

Power of Generalized Smoothness in Stochastic Convex Optimization: First- and Zero-Order Algorithms

Aleksandr Lobanov
MIPT, Skoltech, HSE
lobbsasha@mail.ru

Alexander Gasnikov
Innopolis, MIPT, ISP RAS
gasnikov@yandex.ru

Abstract

This paper is devoted to the study of stochastic optimization problems under the generalized smoothness assumption. By considering the unbiased gradient oracle in *Stochastic Gradient Descent*, we provide strategies to achieve in bounds the summands describing linear rate. In particular, in the case $L_0 = 0$, we obtain in the **convex setup** the iteration complexity: $N = \mathcal{O}\left(L_1 R \log \frac{1}{\varepsilon} + \frac{L_1 c R^2}{\varepsilon}\right)$ for *Clipped Stochastic Gradient Descent* and $N = \mathcal{O}\left(L_1 R \log \frac{1}{\varepsilon}\right)$ for *Normalized Stochastic Gradient Descent*. Furthermore, we generalize the convergence results to the case with a biased gradient oracle, and show that the power of (L_0, L_1) -smoothness extends to *zero-order algorithms*. Finally, we demonstrate the possibility of linear convergence in the convex setup through numerical experimentation, which has aroused some interest in the machine learning community.

1 Introduction

In many real-world scenarios, systems are often noisy and complex, making deterministic optimization infeasible. Therefore, this work focuses on a stochastic optimization problem:

$$\min_{x \in \mathbb{R}^d} \{f(x) := \mathbb{E}_{\xi \sim \mathcal{D}} [f(x, \xi)]\}, \quad (1)$$

where $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is a convex function and where we assume that optimization algorithms only have access to the gradient oracle $\mathbf{g} : \mathbb{R}^d \times \mathcal{D} \rightarrow \mathbb{R}^d$ with stochastic gradient $\mathbb{E} [\nabla f(x, \xi)] = \nabla f(x)$ and bias $\mathbf{b}(x)$ terms:

$$\mathbf{g}(x, \xi) = \nabla f(x, \xi) + \mathbf{b}(x). \quad (2)$$

Frequently, to solve problem (1) one uses what is likely already a classic optimization algorithm, namely Stochastic Gradient Descent (SGD) [10] or its variations, which have demonstrated their effectiveness in different settings, for instance, federated learning [60, 33, 58], deep learning [15, 64, 18], reinforcement learning [7, 39] and others. Among the variants of SGD, it is worth noting the Normalized Stochastic Gradient Descent (NSGD) [27, 66] which has received widely attention from the community because it addresses challenges in optimization for machine learning [8]. And it's also worth noting the Clipped Stochastic Gradient Descent (ClipSGD) [24], which is commonly used to stabilize the training of deep learning models [48, 25].

Many standard literatures analyze stochastic optimization algorithms with unbiased gradient oracle (2). In particular, SGD [36, 11], NSGD [65, 30], ClipSGD [25, 34]. However, there are a number of applications where gradient oracle (2) is biased. For example, sparsified SGD [4], delayed SGD [53], etc. Zero-order algorithms [46, 16] occupy a special place in the class of stochastic methods with biased gradient oracle (2). They are motivated by various applications, including multi-armed bandit [52, 38], online optimization [1, 6, 3], hyperparameter tuning [28, 47].

In our work, we investigate the convergence of first-order algorithms: ClipSGD, NSGD, and zero-order algorithms: ZO-ClipSGD, ZO-NSGD, assuming convexity and (L_0, L_1) -smoothness.

We emphasize the following points:

Algorithm step size. Zero-order algorithms do not have access to the exact (stochastic) gradient in particular, as well as every algorithm with a biased gradient oracle, so we focus on creating first-order methods whose step size does not depend on knowledge of the gradient at a given point. We use the developed first-order algorithms as a basis for creating zero-order methods [22].

Linear convergence. Historically [45], stochastic optimization first- and zero-order algorithms have achieved the desired accuracy with a linear rate of convergence only in strongly convex case and under assumption of standard smoothness. However, the work of [44] showed that if the generalized smoothness assumption is satisfied in a *deterministic convex* optimization problem, then gradient descent has two regimes: linear convergence rate as long as $\|\nabla f(x^k)\| \geq \frac{L_0}{L_1}$, and a sublinear convergence rate in the other case (see the example of the power of norm function in Figure 1). Considering these points, our work answers the following question:

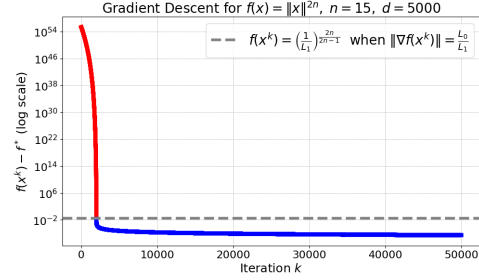


Figure 1: Changing regimes demonstration

Can linear convergence rate in stochastic convex optimization be achieved for first- and zero-order algorithms with constant step size?

1.1 Main Contributions

More specifically, our contributions are the following:

- We provide strategies to obtain summands that describe the linear convergence rate. In particular, we show that using clipping or normalization techniques can achieve the desired results.
- We improve convergence results for ClipSGD and NSGD with unbiased gradient oracle (2) in the convex setting assuming (L_0, L_1) -smoothness (see Table 1). Moreover, we show that in the case $L_0 = 0$, NSGD can converge in the convex setup with a linear convergence rate to the desired accuracy, requiring $N = \tilde{\mathcal{O}}(L_1 R)$ iterations and $B = \mathcal{O}\left(\frac{\sigma^2 M R^3}{\varepsilon^3}\right)$ batch size.
- We generalize ClipSGD, and NSGD to the case of a biased gradient oracle, showing how the bias accumulates over iterations (this result may be of independent interest).
- We provide the first convergence results for the zero-order algorithms ZO-ClipSGD (Algorithm 3), and ZO-NSGD (Algorithm 4) in the convex and (L_0, L_1) -smooth setting. We show that the power of generalized smoothness extends to zero-order methods as well, achieving summands characterizing the linear convergence rate (see Table 1).
- We demonstrate on a numerical example of logistic regression (which is of particular interest to the machine learning community) that indeed, zero- and first-order stochastic algorithms can converge with linear rates in a convex setup.

1.2 Formal Setting and Assumptions

In this subsection, we introduce and discuss main assumptions and notations used throughout paper.

Notations. We use $\langle x, y \rangle := \sum_{i=1}^d x_i y_i$ to denote standard inner product of $x, y \in \mathbb{R}^d$. We denote Euclidean norm in $\mathbb{R}^d$ as $\|x\| := \sqrt{\sum_{i=1}^d x_i^2}$. In particular, this norm $\|x\| := \sqrt{\langle x, x \rangle}$ is related to the inner product. We use $\mathcal{P}[\cdot]$ to define probability measure which is always known from the context, $\mathbb{E}[\cdot]$ denotes mathematical expectation. We use the following notation $B^d(r) := \{x \in \mathbb{R}^d : \|x\| \leq r\}$ to denote Euclidean ball (l_2 -ball) and $S^d(r) := \{x \in \mathbb{R}^d : \|x\| = r\}$ to denote Euclidean sphere. We denote $0 \leq M < \infty$ as the upper bound of the gradient norm $\|\nabla f(x^k)\|$. For simplicity, we denote $f^* := f(x^*)$ and $R = \|x^0 - x^*\|$. We use $\tilde{\mathcal{O}}(\cdot)$ to hide the logarithmic coefficients.

Table 1: Comparison of convergence results of SGD variants to the most related work [20] in the convex and (L_0, L_1) -smooth setup. Notation: $\eta \leq (L_0 + L_1 c)^{-1}$ – step size; $c > 0$ – clipping radius; $\mathcal{R} = (\eta + \frac{MR}{c^2} + \frac{R}{c})$; ε = desired accuracy; d = dimension; SLCR = summand with linear convergence rate.

Algorithm	Number of Iterations #N	Batch Size #B	Maximum Noise Level # Δ	SLCR?	Reference
ClipSGD	$\mathcal{O}\left(\frac{L_0 R^2}{\varepsilon} + \frac{\sigma^2 R^2}{\varepsilon^2} + L_1^2 R^2\right)$	✗	✗	✗	Gaash et al. [20]
	$\mathcal{O}\left(\frac{R}{\eta c} \log \frac{1}{\varepsilon} + \frac{R^2}{\eta \varepsilon}\right)$	$\mathcal{O}\left(\frac{\sigma^2 R}{\varepsilon}\right)$	✗	✓	Theorem 3.1 (Ours)
NSGD	$\mathcal{O}\left((L_1 R + \frac{L_0 R^2}{\varepsilon}) \log \frac{1}{\varepsilon}\right)$	$\mathcal{O}\left(\frac{\sigma^2 M R^3}{\varepsilon^3}\right)$	✗	✓	Theorem 4.1 (Ours)
ZO-ClipSGD	$\mathcal{O}\left(\frac{R}{\eta c} \log \frac{1}{\varepsilon} + \frac{R^2}{\eta \varepsilon}\right)$	$\mathcal{O}\left(\frac{d M R \sigma^2}{\varepsilon c^2}\right)$	$\mathcal{O}\left(\frac{\varepsilon}{\sqrt{d R (L_0 + L_1 M)}} \min\left\{\tilde{\sigma}, \frac{\varepsilon}{\sqrt{d R}}\right\}\right)$	✓	Theorem 5.2 (Ours)
ZO-NSGD	$\mathcal{O}\left((L_1 R + \frac{L_0 R^2}{\varepsilon}) \log \frac{1}{\varepsilon}\right)$	$\mathcal{O}\left(\frac{d M R^3 \sigma^2}{\varepsilon^3}\right)$	$\mathcal{O}\left(\frac{\varepsilon}{\sqrt{d R^{3/2} (L_0 + L_1 M)}} \min\left\{\tilde{\sigma}, \frac{\varepsilon^{3/2}}{\sqrt{d R^{3/2}}}\right\}\right)$	✓	Theorem 5.4 (Ours)

Assumptions on objective function. Throughout this paper, we refer to the standard L -smoothness assumption, which is widely used in the literature [e.g. 51] and has the following form:

Assumption 1.1 (L -smoothness). Function f is L -smooth if for any $x, y \in \mathbb{R}^d$ is satisfied:

$$\|\nabla f(y) - \nabla f(x)\| \leq L \|y - x\|.$$

Despite the widespread use of Assumption 1.1, our work focuses on the more general smoothness assumption, which has recently attracted increased interest. In particular, in [62] it was shown that norm of Hesse matrix correlates with norm of gradient function when training neural networks, and in [44] it was shown that using generalized smoothness it is possible to significantly improve the convergence of algorithms. (L_0, L_1) -smoothness [62, 63] has been proposed as a natural relaxation of standard smoothness assumption.

Assumption 1.2 ((L_0, L_1) -smoothness). A function $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is (L_0, L_1) -smooth if the following inequality is satisfied for any $x, y \in \mathbb{R}^d$ with $\|y - x\| \leq \frac{1}{L_1}$:

$$\|\nabla f(y) - \nabla f(x)\| \leq (L_0 + L_1 \|\nabla f(x)\|) \|y - x\|.$$

Assumption 1.2 in the case $L_1 = 0$ covers the standard Assumption 1.1. Moreover, (L_0, L_1) -smoothness is strictly more general than L -smoothness, see the examples in [62, 12, 34, 26].

Remark 1.3 (Clarification regarding $L_0 = 0$). *In this paper we often emphasize the case $L_0 = 0$ in Assumption 1.2. It is worth noting that the class of functions that do not reach their infimum x^* (converge to an asymptote) satisfies this case. Explicit examples of functions with $L_0 = 0$ are the exponent of the inner product and the logistic function (see [26] for details).*

Assumptions on gradient oracle. In our analysis, we consider cases with both unbiased and biased gradient oracle (2). Therefore, we assume that the bias and variance of gradient oracle (2) are bounded:

Assumption 1.4 (Bounded bias). There exists constant $\zeta \geq 0$ such that the bias is bounded if $\forall x \in \mathbb{R}^d$:

$$\|\mathbf{b}(x)\| \leq \zeta.$$

Assumption 1.5 (Bounded variance). There exists constant $\sigma^2 \geq 0$ such that the variance is bounded if $\forall x \in \mathbb{R}^d$:

$$\mathbb{E} \left[\|\mathbf{g}(x, \xi) - \mathbb{E}[\mathbf{g}(x, \xi)]\|^2 \right] \leq \sigma^2.$$

Assumption 1.4 is organic [see, e.g. 43], and the case $\zeta = 0$ corresponds to the unbiased gradient oracle (2). Assumption 1.5 is often used by the community [e.g. 32, 37], and is sometimes called heavy-tailed noise [25].

1.3 Paper Organization

Next, our paper has the following structure. In Section 2, we discuss related work. In Section 3, We start to present the main results of our work, in particular, we provide the first strategy for obtaining a summand characterizing the linear rate in the convergence estimate. In Section 4, we analyze NSGD in the convex setting, showing in which regime linear convergence can be observed. We provide a first analysis of zero-order algorithms under (L_0, L_1) -smoothness in Section 5. In Section 6, we discuss the results obtained. While, in Section 7, we show experimentally about the possibility of linear convergence in the convex setting. Finally, Section 8 concludes our paper. All missing proofs of Lemmas and Theorems are provided in the supplementary materials (Appendix).

2 Related Works

In this section, we will discuss the most related works.

Algorithms under (L_0, L_1) -smoothness. Generalized smoothness was first introduced in [62], which analyzed ClipSGD in the non-convex setting. A number of works [56, 41, 19, 40, 29, 59, 57] followed that also focused on the non-convex setup, including ClipSGD [63, 61, 34], NSGD [65, 31]. After that, there was interest in research on algorithms in the convex deterministic setting: Clipped Gradient Descent [34], Normalized Gradient Descent [13], Gradient Descent with Polyak step size $\eta_k = \frac{f(x^k) - f^*}{\|\nabla f(x^k)\|^2}$ [54], and $\eta_k = \frac{1}{L_0 + L_1 \|\nabla f(x^k)\|}$ [26, 55]. Moreover, in [44], it was theoretically shown that it is possible to significantly improve the convergence of algorithms in the (strongly) convex setting by achieving linear convergence rate. However, much less attention has been paid to the stochastic convex setting. Perhaps the only results are [26, 20], which considers SGD and ClipSGD achieving only a sublinear convergence rate. Moreover, in the case $L_0 = 0$ in the Assumption, the algorithms from [26, 20] cannot converge to the desired accuracy. *In our work, we focus on the stochastic convex setup, showing that existing convergence results can be significantly improved.*

Zero-order algorithms. The work of [21] showed that to achieve optimal estimates of iteration N and oracle T complexity in zero-order algorithms, one should base it on a first-order algorithm using a gradient approximation as the biased gradient oracle (2), which uses only information about the objective function f . Using this technique a number of works have achieved the best convergence results in various settings including distributed optimization [2], federated optimization [49], overparameterization [42], Polyak-Lojasiewicz condition [23], etc. However, all these works assumed standard smoothness (Assumption 1.1) and achieved only sublinear convergence rates. *In our work, we present convergence results for zero-order algorithms under (L_0, L_1) -smoothness.*

3 Clipped Stochastic Gradient Descent

In this section we begin to present the main results of our work. In particular, we analyze the convergence of SGD variants under convexity and (L_0, L_1) -smoothness with step size independent of the gradient norm. We assume that the gradient oracle (2) is unbiased $\zeta = 0$, i.e., Assumption 1.5 takes:

$$\mathbb{E} [\|\nabla f(x, \xi) - \nabla f(x)\|^2] \leq \sigma^2.$$

As a first strategy to obtain the summands that characterize the linear rate, we consider the clipping technique. Applying this technique we produce the ClipSGD, which has the following form:

Algorithm 1 Clipped Stochastic Gradient Descent Method (ClipSGD)

Input: initial point $x_0 \in \mathbb{R}^d$, iterations N , batch size B , step size $\eta_k > 0$ and clipping radius $c > 0$
for $k = 0$ **to** $N - 1$ **do**
 1. Draw fresh i.i.d. samples $\xi_1^k, \dots, \xi_B^k$
 2. $\nabla f(x^k, \xi^k) = \frac{1}{B} \sum_{i=1}^B \nabla f(x^k, \xi_i^k)$
 3. $\text{clip}_c(\nabla f(x^k, \xi^k)) = \min \left\{ 1, \frac{c}{\|\nabla f(x^k, \xi^k)\|} \right\} \nabla f(x^k, \xi^k)$
 4. $x^{k+1} \leftarrow x^k - \eta_k \cdot \text{clip}_c(\nabla f(x^k, \xi^k))$
end for
Return: x^N

Algorithm 1 uses the clipped stochastic gradient $\text{clip}_c(\nabla f(x, \xi))$, which normalizes the gradient only if $\|\nabla f(x, \xi)\| > c$. Next theorem provides the convergence result for ClipSGD.

Theorem 3.1. *Let function f satisfy Assumption 1.2 ((L_0, L_1) -smoothness) and unbiased gradient oracle (2) satisfy Assumption 1.5 (bounded variance), then Algorithm 1 with constant step size $\eta_k = \eta \leq [4(L_0 + L_1 c)]^{-1}$ and arbitrary clipping radius c guarantees error:*

$$\mathbb{E} [f(x^N)] - f^* \lesssim \left(1 - \frac{\eta c}{R}\right)^K (f(x^0) - f^*) + \frac{R^2}{\eta(N-K)} + \frac{\sigma^2}{B} \left(\eta + \frac{MR}{c^2} + \frac{R}{c}\right),$$

where $0 \leq K < N$ is number of iterations for which $\|\nabla f(x^k)\| \leq \frac{c}{2}$ is satisfied.

It should be noted that the results of Theorem 3.1 are given with a choice of step size independent of the gradient at the current point. This choice of step allows us to separate the constants L_0 and L_1 in the final estimates. The summand $\frac{R^2}{\eta(N-K)}$ is a typical ClipSGD characterizing the sublinear rate (see e.g. [25]). However, it is worth noting that by substituting $\eta = (L_0 + L_1 c)^{-1}$, then the summand with $L_1 : \frac{L_1 c R^2}{N-K}$ already improves existing results both assuming standard smoothness [25] and generalized smoothness [20]. The first summand $(1 - \frac{\eta c}{R})^K (f(x^0) - f^*)$, which characterizes the linear rate, deserves special attention. *To the best of our knowledge, this is the first result for ClipSGD showing such a summand in a convex setting.* Moreover, by substituting $\eta = (L_0 + L_1 c)^{-1}$, it is not hard to see that in the regime $L_0 = 0$ (see Remark 1.3), Algorithm 1 at a batch size $B = \mathcal{O}\left(\frac{\sigma^2}{\varepsilon} \left(\eta + \frac{MR}{c^2} + \frac{R}{c}\right)\right)$ requires only $N = \mathcal{O}\left(L_1 R \log \frac{1}{\varepsilon} + \frac{L_1 c R^2}{\varepsilon}\right)$ iterations. This iteration complexity significantly outperforms standard results in the L -smoothness setting (Assumption 1.1), since [37] shows a lower bound consisting only of a sublinear convergence rate. Moreover, comparing to the closest work to the problem setting, then even when $\sigma = 0$ [20] does not guarantee convergence to the desired accuracy, offering an estimate of $N = \mathcal{O}(L_1^2 R^2)$ that is independent of accuracy.

4 Normalized Stochastic Gradient Descent

In the previous section, we showed that it is possible to obtain in the final convergence estimate a summand characterizing the linear rate. In addition, we highlighted the regime $L_0 = 0$, in which ClipSGD has the following iteration complexity: $N = \mathcal{O}\left(L_1 R \log \frac{1}{\varepsilon} + \frac{L_1 c R^2}{\varepsilon}\right)$. However, with this iteration complexity, it cannot be said that the algorithm can converge with linear rate to the desired accuracy. Such an estimate can characterize that the algorithm converges with linear rate as long as the gradient norm is large $\|\nabla f(x^k)\| \geq c$, and then slows down to the sublinear rate. However, it is worth noting that the summand responsible for the sublinear rate depends on the clipping radius: $\frac{L_1 c R^2}{\varepsilon}$. That is, if we take c smaller, the ClipSGD will take longer to converge to the linear rate. Thus, noticing that the regime $\|\nabla f(x^k)\| \geq c$ is a gradient normalization, then considering NSGD, it seems that one can achieve a true linear convergence rate to the desired accuracy.

In this section, we consider a normalization technique to obtain the summand characterizing the linear rate. Applying this technique we produce the NSGD, which has the following form (see Algorithm 2):

Algorithm 2 Normalized Stochastic Gradient Descent Method (NSGD)

Input: initial point $x_0 \in \mathbb{R}^d$, iterations number N , batch size B and step size $\eta_k > 0$
for $k = 0$ **to** $N - 1$ **do**
 1. Draw fresh i.i.d. samples $\xi_1^k, \dots, \xi_B^k$
 2. $\nabla f(x^k, \xi^k) = \frac{1}{B} \sum_{i=1}^B \nabla f(x^k, \xi_i^k)$
 3. $x^{k+1} \leftarrow x^k - \eta_k \cdot \frac{\nabla f(x^k, \xi^k)}{\|\nabla f(x^k, \xi^k)\|}$
end for
Return: x^N

The following theorem provides a convergence result for Algorithm 2.

Theorem 4.1. *Let function f satisfy Assumption 1.2 ((L_0, L_1) -smoothness) and unbiased gradient oracle (2) satisfy Assumption 1.5 (bounded variance), then Algorithm 2 with hyperparameter $\lambda > 0$ and constant step size $\eta_k = \eta \leq \lambda / [2(L_0 + L_1 \lambda)]$ guarantees:*

$$\mathbb{E}[f(x^N)] - f^* \lesssim \left(1 - \frac{\eta}{R}\right)^N (f(x^0) - f^*) + \frac{\sigma^2 MR}{B\lambda^2} + \lambda R.$$

From Theorem 4.1 we can see that we have indeed got rid of the summand characterizing the sublinear rate from the deterministic part. *Thus, we see that it is normalization that allows us to achieve the summand characterizing the linear rate* $(1 - \frac{\eta}{R})^N (f(x^0) - f^*)$. However, note that by substituting $\eta = \lambda / [2(L_0 + L_1 \lambda)]$, we obtain a summand with $L_0 : \left(1 - \frac{\lambda}{RL_0}\right)^N (f(x^0) - f^*)$, which is in fact

sublinear since it depends on the hyperparameter λ (it follows from the third summand that $\lambda \sim \varepsilon/R$), and with $L_1 : \left(1 - \frac{1}{RL_1}\right)^N (f(x^0) - f^*)$, which is indeed linear since it does not depend on λ in any way. That is, Algorithm 2, which uses batch parallelization, requires $N = \mathcal{O}\left((L_1 R + \frac{L_0 R^2}{\varepsilon}) \log \frac{1}{\varepsilon}\right)$ iterations. Similar to the reasoning in the previous section, it is worth highlighting the regime $L_0 = 0$ (see Remark 1.3). Then we obtain a very surprising result on iteration complexity, namely, to achieve the desired accuracy NSGD converge with a linear rate of $N = \mathcal{O}\left(L_1 R \log \frac{1}{\varepsilon}\right)$ iterations. *This estimate breaks all existing bounds on first-order algorithms [37], given the specificity of the problem formulation, namely convexity.* However, to achieve this rate over iterations, NSGD requires a batch size $B = \mathcal{O}\left(\frac{\sigma^2 M R^3}{\varepsilon^3}\right)$. We emphasize that the fact that NSGD requires a large batch size is not surprising (see, e.g., [14]), in contrast to the true linear convergence rate in the convex setup.

5 Zero-Order Algorithms

In this section, we consider another class of algorithms: optimization algorithms that have access only to an objective function value f possibly with some bounded adversarial noise $|\delta(x)| \leq \Delta$:

$$\tilde{f}(x, \xi) = f(x, \xi) + \delta(x). \quad (3)$$

In (3), Δ means the maximum possible allowable noise level at which the desired accuracy can still be achieved. In [5], the importance of considering Δ as a third optimality criterion for zero-order algorithms was shown. In particular, in some applications [9], the larger noise level Δ is, the cheaper the call to the inexact oracle $\tilde{f}$ in (3).

Since this class of algorithms does not have access to the stochastic gradient $\nabla f(x, \xi)$, the gradient oracle (2) will be the gradient approximation with L_2 randomization [52, 46]:

$$\mathbf{g}(x, \{e, \xi\}) = \frac{d}{2\gamma} \left(\tilde{f}(x + \gamma e, \xi) - \tilde{f}(x - \gamma e, \xi) \right) e, \quad (4)$$

where $\gamma > 0$ is a smoothing parameter, e is a random vector uniformly distributed in $S^d(1)$.

Due to the fact that the gradient approximation is the biased gradient oracle (2), in order to create zero-order algorithms by basing on the results in Sections 3 and 4, it is necessary to first generalize the results of Theorems 3.1, 4.1 (note that in these regimes it is not necessary to know the (stochastic) norm of the gradient with step size) to the case of gradient oracle with bias.

Next, we present convergence results for the following two zero-order algorithms: ZO-ClipSGD (Algorithm 3) and ZO-NSGD (Algorithm 4).

5.1 ZO-ClipSGD Method

The first algorithm we consider in this section is ZO-ClipSGD. This algorithm is a modification of ClipSGD (Algorithm 1), which uses instead of the original $\|\nabla f(x, \xi)\|$, the stochastic gradient approximation (4), which is the biased gradient oracle (2). The ZO-ClipSGD has the following form:

Algorithm 3 Zero-Order Clipped Stochastic Gradient Descent Method (ZO-ClipSGD)

Input: initial point $x_0 \in \mathbb{R}^d$, iterations N , batch size B , step size $\eta_k > 0$ and clipping radius $c > 0$
for $k = 0$ **to** $N - 1$ **do**
 1. Draw fresh i.i.d. samples $\xi_1^k, \dots, \xi_B^k$ and $e_1^k, \dots, e_B^k$
 2. $\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\}) = \frac{1}{B} \sum_{i=1}^B \mathbf{g}(x^k, \{e_i^k, \xi_i^k\})$ via (4)
 3. $\text{clip}_c(\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\})) = \min \left\{ 1, \frac{c}{\|\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\})\|} \right\} \mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\})$
 4. $x^{k+1} \leftarrow x^k - \eta_k \cdot \text{clip}_c(\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\}))$
end for
Return: x^N

Before proceeding to present the convergence results of Algorithm 3, we note that the gradient approximation (4) is a biased gradient oracle, so we cannot directly use the results obtained in

Theorem 3.1. Thus, in order to obtain estimates for the iteration complexity N , oracle complexity T and maximum noise Δ , we first need to generalize the results of Theorem 3.1 to the case with a biased gradient oracle (2).

Lemma 5.1. *Let function f satisfy Assumption 1.2 and biased gradient oracle (2) ($\zeta > 0$) satisfy Assumption 1.5, then Algorithm 1 with step size $\eta_k = \eta \leq [4(L_0 + L_1c)]^{-1}$ guarantees error:*

$$\mathbb{E}[f(x^N)] - f^* \lesssim \left(1 - \frac{\eta c}{R}\right)^K (f(x^0) - f^*) + \frac{R^2}{\eta(N-K)} + \mathcal{R} \cdot \left(\frac{\sigma^2}{B} + \zeta^2\right) + R\zeta,$$

where c is arbitrary clipping radius, $\mathcal{R} = \left(\eta + \frac{MR}{c^2} + \frac{R}{c}\right)$, $0 \leq K < N$ is the number of iterations for which the condition $\|\nabla f(x^k)\| \leq \frac{c}{3}$ is satisfied.

Lemma 5.1 shows how the bias accumulates over iterations, thus converging to the error floor. By estimating the second moment and the bias of the gradient approximation (4) and substituting them into Lemma 5.1, we find three optimality criteria for ZO-ClipSGD.

Theorem 5.2. *Let function f satisfy Assumption 1.2, gradient approximation (4) satisfy Assumption 1.5, then Algorithm 3 with step size $\eta_k = \eta \leq [4(L_0 + L_1c)]^{-1}$ converges to desired ε accuracy after:*

$$N = \mathcal{O}\left(\frac{R}{\eta c} \log \frac{1}{\varepsilon} + \frac{R^2}{\eta \varepsilon}\right); \quad T = \mathcal{O}\left(\frac{d\sigma^2 MR^2}{\varepsilon c^2 \eta} \left(\frac{1}{c} \log \frac{1}{\varepsilon} + \frac{R}{\varepsilon}\right)\right)$$

number of iterations and zero-order oracle calls at

$$\Delta \lesssim \frac{\varepsilon}{\sqrt{dR}(L_0 + L_1M)} \min\left\{\tilde{\sigma}, \frac{\varepsilon}{\sqrt{dR}}\right\}$$

maximum noise level, where $c > 0$ is clipping radius, $\mathbb{E}[\|\nabla f(x, \xi)\|^2] \leq \tilde{\sigma}^2$.

It is not hard to see from Theorem 5.2 that in the generalized smoothness condition, the iteration complexity at ZO-ClipSGD is the same as the first-order method of Algorithm 1. This effect is similar in the standard smoothness condition as well. The oracle complexity is d times worse than its first-order counterpart due to the restriction to the oracle (Algorithm 3 uses only the zero-order oracle (3)). It is worth noting that the maximum noise level Δ outperforms [35], showing that generalized smoothness not only allows us to reach the summands characterizing the linear rate, but also improves the estimate on the maximum noise level (it is Δ that affects the error floor, in other words, the accuracy of the solution, to control the asymptote). See the proof in the Appendix D.

5.2 ZO-NSGD Method

Similar to the first-order algorithms, in this subsection we answer the question whether linear convergence can be achieved by the zero-order method in a convex setup. To answer this question, we consider the Zero-Order Normalized Stochastic Gradient Descent Method.

Algorithm 4 Zero-Order Normalized Stochastic Gradient Descent Method (ZO-NSGD)

Input: initial point $x_0 \in \mathbb{R}^d$, iterations number N , batch size B , step size $\eta_k > 0$

for $k = 0$ **to** $N - 1$ **do**

1. Draw fresh i.i.d. samples $\xi_1^k, \dots, \xi_B^k$ and $e_1^k, \dots, e_B^k$

2. $\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\}) = \frac{1}{B} \sum_{i=1}^B \mathbf{g}(x^k, \{e_i^k, \xi_i^k\})$ via (4)

3. $x^{k+1} \leftarrow x^k - \eta_k \cdot \frac{\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\})}{\|\mathbf{g}(x^k, \{\mathbf{e}^k, \boldsymbol{\xi}^k\})\|}$

end for

Return: x^N

In a similar way as in the previous subsection, we first generalize Theorem 4.1 to the case with a biased gradient oracle to substitute estimates on the bias and the second moment of the gradient approximation to find optimality criteria for the gradient-free algorithm ZO-NSGD.

Lemma 5.3. *Let function f satisfy Assumption 1.2 $((L_0, L_1)$ -smoothness) and biased gradient oracle (2) ($\zeta > 0$) satisfy Assumption 1.5 (bounded variance), then Algorithm 4 with hyperparameter $\lambda > 0$ and step size $\eta_k = \eta \leq \lambda / [2(L_0 + L_1\lambda)]$ guarantees:*

$$\mathbb{E}[f(x^N)] - f^* \lesssim \left(1 - \frac{\eta}{R}\right)^N (f(x^0) - f^*) + \frac{MR}{\lambda^2} \left(\frac{\sigma^2}{B} + \zeta^2\right) + \lambda R.$$

From Lemma 5.3 we can see how the inaccuracy accumulates over the iteration. The summand with ζ^2 is unimprovable for first-order unaccelerated algorithms (see, e.g., [17, 23]). Now, having obtained the results for the biased NSGD we can use them to derive convergence results for Algorithm 4.

Theorem 5.4. *Let function f satisfy Assumption 1.2 $((L_0, L_1)$ -smoothness) and gradient approximation (4) satisfy Assumption 1.5 (bounded variance), then ZO-NSGD with step size $\eta_k = \eta \leq \lambda / [2(L_0 + L_1\lambda)]$ converges to desired ε accuracy ($\mathbb{E}[f(x^N)] - f^* \leq \varepsilon$) after:*

$$N = \mathcal{O}\left(\frac{R}{\eta} \log \frac{1}{\varepsilon}\right); \quad T = \tilde{\mathcal{O}}\left(\frac{d\tilde{\sigma}^2 MR^4}{\varepsilon^3 \eta}\right)$$

number of iterations and zero-order oracle calls (3) at

$$\Delta \lesssim \frac{\varepsilon^{3/2}}{\sqrt{d}R^{3/2}(L_0 + L_1M)} \min\left\{\tilde{\sigma}, \frac{\varepsilon^{3/2}}{\sqrt{d}R^{3/2}}\right\}$$

maximum noise level, where $\lambda > 0$ is hyperparameter, $\mathbb{E}[\|\nabla f(x, \xi)\|^2] \leq \tilde{\sigma}^2$.

It is not hard to see that given a restricted oracle (3), Theorem 5.4 shows that ZO-NSGD requires $N = \mathcal{O}\left(\frac{R}{\eta} \log \frac{1}{\varepsilon}\right)$ iterations to achieve the desired accuracy, which corresponds to a linear rate. However, it is worth noting that, as in Theorem 4.1, if we take the maximum step size $\eta = \lambda / [2(L_0 + L_1\lambda)]$, then the summand with L_0 still shows sublinear convergence $\tilde{\mathcal{O}}\left(\frac{L_0 R^2}{\varepsilon}\right)$, despite the presence of the summand with L_1 . But in the $L_0 = 0$ regime, we can say unambiguously that the zero-order algorithm ZO-NSGD can converge with a linear rate to the desired accuracy in the convex setup if we take the batch size $B = \mathcal{O}\left(\frac{d\tilde{\sigma}^2 MR^3}{\varepsilon^3}\right)$. Theorem 5.4 shows that the power of generalized smoothness (together with Batch parallelization) extends to zero-order algorithms. Comparing the result of ZO-NSGD with Theorem 5.2, we can see that for a significant improvement in iteration complexity N we “pay” for a deterioration in both oracle complexity T and maximum noise level Δ , which seems quite natural. See the proof of Lemma 5.3 and Theorem 5.4 in Appendix E.

6 Discussion and Future Works

In Sections 3-5, we gave two strategies for obtaining summands that characterize the linear rate in convergence estimates of algorithms for the convex setting: clipping (see Section 3) and normalization (see Section 4) techniques. Although these summands are quite unexpected for the convex setup and improve the estimates on the iteration complexity, we cannot claim linear convergence in general, since the convergence is dominated by the summands characterizing the sublinear rate. However, as we noted in Theorems 4.1 and 5.4, in the regime $L_0 = 0$, the NSGD and ZO-NSGD methods break all existing bounds on iteration complexity, demonstrating that it is possible to converge with linear rate to the desired accuracy with the condition of using Batch parallelization. This result is pleasantly surprising and opens up a number of directions for future research.

In this paper we have focused on iteration complexity, so we see a careful analysis of the optimality criterion in the aggregate as future work. In particular, it seems interesting to show that M is indeed bounded, e.g., using the technique from [40] and evaluating with respect to the smoothness constants L_0 and L_1 . The existence of the regime $L_0 = 0$ which allows one to achieve a linear convergence rate in the convex setting prompts the following question: can the iteration complexity be improved by assuming, for example, strong convexity, the PL condition, etc.? It is also interesting to see if similar effects are found in accelerated, adaptive algorithms, variational inequalities, distributed learning, nonsmooth (or increased smoothness) problems, overparameterization, online optimization, etc.

7 Numerical Experiments

In this section, we numerically analyze the algorithms presented in this paper and show that linear convergence in stochastic convex optimization is possible. For this illustration, we have chosen a problem that is of particular interest in the machine learning community: the logistic regression problem on w1a dataset [50]. We consider the following convex problem statement (1):

$$\min_{x \in \mathbb{R}^d} f(x) = \frac{1}{M} \sum_{i=1}^M f_i(x), \quad f_i(x) = \log(1 + \exp(-y_i \cdot (Ax)_i)),$$

where $f_i(x)$ is the loss on the i -th data point, $A \in \mathbb{R}^{M \times d}$ is an instances matrix, $y \in \{-1, 1\}^M$ is a label vector and $x \in \mathbb{R}^d$ is a vector of weights. It is easy to show, that logistic regression function is L -smooth (see Assumption 1.1) with $L = \frac{1}{4M} \sqrt{\lambda_{\max}(A^T A)}$, where $\lambda_{\max}(A^T A)$ denotes the largest eigenvalue of the matrix $A^T A$. Moreover, such a problem statement is a special case of (1) with ξ being a random variable with the uniform distribution on $\{1, \dots, M\}$.

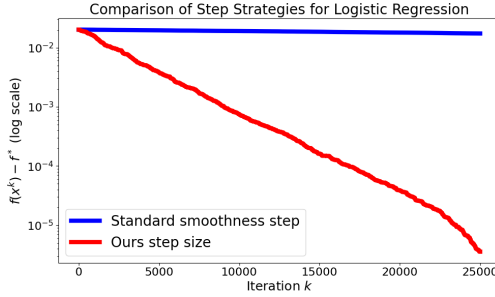


Figure 2: Linear convergence demonstration.

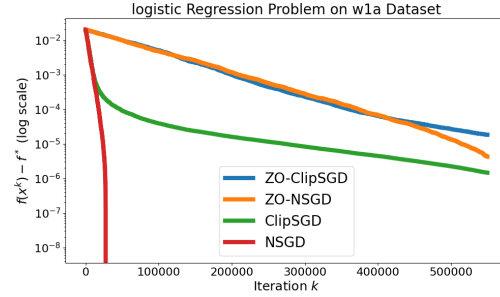


Figure 3: Comparison of the considered algorithms.

In all tests we used the following parameters: $M = 2477$ - number of data, $d = 300$ - problem dimension, $B = 10$ - batch size, $h = 10^{-5}$ - smoothing parameter, $\Delta = 10^{-9}$ - noise level. Figure 2 shows a comparison of two step strategies of the NSGD algorithm. The **blue line** (see "Standard smoothness step"), which corresponds to the theoretical step size in the L -smoothness setting demonstrates slow (sublinear) convergence. In particular, NSGD with this choice of step advanced the function from 0.02014151259 to 0.01731953499 in 25000 iterations. The **red line** (see "Ours step size"), which corresponds to the next step size $\eta = \frac{1}{\|A\|_1}$, demonstrates linear convergence that significantly outperforms the strategy with the theoretical step size. Moreover, *Figure 2 demonstrates that indeed the first-order algorithm can converge with linear rate to the desired accuracy in a convex setup!*

Figure 3 demonstrates the convergence dynamics of all the algorithms considered in this paper. In particular, as in Figure 2, NSGD (see **red line**) shows a real linear convergence. ClipSGD (see **green line**), which used a step size $\eta = \frac{1}{c \cdot \|A\|_1}$, where $c = 10^{-1}$ is the clipping radius, shows two modes of convergence: as long as $\|\nabla f(x^k)\| \geq c$ the algorithm converges with a linear rate matching the NSGD, as soon as $\|\nabla f(x^k)\| < c$ the algorithm slows down to the sublinear rate. The dynamics are similar for zero-order algorithms: ZO-NSGD (see **orange line**), ZO-ClipSGD (see **blue line**). Expectedly, these algorithms converge slower on the first iterations than their first-order counterparts due to restricted access to the oracle (3). However, it is worth noting that *ZO-NSGD also exhibits linear convergence, thereby outperforming the first-order ClipSGD algorithm after 55000 iterations.*

8 Conclusion

In this paper, we considered a stochastic convex optimization problem under the generalized smoothness condition of the objective function. We are the first who have provided strategies to achieve summands characterizing linear rate, thereby improving the iteration complexity (see Sections 3 and 4). In Section 5, we showed that this effect of generalized smoothness extends to zero-order algorithms as well. Moreover, we highlight the regime $L_0 = 0$, under which we theoretically guarantee linear convergence for the NSGD and ZO-NSGD algorithms (subject to the use of batch parallelization). This is the first result demonstrating linear convergence in such a problem setting, thus opening up a number of future works (see Section 6). Finally, in Section 7 we show using numerical experiments that linear convergence in such a problem formulation is also possible in practice.

References

- [1] Alekh Agarwal, Ofer Dekel, and Lin Xiao. Optimal algorithms for online convex optimization with multi-point bandit feedback. In *Colt*, pages 28–40. Citeseer, 2010.
- [2] Arya Akhavan, Massimiliano Pontil, and Alexandre Tsybakov. Distributed zero-order optimization under adversarial noise. *Advances in Neural Information Processing Systems*, 34:10209–10220, 2021.
- [3] Arya Akhavan, Evgenii Chzhen, Massimiliano Pontil, and Alexandre Tsybakov. A gradient estimator via ℓ_1 -randomization for online zero-order optimization with two point feedback. *Advances in Neural Information Processing Systems*, 35:7685–7696, 2022.
- [4] Dan Alistarh, Torsten Hoefer, Mikael Johansson, Nikola Konstantinov, Sarit Khirirat, and Cédric Renggli. The convergence of sparsified gradient methods. *Advances in Neural Information Processing Systems*, 31, 2018.
- [5] Authors Anonymous. Maximum noise level as third optimality criterion in black-box optimization problem, 2025.
- [6] Francis Bach and Vianney Perchet. Highly-smooth zero-th order online optimization. In *Conference on Learning Theory*, pages 257–283. PMLR, 2016.
- [7] Irwan Bello, Barret Zoph, Vijay Vasudevan, and Quoc V Le. Neural optimizer search with reinforcement learning. In *International Conference on Machine Learning*, pages 459–468. PMLR, 2017.
- [8] Yoshua Bengio, Patrice Simard, and Paolo Frasconi. Learning long-term dependencies with gradient descent is difficult. *IEEE transactions on neural networks*, 5(2):157–166, 1994.
- [9] Lev Bogolubsky, Pavel Dvurechenskii, Alexander Gasnikov, Gleb Gusev, Yurii Nesterov, Andrei M Raigorodskii, Aleksey Tikhonov, and Maksim Zhukovskii. Learning supervised pagerank with gradient-based and gradient-free optimization methods. *Advances in neural information processing systems*, 29, 2016.
- [10] Léon Bottou. Online algorithms and stochastic approximations. *Online learning in neural networks*, 1998.
- [11] Léon Bottou, Frank E Curtis, and Jorge Nocedal. Optimization methods for large-scale machine learning. *SIAM review*, 60(2):223–311, 2018.
- [12] Ziyi Chen, Yi Zhou, Yingbin Liang, and Zhaosong Lu. Generalized-smooth nonconvex optimization is as efficient as smooth nonconvex optimization. In *International Conference on Machine Learning*, pages 5396–5427. PMLR, 2023.
- [13] Ziyi Chen, Yi Zhou, Yingbin Liang, and Zhaosong Lu. Generalized-smooth nonconvex optimization is as efficient as smooth nonconvex optimization. In *International Conference on Machine Learning*, pages 5396–5427. PMLR, 2023.
- [14] Ashok Cutkosky and Harsh Mehta. Momentum improves normalized sgd. In *International conference on machine learning*, pages 2260–2268. PMLR, 2020.
- [15] Jeffrey Dean, Greg Corrado, Rajat Monga, Kai Chen, Matthieu Devin, Mark Mao, Marc’aurélio Ranzato, Andrew Senior, Paul Tucker, Ke Yang, et al. Large scale distributed deep networks. *Advances in neural information processing systems*, 25, 2012.
- [16] Yury Demidovich, Grigory Malinovsky, Igor Sokolov, and Peter Richtárik. A guide through the zoo of biased sgd. *Advances in Neural Information Processing Systems*, 36:23158–23171, 2023.
- [17] Olivier Devolder. Exactness, inexactness and stochasticity in first-order methods for large-scale convex optimization. *Candidate’s Dissertation (CORE UCLouvain Louvain-la-Neuve, Belgium)*, 2013.
- [18] Tolga Dimlioglu and Anna Choromanska. Grawa: Gradient-based weighted averaging for distributed training of deep learning models. In *International Conference on Artificial Intelligence and Statistics*, pages 2251–2259. PMLR, 2024.
- [19] Matthew Faw, Litu Rout, Constantine Caramanis, and Sanjay Shakkottai. Beyond uniform smoothness: A stopped analysis of adaptive sgd. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 89–160. PMLR, 2023.
- [20] Ofir Gaash, Kfir Yehuda Levy, and Yair Carmon. Convergence of clipped sgd on convex (ℓ_0, ℓ_1) -smooth functions. *arXiv preprint arXiv:2502.16492*, 2025.

- [21] Alexander Gasnikov, Anton Novitskii, Vasilii Novitskii, Farshed Abdukhakimov, Dmitry Kamzolov, Aleksandr Beznosikov, Martin Takac, Pavel Dvurechensky, and Bin Gu. The power of first-order smooth optimization for black-box non-smooth problems. In *International Conference on Machine Learning*, pages 7241–7265. PMLR, 2022.
- [22] Alexander Gasnikov, Darina Dvinskikh, Pavel Dvurechensky, Eduard Gorbunov, Aleksandr Beznosikov, and Alexander Lobanov. Randomized gradient-free methods in convex optimization. In *Encyclopedia of Optimization*, pages 1–15. Springer, 2023.
- [23] AV Gasnikov, AV Lobanov, and FS Stonyakin. Highly smooth zeroth-order methods for solving optimization problems under the pl condition. *Computational Mathematics and Mathematical Physics*, 64(4): 739–770, 2024.
- [24] Ian Goodfellow. Deep learning, 2016.
- [25] Eduard Gorbunov, Marina Danilova, and Alexander Gasnikov. Stochastic optimization with heavy-tailed noise via accelerated gradient clipping. *Advances in Neural Information Processing Systems*, 33:15042–15053, 2020.
- [26] Eduard Gorbunov, Nazarii Tupitsa, Sayantan Choudhury, Alen Aliev, Peter Richtárik, Samuel Horváth, and Martin Takáč. Methods for convex (l_0, l_1) -smooth optimization: Clipping, acceleration, and adaptivity. *arXiv preprint arXiv:2409.14989*, 2024.
- [27] Elad Hazan, Kfir Levy, and Shai Shalev-Shwartz. Beyond convexity: Stochastic quasi-convex optimization. *Advances in neural information processing systems*, 28, 2015.
- [28] José Miguel Hernández-Lobato, Matthew W Hoffman, and Zoubin Ghahramani. Predictive entropy search for efficient global optimization of black-box functions. *Advances in neural information processing systems*, 27, 2014.
- [29] Yusu Hong and Junhong Lin. On convergence of adam for stochastic optimization under relaxed assumptions. *arXiv preprint arXiv:2402.03982*, 2024.
- [30] Florian Hübler, Ilyas Fatkhullin, and Niao He. From gradient clipping to normalization for heavy tailed sgd. *arXiv preprint arXiv:2410.13849*, 2024.
- [31] Florian Hübler, Junchi Yang, Xiang Li, and Niao He. Parameter-agnostic optimization under relaxed smoothness. In *International Conference on Artificial Intelligence and Statistics*, pages 4861–4869. PMLR, 2024.
- [32] Anatoli B Juditsky and Arkadii S Nemirovski. First order methods for nonsmooth convex large-scale optimization, i: General purpose methods. *Optimization for Machine Learning*, pages 1–28, 2010.
- [33] Peter Kairouz, H Brendan McMahan, Brendan Avent, Aurélien Bellet, Mehdi Bennis, Arjun Nitin Bhagoji, Kallista Bonawitz, Zachary Charles, Graham Cormode, Rachel Cummings, et al. Advances and open problems in federated learning. *Foundations and trends® in machine learning*, 14(1–2):1–210, 2021.
- [34] Anastasia Koloskova, Hadrien Hendrikx, and Sebastian U Stich. Revisiting gradient clipping: Stochastic bias and tight convergence guarantees. In *International Conference on Machine Learning*, pages 17343–17363. PMLR, 2023.
- [35] Nikita Kornilov, Ohad Shamir, Aleksandr Lobanov, Darina Dvinskikh, Alexander Gasnikov, Innokentii Shibaev, Eduard Gorbunov, and Samuel Horváth. Accelerated zeroth-order method for non-smooth stochastic convex optimization problem with infinite variance. *Advances in Neural Information Processing Systems*, 36:64083–64102, 2023.
- [36] Simon Lacoste-Julien, Mark Schmidt, and Francis Bach. A simpler approach to obtaining an $\mathcal{O}(1/t)$ convergence rate for the projected stochastic subgradient method. *arXiv preprint arXiv:1212.2002*, 2012.
- [37] Guanghui Lan. An optimal method for stochastic composite optimization. *Mathematical Programming*, 133(1):365–397, 2012.
- [38] Tor Lattimore and Andras Gyorgy. Improved regret for zeroth-order stochastic convex bandits. In *Conference on Learning Theory*, pages 2938–2964. PMLR, 2021.
- [39] Hojoon Lee, Hanseul Cho, Hyunseung Kim, Daehoon Gwak, Joonkee Kim, Jaegul Choo, Se-Young Yun, and Chulhee Yun. Plastic: Improving input and label plasticity for sample efficient reinforcement learning. *Advances in Neural Information Processing Systems*, 36, 2024.

- [40] Haochuan Li, Jian Qian, Yi Tian, Alexander Rakhlin, and Ali Jadbabaie. Convex and non-convex optimization under generalized smoothness. *Advances in Neural Information Processing Systems*, 36: 40238–40271, 2023.
- [41] Haochuan Li, Alexander Rakhlin, and Ali Jadbabaie. Convergence of adam under relaxed assumptions. *Advances in Neural Information Processing Systems*, 36:52166–52196, 2023.
- [42] Aleksandr Lobanov and Alexander Gasnikov. Accelerated zero-order sgd method for solving the black box optimization problem under “overparametrization” condition. In *International Conference on Optimization and Applications*, pages 72–83. Springer, 2023.
- [43] Aleksandr Lobanov, Nail Bashirov, and Alexander Gasnikov. The “black-box” optimization problem: Zero-order accelerated stochastic method via kernel approximation. *Journal of Optimization Theory and Applications*, pages 1–36, 2024.
- [44] Aleksandr Lobanov, Alexander Gasnikov, Eduard Gorbunov, and Martin Takáč. Linear convergence rate in convex setup is possible! gradient descent method variants under (l_0, l_1) -smoothness. *arXiv preprint arXiv:2412.17050*, 2024.
- [45] Yurii Nesterov. *Lectures on convex optimization*, volume 137. Springer, 2018.
- [46] Yurii Nesterov and Vladimir Spokoiny. Random gradient-free minimization of convex functions. *Foundations of Computational Mathematics*, 17(2):527–566, 2017.
- [47] Anthony Nguyen and Krishnakumar Balasubramanian. Stochastic zeroth-order functional constrained optimization: Oracle complexity and applications. *INFORMS Journal on Optimization*, 2022.
- [48] Razvan Pascanu, Tomas Mikolov, and Yoshua Bengio. On the difficulty of training recurrent neural networks. 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 1310–1318, Atlanta, Georgia, USA, 17–19 Jun 2013. PMLR.
- [49] Kumar Kshitij Patel, Aadirupa Saha, Lingxiao Wang, and Nathan Srebro. Distributed online and bandit convex optimization. In *OPT 2022: Optimization for Machine Learning (NeurIPS 2022 Workshop)*, 2022.
- [50] John C Platt. Fast training of support vector machines using sequential minimal optimization. *Advances in Kernel Methods - Support Vector Learning*, MIT Press, 1998.
- [51] Boris T Polyak. *Introduction to optimization*. Optimization Software, Inc. Publications Division, New York, 1987.
- [52] Ohad Shamir. An optimal algorithm for bandit and zero-order convex optimization with two-point feedback. *The Journal of Machine Learning Research*, 18(1):1703–1713, 2017.
- [53] Sebastian U Stich and Sai Praneeth Karimireddy. The error-feedback framework: Better rates for sgd with delayed gradients and compressed communication. *arXiv preprint arXiv:1909.05350*, 2019.
- [54] Yuki Takezawa, Han Bao, Ryoma Sato, Kenta Niwa, and Makoto Yamada. Parameter-free clipped gradient descent meets polyak. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [55] Daniil Vankov, Anton Rodomanov, Angelia Nedich, Lalitha Sankar, and Sebastian U Stich. Optimizing (l_0, l_1) -smooth functions by gradient methods. *arXiv preprint arXiv:2410.10800*, 2024.
- [56] Bohan Wang, Huishuai Zhang, Zhiming Ma, and Wei Chen. Convergence of adagrad for non-convex objectives: Simple proofs and relaxed assumptions. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 161–190. PMLR, 2023.
- [57] Bohan Wang, Huishuai Zhang, Qi Meng, Ruoyu Sun, Zhi-Ming Ma, and Wei Chen. On the convergence of adam under non-uniform smoothness: Separability from sgd and beyond. *arXiv preprint arXiv:2403.15146*, 2024.
- [58] Blake E Woodworth, Brian Bullins, Ohad Shamir, and Nathan Srebro. The min-max complexity of distributed stochastic convex optimization with intermittent communication. In *Conference on Learning Theory*, pages 4386–4437. PMLR, 2021.
- [59] Chenghan Xie, Chenxi Li, Chuwen Zhang, Qi Deng, Dongdong Ge, and Yinyu Ye. Trust region methods for nonconvex stochastic optimization beyond lipschitz smoothness. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 16049–16057, 2024.

- [60] Honglin Yuan and Tengyu Ma. Federated accelerated stochastic gradient descent. *Advances in Neural Information Processing Systems*, 33:5332–5344, 2020.
- [61] Bohang Zhang, Jikai Jin, Cong Fang, and Liwei Wang. Improved analysis of clipping algorithms for non-convex optimization. *Advances in Neural Information Processing Systems*, 33:15511–15521, 2020.
- [62] Jingzhao Zhang, Tianxing He, Suvrit Sra, and Ali Jadbabaie. Why gradient clipping accelerates training: A theoretical justification for adaptivity. In *International Conference on Learning Representations*, 2019.
- [63] Jingzhao Zhang, Sai Praneeth Karimireddy, Andreas Veit, Seungyeon Kim, Sashank Reddi, Sanjiv Kumar, and Suvrit Sra. Why are adaptive methods good for attention models? *Advances in Neural Information Processing Systems*, 33:15383–15393, 2020.
- [64] Sixin Zhang, Anna E Choromanska, and Yann LeCun. Deep learning with elastic averaging sgd. *Advances in neural information processing systems*, 28, 2015.
- [65] Shen-Yi Zhao, Yin-Peng Xie, and Wu-Jun Li. On the convergence and improvement of stochastic normalized gradient descent. *Science China Information Sciences*, 64:1–13, 2021.
- [66] Shen-Yi Zhao, Chang-Wei Shi, Yin-Peng Xie, and Wu-Jun Li. Stochastic normalized gradient descent with momentum for large-batch training. *Science China Information Sciences*, 67(11):212101, 2024.
- [67] Vladimir Antonovich Zorich and Octavio Paniagua. *Mathematical analysis II*, volume 220. Springer, 2016.

APPENDIX

Power of Generalized Smoothness in Stochastic Convex Optimization: First- and Zero-Order Algorithms

A Auxiliary Results

In this section we provide auxiliary materials that are used in the proof of Theorems.

A.1 Basic inequalities and assumptions

Basic inequalities. For all $a, b \in \mathbb{R}^d$ ($d \geq 1$) the following equality holds:

$$2 \langle a, b \rangle - \|b\|^2 = \|a\|^2 - \|a - b\|^2, \quad (5)$$

$$\langle a, b \rangle \leq \|a\| \cdot \|b\|. \quad (6)$$

Squared norm of the sum For all $a_1, \dots, a_n \in \mathbb{R}^d$, where $n = \{2, 3\}$

$$\|a_1 + \dots + a_n\|_2^2 \leq n\|a_1\|_2^2 + \dots + n\|a_n\|_2^2. \quad (7)$$

Generalized-Lipschitz-smoothness. Throughout this paper, we assume that the (L_0, L_1) -smoothness condition (Assumption 1.2) is satisfied. This inequality can be represented in the equivalent form for any $x, y \in \mathbb{R}^d$:

$$f(y) - f(x) \leq \langle \nabla f(x), y - x \rangle + \frac{L_0 + L_1 \|\nabla f(x)\|}{2} \|y - x\|^2, \quad (8)$$

where $L_0, L_1 \geq 0$ for any $x \in \mathbb{R}^d$ and $\|y - x\| \leq \frac{1}{L_1}$.

Variance decomposition. If ξ is random vector in $\mathbb{R}^d$ with bounded second moment, then

$$\mathbb{E} [\|\xi + a\|^2] = \mathbb{E} [\|\xi - \mathbb{E}[\xi]\|^2] + \mathbb{E} [\|\mathbb{E}[\xi] - a\|^2], \quad (9)$$

for any deterministic vector $a \in \mathbb{R}^d$.

A.2 Auxiliary Lemma about Generalized Smoothness

If Assumption 1.2 holds, then it also holds that $\forall x \in \mathbb{R}^d$:

$$\|\nabla f(x)\|^2 \leq 2(L_0 + L_1 \|\nabla f(x)\|)(f(x) - f^*), \quad (10)$$

where $f^* = \inf_x f(x)$.

Proof. We start the proof by applying (8) for $y = x - \frac{1}{L_0 + L_1 \|\nabla f(x)\|} \nabla f(x)$, where $\|y - x\| = \frac{\|\nabla f(x)\|}{L_0 + L_1 \|\nabla f(x)\|} \leq \frac{1}{L_1}$. Then we can obtain:

$$f^* \leq f \left(x - \frac{1}{L_0 + L_1 \|\nabla f(x)\|} \nabla f(x) \right) \stackrel{(8)}{\leq} f(x) - \frac{1}{2(L_0 + L_1 \|\nabla f(x)\|)} \|\nabla f(x)\|^2.$$

□

A.3 Wirtinger-Poincare inequality

Let f is differentiable, then for all $x \in \mathbb{R}^d, \gamma e \in S^d(\gamma)$:

$$\mathbb{E} [f(x + \gamma e)^2] \leq \frac{\gamma^2}{d} \mathbb{E} [\|\nabla f(x + \gamma e)\|^2]. \quad (11)$$

B Clipped Stochastic Gradient Descent (Proof of the Theorem 3.1)

We start by using (L_0, L_1) -smoothness (see Assumption 1.2):

$$\begin{aligned} f(x^{k+1}) - f(x^k) &\stackrel{(8)}{\leq} \langle \nabla f(x^k), x^{k+1} - x^k \rangle + \frac{L_0 + L_1 \|\nabla f(x^k)\|}{2} \|x^{k+1} - x^k\|^2 \\ &= -\eta \langle \nabla f(x^k), \text{clip}_c(\nabla f(x^k, \xi^k)) \rangle \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2. \end{aligned} \quad (12)$$

Next, we consider three cases depending on the gradient norm: $\|\nabla f(x^k)\| \geq c$ – the full gradient is clipped and $\frac{c}{2} \leq \|\nabla f(x^k)\| \leq c$ and $\|\nabla f(x^k)\| \leq \frac{c}{2}$ – the full gradient is not clipped.

B.1 First case: $\|\nabla f(x^k)\| \geq c$

In this case $\alpha \nabla f(x^k) = \text{clip}_c(\nabla f(x^k))$ with $\alpha = \min \left\{ 1, \frac{c}{\|\nabla f(x^k)\|} \right\} = \frac{c}{\|\nabla f(x^k)\|}$, therefore we have the following

$$\begin{aligned} -\eta \langle \nabla f(x^k), \text{clip}_c(\nabla f(x^k, \xi^k)) \rangle &\stackrel{(5)}{\leq} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \\ &\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k)) - \alpha \nabla f(x^k)\|^2 \\ &= -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \\ &\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k)) - \text{clip}_c(\nabla f(x^k))\|^2 \\ &= -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \\ &\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k)) - \text{clip}_c(\nabla f(x^k))\|^2. \end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius c , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned} -\eta \langle \nabla f(x^k), \mathbb{E} [\text{clip}_c(\nabla f(x^k, \xi^k))] \rangle &\leq -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\quad + \frac{\eta}{2\alpha} \mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\ &\leq -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\quad + \frac{\eta\sigma^2}{2\alpha B} \\ &= -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\quad + \frac{\eta \|\nabla f(x^k)\| \sigma^2}{2cB}. \end{aligned} \quad (13)$$

We now consider the cases depending on the relation between c and σ :

In the case $c \geq \sqrt{2}\sigma$ We have in (13):

$$\begin{aligned} -\eta \langle \nabla f(x^k), \mathbb{E} [\text{clip}_c(\nabla f(x^k, \xi^k))] \rangle &\stackrel{(13)}{\leq} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\quad + \frac{\eta \|\nabla f(x^k)\| \sigma^2}{2cB} \\ &= -\frac{\eta}{2\alpha} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \end{aligned}$$

$$\begin{aligned}
& -\frac{c\eta}{2} \|\nabla f(x^k)\| \left(1 - \frac{\sigma^2}{c^2 B}\right) \\
& \leq -\frac{\eta}{2\alpha} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
& \quad -\frac{c\eta}{4} \|\nabla f(x^k)\| \\
& = -\frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
& \quad -\frac{c\eta}{4} \|\nabla f(x^k)\|.
\end{aligned}$$

Plugging this into (12) and choosing $\eta \leq \frac{1}{4(L_0 + L_1 c)}$ we have:

$$\begin{aligned}
\mathbb{E} [\nabla f(x^{k+1})] - f(x^k) & \stackrel{(12)}{\leq} -\frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] - \frac{c\eta}{4} \|\nabla f(x^k)\| \\
& \quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
& = -\frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] (1 - \eta L_1 c) \\
& \quad -\frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta^2 L_0}{2} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
& \leq -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] (1 - \eta(L_0 + L_1 c)) \\
& \leq -\frac{c\eta}{4} \|\nabla f(x^k)\|. \tag{14}
\end{aligned}$$

Using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f^* & \leq \langle \nabla f(x^k), x^k - x^* \rangle \\
& \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\
& \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{15}$$

Then substituting (15) into (14) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta c}{4} \|\nabla f(x^k)\| \leq -\frac{\eta c}{4R} (f(x^k) - f^*).$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta c}{4R}\right) (f(x^k) - f^*).$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \geq c \geq \sqrt{2}\sigma$, then ClipSGD has linear convergence

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*).$$

In the case $c \leq \sqrt{2}\sigma$ We have in (13):

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \mathbb{E} [\text{clip}_c(f(x^k, \xi^k))] \rangle & \stackrel{(13)}{\leq} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
& \quad + \frac{\eta \|\nabla f(x^k)\| \sigma^2}{2cB}
\end{aligned}$$

$$\begin{aligned}
&= -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
&\quad + \frac{\eta M \sigma^2}{2cB}.
\end{aligned}$$

Plugging this into (12) and choosing $\eta \leq \frac{1}{4(L_0 + L_1 c)}$ we have:

$$\begin{aligned}
\mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(12)}{\leq} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] + \frac{\eta M \sigma^2}{2cB} \\
&= -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] (1 - \eta L_1 c) \\
&\quad + \frac{\eta^2 L_0}{2} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] + \frac{\eta M \sigma^2}{2cB} \\
&\leq -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} \left[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \right] (1 - \eta(L_0 + L_1 c)) \\
&\quad + \frac{\eta M \sigma^2}{2cB} \\
&\leq -\frac{c\eta}{2} \|\nabla f(x^k)\| + \frac{\eta M \sigma^2}{2cB}. \tag{16}
\end{aligned}$$

Using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f^* &\leq \langle \nabla f(x^k), x^k - x^* \rangle \\
&\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\
&\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{17}$$

Then substituting (17) into (16) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta c}{2} \|\nabla f(x^k)\| + \frac{\eta M \sigma^2}{2cB} \leq -\frac{\eta c}{2R} (f(x^k) - f^*) + \frac{\eta M \sigma^2}{2cB}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta c}{2R}\right) (f(x^k) - f^*) + \frac{\eta M \sigma^2}{2cB}.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \geq c$ and $c \leq \sqrt{2}\sigma$, then ClipSGD has linear convergence

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta c}{2R}\right)^N (f(x^0) - f^*) + \frac{MR\sigma^2}{c^2 B}.$$

B.2 Second case: $\frac{c}{2} \leq \|\nabla f(x^k)\| \leq c$

In this case $\nabla f(x^k) = \text{clip}_c(\nabla f(x^k))$ with $\alpha = \min \left\{ 1, \frac{c}{\|\nabla f(x^k)\|} \right\} = 1$, therefore we have the following

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \text{clip}_c(\nabla f(x^k, \xi^k)) \rangle &\stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\nabla f(x^k, \xi^k)) - \alpha \nabla f(x^k)\|^2
\end{aligned}$$

$$\begin{aligned}
&= -\frac{\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \\
&\quad + \frac{\eta}{2} \|\text{clip}_c(\nabla f(x^k, \xi^k)) - \text{clip}_c(\nabla f(x^k))\|^2 \\
&\leq -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2 \\
&\quad + \frac{\eta}{2} \|\text{clip}_c(\nabla f(x^k, \xi^k)) - \text{clip}_c(\nabla f(x^k))\|^2.
\end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius c , and thus is Lipshitz operator with Lipshitz constant 1, we can obtain:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \mathbb{E} [\text{clip}_c(\nabla f(x^k, \xi^k))] \rangle &\leq -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\
&\quad + \frac{\eta}{2} \mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\
&\leq -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] + \frac{\eta\sigma^2}{2B} \\
&= -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] + \frac{\eta\sigma^2}{2B}.
\end{aligned}$$

Plugging this into (12) and choosing $\eta \leq \frac{1}{4(L_0 + L_1 c)}$ we have:

$$\begin{aligned}
\mathbb{E} [f(x^{k+1})] - f(x^k) &\stackrel{(12)}{\leq} -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] + \frac{\eta\sigma^2}{2B} \\
&= -\frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2B} \\
&\quad - \frac{\eta}{2} \mathbb{E} [\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] (1 - \eta(L_0 + L_1 \|\nabla f(x^k)\|)) \\
&\leq -\frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2B}. \tag{18}
\end{aligned}$$

Using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f^* &\leq \langle \nabla f(x^k), x^k - x^* \rangle \\
&\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\
&\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{19}$$

Then substituting (19) into (18) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta c}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2B} \leq -\frac{\eta c}{4R} (f(x^k) - f^*) + \frac{\eta\sigma^2}{2B}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta c}{4R}\right) (f(x^k) - f^*) + \frac{\eta\sigma^2}{2B}.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\frac{c}{2} \leq \|\nabla f(x^k)\| \leq c$, then ClipSGD has linear convergence

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta c}{4R}\right)^N (f(x^0) - f^*) + \frac{2\sigma^2 R}{cB}.$$

Let $\mathcal{T}_1 = \{m_0^{\mathcal{T}_1}, m_1^{\mathcal{T}_1}, m_2^{\mathcal{T}_1}, \dots, m_{K-1}^{\mathcal{T}_1}\} = \{k \in \{0, 1, 2, \dots, N-1\} \mid \|\nabla f(x^k, \xi^k)\| \geq \frac{c}{2}\}$, where $K = |\mathcal{T}_1|$. Then for $k \in \mathcal{T}_1$ ClipSGD shows linear convergence:

$$\begin{aligned} F_N \cdot \mathbb{1}[\mathcal{T}_1] &\leq \left(1 - \frac{1}{4L_1R}\right) F_N \leq \left(1 - \frac{1}{4L_1R}\right) F_{m_{K-1}^{\mathcal{T}_1}} + \frac{\eta M \sigma^2}{2cB} + \frac{\eta \sigma^2}{2B} \\ &\leq \dots \leq \left(1 - \frac{1}{4L_1R}\right)^K F_{m_0^{\mathcal{T}_1}} + \frac{MR\sigma^2}{c^2B} + \frac{2\sigma^2R}{cB} \\ &\leq \left(1 - \frac{1}{4L_1R}\right)^K F_0 + \frac{MR\sigma^2}{c^2B} + \frac{2\sigma^2R}{cB}, \end{aligned}$$

where $F_k = \mathbb{E}[f(x^k)] - f^*$, and we used that $F_k \leq F_{k-1}$.

B.3 Third case: $\|\nabla f(x^k)\| \leq \frac{c}{2}$

We introduce an indicative function:

$$\aleph_k = \mathbb{1}\{\|\nabla f(x^k, \xi^k)\| > c\}. \quad (20)$$

Then the following is true:

$$\mathbb{E}[\aleph_k] = \mathbb{E}[\aleph_k^2] = \mathcal{P}[\|\nabla f(x^k, \xi^k)\| > c] \stackrel{\textcircled{1}}{\leq} \mathcal{P}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\| > \frac{c}{2}] \stackrel{\textcircled{2}}{\leq} \frac{4\sigma^2}{c^2B}, \quad (21)$$

where in $\textcircled{1}$ we used $\|\nabla f(x^k, \xi^k)\| \leq \|\nabla f(x^k, \xi^k) - \nabla f(x^k)\| + \|\nabla f(x^k)\| \leq \|\nabla f(x^k, \xi^k) - \nabla f(x^k)\| + \frac{c}{2}$, and in $\textcircled{2}$ we used Markov's inequality.

Let $r_{k+1} = \mathbb{E}[\|x^{k+1} - x^*\|]$ and $F_{k+1} = \mathbb{E}[f(x^{k+1}) - f^*]$, then given that

$$\begin{aligned} \text{clip}_c(\nabla f(x^k, \xi^k)) &= \nabla f(x^k, \xi^k)(1 - \aleph_k) + \frac{c}{\|\nabla f(x^k, \xi^k)\|} \nabla f(x^k, \xi^k) \aleph_k \\ &= \nabla f(x^k, \xi^k) + \left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1\right) \nabla f(x^k, \xi^k) \aleph_k \end{aligned}$$

we get with $\eta \leq \frac{1}{4(L_0 + L_1c)}$:

$$\begin{aligned} r_{k+1}^2 &= r_k^2 - 2\eta \langle \mathbb{E}[\text{clip}_c(\nabla f(x^k, \xi^k))], x^k - x^* \rangle + \eta^2 \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &= r_k^2 - 2\eta \langle \nabla f(x^k), x^k - x^* \rangle - 2\eta \left\langle \mathbb{E}\left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1\right) \nabla f(x^k, \xi^k) \aleph_k\right], x^k - x^* \right\rangle \\ &\quad + \eta^2 \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\stackrel{(6)}{\leq} r_k^2 - 2\eta \langle \nabla f(x^k), x^k - x^* \rangle + 2\eta \left\| \mathbb{E}\left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1\right) \nabla f(x^k, \xi^k) \aleph_k\right] \right\| \|x^k - x^*\| \\ &\quad + \eta^2 \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\stackrel{\textcircled{1}}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E}\left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1\right) \nabla f(x^k, \xi^k) \aleph_k\right] \right\| \|x^0 - x^*\| \\ &\quad + \eta^2 \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \xi^k))\|^2] \\ &\stackrel{(7)}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E}\left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1\right) \nabla f(x^k, \xi^k) \aleph_k\right] \right\| \|x^0 - x^*\| \\ &\quad + 2\eta^2 \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \xi^k)) - \nabla f(x^k)\|^2] + 2\eta^2 \|\nabla f(x^k)\|^2 \\ &= r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E}\left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1\right) \nabla f(x^k, \xi^k) \aleph_k\right] \right\| R \\ &\quad + 2\eta^2 \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \xi^k)) - \text{clip}_c(\nabla f(x^k))\|^2] + 2\eta^2 \|\nabla f(x^k)\|^2 \end{aligned}$$

$$\begin{aligned}
& \stackrel{\textcircled{2}}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R \\
& \quad + 2\eta^2 \mathbb{E} \left[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2 \right] + 2\eta^2 \|\nabla f(x^k)\|^2 \\
& \leq r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R \\
& \quad + \frac{2\eta^2 \sigma^2}{B} + 2\eta^2 \|\nabla f(x^k)\|^2 \\
& \stackrel{(10)}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R \\
& \quad + \frac{2\eta^2 \sigma^2}{B} + 4\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|) F_k \\
& \leq r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R \\
& \quad + \frac{2\eta^2 \sigma^2}{B} + 4\eta^2 (L_0 + L_1 c) F_k \\
& = r_k^2 - 2\eta F_k (1 - 2\eta (L_0 + L_1 c)) + \frac{2\eta^2 \sigma^2}{B} \\
& \quad + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R \\
& \leq r_k^2 - \eta F_k + \frac{2\eta^2 \sigma^2}{B} + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R. \tag{22}
\end{aligned}$$

Let's find the upper bound of the last summand:

$$\begin{aligned}
& 2\eta R \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| \\
& \stackrel{(20)}{\leq} 2\eta R \mathbb{E} \left[\|\nabla f(x^k, \xi^k)\| \cdot \left(1 - \frac{c}{\|\nabla f(x^k, \xi^k)\|} \right) \aleph_k \right] \\
& \leq 2\eta R \mathbb{E} [\|\nabla f(x^k, \xi^k)\| \cdot \aleph_k] \\
& \leq 2\eta R (\mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\| \cdot \aleph_k] + \|\nabla f(x^k)\| \mathbb{E} [\aleph_k]) \\
& \leq 2\eta R \left(\sqrt{\mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \cdot \mathbb{E} [\aleph_k^2]} + \|\nabla f(x^k)\| \mathbb{E} [\aleph_k] \right) \\
& \stackrel{(21)}{\leq} 2\eta R \left(\frac{2\sigma^2}{cB} + \frac{c}{2} \cdot \frac{4\sigma^2}{c^2 B} \right) \\
& = \frac{8\eta\sigma^2 R}{cB}. \tag{23}
\end{aligned}$$

Substituting into the initial formula and rearrange the summands, we obtain

$$\begin{aligned}
\eta F_k & \stackrel{(22)}{\leq} r_k^2 - r_{k+1}^2 + \frac{2\eta^2 \sigma^2}{B} + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \aleph_k \right] \right\| R \\
& \stackrel{(23)}{\leq} r_k^2 - r_{k+1}^2 + \frac{2\eta^2 \sigma^2}{B} + \frac{8\eta\sigma^2 R}{cB}
\end{aligned}$$

Let $\mathcal{T}_2 = \{m_0^{\mathcal{T}_2}, m_1^{\mathcal{T}_2}, m_2^{\mathcal{T}_2}, \dots, m_{N-K}^{\mathcal{T}_2}\} = \{k \in \{0, 1, 2, \dots, N-1\} \mid \|\nabla f(x^k)\| < \frac{c}{2}\}$, where $|\mathcal{T}_2| = N - K$. Then rearranging and summing over all $k \in \mathcal{T}_2$ we obtain

$$F_N \cdot \mathbb{1}[\mathcal{T}_2] \leq \frac{1}{N-K} \sum_{k \in \mathcal{T}_2} F_k \leq \frac{1}{\eta(N-K)} \sum_{k \in \mathcal{T}_2} (r_k^2 - r_{k+1}^2) + \frac{1}{N-K} \sum_{k \in \mathcal{T}_2} \left(\frac{2\eta\sigma^2}{B} + \frac{8\sigma^2 R}{cB} \right)$$

$$= \frac{r_0^2 - r_N^2}{\eta(N-K)} + \frac{2\eta\sigma^2}{B} + \frac{8\sigma^2 R}{cB} \leq \frac{r_0^2}{\eta(N-K)} + \frac{2\eta\sigma^2}{B} + \frac{8\sigma^2 R}{cB}.$$

Hence we obtain:

$$F_N \cdot \mathbb{1}[\mathcal{T}_2] \leq \frac{R^2}{\eta(N-K)} + \frac{2\eta\sigma^2}{B} + \frac{8\sigma^2 R}{cB}.$$

Combining all cases we have:

$$\begin{aligned} \mathbb{E}[f(x^N)] - f^* &\leq F_N \cdot \mathbb{1}[\mathcal{T}_1] + F_N \cdot \mathbb{1}[\mathcal{T}_2] \\ &\leq \left(1 - \frac{\eta c}{2R}\right)^K F_0 + \frac{R^2}{\eta(N-K)} + \frac{\sigma^2 MR}{c^2 B} + \frac{2\eta\sigma^2}{B} + \frac{8\sigma^2 R}{cB}. \end{aligned}$$

C Normalized Stochastic Gradient Descent (Proof of the Theorem 4.1)

Let's introduce the notation $G(x^k, \xi^k) = \frac{\nabla f(x^k, \xi^k)}{\|\nabla f(x^k, \xi^k)\|}$, then using (L_0, L_1) -smoothness (see Assumption 1.2):

$$\begin{aligned} f(x^{k+1}) - f(x^k) &\stackrel{(8)}{\leq} \langle \nabla f(x^k), x^{k+1} - x^k \rangle + \frac{L_0 + L_1 \|\nabla f(x^k)\|}{2} \|x^{k+1} - x^k\|^2 \\ &= -\eta \langle \nabla f(x^k), G(x^k, \xi^k) \rangle + \frac{\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|)}{2} \|G(x^k, \xi^k)\|^2. \quad (24) \end{aligned}$$

Next, we consider 4 cases of the relation $\|\nabla f(x^k)\|$ and $\|\nabla f(x^k, \xi^k)\|$ with respect to the hyperparameter λ .

C.1 First case: $\|\nabla f(x^k)\| \geq \lambda$ and $\|\nabla f(x^k, \xi^k)\| \geq \lambda$

Let us evaluate first summand of (24) with $\alpha = \|\nabla f(x^k)\|^{-1}$:

$$\begin{aligned} -\eta \langle \nabla f(x^k), G(x^k, \xi^k) \rangle &\stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\ &\quad + \frac{\eta}{2\alpha} \|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2 \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\ &\quad + \frac{\eta}{2\lambda^2\alpha} \|\lambda G(x^k, \xi^k) - \lambda \alpha \nabla f(x^k)\|^2 \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\ &\quad + \frac{\eta}{2\lambda^2\alpha} \|\text{clip}_\lambda(\nabla f(x^k, \xi^k)) - \text{clip}_\lambda(\nabla f(x^k))\|^2 \end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius λ , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &\quad + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2]. \quad (25) \end{aligned}$$

In the case: $0 \leq \sigma \leq \frac{\lambda}{\sqrt{2}}$. Using this in (25), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\begin{aligned} \mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(24)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle + \frac{\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &\stackrel{(25)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\ &\quad + \frac{\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \end{aligned}$$

$$\begin{aligned}
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E} \left[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2 \right] \\
&\quad - \frac{\eta}{2} \mathbb{E} \left[\|G(x^k, \xi^k)\|^2 \right] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|} \right) \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda^2\alpha} \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{4} \|\nabla f(x^k)\| \\
&= -\frac{\eta}{4} \|\nabla f(x^k)\|. \tag{26}
\end{aligned}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f^* &\leq \langle \nabla f(x^k), x^k - x^* \rangle \\
&\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\
&\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{27}$$

Then substituting (27) into (26) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{4} \|\nabla f(x^k)\| \leq -\frac{\eta}{4R} (f(x^k) - f^*).$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{4R}\right) (f(x^k) - f^*).$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k, \xi^k)\| \geq \sqrt{2}\sigma$ and $\|\nabla f(x^k)\| \geq \sqrt{2}\sigma$ NSGD shows linear convergence:

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta}{4R}\right)^N (f(x^0) - f^*).$$

In the case: $\frac{\lambda}{\sqrt{2}} \leq \sigma$. Using this in (25), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\begin{aligned}
\mathbb{E} [f(x^{k+1})] - f(x^k) &\stackrel{(24)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E} [G(x^k, \xi^k)] \rangle + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \\
&\stackrel{(25)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} [\|G(x^k, \xi^k)\|^2] + \frac{\eta}{2\lambda^2\alpha} \mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2]
\end{aligned}$$

$$\begin{aligned}
& -\frac{\eta}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|} \right) \\
& \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda^2\alpha B} \\
& \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B}.
\end{aligned} \tag{28}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f^* & \leq \langle \nabla f(x^k), x^k - x^* \rangle \\
& \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\
& \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{29}$$

Then substituting (29) into (28) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B} \leq -\frac{\eta}{2R} (f(x^k) - f^*) + \frac{\eta\sigma^2 M}{2\lambda^2 B}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{2R}\right) (f(x^k) - f^*) + \frac{\eta\sigma^2 M}{2\lambda^2 B}.$$

Then for $k = 0, 1, 2, \dots, N - 1$ iterations that satisfy the conditions $\|\nabla f(x^k, \xi^k)\| \geq \lambda$ and $\|\nabla f(x^k)\| \geq \lambda$ and $\sigma \geq \sqrt{2}\lambda$ NSGD shows linear convergence:

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \frac{\sigma^2 MR}{\lambda^2 B}.$$

C.2 Second case: $\|\nabla f(x^k)\| \leq \lambda$ and $\|\nabla f(x^k, \xi^k)\| \geq \lambda$

Let us evaluate first summand of (24) with $\alpha = \lambda^{-1}$:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), G(x^k, \xi^k) \rangle & \stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
& \quad + \frac{\eta}{2\alpha} \|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2 \\
& \leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
& \quad + \frac{\eta}{2\lambda} \|\lambda G(x^k, \xi^k) - \nabla f(x^k)\|^2 \\
& = -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
& \quad + \frac{\eta}{2\lambda} \|\text{clip}_\lambda(\nabla f(x^k, \xi^k)) - \text{clip}_\lambda(\nabla f(x^k))\|^2
\end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius λ , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &\quad + \frac{\eta}{2\lambda} \mathbb{E}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2]. \end{aligned} \quad (30)$$

Using this, we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\begin{aligned} \mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(24)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &\stackrel{(30)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] + \frac{\eta}{2\lambda} \mathbb{E}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda} \mathbb{E}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\ &\quad - \frac{\eta}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|}\right) \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda B} \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda B}. \end{aligned} \quad (31)$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$\begin{aligned} f(x^k) - f^* &\leq \langle \nabla f(x^k), x^k - x^* \rangle \\ &\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\ &\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R. \end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \quad (32)$$

Then substituting (32) into (31) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda B} \leq -\frac{\eta}{2R} (f(x^k) - f^*) + \frac{\eta\sigma^2}{2\lambda B}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{2R}\right) (f(x^k) - f^*) + \frac{\eta\sigma^2}{2\lambda B}.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \leq \lambda$ and $\|\nabla f(x^k, \xi^k)\| \geq \lambda$ NSGD shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \frac{\sigma^2 R}{\lambda B}.$$

C.3 Third case: $\|\nabla f(x^k)\| \leq \lambda$ and $\|\nabla f(x^k, \xi^k)\| \leq \lambda$

Using this in (24), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$ and $\alpha = \|\nabla f(x^k)\|^{-1}$:

$$\begin{aligned}
\mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(24)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&= -\frac{\eta\alpha}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&\quad + \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2] \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2] \\
&\quad - \frac{\eta}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|}\right) \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2] \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2 + \|\alpha \nabla f(x^k)\|^2] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{\alpha} \mathbb{E}\left[\left\|\frac{\nabla f(x^k, \xi^k)}{\|\nabla f(x^k, \xi^k)\|}\right\|^2 + \left\|\frac{\nabla f(x^k)}{\|\nabla f(x^k)\|}\right\|^2\right] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{2\eta\lambda \|\nabla f(x^k)\|}{\lambda} \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + 2\eta\lambda.
\end{aligned} \tag{33}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f^* &\leq \langle \nabla f(x^k), x^k - x^* \rangle \\
&\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\
&\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{34}$$

Then substituting (34) into (33) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + 2\eta\lambda \leq -\frac{\eta}{2R} (f(x^k) - f^*) + 2\eta\lambda.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{2R}\right) (f(x^k) - f^*) + 2\eta\lambda.$$

Then for $k = 0, 1, 2, \dots, N - 1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \leq \lambda$ NSGD shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \lambda R.$$

C.4 Fourth case: $\|\nabla f(x^k)\| \geq \lambda$ and $\|\nabla f(x^k, \xi^k)\| \leq \lambda$

Using this in (24), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$ and $\alpha = \lambda^{-1}$:

$$\begin{aligned} \mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(24)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &= -\frac{\eta\alpha}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\mathbb{E}[G(x^k, \xi^k)]\|^2 \\ &\quad + \frac{\eta}{2\alpha} \|\mathbb{E}[G(x^k, \xi^k)] - \alpha \nabla f(x^k)\|^2 \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\ &= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \|\mathbb{E}[\lambda G(x^k, \xi^k)] - \nabla f(x^k)\|^2 \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\ &= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \left\| \mathbb{E} \left[\frac{\lambda \nabla f(x^k, \xi^k)}{\|\nabla f(x^k, \xi^k)\|} - \nabla f(x^k, \xi^k) \right] \right\|^2 \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\ &= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \left\| \mathbb{E} \left[\left(\frac{\lambda}{\|\nabla f(x^k, \xi^k)\|} - 1 \right) \nabla f(x^k, \xi^k) \right] \right\|^2 \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\ &\leq -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \mathbb{E} \left[\left(\frac{\lambda}{\|\nabla f(x^k, \xi^k)\|} - 1 \right)^2 \|\nabla f(x^k, \xi^k)\|^2 \right] \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\ &\leq -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \mathbb{E} \left[\frac{\lambda^2}{\|\nabla f(x^k, \xi^k)\|^2} \|\nabla f(x^k, \xi^k)\|^2 \right] \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\ &= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} + \frac{\eta\lambda}{2} \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} + \frac{\eta\lambda}{2} \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|} \right) + \frac{\eta\lambda}{2} \\ &\leq -\frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\lambda}{2}. \end{aligned} \tag{35}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1 \lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1 \lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1 \lambda)}$.

Using the convexity assumption of the function, we have the following:

$$\begin{aligned} f(x^k) - f^* &\leq \langle \nabla f(x^k), x^k - x^* \rangle \\ &\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \\ &\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R. \end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \quad (36)$$

Then substituting (36) into (35) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\lambda}{2} \leq -\frac{\eta}{4R} (f(x^k) - f^*) + \frac{\eta\lambda}{2}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{4R}\right) (f(x^k) - f^*) + \frac{\eta\lambda}{2}.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \geq \lambda$ and $\|\nabla f(x^k, \xi^k)\| \leq \lambda$ NSGD shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{4R}\right)^N (f(x^0) - f^*) + 2\lambda R.$$

Combining all the cases considered, we obtain the convergence rate of NSGD:

$$\mathbb{E}[f(x^N)] - f^* \lesssim \left(1 - \frac{\eta}{R}\right)^N (f(x^0) - f^*) + \frac{\sigma^2 MR}{B\lambda^2} + \lambda R.$$

D Zero-Order Clipped Stochastic Gradient Descent Method

This section consists of two parts: 1) a generalization of the convergence result of ClipSGD (Algorithm 1) to the biased gradient oracle $\mathbf{g}(x^k, \xi^k) = \nabla f(x^k, \xi^k) + \mathbf{b}(x^k)$, where $\mathbf{b}(x^k)$ is biased bounded by $\zeta \geq 0 : \|\mathbf{b}(x^k)\| \leq \zeta$; 2) deriving convergence estimates of ZO-ClipSGD directly.

D.1 Biased Clipped Stochastic Gradient Descent Method (Proof of the Lemma 5.1)

We start by using (L_0, L_1) -smoothness (see Assumption 1.2):

$$\begin{aligned} f(x^{k+1}) - f(x^k) &\stackrel{(8)}{\leq} \langle \nabla f(x^k), x^{k+1} - x^k \rangle + \frac{L_0 + L_1 \|\nabla f(x^k)\|}{2} \|x^{k+1} - x^k\|^2 \\ &= -\eta \langle \nabla f(x^k), \text{clip}_c(\mathbf{g}(x^k, \xi^k)) \rangle \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \|\text{clip}_c(\mathbf{g}(x^k, \xi^k))\|^2. \end{aligned} \quad (37)$$

Next, we consider three cases depending on the gradient norm: $\|\nabla f(x^k)\| \geq c$ – the full gradient is clipped and $\frac{c}{3} \leq \|\nabla f(x^k)\| \leq c$ and $\|\nabla f(x^k)\| \leq \frac{c}{3}$ – the full gradient is not clipped.

D.1.1 First case: $\|\nabla f(x^k)\| \geq c$

In this case $\alpha \nabla f(x^k) = \text{clip}_c(\nabla f(x^k))$ with $\alpha = \min \left\{ 1, \frac{c}{\|\nabla f(x^k)\|} \right\} = \frac{c}{\|\nabla f(x^k)\|}$, therefore we have the following

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \rangle &\stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \alpha \nabla f(x^k)\|^2 \\
&= -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \text{clip}_c(\nabla f(x^k))\|^2 \\
&= -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \text{clip}_c(\nabla f(x^k))\|^2.
\end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius c , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \mathbb{E}[\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))] \rangle &\leq -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] \\
&\quad + \frac{\eta}{2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] \\
&\stackrel{(9)}{=} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] \\
&\quad + \frac{\eta}{2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E}[\mathbf{g}(x^k, \boldsymbol{\xi}^k)]\|^2] \\
&\quad + \frac{\eta}{2\alpha} \|\mathbf{b}(x^k)\| \\
&\leq -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \boldsymbol{\xi}^k))\|^2] \\
&\quad + \frac{\eta\sigma^2 M}{2cB} + \frac{\eta \|\nabla f(x^k)\| \zeta^2}{2c}. \tag{38}
\end{aligned}$$

We now consider the cases depending on the relation between c and ζ :

In the case $c \geq \sqrt{2}\zeta$ We have in (38):

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \mathbb{E}[\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))] \rangle &\stackrel{(38)}{\leq} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] \\
&\quad + \frac{\eta\sigma^2 M}{2cB} + \frac{\eta \|\nabla f(x^k)\| \zeta^2}{2c} \\
&= -\frac{\eta}{2\alpha} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] \\
&\quad - \frac{c\eta}{2} \|\nabla f(x^k)\| \left(1 - \frac{\zeta^2}{c^2}\right) + \frac{\eta\sigma^2 M}{2cB} \\
&\leq -\frac{\eta}{2\alpha} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] - \frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2cB} \\
&= -\frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E}[\|\text{clip}_c(\nabla f(x^k, \boldsymbol{\xi}^k))\|^2] \\
&\quad - \frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2cB}.
\end{aligned}$$

Plugging this into (37) and choosing $\eta \leq \frac{1}{4(L_0 + L_1 c)}$ we have:

$$\mathbb{E}[\nabla f(x^{k+1})] - f(x^k) \stackrel{(12)}{\leq} -\frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] - \frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2cB}$$

$$\begin{aligned}
& + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] \\
& = -\frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] (1 - \eta L_1 c) - \frac{c\eta}{4} \|\nabla f(x^k)\| \\
& \quad + \frac{\eta^2 L_0}{2} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] + \frac{\eta \sigma^2 M}{2cB} \\
& \leq -\frac{c\eta}{4} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] (1 - \eta(L_0 + L_1 c)) \\
& \quad + \frac{\eta \sigma^2 M}{2cB} \\
& \leq -\frac{c\eta}{4} \|\nabla f(x^k)\| + \frac{\eta \sigma^2 M}{2cB}. \tag{39}
\end{aligned}$$

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{40}$$

Then substituting (40) into (39) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta c}{4} \|\nabla f(x^k)\| + \frac{\eta \sigma^2 M}{2cB} \leq -\frac{\eta c}{4R} (f(x^k) - f^*) + \frac{\eta \sigma^2 M}{2cB}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta c}{4R}\right) (f(x^k) - f^*) + \frac{\eta \sigma^2 M}{2cB}.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \geq c \geq \sqrt{2}\zeta$, then ClipSGD with biased gradient oracle has linear convergence

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \frac{\sigma^2 MR}{cB}.$$

In the case $c \leq \sqrt{2}\zeta$

We have in (38):

$$\begin{aligned}
& -\eta \langle \nabla f(x^k), \mathbb{E} [\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))] \rangle \stackrel{(38)}{\leq} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] \\
& \quad + \frac{\eta \sigma^2 M}{2cB} + \frac{\eta \|\nabla f(x^k)\| \zeta^2}{2c} \\
& = -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] \\
& \quad + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right).
\end{aligned}$$

Plugging this into (37) and choosing $\eta \leq \frac{1}{4(L_0 + L_1 c)}$ we have:

$$\begin{aligned}
\mathbb{E} [f(x^{k+1})] - f(x^k) & \stackrel{(37)}{\leq} -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] \\
& \quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \\
& = -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta \|\nabla f(x^k)\|}{2c} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] (1 - \eta L_1 c)
\end{aligned}$$

$$\begin{aligned}
& + \frac{\eta^2 L_0}{2} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \\
& \leq -\frac{c\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E} \left[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \right] (1 - \eta(L_0 + L_1 c)) \\
& \quad + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \\
& \leq -\frac{c\eta}{2} \|\nabla f(x^k)\| + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right). \tag{41}
\end{aligned}$$

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{42}$$

Then substituting (42) into (41) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta c}{2} \|\nabla f(x^k)\| + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \leq -\frac{\eta c}{2R} (f(x^k) - f^*) + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right).$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta c}{2R}\right) (f(x^k) - f^*) + \frac{\eta M}{2c} \left(\frac{\sigma^2}{B} + \zeta^2 \right).$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \geq c$ and $c \leq \sqrt{2}\zeta$, then ClipSGD has linear convergence

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta c}{2R}\right)^N (f(x^0) - f^*) + \frac{MR}{c^2} \left(\frac{\sigma^2}{B} + \zeta^2 \right).$$

D.1.2 Second case: $\frac{c}{3} \leq \|\nabla f(x^k)\| \leq c$

In this case $\nabla f(x^k) = \text{clip}_c(\nabla f(x^k))$ with $\alpha = \min \left\{ 1, \frac{c}{\|\nabla f(x^k)\|} \right\} = 1$, therefore we have the following

$$\begin{aligned}
& -\eta \langle \nabla f(x^k), \text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \rangle \stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \\
& \quad + \frac{\eta}{2\alpha} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \alpha \nabla f(x^k)\|^2 \\
& = -\frac{\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \\
& \quad + \frac{\eta}{2} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \text{clip}_c(\nabla f(x^k))\|^2 \\
& \leq -\frac{c\eta}{6} \|\nabla f(x^k)\| - \frac{\eta}{2} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2 \\
& \quad + \frac{\eta}{2} \|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \text{clip}_c(\nabla f(x^k))\|^2.
\end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius c , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned}
& -\eta \langle \nabla f(x^k), \mathbb{E}[\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))] \rangle \leq -\frac{c\eta}{6} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] \\
& \quad + \frac{\eta}{2} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] \\
& \stackrel{(9)}{=} -\frac{c\eta}{6} \|\nabla f(x^k)\| - \frac{\eta}{2} \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2]
\end{aligned}$$

$$\begin{aligned}
& + \frac{\eta}{2} \mathbb{E} \left[\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E} [\mathbf{g}(x^k, \boldsymbol{\xi}^k)] \right\|^2 \right] \\
& + \frac{\eta}{2} \left\| \mathbf{b}(x^k) \right\|^2 \\
& \leq -\frac{c\eta}{6} \left\| \nabla f(x^k) \right\| - \frac{\eta}{2} \mathbb{E} \left[\left\| \text{clip}_c (\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \right\|^2 \right] \\
& \quad + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \\
& = -\frac{c\eta}{6} \left\| \nabla f(x^k) \right\| - \frac{\eta}{2} \mathbb{E} \left[\left\| \text{clip}_c (\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \right\|^2 \right] \\
& \quad + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right).
\end{aligned}$$

Plugging this into (37) and choosing $\eta \leq \frac{1}{4(L_0 + L_1 c)}$ we have:

$$\begin{aligned}
\mathbb{E} [f(x^{k+1})] - f(x^k) & \stackrel{(37)}{\leq} -\frac{c\eta}{6} \left\| \nabla f(x^k) \right\| - \frac{\eta}{2} \mathbb{E} \left[\left\| \text{clip}_c (\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \right\|^2 \right] \\
& \quad + \frac{\eta^2 (L_0 + L_1 \left\| \nabla f(x^k) \right\|)}{2} \mathbb{E} \left[\left\| \text{clip}_c (\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \right\|^2 \right] + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \\
& = -\frac{c\eta}{6} \left\| \nabla f(x^k) \right\| - \frac{\eta}{2} \mathbb{E} \left[\left\| \text{clip}_c (\mathbf{g}(x^k, \boldsymbol{\xi}^k)) \right\|^2 \right] (1 - \eta(L_0 + L_1 \left\| \nabla f(x^k) \right\|)) \\
& \quad + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \\
& \leq -\frac{c\eta}{6} \left\| \nabla f(x^k) \right\| + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right). \tag{43}
\end{aligned}$$

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \left\| \nabla f(x^k) \right\| \left\| x^k - x^* \right\| \leq \left\| \nabla f(x^k) \right\| \underbrace{\left\| x^0 - x^* \right\|}_R.$$

Hence we have:

$$\left\| \nabla f(x^k) \right\| \geq \frac{f(x^k) - f^*}{R}. \tag{44}$$

Then substituting (44) into (43) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta c}{6} \left\| \nabla f(x^k) \right\| + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right) \leq -\frac{\eta c}{6R} (f(x^k) - f^*) + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right).$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta c}{6R} \right) (f(x^k) - f^*) + \frac{\eta}{2} \left(\frac{\sigma^2}{B} + \zeta^2 \right).$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\frac{c}{2} \leq \left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\| \leq c$, then ClipSGD with biased gradient oracle has linear convergence

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta c}{6R} \right)^N (f(x^0) - f^*) + \frac{3R}{c} \left(\frac{\sigma^2}{B} + \zeta^2 \right).$$

Let $\mathcal{T}_1 = \{m_0^{\mathcal{T}_1}, m_1^{\mathcal{T}_1}, m_2^{\mathcal{T}_1}, \dots, m_{K-1}^{\mathcal{T}_1}\} = \{k \in \{0, 1, 2, \dots, N-1\} \mid \left\| \nabla f(x^k, \boldsymbol{\xi}^k) \right\| \geq \frac{c}{3}\}$, where $K = |\mathcal{T}_1|$. Then for $k \in \mathcal{T}_1$ ClipSGD with biased gradient oracle shows linear convergence:

$$F_N \cdot \mathbb{1}[\mathcal{T}_1] \lesssim \left(1 - \frac{\eta c}{6R} \right) F_N \leq \left(1 - \frac{\eta c}{6R} \right) F_{m_{K-1}^{\mathcal{T}_1}} + \left(\frac{\eta M}{c} + \eta \right) \cdot \left(\frac{\sigma^2}{B} + \zeta^2 \right) \leq \dots \leq$$

$$\begin{aligned}
&\leq \left(1 - \frac{\eta c}{R}\right)^K F_{m_0 T_1} + \left(\frac{MR}{c^2} + \frac{R}{c}\right) \cdot \left(\frac{\sigma^2}{B} + \zeta^2\right) \\
&\leq \left(1 - \frac{\eta c}{R}\right)^K F_0 + \left(\frac{MR}{c^2} + \frac{R}{c}\right) \cdot \left(\frac{\sigma^2}{B} + \zeta^2\right),
\end{aligned}$$

where $F_k = \mathbb{E}[f(x^k)] - f^*$, and we used that $F_k \leq F_{k-1}$.

D.1.3 Third case: $\|\nabla f(x^k)\| \leq \frac{c}{3}$

We introduce an indicative function:

$$\aleph_k = \mathbf{1} \{ \|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| > c \}. \quad (45)$$

Then the following is true:

$$\mathbb{E}[\aleph_k] = \mathbb{E}[\aleph_k^2] = \mathcal{P}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| > c] \stackrel{\textcircled{1}}{\leq} \mathcal{P}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E}[\mathbf{g}(x^k, \boldsymbol{\xi}^k)]\| > \frac{c}{3}] \stackrel{\textcircled{2}}{\leq} \frac{9\sigma^2}{c^2 B}, \quad (46)$$

where in $\textcircled{1}$ we used $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| \leq \|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E}[\mathbf{g}(x^k, \boldsymbol{\xi}^k)]\| + \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|] \leq \|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E}[\mathbf{g}(x^k, \boldsymbol{\xi}^k)]\| + \frac{c}{2}$, where assume that $\zeta \leq \frac{c}{3}$: and in $\textcircled{2}$ we used Markov's inequality.

Let $r_{k+1} = \mathbb{E}[\|x^{k+1} - x^*\|]$ and $F_{k+1} = \mathbb{E}[f(x^{k+1}) - f^*]$, then given that

$$\begin{aligned}
\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) &= \mathbf{g}(x^k, \boldsymbol{\xi}^k)(1 - \aleph_k) + \frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \\
&= \mathbf{g}(x^k, \boldsymbol{\xi}^k) + \left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k
\end{aligned}$$

we get with $\eta \leq \frac{1}{4(L_0 + L_1 c)}$:

$$\begin{aligned}
r_{k+1}^2 &= r_k^2 - 2\eta \langle \mathbb{E}[\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))], x^k - x^* \rangle + \eta^2 \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] \\
&= r_k^2 - 2\eta \langle \nabla f(x^k), x^k - x^* \rangle - 2\eta \left\langle \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right], x^k - x^* \right\rangle \\
&\quad + \eta^2 \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] + 2\eta \langle \mathbf{b}(x^k), x^k - x^* \rangle \\
&\stackrel{(6)}{\leq} r_k^2 - 2\eta \langle \nabla f(x^k), x^k - x^* \rangle + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| \|x^k - x^*\| \\
&\quad + \eta^2 \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] + 2\eta \|\mathbf{b}(x^k)\| \|x^k - x^*\| \\
&\stackrel{\textcircled{1}}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| \|x^0 - x^*\| \\
&\quad + \eta^2 \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k))\|^2] + 2\eta \|\mathbf{b}(x^k)\| \|x^0 - x^*\| \\
&\stackrel{(7)}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| \|x^0 - x^*\| \\
&\quad + 2\eta^2 \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \nabla f(x^k)\|^2] + 2\eta^2 \|\nabla f(x^k)\|^2 + 2\eta \zeta \|x^0 - x^*\| \\
&= r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
&\quad + 2\eta^2 \mathbb{E}[\|\text{clip}_c(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \text{clip}_c(\nabla f(x^k))\|^2] + 2\eta^2 \|\nabla f(x^k)\|^2 + 2\eta \zeta R \\
&\stackrel{\textcircled{2}}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
&\quad + 2\eta^2 \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] + 2\eta^2 \|\nabla f(x^k)\|^2 + 2\eta \zeta R \\
&\stackrel{(9)}{=} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R
\end{aligned}$$

$$\begin{aligned}
& + 2\eta^2 \mathbb{E} \left[\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E} [\mathbf{g}(x^k, \boldsymbol{\xi}^k)] \right\|^2 \right] + 2\eta^2 \left\| \nabla f(x^k) \right\|^2 + 2\eta\zeta R + 2\eta^2 \left\| \mathbf{b}(x^k) \right\| \\
& \leq r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
& \quad + \frac{2\eta^2 \sigma^2}{B} + 2\eta^2 \left\| \nabla f(x^k) \right\|^2 + 2\eta\zeta R + 2\eta^2 \zeta^2 \\
& \stackrel{(10)}{\leq} r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
& \quad + \frac{2\eta^2 \sigma^2}{B} + 4\eta^2 (L_0 + L_1 \left\| \nabla f(x^k) \right\|) F_k + 2\eta\zeta R + 2\eta^2 \zeta^2 \\
& \leq r_k^2 - 2\eta F_k + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
& \quad + \frac{2\eta^2 \sigma^2}{B} + 4\eta^2 (L_0 + L_1 c) F_k + 2\eta\zeta R + 2\eta^2 \zeta^2 \\
& = r_k^2 - 2\eta F_k (1 - 2\eta (L_0 + L_1 c)) + \frac{2\eta^2 \sigma^2}{B} + 2\eta\zeta R + 2\eta^2 \zeta^2 \\
& \quad + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
& \leq r_k^2 - \eta F_k + \frac{2\eta^2 \sigma^2}{B} + 2\eta\zeta R + 2\eta^2 \zeta^2 \\
& \quad + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R. \tag{47}
\end{aligned}$$

Let's find the upper bound of the last summand:

$$\begin{aligned}
& 2\eta R \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| \\
& \stackrel{(45)}{\leq} 2\eta R \mathbb{E} \left[\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\| \cdot \left(1 - \frac{c}{\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\|} \right) \aleph_k \right] \\
& \leq 2\eta R \mathbb{E} \left[\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right\| \cdot \aleph_k \right] \\
& \leq 2\eta R (\mathbb{E} [\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E} [\mathbf{g}(x^k, \boldsymbol{\xi}^k)] \right\| \cdot \aleph_k] + \left\| \nabla f(x^k) \right\| \mathbb{E} [\aleph_k] + \left\| \mathbf{b}(x^k) \right\| \mathbb{E} [\aleph_k]) \\
& \leq 2\eta R \left(\sqrt{\mathbb{E} [\left\| \mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E} [\mathbf{g}(x^k, \boldsymbol{\xi}^k)] \right\|^2] \cdot \mathbb{E} [\aleph_k^2]} + \frac{2c}{3} \mathbb{E} [\aleph_k] \right) \\
& \stackrel{(46)}{\leq} 2\eta R \left(\frac{3\sigma^2}{cB} + \frac{2c}{3} \cdot \frac{9\sigma^2}{c^2 B} \right) \\
& = \frac{18\eta\sigma^2 R}{cB}. \tag{48}
\end{aligned}$$

Substituting into the initial formula and rearrange the summands, we obtain

$$\begin{aligned}
\eta F_k & \stackrel{(47)}{\leq} r_k^2 - r_{k+1}^2 + \frac{2\eta^2 \sigma^2}{B} + 2\eta\zeta R + 2\eta^2 \zeta^2 + 2\eta \left\| \mathbb{E} \left[\left(\frac{c}{\left\| \nabla f(x^k, \boldsymbol{\xi}^k) \right\|} - 1 \right) \nabla f(x^k, \boldsymbol{\xi}^k) \aleph_k \right] \right\| R \\
& \stackrel{(48)}{\leq} r_k^2 - r_{k+1}^2 + \frac{2\eta^2 \sigma^2}{B} + \frac{18\eta\sigma^2 R}{cB} + 2\eta\zeta R + 2\eta^2 \zeta^2
\end{aligned}$$

Let $\mathcal{T}_2 = \{m_0^{\mathcal{T}_2}, m_1^{\mathcal{T}_2}, m_2^{\mathcal{T}_2}, \dots, m_{N-K}^{\mathcal{T}_2}\} = \{k \in \{0, 1, 2, \dots, N-1\} \mid \left\| \nabla f(x^k) \right\| < \frac{c}{3}\}$, where $|\mathcal{T}_2| = N - K$. Then rearranging and summing over all $k \in \mathcal{T}_2$ we obtain

$$F_N \cdot \mathbb{1}[\mathcal{T}_2] \leq \frac{1}{N-K} \sum_{k \in \mathcal{T}_2} F_k \leq \frac{1}{\eta(N-K)} \sum_{k \in \mathcal{T}_2} (r_k^2 - r_{k+1}^2)$$

$$\begin{aligned}
& + \frac{1}{N-K} \sum_{k \in \mathcal{T}_2} \left[\eta \left(\frac{\sigma^2}{B} + \zeta^2 \right) + 2R \left(\frac{9\sigma^2}{cB} + \zeta \right) \right] \\
& = \frac{r_0^2 - r_N^2}{\eta(N-K)} \eta \left(\frac{\sigma^2}{B} + \zeta^2 \right) + 2R \left(\frac{9\sigma^2}{cB} + \zeta \right) \leq \frac{r_0^2}{\eta(N-K)} \eta \left(\frac{\sigma^2}{B} + \zeta^2 \right) + 2R \left(\frac{9\sigma^2}{cB} + \zeta \right).
\end{aligned}$$

Hence we obtain:

$$F_N \cdot \mathbb{1}[\mathcal{T}_2] \leq \frac{R^2}{\eta(N-K)} \eta \left(\frac{\sigma^2}{B} + \zeta^2 \right) + 2R \left(\frac{9\sigma^2}{cB} + \zeta \right).$$

Combining all the cases considered, we obtain the convergence rate of ClipSGD with biased gradient oracle:

$$\begin{aligned}
\mathbb{E}[f(x^N)] - f^* & \leq F_N \cdot \mathbb{1}[\mathcal{T}_1] + F_N \cdot \mathbb{1}[\mathcal{T}_2] \\
& \lesssim \left(1 - \frac{\eta c}{R}\right)^K F_0 + \frac{R^2}{\eta(N-K)} + \left(\frac{MR}{c^2} + \frac{R}{c} + \eta\right) \cdot \left(\frac{\sigma^2}{B} + \zeta^2\right) + R\zeta. \quad (49)
\end{aligned}$$

D.2 Convergence Results for ZO-ClipSGD

In order to obtain convergence results for ZO-ClipSGD it is necessary to estimate the bias and variance of the gradient approximation (4).

Bias of gradient approximation Using the variational representation of the Euclidean norm, and definition of gradient approximation (4) we can write:

$$\begin{aligned}
\|\mathbb{E}[\mathbf{g}(x, \{\xi, e\})] - \nabla f(x)\| & = \left\| \mathbb{E} \left[\frac{d}{2\gamma} \left(\tilde{f}(x + \gamma e, \xi) - \tilde{f}(x - \gamma e, \xi) \right) e \right] - \nabla f(x) \right\| \\
& \stackrel{\textcircled{1}}{=} \left\| \mathbb{E} \left[\frac{d}{\gamma} (f(x + \gamma e, \xi) + \delta(x + \gamma e)) e \right] - \nabla f(x) \right\| \\
& \stackrel{\textcircled{2}}{\leq} \left\| \mathbb{E} \left[\frac{d}{\gamma} f(x + \gamma e, \xi) e \right] - \nabla f(x) \right\| + \frac{d\Delta}{\gamma} \\
& \stackrel{\textcircled{3}}{=} \|\mathbb{E}[\nabla f(x + \gamma u, \xi)] - \nabla f(x)\| + \frac{d\Delta}{\gamma} \\
& = \sup_{z \in S^d(1)} \mathbb{E}[\|\nabla_z f(x + \gamma u, \xi) - \nabla_z f(x)\|] + \frac{d\Delta}{\gamma} \\
& \stackrel{\textcircled{8}}{\leq} (L_0 + L_1 \|\nabla f(x^k)\|) \gamma \mathbb{E}[\|u\|] + \frac{d\Delta}{\gamma} \\
& \leq (L_0 + L_1 M) \gamma + \frac{d\Delta}{\gamma}, \quad (50)
\end{aligned}$$

where $u \in B^d(1)$, $\textcircled{1}$ = the equality is obtained from the fact, namely, distribution of e is symmetric, $\textcircled{2}$ = the inequality is obtain from bounded noise $|\delta(x)| \leq \Delta$, $\textcircled{3}$ = the equality is obtained from a version of Stokes' theorem [see Section 13.3.5, Exercise 14a, 67].

Bounding second moment (variance) of gradient approximation By definition gradient approximation (4) and Wirtinger-Poincare inequality (11) we have

$$\begin{aligned}
& \mathbb{E} \left[\|\mathbf{g}(x, \{\xi, e\}) - \mathbb{E}[\mathbf{g}(x, \{\xi, e\})]\|^2 \right] \\
& \leq \mathbb{E} \left[\|\mathbf{g}(x, \{\xi, e\})\|^2 \right] \\
& = \frac{d^2}{4\gamma^2} \mathbb{E} \left[\left\| \left(\tilde{f}(x + \gamma e, \xi) - \tilde{f}(x - \gamma e, \xi) \right) e \right\|^2 \right] \\
& = \frac{d^2}{4\gamma^2} \mathbb{E} \left[(f(x + \gamma e, \xi) - f(x - \gamma e, \xi) + \delta(x + \gamma e) - \delta(x - \gamma e))^2 \right]
\end{aligned}$$

$$\begin{aligned}
&\stackrel{(7)}{\leq} \frac{d^2}{2\gamma^2} \left(\mathbb{E} \left[(f(x + \gamma e, \xi) - f(x - \gamma e, \xi))^2 \right] + 2\Delta^2 \right) \\
&\stackrel{(11)}{\leq} \frac{d^2}{2\gamma^2} \left(\frac{\gamma^2}{d} \mathbb{E} \left[\|\nabla f(x + \gamma e, \xi) + \nabla f(x - \gamma e, \xi)\|^2 \right] + 2\Delta^2 \right) \\
&= \frac{d^2}{2\gamma^2} \left(\frac{\gamma^2}{d} \mathbb{E} \left[\|\nabla f(x + \gamma e, \xi) + \nabla f(x - \gamma e, \xi) \pm 2\nabla f(x, \xi)\|^2 \right] + 2\Delta^2 \right) \\
&\stackrel{(8)}{\leq} 4d\mathbb{E} \left[\|\nabla f(x, \xi)\|^2 \right] + 4dL^2\gamma^2\mathbb{E} \left[\|e\|^2 \right] + \frac{d^2\Delta^2}{\gamma^2} \\
&\stackrel{\textcircled{1}}{\leq} 4d\tilde{\sigma}^2 + 4d(L_0 + L_1 \|\nabla f(x^k)\|)^2 \gamma^2 \mathbb{E} \left[\|e\|^2 \right] + \frac{d^2\Delta^2}{\gamma^2} \\
&\leq 4d\tilde{\sigma}^2 + 4d(L_0 + L_1 M)^2 \gamma^2 + \frac{d^2\Delta^2}{\gamma^2}, \tag{51}
\end{aligned}$$

where $\textcircled{1}$ = the inequality is obtain from $\mathbb{E} \left[\|\nabla f(x, \xi)\|^2 \right] \leq \tilde{\sigma}^2$.

D.2.1 Proof of Theorem 5.2

In order to obtain the convergence rate of ZO-ClipSGD in the convex setting, we need to substitute the obtained estimates (50) and (51) into the convergence rate of ClipSGD (49) instead of ζ and σ^2 , respectively. Given that $\frac{MR}{c^2} + \frac{R}{c} + \eta \lesssim \frac{MR}{c^2}$ at small c , then the convergence of ZO-ClipSGD in the convex setup is as follows:

$$\begin{aligned}
\mathbb{E} [f(x^N)] - f^* &\lesssim \underbrace{\left(1 - \frac{\eta}{R}\right)^K (f(x^0) - f^*)}_{\textcircled{1}} + \underbrace{\frac{R^2}{\eta(N-K)}}_{\textcircled{2}} + \underbrace{\frac{dMR\tilde{\sigma}^2}{c^2B}}_{\textcircled{3}} + \underbrace{\frac{dMR(L_0 + L_1M)^2\gamma^2}{c^2B}}_{\textcircled{4}} \\
&\quad + \underbrace{\frac{d^2MR\Delta^2}{c^2B\gamma^2}}_{\textcircled{5}} + \underbrace{\frac{MR(L_0 + L_1M)^2\gamma^2}{c^2}}_{\textcircled{6}} + \underbrace{\frac{d^2MR\Delta^2}{c^2\gamma^2}}_{\textcircled{7}} \\
&\quad + \underbrace{(L_0 + L_1M)\gamma R}_{\textcircled{8}} + \underbrace{\frac{d\Delta R}{\gamma}}_{\textcircled{9}}.
\end{aligned}$$

From term $\textcircled{1}$, we find the K :

$$\textcircled{1}: \quad \left(1 - \frac{\eta c}{R}\right)^K (f(x^0) - f^*) \leq \varepsilon \quad \Rightarrow \quad K \geq \frac{R}{\eta c} \log \frac{f(x^0) - f^*}{\varepsilon}. \tag{52}$$

From term $\textcircled{2}$, we find the number of iterations N required for Algorithm 3 in convex setup to achieve ε -accuracy:

$$\begin{aligned}
\textcircled{2}: \quad \frac{R^2}{\eta(N-K)} \leq \varepsilon \quad &\Rightarrow \quad N \stackrel{(52)}{\geq} \frac{R^2}{\eta\varepsilon} + \frac{R}{\eta c} \log \frac{f(x^0) - f^*}{\varepsilon}; \\
N &= \mathcal{O} \left(\frac{R^2}{\eta\varepsilon} + \frac{R}{\eta c} \log \frac{1}{\varepsilon} \right). \tag{53}
\end{aligned}$$

From terms $\textcircled{3}$, we find the batch size B :

$$\begin{aligned}
\textcircled{3}: \quad \frac{dMR\tilde{\sigma}^2}{c^2B} \leq \varepsilon \quad &\Rightarrow \quad B \geq \frac{dMR\tilde{\sigma}^2}{\varepsilon c^2}; \\
B &= \mathcal{O} \left(\frac{dMR\tilde{\sigma}^2}{\varepsilon c^2} \right). \tag{54}
\end{aligned}$$

From terms $\textcircled{4}$, $\textcircled{6}$ and $\textcircled{8}$ we find the smoothing parameter γ :

$$\textcircled{4}: \quad \frac{dMR(L_0 + L_1M)^2\gamma^2}{c^2B} \leq \varepsilon \quad \Rightarrow \quad \gamma \leq \sqrt{\frac{\varepsilon c^2B}{dMR(L_0 + L_1M)^2}} \stackrel{(54)}{=} \frac{\tilde{\sigma}}{(L_0 + L_1M)};$$

$$\begin{aligned}
\textcircled{6} : \quad & \frac{MR(L_0 + L_1M)^2 \gamma^2}{c^2} \leq \varepsilon \quad \Rightarrow \quad \gamma \leq \frac{\sqrt{\varepsilon}c}{\sqrt{MR}(L_0 + L_1M)}; \\
\textcircled{8} : \quad & (L_0 + L_1M) R \gamma \leq \varepsilon \quad \Rightarrow \quad \gamma \leq \frac{\varepsilon}{R(L_0 + L_1M)}; \\
& \gamma \leq \frac{1}{(L_0 + L_1M)} \min \left\{ \tilde{\sigma}, \frac{\sqrt{\varepsilon}c}{\sqrt{MR}}, \frac{\varepsilon}{R} \right\} = \frac{\varepsilon}{R(L_0 + L_1M)}. \tag{55}
\end{aligned}$$

From the remaining terms $\textcircled{5}$, $\textcircled{7}$ and $\textcircled{9}$, we find the maximum allowable level of adversarial noise Δ that still guarantees the convergence of the ZO-ClipSGD to desired accuracy ε in convex setup:

$$\begin{aligned}
\textcircled{5} : \quad & \frac{d^2 MR \Delta^2}{c^2 B \gamma^2} \leq \varepsilon \quad \Rightarrow \quad \Delta \leq \frac{\sqrt{\varepsilon}c\gamma\sqrt{B}}{d\sqrt{MR}} \stackrel{(54),(55)}{=} \frac{\varepsilon \tilde{\sigma}}{\sqrt{d}(L_0 + L_1M) R}; \\
\textcircled{7} : \quad & \frac{d^2 MR \Delta^2}{\gamma^2 c^2} \leq \varepsilon \quad \Rightarrow \quad \Delta \leq \sqrt{\frac{\gamma^2 c^2 \varepsilon}{d^2 MR}} \stackrel{(55)}{=} \frac{\varepsilon^{3/2} c}{d(L_0 + L_1M) \sqrt{MR}^{3/2}}; \\
\textcircled{9} : \quad & \frac{d\Delta R}{\gamma} \leq \varepsilon \quad \Rightarrow \quad \Delta \leq \sqrt{\frac{\gamma \varepsilon}{dR}} \stackrel{(55)}{=} \frac{\varepsilon^2}{d(L_0 + L_1M) R^2}; \\
& \Delta \leq \frac{\varepsilon}{\sqrt{d}(L_0 + L_1M) R} \min \left\{ \tilde{\sigma}, \frac{\sqrt{\varepsilon}c}{\sqrt{d}\sqrt{MR}}, \frac{\varepsilon}{\sqrt{d}R} \right\} \\
& = \frac{\varepsilon}{\sqrt{d}(L_0 + L_1M) R} \min \left\{ \tilde{\sigma}, \frac{\varepsilon}{\sqrt{d}R} \right\}. \tag{56}
\end{aligned}$$

In this way, the ZO-ClipSGD achieves ε -accuracy: $\mathbb{E}[f(x^N) - f^*] \leq \varepsilon$ in convex setup after

$$N \stackrel{(53)}{=} \mathcal{O} \left(\frac{R^2}{\eta \varepsilon} + \frac{R}{\eta c} \log \frac{1}{\varepsilon} \right), \quad T = N \cdot B \stackrel{(53),(54)}{=} \mathcal{O} \left(\frac{d\tilde{\sigma}^2 MR^2}{\varepsilon c^2 \eta} \left(\frac{1}{c} \log \frac{1}{\varepsilon} + \frac{R}{\varepsilon} \right) \right)$$

number of iterations, total number of zero-order oracle calls and at

$$\Delta \stackrel{(56)}{\lesssim} \frac{\varepsilon}{\sqrt{d}(L_0 + L_1M) R} \min \left\{ \tilde{\sigma}, \frac{\varepsilon}{\sqrt{d}R} \right\}$$

the maximum level of noise with smoothing parameter $\frac{\varepsilon}{(L_0 + L_1M)R}$ (55).

E Zero-Order Normalized Stochastic Gradient Descent Method

This section consists of two parts: 1) a generalization of the convergence result of NSGD (Algorithm 2) to the biased gradient oracle $\mathbf{g}(x^k, \boldsymbol{\xi}^k) = \nabla f(x^k, \boldsymbol{\xi}^k) + \mathbf{b}(x^k)$, where $\mathbf{b}(x^k)$ is biased bounded by $\zeta \geq 0 : \|\mathbf{b}(x^k)\| \leq \zeta$; 2) deriving convergence estimates of ZO-NSGD directly.

E.1 Biased Normalized Stochastic Gradient Descent Method (Proof of the Lemma 5.3)

Let's introduce the notation $G(x^k, \boldsymbol{\xi}^k) = \frac{\mathbf{g}(x^k, \boldsymbol{\xi}^k)}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|}$, then using (L_0, L_1) -smoothness (see Assumption 1.2):

$$\begin{aligned}
f(x^{k+1}) - f(x^k) & \stackrel{(8)}{\leq} \langle \nabla f(x^k), x^{k+1} - x^k \rangle + \frac{L_0 + L_1 \|\nabla f(x^k)\|}{2} \|x^{k+1} - x^k\|^2 \\
& = -\eta \langle \nabla f(x^k), G(x^k, \boldsymbol{\xi}^k) \rangle + \frac{\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|)}{2} \|G(x^k, \boldsymbol{\xi}^k)\|^2. \tag{57}
\end{aligned}$$

Next, we consider 4 cases of the relation $\|\nabla f(x^k)\|$ and $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|$ with respect to the hyperparameter λ .

E.1.1 First case: $\|\nabla f(x^k)\| \geq \lambda$ and $\|\mathbf{g}(x^k, \xi^k)\| \geq \lambda$

Let us evaluate first summand of (57) with $\alpha = \|\nabla f(x^k)\|^{-1}$:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), G(x^k, \xi^k) \rangle &\stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2 \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
&\quad + \frac{\eta}{2\lambda^2\alpha} \|\lambda G(x^k, \xi^k) - \lambda\alpha \nabla f(x^k)\|^2 \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
&\quad + \frac{\eta}{2\lambda^2\alpha} \|\text{clip}_\lambda(\mathbf{g}(x^k, \xi^k)) - \text{clip}_\lambda(\nabla f(x^k))\|^2
\end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius λ , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&\quad + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \xi^k) - \nabla f(x^k)\|^2]. \tag{58}
\end{aligned}$$

In the case: $0 \leq \zeta \leq \frac{\lambda}{\sqrt{2}}$. Using this in (58), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\begin{aligned}
\mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(57)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&\stackrel{(58)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \xi^k) - \nabla f(x^k)\|^2] \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \xi^k) - \nabla f(x^k)\|^2] \\
&\quad - \frac{\eta}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|}\right) \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \xi^k) - \nabla f(x^k)\|^2] \\
&\stackrel{(9)}{=} -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \xi^k) - \mathbb{E}[\mathbf{g}(x^k, \xi^k)]\|^2] + \frac{\eta}{2\lambda^2\alpha} \|\mathbf{b}(x^k)\|^2 \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda^2\alpha B} + \frac{\eta\zeta^2}{2\lambda^2\alpha} \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B} \\
&= -\frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B}. \tag{59}
\end{aligned}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \quad (60)$$

Then substituting (60) into (59) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B} \leq -\frac{\eta}{4R} (f(x^k) - f^*) + \frac{\eta\sigma^2 M}{2\lambda^2 B}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{4R}\right) (f(x^k) - f^*) + \frac{\eta\sigma^2 M}{2\lambda^2 B}.$$

Then for $k = 0, 1, 2, \dots, N - 1$ iterations that satisfy the conditions $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| \geq \sqrt{2}\zeta$ and $\|\nabla f(x^k)\| \geq \sqrt{2}\zeta$ NSGD with biased gradient oracle shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{4R}\right)^N (f(x^0) - f^*) + \frac{2\sigma^2 MR}{\lambda^2 B}.$$

In the case: $\frac{\lambda}{\sqrt{2}} \leq \zeta$. Using this in (58), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\begin{aligned} \mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(57)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \boldsymbol{\xi}^k)] \rangle + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\ &\stackrel{(58)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] \\ &\quad - \frac{\eta}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|}\right) \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] \\ &\stackrel{(9)}{=} -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E}[\mathbf{g}(x^k, \boldsymbol{\xi}^k)]\|^2] + \frac{\eta}{2\lambda^2\alpha} \|\mathbf{b}(x^k)\|^2 \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda^2\alpha B} + \frac{\eta\zeta^2}{2\lambda^2\alpha} \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B} + \frac{\eta\zeta^2 M}{2\lambda^2} \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B} + \frac{\eta\zeta^2 M}{2\lambda^2}. \end{aligned} \quad (61)$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \underbrace{\|\nabla f(x^k)\|}_R \|x^0 - x^*\|.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \quad (62)$$

Then substituting (62) into (61) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2 M}{2\lambda^2 B} + \frac{\eta\zeta^2 M}{2\lambda^2} \leq -\frac{\eta}{2R}(f(x^k) - f^*) + \frac{\eta\sigma^2 M}{2\lambda^2 B} + \frac{\eta\zeta^2 M}{2\lambda^2}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{2R}\right)(f(x^k) - f^*) + \frac{\eta\sigma^2 M}{2\lambda^2 B} + \frac{\eta\zeta^2 M}{2\lambda^2}.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| \geq \lambda$ and $\|\nabla f(x^k)\| \geq \lambda$ and $\zeta \geq \sqrt{2}\lambda$ NSGD with biased gradient oracle shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \frac{\sigma^2 MR}{\lambda^2 B} + \frac{\zeta^2 MR}{\lambda^2}.$$

E.1.2 Second case: $\|\nabla f(x^k)\| \leq \lambda$ and $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| \geq \lambda$

Let us evaluate first summand of (57) with $\alpha = \lambda^{-1}$:

$$\begin{aligned} -\eta \langle \nabla f(x^k), G(x^k, \boldsymbol{\xi}^k) \rangle &\stackrel{(5)}{=} -\frac{\alpha\eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|G(x^k, \boldsymbol{\xi}^k)\|^2 \\ &\quad + \frac{\eta}{2\alpha} \|G(x^k, \boldsymbol{\xi}^k) - \alpha \nabla f(x^k)\|^2 \\ &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \boldsymbol{\xi}^k)\|^2 \\ &\quad + \frac{\eta}{2\lambda} \|\lambda G(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2 \\ &= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \boldsymbol{\xi}^k)\|^2 \\ &\quad + \frac{\eta}{2\lambda} \|\text{clip}_\lambda(\mathbf{g}(x^k, \boldsymbol{\xi}^k)) - \text{clip}_\lambda(\nabla f(x^k))\|^2 \end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius λ , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \boldsymbol{\xi}^k)] \rangle &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\ &\quad + \frac{\eta}{2\lambda} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2]. \end{aligned} \quad (63)$$

Using this, we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\begin{aligned} \mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(57)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \boldsymbol{\xi}^k)] \rangle + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\ &\stackrel{(63)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] + \frac{\eta}{2\lambda} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \nabla f(x^k)\|^2] \\ &\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\ &\stackrel{(9)}{=} -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda} \mathbb{E}[\|\mathbf{g}(x^k, \boldsymbol{\xi}^k) - \mathbb{E}[\mathbf{g}(x^k, \boldsymbol{\xi}^k)]\|^2] + \frac{\eta}{2\lambda} \|\mathbf{b}(x^k)\|^2 \end{aligned}$$

$$\begin{aligned}
& -\frac{\eta}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|} \right) \\
& \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda B} + \frac{\eta\zeta^2}{2\lambda}.
\end{aligned} \tag{64}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{65}$$

Then substituting (65) into (64) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda B} + \frac{\eta\zeta^2}{2\lambda} \leq -\frac{\eta}{2R} (f(x^k) - f^*) + \frac{\eta\sigma^2}{2\lambda B} + \frac{\eta\zeta^2}{2\lambda}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{2R}\right) (f(x^k) - f^*) + \frac{\eta}{2\lambda} \left(\frac{\sigma^2}{B} + \zeta^2\right).$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \leq \lambda$ and $\|g(x^k, \xi^k)\| \geq \lambda$ NSGD with biased gradient oracle shows linear convergence:

$$\mathbb{E} [f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \frac{R}{\lambda} \left(\frac{\sigma^2}{B} + \zeta^2\right).$$

E.1.3 Third case: $\|\nabla f(x^k)\| \leq \lambda$ and $\|g(x^k, \xi^k)\| \leq \lambda$

Using this in (57), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$ and $\alpha = \|\nabla f(x^k)\|^{-1}$:

$$\begin{aligned}
\mathbb{E} [f(x^{k+1})] - f(x^k) & \stackrel{(57)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E} [G(x^k, \xi^k)] \rangle + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \\
& \stackrel{(5)}{=} -\frac{\eta\alpha}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \mathbb{E} [\|G(x^k, \xi^k)\|^2] + \frac{\eta}{2\alpha} \mathbb{E} [\|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2] \\
& \quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \\
& = -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\alpha} \mathbb{E} [\|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2] \\
& \quad - \frac{\eta}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|} \right) \\
& \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\alpha} \mathbb{E} [\|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2] \\
& \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{\alpha} \mathbb{E} [\|G(x^k, \xi^k)\|^2 + \|\alpha \nabla f(x^k)\|^2]
\end{aligned}$$

$$\begin{aligned}
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{\alpha} \mathbb{E} \left[\left\| \frac{\mathbf{g}(x^k, \boldsymbol{\xi}^k)}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} \right\|^2 + \left\| \frac{\nabla f(x^k)}{\|\nabla f(x^k)\|} \right\|^2 \right] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{2\eta\lambda \|\nabla f(x^k)\|}{\lambda} \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + 2\eta\lambda.
\end{aligned} \tag{66}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{67}$$

Then substituting (67) into (66) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{2} \|\nabla f(x^k)\| + 2\eta\lambda \leq -\frac{\eta}{2R} (f(x^k) - f^*) + 2\eta\lambda.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{2R}\right) (f(x^k) - f^*) + 2\eta\lambda.$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \leq \lambda$ NSGD with biased gradient oracle shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{2R}\right)^N (f(x^0) - f^*) + \lambda R.$$

E.1.4 Fourth case: $\|\nabla f(x^k)\| \geq \lambda$ and $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| \leq \lambda$

Using this in (57), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$ and $\alpha = \lambda^{-1}$:

$$\begin{aligned}
\mathbb{E}[f(x^{k+1})] - f(x^k) &\stackrel{(57)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \boldsymbol{\xi}^k)] \rangle \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\
&\stackrel{(5)}{=} -\frac{\eta\alpha}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|\mathbb{E}[G(x^k, \boldsymbol{\xi}^k)]\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|\mathbb{E}[G(x^k, \boldsymbol{\xi}^k)] - \alpha \nabla f(x^k)\|^2 \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \boldsymbol{\xi}^k)\|^2] \\
&= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \|\mathbb{E}[\lambda G(x^k, \boldsymbol{\xi}^k)] - \nabla f(x^k)\|^2 \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2}
\end{aligned}$$

$$\begin{aligned}
&= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{\lambda} \left\| \mathbb{E} \left[\frac{\lambda \mathbf{g}(x^k, \boldsymbol{\xi}^k)}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right] \right\|^2 \\
&\quad + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\
&= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \left\| \mathbb{E} \left[\left(\frac{\lambda}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right) \mathbf{g}(x^k, \boldsymbol{\xi}^k) \right] \right\|^2 \\
&\quad + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\
&\leq -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \mathbb{E} \left[\left(\frac{\lambda}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|} - 1 \right)^2 \|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|^2 \right] \\
&\quad + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\
&\leq -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta}{2\lambda} \mathbb{E} \left[\frac{\lambda^2}{\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|^2} \|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\|^2 \right] \\
&\quad + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \\
&= -\frac{\eta}{2\lambda} \|\nabla f(x^k)\|^2 + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} + \frac{\eta\lambda}{2} + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} + \frac{\eta\lambda}{2} + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|} \right) + \frac{\eta\lambda}{2} + \frac{\eta}{\lambda} \|\mathbf{b}(x^k)\|^2 \\
&\leq -\frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\lambda}{2} + \frac{\eta\zeta^2}{\lambda}. \tag{68}
\end{aligned}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

Using the convexity assumption of the function, we have the following:

$$f(x^k) - f^* \leq \langle \nabla f(x^k), x^k - x^* \rangle \stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - x^*\| \leq \|\nabla f(x^k)\| \underbrace{\|x^0 - x^*\|}_R.$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f^*}{R}. \tag{69}$$

Then substituting (69) into (68) we obtain:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{4} \|\nabla f(x^k)\| + \frac{\eta\lambda}{2} + \frac{\eta\zeta^2}{\lambda} \leq -\frac{\eta}{4R} (f(x^k) - f^*) + \frac{\eta\lambda}{2} + \frac{\eta\zeta^2}{\lambda}.$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E}[f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{4R}\right) (f(x^k) - f^*) + \frac{\eta\lambda}{2} + \frac{\eta\zeta^2}{\lambda}.$$

Then for $k = 0, 1, 2, \dots, N - 1$ iterations that satisfy the conditions $\|\nabla f(x^k)\| \geq \lambda$ and $\|\mathbf{g}(x^k, \boldsymbol{\xi}^k)\| \leq \lambda$ NSGD with biased gradient oracle shows linear convergence:

$$\mathbb{E}[f(x^N)] - f^* \leq \left(1 - \frac{\eta}{4R}\right)^N (f(x^0) - f^*) + 2\lambda R + \frac{2\zeta^2 R}{\lambda}.$$

Combining all the cases considered, we obtain the convergence rate of NSGD with biased gradient oracle:

$$\mathbb{E}[f(x^N)] - f^* \lesssim \left(1 - \frac{\eta}{R}\right)^N (f(x^0) - f^*) + \frac{MR}{\lambda^2} \left(\frac{\sigma^2}{B} + \zeta^2\right) + \lambda R. \quad (70)$$

E.2 Convergence Results for ZO-NSGD (Proof of the Theorem 5.4)

In order to obtain the convergence rate of ZO-NSGD in the convex setting, we need to substitute the obtained estimates (50) and (51) into the convergence rate of NSGD (70) instead of ζ and σ^2 , respectively. Then the convergence of ZO-NSGD in the convex setup is as follows:

$$\begin{aligned} \mathbb{E}[f(x^N)] - f^* \lesssim & \underbrace{\left(1 - \frac{\eta}{R}\right)^N (f(x^0) - f^*)}_{\textcircled{1}} + \underbrace{\frac{dMR\tilde{\sigma}^2}{\lambda^2 B}}_{\textcircled{2}} + \underbrace{\frac{dMR(L_0 + L_1 M)^2 \gamma^2}{\lambda^2 B}}_{\textcircled{3}} + \underbrace{\frac{d^2 MR \Delta^2}{\lambda^2 B \gamma^2}}_{\textcircled{4}} \\ & + \underbrace{\frac{MR(L_0 + L_1 M)^2 \gamma^2}{\lambda^2}}_{\textcircled{5}} + \underbrace{\frac{d^2 MR \Delta^2}{\lambda^2 \gamma^2}}_{\textcircled{6}} + \underbrace{\lambda R}_{\textcircled{7}}. \end{aligned}$$

From term $\textcircled{7}$, we find the hyperparameter λ :

$$\textcircled{1} : \quad \lambda R \leq \varepsilon \quad \Rightarrow \quad \lambda \leq \frac{\varepsilon}{R}. \quad (71)$$

From term $\textcircled{1}$, we find the number of iterations N required for Algorithm 4 in convex setup to achieve ε -accuracy:

$$\begin{aligned} \textcircled{1} : \quad \left(1 - \frac{\eta}{R}\right)^N (f(x^0) - f^*) \leq \varepsilon \quad & \Rightarrow \quad N \geq \frac{R}{\eta} \log \frac{(f(x^0) - f^*)}{\varepsilon}; \\ N = \tilde{\mathcal{O}}\left(\frac{R}{\eta}\right). \end{aligned} \quad (72)$$

From terms $\textcircled{2}$, we find the batch size B :

$$\begin{aligned} \textcircled{2} : \quad \frac{dMR\tilde{\sigma}^2}{\lambda^2 B} \leq \varepsilon \quad & \Rightarrow \quad B \stackrel{(71)}{\geq} \frac{dMR^3 \tilde{\sigma}^2}{\varepsilon^3}; \\ B = \mathcal{O}\left(\frac{dMR^3 \tilde{\sigma}^2}{\varepsilon^3}\right). \end{aligned} \quad (73)$$

From terms $\textcircled{3}$ and $\textcircled{5}$ we find the smoothing parameter γ :

$$\begin{aligned} \textcircled{3} : \quad \frac{dMR(L_0 + L_1 M)^2 \gamma^2}{\lambda^2 B} \leq \varepsilon \quad & \Rightarrow \quad \gamma \leq \sqrt{\frac{\varepsilon \lambda^2 B}{dMR(L_0 + L_1 M)^2}} \stackrel{(73), (71)}{=} \frac{\tilde{\sigma}}{(L_0 + L_1 M)}; \\ \textcircled{5} : \quad \frac{MR(L_0 + L_1 M)^2 \gamma^2}{\lambda^2} \leq \varepsilon \quad & \Rightarrow \quad \gamma \leq \frac{\sqrt{\varepsilon^3}}{\sqrt{MR^{3/2}}(L_0 + L_1 M)}; \\ \gamma \leq \frac{1}{(L_0 + L_1 M)} \min\left\{\tilde{\sigma}, \frac{\varepsilon^{3/2}}{\sqrt{MR^{3/2}}}\right\} & = \frac{\varepsilon^{3/2}}{(L_0 + L_1 M) \sqrt{MR^{3/2}}}. \end{aligned} \quad (74)$$

From the remaining terms $\textcircled{4}$ and $\textcircled{6}$, we find the maximum allowable level of adversarial noise Δ that still guarantees the convergence of the ZO-NSGD to desired accuracy ε in convex setup:

$$\textcircled{4} : \quad \frac{d^2 MR \Delta^2}{\lambda^2 B \gamma^2} \leq \varepsilon \quad \Rightarrow \quad \Delta \leq \frac{\sqrt{\varepsilon} \lambda \gamma \sqrt{B}}{d \sqrt{MR}} \stackrel{(73), (74), (71)}{=} \frac{\varepsilon^{3/2} \tilde{\sigma}}{\sqrt{d} (L_0 + L_1 M) R^{3/2}};$$

$$\begin{aligned}
\textcircled{6} : \quad \frac{d^2 MR \Delta^2}{\gamma^2 \lambda^2} \leq \varepsilon \quad \Rightarrow \quad \Delta \leq \sqrt{\frac{\gamma^2 \lambda^2 \varepsilon}{d^2 MR}} \stackrel{(71),(74)}{=} \frac{\varepsilon^3}{d(L_0 + L_1 M) R^3}; \\
\Delta \leq \frac{\varepsilon^{3/2}}{\sqrt{d}(L_0 + L_1 M) R^{3/2}} \min \left\{ \tilde{\sigma}, \frac{\varepsilon^{3/2}}{\sqrt{d} R^{3/2}} \right\}.
\end{aligned} \tag{75}$$

In this way, the ZO-NSGD achieves ε -accuracy: $\mathbb{E}[f(x^N) - f^*] \leq \varepsilon$ in convex setup after

$$N \stackrel{(72)}{=} \tilde{\mathcal{O}}\left(\frac{R}{\eta}\right), \quad T = N \cdot B \stackrel{(72),(73)}{=} \mathcal{O}\left(\frac{d\tilde{\sigma}^2 MR^4}{\varepsilon^3 \eta}\right)$$

number of iterations, total number of zero-order oracle calls and at

$$\Delta \stackrel{(75)}{\lesssim} \frac{\varepsilon^{3/2}}{\sqrt{d}(L_0 + L_1 M) R^{3/2}} \min \left\{ \tilde{\sigma}, \frac{\varepsilon^{3/2}}{\sqrt{d} R^{3/2}} \right\}$$

the maximum level of noise with smoothing parameter $\frac{\varepsilon^{3/2}}{(L_0 + L_1 M)\sqrt{M} R^{3/2}}$ (74).

F Additional Clarification

In this section, we would like to clarify the convergence in the case $L_0 = 0$ (Remark 1.3). In this case the problem does not reach a minimum (hence $R = \arg \inf f(x) = +\infty$). Therefore, we exemplify the special case of NSGD (when $\|\nabla f(x^k, \xi^k)\| \geq \sqrt{2}\sigma$ and $\|\nabla f(x^k)\| \geq \sqrt{2}\sigma$), shows that it is possible to achieve the desired accuracy ε in a finite number of iterations.

Let's introduce the notation $G(x^k, \xi^k) = \frac{\nabla f(x^k, \xi^k)}{\|\nabla f(x^k, \xi^k)\|}$, then using (L_0, L_1) -smoothness (see Assumption 1.2):

$$\begin{aligned}
f(x^{k+1}) - f(x^k) &\stackrel{(8)}{\leq} \langle \nabla f(x^k), x^{k+1} - x^k \rangle + \frac{L_0 + L_1 \|\nabla f(x^k)\|}{2} \|x^{k+1} - x^k\|^2 \\
&= -\eta \langle \nabla f(x^k), G(x^k, \xi^k) \rangle + \frac{\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|)}{2} \|G(x^k, \xi^k)\|^2. \tag{76}
\end{aligned}$$

Let us evaluate first summand of (76) with $\alpha = \|\nabla f(x^k)\|^{-1}$:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), G(x^k, \xi^k) \rangle &\stackrel{(5)}{=} -\frac{\alpha \eta}{2} \|\nabla f(x^k)\|^2 - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
&\quad + \frac{\eta}{2\alpha} \|G(x^k, \xi^k) - \alpha \nabla f(x^k)\|^2 \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
&\quad + \frac{\eta}{2\lambda^2 \alpha} \|\lambda G(x^k, \xi^k) - \lambda \alpha \nabla f(x^k)\|^2 \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \|G(x^k, \xi^k)\|^2 \\
&\quad + \frac{\eta}{2\lambda^2 \alpha} \|\text{clip}_\lambda(\nabla f(x^k, \xi^k)) - \text{clip}_\lambda(\nabla f(x^k))\|^2
\end{aligned}$$

Using that clipping is a projection on onto a convex set, namely ball with radius λ , and thus is Lipschitz operator with Lipschitz constant 1, we can obtain:

$$\begin{aligned}
-\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle &\leq -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E}[\|G(x^k, \xi^k)\|^2] \\
&\quad + \frac{\eta}{2\lambda^2 \alpha} \mathbb{E}[\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2]. \tag{77}
\end{aligned}$$

Using this in (77), we have the following with $\eta_k \leq \frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)}$:

$$\mathbb{E}[f(x^{k+1})] - f(x^k) \stackrel{(76)}{\leq} -\eta \langle \nabla f(x^k), \mathbb{E}[G(x^k, \xi^k)] \rangle + \frac{\eta^2 (L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E}[\|G(x^k, \xi^k)\|^2]$$

$$\begin{aligned}
&\stackrel{(77)}{\leq} -\frac{\eta}{2} \|\nabla f(x^k)\| - \frac{\eta}{2\alpha} \mathbb{E} [\|G(x^k, \xi^k)\|^2] + \frac{\eta}{2\lambda^2\alpha} \mathbb{E} [\|\nabla f(x^k, \xi) - \nabla f(x^k)\|^2] \\
&\quad + \frac{\eta^2(L_0 + L_1 \|\nabla f(x^k)\|)}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \\
&= -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{2\lambda^2\alpha} \mathbb{E} [\|\nabla f(x^k, \xi^k) - \nabla f(x^k)\|^2] \\
&\quad - \frac{\eta}{2} \mathbb{E} [\|G(x^k, \xi^k)\|^2] \left(1 - \frac{\eta(L_0 + L_1 \|\nabla f(x^k)\|)}{\|\nabla f(x^k)\|}\right) \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta\sigma^2}{2\lambda^2\alpha} \\
&\leq -\frac{\eta}{2} \|\nabla f(x^k)\| + \frac{\eta}{4} \|\nabla f(x^k)\| \\
&= -\frac{\eta}{4} \|\nabla f(x^k)\|. \tag{78}
\end{aligned}$$

The step size will be constant, depending on the hyperparameter λ :

$$\frac{\|\nabla f(x^k)\|}{2(L_0 + L_1 \|\nabla f(x^k)\|)} = \frac{1}{2\left(L_0 \frac{1}{\|\nabla f(x^k)\|} + L_1\right)} = \frac{\lambda}{2\left(L_0 \frac{\lambda}{\|\nabla f(x^k)\|} + L_1\lambda\right)} \geq \frac{\lambda}{2(L_0 + L_1\lambda)}.$$

Thus, $\eta_k = \eta \leq \frac{\lambda}{2(L_0 + L_1\lambda)}$.

We introduce the hyperparameter of the algorithm $R_s = \|x^0 - s\|$. Then using the convexity assumption of the function, we have the following:

$$\begin{aligned}
f(x^k) - f(s) &\leq \langle \nabla f(x^k), x^k - s \rangle \\
&\stackrel{(6)}{\leq} \|\nabla f(x^k)\| \|x^k - s\| \\
&\leq \|\nabla f(x^k)\| \underbrace{\|x^0 - s\|}_{R_s}.
\end{aligned}$$

Hence we have:

$$\|\nabla f(x^k)\| \geq \frac{f(x^k) - f(s)}{R_s}. \tag{79}$$

Then substituting (79) into (78) we obtain:

$$\mathbb{E} [f(x^{k+1})] - f(x^k) \leq -\frac{\eta}{4} \|\nabla f(x^k)\| \leq -\frac{\eta}{4R_s} (f(x^k) - f(s)).$$

This inequality is equivalent to the trailing inequality:

$$\mathbb{E} [f(x^{k+1})] - f^* \leq \left(1 - \frac{\eta}{4R_s}\right) (f(x^k) - f^*) + \frac{\eta}{4R_s} (f(s) - f^*).$$

Then for $k = 0, 1, 2, \dots, N-1$ iterations that satisfy the conditions $\|\nabla f(x^k, \xi^k)\| \geq \sqrt{2}\sigma$ and $\|\nabla f(x^k)\| \geq \sqrt{2}\sigma$ NSGD shows linear convergence:

$$f(x^N) - f^* \leq \left(1 - \frac{\eta}{4R_s}\right)^N (f(x^0) - f^*) + f(s) - f^*.$$

Thus, we have shown that it is indeed possible to converge to a linear rate of convergence on logistic regression using the hyperparameter R_s .