Dudík, editors, *Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics*, volume 15 of *Proceedings of Machine Learning Research*, pages 525–533, Fort Lauderdale, FL, USA, 11–13 Apr 2011. PMLR. URL <https://proceedings.mlr.press/v15/mcmahan11b.html>.
- H. Brendan McMahan and Matthew J. Streeter. Adaptive bound optimization for online convex optimization. In *Conference on Learning Theory (COLT)*, pages 244–256. Omnipress, 2010.
- Arkadi Nemirovski and David Yudin. Problem complexity and method efficiency in optimization. *Wiley-Interscience*, 1983.
- Yurii Nesterov. Primal-dual subgradient methods for convex problems. *Mathematical programming*, 120(1):221–259, 2009.
- Ta Duy Nguyen, Thien H Nguyen, Alina Ene, and Huy Nguyen. Improved convergence in high probability of clipped gradient methods with heavy tailed noise. In A. Oh, T. Naumann, A. Globerson, K. Saenko, M. Hardt, and S. Levine, editors, *Advances in Neural Information Processing Systems*, volume 36, pages 24191–24222. Curran Associates, Inc., 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/file/4c454d34f3a4c8d6b4ca85a918e5d7ba-Paper-Conference.pdf.
- Francesco Orabona. A modern introduction to online learning. *arXiv preprint arXiv:1912.13213*, 2019.
- Francesco Orabona. Last iterate of sgd converges (even in unbounded domains). 2020. URL <https://parameterfree.com/2020/08/07/last-iterate-of-sgd-converges-even-in-unbounded-domains/>.
- Daniela Angela Parletta, Andrea Paudice, and Saverio Salzo. An improved analysis of the clipped stochastic subgradient method under heavy-tailed noise, 2025. URL <https://arxiv.org/abs/2410.00573>.

- Maxim Raginsky and Alexander Rakhlin. Information complexity of black-box convex optimization: A new look via feedback information theory. In *2009 47th Annual Allerton Conference on Communication, Control, and Computing (Allerton)*, pages 803–510, 2009. doi: 10.1109/ALLERTON.2009.5394945.
- Herbert Robbins and Sutton Monro. A Stochastic Approximation Method. *The Annals of Mathematical Statistics*, 22(3):400 – 407, 1951. doi: 10.1214/aoms/1177729586. URL <https://doi.org/10.1214/aoms/1177729586>.
- R Tyrrell Rockafellar. *Convex analysis*, volume 28. Princeton university press, 1997.
- Abdurakhmon Sadiev, Marina Danilova, Eduard Gorbunov, Samuel Horváth, Gauthier Gidel, Pavel Dvurechensky, Alexander Gasnikov, and Peter Richtárik. High-probability bounds for stochastic optimization and variational inequalities: the case of unbounded variance. In Andreas Krause, Emma Brunskill, Kyunghyun Cho, Barbara Engelhardt, Sivan Sabato, and Jonathan Scarlett, editors, *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pages 29563–29648. PMLR, 23–29 Jul 2023. URL <https://proceedings.mlr.press/v202/sadiev23a.html>.
- Shai Shalev-Shwartz. Online learning and online convex optimization. *Foundations and Trends in Machine Learning*, 4(2):107–194, 2012. ISSN 1935-8237. doi: 10.1561/22000000018. URL <http://dx.doi.org/10.1561/22000000018>.
- Ohad Shamir and Tong Zhang. Stochastic gradient descent for non-smooth optimization: Convergence results and optimal averaging schemes. In Sanjoy Dasgupta and David McAllester, editors, *Proceedings of the 30th International Conference on Machine Learning*, volume 28 of *Proceedings of Machine Learning Research*, pages 71–79, Atlanta, Georgia, USA, 17–19 Jun 2013. PMLR. URL <https://proceedings.mlr.press/v28/shamir13.html>.
- Umut Simsekli, Levent Sagun, and Mert Gurbuzbalaban. A tail-index analysis of stochastic gradient noise in deep neural networks. In Kamalika Chaudhuri and Ruslan Salakhutdinov, editors, *Proceedings of the 36th International Conference on Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pages 5827–5837. PMLR, 09–15 Jun 2019. URL <https://proceedings.mlr.press/v97/simsekli19a.html>.
- Lai Tian and Anthony Man-Cho So. No dimension-free deterministic algorithm computes approximate stationarities of lipschitzians. *Mathematical Programming*, 208(1):51–74, 2024.
- Lai Tian, Kaiwen Zhou, and Anthony Man-Cho So. On the finite-time complexity and practical computation of approximate stationarity concepts of Lipschitz functions. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvari, Gang Niu, and Sivan Sabato, editors, *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, pages 21360–21379. PMLR, 17–23 Jul 2022. URL <https://proceedings.mlr.press/v162/tian22a.html>.
- Tijmen Tieleman, Geoffrey Hinton, et al. Lecture 6.5-rmsprop: Divide the gradient by a running average of its recent magnitude. *COURSERA: Neural networks for machine learning*, 4(2):26–31, 2012.

- Nuri Mert Vural, Lu Yu, Krishna Balasubramanian, Stanislav Volgushev, and Murat A Erdogdu. Mirror descent strikes again: Optimal stochastic convex optimization under infinite noise variance. In Po-Ling Loh and Maxim Raginsky, editors, *Proceedings of Thirty Fifth Conference on Learning Theory*, volume 178 of *Proceedings of Machine Learning Research*, pages 65–102. PMLR, 02–05 Jul 2022. URL <https://proceedings.mlr.press/v178/vural22a.html>.
- Rachel Ward, Xiaoxia Wu, and Leon Bottou. AdaGrad stepsizes: Sharp convergence over nonconvex landscapes. In Kamalika Chaudhuri and Ruslan Salakhutdinov, editors, *Proceedings of the 36th International Conference on Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pages 6677–6686. PMLR, 09–15 Jun 2019. URL <https://proceedings.mlr.press/v97/ward19a.html>.
- Lin Xiao. Dual averaging method for regularized stochastic learning and online optimization. In Y. Bengio, D. Schuurmans, J. Lafferty, C. Williams, and A. Culotta, editors, *Advances in Neural Information Processing Systems*, volume 22. Curran Associates, Inc., 2009. URL https://proceedings.neurips.cc/paper_files/paper/2009/file/7cce53cf90577442771720a370c3c723-Paper.pdf.
- Moslem Zamani and François Glineur. Exact convergence rate of the last iterate in subgradient methods. *SIAM Journal on Optimization*, 35(3):2182–2201, 2025. doi: 10.1137/24M1717762. URL <https://doi.org/10.1137/24M1717762>.
- Jingzhao Zhang, Sai Praneeth Karimireddy, Andreas Veit, Seungyeon Kim, Sashank Reddi, Sanjiv Kumar, and Suvrit Sra. Why are adaptive methods good for attention models? In H. Larochelle, M. Ranzato, R. Hadsell, M.F. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems*, volume 33, pages 15383–15393. Curran Associates, Inc., 2020a. URL https://proceedings.neurips.cc/paper_files/paper/2020/file/b05b57f6add810d3b7490866d74c0053-Paper.pdf.
- Jingzhao Zhang, Hongzhou Lin, Stefanie Jegelka, Suvrit Sra, and Ali Jadbabaie. Complexity of finding stationary points of nonconvex nonsmooth functions. In Hal DaumÃ III and Aarti Singh, editors, *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 11173–11182. PMLR, 13–18 Jul 2020b. URL <https://proceedings.mlr.press/v119/zhang20p.html>.
- Jiujia Zhang and Ashok Cutkosky. Parameter-free regret in high probability with heavy tails. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh, editors, *Advances in Neural Information Processing Systems*, volume 35, pages 8000–8012. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/349956dee974cfdcbbb2d06afad5dd4a-Paper-Conference.pdf.
- Qinzi Zhang and Ashok Cutkosky. Random scaling and momentum for non-smooth non-convex optimization. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 58780–58799. PMLR, 21–27 Jul 2024. URL <https://proceedings.mlr.press/v235/zhang24k.html>.

Tong Zhang. Solving large scale linear prediction problems using stochastic gradient descent algorithms. In *Proceedings of the twenty-first international conference on Machine learning*, page 116, 2004.

Martin Zinkevich. Online convex programming and generalized infinitesimal gradient ascent. In *Proceedings of the 20th international conference on machine learning (icml-03)*, pages 928–936, 2003.

Appendix A. Hardness of High-Probability Bounds Even on Bounded Domains

Without additional assumptions on noise or algorithmic changes to the classical methods studied in the paper, even on a bounded domain, we are inclined to believe that a high-probability bound is unlikely (at least within the current analysis framework).

An upper-bound perspective. When deriving regret bounds for OGD/DA/AdaGrad in the current analysis, one needs to upper bound the term $\sum_{t=1}^T \|\epsilon_t\|^p$. Under Assumption 1, this can be easily done in expectation. However, if one instead wants to obtain a high-probability bound with a polylogarithmic dependency on the failure probability δ , this is unlikely to be achieved.

As an example, we consider $d = 1$ for simplicity. Let $\epsilon_t = e_t \xi_t$ where e_t and ξ_t are independent of each other and the history, satisfying that $\Pr[e_t = \pm 1] = \frac{1}{2}$ and $\Pr[\xi_t > z] = \left(\frac{z_\star}{z}\right)^{p+a} \mathbb{1}[z \geq z_\star] + \mathbb{1}[z < z_\star]$, where $a > 0$ can be arbitrarily small and $z_\star > 0$ can be any number. In other words, e_t follows the Rademacher distribution and ξ_t follows the Pareto distribution. As one can check, there are $\mathbb{E}[\epsilon_t | \mathcal{F}_{t-1}] = 0$ and $\mathbb{E}[\|\epsilon_t\|^p] < \infty$. Crucially, we can find that, for $z \geq z_\star^p$ and any $T \geq 1$,

$$\Pr \left[\sum_{t=1}^T \|\epsilon_t\|^p > z \right] = \Pr \left[\sum_{t=1}^T \xi_t^p > z \right] \geq \Pr \left[\xi_1 > z^{\frac{1}{p}} \right] = \frac{z_\star^{p+a}}{z^{1+\frac{a}{p}}},$$

indicating that the tail of $\sum_{t=1}^T \|\epsilon_t\|^p$ decays at least polynomially, which further suggests that a polylogarithmic dependency on the failure probability δ (equivalently, an exponentially decaying tail) is impossible.

A lower-bound perspective. We consider the special case of OCO, stochastic convex optimization (i.e., $\ell_t = F$), and show a negative result for OGD/SGD when $d = 1$ and $p = 2$.

We set $F(\mathbf{x}) = |\mathbf{x}|$ with $\mathcal{X} \triangleq [-1, 1]$ and $\mathbf{g}_t = \mathbf{g}(\mathbf{x}_t; \xi_t) = \text{sgn}(\mathbf{x}_t) \xi_t$ (with $\text{sgn}(0) = 1$), where ξ_t is independent of each other and follows the Pareto distribution satisfying $\Pr[\xi_t > z] = \left(\frac{b-1}{bz}\right)^b \mathbb{1}[z \geq \frac{b-1}{b}] + \mathbb{1}[z < \frac{b-1}{b}]$, in which $b \triangleq 2 + 2a$ and $a > 0$ can be arbitrarily small. As one can check, $\mathbb{E}[\mathbf{g}_t | \mathcal{F}_{t-1}] = \nabla F(\mathbf{x}_t)$ due to $\mathbb{E}[\xi_t] = 1$ and $\mathbb{E}[\|\epsilon_t\|^2] < \infty$ due to $\mathbb{E}[\xi_t^2] < \infty$.

Proposition 1 *Consider the setting described above, if OGD/SGD with $\eta_t = \frac{\eta}{\sqrt{t}}$ (the standard time-varying stepsize) for $\eta \leq 1$ guarantees a high-probability rate $\Pr \left[F(\mathbf{x}_t) > \frac{\text{polylog}(1/\delta)}{t^c} \right] \leq \delta, \forall \delta \in (0, 1], \forall t \in \mathbb{N}$, then there must be $c \leq \frac{a}{b} = \frac{a}{2+2a}$. Since $a > 0$ can be arbitrarily small, this suggests that no meaningful high-probability rates (i.e., decay polynomially in t) exist, at least for the last iterate.*

Remark 11 *We clarify that the above result does not rule out the possibility that OGD/SGD may admit a meaningful high-probability bound for more sophisticated stepsize or the average iterate.*

Proof For $t \in \mathbb{N}$, we take $\delta = \frac{1}{t^2}$ to have $\Pr \left[F(\mathbf{x}_t) > \frac{\text{polylog}(t^2)}{t^c} \right] \leq \frac{1}{t^2}$. By the Borel-Cantelli Lemma, we obtain

$$\Pr \left[F(\mathbf{x}_t) > \frac{\text{polylog}(t^2)}{t^c} \text{ i.o.} \right] = 0 \Leftrightarrow \Pr \left[F(\mathbf{x}_t) \leq \frac{\text{polylog}(t^2)}{t^c} \text{ eventually} \right] = 1. \quad (10)$$

Next, we define the event $A_t \triangleq \left\{ \xi_t > \frac{1}{\eta_t t^{a/b}} \right\}$. Note that $\frac{1}{\eta_t t^{a/b}} = \frac{t^{1/2-a/b}}{\eta} \geq 1 > \frac{b-1}{b}$, implying that $\Pr[A_t] = \left(\frac{(b-1)\eta}{b}\right)^b \frac{1}{t} \Rightarrow \sum_{t=1}^{\infty} \Pr[A_t] < \infty$. Since A_t is independent of each

other, then by the second Borel-Cantelli Lemma,

$$\Pr[A_t \text{ i.o.}] = 1. \quad (11)$$

Now, let us write $A_t = A_{t,1} \cup A_{t,2}$, where

$$A_{t,1} \triangleq A_t \cap \{|\mathbf{x}_t - \eta_t \mathbf{g}_t| > 1\} \quad \text{and} \quad A_{t,2} \triangleq A_t \cap \{|\mathbf{x}_t - \eta_t \mathbf{g}_t| \leq 1\}$$

are two disjoint events. Moreover, we introduce another event $B_t \triangleq \left\{|\mathbf{x}_{t+1}| + |\mathbf{x}_t| \geq \frac{1}{t^{a/b}}\right\}$ and observe that:

- There is $A_{t,1} \subseteq \{|\mathbf{x}_t - \eta_t \mathbf{g}_t| > 1\} \subseteq \{|\mathbf{x}_{t+1}| = 1\} \subseteq B_t$.
- There is $A_{t,2} = A_t \cap \{\mathbf{x}_{t+1} = \mathbf{x}_t - \eta_t \mathbf{g}_t\} \subseteq A_t \cap \{|\mathbf{x}_{t+1}| + |\mathbf{x}_t| \geq \eta_t |\mathbf{g}_t| = \eta_t \xi_t\} \subseteq B_t$.

Therefore, we always have

$$A_t \subseteq B_t \stackrel{(11)}{\Rightarrow} \Pr[B_t \text{ i.o.}] = 1. \quad (12)$$

Finally, for a sample path satisfying $|\mathbf{x}_t| = F(\mathbf{x}_t) \leq \frac{\text{polylog}(t^2)}{t^c}$ eventually and $|\mathbf{x}_{t+1}| + |\mathbf{x}_t| \geq \frac{1}{t^{a/b}}$ i.o. (the existence is due to (10) and (12)), there must be $\frac{1}{t^{a/b}} \lesssim \frac{\text{polylog}(t^2)}{t^c}$ for infinitely many $t \in \mathbb{N}$, which implies that $c \leq \frac{a}{b} = \frac{a}{2+2a}$. $\blacksquare$

Appendix B. Missing Proofs for Online Gradient Descent

This section provides missing proofs for regret bounds of OGD. Before showing the formal proof, we recall the following core inequality that holds for any $\mathbf{x} \in \mathcal{X}$ given in (6):

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x}_t - \mathbf{x}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2\eta_t} + \eta_t G^2 + C(p)\eta_t^{p-1} \|\epsilon_t\|^p D^{2-p}. \quad (13)$$

The key to establishing the above result is showing

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_t} \leq \eta_t G^2 + C(p)\eta_t^{p-1} \|\epsilon_t\|^p D^{2-p}, \quad (14)$$

the proof of which is by combining (3), (4), and (5) established in the main text.

B.1. Proof of Theorem 1

Proof For any $\mathbf{x} \in \mathcal{X}$, sum up (13) from $t = 1$ to T and drop the term $-\frac{\|\mathbf{x}_{T+1} - \mathbf{x}\|^2}{2\eta_T}$ to obtain

$$\begin{aligned} & \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \\ & \leq \frac{\|\mathbf{x}_1 - \mathbf{x}\|^2}{2\eta_1} + \sum_{t=1}^{T-1} \left(\frac{1}{\eta_{t+1}} - \frac{1}{\eta_t} \right) \frac{\|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2} + \sum_{t=1}^T \eta_t G^2 + C(p)\eta_t^{p-1} \|\epsilon_t\|^p D^{2-p} \end{aligned} \quad (15)$$

$$\leq \frac{D^2}{\eta_T} + \sum_{t=1}^T \eta_t G^2 + C(p)\eta_t^{p-1} \|\epsilon_t\|^p D^{2-p}, \quad (16)$$

where the last step is due to $\|\mathbf{x}_t - \mathbf{x}\| \leq D, \forall t \in [T]$ and $\eta_{t+1} \leq \eta_t, \forall t \in [T-1]$.

Taking expectations on both sides of (16) yields that

$$\mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] \leq \frac{D^2}{\eta_T} + \sum_{t=1}^T \eta_t G^2 + C(p) \eta_t^{p-1} \sigma^p D^{2-p}, \quad (17)$$

where for the L.H.S., we use $\mathbb{E} [\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle] = \mathbb{E} [\mathbb{E} [\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \mid \mathcal{F}_{t-1}]]$ and

$$\mathbb{E} [\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \mid \mathcal{F}_{t-1}] = \langle \mathbb{E} [\mathbf{g}_t \mid \mathcal{F}_{t-1}], \mathbf{x}_t - \mathbf{x} \rangle = \langle \nabla \ell_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x} \rangle \geq \ell_t(\mathbf{x}_t) - \ell_t(\mathbf{x}), \quad (18)$$

for the R.H.S., we use $\mathbb{E} [\|\epsilon_t\|^p] \leq \sigma^p$.

Finally, we plug $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/p}}, \forall t \in [T]$ into (17), then use $\sum_{t=1}^T \frac{1}{\sqrt{t}} \lesssim \sqrt{T}$ and $\sum_{t=1}^T \frac{1}{t^{1-1/p}} \lesssim T^{1/p}$ to conclude

$$\mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}.$$

■

B.2. Extension to Online Strongly Convex Optimization

Next, we extend Theorem 1 to the strongly convex case, i.e., $\exists \mu > 0$ such that for all $t \in [T]$,

$$\frac{\mu}{2} \|\mathbf{x} - \mathbf{y}\|^2 + \langle \nabla \ell_t(\mathbf{y}), \mathbf{x} - \mathbf{y} \rangle + \ell_t(\mathbf{y}) \leq \ell_t(\mathbf{x}), \forall \mathbf{x}, \mathbf{y} \in \mathcal{X}, \nabla \ell_t(\mathbf{y}) \in \partial \ell_t(\mathbf{y}). \quad (19)$$

In this setting, it is well known that OGD achieves a logarithmic regret bound when $p = 2$ (Hazan, 2016; Orabona, 2019). Theorem 7 below provides the first provable result for $p < 2$.

Theorem 7 *Under Assumption 1 and additionally assuming (19), taking $\eta_t = \frac{1}{\mu t}$ in OGD (Algorithm 1), we have*

$$\mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] \lesssim \frac{G^2 (1 + \log T)}{\mu} + \frac{\sigma^p D^{2-p}}{\mu^{p-1}} \times \begin{cases} T^{2-p} & p \in (1, 2) \\ 1 + \log T & p = 2 \end{cases}, \forall \mathbf{x} \in \mathcal{X}.$$

Theorem 7 shows that under strong convexity, OGD for $p \in (1, 2)$ achieves a better sublinear regret T^{2-p} than $T^{1/p}$ in Theorem 1 as $2 - p \leq 1/p, \forall p > 0$. One point we highlight here is that the stepsize $\eta_t = \frac{1}{\mu t}$ is commonly used in the OCO literature and is independent of the tail index p .

However, in contrast to Theorem 1, we suspect Theorem 7 is not tight in T for $p \in (1, 2)$. The reason is that for nonsmooth strongly convex optimization with heavy tails (i.e., $\ell_t = F, \forall t \in [T]$ where F is strongly convex), Theorem 7 can convert to a convergence rate only in the order of $1/T^{p-1}$, which is worse than the lower bound $1/T^{2-2/p}$ (Zhang et al., 2020a). Therefore, we conjecture that a way to obtain a better regret bound than T^{2-p} exists, which we leave as future work.

Proof [Proof of Theorem 7] For any $\mathbf{x} \in \mathcal{X}$, we take expectations on both sides of (15) to have

$$\begin{aligned} \mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] &\leq \left(\frac{1}{\eta_1} - \mu \right) \frac{\|\mathbf{x}_1 - \mathbf{x}\|^2}{2} + \sum_{t=1}^{T-1} \left(\frac{1}{\eta_{t+1}} - \frac{1}{\eta_t} - \mu \right) \frac{\mathbb{E} [\|\mathbf{x}_{t+1} - \mathbf{x}\|^2]}{2} \\ &\quad + \sum_{t=1}^T \eta_t G^2 + C(p) \eta_t^{p-1} \sigma^p D^{2-p}, \end{aligned} \quad (20)$$

where for the L.H.S., we follow a similar step of reasoning out (18) but instead using

$$\langle \nabla \ell_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x} \rangle \geq \ell_t(\mathbf{x}_t) - \ell_t(\mathbf{x}) + \frac{\mu}{2} \|\mathbf{x}_t - \mathbf{x}\|^2,$$

for the R.H.S., we use $\mathbb{E}[\|\epsilon_t\|^p] \leq \sigma^p$.

Next, we plug $\eta_t = \frac{1}{\mu t}$, $\forall t \in [T]$ into (20) to obtain

$$\mathbb{E} \left[R_T^{\text{OGD}}(\mathbf{x}) \right] \lesssim \sum_{t=1}^T \frac{G^2}{\mu t} + \frac{\sigma^p D^{2-p}}{\mu^{p-1} t^{p-1}} \lesssim \frac{G^2 (1 + \log T)}{\mu} + \frac{\sigma^p D^{2-p}}{\mu^{p-1}} \times \begin{cases} T^{2-p} & p \in (1, 2) \\ 1 + \log T & p = 2 \end{cases}.$$

■

Appendix C. Missing Proofs for Dual Averaging

This section provides missing proofs for regret bounds of DA.

C.1. Proof of Theorem 2

Proof Let $L_t(\mathbf{x}) \triangleq \frac{\|\mathbf{x} - \mathbf{x}_1\|^2}{2\eta_{t-1}} + \sum_{s=1}^{t-1} \langle \mathbf{g}_s, \mathbf{x} \rangle$, $\forall t \in [T+1]$, where $\eta_0 \triangleq \eta_1$. Then DA can be equivalently written as

$$\mathbf{x}_t = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} L_t(\mathbf{x}), \forall t \in [T+1].$$

By Lemma 7.1 of Orabona (2019), for any $\mathbf{x} \in \mathcal{X}$,

$$\begin{aligned} \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle &= \frac{\|\mathbf{x} - \mathbf{x}_1\|^2}{2\eta_T} + L_{T+1}(\mathbf{x}_{T+1}) - L_{T+1}(\mathbf{x}) + \sum_{t=1}^T L_t(\mathbf{x}_t) + \langle \mathbf{g}_t, \mathbf{x}_t \rangle - L_{t+1}(\mathbf{x}_{t+1}) \\ &\leq \frac{\|\mathbf{x} - \mathbf{x}_1\|^2}{2\eta_T} + \sum_{t=1}^T L_t(\mathbf{x}_t) - L_{t+1}(\mathbf{x}_{t+1}) + \langle \mathbf{g}_t, \mathbf{x}_t \rangle, \end{aligned}$$

where the inequality holds by $L_{T+1}(\mathbf{x}_{T+1}) \leq L_{T+1}(\mathbf{x})$, $\forall \mathbf{x} \in \mathcal{X}$ due to $\mathbf{x}_{T+1} = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} L_{T+1}(\mathbf{x})$. Note that for any $t \in [T]$,

$$\begin{aligned} &L_t(\mathbf{x}_t) - L_{t+1}(\mathbf{x}_{t+1}) + \langle \mathbf{g}_t, \mathbf{x}_t \rangle \\ &= L_t(\mathbf{x}_t) - L_t(\mathbf{x}_{t+1}) + \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle + \frac{\|\mathbf{x}_{t+1} - \mathbf{x}_1\|^2}{2\eta_{t-1}} - \frac{\|\mathbf{x}_{t+1} - \mathbf{x}_1\|^2}{2\eta_t} \\ &\stackrel{(a)}{\leq} L_t(\mathbf{x}_t) - L_t(\mathbf{x}_{t+1}) + \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle \\ &\stackrel{(b)}{\leq} \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}}, \end{aligned}$$

where (a) is by $\eta_t \leq \eta_{t-1}$, $\forall t \in [T]$ and (b) holds because L_t is $\frac{1}{\eta_{t-1}}$ -strongly convex and $\mathbf{x}_t = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} L_t(\mathbf{x})$, which together imply

$$L_t(\mathbf{x}_t) - L_t(\mathbf{x}_{t+1}) \leq \langle \nabla L_t(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}} \leq -\frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}}.$$

Therefore, we have

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x} - \mathbf{x}_1\|^2}{2\eta_T} + \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}}. \quad (21)$$

By the same argument as proving (14) but replacing η_t with η_{t-1} , there is

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}} \leq \eta_{t-1} G^2 + C(p) \eta_{t-1}^{p-1} \|\epsilon_t\|^p D^{2-p}.$$

As such, we know

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x} - \mathbf{x}_1\|^2}{2\eta_T} + \sum_{t=1}^T \eta_{t-1} G^2 + C(p) \eta_{t-1}^{p-1} \|\epsilon_t\|^p D^{2-p}. \quad (22)$$

Finally, following similar steps in proving Theorem 1 in Appendix B, we conclude

$$\mathbb{E} \left[R_T^{\text{DA}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}.$$

■

C.2. Dual Averaging with an Adaptive Stepsize

We show that DA with an adaptive stepsize can also achieve the optimal regret $GD\sqrt{T} + \sigma DT^{1/p}$.

Theorem 8 *Under Assumption 1, taking $\eta_t = 2DV_t^{-1/2}$ and $V_t = \sum_{s=1}^t \|\mathbf{g}_s\|^2$ in DA (Algorithm 2), we have*

$$\mathbb{E} \left[R_T^{\text{DA}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}, \forall \mathbf{x} \in \mathcal{X}.$$

Proof For any $\mathbf{x} \in \mathcal{X}$, we have

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \stackrel{(21)}{\leq} \frac{\|\mathbf{x} - \mathbf{x}_1\|^2}{2\eta_T} + \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}}, \quad (23)$$

where $\eta_0 \triangleq \eta_1$. On the one hand, we can use AM-GM inequality to bound

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}} \leq \frac{\eta_{t-1} \|\mathbf{g}_t\|^2}{2}.$$

On the other hand, we know

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}} \leq \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle \leq \|\mathbf{g}_t\| \|\mathbf{x}_t - \mathbf{x}_{t+1}\| \leq \|\mathbf{g}_t\| D, \quad (24)$$

where the second step is by Cauchy-Schwarz inequality. Therefore, for any $t \geq 2$,

$$\begin{aligned} \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_{t-1}} &\leq \frac{\eta_{t-1} \|\mathbf{g}_t\|^2}{2} \wedge \|\mathbf{g}_t\| D \stackrel{(a)}{\leq} \frac{2}{\frac{2}{\eta_{t-1} \|\mathbf{g}_t\|^2} + \frac{1}{\|\mathbf{g}_t\| D}} \\ &\stackrel{(b)}{=} \frac{2D \|\mathbf{g}_t\|^2}{\sqrt{\sum_{s=1}^{t-1} \|\mathbf{g}_s\|^2} + \|\mathbf{g}_t\|} \stackrel{(c)}{\leq} \frac{2D \|\mathbf{g}_t\|^2}{\sqrt{\sum_{s=1}^t \|\mathbf{g}_s\|^2}}, \end{aligned} \quad (25)$$

where (a) is due to $x \wedge y \leq \frac{2}{x^{-1}+y^{-1}}, \forall x, y > 0$, (b) is by $\eta_{t-1} = \frac{2D}{\sqrt{\sum_{s=1}^{t-1} \|\mathbf{g}_s\|^2}}$, and (c) holds because of $\sqrt{\sum_{s=1}^t \|\mathbf{g}_s\|^2} \leq \sqrt{\sum_{s=1}^{t-1} \|\mathbf{g}_s\|^2} + \|\mathbf{g}_t\|$. Note that (25) is also true for $t = 1$ by (24). Combine (23) and (25) and use $\|\mathbf{x} - \mathbf{x}_1\| \leq D$ to obtain

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{D^2}{2\eta_T} + \sum_{t=1}^T \frac{2D \|\mathbf{g}_t\|^2}{\sqrt{\sum_{s=1}^t \|\mathbf{g}_s\|^2}} = \frac{D^2}{2\eta_T} + \sum_{t=1}^T \eta_t \|\mathbf{g}_t\|^2,$$

which only differs from (26) by a constant. Hence, by a similar proof for (28), there is

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \lesssim D \left[\sqrt{\sum_{t=1}^T \|\nabla \ell_t(\mathbf{x}_t)\|^2} + \left(\sum_{t=1}^T \|\epsilon_t\|^p \right)^{\frac{1}{p}} \right],$$

implying

$$\mathbb{E} \left[\mathbf{R}_T^{\text{DA}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}.$$

■

Appendix D. Missing Proofs for AdaGrad

This section provides missing proofs for regret bounds of AdaGrad.

D.1. Proof of Theorem 3

Proof As mentioned, AdaGrad can be viewed as OGD with a stepsize $\eta_t = \frac{\eta}{\sqrt{V_t}} = \frac{\eta}{\sqrt{\sum_{s=1}^t \|\mathbf{g}_s\|^2}}$. Therefore, we can use (1) for AdaGrad to know for any $\mathbf{x} \in \mathcal{X}$,

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{\|\mathbf{x}_t - \mathbf{x}\|^2 - \|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2\eta_t} + \frac{\eta_t \|\mathbf{g}_t\|^2}{2}.$$

Sum up the above inequality from $t = 1$ to T and drop the term $-\frac{\|\mathbf{x}_{T+1} - \mathbf{x}\|^2}{2\eta_T}$ to have

$$\begin{aligned} \sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle &\leq \frac{\|\mathbf{x}_1 - \mathbf{x}\|^2}{2\eta_1} + \sum_{t=1}^{T-1} \left(\frac{1}{\eta_{t+1}} - \frac{1}{\eta_t} \right) \frac{\|\mathbf{x}_{t+1} - \mathbf{x}\|^2}{2} + \sum_{t=1}^T \frac{\eta_t \|\mathbf{g}_t\|^2}{2} \\ &\leq \frac{D^2}{2\eta_T} + \sum_{t=1}^T \frac{\eta_t \|\mathbf{g}_t\|^2}{2}, \end{aligned} \tag{26}$$

where the last step is by $\|\mathbf{x}_t - \mathbf{x}\| \leq D, \forall t \in [T]$ and $\eta_{t+1} \leq \eta_t, \forall t \in [T-1]$.

Next, observe that for any $t \in [T]$,

$$\|\mathbf{g}_t\|^2 = \frac{\eta^2}{\eta_t^2} - \frac{\eta^2}{\eta_{t-1}^2} = \eta^2 \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} \right) \left(\frac{1}{\eta_t} + \frac{1}{\eta_{t-1}} \right) \leq \frac{2\eta^2}{\eta_t} \left(\frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} \right),$$

where $1/\eta_0$ should be read as 0. The above inequality implies

$$\sum_{t=1}^T \frac{\eta_t \|\mathbf{g}_t\|^2}{2} \leq \eta^2 \sum_{t=1}^T \frac{1}{\eta_t} - \frac{1}{\eta_{t-1}} = \frac{\eta^2}{\eta_T}. \quad (27)$$

Combine (26) and (27) to have

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \frac{D^2}{2\eta_T} + \frac{\eta^2}{\eta_T} = \left(\frac{D^2}{2\eta} + \eta \right) \sqrt{\sum_{t=1}^T \|\mathbf{g}_t\|^2}.$$

Note that there is

$$\begin{aligned} \sqrt{\sum_{t=1}^T \|\mathbf{g}_t\|^2} &\leq \sqrt{\sum_{t=1}^T 2 \|\nabla \ell_t(\mathbf{x}_t)\|^2 + 2 \|\epsilon_t\|^2} \leq \sqrt{2 \sum_{t=1}^T \|\nabla \ell_t(\mathbf{x}_t)\|^2} + \sqrt{2 \sum_{t=1}^T \|\epsilon_t\|^2} \\ &\leq \sqrt{2 \sum_{t=1}^T \|\nabla \ell_t(\mathbf{x}_t)\|^2} + \sqrt{2} \left(\sum_{t=1}^T \|\epsilon_t\|^p \right)^{\frac{1}{p}}, \end{aligned}$$

where the last step is due to $\|\cdot\|_2 \leq \|\cdot\|_p$ for any $p \in [1, 2]$. Hence, we obtain

$$\sum_{t=1}^T \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x} \rangle \leq \sqrt{2} \left(\frac{D^2}{2\eta} + \eta \right) \left[\sqrt{\sum_{t=1}^T \|\nabla \ell_t(\mathbf{x}_t)\|^2} + \left(\sum_{t=1}^T \|\epsilon_t\|^p \right)^{\frac{1}{p}} \right]. \quad (28)$$

We take expectations on both sides of (28), then apply Hölder's inequality to have

$$\mathbb{E} \left[\left(\sum_{t=1}^T \|\epsilon_t\|^p \right)^{\frac{1}{p}} \right] \leq \left(\sum_{t=1}^T \mathbb{E} [\|\epsilon_t\|^p] \right)^{\frac{1}{p}} \leq \sigma T^{\frac{1}{p}},$$

and finally plug in $\eta = D/\sqrt{2}$ to conclude

$$\mathbb{E} \left[R_T^{\text{AdaGrad}}(\mathbf{x}) \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}.$$

■

Appendix E. Missing Proofs for Applications: Nonsmooth Convex Optimization

We prove the following last-iterate convergence result for SGD (i.e., OGD for stochastic optimization) under heavy-tailed noise. The proof of Theorem 9 is inspired by [Zamani and Glineur \(2025\)](#); [Liu and Zhou \(2024\)](#).

Theorem 9 *Under Assumption 1 and additionally assuming (19) (possibly $\mu = 0$) for $\ell_t(\mathbf{x}) = F(\mathbf{x})$, for any stepsize $0 < \eta_t \leq \frac{1}{\mu}$ in OGD (Algorithm 1), let $\gamma_t \triangleq \frac{\eta_t}{\prod_{s=2}^t (1 - \mu\eta_s)}$, we have*

$$\mathbb{E} [F(\mathbf{x}_T) - F(\mathbf{x})] \lesssim \frac{(1 - \mu\eta_1)D^2}{\sum_{t=1}^{T-1} \gamma_t} + G^2 \sum_{t=1}^{T-1} \frac{\gamma_t \eta_t}{\sum_{s=t}^{T-1} \gamma_s} + \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\gamma_t \eta_t^{p-1}}{\sum_{s=t}^{T-1} \gamma_s}, \forall \mathbf{x} \in \mathcal{X}.$$

Proof Given $\mathbf{x} \in \mathcal{X}$, we recursively define

$$\mathbf{y}_0 \triangleq \mathbf{x} \quad \text{and} \quad \mathbf{y}_t \triangleq \left(1 - \frac{w_{t-1}}{w_t}\right) \mathbf{x}_t + \frac{w_{t-1}}{w_t} \mathbf{y}_{t-1}, \forall t \in [T], \quad (29)$$

in which

$$w_t \triangleq \frac{\gamma_T}{\sum_{s=t}^T \gamma_s}, \forall t \in [T] \quad \text{and} \quad w_0 \triangleq w_1. \quad (30)$$

Equivalently, $\mathbf{y}_t$ can be written into a convex combination of $\mathbf{x}, \mathbf{x}_1, \dots, \mathbf{x}_t$ as

$$\mathbf{y}_t = \frac{w_0}{w_t} \mathbf{x} + \sum_{s=1}^t \frac{w_s - w_{s-1}}{w_t} \mathbf{x}_s, \forall t \in \{0\} \cup [T]. \quad (31)$$

Therefore, $\mathbf{y}_t$ also falls into $\mathcal{X}$ and satisfies $\mathbf{y}_t \in \mathcal{F}_{t-1}$.

We invoke (2) for $\mathbf{y}_t$ to obtain

$$\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{y}_t \rangle \leq \frac{\|\mathbf{x}_t - \mathbf{y}_t\|^2 - \|\mathbf{x}_{t+1} - \mathbf{y}_t\|^2}{2\eta_t} + \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle - \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_t},$$

and then bound

$$\begin{aligned} \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle &= \langle \nabla F(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x}_{t+1} \rangle + \langle \boldsymbol{\epsilon}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle \\ &\leq F(\mathbf{x}_t) - F(\mathbf{x}_{t+1}) + 4\eta_t G^2 + C(p)\eta_t^{p-1} \|\boldsymbol{\epsilon}_t\|^p D^{2-p} + \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{2\eta_t}, \end{aligned}$$

where, in the second step, we use (5) to bound $\langle \boldsymbol{\epsilon}_t, \mathbf{x}_t - \mathbf{x}_{t+1} \rangle$ and the following step to bound $\langle \nabla F(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x}_{t+1} \rangle$,

$$\begin{aligned} \langle \nabla F(\mathbf{x}_t), \mathbf{x}_t - \mathbf{x}_{t+1} \rangle &= \langle \nabla F(\mathbf{x}_{t+1}), \mathbf{x}_t - \mathbf{x}_{t+1} \rangle + \langle \nabla F(\mathbf{x}_t) - \nabla F(\mathbf{x}_{t+1}), \mathbf{x}_t - \mathbf{x}_{t+1} \rangle \\ &\leq F(\mathbf{x}_t) - F(\mathbf{x}_{t+1}) + 2G \|\mathbf{x}_t - \mathbf{x}_{t+1}\| \\ &\leq F(\mathbf{x}_t) - F(\mathbf{x}_{t+1}) + 4\eta_t G^2 + \frac{\|\mathbf{x}_t - \mathbf{x}_{t+1}\|^2}{4\eta_t}. \end{aligned}$$

Hence, we know

$$\begin{aligned} \langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{y}_t \rangle &\leq F(\mathbf{x}_t) - F(\mathbf{x}_{t+1}) + \frac{\|\mathbf{x}_t - \mathbf{y}_t\|^2 - \|\mathbf{x}_{t+1} - \mathbf{y}_t\|^2}{2\eta_t} \\ &\quad + 4\eta_t G^2 + C(p)\eta_t^{p-1} \|\boldsymbol{\epsilon}_t\|^p D^{2-p}. \end{aligned} \quad (32)$$

Since $\mathbf{x}_t, \mathbf{y}_t \in \mathcal{F}_{t-1}$, there is

$$\begin{aligned} \mathbb{E}[\langle \mathbf{g}_t, \mathbf{x}_t - \mathbf{y}_t \rangle] &= \mathbb{E}[\langle \mathbb{E}[\mathbf{g}_t \mid \mathcal{F}_{t-1}], \mathbf{x}_t - \mathbf{y}_t \rangle] = \mathbb{E}[\langle \nabla F(\mathbf{x}_t), \mathbf{x}_t - \mathbf{y}_t \rangle] \\ &\stackrel{(19)}{\geq} \mathbb{E}\left[F(\mathbf{x}_t) - F(\mathbf{y}_t) + \frac{\mu}{2} \|\mathbf{x}_t - \mathbf{y}_t\|^2\right]. \end{aligned}$$

As such, we can take expectations on both sides of (32) and rearrange terms to have

$$\begin{aligned}
 & \mathbb{E} [F(\mathbf{x}_{t+1}) - F(\mathbf{y}_t)] \\
 & \leq \frac{(1 - \mu\eta_t) \mathbb{E} [\|\mathbf{x}_t - \mathbf{y}_t\|^2] - \mathbb{E} [\|\mathbf{x}_{t+1} - \mathbf{y}_t\|^2]}{2\eta_t} + 4\eta_t G^2 + C(p)\eta_t^{p-1} \|\epsilon_t\|^p D^{2-p} \\
 & \leq \frac{(1 - \mu\eta_t) \mathbb{E} [\frac{w_{t-1}}{w_t} \|\mathbf{x}_t - \mathbf{y}_{t-1}\|^2] - \mathbb{E} [\|\mathbf{x}_{t+1} - \mathbf{y}_t\|^2]}{2\eta_t} + 4\eta_t G^2 + C(p)\eta_t^{p-1} \|\epsilon_t\|^p D^{2-p},
 \end{aligned} \tag{33}$$

where the second step is due to $\|\mathbf{x}_t - \mathbf{y}_t\|^2 \leq \left(1 - \frac{w_{t-1}}{w_t}\right) \|\mathbf{x}_t - \mathbf{x}_t\|^2 + \frac{w_{t-1}}{w_t} \|\mathbf{x}_t - \mathbf{y}_{t-1}\|^2 = \frac{w_{t-1}}{w_t} \|\mathbf{x}_t - \mathbf{y}_{t-1}\|^2$ by (29) and the convexity of $\|\mathbf{x}_t - \cdot\|^2$. Multiply both sides of (33) by $w_t\gamma_t$ and sum up from $t = 1$ to T to obtain (note that $\frac{\gamma_t}{\eta_t} = \frac{1}{\prod_{s=2}^t (1 - \mu\eta_s)} = \frac{\gamma_{t+1}(1 - \mu\eta_{t+1})}{\eta_{t+1}}, \forall t \in [T - 1]$)

$$\begin{aligned}
 & \mathbb{E} \left[\sum_{t=1}^T w_t \gamma_t (F(\mathbf{x}_{t+1}) - F(\mathbf{y}_t)) \right] \\
 & \leq \frac{\gamma_1(1 - \mu\eta_1)w_0 \|\mathbf{x}_1 - \mathbf{y}_0\|^2}{2\eta_1} - \frac{\gamma_T \mathbb{E} [w_T \|\mathbf{x}_{T+1} - \mathbf{y}_T\|^2]}{2\eta_T} + \sum_{t=1}^T 4w_t \gamma_t \eta_t G^2 + C(p)w_t \gamma_t \eta_t^{p-1} \sigma^p D^{2-p} \\
 & \leq \frac{(1 - \mu\eta_1)w_0 D^2}{2} + \sum_{t=1}^T 4w_t \gamma_t \eta_t G^2 + C(p)w_t \gamma_t \eta_t^{p-1} \sigma^p D^{2-p}.
 \end{aligned} \tag{34}$$

Now observe that

$$\begin{aligned}
 F(\mathbf{y}_t) - F(\mathbf{x}) & \stackrel{(31)}{\leq} \frac{w_0}{w_t} (F(\mathbf{x}) - F(\mathbf{x})) + \sum_{s=1}^t \frac{w_s - w_{s-1}}{w_t} (F(\mathbf{x}_s) - F(\mathbf{x})) \\
 & = \sum_{s=1}^t \frac{w_s - w_{s-1}}{w_t} (F(\mathbf{x}_s) - F(\mathbf{x})),
 \end{aligned}$$

which implies that

$$\begin{aligned}
 \sum_{t=1}^T w_t \gamma_t (F(\mathbf{y}_t) - F(\mathbf{x})) & \leq \sum_{t=1}^T \sum_{s=1}^t (w_s - w_{s-1}) \gamma_t (F(\mathbf{x}_s) - F(\mathbf{x})) \\
 & = \sum_{t=1}^T (w_t - w_{t-1}) \left(\sum_{s=t}^T \gamma_s \right) (F(\mathbf{x}_t) - F(\mathbf{x})).
 \end{aligned}$$

Thus, we can lower bound the L.H.S. of (34) by

$$\begin{aligned}
 \sum_{t=1}^T w_t \gamma_t (F(\mathbf{x}_{t+1}) - F(\mathbf{y}_t)) &= \sum_{t=1}^T w_t \gamma_t (F(\mathbf{x}_{t+1}) - F(\mathbf{x})) - w_t \gamma_t (F(\mathbf{y}_t) - F(\mathbf{x})) \\
 &\geq w_T \gamma_T (F(\mathbf{x}_{T+1}) - F(\mathbf{x})) - (w_1 - w_0) \left(\sum_{s=1}^T \gamma_s \right) (F(\mathbf{x}_1) - F(\mathbf{x})) \\
 &\quad + \sum_{t=2}^T \left[w_{t-1} \gamma_{t-1} - (w_t - w_{t-1}) \left(\sum_{s=t}^T \gamma_s \right) \right] (F(\mathbf{x}_t) - F(\mathbf{x})) \\
 &= w_T \gamma_T (F(\mathbf{x}_{T+1}) - F(\mathbf{x})), \tag{35}
 \end{aligned}$$

where the last step is due to, for $t \geq 2$,

$$\begin{aligned}
 w_{t-1} \gamma_{t-1} - (w_t - w_{t-1}) \left(\sum_{s=t}^T \gamma_s \right) &\stackrel{(30)}{=} \frac{\gamma_T}{\sum_{s=t-1}^T \gamma_s} \cdot \gamma_{t-1} - \left(\frac{\gamma_T}{\sum_{s=t}^T \gamma_s} - \frac{\gamma_T}{\sum_{s=t-1}^T \gamma_s} \right) \left(\sum_{s=t}^T \gamma_s \right) \\
 &= \frac{\gamma_T}{\sum_{s=t-1}^T \gamma_s} \cdot \gamma_{t-1} - \frac{\gamma_T}{\sum_{s=t-1}^T \gamma_s} \cdot \gamma_{t-1} = 0,
 \end{aligned}$$

and $w_1 \stackrel{(30)}{=} w_0$.

We plug (35) back into (34) and divide both sides by $w_T \gamma_T$ to obtain

$$\begin{aligned}
 \mathbb{E} [F(\mathbf{x}_{T+1}) - F(\mathbf{x})] &\leq \frac{(1 - \mu \eta_1) w_0 D^2}{2 w_T \gamma_T} + \sum_{t=1}^T \frac{4 w_t \gamma_t \eta_t}{w_T \gamma_T} G^2 + \frac{C(\mathbf{p}) w_t \gamma_t \eta_t^{\mathbf{p}-1}}{w_T \gamma_T} \sigma^{\mathbf{p}} D^{2-\mathbf{p}} \\
 &\stackrel{(30)}{\lesssim} \frac{(1 - \mu \eta_1) D^2}{\sum_{t=1}^T \gamma_t} + G^2 \sum_{t=1}^T \frac{\gamma_t \eta_t}{\sum_{s=t}^T \gamma_s} + \sigma^{\mathbf{p}} D^{2-\mathbf{p}} \sum_{t=1}^T \frac{\gamma_t \eta_t^{\mathbf{p}-1}}{\sum_{s=t}^T \gamma_s}.
 \end{aligned}$$

Finally, relabel T by $T - 1$ to complete the proof. $\blacksquare$

Equipped with Theorem 9, we show the following last-iterate convergence rate for SGD/OGD. As far as we know, this is the first and only provable result demonstrating that the last iterate of SGD can converge in heavy-tailed stochastic optimization without gradient clipping. When T is unknown, we only incur an extra logarithmic factor compared to the best possible rate. When T is known, we employ the linear decay rate proposed by [Zamani and Glineur \(2025\)](#) to give an optimal rate.

Corollary 4 *Under Assumption 1 for $\ell_t(\mathbf{x}) = F(\mathbf{x})$:*

- taking $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/\mathbf{p}}}$ in OGD (Algorithm 1), we have

$$\mathbb{E} [F(\mathbf{x}_T) - F(\mathbf{x})] \lesssim \frac{GD(1 + \log T)}{\sqrt{T}} + \frac{\sigma D(1 + \log T)}{T^{1-\frac{1}{\mathbf{p}}}}, \forall \mathbf{x} \in \mathcal{X}.$$

- taking $\eta_t = \frac{D(T-t)}{GT^{3/2}} \wedge \frac{D(T-t)}{\sigma T^{1+1/\mathbf{p}}}$ in OGD (Algorithm 1), we have

$$\mathbb{E} [F(\mathbf{x}_T) - F(\mathbf{x})] \lesssim \frac{GD}{\sqrt{T}} + \frac{\sigma D}{T^{1-\frac{1}{\mathbf{p}}}}, \forall \mathbf{x} \in \mathcal{X}.$$

Proof Since both kinds of stepsize satisfy $\eta_t \leq \frac{1}{\mu} \stackrel{\mu=0}{=} +\infty$, by Theorem 9, we have

$$\begin{aligned} \mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] &\lesssim \frac{(1 - \mu\eta_1)D^2}{\sum_{t=1}^{T-1} \gamma_t} + G^2 \sum_{t=1}^{T-1} \frac{\gamma_t \eta_t}{\sum_{s=t}^{T-1} \gamma_s} + \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\gamma_t \eta_t^{p-1}}{\sum_{s=t}^{T-1} \gamma_s} \\ &= \frac{D^2}{\sum_{t=1}^{T-1} \eta_t} + G^2 \sum_{t=1}^{T-1} \frac{\eta_t^2}{\sum_{s=t}^{T-1} \eta_s} + \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\eta_t^p}{\sum_{s=t}^{T-1} \eta_s}, \end{aligned} \quad (36)$$

where the second step is due $\gamma_t = \frac{\eta_t}{\prod_{s=2}^t (1 - \mu\eta_s)} = \eta_t$ when $\mu = 0$.

For $\eta_t = \frac{D}{G\sqrt{t}} \wedge \frac{D}{\sigma t^{1/p}}$, we observe that by Cauchy-Schwarz inequality

$$(T - t)^2 \leq \left(\sum_{s=t}^{T-1} \frac{1}{\eta_s} \right) \left(\sum_{s=t}^{T-1} \eta_s \right) \Rightarrow \frac{1}{\sum_{s=t}^{T-1} \eta_s} \leq \frac{\sum_{s=t}^{T-1} \frac{1}{\eta_s}}{(T - t)^2}.$$

Thus, there is

$$\mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] \stackrel{(36)}{\lesssim} \frac{D^2}{(T-1)^2} \sum_{t=1}^{T-1} \frac{1}{\eta_t} + G^2 \sum_{t=1}^{T-1} \frac{\eta_t^2 \sum_{s=t}^{T-1} \frac{1}{\eta_s}}{(T-t)^2} + \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\eta_t^p \sum_{s=t}^{T-1} \frac{1}{\eta_s}}{(T-t)^2}. \quad (37)$$

We first bound

$$\sum_{t=1}^{T-1} \frac{1}{\eta_t} = \sum_{t=1}^{T-1} \frac{G\sqrt{t}}{D} \vee \frac{\sigma t^{1/p}}{D} \leq \sum_{t=1}^{T-1} \frac{G\sqrt{t}}{D} + \frac{\sigma t^{1/p}}{D} \lesssim \frac{G}{D} T^{3/2} + \frac{\sigma}{D} T^{1+1/p},$$

which implies

$$\frac{D^2}{(T-1)^2} \sum_{t=1}^{T-1} \frac{1}{\eta_t} \lesssim \frac{GD}{\sqrt{T}} + \frac{\sigma D}{T^{1-\frac{1}{p}}}. \quad (38)$$

Next, we know

$$\begin{aligned} \sum_{t=1}^{T-1} \frac{\eta_t^2 \sum_{s=t}^{T-1} \frac{1}{\eta_s}}{(T-t)^2} &\stackrel{(a)}{\leq} \sum_{t=1}^{T-1} \left[\frac{D}{G} \cdot \frac{\sum_{s=t}^{T-1} \sqrt{s}}{t(T-t)^2} + \frac{\sigma D}{G^2} \cdot \frac{\sum_{s=t}^{T-1} s^{1/p}}{t(T-t)^2} \right] \\ &\stackrel{\text{Fact 1}}{\lesssim} \frac{D(1 + \log T)}{G\sqrt{T}} + \frac{\sigma D(1 + \log T)}{G^2 T^{1-\frac{1}{p}}}, \end{aligned}$$

where (a) is by $\eta_t \leq \frac{D}{G\sqrt{t}}$ and $\frac{1}{\eta_s} \leq \frac{G\sqrt{s}}{D} + \frac{\sigma s^{1/p}}{D}$. Hence, there is

$$G^2 \sum_{t=1}^{T-1} \frac{\eta_t^2 \sum_{s=t}^{T-1} \frac{1}{\eta_s}}{(T-t)^2} \lesssim \frac{GD(1 + \log T)}{\sqrt{T}} + \frac{\sigma D(1 + \log T)}{T^{1-\frac{1}{p}}}. \quad (39)$$

Similarly, we can bound

$$\sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\eta_t^p \sum_{s=t}^{T-1} \frac{1}{\eta_s}}{(T-t)^2} \lesssim \frac{GD(1 + \log T)}{\sqrt{T}} + \frac{\sigma D(1 + \log T)}{T^{1-\frac{1}{p}}}. \quad (40)$$

Finally, we plug (38), (39) and (40) back into (37) to conclude.

For $\eta_t = \frac{D(T-t)}{GT^{3/2}} \wedge \frac{D(T-t)}{\sigma T^{1+1/p}}$, we denote by $\eta_t = \eta(T-t)$ for $\eta \triangleq \frac{D}{GT^{3/2}} \wedge \frac{D}{\sigma T^{1+1/p}}$ to have

$$\begin{aligned} \mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] &\stackrel{(36)}{\lesssim} \frac{D^2}{\eta \sum_{t=1}^{T-1} T-t} + \eta G^2 \sum_{t=1}^{T-1} \frac{(T-t)^2}{\sum_{s=t}^{T-1} T-s} + \eta^{p-1} \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{(T-t)^p}{\sum_{s=t}^{T-1} T-s} \\ &= \frac{2D^2}{\eta T(T-1)} + 2\eta G^2 \sum_{t=1}^{T-1} \frac{T-t}{T-t+1} + 2\eta^{p-1} \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{(T-t)^{p-1}}{T-t+1} \\ &\lesssim \frac{D^2}{\eta T^2} + \eta G^2 T + \eta^{p-1} \sigma^p D^{2-p} T^{p-1} \lesssim \frac{GD}{\sqrt{T}} + \frac{\sigma D}{T^{1-\frac{1}{p}}}. \end{aligned}$$

■

Next, Corollary 5 provides the first last-iterate convergence results in the strongly convex case. The last-iterate rate $\frac{G^2(1+\log T)}{\mu T} + \frac{\sigma^p D^{2-p}(1+\log T)}{\mu^{p-1} T^{p-1}}$ matches the average-iterate rate (up to an extra logarithmic factor) implied by the online-to-batch conversion for Theorem 7.

Corollary 5 *Under Assumption 1 and additionally assuming (19) for $\ell_t(\mathbf{x}) = F(\mathbf{x})$, taking $\eta_t = \frac{1}{\mu t}$ in OGD (Algorithm 1), we have*

$$\mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] \lesssim \frac{G^2(1+\log T)}{\mu T} + \frac{\sigma^p D^{2-p}(1+\log T)}{\mu^{p-1} T^{p-1}}, \forall \mathbf{x} \in \mathcal{X}.$$

Proof Since the stepsize satisfies $\eta_t \leq \frac{1}{\mu}$, by Theorem 9, we have

$$\mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] \lesssim \frac{(1-\mu\eta_1)D^2}{\sum_{t=1}^{T-1} \gamma_t} + G^2 \sum_{t=1}^{T-1} \frac{\gamma_t \eta_t}{\sum_{s=t}^{T-1} \gamma_s} + \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\gamma_t \eta_t^{p-1}}{\sum_{s=t}^{T-1} \gamma_s}. \quad (41)$$

For $\eta_t = \frac{1}{\mu t}$, we can find

$$\gamma_t = \frac{\eta_t}{\prod_{s=2}^t (1-\mu\eta_s)} = \frac{1}{\mu t \prod_{s=2}^t \frac{s-1}{s}} = \frac{1}{\mu} = \eta_1. \quad (42)$$

Hence, there is

$$\begin{aligned} \mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] &\stackrel{(41),(42)}{\lesssim} \frac{(\eta_1^{-1} - \mu)D^2}{T-1} + G^2 \sum_{t=1}^{T-1} \frac{\eta_t}{T-t} + \sigma^p D^{2-p} \sum_{t=1}^{T-1} \frac{\eta_t^{p-1}}{T-t} \\ &= \frac{G^2}{\mu} \sum_{t=1}^{T-1} \frac{1}{(T-t)t} + \frac{\sigma^p D^{2-p}}{\mu^{p-1}} \sum_{t=1}^{T-1} \frac{1}{(T-t)t^{p-1}}. \end{aligned} \quad (43)$$

Given $a \in (0, 1]$, we can bound

$$\sum_{t=1}^{T-1} \frac{1}{(T-t)t^a} = \frac{1}{T} \sum_{t=1}^{T-1} \frac{t^{1-a}}{T-t} + \frac{1}{t^a} \lesssim \frac{1+\log T}{T^a},$$

where the last step is due to

$$\sum_{t=1}^{T-1} \frac{t^{1-a}}{T-t} \lesssim T^{1-a}(1 + \log T) \quad \text{and} \quad \sum_{t=1}^{T-1} \frac{1}{t^a} \lesssim \begin{cases} T^{1-a} & a \in (0, 1) \\ 1 + \log T & a = 1 \end{cases}.$$

Finally, we apply the above inequality to (43) with $a = 1$ and $a = p - 1$ to obtain

$$\mathbb{E}[F(\mathbf{x}_T) - F(\mathbf{x})] \lesssim \frac{G^2(1 + \log T)}{\mu T} + \frac{\sigma^p D^{2-p}(1 + \log T)}{\mu^{p-1} T^{p-1}}.$$

■

Appendix F. Missing Proofs for Applications: Nonsmooth Nonconvex Optimization

F.1. (δ, ϵ) -Stationary Points

Definition 2 (Definition 4 of Cutkosky et al. (2023)) A point $\mathbf{x} \in \mathbb{R}^d$ is a (δ, ϵ) -stationary point of an almost-everywhere differentiable function F if there is a finite subset $S \subset \mathcal{B}^d(\mathbf{x}, \delta)$ such that for $\mathbf{y}$ selected uniformly at random from S , $\mathbb{E}[\mathbf{y}] = \mathbf{x}$ and $\|\mathbb{E}[\nabla F(\mathbf{y})]\| \leq \epsilon$.

The concept of the (δ, ϵ) -stationary point presented here is due to Cutkosky et al. (2023), which is mildly more stringent than the notion of Zhang et al. (2020b), since the latter does not require $\mathbb{E}[\mathbf{y}] = \mathbf{x}$. For more discussions, see Section 2.1 of Cutkosky et al. (2023).

F.2. Proof of Theorem 4

In this section, our ultimate goal is to prove Theorem 4 for the O2NC algorithm, extending Theorem 8 of Cutkosky et al. (2023) from $p = 2$ to any $p \in (1, 2]$. Notably, our new result does not require any modification to the O2NC method, but is obtained only from a more careful analysis, indicating that O2NC is a robust and powerful algorithmic framework.

We begin with Lemma 2, which lies as the cornerstone for establishing the convergence of O2NC.

Lemma 2 (Theorem 7 of Cutkosky et al. (2023)) Under Assumption 2 (only need the second point and the unbiased part in the fourth point), for any sequence of vectors $\mathbf{u}_1, \dots, \mathbf{u}_{KT} \in \mathbb{R}^d$, O2NC (Algorithm 4) guarantees

$$\mathbb{E}[F(\mathbf{y}_{KT})] = F(\mathbf{y}_0) + \mathbb{E}\left[\sum_{n=1}^{KT} \langle \mathbf{g}_n, \mathbf{x}_n - \mathbf{u}_n \rangle\right] + \mathbb{E}\left[\sum_{n=1}^{KT} \langle \mathbf{g}_n, \mathbf{u}_n \rangle\right]. \quad (44)$$

To relate Lemma 2 to the concept of K -shifting regret introduced before (see (9)), suppose now a sequence of vectors $\mathbf{v}_1, \dots, \mathbf{v}_K$ is given, if we set $\mathbf{u}_n = \mathbf{v}_k$ for all $n \in \{(k-1)T+1, \dots, kT\}$ and $k \in [K]$, then the second term on the R.H.S. of (44) can be written as $\mathbb{E}[\mathbf{R}_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)]$, and the third term can be simplified into $\sum_{k=1}^K \mathbb{E}\left[\left\langle \sum_{n=(k-1)T+1}^{kT} \mathbf{g}_n, \mathbf{v}_k \right\rangle\right]$.

Same as [Cutkosky et al. \(2023\)](#), we pick $\mathbf{v}_k \triangleq -D \frac{\sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n)}{\left\| \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\|}$ for some constant $D > 0$, which gives us

$$\begin{aligned} \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \mathbf{g}_n, \mathbf{v}_k \right\rangle \right] &= \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n, \mathbf{v}_k \right\rangle \right] - D \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\| \right] \\ &\leq D \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\| \right] - D \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\| \right]. \end{aligned}$$

If $\boldsymbol{\epsilon}_n$ has a finite variance (i.e., $p = 2$), then like [Cutkosky et al. \(2023\)](#), one can invoke Hölder's inequality and use the fact $\mathbb{E}[\langle \boldsymbol{\epsilon}_m, \boldsymbol{\epsilon}_n \rangle] = 0, \forall m \neq n \in [KT]$ to obtain for any $k \in [K]$,

$$\mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\| \right] \leq \sqrt{\mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\|^2 \right]} = \sqrt{\sum_{n=(k-1)T+1}^{kT} \mathbb{E} [\|\boldsymbol{\epsilon}_n\|^2]} \leq \sigma \sqrt{T}.$$

However, this argument immediately fails when $p < 2$ as $\mathbb{E} [\|\boldsymbol{\epsilon}_n\|^2]$ can be $+\infty$. To handle this potential issue, we require the following Lemma 3.

Lemma 3 (Lemma 4.3 of [Liu and Zhou \(2025\)](#)) *Given a vector-valued martingale difference sequence $\mathbf{w}_1, \dots, \mathbf{w}_T$, there is*

$$\mathbb{E} \left[\left\| \sum_{t=1}^T \mathbf{w}_t \right\| \right] \leq 2\sqrt{2} \mathbb{E} \left[\left(\sum_{t=1}^T \|\mathbf{w}_t\|^p \right)^{\frac{1}{p}} \right], \forall p \in [1, 2].$$

Equipped with Lemmas 2 and 3, we are ready to formally prove Theorem 4, demonstrating that the O2NC framework provably works under heavy-tailed noise.

Proof [Proof of Theorem 4] We invoke Lemma 2 with $\mathbf{u}_n = \mathbf{v}_{\lceil n/T \rceil}, \forall n \in [KT]$ (equivalently, $\mathbf{u}_n = \mathbf{v}_k$ if $n \in \{(k-1)T+1, \dots, kT\}$) and use the definition of K -shifting regret (see (9)) to obtain

$$\mathbb{E} [F(\mathbf{y}_{KT})] = F(\mathbf{y}_0) + \mathbb{E} [\mathbf{R}_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)] + \sum_{k=1}^K \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \mathbf{g}_n, \mathbf{v}_k \right\rangle \right]. \quad (45)$$

Recall that $\mathbf{g}_n = \nabla F(\mathbf{z}_n) + \boldsymbol{\epsilon}_n$, which implies for any $k \in [K]$,

$$\begin{aligned} \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \mathbf{g}_n, \mathbf{v}_k \right\rangle \right] &= \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n, \mathbf{v}_k \right\rangle \right] + \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n), \mathbf{v}_k \right\rangle \right] \\ &\leq \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\| \|\mathbf{v}_k\| \right] + \mathbb{E} \left[\left\langle \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n), \mathbf{v}_k \right\rangle \right] \\ &= D \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\| \right] - D \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\| \right], \quad (46) \end{aligned}$$

where the second step is by Cauchy-Schwarz inequality and the last equation holds due to

$$\mathbf{v}_k = -D \frac{\sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n)}{\left\| \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\|}, \forall k \in [K]. \quad (47)$$

Combine (45) and (46), apply $F(\mathbf{y}_{KT}) \geq F_\star$, and rearrange terms to have

$$\begin{aligned} & \mathbb{E} \left[\sum_{k=1}^K \frac{1}{K} \left\| \frac{1}{T} \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\|^2 \right] \\ & \leq \frac{F(\mathbf{y}_0) - F_\star}{DKT} + \frac{\mathbb{E} [R_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)]}{DKT} + \frac{\sum_{k=1}^K \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\|^2 \right]}{KT}. \end{aligned} \quad (48)$$

For any fixed $k \in [K]$, we apply Lemma 3 with $\mathbf{w}_t = \boldsymbol{\epsilon}_{(k-1)T+t}$, $\forall t \in [T]$ to know

$$\begin{aligned} \mathbb{E} \left[\left\| \sum_{n=(k-1)T+1}^{kT} \boldsymbol{\epsilon}_n \right\|^2 \right] & \leq 2\sqrt{2} \mathbb{E} \left[\left(\sum_{n=(k-1)T+1}^{kT} \|\boldsymbol{\epsilon}_n\|^p \right)^{\frac{1}{p}} \right] \\ & \leq 2\sqrt{2} \left(\sum_{n=(k-1)T+1}^{kT} \mathbb{E} [\|\boldsymbol{\epsilon}_n\|^p] \right)^{\frac{1}{p}} \leq 2\sqrt{2} \sigma T^{\frac{1}{p}}, \end{aligned} \quad (49)$$

where the second step is by Hölder's inequality (note that $p > 1$). Finally, we conclude the proof after plugging (49) back into (48). $\blacksquare$

F.3. Proof of Theorem 5

Proof By Theorem 4, there is

$$\mathbb{E} \left[\sum_{k=1}^K \frac{1}{K} \left\| \frac{1}{T} \sum_{n=(k-1)T+1}^{kT} \nabla F(\mathbf{z}_n) \right\|^2 \right] \lesssim \frac{F(\mathbf{y}_0) - F_\star}{DKT} + \frac{\mathbb{E} [R_T^A(\mathbf{v}_1, \dots, \mathbf{v}_K)]}{DKT} + \frac{\sigma}{T^{1-\frac{1}{p}}}. \quad (50)$$

Note that $\mathbf{A}$ has the domain $\mathcal{X} = \mathcal{B}^d(D)$ and $s_n \sim \text{Uniform}[0, 1]$. Thus, for any $n \in [KT]$,

$$\|\mathbf{x}_n\| \leq D \quad \text{and} \quad s_n \in [0, 1]. \quad (51)$$

We first lower bound the L.H.S. of (50). Given $k \in [K]$, for any $m < n \in \{(k-1)T+1, \dots, kT\}$, observe that

$$\begin{aligned} \|\mathbf{z}_n - \mathbf{z}_m\| & = \|\mathbf{y}_{n-1} + s_n \mathbf{x}_n - \mathbf{y}_{m-1} - s_m \mathbf{x}_m\| = \left\| s_n \mathbf{x}_n - s_m \mathbf{x}_m + \sum_{i=m}^{n-1} \mathbf{x}_i \right\| \\ & \leq s_n \|\mathbf{x}_n\| + (1 - s_m) \|\mathbf{x}_m\| + \sum_{i=m+1}^{n-1} \|\mathbf{x}_i\| \stackrel{(51)}{\leq} (n - m + 1) D \leq DT. \end{aligned}$$

Recall that $\bar{z}_k = \frac{1}{T} \sum_{n=(k-1)T+1}^{kT} z_n$ and $D = \delta/T$ now, then the above inequality implies

$$\|z_n - \bar{z}_k\| \leq DT = \delta, \forall n \in \{(k-1)T+1, \dots, kT\}, \quad (52)$$

which means

$$z_n \in \mathcal{B}^d(\bar{z}_k, \delta), \forall n \in \{(k-1)T+1, \dots, kT\}.$$

By the definition of $\|\nabla F(\bar{z}_k)\|_\delta$ (see Definition 1), there is

$$\|\nabla F(\bar{z}_k)\|_\delta \leq \left\| \frac{1}{T} \sum_{n=(k-1)T+1}^{kT} \nabla F(z_n) \right\|. \quad (53)$$

Next, we upper bound the R.H.S. of (50). By the definition of K -shifting regret (see (9)), there is

$$\mathbb{E} \left[R_T^A(v_1, \dots, v_K) \right] = \sum_{k=1}^K \mathbb{E} \left[\sum_{n=(k-1)T+1}^{kT} \langle g_n, x_n - v_k \rangle \right].$$

Note that we reset the stepsize in A after every T iterations and $v_k \in \mathcal{B}^d(D)$ by its definition (see (47)). Then for any $A \in \{\text{OGD}, \text{DA}, \text{AdaGrad}\}$, we can invoke its regret bound³ (i.e., Theorems 1, 2 and 3) to obtain

$$\mathbb{E} \left[\sum_{n=(k-1)T+1}^{kT} \langle g_n, x_n - v_k \rangle \right] \lesssim GD\sqrt{T} + \sigma DT^{1/p}, \forall k \in [K],$$

which implies

$$\mathbb{E} \left[R_T^A(v_1, \dots, v_K) \right] \lesssim GDK\sqrt{T} + \sigma DK T^{1/p}. \quad (54)$$

Finally, we plug (53) and (54) back into (50), then use $D = \delta/T$ and $\Delta = F(y_0) - F_\star$ to have

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{z}_k)\|_\delta \right] \lesssim \frac{\Delta}{\delta K} + \frac{G}{\sqrt{T}} + \frac{\sigma}{T^{1-\frac{1}{p}}}.$$

■

F.4. Proof of Corollary 3

Proof Recall that we pick

$$K = \left\lfloor \frac{N}{T} \right\rfloor \quad \text{and} \quad T = \left\lfloor \frac{N}{2} \right\rfloor \wedge \left(\left\lceil \left(\frac{\delta G N}{\Delta} \right)^{\frac{2}{3}} \right\rceil \vee \left\lceil \left(\frac{\delta \sigma N}{\Delta} \right)^{\frac{p}{2p-1}} \right\rceil \right),$$

3. A minor point here is that the current function $\ell_n(x) = \langle g_n, x \rangle$ does not entirely fit Assumption 1. We clarify that one does not need to worry about it, since all results proved in Section 3 hold under this change. For example, in the proof of Theorem 1, we can safely replace the L.H.S. of (17) with $\mathbb{E} \left[\sum_{t=1}^T \langle g_t, x_t - x \rangle \right]$.

where $\Delta = F(\mathbf{y}_0) - F_\star$. We invoke Theorem 5 and use $KT \geq N/4$ (see Fact 2) to obtain

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{\mathbf{z}}_k)\|_\delta \right] \lesssim \frac{\Delta T}{\delta N} + \frac{G}{\sqrt{T}} + \frac{\sigma}{T^{1-\frac{1}{p}}}.$$

By the definition of T , we know

$$\frac{\Delta T}{\delta N} \lesssim \frac{\Delta}{\delta N} \left[1 + \left(\frac{\delta G N}{\Delta} \right)^{\frac{2}{3}} + \left(\frac{\delta \sigma N}{\Delta} \right)^{\frac{p}{2p-1}} \right] = \frac{\Delta}{\delta N} + \frac{G^{\frac{2}{3}} \Delta^{\frac{1}{3}}}{(\delta N)^{\frac{1}{3}}} + \frac{\sigma^{\frac{p}{2p-1}} \Delta^{\frac{p-1}{2p-1}}}{(\delta N)^{\frac{p-1}{2p-1}}},$$

and

$$\frac{G}{\sqrt{T}} \lesssim \frac{G}{\sqrt{N}} + \frac{G^{\frac{2}{3}} \Delta^{\frac{1}{3}}}{(\delta N)^{\frac{1}{3}}}, \quad \frac{\sigma}{T^{1-\frac{1}{p}}} \lesssim \frac{\sigma}{N^{1-\frac{1}{p}}} + \frac{\sigma^{\frac{p}{2p-1}} \Delta^{\frac{p-1}{2p-1}}}{(\delta N)^{\frac{p-1}{2p-1}}}.$$

Therefore, there is

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{\mathbf{z}}_k)\|_\delta \right] \lesssim \frac{G}{\sqrt{N}} + \frac{\sigma}{N^{1-\frac{1}{p}}} + \frac{\Delta}{\delta N} + \frac{G^{\frac{2}{3}} \Delta^{\frac{1}{3}}}{(\delta N)^{\frac{1}{3}}} + \frac{\sigma^{\frac{p}{2p-1}} \Delta^{\frac{p-1}{2p-1}}}{(\delta N)^{\frac{p-1}{2p-1}}}.$$

■

F.5. Extension to the Case of Unknown Problem-Dependent Parameters

In Corollary 6, we show how to set K and T when all problem-dependent parameters are unknown. It is particularly meaningful for AdaGrad. As in that case, the rate is achieved without knowing any problem-dependent parameter. This kind of result is the first to appear for nonsmooth nonconvex optimization with heavy tails. However, the rate is not as good as Corollary 3. It is currently unclear whether the same bound $1/(\delta N)^{\frac{p-1}{2p-1}}$ as in Corollary 3 can be obtained when no information about the problem is known.

Corollary 6 *Under the same setting of Theorem 5, suppose we have $N \geq 2$ stochastic gradient budgets, taking $K = \lfloor N/T \rfloor$ and $T = \lceil N/2 \rceil \wedge \left\lceil (\delta N)^{\frac{2}{3}} \right\rceil$, we have*

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{\mathbf{z}}_k)\|_\delta \right] \lesssim \frac{\Delta}{(\delta N) \wedge (\delta N)^{\frac{1}{3}}} + \frac{G}{\sqrt{N} \wedge (\delta N)^{\frac{1}{3}}} + \frac{\sigma}{N^{1-\frac{1}{p}} \wedge (\delta N)^{\frac{2(p-1)}{3p}}}.$$

Proof We invoke Theorem 5 and use $KT \geq N/4$ (see Fact 2) to obtain

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{\mathbf{z}}_k)\|_\delta \right] \lesssim \frac{\Delta T}{\delta N} + \frac{G}{\sqrt{T}} + \frac{\sigma}{T^{1-\frac{1}{p}}}.$$

By the definition of T , we know

$$\frac{\Delta T}{\delta N} \lesssim \frac{\Delta}{\delta N} \left[1 + (\delta N)^{\frac{2}{3}} \right] \lesssim \frac{\Delta}{(\delta N) \wedge (\delta N)^{\frac{1}{3}}}.$$

and

$$\begin{aligned}\frac{G}{\sqrt{T}} &\lesssim \frac{G}{\sqrt{N}} + \frac{G}{(\delta N)^{\frac{1}{3}}} \lesssim \frac{G}{\sqrt{N} \wedge (\delta N)^{\frac{1}{3}}}, \\ \frac{\sigma}{T^{1-\frac{1}{p}}} &\lesssim \frac{\sigma}{N^{1-\frac{1}{p}}} + \frac{\sigma}{(\delta N)^{\frac{2(p-1)}{3p}}} \lesssim \frac{\sigma}{N^{1-\frac{1}{p}} \wedge (\delta N)^{\frac{2(p-1)}{3p}}}.\end{aligned}$$

Therefore, there is

$$\mathbb{E} \left[\frac{1}{K} \sum_{k=1}^K \|\nabla F(\bar{z}_k)\|_\delta \right] \lesssim \frac{\Delta}{(\delta N) \wedge (\delta N)^{\frac{1}{3}}} + \frac{G}{\sqrt{N} \wedge (\delta N)^{\frac{1}{3}}} + \frac{\sigma}{N^{1-\frac{1}{p}} \wedge (\delta N)^{\frac{2(p-1)}{3p}}}.$$

■

F.6. Proof of Theorem 6

In this section, our ultimate goal is to prove Theorem 6. The analysis follows the framework first established in Arjevani et al. (2023) and later developed by Cutkosky et al. (2023) but with some necessary (though minor) variation to make it compatible with heavy-tailed noise. In the following, we will slightly abuse the notation A to denote any possible randomized first-order algorithm instead of an online learning algorithm.

F.6.1. BASIC DEFINITIONS

To begin with, we introduce some basic definitions given in Arjevani et al. (2023).

Definition 3 (stochastic first-order oracle) *Given a differentiable function $F : \mathbb{R}^d \rightarrow \mathbb{R}$, a tuple $(\mathbf{g}, \mathcal{R}, \mathbb{P}_r)$ is called a stochastic first-order oracle of F if $\mathbb{P}_r$ is a probability distribution on the measurable space $\mathcal{R}$ and $\mathbf{g} : \mathbb{R}^d \times \mathcal{R} \rightarrow \mathbb{R}^d$ satisfies $\mathbb{E}_{r \sim \mathbb{P}_r} [\mathbf{g}(\mathbf{x}, r)] = \nabla F(\mathbf{x}), \forall \mathbf{x} \in \mathbb{R}^d$.*

Remark 12 *When the context is clear, we will omit F , $\mathcal{R}$ and $\mathbb{P}_r$, and simply call $\mathbf{g}$ as a stochastic first-order oracle.*

Definition 4 (randomized algorithm) *A randomized algorithm A consists of a probability distribution $\mathbb{P}_s$ over a measurable space $\mathcal{S}$ and a sequence of measurable mappings $A_t, \forall t \in \mathbb{N}$ such that every A_t takes a common random seed $s \in \mathcal{S}$ and the first $t - 1$ oracle responses to produce the t -th query. Concretely, given a differentiable function F equipped with a stochastic first-order oracle $(\mathbf{g}, \mathcal{R}, \mathbb{P}_r)$, the sequence $\mathbf{x}_t, \forall t \in \mathbb{N}$ produced by A to optimize F is recursively defined as*

$$\mathbf{x}_t = A_t(s, \mathbf{g}(\mathbf{x}_{t-1}, r_{t-1}), \dots, \mathbf{g}(\mathbf{x}_1, r_1)), \forall t \in \mathbb{N},$$

where $s \sim \mathbb{P}_s$ is drawn a single time at the beginning of the algorithm and $r_t \sim \mathbb{P}_r, \forall t \in \mathbb{N}$ is a sequence of i.i.d. random variables. Moreover, A_{rand} denotes the set containing all randomized algorithms.

Next, we require a useful concept named probability- q zero-chain. Before formally stating what it is, we need some notations. Given $\mathbf{x} \in \mathbb{R}^d$ and $\alpha \in [0, 1]$, $\text{prog}_\alpha(\mathbf{x})$ denotes the largest index whose entry is α -far from 0, i.e.,

$$\text{prog}_\alpha(\mathbf{x}) \triangleq \max \{i \in [d] : |\mathbf{x}[i]| > \alpha\} \quad \text{where} \quad \max \emptyset \triangleq 0.$$

In addition, for any $j \in \{0\} \cup \mathbb{N}$, let $\mathbf{x}_{\leq j}[i] \triangleq \mathbf{x}[i] \mathbb{1}[i \leq j]$, $\forall i \in [d]$ be the truncated version of $\mathbf{x}$. Now we are ready to provide the definition of the probability- q zero-chain.

Definition 5 (probability- q zero-chain) A stochastic first-order oracle $(\mathbf{g}, \mathcal{R}, \mathbb{P}_r)$ is called a probability- q zero-chain if and only if

$$\mathbb{P}_r \left[\forall \mathbf{x} \in \mathbb{R}^d : \text{prog}_0(\mathbf{g}(\mathbf{x}, r)) \leq \text{prog}_{1/4}(\mathbf{x}) \text{ and } \mathbf{g}(\mathbf{x}, r) = \mathbf{g}(\mathbf{x}_{\leq \text{prog}_{1/4}(\mathbf{x})}, r) \right] \geq 1 - q,$$

and

$$\mathbb{P}_r \left[\forall \mathbf{x} \in \mathbb{R}^d : \text{prog}_0(\mathbf{g}(\mathbf{x}, r)) \leq \text{prog}_{1/4}(\mathbf{x}) + 1 \text{ and } \mathbf{g}(\mathbf{x}, r) = \mathbf{g}(\mathbf{x}_{\leq \text{prog}_{1/4}(\mathbf{x})+1}, r) \right] = 1.$$

F.6.2. USEFUL EXISTING RESULTS

In this part, we list some useful existing results from [Arjevani et al. \(2023\)](#) and [Cutkosky et al. \(2023\)](#).

Given $d \geq T \in \mathbb{N}$ and a differentiable function $F_T : \mathbb{R}^T \rightarrow \mathbb{R}$ with a stochastic first-order oracle $\mathbf{g}_T$ that is a probability- q zero-chain, their rotated variants parametrized by a matrix $U \in \text{Ortho}(d, T) \triangleq \{U \in \mathbb{R}^{d \times T} : U^\top U = I_T\}$ (where I_T is the T -dimensional identity matrix) are defined as

$$\bar{F}_{T,U}(\mathbf{x}) \triangleq F_T(U^\top \mathbf{x}) \quad \text{and} \quad \bar{\mathbf{g}}_{T,U}(\mathbf{x}, r) \triangleq U \mathbf{g}_T(U^\top \mathbf{x}, r). \quad (55)$$

Clearly, $\bar{\mathbf{g}}_{T,U}$ is a stochastic first-order oracle of $\bar{F}_{T,U}$. In addition, we emphasize that the input variable $\mathbf{x}$ is from $\mathbb{R}^d$ instead of $\mathbb{R}^T$ now.

With $\bar{F}_{T,U}$ and $\bar{\mathbf{g}}_{T,U}$, we state the following lemma due to [Arjevani et al. \(2023\)](#).

Lemma 4 (Lemma 5 of [Arjevani et al. \(2023\)](#)) Given $q, \iota \in (0, 1)$, $R > 0$, $T \in \mathbb{N}$, $d \in \mathbb{N}$ satisfying $d \geq T + 32R^2 \log \frac{2T^2}{q\iota}$ and a randomized algorithm $A \in \mathcal{A}_{\text{rand}}$, suppose the output of A always has a norm bounded by R , let $\mathbf{x}_t, \forall t \in \mathbb{N}$ be the trajectory produced by applying A to optimize $\bar{F}_{T,U}$ interacting with the stochastic first-order oracle $\bar{\mathbf{g}}_{T,U}$, where U is drawn from $\text{Ortho}(d, T)$ uniformly, then there is

$$\Pr \left[\text{prog}_{1/4} \left(U^\top \mathbf{x}_t \right) < T, \forall t \leq \frac{T - \log(2/\iota)}{2q} \right] \geq 1 - \iota.$$

Here $\Pr$ takes into account all randomness over A , $\mathbf{g}_T$, and U .

Note that Lemma 4 can be only applied to a randomized algorithm with bounded outputs. To overcome this issue, we need the following variants of $\bar{F}_{T,U}$ and $\bar{\mathbf{g}}_{T,U}$ introduced by [Cutkosky et al. \(2023\)](#):

$$\hat{F}_{T,U}(\mathbf{x}) \triangleq \bar{F}_{T,U}(\boldsymbol{\rho}_{R,d}(\mathbf{x})) + \eta \mathbf{x}^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}), \quad (56)$$

$$\hat{\mathbf{g}}_{T,U}(\mathbf{x}, r) \triangleq \mathcal{J}_{R,d}(\mathbf{x})^\top \bar{\mathbf{g}}_{T,U}(\boldsymbol{\rho}_{R,d}(\mathbf{x}), r) + \eta \nabla(\mathbf{x}^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})), \quad (57)$$

where $\rho_{R,d}(\mathbf{x}) \triangleq \frac{\mathbf{x}}{\sqrt{1+\|\mathbf{x}\|^2/R^2}}$ is a bijection from $\mathbb{R}^d$ to $\text{int}\mathcal{B}^d(R)$ ⁴, $\mathcal{J}_{R,d}$ denotes the Jacobian of $\rho_{R,d}$, and $R, \eta > 0$ are two constants being determined later.

Before moving on, we state some useful properties of $\rho_{R,d}$ here.

Lemma 5 (Lemma 13 of Arjevani et al. (2023) and Proposition 29 of Cutkosky et al. (2023)) *Let $\|\cdot\|_{\text{op}}$ denotes the operator norm, then for $\rho_{R,d} : \mathbb{R}^d \rightarrow \mathbb{R}^d, \mathbf{x} \mapsto \frac{\mathbf{x}}{\sqrt{1+\|\mathbf{x}\|^2/R^2}}$, there are*

$$1. \nabla(\mathbf{x}^\top \rho_{R,d}(\mathbf{x})) = \left(2 - \|\rho_{R,d}(\mathbf{x})\|^2 / R^2\right) \rho_{R,d}(\mathbf{x}),$$

$$2. \|\rho_{R,d}(\mathbf{x}) - \rho_{R,d}(\mathbf{y})\| \leq \|\mathbf{x} - \mathbf{y}\|,$$

$$3. \mathcal{J}_{R,d}(\mathbf{x}) = \frac{I_d - \rho_{R,d}(\mathbf{x})\rho_{R,d}(\mathbf{x})^\top / R^2}{\sqrt{1+\|\mathbf{x}\|^2/R^2}},$$

$$4. \|\mathcal{J}_{R,d}(\mathbf{x})\|_{\text{op}} = \frac{1}{\sqrt{1+\|\mathbf{x}\|^2/R^2}} \leq 1,$$

$$5. \|\mathcal{J}_{R,d}(\mathbf{x}) - \mathcal{J}_{R,d}(\mathbf{y})\|_{\text{op}} \leq \frac{3}{R} \|\mathbf{x} - \mathbf{y}\|.$$

We then can combine (55), (56), (57) and Lemma 5 to have

$$\hat{F}_{T,U}(\mathbf{x}) = \bar{F}_{T,U}(\rho_{R,d}(\mathbf{x})) + \eta \mathbf{x}^\top \rho_{R,d}(\mathbf{x}) = F_T(U^\top \rho_{R,d}(\mathbf{x})) + \eta \mathbf{x}^\top \rho_{R,d}(\mathbf{x}), \quad (58)$$

$$\nabla \hat{F}_{T,U}(\mathbf{x}) = \mathcal{J}_{R,d}(\mathbf{x})^\top \nabla \bar{F}_{T,U}(\rho_{R,d}(\mathbf{x})) + \eta \left(2 - \|\rho_{R,d}(\mathbf{x})\|^2 / R^2\right) \rho_{R,d}(\mathbf{x}) \quad (59)$$

$$= \mathcal{J}_{R,d}(\mathbf{x})^\top U \nabla F_T(U^\top \rho_{R,d}(\mathbf{x})) + \eta \left(2 - \|\rho_{R,d}(\mathbf{x})\|^2 / R^2\right) \rho_{R,d}(\mathbf{x}), \quad (60)$$

$$\hat{\mathbf{g}}_{T,U}(\mathbf{x}, r) = \mathcal{J}_{R,d}(\mathbf{x})^\top \bar{\mathbf{g}}_{T,U}(\rho_{R,d}(\mathbf{x}), r) + \eta \left(2 - \|\rho_{R,d}(\mathbf{x})\|^2 / R^2\right) \rho_{R,d}(\mathbf{x}) \quad (61)$$

$$= \mathcal{J}_{R,d}(\mathbf{x})^\top U \mathbf{g}_T(U^\top \rho_{R,d}(\mathbf{x}), r) + \eta \left(2 - \|\rho_{R,d}(\mathbf{x})\|^2 / R^2\right) \rho_{R,d}(\mathbf{x}). \quad (62)$$

Now we are ready to give the following Lemma 6, which is almost identical to Lemma 25 in Cutkosky et al. (2023) by differing up to numerical constants. Since Lemma 6 is particularly important, we therefore provide its full proof here.

Lemma 6 (Lemma 25 of Cutkosky et al. (2023)) *For $T \in \mathbb{N}$, suppose that F_T satisfies the following two requirements:*

1. *For all $\mathbf{x} \in \mathbb{R}^T$, $\|\nabla F_T(\mathbf{x})\| \leq \gamma\sqrt{T}$ where $\gamma > 0$ is a constant satisfying $\frac{1}{\sqrt{5}} - \frac{39}{140} \geq \frac{10}{63} + \frac{5\sqrt{5}}{126\gamma} \Leftrightarrow \gamma \geq \frac{50\sqrt{5}}{252\sqrt{5}-551} \approx 8.95$.*
2. *For all $\mathbf{x} \in \mathbb{R}^T$, if $\text{prog}_1(\mathbf{x}) < T$, then $\|\nabla F_T(\mathbf{x})\| \geq \|\nabla F_T(\mathbf{x})\| [\text{prog}_1(\mathbf{x}) + 1] > 1$.*

Given $q, \iota \in (0, 1)$, $\eta = \frac{1}{\sqrt{5}} - \frac{39}{140}$, $R = 7\gamma\sqrt{T}$, $d \in \mathbb{N}$ satisfying $d \geq T + 32R^2 \log \frac{2T^2}{q\iota}$ and a randomized algorithm $\mathbf{A} \in \mathbf{A}_{\text{rand}}$, let $\mathbf{x}_t, \forall t \in \mathbb{N}$ be the trajectory produced by applying $\mathbf{A}$ to optimize $\hat{F}_{T,U}$ interacting with the stochastic first-order oracle $\hat{\mathbf{g}}_{T,U}$, where U is drawn from $\text{Orth}(d, T)$ uniformly, then there is

$$\Pr \left[\left\| \nabla \hat{F}_{T,U}(\mathbf{x}_t) \right\| \geq \frac{1}{2}, \forall t \leq \frac{T - \log(2/\iota)}{2q} \right] \geq 1 - \iota.$$

Here $\Pr$ takes into account all randomness over $\mathbf{A}$, $\mathbf{g}_T$ and U .

4. As one can check, $\rho_{R,d}^{-1}(\mathbf{x}) : \text{int}\mathcal{B}^d(R) \rightarrow \mathbb{R}^d$ exists and equals $\frac{\mathbf{x}}{\sqrt{1-\|\mathbf{x}\|^2/R^2}}$.

Proof By definition of randomized algorithms (see Definition 4), we have

$$\mathbf{x}_t = \mathbf{A}_t(s, \hat{\mathbf{g}}_{T,U}(\mathbf{x}_{t-1}, r_{t-1}), \dots, \hat{\mathbf{g}}_{T,U}(\mathbf{x}_1, r_1)), \forall t \in \mathbb{N}. \quad (63)$$

Now let us consider another sequence

$$\mathbf{y}_t \triangleq \boldsymbol{\rho}_{R,d}(\mathbf{x}_t) = \frac{\mathbf{x}_t}{\sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} \in \text{int}\mathcal{B}^d(R), \forall t \in \mathbb{N}. \quad (64)$$

Note that

$$\begin{aligned} \hat{\mathbf{g}}_{T,U}(\mathbf{x}_t, r_t) &\stackrel{(61)}{=} \mathcal{J}_{R,d}(\mathbf{x}_t)^\top \bar{\mathbf{g}}_{T,U}(\boldsymbol{\rho}_{R,d}(\mathbf{x}_t), r_t) + \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x}_t)\|^2 / R^2\right) \boldsymbol{\rho}_{R,d}(\mathbf{x}_t) \\ &\stackrel{(64)}{=} \mathcal{J}_{R,d}(\boldsymbol{\rho}_{R,d}^{-1}(\mathbf{y}_t))^\top \bar{\mathbf{g}}_{T,U}(\mathbf{y}_t, r_t) + \eta \left(2 - \|\mathbf{y}_t\|^2 / R^2\right) \mathbf{y}_t \\ &= \mathcal{G}(\mathbf{y}_t, \bar{\mathbf{g}}_{T,U}(\mathbf{y}_t, r_t)), \end{aligned} \quad (65)$$

where $\mathcal{G}(\mathbf{u}, \mathbf{v}) : \text{int}\mathcal{B}^d(R) \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ is a measurable mapping defined as

$$\mathcal{G}(\mathbf{u}, \mathbf{v}) \triangleq \mathcal{J}_{R,d}(\boldsymbol{\rho}_{R,d}^{-1}(\mathbf{u}))^\top \mathbf{v} + \eta \left(2 - \|\mathbf{u}\|^2 / R^2\right) \mathbf{u}.$$

Thus, we can write

$$\mathbf{y}_t \stackrel{(63),(64),(65)}{=} \boldsymbol{\rho}_{R,d} \circ \mathbf{A}_t(s, \mathcal{G}(\mathbf{y}_{t-1}, \bar{\mathbf{g}}_{T,U}(\mathbf{y}_{t-1}, r_{t-1})), \dots, \mathcal{G}(\mathbf{y}_1, \bar{\mathbf{g}}_{T,U}(\mathbf{y}_1, r_1))), \forall t \in \mathbb{N}.$$

By a simple induction, there exists a sequence of measurable mappings $\mathbf{A}_t^{\mathbf{y}}, \forall t \in \mathbb{N}$ such that

$$\mathbf{y}_t = \mathbf{A}_t^{\mathbf{y}}(s, \bar{\mathbf{g}}_{T,U}(\mathbf{y}_{t-1}, r_{t-1}), \dots, \bar{\mathbf{g}}_{T,U}(\mathbf{y}_1, r_1)), \forall t \in \mathbb{N}.$$

The above reformulation implies $\mathbf{y}_t, \forall t \in \mathbb{N}$ can be viewed as a sequence produced by a randomized algorithm $\mathbf{A}^{\mathbf{y}} \in \mathbf{A}_{\text{rand}}$ interacting with stochastic first-order oracle $\bar{\mathbf{g}}_{T,U}$. Note that $d \geq T + 32R^2 \log \frac{2T^2}{q\epsilon}$ and $\|\mathbf{y}_t\| \stackrel{(64)}{\leq} R$, we hence have by Lemma 4,

$$\Pr \left[\text{prog}_{1/4} \left(U^\top \mathbf{y}_t \right) < T, \forall t \leq \frac{T - \log(2/\iota)}{2q} \right] \geq 1 - \iota. \quad (66)$$

Now we recall

$$\begin{aligned} \nabla \hat{F}_{T,U}(\mathbf{x}_t) &\stackrel{(59)}{=} \mathcal{J}_{R,d}(\mathbf{x}_t)^\top \nabla \bar{F}_{T,U}(\boldsymbol{\rho}_{R,d}(\mathbf{x}_t)) + \eta \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x}_t)\|^2 / R^2\right) \boldsymbol{\rho}_{R,d}(\mathbf{x}_t) \\ &\stackrel{(64)}{=} \mathcal{J}_{R,d}(\mathbf{x}_t)^\top \nabla \bar{F}_{T,U}(\mathbf{y}_t) + \eta \left(2 - \|\mathbf{y}_t\|^2 / R^2\right) \mathbf{y}_t. \end{aligned} \quad (67)$$

Moreover, by the first requirement on F_T , there is

$$\|\nabla \bar{F}_{T,U}(\mathbf{y}_t)\| \stackrel{(55)}{=} \left\| U \nabla F_T(U^\top \mathbf{y}_t) \right\| \stackrel{U^\top U = I_T}{=} \left\| \nabla F_T(U^\top \mathbf{y}_t) \right\| \leq \gamma \sqrt{T}. \quad (68)$$

We fix $t \leq \frac{T - \log(2/\iota)}{2q}$ and consider the following two cases:

Case 1. $\|\mathbf{x}_t\| \geq \frac{R}{2}$. In this case, we first have

$$\left\| \nabla \hat{F}_{T,U}(\mathbf{x}_t) \right\| \stackrel{(67)}{\geq} \eta \left(2 - \|\mathbf{y}_t\|^2 / R^2 \right) \|\mathbf{y}_t\| - \left\| \mathcal{J}_{R,d}(\mathbf{x}_t)^\top \nabla \bar{F}_{T,U}(\mathbf{y}_t) \right\|.$$

Note that $\frac{\|\mathbf{y}_t\|}{R} \stackrel{(64)}{=} \frac{\|\mathbf{x}_t\|/R}{\sqrt{1+\|\mathbf{x}_t\|^2/R^2}} \in \left[\frac{1}{\sqrt{5}}, 1 \right]$ when $\|\mathbf{x}_t\| \geq \frac{R}{2}$, implying

$$\eta \left(2 - \|\mathbf{y}_t\|^2 / R^2 \right) \|\mathbf{y}_t\| \geq \eta R \min_{\frac{1}{\sqrt{5}} \leq a \leq 1} (2 - a^2)a = \frac{9\eta R}{5\sqrt{5}}.$$

In addition, there is

$$\begin{aligned} \left\| \mathcal{J}_{R,d}(\mathbf{x}_t)^\top \nabla \bar{F}_{T,U}(\mathbf{y}_t) \right\| &\leq \|\mathcal{J}_{R,d}(\mathbf{x}_t)\|_{\text{op}} \|\nabla \bar{F}_{T,U}(\mathbf{y}_t)\| \stackrel{\text{Lemma 5}}{=} \frac{\|\nabla \bar{F}_{T,U}(\mathbf{y}_t)\|}{\sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} \\ &\stackrel{\|\mathbf{x}_t\| \geq \frac{R}{2}}{\leq} \frac{2 \|\nabla \bar{F}_{T,U}(\mathbf{y}_t)\|}{\sqrt{5}} \stackrel{(68)}{\leq} \frac{2\gamma\sqrt{T}}{\sqrt{5}}. \end{aligned}$$

As such, by our choices of η and R ,

$$\left\| \nabla \hat{F}_{T,U}(\mathbf{x}_t) \right\| \geq \frac{9\eta R}{5\sqrt{5}} - \frac{2\gamma\sqrt{T}}{\sqrt{5}} \geq \frac{1}{2}.$$

Case 2. $\|\mathbf{x}_t\| < \frac{R}{2}$. In this case, we introduce

$$j_t \triangleq \text{prog}_1 \left(U^\top \mathbf{y}_t \right) + 1 \leq \text{prog}_{1/4} \left(U^\top \mathbf{y}_t \right) + 1 \stackrel{(66)}{\leq} T.$$

Let $\mathbf{u}_{j_t} \in \mathbb{R}^d$ denotes the j_t -th column of U . By the definition of j_t , there is

$$|\langle \mathbf{u}_{j_t}, \mathbf{y}_t \rangle| = \left| \left(U^\top \mathbf{y}_t \right) [j_t] \right| < 1. \quad (69)$$

In addition, we have

$$|\langle \mathbf{u}_{j_t}, \nabla \bar{F}_{T,U}(\mathbf{y}_t) \rangle| = \left| \left(U^\top \nabla \bar{F}_{T,U}(\mathbf{y}_t) \right) [j_t] \right| \stackrel{(55), U^\top U = I_T}{=} \left| \nabla F_T(U^\top \mathbf{y}_t) [j_t] \right| > 1, \quad (70)$$

where the last step is by the second requirement on F_T . Now we compute

$$\begin{aligned} \left\| \nabla \hat{F}_{T,U}(\mathbf{x}_t) \right\| &\stackrel{\|\mathbf{u}_{j_t}\|=1}{\geq} \left| \langle \mathbf{u}_{j_t}, \nabla \hat{F}_{T,U}(\mathbf{x}_t) \rangle \right| \\ &\stackrel{(67)}{=} \left| \langle \mathbf{u}_{j_t}, \mathcal{J}_{R,d}(\mathbf{x}_t)^\top \nabla \bar{F}_{T,U}(\mathbf{y}_t) \rangle + \eta \left(2 - \|\mathbf{y}_t\|^2 / R^2 \right) \langle \mathbf{u}_{j_t}, \mathbf{y}_t \rangle \right| \\ &\stackrel{\text{Lemma 5}}{=} \left| \frac{\langle \mathbf{u}_{j_t}, \nabla \bar{F}_{T,U}(\mathbf{y}_t) \rangle}{\sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} - \frac{\langle \mathbf{u}_{j_t}, \mathbf{y}_t \rangle \langle \mathbf{y}_t, \nabla \bar{F}_{T,U}(\mathbf{y}_t) \rangle}{R^2 \sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} + \eta \left(2 - \|\mathbf{y}_t\|^2 / R^2 \right) \langle \mathbf{u}_{j_t}, \mathbf{y}_t \rangle \right| \\ &\geq \frac{|\langle \mathbf{u}_{j_t}, \nabla \bar{F}_{T,U}(\mathbf{y}_t) \rangle|}{\sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} - |\langle \mathbf{u}_{j_t}, \mathbf{y}_t \rangle| \left[\frac{|\langle \mathbf{y}_t, \nabla \bar{F}_{T,U}(\mathbf{y}_t) \rangle|}{R^2 \sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} + \eta \left(2 - \|\mathbf{y}_t\|^2 / R^2 \right) \right] \\ &\stackrel{(a)}{\geq} \frac{1}{\sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} - \frac{\|\mathbf{y}_t\| \|\nabla \bar{F}_{T,U}(\mathbf{y}_t)\|}{R^2 \sqrt{1 + \|\mathbf{x}_t\|^2 / R^2}} - 2\eta \stackrel{(b)}{\geq} \frac{2}{\sqrt{5}} - \frac{2\gamma\sqrt{T}}{5R} - 2\eta \stackrel{(c)}{=} \frac{1}{2}, \end{aligned}$$

where (a) is due to (69), (70), Cauchy-Schwarz inequality and $\eta \geq 0$, (b) is by $\frac{1}{\sqrt{1+\|\mathbf{x}_t\|^2/R^2}} > \frac{2}{\sqrt{5}}$ and $\frac{\|\mathbf{y}_t\|}{\sqrt{1+\|\mathbf{x}_t\|^2/R^2}} \stackrel{(64)}{=} \frac{\|\mathbf{x}_t\|}{1+\|\mathbf{x}_t\|^2/R^2} < \frac{2R}{5}$ when $\|\mathbf{x}_t\| < \frac{R}{2}$ and (68), and (c) holds under our choices of η and R .

We combine two cases to conclude

$$\Pr \left[\left\| \nabla \hat{F}_{T,U}(\mathbf{x}_t) \right\| \geq \frac{1}{2}, \forall t \leq \frac{T - \log(2/\iota)}{2q} \right] \geq 1 - \iota.$$

■

Lastly, we recall the following hard instance and its stochastic first-order oracle studied in Arjevani et al. (2023).

Lemma 7 (Lemma 2 of Arjevani et al. (2023)) *For any $T \in \mathbb{N}$, there exists a differentiable function $F_T : \mathbb{R}^T \rightarrow \mathbb{R}$ satisfying the following properties:*

1. $F_T(\mathbf{0}) = 0$ and $\inf_{\mathbf{x} \in \mathbb{R}^T} F_T(\mathbf{x}) \geq -fT$, where $f = 12$.
2. F_T is ℓ -smooth, where $\ell = 152$.
3. For all $\mathbf{x} \in \mathbb{R}^T$, $\|\nabla F_T(\mathbf{x})\|_\infty \leq \gamma$, where $\gamma = 23$.
4. For all $\mathbf{x} \in \mathbb{R}^T$, $\text{prog}_0(\nabla F_T(\mathbf{x})) \leq \text{prog}_{1/2}(\mathbf{x}) + 1$.
5. For all $\mathbf{x} \in \mathbb{R}^T$ and $i \triangleq \text{prog}_{1/2}(\mathbf{x})$, $\nabla F_T(\mathbf{x}) = \nabla F_T(\mathbf{x}_{\leq i+1})$ and $[\nabla F_T(\mathbf{x})]_{\leq i} = [\nabla F_T(\mathbf{x}_{\leq i})]_{\leq i}$.
6. For all $\mathbf{x} \in \mathbb{R}^T$, if $\text{prog}_1(\mathbf{x}) < T$, then $\|\nabla F_T(\mathbf{x})\| \geq |\nabla F_T(\mathbf{x})[\text{prog}_1(\mathbf{x}) + 1]| > 1$.

Lemma 8 (Lemma 3 of Arjevani et al. (2023)) *For any $T \in \mathbb{N}$ and F_T in Lemma 7, the following $\mathbf{g}_T : \mathbb{R}^T \times \mathcal{R} \rightarrow \mathbb{R}^T$, where $\mathcal{R} \triangleq \{0, 1\}$, is a stochastic first-order oracle of F_T :*

$$\mathbf{g}_T(\mathbf{x}, r)[i] \triangleq \begin{cases} \nabla F_T(\mathbf{x})[i] & i \neq \text{prog}_{1/4}(\mathbf{x}) + 1 \\ \frac{r}{q} \nabla F_T(\mathbf{x})[i] & i = \text{prog}_{1/4}(\mathbf{x}) + 1 \end{cases}, \forall i \in [T],$$

where $r = \text{Bernoulli}(q)$ for some $q \in (0, 1)$. Moreover, $\mathbf{g}_T$ is a probability- q zero-chain.

F.6.3. ANALYSIS UNDER HEAVY-TAILED NOISE

From now on, we need to diverge from Arjevani et al. (2023) and Cutkosky et al. (2023), since both of which are under the finite variance case (i.e., $p = 2$) instead of heavy-tailed noise (i.e., $p \in (1, 2]$).

Lemma 9 (variation of Lemma 7 in Arjevani et al. (2023) and Lemma 26 in Cutkosky et al. (2023))

The instances $\hat{F}_{T,U}$ and $\hat{\mathbf{g}}_{T,U}$ (see (56) and (57)) constructed based on F_T in Lemma 7 and $\mathbf{g}_T$ in Lemma 8 under $\eta = \frac{1}{\sqrt{5}} - \frac{39}{140}$ and $R = 7\gamma\sqrt{T}$ for $\gamma = 23$ satisfy the following properties:

1. $\hat{F}_{T,U}(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^d} \hat{F}_{T,U}(\mathbf{x}) \leq \hat{f}T$, where $\hat{f} = 12$.

2. $\widehat{F}_{T,U}$ is $\widehat{\ell}$ -smooth, where $\widehat{\ell} = 154$.

3. For all $\mathbf{x} \in \mathbb{R}^d$, $\|\nabla \widehat{F}_{T,U}(\mathbf{x})\| \leq \widehat{\gamma} \sqrt{T}$, where $\widehat{\gamma} = 53$.

4. For all $\mathbf{x} \in \mathbb{R}^d$ and $p \in (1, 2]$, $\mathbb{E}_r \left[\left\| \nabla \widehat{F}_{T,U}(\mathbf{x}) - \widehat{\mathbf{g}}_{T,U}(\mathbf{x}, r) \right\|^p \right] \leq \frac{(2\gamma)^p (1 - q^{p-1})}{(p-1)q^{p-1}}.$

Proof First, we know

$$\begin{aligned} \widehat{F}_{T,U}(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^d} \widehat{F}_{T,U}(\mathbf{x}) &\stackrel{(58)}{=} F_T(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^d} \left(F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) + \eta \mathbf{x}^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}) \right) \\ &\stackrel{\mathbf{x}^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}) \geq 0}{\leq} F_T(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^T} F_T(\mathbf{x}) \stackrel{\text{Lemma 7}}{\leq} fT = \widehat{f}T. \end{aligned}$$

Next, for any $\mathbf{x} \in \mathbb{R}^d$,

$$\nabla \widehat{F}_{T,U}(\mathbf{x}) \stackrel{(60)}{=} \mathcal{J}_{R,d}(\mathbf{x})^\top U \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) + \eta \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\|^2 / R^2 \right) \boldsymbol{\rho}_{R,d}(\mathbf{x}).$$

By Lemma 14 in (Arjevani et al., 2023) and Lemma 7, $\mathcal{J}_{R,d}(\mathbf{x})^\top U \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}))$ is $\ell + \frac{3\gamma\sqrt{T}}{R} \stackrel{R=7\sqrt{T}}{=} \ell + \frac{3}{7}$ -Lipschitz for $\ell = 152$. Moreover, we have

$$\begin{aligned} &\left\| \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\|^2 / R^2 \right) \boldsymbol{\rho}_{R,d}(\mathbf{x}) - \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{y})\|^2 / R^2 \right) \boldsymbol{\rho}_{R,d}(\mathbf{y}) \right\| \\ &\leq 2 \|\boldsymbol{\rho}_{R,d}(\mathbf{x}) - \boldsymbol{\rho}_{R,d}(\mathbf{y})\| + \frac{\|\boldsymbol{\rho}_{R,d}(\mathbf{y})\|^2}{R^2} \|\boldsymbol{\rho}_{R,d}(\mathbf{y}) - \boldsymbol{\rho}_{R,d}(\mathbf{x})\| \\ &\quad + \frac{\left| \|\boldsymbol{\rho}_{R,d}(\mathbf{y})\|^2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\|^2 \right|}{R^2} \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\| \\ &\stackrel{(a)}{\leq} 2 \|\mathbf{x} - \mathbf{y}\| + \|\mathbf{y} - \mathbf{x}\| + 2 \|\mathbf{y} - \mathbf{x}\| = 5 \|\mathbf{x} - \mathbf{y}\|, \end{aligned}$$

where (a) is by $\|\boldsymbol{\rho}_{R,d}(\mathbf{x}) - \boldsymbol{\rho}_{R,d}(\mathbf{y})\| \stackrel{\text{Lemma 5}}{\leq} \|\mathbf{x} - \mathbf{y}\|$ and $\|\boldsymbol{\rho}_{R,d}(\cdot)\| \leq R$. So $\widehat{F}_{T,U}(\mathbf{x})$ is $\ell + \frac{3}{7} + 5\eta \stackrel{\eta = \frac{1}{\sqrt{5}} - \frac{39}{140}}{=} \frac{4229}{28} + \sqrt{5} \leq 154 = \widehat{\ell}$ -smooth.

Moreover, we observe that

$$\left\| \nabla \widehat{F}_{T,U}(\mathbf{x}) \right\| \stackrel{(60)}{\leq} \|\mathcal{J}_{R,d}(\mathbf{x})\|_{\text{op}} \left\| U \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) \right\| + \eta \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\|^2 / R^2 \right) \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\|.$$

Note that for $\gamma = 23$,

$$\begin{aligned} \|\mathcal{J}_{R,d}(\mathbf{x})\|_{\text{op}} \left\| U \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) \right\| &\stackrel{\text{Lemma 5}}{\leq} \left\| U \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) \right\| \\ &\stackrel{U^\top U = I_T}{=} \left\| \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) \right\| \stackrel{\text{Lemma 7}}{\leq} \gamma \sqrt{T}, \end{aligned}$$

and

$$\begin{aligned} \eta \left(2 - \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\|^2 / R^2 \right) \|\boldsymbol{\rho}_{R,d}(\mathbf{x})\| &\stackrel{\|\boldsymbol{\rho}_{R,d}(\cdot)\| \leq R}{\leq} \eta R \max_{0 \leq a \leq 1} (2 - a^2) a = \eta R \cdot \frac{4}{3} \sqrt{\frac{2}{3}} \\ &= \left(\frac{1}{\sqrt{5}} - \frac{39}{140} \right) \cdot 7 \cdot \frac{4}{3} \sqrt{\frac{2}{3}} \cdot \gamma \sqrt{T} \leq 1.3 \gamma \sqrt{T}. \end{aligned}$$

We hence have $\|\nabla \hat{F}_{T,U}(\mathbf{x})\| \leq 2.3\gamma\sqrt{T} \leq 53\sqrt{T} = \hat{\gamma}\sqrt{T}$.

Lastly, still for $\gamma = 23$, we compute

$$\begin{aligned} \|\nabla \hat{F}_{T,U}(\mathbf{x}) - \hat{\mathbf{g}}_{T,U}(\mathbf{x}, r)\| &\stackrel{(60),(62)}{=} \|\mathcal{J}_{R,d}(\mathbf{x})^\top U \nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) - \mathcal{J}_{R,d}(\mathbf{x})^\top U \mathbf{g}_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}), r)\| \\ &\leq \|\mathcal{J}_{R,d}(\mathbf{x})\|_{\text{op}} \|U (\nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) - \mathbf{g}_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}), r))\| \\ &\stackrel{\text{Lemma 5}}{\leq} \|U (\nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) - \mathbf{g}_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}), r))\| \\ &\stackrel{U^\top U = I_T}{=} \|\nabla F_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x})) - \mathbf{g}_T(U^\top \boldsymbol{\rho}_{R,d}(\mathbf{x}), r)\| \\ &\stackrel{\text{Lemma 8}}{=} \left|1 - \frac{r}{q}\right| \|\nabla F_T(\mathbf{x}) [\text{prog}_{1/4}(\mathbf{x}) + 1]\| \stackrel{\text{Lemma 7}}{\leq} \gamma \left|1 - \frac{r}{q}\right|, \end{aligned}$$

which implies

$$\begin{aligned} \mathbb{E}_r \left[\|\nabla \hat{F}_{T,U}(\mathbf{x}) - \hat{\mathbf{g}}_{T,U}(\mathbf{x}, r)\|^p \right] &\leq \gamma^p \left(1 - q + \frac{(1-q)^p}{q^{p-1}}\right) = \gamma^p (1-q) \left(\frac{q^{p-1} + (1-q)^{p-1}}{q^{p-1}}\right) \\ &\stackrel{\text{Fact 3}}{\leq} \frac{\gamma^p 2^{2-p} (1 - q^{p-1})}{(p-1)q^{p-1}} \stackrel{p \geq 1}{\leq} \frac{(2\gamma)^p (1 - q^{p-1})}{(p-1)q^{p-1}}. \end{aligned}$$

■

Next, inspired by [Cutkosky et al. \(2023\)](#), we first prove a lower bound for heavy-tailed smooth nonconvex optimization, as presented in the following Theorem 10.

Theorem 10 For any $\Delta > 0$, $H > 0$, $p \in (1, 2]$, $\sigma \geq 0$, $0 < \epsilon \leq \sqrt{\frac{\Delta H}{96 \cdot 12 \cdot 154}}$, let d be in the order of $\frac{\Delta H}{\epsilon^2} \log \left(\frac{\Delta H}{\epsilon^2} \left(1 + \left(\frac{\sigma}{\epsilon}\right)^{\frac{p}{2(p-1)}}\right) \right)$ (see proof for the precise definition), there exists a distribution over functions $F : \mathbb{R}^d \rightarrow \mathbb{R}$ and stochastic first-order oracles $\mathbf{g}$ such that with probability 1, $F(\mathbf{0}) - F_\star \leq \Delta$, F is H -smooth and $\frac{3}{2}\sqrt{H\Delta}$ -Lipschitz, and $\mathbf{g}$ has a finite p -th centered moment σ^p . Moreover, for any randomized $\mathbf{A} \in \mathbf{A}_{\text{rand}}$ employed to optimize a randomly selected F interacting with $\mathbf{g}$, to output a point $\mathbf{x}$ such that $\mathbb{E}[\|\nabla F(\mathbf{x})\|] \leq \epsilon$, the number of queries of $\mathbf{g}$ by $\mathbf{A}$ satisfies $\gtrsim \Delta H \epsilon^{-2} + \Delta H \sigma^{\frac{p}{p-1}} \epsilon^{-\frac{3p-2}{p-1}}$.

Remark 13 The reader familiar with the literature may find that a similar lower bound (in fact, exactly the same order) has been established by [Zhang et al. \(2020a\)](#); [Liu and Zhou \(2025\)](#) before, and hence may wonder about the difference. Here, we note that the lower bounds in [Zhang et al. \(2020a\)](#); [Liu and Zhou \(2025\)](#) are shown for a special algorithmic class known as zero-respecting algorithms⁵. However, our Theorem 10 is proved for a broader family, i.e., randomized algorithms. This fact is important because Algorithm 4 is a randomized algorithm but not a zero-respecting algorithm.

Proof For $T \in \mathbb{N}$ and $q \in (0, 1)$ being determined later, let $\iota = 1/2$, $\eta = \frac{1}{\sqrt{5}} - \frac{39}{140}$, $R = 7\gamma\sqrt{T}$ where $\gamma = 23$ and $d = \left\lceil T + 32R^2 \log \frac{2T^2}{qu} \right\rceil = \left\lceil T + 32R^2 \log \frac{4T^2}{q} \right\rceil = \Theta(T \log \frac{T^2}{q})$, we construct $\hat{F}_{T,U} : \mathbb{R}^d \rightarrow \mathbb{R}$ and $\hat{\mathbf{g}}_{T,U}(\mathbf{x}, r) : \mathbb{R}^d \times \mathcal{R} \rightarrow \mathbb{R}^d$ based on F_T in Lemma 7 and $\mathbf{g}_T$ in Lemma 8.

5. A first-order algorithm is called zero-respecting if it satisfies $\mathbf{x}_t \in \cup_{s < t} \text{support}(\mathbf{g}_s)$, $\forall t \in \mathbb{N}$. For more details, see Definition 1 of [Arjevani et al. \(2023\)](#).

By Lemma 6 (note that F_T satisfies the requirements of Lemma 6 due to Lemma 7), for any $A \in A_{\text{rand}}$ employed to optimize $\widehat{F}_{T,U}$ interacting with the stochastic first-order oracle $\widehat{\mathbf{g}}_{T,U}$, where U is drawn from $\text{Ortho}(d, T)$ uniformly, we have

$$\Pr \left[\left\| \nabla \widehat{F}_{T,U}(\mathbf{x}_t) \right\| \geq \frac{1}{2}, \forall t \leq \frac{T - \log 4}{2q} \right] \geq \frac{1}{2}. \quad (71)$$

By Lemma 9, $\widehat{F}_{T,U}(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^d} \widehat{F}_{T,U}(\mathbf{x}) \leq \widehat{f}T$ for $\widehat{f} = 12$, $\widehat{F}_{T,U}$ is $\widehat{\ell}$ -smooth for $\widehat{\ell} = 154$ and $\widehat{\gamma}\sqrt{T}$ -Lipschitz for $\widehat{\gamma} = 53$, $\mathbb{E}_r \left[\left\| \nabla \widehat{F}_{T,U}(\mathbf{x}) - \widehat{\mathbf{g}}_{T,U}(\mathbf{x}, r) \right\|^p \right] \leq \frac{(2\gamma)^p (1 - q^{p-1})}{(p-1)q^{p-1}}, \forall \mathbf{x} \in \mathbb{R}^d$.

Now we set

$$\lambda \triangleq \frac{4\widehat{\ell}\epsilon}{H}, \quad T \triangleq \left\lfloor \frac{\widehat{\ell}\Delta}{\widehat{f}H\lambda^2} \right\rfloor, \quad q \triangleq \frac{1}{\left[1 + (p-1) \left(\frac{\widehat{\ell}\sigma}{2\gamma H\lambda} \right)^p \right]^{\frac{1}{p-1}}}. \quad (72)$$

For any $U \in \text{Ortho}(d, T)$, we introduce

$$F_U(\mathbf{x}) \triangleq \frac{H\lambda^2}{\widehat{\ell}} \widehat{F}_{T,U}(\mathbf{x}/\lambda) \quad \text{and} \quad \mathbf{g}_U(\mathbf{x}) = \frac{H\lambda}{\widehat{\ell}} \widehat{\mathbf{g}}_{T,U}(\mathbf{x}/\lambda, r).$$

We observe that

$$F_U(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^d} F_U(\mathbf{x}) = \frac{H\lambda^2}{\widehat{\ell}} \left(\widehat{F}_{T,U}(\mathbf{0}) - \inf_{\mathbf{x} \in \mathbb{R}^d} \widehat{F}_{T,U}(\mathbf{x}/\lambda) \right) \leq \frac{H\lambda^2}{\widehat{\ell}} \widehat{f}T \stackrel{(72)}{\leq} \Delta.$$

Moreover, $\nabla F_U(\mathbf{x}) = \frac{H\lambda}{\widehat{\ell}} \nabla \widehat{F}_{T,U}(\mathbf{x}/\lambda)$, which implies F_U is H -smooth and $\frac{H\lambda\widehat{\gamma}\sqrt{T}}{\widehat{\ell}} \stackrel{(72)}{\leq} \frac{3\sqrt{H\Delta}}{2}$ -Lipschitz. In addition, there is

$$\begin{aligned} \mathbb{E}_r [\| \nabla F_U(\mathbf{x}) - \mathbf{g}_U(\mathbf{x}) \|^p] &= \left(\frac{H\lambda}{\widehat{\ell}} \right)^p \mathbb{E}_r [\| \nabla \widehat{F}_{T,U}(\mathbf{x}/\lambda) - \widehat{\mathbf{g}}_{T,U}(\mathbf{x}/\lambda, r) \|^p] \\ &\leq \left(\frac{2\gamma H\lambda}{\widehat{\ell}} \right)^p \frac{1 - q^{p-1}}{(p-1)q^{p-1}} \stackrel{(72)}{=} \sigma^p. \end{aligned}$$

For any $A \in A_{\text{rand}}$ used to optimize F_U with $\mathbf{g}_U$ when U is drawn from $\text{Ortho}(d, T)$ uniformly, we can view as $\mathbf{x}_t/\lambda$ as the output of another $A^\lambda \in A_{\text{rand}}$ interacting with $\widehat{F}_{T,U}$ and $\widehat{\mathbf{g}}_{T,U}$ (a similar argument is used in the proof of Lemma 6). As such, by (71),

$$\Pr \left[\left\| \nabla \widehat{F}_{T,U}(\mathbf{x}_t/\lambda) \right\| \geq \frac{1}{2}, \forall t \leq \frac{T - \log 4}{2q} \right] \geq \frac{1}{2},$$

which implies with probability at least $1/2$

$$\| \nabla F_U(\mathbf{x}_t) \| = \frac{H\lambda}{\widehat{\ell}} \left\| \nabla \widehat{F}_{T,U}(\mathbf{x}_t/\lambda) \right\| \geq \frac{H\lambda}{2\widehat{\ell}}, \forall t \leq \frac{T - \log 4}{2q}.$$

Therefore, we have

$$\mathbb{E} [\| \nabla F_U(\mathbf{x}_t) \|] \geq \frac{H\lambda}{2\widehat{\ell}} \Pr \left[\| \nabla F_U(\mathbf{x}_t) \| \geq \frac{H\lambda}{2\widehat{\ell}} \right] \geq \frac{H\lambda}{4\widehat{\ell}} \stackrel{(72)}{=} \epsilon, \forall t \leq \frac{T - \log 4}{2q}.$$

Lastly, we compute

$$\begin{aligned}
 \frac{T - \log 4}{2q} &\stackrel{(72)}{=} \frac{1}{2} \left(\left\lfloor \frac{\Delta H}{16\widehat{f}\widehat{\ell}\epsilon^2} \right\rfloor - \log 4 \right) \left[1 + (\mathfrak{p} - 1) \left(\frac{\widehat{\ell}\sigma}{2\gamma H\lambda} \right)^{\mathfrak{p}} \right]^{\frac{1}{\mathfrak{p}-1}} \\
 &\stackrel{(a)}{\geq} \frac{1}{2} \left(\frac{\Delta H}{16\widehat{f}\widehat{\ell}\epsilon^2} - 3 \right) \left[1 + (\mathfrak{p} - 1) \left(\frac{\widehat{\ell}\sigma}{2\gamma H\lambda} \right)^{\mathfrak{p}} \right]^{\frac{1}{\mathfrak{p}-1}} \\
 &\stackrel{(b)}{\geq} \frac{\Delta H}{64\widehat{f}\widehat{\ell}\epsilon^2} \left(1 + (\mathfrak{p} - 1)^{\frac{1}{\mathfrak{p}-1}} \left(\frac{\widehat{\ell}\sigma}{2\gamma H\lambda} \right)^{\frac{\mathfrak{p}}{\mathfrak{p}-1}} \right) \\
 &\stackrel{(72)}{\gtrsim} \Delta H \epsilon^{-2} + \Delta H \sigma^{\frac{\mathfrak{p}}{\mathfrak{p}-1}} \epsilon^{-\frac{3\mathfrak{p}-2}{\mathfrak{p}-1}},
 \end{aligned}$$

where the (a) is by $\lfloor \cdot \rfloor \geq \cdot - 1$ and $\log 4 \leq 2$, and (b) holds due to the condition $\epsilon^2 \leq \frac{\Delta H}{96 \cdot 12 \cdot 154} = \frac{\Delta H}{96\widehat{f}\widehat{\ell}}$, implying $\frac{\Delta H}{16\widehat{f}\widehat{\ell}\epsilon^2} - 3 \geq \frac{\Delta H}{32\widehat{f}\widehat{\ell}\epsilon^2}$, and $(1 + x)^{\frac{1}{\mathfrak{p}-1}} \geq 1 + x^{\frac{1}{\mathfrak{p}-1}}$ for $x \geq 0$ when $\mathfrak{p} \in (1, 2]$. So to achieve $\mathbb{E}[\|\nabla F_U(\mathbf{x}_t)\|] < \epsilon$, it requires at least $\Delta H \epsilon^{-2} + \Delta H \sigma^{\frac{\mathfrak{p}}{\mathfrak{p}-1}} \epsilon^{-\frac{3\mathfrak{p}-2}{\mathfrak{p}-1}}$ many iterations. In particular, we note that $d = \Theta(T \log \frac{T^2}{q}) = \Theta\left(\frac{\Delta H}{\epsilon^2} \log\left(\frac{\Delta H}{\epsilon^2} \left(1 + \left(\frac{\sigma}{\epsilon}\right)^{\frac{\mathfrak{p}}{2(\mathfrak{p}-1)}}\right)\right)\right)$. $\blacksquare$

Equipped with Theorem 10, we can show the lower bound on the distributional complexity (Nemirovski and Yudin, 1983) for nonsmooth nonconvex optimization with heavy tails in the following Theorem 11.

Theorem 11 *For any given $\Delta > 0$, $G > 0$, $\mathfrak{p} \in (1, 2]$, $\sigma \geq 0$, $\delta > 0$ and $0 < \epsilon \leq \frac{\Delta}{96 \cdot 12 \cdot 154 \delta} \wedge \frac{4G^2\delta}{9\Delta}$, let d be in the order of $\frac{\Delta}{\delta\epsilon} \log\left(\frac{\Delta}{\delta\epsilon} \left(1 + \left(\frac{\sigma}{\epsilon}\right)^{\frac{\mathfrak{p}}{2(\mathfrak{p}-1)}}\right)\right)$, there exists a distribution over functions $F : \mathbb{R}^d \rightarrow \mathbb{R}$, and stochastic first-order oracles $\mathbf{g}$ such that with probability 1, $F(\mathbf{0}) - F_\star \leq \Delta$, F is G -Lipschitz and $\mathbf{g}$ has a finite $\mathfrak{p}$ -th centered moment $\sigma^{\mathfrak{p}}$. Moreover, for any randomized $\mathbf{A} \in \mathbf{A}_{\text{rand}}$ employed to optimize a randomly selected F interacting with $\mathbf{g}$, to output a point $\mathbf{x}$ such that $\mathbb{E}[\|\nabla F(\mathbf{x})\|_\delta] \leq \epsilon$, the number of queries of $\mathbf{g}$ by $\mathbf{A}$ satisfies $\gtrsim \Delta \delta^{-1} \epsilon^{-1} + \Delta \sigma^{\frac{\mathfrak{p}}{\mathfrak{p}-1}} \delta^{-1} \epsilon^{-\frac{2\mathfrak{p}-1}{\mathfrak{p}-1}}$.*

Proof Let $H \triangleq \epsilon/\delta$ in the following. Note that $\epsilon \leq \sqrt{\frac{\Delta H}{96 \cdot 12 \cdot 154}}$ due to our requirement on ϵ . So we can consider the same distribution on F and $\mathbf{g}$ as in Theorem 10. Importantly, F is H -smooth and $\frac{3}{2}\sqrt{H\Delta} = \frac{3}{2}\sqrt{\frac{\epsilon\Delta}{\delta}} \leq G$ -Lipschitz (again due to the condition on ϵ) with probability 1.

If $\mathbf{A} \in \mathbf{A}_{\text{rand}}$ finds a point $\mathbf{x}$ such that $\mathbb{E}[\|\nabla F(\mathbf{x})\|_\delta] \leq \epsilon$ (i.e., a (δ, ϵ) -stationary point of F), then by Proposition 14 of Cutkosky et al. (2023), it also satisfies $\mathbb{E}[\|\nabla F(\mathbf{x})\|] \leq \epsilon + H\delta = 2\epsilon$. Therefore, Theorem 10 implies the number of queries of $\mathbf{g}$ by $\mathbf{A}$ is at least $\gtrsim \Delta H \epsilon^{-2} + \Delta H \sigma^{\frac{\mathfrak{p}}{\mathfrak{p}-1}} \epsilon^{-\frac{3\mathfrak{p}-2}{\mathfrak{p}-1}} = \Delta \delta^{-1} \epsilon^{-1} + \Delta \sigma^{\frac{\mathfrak{p}}{\mathfrak{p}-1}} \delta^{-1} \epsilon^{-\frac{2\mathfrak{p}-1}{\mathfrak{p}-1}}$. $\blacksquare$

Finally, we are able to prove Theorem 6.

Proof [Proof of Theorem 6] We recall the fact that lower bounds on the distributional complexity imply lower bounds on the minimax complexity (Nemirovski and Yudin, 1983). Thus, Theorem 6 can be directly concluded from Theorem 11. $\blacksquare$

Appendix G. Algebraic Facts

We give three useful algebraic facts in this section.

Fact 1 For any $T \in \mathbb{N}$ and $a \in (0, 1)$, there is

$$\sum_{t=1}^{T-1} \frac{\sum_{s=t+1}^T s^a}{t(T-t)^2} \lesssim \frac{1 + \log T}{T^{1-a}}.$$

Proof Note that $\sum_{s=t+1}^T s^a \leq (T-t)T^a$, which implies

$$\sum_{t=1}^{T-1} \frac{\sum_{s=t+1}^T s^a}{t(T-t)^2} \leq \sum_{t=1}^{T-1} \frac{T^a}{t(T-t)} = \frac{1}{T^{1-a}} \sum_{t=1}^{T-1} \frac{1}{t} + \frac{1}{T-t} = \frac{2}{T^{1-a}} \sum_{t=1}^{T-1} \frac{1}{t} \lesssim \frac{1 + \log T}{T^{1-a}}.$$

■

Fact 2 Given $2 \leq N \in \mathbb{N}$, $K = \lfloor N/T \rfloor$ and $T \in \mathbb{N}$ satisfying $T \leq \lceil N/2 \rceil$, there is $KT \geq N/4$.

Proof Note that $KT = \lfloor N/T \rfloor T \geq N - T \geq (N-1)/2 \geq N/4$.

■

Fact 3 Given $p \in (1, 2]$ and $q \in (0, 1)$, there are $q^{p-1} + (1-q)^{p-1} \leq 2^{2-p}$ and $1-q \leq \frac{1-q^{p-1}}{p-1}$.

Proof Note that x^{p-1} is concave when $p \in (1, 2]$, we hence have $\frac{q^{p-1} + (1-q)^{p-1}}{2} \leq \left(\frac{q+1-q}{2}\right)^{p-1} \Rightarrow q^{p-1} + (1-q)^{p-1} \leq 2^{2-p}$. Next, let $x = 1-q \in (0, 1)$, we have

$$1-q \leq \frac{1-q^{p-1}}{p-1} \Leftrightarrow (1-x)^{p-1} \leq 1-(p-1)x,$$

which is true by Bernoulli's inequality.

■