
SMG: A Shuffling Gradient-Based Method with Momentum

Trang H. Tran¹ Lam M. Nguyen² Quoc Tran-Dinh³

Abstract

We combine two advanced ideas widely used in optimization for machine learning: *shuffling* strategy and *momentum* technique to develop a novel shuffling gradient-based method with momentum, coined **Shuffling Momentum Gradient (SMG)**, for non-convex finite-sum optimization problems. While our method is inspired by momentum techniques, its update is fundamentally different from existing momentum-based methods. We establish state-of-the-art convergence rates of SMG for any shuffling strategy using either constant or diminishing learning rate under standard assumptions (i.e. *L-smoothness* and *bounded variance*). When the shuffling strategy is fixed, we develop another new algorithm that is similar to existing momentum methods, and prove the same convergence rates for this algorithm under the *L-smoothness* and *bounded gradient* assumptions. We demonstrate our algorithms via numerical simulations on standard datasets and compare them with existing shuffling methods. Our tests have shown encouraging performance of the new algorithms.

1. Introduction

Most training tasks in supervised learning are boiled down to solving the following finite-sum minimization:

$$\min_{w \in \mathbb{R}^d} \left\{ F(w) := \frac{1}{n} \sum_{i=1}^n f(w; i) \right\}, \quad (1)$$

where $f(\cdot; i) : \mathbb{R}^d \rightarrow \mathbb{R}$ is a given smooth and possibly nonconvex function for $i \in [n] := \{1, \dots, n\}$.

Problem (1) looks simple, but covers various convex and nonconvex applications in machine learning and statistical

learning, including, but not limited to, logistic regression, multi-kernel learning, conditional random fields, and neural networks. Especially, (1) covers the *empirical risk minimization* as a special case. Solution methods for approximately solving (1) have been widely studied in the literature under different sets of assumptions. The most common approach is perhaps stochastic gradient-type (SGD) methods (Robbins & Monro, 1951; Ghadimi & Lan, 2013; Bottou et al., 2018; Nguyen et al., 2018; 2019) and their variants.

Motivation. While SGD and its variants rely on randomized sampling strategies with replacement, gradient-based methods using without-replacement strategies are often easier and faster to implement. Moreover, practical evidence (Bottou, 2009) has shown that they usually produce a faster decrease of the training loss. Randomized shuffling strategies (also viewed as sampling without replacement) allow the algorithm to use exactly one function component $f(\cdot; i)$ at each epoch compared to SGD, which has only statistical convergence guarantees (e.g., in expectation or with high probability). However, very often, the analysis of shuffling methods is more challenging than SGD due to the lack of statistical independence.

In the deterministic case, single permutation (also called shuffle once, or single shuffling) and incremental gradient methods can be considered as special cases of the shuffling gradient-based methods we study in this paper. One special shuffling strategy is randomized reshuffling, which is broadly used in practice, where we use a different random permutation at each epoch. Alternatively, in recent years, it has been shown that many gradient-based methods with momentum update can notably boost the convergence speed both in theory and practice (Nesterov, 2004; Dozat, 2016; Wang et al., 2020). These methods have been widely used in both convex and nonconvex settings, especially, in deep learning community. Remarkably, Nesterov’s accelerated method (Nesterov, 1983) has made a revolution in large-scale convex optimization in the last two decades, and has been largely exploited in nonconvex problems. The developments we have discussed here motivate us to raise the following research question:

Can we combine both shuffling strategy and momentum scheme to develop new provable gradient-based algorithms for handling (1)?

¹School of Operations Research and Information Engineering, Cornell University, Ithaca, NY, USA. ²IBM Research, Thomas J. Watson Research Center, Yorktown Heights, NY, USA. ³Department of Statistics and Operations Research, The University of North Carolina at Chapel Hill, NC, USA. Correspondence to: Lam M. Nguyen <LamNguyen.MLTD@ibm.com>.

In this paper, we answer this question affirmatively by proposing a novel algorithm called Shuffling Momentum Gradient (SMG). We establish its convergence guarantees for different shuffling strategies, and in particular, randomized reshuffling strategy. We also investigate different variants of our method.

Our contribution. To this end, our contributions in this paper can be summarized as follows.

- (a) We develop a novel shuffling gradient-based method with momentum (Algorithm 1 in Section 2) for approximating a stationary point of the nonconvex minimization problem (1). Our algorithm covers any shuffling strategy ranging from deterministic to randomized, including incremental, single shuffling, and randomized reshuffling variants.
- (b) We establish the convergence of our method in the nonconvex setting and achieve the state-of-the-art $\mathcal{O}(1/T^{2/3})$ convergence rate under standard assumptions (i.e. the L -smoothness and bounded variance conditions), where T is the number of epochs. For randomized reshuffling strategy, we can improve our convergence rate up to $\mathcal{O}(1/(n^{1/3}T^{2/3}))$.
- (c) We study different strategies for selecting learning rates (LR), including constant, diminishing, exponential, and cosine scheduled learning rates. In all cases, we prove the same convergence rate of the corresponding variants without any additional assumption.
- (d) When a single shuffling strategy is used, we show that a momentum strategy can be incorporated directly at each iteration of the shuffling gradient method to obtain a different variant as presented in Algorithm 2. We analyze the convergence of this algorithm and achieve the same $\mathcal{O}(1/T^{2/3})$ epoch-wise convergence rate, but under a bounded gradient assumption instead of the bounded variance as for the SMG algorithm.

Our $\mathcal{O}(1/T^{2/3})$ convergence rate is the best known so far for shuffling gradient-type methods in nonconvex optimization (Nguyen et al., 2020; Mishchenko et al., 2020). However, like (Mishchenko et al., 2020), our SMG method only requires a generalized bounded variance assumption (Assumption 1(c)), which is weaker and more standard than the bounded component gradient assumption used in existing works. Algorithm 2 uses the same set of assumptions as in (Nguyen et al., 2020) to achieve the same rate, but has a momentum update. For the randomized reshuffling strategy, our $\mathcal{O}(1/(n^{1/3}T^{2/3}))$ convergence rate also matches the rate of the without-momentum algorithm in (Mishchenko et al., 2020). It leads to the total of iterations $nT = \mathcal{O}(\sqrt{n}\varepsilon^{-3})$.

We emphasize that, in many existing momentum variants, the momentum $m_i^{(t)}$ is updated recursively at each iteration as $m_{i+1}^{(t)} := \beta m_i^{(t)} + (1 - \beta)g_i^{(t)}$ for a given weight $\beta \in$

$(0, 1)$. This update shows that the momentum $m_{i+1}^{(t)}$ incorporates all the past gradient terms $g_i^{(t)}, g_{i-1}^{(t)}, g_{i-2}^{(t)}, \dots, g_0^{(t)}$ with exponential decay weights $1, \beta, \beta^2, \dots, \beta^i$, respectively. However, in shuffling methods, the convergence guarantee is often obtained in epoch. Based on this observation, we modify the classical momentum update in the shuffling method as shown in Algorithm 1. More specifically, the momentum term $m_0^{(t)}$ is fixed at the beginning of each epoch, and an auxiliary sequence $\{v_i^{(t)}\}$ is introduced to keep track of the gradient average to update the momentum term in the next epoch. This modification makes Algorithm 1 fundamentally different from existing momentum-based methods. This new algorithm still achieves $\mathcal{O}(1/T^{2/3})$ epoch-wise convergence rate under standard assumptions. To the best of our knowledge, our work is the first analyzing convergence rate guarantees of shuffling-type gradient methods with momentum under standard assumptions.

Besides Algorithm 1, we also exploit recent advanced strategies for selecting learning rates, including exponential and cosine scheduled learning rates. These two strategies have shown state-of-the-art performance in practice (Smith, 2017; Loshchilov & Hutter, 2017; Li et al., 2020). Therefore, it is worth incorporating them in shuffling methods.

Related work. Let us briefly review the most related works to our methods studied in this paper.

Shuffling gradient-based methods. Shuffling gradient-type methods for solving (1) have been widely studied in the literature in recent years (Bottou, 2009; Gürbüzbalaban et al., 2019; Shamir, 2016; Haochen & Sra, 2019; Nguyen et al., 2020) for both convex and nonconvex settings. It was empirically investigated in a short note (Bottou, 2009) and also discussed in (Bottou, 2012). These methods have also been implemented in several software packages such as TensorFlow and PyTorch, broadly used in machine learning (Abadi et al., 2015; Paszke et al., 2019).

In the strongly convex case, shuffling methods have been extensively studied in (Ahn et al., 2020; Gürbüzbalaban et al., 2019; Haochen & Sra, 2019; Safran & Shamir, 2020; Nagaraj et al., 2019; Rajput et al., 2020; Nguyen et al., 2020; Mishchenko et al., 2020) under different assumptions. The best known convergence rate in this case is $\mathcal{O}(1/(nT)^2 + 1/(nT^3))$, which matches the lower bound rate studied in (Safran & Shamir, 2020) up to some constant factor. Most results in the convex case are for the incremental gradient variant, which are studied in (Nedic & Bertsekas, 2001; Nedić & Bertsekas, 2001). Convergence results of shuffling methods on the general convex case are investigated in (Shamir, 2016; Mishchenko et al., 2020), where (Mishchenko et al., 2020) provides a unified approach to cover different settings. The authors in (Ying et al., 2017) combine a randomized shuffling and a variance reduction

technique (e.g., SAGA (Defazio et al., 2014) and SVRG (Johnson & Zhang, 2013)) to develop a new variant. They show a linear convergence rate for strongly convex problems but using an energy function, which is unclear how to convert it into known convergence criteria.

In the nonconvex case, (Nguyen et al., 2020) first shows $\mathcal{O}(1/T^{2/3})$ convergence rate for a general class of shuffling gradient methods under the L -smoothness and bounded gradient assumptions on (1). This analysis is then extended in (Mishchenko et al., 2020) to a more relaxed assumption. The authors in (Meng et al., 2019) study different distributed SGD variants with shuffling for strongly convex, general convex, and nonconvex problems. An incremental gradient method for weakly convex problems is investigated in (Li et al., 2019), where the authors show $\mathcal{O}(1/T^{1/2})$ convergence rate as in standard SGD. To the best of our knowledge, the best known rate of shuffling gradient methods for the nonconvex case under standard assumptions is $\mathcal{O}(1/T^{2/3})$ as shown in (Nguyen et al., 2020; Mishchenko et al., 2020).

Our Algorithm 1 developed in this paper is a nontrivial momentum variant of the general shuffling method in (Nguyen et al., 2020) but our analysis uses a standard bounded variance assumption instead of bounded gradient one.

Momentum-based methods. Gradient methods with momentum were studied in early works for convex problems such as heavy-ball, inertial, and Nesterov’s accelerated gradient methods (Polyak, 1964; Nesterov, 2004). Nesterov’s accelerated method is the most influent scheme and achieves optimal convergence rate for convex problems. While momentum-based methods are not yet known to improve theoretical convergence rates in the nonconvex setting, they show significantly encouraging performance in practice (Dozat, 2016; Wang et al., 2020), especially in the deep learning community. However, the momentum strategy has not yet been exploited in shuffling methods.

Adaptive learning rate schemes. Gradient-type methods with adaptive learning rates such as AdaGrad (Duchi et al., 2011) and Adam (Kingma & Ba, 2014) have shown state-of-the-art performance in several optimization applications. Recently, many adaptive schemes for learning rates have been proposed such as diminishing (Nguyen et al., 2020), exponential scheduled (Li et al., 2020), and cosine scheduled (Smith, 2017; Loshchilov & Hutter, 2017). In (Li et al., 2020), the authors analyze convergence guarantees for the exponential and cosine learning rates in SGD. These adaptive learning rates have also been empirically studied in the literature, especially in machine learning tasks. Their convergence guarantees have also been investigated accordingly under certain assumptions.

Although the analyses for adaptive learning rates are generally non-trivial, our theoretical results in this paper are

flexible enough to cover various different learning rates. However, we only exploit the diminishing, exponential scheduled, and cosine scheduled schemes for our shuffling methods with momentum. We establish that in the last two cases, our algorithm still achieves state-of-the-art $\mathcal{O}(T^{-2/3})$ (and possibly up to $\mathcal{O}(n^{-1/3}T^{-2/3})$) epoch-wise rates.

Content. The rest of this paper is organized as follows. Section 2 describes our novel method, Shuffling Momentum Gradient (Algorithm 1). Section 3 investigates its convergence rate under different shuffling-type strategies and different learning rates. Section 4 proposes an algorithm with traditional momentum update for single shuffling strategy. Section 5 presents our numerical simulations. Due to space limit, the convergence analysis of our methods, all technical proofs, and additional experiments are deferred to the Supplementary Document (Supp. Doc.).

2. Shuffling Gradient-Based Method with Momentum

In this section, we describe our new shuffling gradient algorithm with momentum in Algorithm 1. We also compare our algorithm with existing methods and discuss its per-iteration complexity and possible modifications, e.g., mini-batch.

Algorithm 1 Shuffling Momentum Gradient (SMG)

- 1: **Initialization:** Choose $\tilde{w}_0 \in \mathbb{R}^d$ and set $\tilde{m}_0 := \mathbf{0}$.
 - 2: **for** $t := 1, 2, \dots, T$ **do**
 - 3: Set $w_0^{(t)} := \tilde{w}_{t-1}$; $m_0^{(t)} := \tilde{m}_{t-1}$; and $v_0^{(t)} := \mathbf{0}$;
 - 4: Generate an arbitrarily deterministic or random permutation $\pi^{(t)}$ of $[n]$;
 - 5: **for** $i := 0, \dots, n-1$ **do**
 - 6: Query $g_i^{(t)} := \nabla f(w_i^{(t)}; \pi^{(t)}(i+1))$;
 - 7: Choose $\eta_i^{(t)} := \frac{\eta_t}{n}$ and update

$$\begin{cases} m_{i+1}^{(t)} &:= \beta m_0^{(t)} + (1-\beta)g_i^{(t)} \\ v_{i+1}^{(t)} &:= v_i^{(t)} + \frac{1}{n}g_i^{(t)} \\ w_{i+1}^{(t)} &:= w_i^{(t)} - \eta_i^{(t)}m_{i+1}^{(t)} \end{cases} \quad (2)$$
 - 8: **end for**
 - 9: Set $\tilde{w}_t := w_n^{(t)}$ and $\tilde{m}_t := v_n^{(t)}$;
 - 10: **end for**
 - 11: **Output:** Choose $\hat{w}_T \in \{\tilde{w}_0, \dots, \tilde{w}_{T-1}\}$ at random with probability $\mathbb{P}[\hat{w}_T = \tilde{w}_{t-1}] = \frac{\eta_t}{\sum_{t=1}^T \eta_t}$.
-

Clearly, if $\beta = 0$, then Algorithm 1 reduces to the standard shuffling gradient method as in (Nguyen et al., 2020; Mishchenko et al., 2020). Since each inner iteration of our method uses one component $f(\cdot, i)$, we use $\eta_i^{(t)} = \frac{\eta_t}{n}$ in Algorithm 1 to make it consistent with one full gradient, which consists of n gradient components. This form looks

different from the learning rates used in previous work for SGD (Shamir, 2016; Mishchenko et al., 2020), however, it does not necessary make our learning rate smaller. In the same order of training samples n , our learning rate matches the one in (Mishchenko et al., 2020) as well as the state-of-the-art complexity results. In fact, the detailed learning rate η_t used in Algorithm 1 will be specified in Section 3.

Comparison. Unlike existing momentum methods where $m_i^{(t)}$ in (2) is updated recursively as $m_{i+1}^{(t)} := \beta m_i^{(t)} + (1 - \beta)g_i^{(t)}$ for $\beta \in (0, 1)$, we instead fix the first term $m_0^{(i)}$ in (2) at each epoch. It is only updated at the end of each epoch by averaging all the gradient components $\{g_i^{(t)}\}_{i=0}^{n-1}$ evaluated in such an epoch. To avoid storing these gradients, we introduce an auxiliary variable $v_i^{(t)}$ to keep track of the gradient average. Here, we fix the weight β for the sake of our analysis, but it is possible to make β adaptive.

Our new method is based on the following observations. First, while SGD generally uses an unbiased estimator for the full gradient, shuffling gradient methods often do not have such a nice property, making them difficult to analyze. Due to this fact, updating the momentum at each inner iteration does not seem preferable since it could make the estimator deviate from the true gradient. Therefore, we consider updating the momentum after each epoch. Second, unlike the traditional momentum with exponential decay weights, our momentum term $m_0^{(t)}$ is an equal-weighted average of all the past gradients in an epoch, but evaluated at different points $w_i^{(t-1)}$ in the inner loop. Based on these observations, the momentum term should aid the full gradient ∇F instead of its component $\nabla f(\cdot; i)$, leading to the update at the end of each epoch.

It is also worth noting that our algorithm is fundamentally different from variance reduction methods. For instance, SVRG and SARAH variants require evaluating full gradient at each snapshot point while our method always uses single component gradient. Hence, our algorithm does not require a full gradient evaluation at each outer iteration, and our momentum term does not require full gradient of f .

When the learning rate $\eta_i^{(t)}$ is fixed at each epoch as $\eta_i^{(t)} := \frac{\eta_t}{n}$, where $\eta_t > 0$, we can derive from (2) that

$$w_{i+1}^{(t)} = w_i^{(t)} - \frac{(1 - \beta)\eta_t}{n} g_i^{(t)} + \frac{\beta\eta_t}{n(1 - \beta)\eta_{t-1}} e_t,$$

where $e_t := \tilde{w}_{t-1} - \tilde{w}_{t-2} + \beta\eta_{t-1}\tilde{m}_{t-2}$. Here, e_t plays a role as a momentum or an inertial term, but it is different from the usual momentum term. However, we still name Algorithm 1 the *Shuffling Momentum Gradient* since it is inspired by momentum methods.

Per-iteration complexity. The per-iteration complexity of Algorithm 1 is almost the same as in standard shuffling

gradient schemes, see, e.g., (Shamir, 2016). It needs to store two additional vectors $m_{i+1}^{(t)}$ and $v_i^{(t)}$, and performs two vector additions and three scalar-vector multiplications. Such additional costs are very mild.

Shuffling strategies. Our convergence guarantee in Section 3 holds for any permutation $\pi^{(t)}$ of $\{1, 2, \dots, n\}$, including deterministic and randomized ones. Therefore, our method covers any shuffling strategy, including incremental, single shuffling, and randomized reshuffling variants as special cases. We highlight that our convergence result for the randomized reshuffling variant is significantly better than the general ones as we will explain in detail in Section 3.

Mini-batch. Algorithm 1 also works with mini-batches, and our theory can be slightly adapted to establish the same convergence rate for mini-batch variants. However, it remains unclear to us how mini-batches can improve the convergence rate guarantee of Algorithm 1 in the general case where the shuffling strategy is not specified.

3. Convergence Analysis

We analyze the convergence of Algorithm 1 under standard assumptions, which is organized as follows.

3.1. Technical Assumptions

Our analysis relies on the following assumptions:

Assumption 1. *Problem (1) satisfies:*

- (a) (**Boundedness from below**) $\text{dom}(F) \neq \emptyset$ and F is bounded from below, i.e. $F_* := \inf_{w \in \mathbb{R}^d} F(w) > -\infty$.
- (b) (**L -Smoothness**) $f(\cdot; i)$ is L -smooth for all $i \in [n]$, i.e. there exists a universal constant $L > 0$ such that, for all $w, w' \in \text{dom}(F)$, it holds that

$$\|\nabla f(w; i) - \nabla f(w'; i)\| \leq L\|w - w'\|. \quad (3)$$

- (c) (**Generalized bounded variance**) There exist two non-negative and finite constants Θ and σ such that for any $w \in \text{dom}(F)$ we have

$$\frac{1}{n} \sum_{i=1}^n \|\nabla f(w; i) - \nabla F(w)\|^2 \leq \Theta \|\nabla F(w)\|^2 + \sigma^2. \quad (4)$$

Assumption 1(a) is required in any algorithm to guarantee the well-definedness of (1). The L -smoothness (3) is standard in gradient-type methods for both stochastic and deterministic algorithms. From this assumption, we have for any $w, w' \in \text{dom}(F)$ (see (Nesterov, 2004)):

$$F(w) \leq F(w') + \langle \nabla F(w'), w - w' \rangle + \frac{L}{2} \|w - w'\|^2. \quad (5)$$

The condition (4) in Assumption 1(c) reduces to the standard bounded variance condition if $\Theta = 0$. Therefore, (4) is

more general than the bounded variance assumption, which is often used in stochastic optimization. Unlike recent existing works on momentum SGD and shuffling (Chen et al., 2018; Nguyen et al., 2020), we do not assume the bounded gradient assumption on each $f(\cdot; i)$ in Algorithm 1 (see Assumption 2). This condition is stronger than (4).

3.2. Main Result 1 and Its Consequences

Our first main result is the following convergence theorem for Algorithm 1 under any shuffling strategy.

Theorem 1. Suppose that Assumption 1 holds for (1). Let $\{w_i^{(t)}\}_{t=1}^T$ be generated by Algorithm 1 with a fixed momentum weight $0 \leq \beta < 1$ and an epoch learning rate $\eta_i^{(t)} := \frac{\eta_t}{n}$ for every $t \geq 1$. Assume that $\eta_0 = \eta_1$, $\eta_t \geq \eta_{t+1}$, and $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$ for $t \geq 1$, where $K := \max\left\{\frac{5}{2}, \frac{9(5-3\beta)(\Theta+1)}{1-\beta}\right\}$. Then, it holds that

$$\begin{aligned} \mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] &= \frac{1}{\sum_{t=1}^T \eta_t} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 \quad (6) \\ &\leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\sum_{t=1}^T \eta_t} + \frac{9\sigma^2 L^2(5-3\beta)}{(1-\beta)} \left(\frac{\sum_{t=1}^T \eta_t^3}{\sum_{t=1}^T \eta_t} \right). \end{aligned}$$

Remark 1 (Convergence guarantee). When a deterministic permutation $\pi^{(t)}$ is used, our convergence rate can be achieved in a deterministic sense. However, to unify our analysis, we will express our convergence guarantees in expectation, where the expectation is taken over all the randomness generated by $\pi^{(t)}$ and $\hat{w}_T$ up to T iterations. Since we can choose permutations $\pi^{(t)}$ either deterministically or randomly, our bounds in the sequel will hold either deterministically or with probability 1 (w.p.1), respectively. Without loss of generality, we write these results in expectation.

Next, we derive two direct consequences of Theorem 1 by choosing constant and diminishing learning rates.

Corollary 1 (Constant learning rate). Let us fix the number of epochs $T \geq 1$, and choose a constant learning rate $\eta_t := \frac{\gamma}{T^{1/3}}$ for some $\gamma > 0$ such that $\frac{\gamma}{T^{1/3}} \leq \frac{1}{L\sqrt{K}}$ for $t \geq 1$ in Algorithm 1. Then, under the conditions of Theorem 1, $\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2]$ is upper bounded by

$$\frac{1}{T^{2/3}} \left(\frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{9\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)} \right).$$

Consequently, the convergence rate of Algorithm 1 is $\mathcal{O}(T^{-2/3})$ in epoch.

With a constant LR as in Corollary 1, the convergence rate of Algorithm 1 is exactly expressed as

$$\mathcal{O}\left(\frac{[F(\tilde{w}_0) - F_*] + \sigma^2}{T^{2/3}}\right),$$

which matches the best known rate in the literature (Mishchenko et al., 2020; Nguyen et al., 2020) in term of T for general shuffling-type strategies.

Corollary 2 (Diminishing learning rate). Let us choose a diminishing learning rate $\eta_t := \frac{\gamma}{(t+\lambda)^{1/3}}$ for some $\gamma > 0$ and $\lambda \geq 0$ for $t \geq 1$ such that $\eta_1 := \frac{\gamma}{(1+\lambda)^{1/3}} \leq \frac{1}{L\sqrt{K}}$ in Algorithm 1. Then, under the conditions of Theorem 1, after T epochs with $T \geq 2$, we have

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{C_1 + C_2 \log(T-1+\lambda)}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}},$$

where C_1 and C_2 are respectively given by

$$\begin{aligned} C_1 &:= \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{18\sigma^2 L^2(5-3\beta)\gamma^2}{(1-\beta)(1+\lambda)} \quad \text{and} \\ C_2 &:= \frac{9\sigma^2 L^2(5-3\beta)\gamma^2}{(1-\beta)}. \end{aligned}$$

Consequently, the convergence rate of Algorithm 1 is $\mathcal{O}(T^{-2/3} \log(T))$ in epoch.

The diminishing LR $\eta_t := \frac{\gamma}{(t+\lambda)^{1/3}}$ allows us to use large learning rate values at early epochs compared to the constant case. However, we loose a $\log(T)$ factor in the second term of our worst-case convergence rate bound.

We also derive the following two consequences of Theorem 1 for exponential and cosine scheduled learning rates.

Exponential scheduled learning rate. Given an epoch budget $T \geq 1$, and two positive constants $\gamma > 0$ and $\rho > 0$, we consider the following exponential LR, see (Li et al., 2020):

$$\eta_t := \frac{\gamma \alpha^t}{T^{1/3}}, \quad \text{where } \alpha := \rho^{1/T} \in (0, 1). \quad (7)$$

The following corollary shows the convergence of Algorithm 1 using this LR without any additional assumption.

Corollary 3. Let $\{w_i^{(t)}\}_{t=1}^T$ be generated by Algorithm 1 with $\eta_i^{(t)} := \frac{\eta_t}{n}$, where η_t is in (7) such that $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$. Then, under Assumption 1, we have

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma\rho T^{2/3}} + \frac{9\sigma^2 L^2(5-3\beta)\gamma^2}{(1-\beta)\rho T^{2/3}}.$$

Cosine annealing learning rate: Alternatively, given an epoch budget $T \geq 1$, and a positive constant $\gamma > 0$, we consider the following cosine LR for Algorithm 1:

$$\eta_t := \frac{\gamma}{T^{1/3}} \left(1 + \cos \frac{t\pi}{T} \right), \quad t = 1, 2, \dots, T. \quad (8)$$

This LR is adopted from (Loshchilov & Hutter, 2017; Smith, 2017). However, different from these works, we fix our learning rate at each epoch instead of updating it at every iteration as in (Loshchilov & Hutter, 2017; Smith, 2017).

Corollary 4. Let $\{w_i^{(t)}\}_{t=1}^T$ be generated by Algorithm 1 with $\eta_i^{(t)} := \frac{\eta_t}{n}$, where η_t is given by (8) such that $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$. Then, under Assumption 1, and for $T \geq 2$, $\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2]$ is upper bounded by

$$\frac{1}{T^{2/3}} \left(\frac{8[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{144\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)} \right).$$

The scheduled LR (7) and (8) still preserve our best known convergence rate $\mathcal{O}(T^{-2/3})$. Note that the exponential learning rates are available in both TensorFlow and PyTorch, while cosine learning rates are also used in PyTorch.

3.3. Main Result 2: Randomized Reshuffling Variant

A variant of Algorithm 1 is called a **randomized reshuffling variant** if at each iteration t , the generated permutation $\pi^{(t)} = (\pi^{(t)}(1), \dots, \pi^{(t)}(n))$ is uniformly sampled at random without replacement from $\{1, \dots, n\}$. Since the randomized reshuffling strategy is extremely popular in practice, we analyze Algorithm 1 under this strategy.

Theorem 2. Suppose that Assumption 1 holds for (1). Let $\{w_i^{(t)}\}_{t=1}^T$ be generated by Algorithm 1 under a **randomized reshuffling strategy**, a fixed momentum weight $0 \leq \beta < 1$, and an epoch learning rate $\eta_i^{(t)} := \frac{\eta_t}{n}$ for every $t \geq 1$. Assume that $\eta_t \geq \eta_{t+1}$ and $0 < \eta_t \leq \frac{1}{L\sqrt{D}}$ for $t \geq 1$, where $D = \max\left(\frac{5}{3}, \frac{6(5-3\beta)(\Theta+n)}{n(1-\beta)}\right)$ and $\eta_0 = \eta_1$. Then

$$\begin{aligned} \mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] &= \frac{1}{\sum_{t=1}^T \eta_t} \sum_{t=1}^T \eta_t \mathbb{E}[\|\nabla F(\tilde{w}_{t-1})\|^2] \quad (9) \\ &\leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta) \sum_{t=1}^T \eta_t} + \frac{6\sigma^2(5-3\beta)L^2}{n(1-\beta)} \left(\frac{\sum_{t=1}^T \eta_{t-1}^3}{\sum_{t=1}^T \eta_t} \right). \end{aligned}$$

We can derive the following two consequences.

Corollary 5 (Constant learning rate). Let us fix the number of epochs $T \geq 1$, and choose a constant learning rate $\eta_t := \frac{\gamma n^{1/3}}{T^{1/3}}$ for some $\gamma > 0$ such that $\frac{\gamma n^{1/3}}{T^{1/3}} \leq \frac{1}{L\sqrt{D}}$ for $t \geq 1$ in Algorithm 1. Then, under the conditions of Theorem 2, $\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2]$ is upper bounded by

$$\frac{1}{n^{1/3}T^{2/3}} \left(\frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{6\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)} \right).$$

Consequently, the convergence rate of Algorithm 1 is $\mathcal{O}(n^{-1/3}T^{-2/3})$ in epoch.

With a constant LR as in Corollary 5, the convergence rate of Algorithm 1 is improved to

$$\mathcal{O}\left(\frac{[F(\tilde{w}_0) - F_*] + \sigma^2}{n^{1/3}T^{2/3}}\right),$$

which matches the best known rate as in the randomized reshuffling scheme, see, e.g., (Mishchenko et al., 2020).

In this case, the total number of iterations $\mathcal{T}_{\text{tol}} := nT$ is $\mathcal{T}_{\text{tol}} = \mathcal{O}(\sqrt{n}\varepsilon^{-3})$ to obtain $\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \varepsilon^2$. Compared to the complexity $\mathcal{O}(\varepsilon^{-4})$ of SGD, our randomized reshuffling variant is better than SGD if $n \leq \mathcal{O}(\varepsilon^{-2})$.

Corollary 6 (Diminishing learning rate). Let us choose a diminishing learning rate $\eta_t := \frac{\gamma n^{1/3}}{(t+\lambda)^{1/3}}$ for some $\gamma > 0$ and $\lambda \geq 0$ for $t \geq 1$ such that $\eta_1 := \frac{\gamma n^{1/3}}{(1+\lambda)^{1/3}} \leq \frac{1}{L\sqrt{D}}$ in Algorithm 1. Then, under the conditions of Theorem 2, after T epochs with $T \geq 2$, we have

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{C_3 + C_4 \log(T-1+\lambda)}{n^{1/3} [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]},$$

where C_3 and C_4 are respectively given by

$$\begin{aligned} C_3 &:= \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{12\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)(1+\lambda)} \quad \text{and} \\ C_4 &:= \frac{6\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)}. \end{aligned}$$

Consequently, the convergence rate of Algorithm 1 is $\mathcal{O}(n^{-1/3}T^{-2/3} \log(T))$ in epoch.

Remark 2. Algorithm 1 under randomized reshuffling still works with exponential and cosine scheduled learning rates, and our analysis is similar to the one with general shuffling schemes. However, we omit their analysis here.

4. Single Shuffling Variant

In this section, we modify the single-shuffling gradient method by directly incorporating a momentum term at each iteration. We prove for the first time that this variant still achieves state-of-the-art convergence rate guarantee under the smoothness and bounded gradient (i.e. Assumption 2) assumptions as in existing shuffling methods. This variant, though somewhat special, also covers an incremental gradient method with momentum as a special case.

Our new momentum algorithm using single shuffling strategy for solving (1) is presented in Algorithm 2.

Besides fixing the permutation π , Algorithm 2 is different from Algorithm 1 at the point of updating $m_i^{(t)}$. Here, $m_{i+1}^{(t)}$ is updated from $m_i^{(t)}$ instead of $m_0^{(t)}$ as in Algorithm 1. In addition, we update $\tilde{m}_t := m_n^{(t)}$ instead of the epoch gradient average. Note that we can write the main update of Algorithm 2 as

$$w_{i+1}^{(t)} := w_i^{(t)} - \frac{(1-\beta)\eta_t}{n} g_i^{(t)} + \beta(w_i^{(t)} - w_{i-1}^{(t)}),$$

which exactly reduces to existing momentum updates.

Algorithm 2 Single Shuffling Momentum Gradient

```

1: Initialization: Choose  $\tilde{w}_0 \in \mathbb{R}^d$  and set  $\tilde{m}_0 := \mathbf{0}$ ;
2: Generate a permutation  $\pi$  of  $[n]$ ;
3: for  $t := 1, 2, \dots, T$  do
4:   Set  $w_0^{(t)} := \tilde{w}_{t-1}$  and  $m_0^{(t)} := \tilde{m}_{t-1}$ ;
5:   for  $i := 0, \dots, n-1$  do
6:     Query  $g_i^{(t)} := \nabla f(w_i^{(t)}; \pi(i+1))$ ;
7:     Choose  $\eta_i^{(t)} := \frac{\eta_t}{n}$  and update
           
$$\begin{cases} m_{i+1}^{(t)} := \beta m_i^{(t)} + (1-\beta)g_i^{(t)} \\ w_{i+1}^{(t)} := w_i^{(t)} - \eta_i^{(t)} m_{i+1}^{(t)} \end{cases}$$

8:   end for
9:   Set  $\tilde{w}_t := w_n^{(t)}$  and  $\tilde{m}_t := m_n^{(t)}$ ;
10: end for
11: Output: Choose  $\hat{w}_T \in \{\tilde{w}_0, \dots, \tilde{w}_{T-1}\}$  at random
    with probability  $\mathbb{P}[\hat{w}_T = \tilde{w}_{t-1}] = \frac{\eta_t}{\sum_{t=1}^T \eta_t}$ .

```

Incremental gradient method with momentum. If we choose $\pi := [n]$, then we obtain the well-known incremental gradient variant, but with momentum. Hence, Algorithm 2 is still new compared to the standard incremental gradient algorithm (Bertsekas, 2011). To prove convergence of Algorithm 2, we replace Assumption 1(c) by the following:

Assumption 2 (Bounded gradient). There exists $G > 0$ such that $\|\nabla f(x; i)\| \leq G, \forall x \in \text{dom}(F)$ and $i \in [n]$.

Now, we state the convergence of Algorithm 2 in the following theorem as our third main result.

Theorem 3. Let $\{w_i^{(t)}\}_{t=1}^T$ be generated by Algorithm 2 with a LR $\eta_i^{(t)} := \frac{\eta_t}{n}$ and $0 < \eta_t \leq \frac{1}{L}$ for $t \geq 1$. Then, under Assumption 1(a)-(b) and Assumption 2, we have

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{\Delta_1}{(\sum_{t=1}^T \eta_t)(1-\beta^n)} + L^2 G^2 \left(\frac{\sum_{t=1}^T \xi_t^3}{\sum_{t=1}^T \eta_t} \right) + \frac{4\beta^n G^2}{1-\beta^n}, \quad (10)$$

where $\xi_t := \max(\eta_t, \eta_{t-1})$ for $t \geq 2$, $\xi_1 = \eta_1$, and

$$\Delta_1 := 2[F(\tilde{w}_0) - F_*] + \left(\frac{1}{L} + \eta_1 \right) \|\nabla F(\tilde{w}_0)\|^2 + 2L\eta_1^2 G^2.$$

Compared to (9) in Theorem 2, (10) strongly depends on the weight β as $\frac{4\beta^n G^2}{1-\beta^n}$ appears on the right-hand side of (10). Hence, β must be chosen in a specific form to obtain desired convergence rate.

Theorem 3 provides a key bound to derive concrete convergence rates in the following two corollaries.

Corollary 7 (Constant learning rate). In Algorithm 2, let us fix $T \geq 1$ and choose the parameters:

- $\beta := \left(\frac{\nu}{T^{2/3}}\right)^{1/n}$ for some constant $\nu \geq 0$ such that $\beta \leq \left(\frac{R-1}{R}\right)^{1/n}$ for some $R \geq 1$, and
- $\eta_t := \frac{\gamma}{T^{1/3}}$ for some $\gamma > 0$ such that $\frac{\gamma}{T^{1/3}} \leq \frac{1}{L}$.

Then, under the conditions of Theorem 3, we have

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{D_0}{T^{2/3}} + \frac{R\|\nabla F(\tilde{w}_0)\|^2}{T} + \frac{2LG^2R\gamma}{T^{4/3}},$$

where

$$D_0 := \frac{2R}{\gamma} [F(\tilde{w}_0) - F_*] + \frac{R}{\gamma L} \|\nabla F(\tilde{w}_0)\|^2 + L^2 G^2 \gamma^2 + 4\nu G^2 R.$$

Thus the convergence rate of Algorithm 2 is $\mathcal{O}(T^{-2/3})$.

With a constant learning rate as in Corollary 7, the convergence rate of Algorithm 2 is

$$\mathcal{O}\left(\frac{L[F(\tilde{w}_0) - F_*] + \|\nabla F(\tilde{w}_0)\|^2 + G^2}{T^{2/3}}\right),$$

which depends on $L[F(\tilde{w}_0) - F_*]$, $\|\nabla F(\tilde{w}_0)\|^2$, and G^2 , and is slightly different from Corollary 1.

Corollary 8 (Diminishing learning rate). In Algorithm 2, let us choose the parameters:

- $\beta := \left(\frac{\nu}{T^{2/3}}\right)^{1/n}$ for some constant $\nu \geq 0$ such that $\beta \leq \left(\frac{R-1}{R}\right)^{1/n}$ for some $R \geq 1$, and
- a diminishing learning rate $\eta_t := \frac{\gamma}{(t+\lambda)^{1/3}}$ for all $t \in [T]$ for some $\gamma > 0$ and $\lambda \geq 0$ such that $\eta_1 = \frac{\gamma}{(1+\lambda)^{1/3}} \leq \frac{1}{L}$.

Then, under the conditions of Theorem 3, we have

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{D_1}{[(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]} + \frac{4\nu G^2 R}{T^{2/3}} + \frac{L^2 G^2 \gamma^2 \log(T-1+\lambda)}{[(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]},$$

for $T \geq 2$, where

$$D_1 := \frac{2R}{\gamma} [F(\tilde{w}_0) - F_*] + \left[\frac{R}{\gamma L} + \frac{R}{(1+\lambda)^{1/3}} \right] \|\nabla F(\tilde{w}_0)\|^2 + \frac{2RL\gamma G^2}{(1+\lambda)^{2/3}} + \frac{2}{1+\lambda} L^2 G^2 \gamma^2.$$

Thus the convergence rate of Algorithm 2 is $\mathcal{O}(\frac{\log(T)}{T^{2/3}})$.

Remark 3. Algorithm 2 still works with exponential and cosine scheduled LR, and our analysis is similar to the one in Algorithm 1, which is deferred to the Supp. Doc.

5. Numerical Experiments

In order to examine our algorithms, we present two numerical experiments for different nonconvex problems and compare them with some state-of-the-art SGD-type and shuffling gradient methods.

5.1. Models and Datasets

Neural Networks. We test our Shuffling Momentum Gradient (SMG) algorithm using two standard network architectures: fully connected network (FCN) and convolutional neural network (CNN). For the fully connected setting, we train the classic LeNet-300-100 model (LeCun et al., 1998) on the Fashion-MNIST dataset (Xiao et al., 2017) with 60,000 images.

We also use the convolutional LeNet-5 (LeCun et al., 1998) architecture to train the well-known CIFAR-10 dataset (Krizhevsky & Hinton, 2009) with 50,000 samples. We repeatedly run the experiments for 10 random seeds and report the average results. All the algorithms are implemented and run in Python using the PyTorch package (Paszke et al., 2019).

Nonconvex Logistic Regression. Nonconvex regularizers are widely used in statistical learning such as approximating sparsity or gaining robustness. We consider the following nonconvex binary classification problem:

$$\min_{w \in \mathbb{R}^d} \left\{ F(w) := \frac{1}{n} \sum_{i=1}^n \log(1 + \exp(-y_i x_i^\top w)) + \lambda r(w) \right\},$$

where $\{(x_i, y_i)\}_{i=1}^n$ is a set of training samples; and $r(w) := \frac{1}{2} \sum_{j=1}^d \frac{w_j^2}{1+w_j^2}$ is a nonconvex regularizer, and $\lambda := 0.01$ is a regularization parameter. This example was also previously used in (Wang et al., 2019; Tran-Dinh et al., 2019; Nguyen et al., 2020).

We have conducted the experiments on two classification datasets `w8a` (49,749 samples) and `ijcnn1` (91,701 samples) from LIBSVM (Chang & Lin, 2011). The experiment is repeated with random seeds 10 times and the average result is reported.

5.2. Comparing SMG with Other Methods

We compare our SMG algorithm with Stochastic Gradient Descent (SGD) and two other methods: SGD with Momentum (SGD-M) (Polyak, 1964) and Adam (Kingma & Ba, 2014). For the latter two algorithms, we use the hyper-parameter settings recommended and widely used in practice (i.e. momentum: 0.9 for SGD-M, and two hyper-parameters $\beta_1 := 0.9$, $\beta_2 := 0.999$ for Adam). For our new SMG algorithm, we fixed the parameter $\beta := 0.5$ since it usually performs the best in our experiments.

To have a fair comparison, we apply the randomized reshuffling scheme to all methods. Note that shuffling strategies are broadly used in practice and have been implemented in TensorFlow, PyTorch, and Keras (Chollet et al., 2015). We tune each algorithm using constant learning rate and report the best result obtained.

Results. Our first experiment is presented in Figure 1, where we depict the value of “train loss” (i.e. $F(w)$ in (1)) on the y -axis and the “number of effective passes” (i.e. the number of epochs) on the x -axis. It was observed that SGD-M and Adam work well for machine learning tasks (see, e.g., (Ruder, 2017)). Align with this fact, from Figure 1, we also observe that our SMG algorithm and SGD is slightly worse than SGD-M and Adam at the initial stage when training a neural network. However, SMG quickly catches up to Adam and demonstrates a good performance at the end of the training process.

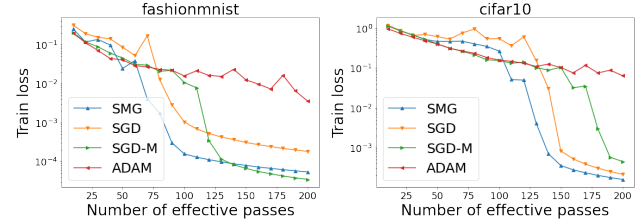


Figure 1. The train loss produced by SMG, SGD, SGD-M, and Adam for Fashion-MNIST and CIFAR-10, respectively.

For the nonconvex logistic regression problem, our result is reported in Figure 2. For two small datasets tested in our experiments, our algorithm performs significantly better than SGD and Adam, and slightly better than SGD with momentum.

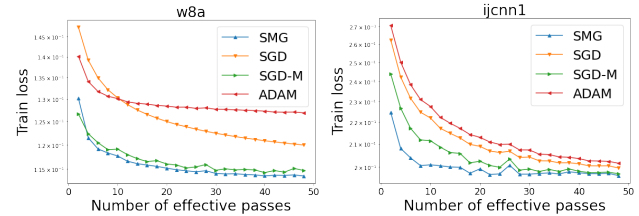


Figure 2. The train loss produced by SMG, SGD, SGD-M, and Adam for the `w8a` and `ijcnn1` datasets, respectively.

5.3. The Choice of Hyper-parameter β

Since the hyper-parameter β plays a critical role in the proposed SMG method, our next experiment is to investigate how this algorithm depends on β , while using the same constant learning rate.

Results. Our result is presented in Figure 3. We can observe from this figure that in the early stage of the training process, the choice $\beta := 0.5$ gives comparably good performance comparing to other smaller values. This choice also results in the best train loss in the end of the training process. However, the difference is not really significant, showing that SMG seems robust to the choice of the momentum weight β in the range of $[0.1, 0.5]$.

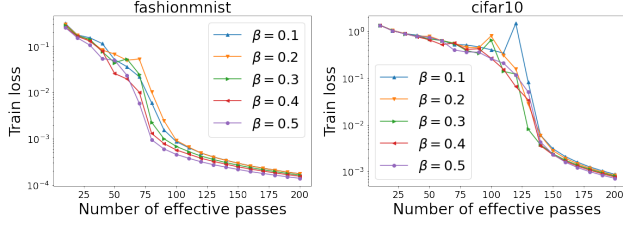


Figure 3. The train loss reported by SMG with different β on the Fashion-MNIST and CIFAR-10 datasets, respectively.

In the nonconvex logistic regression problem, the same choice of β also yields similar outcome for two datasets: w8a and ijcnn1, as shown in Figure 4. We have also experimented with different choices of $\beta \in \{0.6, 0.7, 0.8, 0.9\}$. However, these choices of β do not lead to good performance, and, therefore, we omit to report them here. This finding raises an open question on the optimal choice of β . Our empirical study here shows that the choice of β in $[0.1, 0.5]$ works reasonably well, and $\beta := 0.5$ seems to be the best in our test.

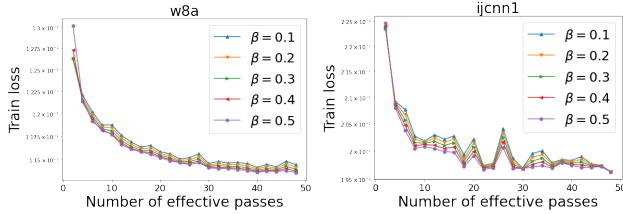


Figure 4. The train loss produced by SMG under different values of β on the w8a and ijcnn1 datasets, respectively.

5.4. Different Learning Rate Schemes

Our last experiment is to examine the effect of different learning rate variants on the performance of our SMG method, i.e. Algorithm 1.

Results. We conduct this test using four different learning rate variants: constant, diminishing, exponential decay, and cosine annealing learning rates. Our results are reported in Figure 5 and Figure 6. From Figure 5, we can observe that the cosine scheduled and the diminishing learning rates converge relatively fast at the early stage. However, the exponential decay and the constant learning rates make faster progress in the last epochs and tend to give the best result at the end of the training process in our neural network training experiments.

For the non-convex logistic regression, Figure 6 shows that the cosine learning rate has certain advantages compared to other choices after a few dozen numbers of epochs.

We note that the detailed settings and additional experiments in this section can be found in the Supp. Doc.

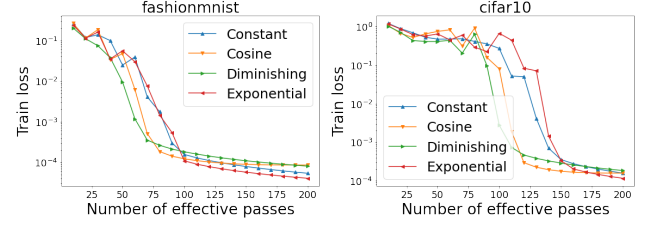


Figure 5. The train loss produced by SMG under four different learning rate schemes on the Fashion-MNIST and CIFAR-10 datasets, respectively.

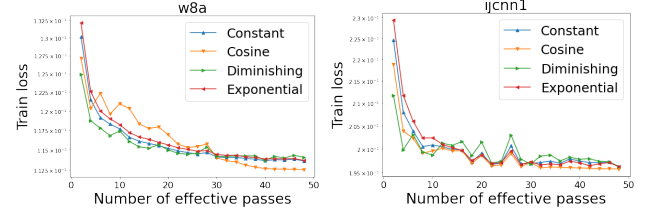


Figure 6. The train loss produced by SMG using four different learning rates on the w8a and ijcnn1 datasets, respectively.

6. Conclusions and Future Work

We have developed two new shuffling gradient algorithms with momentum for solving nonconvex finite-sum minimization problems. Our Algorithm 1 is novel and can work with any shuffling strategy, while Algorithm 2 is similar to existing momentum methods using single shuffling strategy. Our methods achieve the state-of-the-art $\mathcal{O}(1/T^{2/3})$ convergence rate under standard assumptions using different learning rates and shuffling strategies. When a randomized reshuffling strategy is exploited for Algorithm 1, we can further improve our rates by a fraction $n^{1/3}$ of the data size n , matching the best known results in the non-momentum shuffling method. Our numerical results also show encouraging performance of the new methods.

We believe that our analysis framework can be extended to study non-asymptotic rates for some recent adaptive SGD schemes such as Adam (Kingma & Ba, 2014) and AdaGrad (Duchi et al., 2011) as well as variance-reduced methods such as SARAH (Nguyen et al., 2017) under shuffling strategies. It is also interesting to extend our framework to other problems such as minimax and federated learning. We will further investigate these opportunities in the future.

Acknowledgements

The authors would like to thank the reviewers for their useful comments and suggestions which helped to improve the exposition in this paper. The work of Q. Tran-Dinh has partly been supported by the Office of Naval Research under Grant No. ONR-N00014-20-1-2088 (2020-2023).

References

- Abadi, M., Agarwal, A., Barham, P., Brevdo, E., Chen, Z., Citro, C., Corrado, G. S., Davis, A., Dean, J., Devin, M., Ghemawat, S., Goodfellow, I., Harp, A., Irving, G., Isard, M., Jia, Y., Jozefowicz, R., Kaiser, L., Kudlur, M., Levenberg, J., Mané, D., Monga, R., Moore, S., Murray, D., Olah, C., Schuster, M., Shlens, J., Steiner, B., Sutskever, I., Talwar, K., Tucker, P., Vanhoucke, V., Vasudevan, V., Viégas, F., Vinyals, O., Warden, P., Wattenberg, M., Wicke, M., Yu, Y., and Zheng, X. TensorFlow: Large-scale machine learning on heterogeneous systems, 2015. URL <https://www.tensorflow.org/>. Software available from tensorflow.org.
- Ahn, K., Yun, C., and Sra, S. Sgd with shuffling: optimal rates without component convexity and large epoch requirements. *arXiv preprint arXiv:2006.06946*, 2020.
- Bertsekas, D. Incremental proximal methods for large scale convex optimization. *Math. Program.*, 129(2):163–195, 2011.
- Bottou, L. Curiously fast convergence of some stochastic gradient descent algorithms. In *Proceedings of the symposium on learning and data science, Paris*, volume 8, pp. 2624–2633, 2009.
- Bottou, L. Stochastic gradient descent tricks. In *Neural networks: Tricks of the trade*, pp. 421–436. Springer, 2012.
- Bottou, L., Curtis, F. E., and Nocedal, J. Optimization Methods for Large-Scale Machine Learning. *SIAM Rev.*, 60(2):223–311, 2018.
- Chang, C.-C. and Lin, C.-J. LIBSVM: A library for support vector machines. *ACM Transactions on Intelligent Systems and Technology*, 2:27:1–27:27, 2011. Software available at <http://www.csie.ntu.edu.tw/~cjlin/libsvm>.
- Chen, X., Liu, S., Sun, R., and Hong, M. On the convergence of a class of adam-type algorithms for non-convex optimization. *ICLR*, 2018.
- Chollet, F. et al. Keras, 2015. URL <https://github.com/fchollet/keras>.
- Defazio, A., Bach, F., and Lacoste-Julien, S. Saga: A fast incremental gradient method with support for non-strongly convex composite objectives. In *Advances in Neural Information Processing Systems*, pp. 1646–1654, 2014.
- Dozat, T. Incorporating nesterov momentum into ADAM. *ICLR Workshop*, 1:2013—2016, 2016.
- Duchi, J., Hazan, E., and Singer, Y. Adaptive subgradient methods for online learning and stochastic optimization. *Journal of Machine Learning Research*, 12:2121–2159, 2011.
- Ghadimi, S. and Lan, G. Stochastic first-and zeroth-order methods for nonconvex stochastic programming. *SIAM J. Optim.*, 23(4):2341–2368, 2013.
- Gürbüzbalaban, M., Ozdaglar, A., and Parrilo, P. A. Why random reshuffling beats stochastic gradient descent. *Mathematical Programming*, Oct 2019. ISSN 1436-4646. doi: 10.1007/s10107-019-01440-w. URL <http://dx.doi.org/10.1007/s10107-019-01440-w>.
- Haochen, J. and Sra, S. Random shuffling beats sgd after finite epochs. In *International Conference on Machine Learning*, pp. 2624–2633. PMLR, 2019.
- Johnson, R. and Zhang, T. Accelerating stochastic gradient descent using predictive variance reduction. In *NIPS*, pp. 315–323, 2013.
- Kingma, D. P. and Ba, J. ADAM: A Method for Stochastic Optimization. *Proceedings of the 3rd International Conference on Learning Representations (ICLR)*, abs/1412.6980, 2014.
- Krizhevsky, A. and Hinton, G. Learning multiple layers of features from tiny images. Technical report, Citeseer, 2009.
- LeCun, Y., Bottou, L., Bengio, Y., and Haffner, P. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Li, L., Zhuang, Z., and Orabona, F. Exponential step sizes for non-convex optimization. *arXiv preprint arXiv:2002.05273*, 2020.
- Li, X., Zhu, Z., So, A., and Lee, J. D. Incremental methods for weakly convex optimization. *arXiv preprint arXiv:1907.11687*, 2019.
- Loshchilov, I. and Hutter, F. Sgdr: Stochastic gradient descent with warm restarts, 2017.
- Meng, Q., Chen, W., Wang, Y., Ma, Z.-M., and Liu, T.-Y. Convergence analysis of distributed stochastic gradient descent with shuffling. *Neurocomputing*, 337:46–57, 2019.
- Mishchenko, K., Khaled Ragab Bayoumi, A., and Richtárik, P. Random reshuffling: Simple analysis with vast improvements. *Advances in Neural Information Processing Systems*, 33, 2020.

- Nagaraj, D., Jain, P., and Netrapalli, P. Sgd without replacement: Sharper rates for general smooth convex functions. In *International Conference on Machine Learning*, pp. 4703–4711, 2019.
- Nedić, A. and Bertsekas, D. Convergence rate of incremental subgradient algorithms. In *Stochastic optimization: algorithms and applications*, pp. 223–264. Springer, 2001.
- Nedic, A. and Bertsekas, D. P. Incremental subgradient methods for nondifferentiable optimization. *SIAM J. on Optim.*, 12(1):109–138, 2001.
- Nesterov, Y. A method for unconstrained convex minimization problem with the rate of convergence $\mathcal{O}(1/k^2)$. *Doklady AN SSSR*, 269:543–547, 1983. Translated as Soviet Math. Dokl.
- Nesterov, Y. *Introductory lectures on convex optimization: A basic course*, volume 87 of *Applied Optimization*. Kluwer Academic Publishers, 2004.
- Nguyen, L., Nguyen, P. H., van Dijk, M., Richtarik, P., Scheinberg, K., and Takac, M. SGD and Hogwild! convergence without the bounded gradients assumption. In *Proceedings of the 35th International Conference on Machine Learning-Volume 80*, pp. 3747–3755, 2018.
- Nguyen, L. M., Liu, J., Scheinberg, K., and Takáč, M. Sarah: A novel method for machine learning problems using stochastic recursive gradient. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pp. 2613–2621. JMLR. org, 2017.
- Nguyen, L. M., Nguyen, P. H., Richtárík, P., Scheinberg, K., Takáč, M., and van Dijk, M. New convergence aspects of stochastic gradient algorithms. *Journal of Machine Learning Research*, 20(176):1–49, 2019. URL <http://jmlr.org/papers/v20/18-759.html>.
- Nguyen, L. M., Tran-Dinh, Q., Phan, D. T., Nguyen, P. H., and van Dijk, M. A unified convergence analysis for shuffling-type gradient methods. *arXiv preprint arXiv:2002.08246*, 2020.
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimelshein, N., Antiga, L., Desmaison, A., Kopf, A., Yang, E., DeVito, Z., Raison, M., Tejani, A., Chilamkurthy, S., Steiner, B., Fang, L., Bai, J., and Chintala, S. Pytorch: An imperative style, high-performance deep learning library. In *Advances in Neural Information Processing Systems 32*, pp. 8024–8035. Curran Associates, Inc., 2019.
- Polyak, B. T. Some methods of speeding up the convergence of iteration methods. *USSR Computational Mathematics and Mathematical Physics*, 4(5):1–17, 1964.
- Rajput, S., Gupta, A., and Papailiopoulos, D. Closing the convergence gap of sgd without replacement. In *International Conference on Machine Learning*, pp. 7964–7973. PMLR, 2020.
- Robbins, H. and Monro, S. A stochastic approximation method. *The Annals of Mathematical Statistics*, 22(3): 400–407, 1951.
- Ruder, S. An overview of gradient descent optimization algorithms, 2017.
- Safran, I. and Shamir, O. How good is sgd with random shuffling? In *Conference on Learning Theory*, pp. 3250–3284. PMLR, 2020.
- Shamir, O. Without-replacement sampling for stochastic gradient methods. In *Advances in neural information processing systems*, pp. 46–54, 2016.
- Smith, L. N. Cyclical learning rates for training neural networks. In *2017 IEEE Winter Conference on Applications of Computer Vision (WACV)*, pp. 464–472. IEEE, 2017.
- Tran-Dinh, Q., Pham, N. H., Phan, D. T., and Nguyen, L. M. Hybrid stochastic gradient descent algorithms for stochastic nonconvex optimization. *arXiv preprint arXiv:1905.05920*, 2019.
- Wang, B., Nguyen, T. M., Bertozzi, A. L., Baraniuk, R. G., and Osher, S. J. Scheduled restart momentum for accelerated stochastic gradient descent. *arXiv preprint arXiv:2002.10583*, 2020.
- Wang, Z., Ji, K., Zhou, Y., Liang, Y., and Tarokh, V. Spiderboost and momentum: Faster variance reduction algorithms. *Advances in Neural Information Processing Systems*, 32:2406–2416, 2019.
- Xiao, H., Rasul, K., and Vollgraf, R. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. *arXiv preprint arXiv:1708.07747*, 2017.
- Ying, B., Yuan, K., and Sayed, A. H. Convergence of variance-reduced stochastic learning under random reshuffling. *arXiv preprint arXiv:1708.01383*, 2(3):6, 2017.

SUPPLEMENTARY DOCUMENT:

SMG: A Shuffling Gradient-Based Method with Momentum
7. Technical Proofs of Theorem 1 and Its Consequences in Section 3

This section provide the full proofs of Theorem 1 and its consequences in Section 3.

7.1. Auxiliary Notation and Common Expressions

Remark 4. We use the superscript “(t)” for the epoch counter, and the subscript “i” for the counter of the inner loop of the t-th epoch. In addition, at each epoch t, we fix the learning rate $\eta_i^{(t)} := \frac{\eta_t}{n}$ for given $\eta_t > 0$.

In this Supp. Doc. we repeatedly use some common notations to analyze the analysis for three different shuffling strategies, including: Randomized Reshuffling, Shuffle Once and Incremental Gradient. We are ready to introduce the notations here.

For any $t \geq 1$, $i = 0, \dots, n$, $w_0^{(t)} \in \mathbb{R}^d$ generated by Algorithm 1, and two permutations $\pi^{(t)}$ and $\pi^{(t-1)}$ of $[n]$ generated at epochs t and t - 1, we denote

$$A_i^{(t)} := \left\| \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2, \quad (11)$$

$$B_i^{(t)} := \left\| \sum_{j=i}^{n-1} g_j^{(t)} \right\|^2, \quad (12)$$

$$I_t := \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(i+1)) - g_i^{(t)} \right\|^2, \quad (13)$$

$$J_t := \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t-1)}(i+1)) - g_i^{(t-1)} \right\|^2, \quad t \geq 2, \quad (14)$$

$$K_t := n(\Theta + 1) \left\| \nabla F(w_0^{(t)}) \right\|^2 + n\sigma^2, \quad (15)$$

$$N_t = (n + \Theta) \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2, \quad (16)$$

and

$$\xi_t := \begin{cases} \max(\eta_t, \eta_{t-1}) & \text{if } t > 1, \\ \eta_1 & \text{if } t = 1. \end{cases} \quad (17)$$

Note that we adopt the convention $\sum_{j=0}^{-1} g_j^{(t)} = 0$, and $\sum_{j=n}^{n-1} g_j^{(t)} = 0$ in the definitions of $A_i^{(t)}$ and $B_i^{(t)}$.

For each epoch $t = 1, \dots, T$, we denote $\mathcal{F}_t$ by $\sigma(w_0^{(1)}, \dots, w_0^{(t)})$, the σ -algebra generated by the iterates of Algorithm 1 up to the beginning of the epoch t. From Step 4 of Algorithm 1, we observe that the permutation $\pi^{(t)}$ used at time t is independent of the σ -algebra $\mathcal{F}_t$. This will be the key factor of our analysis for Randomized Reshuffling methods. We also denote $\mathbb{E}_t[\cdot]$ by $\mathbb{E}[\cdot | \mathcal{F}_t]$, the conditional expectation on the σ -algebra $\mathcal{F}_t$.

Now, we collect the following expressions derived from Algorithm 1, which are commonly used for every shuffling method. First, from the update rule at Steps 1, 4, and 5 of Algorithm 1, for $t > 1$, we have

$$m_0^{(t)} = \tilde{m}_{t-1} = v_n^{(t-1)} \stackrel{(2)}{=} v_{n-1}^{(t-1)} + \frac{1}{n} g_{n-1}^{(t-1)} = v_0^{(t-1)} + \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t-1)} = \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t-1)}. \quad (18)$$

For the special case when $t = 1$, we set $m_0^{(1)} = \tilde{m}_0 = \mathbf{0} \in \mathbb{R}^d$.

Next, from the update $w_{i+1}^{(t)} := w_i^{(t)} - \eta_i^{(t)} m_{i+1}^{(t)}$ with $\eta_i^{(t)} := \frac{\eta_t}{n}$, for $i = 1, 2, \dots, n-1$, at Step 1 of Algorithm 1, we can

derive

$$w_i^{(t)} \stackrel{(2)}{=} w_{i-1}^{(t)} - \frac{\eta_t}{n} m_i^{(t)} = w_0^{(t)} - \frac{\eta_t}{n} \sum_{j=0}^{i-1} m_{j+1}^{(t)}, \quad t \geq 1, \quad (19)$$

$$w_0^{(t)} = w_n^{(t-1)} \stackrel{(2)}{=} w_{n-1}^{(t-1)} - \frac{\eta_{t-1}}{n} m_n^{(t-1)} = w_i^{(t-1)} - \frac{\eta_{t-1}}{n} \sum_{j=i}^{n-1} m_{j+1}^{(t-1)}, \quad t \geq 2. \quad (20)$$

Note that $\sum_{j=0}^{-1} m_{j+1}^{(t)}$ and $\sum_{j=n}^{n-1} m_{j+1}^{(t-1)} = 0$ by convention, these equations also hold true for $i = 0$ and $i = n$.

7.2. The Proof Sketch of Theorem 1

Due to the technicality of our proofs, we first provide here a proof sketch of our first main result, Theorem 1, while the full proof is given in the next subsections.

The proof of Theorem 1 is divided into the following steps.

- First, from the L -smoothness of f_i and the update of $w_i^{(t)}$, we can bound

$$F(\tilde{w}_t) \leq F(\tilde{w}_{t-1}) - \frac{\eta_t}{2} \|\nabla F(\tilde{w}_{t-1})\|^2 + \mathcal{T}_{[1]},$$

where

$$\mathcal{T}_{[1]} := \frac{\eta_t}{2} \left\| \nabla F(\tilde{w}_{t-1}) - \frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right\|^2.$$

Hence, the key step here is to upper bound $\mathcal{T}_{[1]}$ by $\sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2$.

- Next, we upper bound $\mathcal{T}_{[1]} \leq \frac{\eta_t}{2n} (\beta I_t + (1-\beta)J_t)$, where I_t and J_t are defined by (13) and (14), respectively, which bound the sum of square errors between $g_i^{(t)}$ and the gradient $\nabla f(w_0^{(t)}, \pi^{(t)}(i+1))$ at t and $t-1$, respectively.
- Then, $\beta I_t + (1-\beta)J_t$ can be upper bounded as

$$\beta I_t + (1-\beta)J_t \leq \mathcal{O} \left(\xi_t^2 \sum_{j=0}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \xi_t^2 \sigma^2 \right),$$

as shown in Lemma 4, where $\xi_t := \max\{\eta_{t-1}, \eta_t\}$ defined in (17).

- We further upper bound the sum of the right-hand side of the last estimate as in Lemma 5 in terms $\sum_{t=0}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2$.

Combining these steps together, and using some simplification, we obtain (6) in Theorem 1.

7.3. Technical Lemmas

The proof of Theorem 1 requires the following lemmas as intermediate steps of its proof. Let us first upper bound $\|w_i^{(t)} - w_0^{(t)}\|^2$ and $\|w_i^{(t-1)} - w_0^{(t)}\|^2$ in the following lemma.

Lemma 1. *Let $\{w_i^{(t)}\}$ be generated by Algorithm 1 with $0 \leq \beta < 1$ and $\eta_i^{(t)} := \frac{\eta_t}{n}$ for every $t \geq 1$. Then, for $i = 0, 1, \dots, n$, we have*

$$(a) \quad \|w_i^{(t)} - w_0^{(t)}\|^2 \leq \begin{cases} \frac{\xi_t^2}{n^2} \left[\beta \frac{i^2}{n^2} A_n^{(t-1)} + (1-\beta) A_i^{(t)} \right], & \text{if } t > 1, \\ \frac{\xi_t^2}{n^2} (1-\beta)^2 A_i^{(t)}, & \text{if } t = 1, \end{cases} \quad (21)$$

$$(b) \quad \|w_i^{(t-1)} - w_0^{(t)}\|^2 \leq \begin{cases} \frac{\xi_t^2}{n^2} \left[\beta \frac{(n-i)^2}{n^2} A_n^{(t-2)} + (1-\beta) B_i^{(t-1)} \right], & \text{if } t > 2, \\ \frac{\xi_t^2}{n^2} (1-\beta)^2 B_i^{(t-1)}, & \text{if } t = 2. \end{cases} \quad (22)$$

Next, we upper bound the quantities $A_n^{(t)}$, $\sum_{i=0}^{n-1} A_i^{(t)}$, and $\sum_{i=0}^{n-1} B_i^{(t)}$ defined above in terms of I_t and K_t .

Lemma 2. Under the same setting as of Lemma 1 and Assumption 1(c), for $t \geq 1$, it holds that

$$(a) \quad A_n^{(t)} \leq 3n(I_t + K_t) \quad \text{and} \quad \sum_{i=0}^{n-1} A_i^{(t)} \leq \frac{3n^2}{2}(I_t + K_t), \quad (23)$$

$$(b) \quad \sum_{i=0}^{n-1} B_i^{(t)} \leq 3n^2(I_t + K_t). \quad (24)$$

The results of Lemma 3 below are direct consequences of Lemma 1, Lemma 2, and the fact that $L^2\xi_t^2 \leq \frac{2}{5}$.

Lemma 3. Under the same setting as of Lemma 2 and Assumption 1(b), it holds that

$$\begin{aligned} (a) \quad I_t &\leq \begin{cases} L^2\xi_t^2 \left[\beta(I_{t-1} + K_{t-1}) + \frac{3}{2}(1-\beta)(I_t + K_t) \right], & \text{if } t > 1, \\ \frac{3}{2}L^2\xi_t^2(1-\beta)^2(I_t + K_t), & \text{if } t = 1, \end{cases} \\ (b) \quad J_t &\leq \begin{cases} 3L^2\xi_t^2 \left[\beta(I_{t-2} + K_{t-2}) + (1-\beta)(I_{t-1} + K_{t-1}) \right], & \text{if } t > 2, \\ 3L^2\xi_t^2(I_{t-1} + K_{t-1}), & \text{if } t = 2, \end{cases} \\ (c) \quad \text{If } L^2\xi_t^2 \leq \frac{2}{5}, \text{ then } I_t + K_t &\leq \begin{cases} \beta(I_{t-1} + K_{t-1}) + \frac{5-3\beta}{2}K_t, & \text{if } t > 1, \\ \frac{5-3\beta}{2}K_1, & \text{if } t = 1. \end{cases} \end{aligned}$$

One of our key estimates is to upper bound the quantity $\beta J_t + (1-\beta)I_t$, which is derived in the following lemma.

Lemma 4. Under the same conditions as in Lemma 3, for any $t \geq 2$, we have

$$\beta J_t + (1-\beta)I_t \leq \frac{9(5-3\beta)L^2\xi_t^2}{2} \left(n(\Theta+1) \sum_{j=1}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \frac{n\sigma^2}{1-\beta} \right). \quad (25)$$

Lemma 5. Under the same settings as in Theorem 1 (or Theorem 2, respectively) and $\eta_t \geq \eta_{t+1}$ for $t \geq 1$, we have

$$\sum_{t=1}^T \eta_t \sum_{j=1}^t \beta^{t-j} \|\nabla F(\tilde{w}_{j-1})\|^2 \leq \frac{1}{1-\beta} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2.$$

Lemma 6. Under the same setting as in Theorem 1, we have

$$\frac{\eta_1}{2} \|\nabla F(\tilde{w}_0)\|^2 \leq \frac{F(\tilde{w}_0) - F(\tilde{w}_1)}{(1-\beta)} + \frac{(1-\beta)\eta_1}{4} \|\nabla F(\tilde{w}_0)\|^2 + \frac{9\sigma^2(5-3\beta)L^2}{4(1-\beta)} \cdot \xi_1^3. \quad (26)$$

7.4. The Proof of Theorem 1: Key Estimate for Algorithm 1

First, from the assumption $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$, $t \geq 1$, we have $0 < \eta_t^2 \leq \frac{1}{KL^2}$. Next, from (17), we have $\xi_t = \max(\eta_t; \eta_{t-1})$ for $t > 1$ and $\xi_1 = \eta_1$, which lead to $0 < \xi_t^2 \leq \frac{1}{KL^2}$ for $t \geq 1$. Moreover, from the definition of $K = \max\left(\frac{5}{2}, \frac{9(5-3\beta)(\Theta+1)}{1-\beta}\right)$ in Theorem 1, we have $L^2\xi_t^2 \leq \frac{2}{5}$ and $9L^2\xi_t^2(5-3\beta)(\Theta+1) \leq 1-\beta$ for $t \geq 1$.

Now, letting $i = n$ in Equation (19), for all $t \geq 1$, we have

$$w_n^{(t)} - w_0^{(t)} = -\frac{\eta_t}{n} \sum_{j=0}^{n-1} m_{j+1}^{(t)} \stackrel{(2)}{=} -\frac{\eta_t}{n} \sum_{j=0}^{n-1} (\beta m_0^{(t)} + (1-\beta)g_j^{(t)}) = -\frac{\eta_t}{n} \left(n\beta m_0^{(t)} + (1-\beta) \sum_{j=0}^{n-1} g_j^{(t)} \right).$$

Since $m_0^{(t)} = \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t-1)}$ for $t \geq 2$ (due to (18)), we have the following update

$$w_n^{(t)} - w_0^{(t)} = -\frac{\eta_t}{n} \sum_{j=0}^{n-1} \left(\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)} \right). \quad (27)$$

From the L -smoothness of F in Assumption (1)(b), (27), and the fact that $w_0^{(t+1)} := w_n^{(t)}$, for every epoch $t \geq 2$, we can derive

$$\begin{aligned}
 F(w_0^{(t+1)}) &\stackrel{(5)}{\leq} F(w_0^{(t)}) + \nabla F(w_0^{(t)})^\top (w_0^{(t+1)} - w_0^{(t)}) + \frac{L}{2} \|w_0^{(t+1)} - w_0^{(t)}\|^2 \\
 &\stackrel{(27)}{=} F(w_0^{(t)}) - \eta_t \nabla F(w_0^{(t)})^\top \left(\frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right) \\
 &\quad + \frac{L\eta_t^2}{2} \left\| \frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right\|^2 \\
 &\stackrel{(a)}{=} F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{\eta_t}{2} \left\| \nabla F(w_0^{(t)}) - \frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right\|^2 \\
 &\quad - \frac{\eta_t}{2} (1 - L\eta_t) \left\| \frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right\|^2 \\
 &\leq F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{\eta_t}{2} \left\| \nabla F(w_0^{(t)}) - \frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right\|^2, \tag{28}
 \end{aligned}$$

where (a) follows from the elementary equality $u^\top v = \frac{1}{2} (\|u\|^2 + \|v\|^2 - \|u - v\|^2)$ and the last equality comes from the fact that $\eta_t \leq \frac{\sqrt{2}}{\sqrt{5}L} \leq \frac{1}{L}$, due to the choice of our learning rate η_t .

Next, we use an interesting fact of the permutations $\pi^{(t)}$ and $\pi^{(t-1)}$ to rewrite the full gradient as

$$\begin{aligned}
 \nabla F(w_0^{(t)}) &= \beta \nabla F(w_0^{(t)}) + (1-\beta) \nabla F(w_0^{(t)}) \\
 &= \frac{\beta}{n} \sum_{j=0}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t-1)}(j+1)) + \frac{(1-\beta)}{n} \sum_{j=0}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \\
 &= \frac{1}{n} \sum_{j=0}^{n-1} \left(\beta \nabla f(w_0^{(t)}; \pi^{(t-1)}(j+1)) + (1-\beta) \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right).
 \end{aligned}$$

Let us upper bound the last term of (28) as follows:

$$\begin{aligned}
 \mathcal{T}_{[1]} &:= \frac{\eta_t}{2} \left\| \nabla F(w_0^{(t)}) - \frac{1}{n} \sum_{j=0}^{n-1} (\beta g_j^{(t-1)} + (1-\beta)g_j^{(t)}) \right\|^2 \\
 &= \frac{\eta_t}{2} \left\| \frac{1}{n} \sum_{j=0}^{n-1} \left[\beta \left(\nabla f(w_0^{(t)}; \pi^{(t-1)}(j+1)) - g_j^{(t-1)} \right) + (1-\beta) \left(\nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - g_j^{(t)} \right) \right] \right\|^2 \\
 &\stackrel{(b)}{\leq} \frac{\eta_t}{2n} \sum_{j=0}^{n-1} \left\| \beta \left(\nabla f(w_0^{(t)}; \pi^{(t-1)}(j+1)) - g_j^{(t-1)} \right) + (1-\beta) \left(\nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - g_j^{(t)} \right) \right\|^2 \\
 &\stackrel{(c)}{\leq} \frac{\eta_t}{2n} \sum_{j=0}^{n-1} \left[\beta \left\| \nabla f(w_0^{(t)}; \pi^{(t-1)}(j+1)) - g_j^{(t-1)} \right\|^2 + (1-\beta) \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - g_j^{(t)} \right\|^2 \right] \\
 &\stackrel{(13),(14)}{\leq} \frac{\eta_t}{2n} [\beta J_t + (1-\beta)I_t] \quad (\text{by the definition of } I_t \text{ and } J_t),
 \end{aligned}$$

where (b) is from the Cauchy-Schwarz inequality, (c) follows from the convexity of $\|\cdot\|^2$ for $0 \leq \beta < 1$.

Applying this into (28), we get the following result:

$$F(w_0^{(t+1)}) \leq F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{\eta_t}{2n} [\beta J_t + (1-\beta)I_t]. \tag{29}$$

Remark 5. Note that the derived estimate (29) is true for all shuffling strategies, including Random Reshuffling, Shuffle Once, and Incremental Gradient (under the assumptions of Theorem 1 and 2, respectively).

Applying the result of Lemma 4 for $t \geq 2$, then we obtain

$$F(w_0^{(t+1)}) \leq F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{9(5-3\beta)\eta_t L^2 \xi_t^2}{4n} \left[n(\Theta+1) \sum_{j=1}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \frac{n\sigma^2}{1-\beta} \right],$$

which leads to

$$F(w_0^{(t+1)}) \leq F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{9(5-3\beta)\eta_t L^2 \xi_t^2 (\Theta+1)}{4} \sum_{j=1}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \frac{9\sigma^2(5-3\beta)L^2 \xi_t^3}{4(1-\beta)}.$$

Since ξ_t satisfies $9L^2 \xi_t^2 (5-3\beta)(\Theta+1) \leq 1-\beta$ as proved above, we can deduce from the last estimate that

$$F(w_0^{(t+1)}) \leq F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{(1-\beta)\eta_t}{4} \sum_{j=1}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \frac{9\sigma^2(5-3\beta)L^2}{4(1-\beta)} \cdot \xi_t^3.$$

Rearranging this inequality and noting that $\tilde{w}_{t-1} = w_0^{(t)}$, for $t \geq 2$, we obtain the following estimate:

$$\begin{aligned} \frac{\eta_t}{2} \|\nabla F(\tilde{w}_{t-1})\|^2 &\leq F(\tilde{w}_{t-1}) - F(\tilde{w}_t) + \frac{(1-\beta)\eta_t}{4} \sum_{j=1}^t \beta^{t-j} \|\nabla F(\tilde{w}_{j-1})\|^2 + \frac{9\sigma^2(5-3\beta)L^2}{4(1-\beta)} \cdot \xi_t^3 \\ &\stackrel{1-\beta < 1}{\leq} \frac{F(\tilde{w}_{t-1}) - F(\tilde{w}_t)}{(1-\beta)} + \frac{(1-\beta)\eta_t}{4} \sum_{j=1}^t \beta^{t-j} \|\nabla F(\tilde{w}_{j-1})\|^2 + \frac{9\sigma^2(5-3\beta)L^2}{4(1-\beta)} \cdot \xi_t^3. \end{aligned}$$

For $t = 1$, since $\xi_1 = \eta_1$ as previously defined in (17), from the result of Lemma 6 we have

$$\frac{\eta_1}{2} \|\nabla F(\tilde{w}_0)\|^2 \leq \frac{F(\tilde{w}_0) - F(\tilde{w}_1)}{(1-\beta)} + \frac{(1-\beta)\eta_1}{4} \sum_{j=1}^1 \beta^{1-j} \|\nabla F(\tilde{w}_{j-1})\|^2 + \frac{9\sigma^2(5-3\beta)L^2}{4(1-\beta)} \cdot \xi_1^3.$$

Summing the previous estimate for $t := 2$ to $t := T$, and then using the last one, we obtain

$$\frac{1}{2} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 \leq \frac{F(\tilde{w}_0) - F_*}{(1-\beta)} + \frac{(1-\beta)}{4} \sum_{t=1}^T \eta_t \sum_{j=1}^t \beta^{t-j} \|\nabla F(\tilde{w}_{j-1})\|^2 + \frac{9\sigma^2 L^2 (5-3\beta)}{4(1-\beta)} \sum_{t=1}^T \xi_t^3.$$

Applying Lemma 5 to the last estimate, we get

$$\frac{1}{2} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 \leq \frac{F(\tilde{w}_0) - F_*}{(1-\beta)} + \frac{(1-\beta)}{4} \frac{1}{(1-\beta)} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 + \frac{9\sigma^2 L^2 (5-3\beta)}{4(1-\beta)} \sum_{t=1}^T \xi_t^3,$$

which is equivalent to

$$\sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)} + \frac{9\sigma^2 L^2 (5-3\beta)}{(1-\beta)} \sum_{t=1}^T \xi_t^3. \quad (30)$$

Dividing both sides of the resulting estimate by $\sum_{t=1}^T \eta_t$, we obtain (6). Finally, due to the choice of $\hat{w}_T$ at Step 1 in Algorithm 1, we have $\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] = \frac{1}{\sum_{t=1}^T \eta_t} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2$. $\square$

7.5. The Proof of Corollaries 1 and 2: Constant and Diminishing Learning Rates

The proof of Corollary 1. Since $T \geq 1$ and $\eta_t := \frac{\gamma}{T^{1/3}}$, we have $\eta_t \geq \eta_{t+1}$. We also have $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$ for all $t \geq 1$, and $\sum_{t=1}^T \eta_t = \gamma T^{2/3}$ and $\sum_{t=1}^T \eta_t^3 = \gamma^3$. Substituting these expressions into (6), we obtain

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma T^{2/3}} + \frac{9\sigma^2 L^2 (5-3\beta)}{(1-\beta)} \cdot \frac{\gamma^2}{T^{2/3}}.$$

which is our desired result. $\square$

The proof of Corollary 2. For $t \geq 1$, since $\eta_t = \frac{\gamma}{(t+\lambda)^{1/3}}$, we have $\eta_t \geq \eta_{t+1}$. We also have $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$ for all $t \geq 1$, and

$$\begin{aligned} \sum_{t=1}^T \eta_t &= \gamma \sum_{t=1}^T \frac{1}{(t+\lambda)^{1/3}} \geq \gamma \int_1^T \frac{d\tau}{(\tau+\lambda)^{1/3}} \geq \gamma [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}], \\ \sum_{t=1}^T \eta_{t-1}^3 &= 2\eta_1^3 + \gamma^3 \sum_{t=3}^T \frac{1}{t-1+\lambda} \leq \frac{2\gamma^3}{(1+\lambda)} + \gamma^3 \int_{t=2}^T \frac{d\tau}{\tau-1+\lambda} \leq \gamma^3 \left[\frac{2}{(1+\lambda)} + \log(T-1+\lambda) \right]. \end{aligned}$$

Substituting these expressions into (6), we obtain

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma[(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]} + \frac{9\sigma^2 L^2(5-3\beta)}{(1-\beta)} \cdot \frac{\gamma^2 \left[\frac{2}{(1+\lambda)} + \log(T-1+\lambda) \right]}{[(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]}.$$

Let us define C_1 and C_2 as follows:

$$C_1 := \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{18\sigma^2 L^2(5-3\beta)\gamma^2}{(1-\beta)(1+\lambda)}, \text{ and } C_2 := \frac{9\sigma^2 L^2(5-3\beta)\gamma^2}{(1-\beta)}.$$

Then, we obtain from the last estimate that

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{C_1 + C_2 \log(T-1+\lambda)}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}},$$

which completes our proof. $\square$

7.6. The Proof of Corollaries 3 and 4: Scheduled Learning Rates

The proof of Corollary 3: Exponential LR. For $t \geq 1$, since $\eta_t = \frac{\gamma\alpha^t}{T^{1/3}}$ where $\alpha := \rho^{1/T} \in (0, 1)$, we have $\eta_t \geq \eta_{t+1}$. We also have $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$ for all $t \geq 1$ and

$$\begin{aligned} \sum_{t=1}^T \eta_t &= \frac{\gamma}{T^{1/3}} \sum_{t=1}^T \alpha^t \geq \frac{\gamma}{T^{1/3}} \sum_{t=1}^T \alpha^T = \gamma \rho T^{2/3}, \\ \sum_{t=1}^T \eta_{t-1}^3 &\leq \sum_{t=1}^T \eta_1^3 = \gamma^3 \alpha^3 \leq \gamma^3. \end{aligned}$$

Substituting these expressions into (6), we obtain

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma\rho T^{2/3}} + \frac{9\sigma^2 L^2(5-3\beta)}{(1-\beta)} \cdot \frac{\gamma^2}{\rho T^{2/3}}.$$

which is our desired result. $\square$

The proof of Corollary 4: Cosine LR. First, we would like to show that $\sum_{t=1}^T \cos \frac{t\pi}{T} = -1$. Let us denote $A := \sum_{t=1}^T \cos \frac{t\pi}{T}$. Multiplying this sum by $\sin \frac{\pi}{2T}$, we get

$$\begin{aligned} A \cdot \sin \frac{\pi}{2T} &= \sum_{t=1}^T \cos \frac{t\pi}{T} \sin \frac{\pi}{2T} \stackrel{(a)}{=} \frac{1}{2} \sum_{t=1}^T \left(\sin \frac{(2t+1)\pi}{2T} - \sin \frac{(2t-1)\pi}{2T} \right) \\ &= \frac{1}{2} \left(\sin \frac{(2T+1)\pi}{2T} - \sin \frac{\pi}{2T} \right) = \frac{1}{2} \left(-\sin \frac{\pi}{2T} - \sin \frac{\pi}{2T} \right) = -\sin \frac{\pi}{2T}, \end{aligned}$$

where (a) comes from the identity $\cos a \cdot \sin b = \frac{1}{2}(\sin(a+b) - \sin(a-b))$. Since $\sin \frac{\pi}{2T}$ is nonzero, we obtain $A = \sum_{t=1}^T \cos \frac{t\pi}{T} = -1$.

For $1 \leq t \leq T$, since $\eta_t = \frac{\gamma}{T^{1/3}} (1 + \cos \frac{t\pi}{T})$, we have $\eta_t \geq \eta_{t+1}$. We also have $0 < \eta_t \leq \frac{1}{L\sqrt{K}}$ for all $t \geq 1$ and

$$\begin{aligned} \sum_{t=1}^T \eta_t &= \frac{\gamma}{T^{1/3}} \sum_{t=1}^T (1 + \cos \frac{t\pi}{T}) \stackrel{(b)}{=} \frac{\gamma}{T^{1/3}} (T-1) \geq \frac{\gamma}{T^{1/3}} \cdot \frac{T}{2} = \frac{\gamma T^{2/3}}{2}, \text{ for } T \geq 2, \\ \sum_{t=1}^T \eta_{t-1}^3 &\leq \frac{\gamma^3}{T} \sum_{t=1}^T 2^3 = 8\gamma^3, \text{ since } (1 + \cos \frac{t\pi}{T})^3 \leq 2^3 \text{ for all } t \geq 1, \end{aligned}$$

where (b) follows since $\sum_{t=1}^T \cos \frac{t\pi}{T} = -1$ as shown above. Substituting these expressions into (6), we obtain

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{8[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma T^{2/3}} + \frac{9\sigma^2 L^2(5-3\beta)}{(1-\beta)} \cdot \frac{16\gamma^2}{T^{2/3}}.$$

which is our desired estimate. $\square$

7.7. The Proof of Technical Lemmas

We now provide the proof of six lemmas that serve for the proof of Theorem 1 above.

7.7.1. PROOF OF LEMMA 1: UPPER BOUNDING TWO TERMS $\|w_i^{(t)} - w_0^{(t)}\|^2$ AND $\|w_i^{(t-1)} - w_0^{(t)}\|^2$

Remark 6. Note that the result of Lemma 1 is true for all shuffling strategies including Random Reshuffling, Shuffle Once and Incremental Gradient (under the assumptions of Theorem 1 and 2, respectively).

(a) From Equation (19), for $i = 0, 1, \dots, n$ and $t \geq 1$, we have

$$w_i^{(t)} - w_0^{(t)} = -\frac{\eta_t}{n} \sum_{j=0}^{i-1} m_{j+1}^{(t)} \stackrel{(2)}{=} -\frac{\eta_t}{n} \sum_{j=0}^{i-1} (\beta m_0^{(t)} + (1-\beta)g_j^{(t)}) = -\frac{\eta_t}{n} \left(i\beta m_0^{(t)} + (1-\beta) \sum_{j=0}^{i-1} g_j^{(t)} \right).$$

Moreover, by (17), we have $\eta_t \leq \xi_t$ for all $t \geq 1$. Therefore, for $t > 1$ and $i = 0, 1, \dots, n$, we can derive that

$$\begin{aligned} \|w_i^{(t)} - w_0^{(t)}\|^2 &\leq \frac{\xi_t^2}{n^2} \left\| i\beta m_0^{(t)} + (1-\beta) \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 \\ &\stackrel{(18)}{=} \frac{\xi_t^2}{n^2} \left\| \beta \frac{i}{n} \sum_{j=0}^{n-1} g_j^{(t-1)} + (1-\beta) \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 \\ &\stackrel{(a)}{\leq} \frac{\xi_t^2}{n^2} \left(\beta \left\| \frac{i}{n} \sum_{j=0}^{n-1} g_j^{(t-1)} \right\|^2 + (1-\beta) \left\| \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 \right) \\ &\stackrel{(11)}{=} \frac{\xi_t^2}{n^2} \left(\beta \frac{i^2}{n^2} A_n^{(t-1)} + (1-\beta) A_i^{(t)} \right), \end{aligned}$$

where (a) follows from the convexity of $\|\cdot\|^2$ and $0 \leq \beta < 1$.

For $t = 1$ and $i = 0, 1, \dots, n$, we have $m_0^{(t)} = \mathbf{0}$, which leads to

$$\|w_i^{(t)} - w_0^{(t)}\|^2 \leq \frac{\xi_t^2}{n^2} \left\| i\beta m_0^{(t)} + (1-\beta) \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 = \frac{\xi_t^2}{n^2} \left\| (1-\beta) \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 = \frac{\xi_t^2}{n^2} (1-\beta)^2 A_i^{(t)}.$$

Combining these two cases, we obtain (21).

(b) Using similar argument for the second equation (20), for $t \geq 2$ and $i = 0, 1, \dots, n$; we can derive that

$$\begin{aligned} w_i^{(t-1)} - w_0^{(t)} &= \frac{\eta_{t-1}}{n} \sum_{j=i}^{n-1} m_{j+1}^{(t-1)} \stackrel{(2)}{=} \frac{\eta_{t-1}}{n} \sum_{j=i}^{n-1} [\beta m_0^{(t-1)} + (1-\beta)g_j^{(t-1)}] \\ &= \frac{\eta_{t-1}}{n} \left[(n-i)\beta m_0^{(t-1)} + (1-\beta) \sum_{j=i}^{n-1} g_j^{(t-1)} \right]. \end{aligned}$$

Note again that $\eta_{t-1} \leq \xi_t$, for $t > 2$ and $i = 0, 1, \dots, n$; we can show that

$$\begin{aligned} \|w_i^{(t-1)} - w_0^{(t)}\|^2 &\leq \frac{\xi_t^2}{n^2} \left\| (n-i)\beta m_0^{(t-1)} + (1-\beta) \sum_{j=i}^{n-1} g_j^{(t-1)} \right\|^2 \\ &\stackrel{(18)}{=} \frac{\xi_t^2}{n^2} \left\| \beta \frac{n-i}{n} \sum_{j=0}^{n-1} g_j^{(t-2)} + (1-\beta) \sum_{j=i}^{n-1} g_j^{(t-1)} \right\|^2 \\ &\stackrel{(b)}{\leq} \frac{\xi_t^2}{n^2} \left[\beta \left\| \frac{n-i}{n} \sum_{j=0}^{n-1} g_j^{(t-2)} \right\|^2 + (1-\beta) \left\| \sum_{j=i}^{n-1} g_j^{(t-1)} \right\|^2 \right] \\ &\stackrel{(11),(12)}{=} \frac{\xi_t^2}{n^2} \left[\beta \frac{(n-i)^2}{n^2} A_n^{(t-2)} + (1-\beta) B_i^{(t-1)} \right], \end{aligned}$$

where in (b) we use again the convexity of $\|\cdot\|^2$ and $0 \leq \beta < 1$.

For $t = 2$, we easily get

$$\|w_i^{(t-1)} - w_0^{(t)}\|^2 \leq \frac{\xi_t^2}{n^2} \left\| (n-i)\beta m_0^{(t-1)} + (1-\beta) \sum_{j=i}^{n-1} g_j^{(t-1)} \right\|^2 \leq \frac{\xi_t^2}{n^2} \left\| (1-\beta) \sum_{j=i}^{n-1} g_j^{(t-1)} \right\|^2 = \frac{\xi_t^2}{n^2} (1-\beta)^2 B_i^{(t-1)}.$$

Combining the two cases above, we obtain (22). $\square$

7.7.2. PROOF OF LEMMA 2: UPPER BOUNDING THE TERMS $A_n^{(t)}$, $\sum_{i=0}^{n-1} A_i^{(t)}$, AND $\sum_{i=0}^{n-1} B_i^{(t)}$

(a) Let first upper bound the term $A_i^{(t)}$ defined by (11). Indeed, for $i = 0, \dots, n-1$ and $t \geq 1$, we have

$$\begin{aligned} A_i^{(t)} &= \left\| \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 \\ &\stackrel{(a)}{\leq} 3 \left\| \sum_{j=0}^{i-1} \left(g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right) \right\|^2 + 3 \left\| \sum_{j=0}^{i-1} \left(\nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right) \right\|^2 + 3 \left\| \sum_{j=0}^{i-1} \nabla F(w_0^{(t)}) \right\|^2 \\ &\stackrel{(a)}{\leq} 3i \sum_{j=0}^{i-1} \left\| g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 + 3i \sum_{j=0}^{i-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 + 3i^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 \\ &\leq 3i \sum_{j=0}^{n-1} \left\| g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 + 3i \sum_{j=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 + 3i^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 \\ &\stackrel{(b)}{\leq} 3iI_t + 3in \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] + 3i^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 \\ &\leq 3iI_t + 3(i^2 + in\Theta) \left\| \nabla F(w_0^{(t)}) \right\|^2 + 3in\sigma^2, \end{aligned}$$

where we use the Cauchy-Schwarz inequality in (a), and (b) follows from Assumption 1(c).

Letting $i = n$ in the above estimate, we obtain the first estimate of Lemma 2, i.e.:

$$A_n^{(t)} \leq 3nI_t + 3n^2(\Theta + 1) \left\| \nabla F(w_0^{(t)}) \right\|^2 + 3n^2\sigma^2 \leq 3n(I_t + K_t), \quad t \geq 1.$$

Now we are ready to calculate the sum $\sum_{i=0}^{n-1} A_i^{(t)}$ as

$$\begin{aligned} \sum_{i=0}^{n-1} A_i^{(t)} &\leq 3I_t \sum_{i=0}^{n-1} i + 3 \sum_{i=0}^{n-1} (i^2 + in\Theta) \left\| \nabla F(w_0^{(t)}) \right\|^2 + 3n\sigma^2 \sum_{i=0}^{n-1} i \\ &\leq \frac{3n^2}{2} I_t + \frac{3n^3}{2} (\Theta + 1) \left\| \nabla F(w_0^{(t)}) \right\|^2 + \frac{3n^3\sigma^2}{2} \leq \frac{3n^2}{2} (I_t + K_t), \quad t \geq 1, \end{aligned}$$

where the last line follows from $\sum_{i=0}^{n-1} i \leq \frac{n^2}{2}$ and $\sum_{i=0}^{n-1} i^2 \leq \frac{n^3}{2}$. This proves the second estimate of Lemma 2.

(b) Using similar argument as above, for $i = 0, 1, \dots, n-1$ and $t \geq 1$, we can derive

$$\begin{aligned} B_i^{(t)} &= \left\| \sum_{j=i}^{n-1} g_j^{(t)} \right\|^2 \stackrel{(a)}{\leq} 3 \left\| \sum_{j=i}^{n-1} \left(g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right) \right\|^2 \\ &\quad + 3 \left\| \sum_{j=i}^{n-1} \left(\nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right) \right\|^2 + 3 \left\| \sum_{j=i}^{n-1} \nabla F(w_0^{(t)}) \right\|^2 \\ &\stackrel{(a)}{\leq} 3(n-i) \sum_{j=i}^{n-1} \left\| g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \end{aligned}$$

$$\begin{aligned}
 & + 3(n-i) \sum_{j=i}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 + 3(n-i)^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 \\
 & \stackrel{(13)}{\leq} 3(n-i)I_t + 3(n-i) \sum_{j=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 + 3(n-i)n \left\| \nabla F(w_0^{(t)}) \right\|^2 \\
 & \stackrel{(b)}{\leq} 3(n-i)I_t + 3(n-i)n \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] + 3(n-i)n \left\| \nabla F(w_0^{(t)}) \right\|^2 \\
 & \leq 3(n-i)I_t + 3(n-i)n(\Theta+1) \left\| \nabla F(w_0^{(t)}) \right\|^2 + 3(n-i)n\sigma^2,
 \end{aligned}$$

where we use again the Cauchy-Schwarz inequality in (a), and (b) follows from Assumption 1(c). Finally, for all $t \geq 1$, we calculate the sum $\sum_{i=0}^{n-1} B_i^{(t)}$ as follows:

$$\begin{aligned}
 \sum_{i=0}^{n-1} B_i^{(t)} & \leq 3I_t \sum_{i=0}^{n-1} (n-i) + 3 \sum_{i=0}^{n-1} (n-i)n(\Theta+1) \left\| \nabla F(w_0^{(t)}) \right\|^2 + 3n\sigma^2 \sum_{i=0}^{n-1} (n-i) \\
 & \leq 3n^2 I_t + 3n^3(\Theta+1) \left\| \nabla F(w_0^{(t)}) \right\|^2 + 3n^3 \sigma^2 \\
 & \leq 3n^2 (I_t + K_t),
 \end{aligned}$$

where the last line follows from the fact that $\sum_{i=0}^{n-1} (n-i) \leq n^2$. This proves (24). $\square$

7.7.3. PROOF OF LEMMA 3: UPPER BOUNDING THE TERMS I_t , J_t , AND $I_t + K_t$

(a) First, for every $t \geq 1$ and $g_i^{(t)} := \nabla f(w_i^{(t)}; \pi^{(t)}(i+1))$, we have

$$\begin{aligned}
 I_t & = \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(i+1)) - g_i^{(t)} \right\|^2 \\
 & \stackrel{\text{Step 1}}{=} \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(i+1)) - \nabla f(w_i^{(t)}; \pi^{(t)}(i+1)) \right\|^2 \\
 & \leq L^2 \sum_{i=0}^{n-1} \left\| w_i^{(t)} - w_0^{(t)} \right\|^2,
 \end{aligned} \tag{31}$$

where the last inequality follows from Assumption 1(b).

Applying Lemma 1 to (31), for $t > 1$, we have

$$\begin{aligned}
 I_t & \leq \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} \left[\beta \frac{i^2}{n^2} A_n^{(t-1)} + (1-\beta) A_i^{(t)} \right] \\
 & \leq \frac{L^2 \xi_t^2}{n^2} \left[\beta \frac{A_n^{(t-1)}}{n^2} \sum_{i=0}^{n-1} i^2 + (1-\beta) \sum_{i=0}^{n-1} A_i^{(t)} \right] \\
 & \stackrel{(a)}{\leq} \frac{L^2 \xi_t^2}{n^2} \left[\beta \frac{n}{3} 3n (I_{t-1} + K_{t-1}) + (1-\beta) \frac{3n^2}{2} (I_t + K_t) \right] \\
 & = L^2 \xi_t^2 \left[\beta (I_{t-1} + K_{t-1}) + \frac{3}{2} (1-\beta) (I_t + K_t) \right],
 \end{aligned}$$

where (a) follows from the result (23) in Lemma 2 and the fact that $\sum_{i=0}^{n-1} i^2 \leq \frac{n^3}{3}$.

For $t = 1$, by applying Lemma 1 and 2 consecutively to (31) we have

$$I_t \leq L^2 \sum_{i=0}^{n-1} \left\| w_i^{(t)} - w_0^{(t)} \right\|^2 \stackrel{(21)}{\leq} \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} \left((1-\beta)^2 A_i^{(t)} \right) \leq \frac{L^2 \xi_t^2 (1-\beta)^2}{n^2} \sum_{i=0}^{n-1} A_i^{(t)}$$

$$\stackrel{(23)}{\leq} \frac{L^2 \xi_t^2 (1-\beta)^2}{n^2} \frac{3n^2}{2} (I_t + K_t) = \frac{3}{2} L^2 \xi_t^2 (1-\beta)^2 (I_t + K_t). \quad (32)$$

(b) Using the similar argument as above we can derive that

$$\begin{aligned} J_t &= \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t-1)}(i+1)) - g_i^{(t-1)} \right\|^2 \\ &\stackrel{\text{Step 1}}{=} \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t-1)}(i+1)) - \nabla f(w_i^{(t-1)}; \pi^{(t-1)}(i+1)) \right\|^2 \\ &\stackrel{(3)}{\leq} L^2 \sum_{i=0}^{n-1} \|w_i^{(t-1)} - w_0^{(t)}\|^2. \end{aligned}$$

Applying Lemma 1 for $t > 2$ we get to the above estimate, we get

$$\begin{aligned} J_t &\leq \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} \left[\beta \frac{(n-i)^2}{n^2} A_n^{(t-2)} + (1-\beta) B_i^{(t-1)} \right] \\ &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta \frac{A_n^{(t-2)}}{n^2} \sum_{i=0}^{n-1} (n-i)^2 + (1-\beta) \sum_{i=0}^{n-1} B_i^{(t-1)} \right] \\ &\stackrel{(b)}{\leq} \frac{L^2 \xi_t^2}{n^2} \left[\beta n \cdot 3n (I_{t-2} + K_{t-2}) + (1-\beta) \cdot 3n^2 (I_{t-1} + K_{t-1}) \right] \\ &= 3L^2 \xi_t^2 \left[\beta (I_{t-2} + K_{t-2}) + (1-\beta) (I_{t-1} + K_{t-1}) \right], \end{aligned}$$

where (b) follows from the results (23), (24) in Lemma 2 and the fact that $\sum_{i=0}^{n-1} (n-i)^2 \leq n^3$.

Now applying Lemma 1 for the special case $t = 2$, we obtain

$$J_t \leq \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} (1-\beta)^2 B_i^{(t-1)} \leq \frac{L^2 \xi_t^2}{n^2} (1-\beta)^2 \cdot 3n^2 (I_{t-1} + K_{t-1}) \leq 3L^2 \xi_t^2 (I_{t-1} + K_{t-1}),$$

where we use the result (24) from Lemma 2 and the fact that $1-\beta \leq 1$.

(c) First, from the result of Part (a), for $t > 1$, we have

$$I_t \leq L^2 \xi_t^2 \left[\beta (I_{t-1} + K_{t-1}) + \frac{3}{2} (1-\beta) K_t \right] + \frac{3}{2} (1-\beta) L^2 \xi_t^2 I_t.$$

Since $L^2 \xi_t^2 \leq \frac{2}{5}$ (due to the choice of our learning rate) and $1-\beta \leq 1$, we further get

$$I_t \leq \frac{2}{5} \left[\beta (I_{t-1} + K_{t-1}) + \frac{3}{2} (1-\beta) K_t \right] + \frac{3}{5} I_t,$$

which is equivalent to

$$I_t \leq \beta (I_{t-1} + K_{t-1}) + \frac{3}{2} (1-\beta) K_t.$$

Adding K_t to both sides of this inequality, for $t > 1$, we get

$$I_t + K_t \leq \beta (I_{t-1} + K_{t-1}) + \frac{5-3\beta}{2} K_t. \quad (33)$$

Finally, for $t = 1$, since $L^2 \xi_t^2 \leq \frac{2}{5}$ and $1-\beta \leq 1$, we have

$$I_t \leq \frac{3}{2} L^2 \xi_t^2 (1-\beta)^2 (I_t + K_t) \leq \frac{3}{5} (1-\beta)^2 I_t + \frac{3}{5} (1-\beta)^2 K_t \leq \frac{3}{5} I_t + \frac{3}{5} (1-\beta) K_t,$$

which is equivalent to

$$\frac{2}{5}I_t \leq \frac{3}{5}(1-\beta)K_t. \quad (34)$$

This leads to our desired result for $t = 1$ as

$$I_1 + K_1 \leq \frac{3}{2}(1-\beta)K_1 + K_1 = \frac{5-3\beta}{2}K_1. \quad (35)$$

Combining two cases, we obtain the desired result in Part (c). $\square$

7.7.4. PROOF OF LEMMA 4: UPPER BOUNDING THE KEY QUANTITY $\beta J_t + (1-\beta)I_t$

First, we analyze the case $t > 2$ using the results of Lemma 3 as follows:

$$\begin{aligned} \beta J_t + (1-\beta)I_t &\leq 3\beta L^2 \xi_t^2 \left[\beta(I_{t-2} + K_{t-2}) + (1-\beta)(I_{t-1} + K_{t-1}) \right] \\ &\quad + (1-\beta)L^2 \xi_t^2 \left[\beta(I_{t-1} + K_{t-1}) + \frac{3}{2}(1-\beta)(I_t + K_t) \right] \\ &\leq L^2 \xi_t^2 [2(I_t + K_t) + 4\beta(I_{t-1} + K_{t-1}) + 3\beta^2(I_{t-2} + K_{t-2})] \\ &=: L^2 \xi_t^2 P_t, \end{aligned} \quad (36)$$

where the last two lines follow since $1-\beta \leq 1$ and $P_t := 2(I_t + K_t) + 4\beta(I_{t-1} + K_{t-1}) + 3\beta^2(I_{t-2} + K_{t-2})$ for $t > 2$.

Next, we bound the term P_t in (36) as

$$\begin{aligned} P_t &= 2(I_t + K_t) + 4\beta(I_{t-1} + K_{t-1}) + 3\beta^2(I_{t-2} + K_{t-2}) \\ &\leq 2\left[\beta(I_{t-1} + K_{t-1}) + \frac{5-3\beta}{2}K_t\right] + 4\beta(I_{t-1} + K_{t-1}) + 3\beta^2(I_{t-2} + K_{t-2}) \quad \text{apply (33) for } t \\ &= (5-3\beta)K_t + 6\beta(I_{t-1} + K_{t-1}) + 3\beta^2(I_{t-2} + K_{t-2}) \\ &\leq (5-3\beta)K_t + 6\beta\left[\beta(I_{t-2} + K_{t-2}) + \frac{5-3\beta}{2}K_{t-1}\right] + 3\beta^2(I_{t-2} + K_{t-2}) \quad \text{apply (33) for } t-1 \\ &= (5-3\beta)K_t + 3\beta(5-3\beta)K_{t-1} + 9\beta^2(I_{t-2} + K_{t-2}). \end{aligned}$$

We consider the last term, which can be bounded as

$$\begin{aligned} \beta^2(I_{t-2} + K_{t-2}) &= \beta^3(I_{t-3} + K_{t-3}) + \frac{5-3\beta}{2}\beta^2 K_{t-2} \quad \text{apply (33) recursively for } t-2 > 1 \\ &= \beta^{t-1}(I_1 + K_1) + \frac{5-3\beta}{2} \sum_{j=2}^{t-2} \beta^{t-j} K_j \\ &\leq \beta^{t-1} \frac{5-3\beta}{2} K_1 + \frac{5-3\beta}{2} \sum_{j=2}^{t-2} \beta^{t-j} K_j \quad \text{apply (35)} \\ &\leq \frac{5-3\beta}{2} \sum_{j=1}^{t-2} \beta^{t-j} K_j. \end{aligned}$$

Note that this bound is also true for the case $t-2=1$:

$$\beta^2(I_{t-2} + K_{t-2}) = \beta^2(I_1 + K_1) \leq \beta^2 \frac{5-3\beta}{2} K_1 \quad \text{apply (35)}$$

Substituting this inequality into P_t , for $t > 2$, we get

$$P_t \leq (5-3\beta)K_t + 3\beta(5-3\beta)K_{t-1} + 9\beta^2(I_{t-2} + K_{t-2})$$

$$\begin{aligned}
 &\leq \frac{9}{2}(5-3\beta)K_t + \frac{9}{2}(5-3\beta)\beta K_{t-1} + \frac{9}{2}(5-3\beta) \sum_{j=1}^{t-2} \beta^{t-j} K_j \\
 &\leq \frac{9}{2}(5-3\beta) \sum_{j=1}^t \beta^{t-j} K_j.
 \end{aligned}$$

Now we analyze similarly for the case $t = 2$ as follows:

$$\begin{aligned}
 \beta J_2 + (1-\beta)I_2 &\leq 3\beta L^2 \xi_2^2 (I_1 + K_1) + (1-\beta)L^2 \xi_2^2 \left[\beta(I_1 + K_1) + \frac{3}{2}(1-\beta)(I_2 + K_2) \right] \\
 &\leq L^2 \xi_2^2 \left[2(I_2 + K_2) + 4\beta(I_1 + K_1) \right] \\
 &=: L^2 \xi_2^2 P_2,
 \end{aligned}$$

where the last line follows since $1 - \beta \leq 1$ and $P_2 := 2(I_2 + K_2) + 4\beta(I_1 + K_1)$.

Next, we bound the term P_2 as follows:

$$\begin{aligned}
 P_2 &= 2(I_2 + K_2) + 4\beta(I_1 + K_1) \\
 &\leq 2\left(\beta(I_1 + K_1) + \frac{5-3\beta}{2}K_2\right) + 4\beta(I_1 + K_1) && \text{apply (33) for } t = 2 \\
 &= (5-3\beta)K_2 + 6\beta(I_1 + K_1) \\
 &\leq (5-3\beta)K_2 + 6\beta \frac{5-3\beta}{2}K_1 && \text{apply (35)} \\
 &= (5-3\beta)K_2 + 3\beta(5-3\beta)K_1 \\
 &\leq \frac{9}{2}(5-3\beta) \sum_{j=1}^2 \beta^{2-j} K_j.
 \end{aligned}$$

Hence, the statements $\beta J_t + (1-\beta)I_t \leq L^2 \xi_t^2 P_t$ and $P_t \leq \frac{9}{2}(5-3\beta) \sum_{j=1}^t \beta^{t-j} K_j$ are true for every $t \geq 2$.

Combining these two cases, we have the following estimates:

$$\begin{aligned}
 \beta J_t + (1-\beta)I_t &\leq L^2 \xi_t^2 P_t \leq L^2 \xi_t^2 \cdot \frac{9}{2}(5-3\beta) \sum_{j=1}^t \beta^{t-j} K_j \\
 &\stackrel{(15)}{\leq} L^2 \xi_t^2 \cdot \frac{9}{2}(5-3\beta) \sum_{j=1}^t \beta^{t-j} \left[n(\Theta+1) \|\nabla F(w_0^{(j)})\|^2 + n\sigma^2 \right] \\
 &\leq L^2 \xi_t^2 \cdot \frac{9}{2}(5-3\beta) \left[n(\Theta+1) \sum_{j=1}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \sum_{j=1}^t \beta^{t-j} n\sigma^2 \right] \\
 &\leq \frac{9(5-3\beta)L^2 \xi_t^2}{2} \left[n(\Theta+1) \sum_{j=1}^t \beta^{t-j} \|\nabla F(w_0^{(j)})\|^2 + \frac{n\sigma^2}{1-\beta} \right],
 \end{aligned}$$

where the last line follows since $\sum_{j=1}^t \beta^{t-j} \leq \frac{1}{1-\beta}$ for every $t \geq 2$. Hence, we have proved (25). $\square$

7.7.5. THE PROOF OF LEMMA 5

Using the assumption that $\eta_1 \geq \eta_2 \geq \dots \geq \eta_T$, we can derive the following sum as

$$\sum_{t=1}^T \eta_t \sum_{j=1}^t \beta^{t-j} \|\nabla F(\tilde{w}_{j-1})\|^2 = \sum_{t=1}^T \eta_t \left[\beta^{t-1} \|\nabla F(\tilde{w}_0)\|^2 + \beta^{t-2} \|\nabla F(\tilde{w}_1)\|^2 + \dots + \beta^0 \|\nabla F(\tilde{w}_{t-1})\|^2 \right]$$

$$\begin{aligned}
 &= \eta_1 \beta^0 \|\nabla F(\tilde{w}_0)\|^2 \\
 &\quad + \eta_2 \beta^1 \|\nabla F(\tilde{w}_0)\|^2 + \eta_2 \beta^0 \|\nabla F(\tilde{w}_1)\|^2 \\
 &\quad + \eta_t \beta^{t-1} \|\nabla F(\tilde{w}_0)\|^2 + \eta_t \beta^{t-2} \|\nabla F(\tilde{w}_1)\|^2 + \dots + \eta_t \beta^0 \|\nabla F(\tilde{w}_{t-1})\|^2 \\
 &\quad + \eta_T \beta^{T-1} \|\nabla F(\tilde{w}_0)\|^2 + \eta_T \beta^{T-2} \|\nabla F(\tilde{w}_1)\|^2 + \dots + \eta_T \beta^0 \|\nabla F(\tilde{w}_{T-1})\|^2 \\
 &\leq \eta_1 \beta^0 \|\nabla F(\tilde{w}_0)\|^2 \\
 &\quad + \eta_1 \beta^1 \|\nabla F(\tilde{w}_0)\|^2 + \eta_2 \beta^0 \|\nabla F(\tilde{w}_1)\|^2 \\
 &\quad + \eta_1 \beta^{t-1} \|\nabla F(\tilde{w}_0)\|^2 + \eta_2 \beta^{t-2} \|\nabla F(\tilde{w}_1)\|^2 + \dots + \eta_t \beta^0 \|\nabla F(\tilde{w}_{t-1})\|^2 \\
 &\quad + \eta_1 \beta^{T-1} \|\nabla F(\tilde{w}_0)\|^2 + \eta_2 \beta^{T-2} \|\nabla F(\tilde{w}_1)\|^2 + \dots + \eta_T \beta^0 \|\nabla F(\tilde{w}_{T-1})\|^2 \\
 &\leq \eta_1 \sum_{i=0}^{T-1} \beta^i \|\nabla F(\tilde{w}_0)\|^2 + \eta_2 \sum_{i=0}^{T-2} \beta^i \|\nabla F(\tilde{w}_1)\|^2 + \dots + \eta_T \beta^0 \|\nabla F(\tilde{w}_{T-1})\|^2 \\
 &\leq \frac{1}{1-\beta} \sum_{t=1}^T \eta_t \|\nabla F(\tilde{w}_{t-1})\|^2,
 \end{aligned}$$

where the last inequality follows since $\sum_{i=0}^{t-1} \beta^i \leq \frac{1}{1-\beta}$ for $t = 1, 2, \dots, T$. $\square$

7.7.6. THE PROOF OF LEMMA 6

We analyze the special case $t = 1$ as follows. First, letting $i = n$ in Equation (19), for $t \geq 1$, we have

$$w_n^{(t)} - w_0^{(t)} = -\frac{\eta_t}{n} \sum_{j=0}^{n-1} m_{j+1}^{(t)} \stackrel{(2)}{=} -\frac{\eta_t}{n} \sum_{j=0}^{n-1} [\beta m_0^{(t)} + (1-\beta)g_j^{(t)}] = -\frac{\eta_t}{n} \left[n\beta m_0^{(t)} + (1-\beta) \sum_{j=0}^{n-1} g_j^{(t)} \right].$$

For $t = 1$, since $m_0^{(1)} = \mathbf{0}$, we have

$$w_n^{(t)} - w_0^{(t)} = -\frac{\eta_t}{n} (1-\beta) \sum_{j=0}^{n-1} g_j^{(t)}. \quad (37)$$

Using the L -smoothness of F in Assumption 1(b), $w_0^{(t+1)} := w_n^{(t)}$, and (37), we have

$$\begin{aligned}
 F(w_0^{(t+1)}) &\stackrel{(5)}{\leq} F(w_0^{(t)}) + \nabla F(w_0^{(t)})^\top (w_0^{(t+1)} - w_0^{(t)}) + \frac{L}{2} \|w_0^{(t+1)} - w_0^{(t)}\|^2 \\
 &\stackrel{(37)}{=} F(w_0^{(t)}) - \eta_t (1-\beta) \nabla F(w_0^{(t)})^\top \left(\frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t)} \right) + \frac{L\eta_t^2}{2} (1-\beta)^2 \left\| \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t)} \right\|^2 \\
 &\stackrel{(a)}{=} F(w_0^{(t)}) - \frac{\eta_t}{2} (1-\beta) \|\nabla F(w_0^{(t)})\|^2 + \frac{\eta_t}{2} (1-\beta) \left\| \nabla F(w_0^{(t)}) - \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t)} \right\|^2 \\
 &\quad - \frac{\eta_t}{2} (1-\beta) (1 - L\eta_t(1-\beta)) \left\| \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t)} \right\|^2 \\
 &\leq F(w_0^{(t)}) - \frac{\eta_t}{2} (1-\beta) \|\nabla F(w_0^{(t)})\|^2 + \frac{\eta_t}{2} (1-\beta) \left\| \nabla F(w_0^{(t)}) - \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(t)} \right\|^2, \quad (38)
 \end{aligned}$$

where (a) follows from the equality $u^\top v = \frac{1}{2} (\|u\|^2 + \|v\|^2 - \|u-v\|^2)$ and the last equality comes from the fact that $\eta_t \leq \frac{\sqrt{2}}{\sqrt{5}L} \leq \frac{1}{L(1-\beta)}$.

Next, we bound the last term of (38) as follows:

$$\left\| \nabla F(w_0^{(1)}) - \frac{1}{n} \sum_{j=0}^{n-1} g_j^{(1)} \right\|^2 = \left\| \frac{1}{n} \sum_{j=0}^{n-1} \left(\nabla f(w_0^{(1)}; \pi^{(1)}(j+1)) - g_j^{(1)} \right) \right\|^2$$

$$\stackrel{(b)}{\leq} \frac{1}{n} \sum_{j=0}^{n-1} \left\| \nabla f(w_0^{(1)}; \pi^{(1)}(j+1)) - g_j^{(1)} \right\|^2$$

$$\stackrel{(13)}{=} \frac{1}{n} I_1,$$

where (b) is from the Cauchy-Schwarz inequality.

Applying this into (38) we get the following estimate

$$F(w_0^{(2)}) \leq F(w_0^{(1)}) - \frac{\eta_1}{2} (1 - \beta) \|\nabla F(w_0^{(1)})\|^2 + \frac{\eta_1}{2} (1 - \beta) \frac{1}{n} I_1. \quad (39)$$

Remark 7. Note that the derived estimate (39) is true for all shuffling strategies, including Random Reshuffling, Shuffle Once and Incremental Gradient (under the assumptions of Theorem 1 and 2, respectively).

Applying Lemma 3, we further have

$$\begin{aligned} F(w_0^{(2)}) &\leq F(w_0^{(1)}) - \frac{\eta_1}{2} (1 - \beta) \|\nabla F(w_0^{(1)})\|^2 + \frac{\eta_1}{2} (1 - \beta) \frac{1}{n} I_1 \\ &\stackrel{(32)}{\leq} F(w_0^{(1)}) - \frac{\eta_1}{2} (1 - \beta) \|\nabla F(w_0^{(1)})\|^2 + \frac{\eta_1}{2} (1 - \beta) \frac{3}{2n} L^2 \xi_1^2 (1 - \beta)^2 (I_1 + K_1) \\ &\stackrel{(35)}{=} F(w_0^{(1)}) - \frac{\eta_1}{2} (1 - \beta) \|\nabla F(w_0^{(1)})\|^2 + \frac{\eta_1}{2} (1 - \beta) \frac{3}{2n} L^2 \xi_1^2 (1 - \beta)^2 \frac{5 - 3\beta}{2} K_1. \end{aligned}$$

Noting that $\tilde{w}_{t-1} = w_0^{(t)}$, $1 - \beta < 1$ and $K_t := n(\Theta + 1) \|\nabla F(w_0^{(t)})\|^2 + n\sigma^2$, we get

$$F(\tilde{w}_1) \leq F(\tilde{w}_0) - \frac{\eta_1}{2} (1 - \beta) \|\nabla F(\tilde{w}_0)\|^2 + \frac{3\eta_1 L^2 \xi_1^2 (1 - \beta) (5 - 3\beta)}{8} \left[(\Theta + 1) \|\nabla F(\tilde{w}_0)\|^2 + \sigma^2 \right].$$

Rearranging the terms and dividing both sides by $(1 - \beta)$, we have

$$\frac{\eta_1}{2} \|\nabla F(\tilde{w}_0)\|^2 \leq \frac{F(\tilde{w}_0) - F(\tilde{w}_1)}{(1 - \beta)} + \frac{3\eta_1 L^2 \xi_1^2 (5 - 3\beta)}{8} (\Theta + 1) \|\nabla F(\tilde{w}_0)\|^2 + \frac{3\eta_1 \sigma^2 L^2 \xi_1^2 (5 - 3\beta)}{8}.$$

Since ξ_t satisfies $9L^2 \xi_t^2 (5 - 3\beta) (\Theta + 1) \leq 1 - \beta$ as above, we can deduce from the last estimate that

$$\frac{\eta_1}{2} \|\nabla F(\tilde{w}_0)\|^2 \leq \frac{F(\tilde{w}_0) - F(\tilde{w}_1)}{(1 - \beta)} + \frac{(1 - \beta) \eta_1}{4} \|\nabla F(\tilde{w}_0)\|^2 + \frac{9\sigma^2 L^2 (5 - 3\beta)}{4(1 - \beta)} \cdot \xi_1^3,$$

which proves (26). $\square$

8. Convergence Analysis of Algorithm 1 for Randomized Reshuffling Strategy

In this section, we present the convergence results of Algorithm 1 for Randomized Reshuffling strategy. Since this analysis is similar to the previous bounds in Theorem 1, we will highlight the similarities and discuss the differences between the two analyses.

8.1. The Proof Sketch of Theorem 2

The proof of Theorem 2 is similar to Theorem 1, except for the fact that $\mathbb{E}[\beta I_t + (1 - \beta) J_t]$ is upper bounded by a different term. The proof of Theorem 2 is divided into the following steps.

- From Theorem 1 we have

$$F(\tilde{w}_t) \leq F(\tilde{w}_{t-1}) - \frac{\eta_t}{2} \|\nabla F(\tilde{w}_{t-1})\|^2 + \frac{\eta_t}{2n} (\beta I_t + (1 - \beta) J_t),$$

where I_t and J_t are defined by (13) and (14).

- Then, $\mathbb{E}[\beta I_t + (1 - \beta)J_t]$ can be upper bounded as

$$\mathbb{E}[\beta I_t + (1 - \beta)J_t] \leq \mathcal{O}\left(\xi_t^2 \sum_{j=0}^t \beta^{t-j} \mathbb{E}\left[\|\nabla F(w_0^{(j)})\|^2\right] + \xi_t^2 \sigma^2\right),$$

as shown in Lemma 9, where $\xi_t := \max\{\eta_{t-1}, \eta_t\}$ defined in (17).

- We further upper bound the sum of the right-hand side of the last estimate as in Lemma 5 in terms $\sum_{t=0}^T \eta_t \mathbb{E}\left[\|\nabla F(\tilde{w}_{t-1})\|^2\right]$.

Combining these steps together, and using some simplification, we obtain (9) in Theorem 2.

8.2. Technical Lemmas

In this subsection, we will introduce some intermediate results used in Theorem 2. Lemmas 7 and 8 in Theorem 2 play a similar role as previous Lemmas 2 and 3 in Theorem 1.

Lemma 7. Assume that the Randomized Reshuffling strategy is used in Algorithm 1. Then, under the same setting as of Lemma 1 and Assumption 1(c), for $t \geq 1$, it holds that

$$(a) \quad \mathbb{E}[A_n^{(t)}] \leq 2n \cdot \mathbb{E}[I_t + N_t] \quad \text{and} \quad \sum_{i=0}^{n-1} \mathbb{E}[A_i^{(t)}] \leq n^2 \cdot \mathbb{E}[I_t + N_t], \quad (40)$$

$$(b) \quad \sum_{i=0}^{n-1} \mathbb{E}[B_i^{(t)}] \leq 2n^2 \cdot \mathbb{E}[I_t + N_t]. \quad (41)$$

The results of Lemma 8 below are direct consequences of Lemma 1, Lemma 7, and the fact that $L^2 \xi_t^2 \leq \frac{3}{5}$.

Lemma 8. Under the same setting as of Lemma 7 and Assumption 1(b), it holds that

$$\begin{aligned} (a) \quad \mathbb{E}[I_t] &\leq \begin{cases} L^2 \xi_t^2 \left[\frac{2}{3} \beta \cdot \mathbb{E}[I_{t-1} + N_{t-1}] + (1 - \beta) \cdot \mathbb{E}[I_t + N_t] \right], & \text{if } t > 1, \\ L^2 \xi_t^2 (1 - \beta)^2 \cdot \mathbb{E}[I_t + N_t], & \text{if } t = 1, \end{cases} \\ (b) \quad \mathbb{E}[J_t] &\leq \begin{cases} 2L^2 \xi_t^2 \left[\beta \cdot \mathbb{E}[I_{t-2} + N_{t-2}] + (1 - \beta) \cdot \mathbb{E}[I_{t-1} + N_{t-1}] \right], & \text{if } t > 2, \\ 2L^2 \xi_t^2 \cdot \mathbb{E}[I_{t-1} + N_{t-1}], & \text{if } t = 2, \end{cases} \\ (c) \quad \text{If } L^2 \xi_t^2 \leq \frac{3}{5}, \text{ then } \mathbb{E}[I_t + N_t] &\leq \begin{cases} \beta \cdot \mathbb{E}[I_{t-1} + N_{t-1}] + \frac{5-3\beta}{2} \cdot \mathbb{E}[N_t], & \text{if } t > 1, \\ \frac{5-3\beta}{2} \cdot \mathbb{E}[N_1], & \text{if } t = 1. \end{cases} \end{aligned}$$

We will present below Lemmas 9 and 10, which play a similar role as previous Lemmas 4 and 6 in Theorem 1.

Lemma 9. Under the same conditions as in Lemma 8, for any $t \geq 2$, we have

$$\mathbb{E}[\beta J_t + (1 - \beta)I_t] \leq 3(5 - 3\beta)L^2 \xi_t^2 \left((\Theta + n) \sum_{j=1}^t \beta^{t-j} \mathbb{E}\left[\|\nabla F(w_0^{(j)})\|^2\right] + \frac{\sigma^2}{1 - \beta} \right). \quad (42)$$

Lemma 10. Under the same setting as in Theorem 2, we have

$$\frac{\eta_1}{2} \mathbb{E}\left[\|\nabla F(\tilde{w}_0)\|^2\right] \leq \frac{\mathbb{E}[F(\tilde{w}_0) - F(\tilde{w}_1)]}{(1 - \beta)} + \frac{(1 - \beta)\eta_1}{4} \mathbb{E}\left[\|\nabla F(\tilde{w}_0)\|^2\right] + \frac{3\sigma^2(5 - 3\beta)L^2}{2n(1 - \beta)} \cdot \xi_1^3. \quad (43)$$

8.3. The Proof of Theorem 2: Key Estimate for Algorithm 1

First, from the assumption $0 < \eta_t \leq \frac{1}{L\sqrt{D}}$, $t \geq 1$, we have $0 < \eta_t^2 \leq \frac{1}{DL^2}$. Next, from (17), we have $\xi_t = \max(\eta_t; \eta_{t-1})$ for $t > 1$ and $\xi_1 = \eta_1$, which lead to $0 < \xi_t^2 \leq \frac{1}{DL^2}$ for $t \geq 1$. Moreover, from the definition of $D = \max\left(\frac{5}{3}, \frac{6(5-3\beta)(\Theta+n)}{n(1-\beta)}\right)$ in Theorem 2, we have $L^2 \xi_t^2 \leq \frac{3}{5}$ and $6L^2 \xi_t^2(5 - 3\beta)(\Theta + n) \leq n(1 - \beta)$ for $t \geq 1$.

Similar to the estimate (29) in Theorem 1, we get the following result:

$$F(w_0^{(t+1)}) \leq F(w_0^{(t)}) - \frac{\eta_t}{2} \|\nabla F(w_0^{(t)})\|^2 + \frac{\eta_t}{2n} [\beta J_t + (1 - \beta) I_t].$$

Taking expectation of both sides and applying Lemma 9, we get:

$$\begin{aligned} \mathbb{E} [F(w_0^{(t+1)})] &\leq \mathbb{E} [F(w_0^{(t)})] - \frac{\eta_t}{2} \mathbb{E} [\|\nabla F(w_0^{(t)})\|^2] + \frac{\eta_t}{2n} \mathbb{E} [\beta J_t + (1 - \beta) I_t] \\ &\stackrel{(42)}{\leq} \mathbb{E} [F(w_0^{(t)})] - \frac{\eta_t}{2} \mathbb{E} [\|\nabla F(w_0^{(t)})\|^2] \\ &\quad + \frac{\eta_t}{2n} 3(5 - 3\beta) L^2 \xi_t^2 \left[(\Theta + n) \sum_{j=1}^t \beta^{t-j} \mathbb{E} [\|\nabla F(w_0^{(j)})\|^2] + \frac{\sigma^2}{1 - \beta} \right] \\ &\leq \mathbb{E} [F(w_0^{(t)})] - \frac{\eta_t}{2} \mathbb{E} [\|\nabla F(w_0^{(t)})\|^2] \\ &\quad + \frac{3(5 - 3\beta) \eta_t L^2 \xi_t^2 (\Theta + n)}{2n} \sum_{j=1}^t \beta^{t-j} \mathbb{E} [\|\nabla F(w_0^{(j)})\|^2] + \frac{3\sigma^2(5 - 3\beta) \eta_t L^2 \xi_t^2}{2n(1 - \beta)}. \end{aligned}$$

Since ξ_t satisfies $6L^2 \xi_t^2 (5 - 3\beta)(\Theta + n) \leq n(1 - \beta)$ as proved above, we can deduce from the last estimate that

$$\begin{aligned} \mathbb{E} [F(w_0^{(t+1)})] &\leq \mathbb{E} [F(w_0^{(t)})] - \frac{\eta_t}{2} \mathbb{E} [\|\nabla F(w_0^{(t)})\|^2] + \frac{(1 - \beta) \eta_t}{4} \sum_{j=1}^t \beta^{t-j} \mathbb{E} [\|\nabla F(w_0^{(j)})\|^2] \\ &\quad + \frac{3\sigma^2(5 - 3\beta) \eta_t L^2 \xi_t^2}{2n(1 - \beta)}. \end{aligned}$$

Rearranging this inequality while noting that $\eta_t \leq \xi_t$ and $\tilde{w}_{t-1} = w_0^{(t)}$ we obtain that for $t \geq 2$:

$$\begin{aligned} \frac{\eta_t}{2} \mathbb{E} [\|\nabla F(\tilde{w}_{t-1})\|^2] &\leq \mathbb{E} [F(\tilde{w}_{t-1}) - F(\tilde{w}_t)] + \frac{(1 - \beta) \eta_t}{4} \sum_{j=1}^t \beta^{t-j} \mathbb{E} [\|\nabla F(\tilde{w}_{j-1})\|^2] + \frac{3\sigma^2(5 - 3\beta) L^2}{2n(1 - \beta)} \cdot \xi_t^3 \\ &\stackrel{1 - \beta < 1}{\leq} \frac{\mathbb{E} [F(\tilde{w}_{t-1}) - F(\tilde{w}_t)]}{(1 - \beta)} + \frac{(1 - \beta) \eta_t}{4} \sum_{j=1}^t \beta^{t-j} \mathbb{E} [\|\nabla F(\tilde{w}_{j-1})\|^2] + \frac{3\sigma^2(5 - 3\beta) L^2}{2n(1 - \beta)} \cdot \xi_t^3. \end{aligned}$$

For $t = 1$, since $\xi_1 = \eta_1$ as previously defined in (17), from the result of Lemma 10 we have

$$\frac{\eta_1}{2} \mathbb{E} [\|\nabla F(\tilde{w}_0)\|^2] \leq \frac{\mathbb{E} [F(\tilde{w}_0) - F(\tilde{w}_1)]}{(1 - \beta)} + \frac{(1 - \beta) \eta_1}{4} \sum_{j=1}^1 \beta^{1-j} \mathbb{E} [\|\nabla F(\tilde{w}_{j-1})\|^2] + \frac{3\sigma^2(5 - 3\beta) L^2}{2n(1 - \beta)} \cdot \xi_1^3.$$

Summing the previous estimate for $t := 2$ to $t := T$, and then using the last one, we obtain

$$\sum_{t=1}^T \frac{\eta_t}{2} \mathbb{E} [\|\nabla F(\tilde{w}_{t-1})\|^2] \leq \frac{F(\tilde{w}_0) - F_*}{(1 - \beta)} + \frac{(1 - \beta)}{4} \sum_{t=1}^T \eta_t \sum_{j=1}^t \beta^{t-j} \mathbb{E} [\|\nabla F(\tilde{w}_{j-1})\|^2] + \frac{3\sigma^2(5 - 3\beta) L^2}{2n(1 - \beta)} \sum_{t=1}^T \xi_t^3.$$

Applying the result of Lemma 5 to the last estimate, we get

$$\frac{1}{2} \sum_{t=1}^T \eta_t \mathbb{E} [\|\nabla F(\tilde{w}_{t-1})\|^2] \leq \frac{F(\tilde{w}_0) - F_*}{(1 - \beta)} + \frac{(1 - \beta)}{4} \frac{1}{(1 - \beta)} \sum_{t=1}^T \eta_t \mathbb{E} [\|\nabla F(\tilde{w}_{t-1})\|^2] + \frac{3\sigma^2(5 - 3\beta) L^2}{2n(1 - \beta)} \sum_{t=1}^T \xi_t^3,$$

which is equivalent to

$$\sum_{t=1}^T \eta_t \mathbb{E} [\|\nabla F(\tilde{w}_{t-1})\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1 - \beta)} + \frac{6\sigma^2(5 - 3\beta) L^2}{n(1 - \beta)} \sum_{t=1}^T \xi_t^3.$$

Dividing both sides of the resulting estimate by $\sum_{t=1}^T \eta_t$, we obtain (9). Finally, due to the choice of $\hat{w}_T$ at Step 1 in Algorithm 1, we have $\mathbb{E} [\|\nabla F(\hat{w}_T)\|^2] = \frac{1}{\sum_{t=1}^T \eta_t} \sum_{t=1}^T \eta_t \mathbb{E} [\|\nabla F(\tilde{w}_{t-1})\|^2]$. $\square$

8.4. The Proof of Corollaries 5 and 6: Constant and Diminishing Learning Rates

The proof of Corollary 5. Since $T \geq 1$ and $\eta_t := \frac{\gamma n^{1/3}}{T^{1/3}}$, we have $\eta_t \geq \eta_{t+1}$. We also have $0 < \eta_t \leq \frac{1}{L\sqrt{D}}$ for all $t \geq 1$, and $\sum_{t=1}^T \eta_t = \gamma n^{1/3} T^{2/3}$ and $\sum_{t=1}^T \eta_{t-1}^3 = n\gamma^3$. Substituting these expressions into (9), we obtain

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma n^{1/3} T^{2/3}} + \frac{6\sigma^2(5-3\beta)L^2}{(1-\beta)} \cdot \frac{\gamma^2}{n^{1/3} T^{2/3}},$$

which is our desired result. $\square$

The proof of Corollary 6. For $t \geq 1$, since $\eta_t = \frac{\gamma n^{1/3}}{(t+\lambda)^{1/3}}$, we have $\eta_t \geq \eta_{t+1}$. We also have $0 < \eta_t \leq \frac{1}{L\sqrt{D}}$ for all $t \geq 1$, and

$$\begin{aligned} \sum_{t=1}^T \eta_t &= \gamma n^{1/3} \sum_{t=1}^T \frac{1}{(t+\lambda)^{1/3}} \geq \gamma n^{1/3} \int_1^T \frac{d\tau}{(\tau+\lambda)^{1/3}} \geq \gamma n^{1/3} [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}], \\ \sum_{t=1}^T \eta_{t-1}^3 &= 2\eta_1^3 + \gamma^3 \sum_{t=3}^T \frac{1}{t-1+\lambda} \leq \frac{2n\gamma^3}{(1+\lambda)} + n\gamma^3 \int_{t=2}^T \frac{d\tau}{\tau-1+\lambda} \leq n\gamma^3 \left[\frac{2}{(1+\lambda)} + \log(T-1+\lambda) \right]. \end{aligned}$$

Substituting these expressions into (9), we obtain

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma n^{1/3} [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]} + \frac{6\sigma^2(5-3\beta)L^2}{(1-\beta)} \cdot \frac{\gamma^2 \left[\frac{2}{(1+\lambda)} + \log(T-1+\lambda) \right]}{n^{1/3} [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]},$$

Let C_3 and C_4 be defined respectively as

$$C_3 := \frac{4[F(\tilde{w}_0) - F_*]}{(1-\beta)\gamma} + \frac{12\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)(1+\lambda)}, \quad \text{and} \quad C_4 := \frac{6\sigma^2(5-3\beta)L^2\gamma^2}{(1-\beta)}.$$

Then, the last estimate leads to

$$\mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] \leq \frac{C_3 + C_4 \log(T-1+\lambda)}{n^{1/3} [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]},$$

which completes our proof. $\square$

8.5. The Proof of Technical Lemmas

We now provide the proof of additional Lemmas that serve for the proof of Theorem 2 above.

8.5.1. PROOF OF LEMMA 7: UPPER BOUNDING THE TERMS $\mathbb{E}[A_n^{(t)}]$, $\sum_{i=0}^{n-1} \mathbb{E}[A_i^{(t)}]$, AND $\sum_{i=0}^{n-1} \mathbb{E}[B_i^{(t)}]$

In this proof, we will use (Mishchenko et al., 2020)[Lemma 1] for sampling without replacement. For the sake of references, we recall it here.

Lemma 11 (Lemma 1 in (Mishchenko et al., 2020)). *Let $X_1, \dots, X_n \in \mathbb{R}^d$ be fixed vectors, $\bar{X} := \frac{1}{n} \sum_{i=1}^n X_i$ be their average and $\sigma^2 := \frac{1}{n} \sum_{i=1}^n \|X_i - \bar{X}\|^2$ be the population variance. Fix any $k \in \{1, \dots, n\}$, let $X_{\pi_1}, \dots, X_{\pi_k}$ be sampled uniformly without replacement from $\{X_1, \dots, X_n\}$ and $\bar{X}_\pi$ be their average. Then, the sample average and the variance are given, respectively by*

$$\mathbb{E}[\bar{X}_\pi] = \bar{X} \quad \text{and} \quad \mathbb{E}[\|\bar{X}_\pi - \bar{X}\|^2] = \frac{n-k}{k(n-1)} \sigma^2.$$

Using this result, we now prove Lemma 7 as follows.

(a) Let first upper bound the term $A_i^{(t)}$ defined by (11). Indeed, for $i = 0, \dots, n-1$ and $t \geq 1$, we have

$$A_i^{(t)} = \left\| \sum_{j=0}^{i-1} g_j^{(t)} \right\|^2 \stackrel{(a)}{\leq} 2 \left\| \sum_{j=0}^{i-1} \left(g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right) \right\|^2 + 2 \left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2$$

$$\begin{aligned}
 &\stackrel{(a)}{\leq} 2i \sum_{j=0}^{i-1} \left\| g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 + 2 \left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \\
 &\leq 2i \sum_{j=0}^{n-1} \left\| g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 + 2 \left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \\
 &\stackrel{(13)}{=} 2iI_t + 2 \left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2,
 \end{aligned}$$

where we have used the Cauchy-Schwarz inequality in (a).

Now taking expectation conditioned on $\mathcal{F}_t$, we get

$$\mathbb{E}_t[A_i^{(t)}] \leq 2i \cdot \mathbb{E}_t[I_t] + 2\mathbb{E}_t \left[\left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \right].$$

Note that $\mathbb{E}_t[\nabla f(w_0^{(t)}; \pi^{(t)}(j+1))] = \nabla F(w_0^{(t)})$ for every index $j = 0, 1, \dots, i-1$. Hence, we have

$$\begin{aligned}
 \mathbb{E}_t[A_i^{(t)}] &\leq 2i \cdot \mathbb{E}_t[I_t] + 2\mathbb{E}_t \left[\left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \right] \\
 &= 2i \cdot \mathbb{E}_t[I_t] + 2 \left\| i \nabla F(w_0^{(t)}) \right\|^2 + 2\mathbb{E}_t \left[\left\| \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - i \nabla F(w_0^{(t)}) \right\|^2 \right] \\
 &= 2i \cdot \mathbb{E}_t[I_t] + 2i^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 + 2i^2 \mathbb{E}_t \left[\left\| \frac{1}{i} \sum_{j=0}^{i-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 \right] \\
 &= 2i \cdot \mathbb{E}_t[I_t] + 2i^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 + 2i^2 \frac{n-i}{i(n-1)} \frac{1}{n} \sum_{j=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 \\
 &\stackrel{(b)}{\leq} 2i \cdot \mathbb{E}_t[I_t] + 2i^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 + \frac{2i(n-i)}{(n-1)} \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right],
 \end{aligned}$$

where we apply the sample variance Lemma from (Mishchenko et al., 2020) and (b) follows from Assumption 1(c).

Letting $i = n$ in the above estimate and taking total expectation, we obtain the first estimate of Lemma 7, i.e.:

$$\mathbb{E}[A_n^{(t)}] \leq 2n \cdot \mathbb{E}[I_t] + 2n^2 \mathbb{E} \left[\left\| \nabla F(w_0^{(t)}) \right\|^2 \right] \leq 2n \cdot \mathbb{E}[I_t + N_t].$$

Now we are ready to calculate the sum $\sum_{i=0}^{n-1} \mathbb{E}_t[A_i^{(t)}]$ as

$$\begin{aligned}
 \sum_{i=0}^{n-1} \mathbb{E}_t[A_i^{(t)}] &\leq 2 \sum_{i=0}^{n-1} i \cdot \mathbb{E}_t[I_t] + 2 \sum_{i=0}^{n-1} i^2 \cdot \left\| \nabla F(w_0^{(t)}) \right\|^2 + 2 \sum_{i=0}^{n-1} \frac{i(n-i)}{(n-1)} \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \\
 &\leq n^2 \cdot \mathbb{E}_t[I_t] + \frac{2n^3}{3} \left\| \nabla F(w_0^{(t)}) \right\|^2 + \frac{n(n+1)}{3} \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \\
 &\leq n^2 \cdot \mathbb{E}_t[I_t] + n^3 \left\| \nabla F(w_0^{(t)}) \right\|^2 + n^2 \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \\
 &\leq n^2 \cdot \mathbb{E}_t[I_t] + n^2 \left[(n + \Theta) \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \stackrel{(16)}{\leq} n^2 (\mathbb{E}_t[I_t] + N_t), \quad t \geq 1,
 \end{aligned}$$

where we use the facts that $\sum_{i=0}^{n-1} i = \frac{n(n-1)}{2} \leq \frac{n^2}{2}$ and $\sum_{i=0}^{n-1} i^2 = \frac{n(n-1)(2n-1)}{6} \leq \frac{n^3}{3}$.

Taking total expectation, we have the second estimate of Lemma 7.

(b) Using similar argument as above, for $i = 0, 1, \dots, n-1$ and $t \geq 1$, we can derive

$$\begin{aligned}
 B_i^{(t)} &= \left\| \sum_{j=i}^{n-1} g_j^{(t)} \right\|^2 \stackrel{(a)}{\leq} 2 \left\| \sum_{j=i}^{n-1} (g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1))) \right\|^2 + 2 \left\| \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \\
 &\stackrel{(a)}{\leq} 2(n-i) \sum_{j=i}^{n-1} \left\| g_j^{(t)} - \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 + 2 \left\| \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \\
 &\stackrel{(13)}{\leq} 2(n-i)I_t + 2 \left\| \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2,
 \end{aligned}$$

where we use again the Cauchy-Schwarz inequality in (a).

Taking expectation conditioned on $\mathcal{F}_t$, we get

$$\mathbb{E}_t[B_i^{(t)}] \leq 2(n-i) \cdot \mathbb{E}_t[I_t] + 2\mathbb{E}_t \left[\left\| \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \right].$$

Note that $\mathbb{E}_t[\nabla f(w_0^{(t)}; \pi^{(t)}(j+1))] = \nabla F(w_0^{(t)})$ for every index $j = i, \dots, n-1$. Hence

$$\begin{aligned}
 \mathbb{E}_t[B_i^{(t)}] &\leq 2(n-i) \cdot \mathbb{E}_t[I_t] + 2\mathbb{E}_t \left[\left\| \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) \right\|^2 \right] \\
 &= 2(n-i) \cdot \mathbb{E}_t[I_t] + 2 \left\| (n-i) \nabla F(w_0^{(t)}) \right\|^2 + 2\mathbb{E}_t \left[\left\| \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - (n-i) \nabla F(w_0^{(t)}) \right\|^2 \right] \\
 &= 2(n-i) \cdot \mathbb{E}_t[I_t] + 2(n-i)^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 \\
 &\quad + 2(n-i)^2 \mathbb{E}_t \left[\left\| \frac{1}{n-i} \sum_{j=i}^{n-1} \nabla f(w_0^{(t)}; \pi^{(t)}(j+1)) - \nabla F(w_0^{(t)}) \right\|^2 \right] \\
 &\stackrel{(b)}{\leq} 2(n-i) \cdot \mathbb{E}_t[I_t] + 2(n-i)^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 + \frac{2i(n-i)}{(n-1)} \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right],
 \end{aligned}$$

where we apply again the sample variance Lemma from (Mishchenko et al., 2020) and (b) follows from Assumption 1(c).

Finally, for all $t \geq 1$, we calculate the sum $\sum_{i=0}^{n-1} \mathbb{E}_t[B_i^{(t)}]$ similarly as follows

$$\begin{aligned}
 \sum_{i=0}^{n-1} \mathbb{E}_t[B_i^{(t)}] &\leq 2 \sum_{i=0}^{n-1} (n-i) \cdot \mathbb{E}_t[I_t] + 2 \sum_{i=0}^{n-1} (n-i)^2 \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sum_{i=0}^{n-1} \frac{2i(n-i)}{(n-1)} \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \\
 &\leq 2n^2 \cdot \mathbb{E}_t[I_t] + \frac{n(n+1)(2n+1)}{3} \left\| \nabla F(w_0^{(t)}) \right\|^2 + \frac{n(n+1)}{3} \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \\
 &\leq 2n^2 \cdot \mathbb{E}_t[I_t] + 2n^3 \left\| \nabla F(w_0^{(t)}) \right\|^2 + 2n^2 \left[\Theta \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \\
 &\leq 2n^2 \cdot \mathbb{E}_t[I_t] + 2n^2 \left[(n+\Theta) \left\| \nabla F(w_0^{(t)}) \right\|^2 + \sigma^2 \right] \stackrel{(16)}{\leq} 2n^2 (\mathbb{E}_t[I_t] + N_t).
 \end{aligned}$$

Taking total expectation, we get (41). □

8.5.2. PROOF OF LEMMA 8: UPPER BOUNDING THE TERMS $\mathbb{E}[I_t]$, $\mathbb{E}[J_t]$, AND $\mathbb{E}[I_t + N_t]$

(a) First, for every $t \geq 1$ and $g_i^{(t)} := \nabla f(w_i^{(t)}; \sigma^{(t)}(i+1))$, we have

$$I_t = \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(i+1)) - g_i^{(t)} \right\|^2$$

$$\begin{aligned}
 &\stackrel{\text{Step 1}}{=} \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t)}(i+1)) - \nabla f(w_i^{(t)}; \pi^{(t)}(i+1)) \right\|^2 \\
 &\leq L^2 \sum_{i=0}^{n-1} \left\| w_i^{(t)} - w_0^{(t)} \right\|^2,
 \end{aligned} \tag{44}$$

where the last inequality follows from Assumption 1(b).

Applying Lemma 1 to (44), for $t > 1$, we have

$$\begin{aligned}
 I_t &\leq \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} \left[\beta \frac{i^2}{n^2} A_n^{(t-1)} + (1-\beta) A_i^{(t)} \right] \\
 &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta \frac{A_n^{(t-1)}}{n^2} \sum_{i=0}^{n-1} i^2 + (1-\beta) \sum_{i=0}^{n-1} A_i^{(t)} \right] \\
 &\stackrel{(a)}{\leq} \frac{L^2 \xi_t^2}{n^2} \left[\beta \cdot \frac{n}{3} \cdot A_n^{(t-1)} + (1-\beta) \sum_{i=0}^{n-1} A_i^{(t)} \right]
 \end{aligned}$$

where (a) follows from the fact that $\sum_{i=0}^{n-1} i^2 \leq \frac{n^3}{3}$. Taking expectation and applying result (40) in Lemma 7 we get

$$\begin{aligned}
 \mathbb{E}[I_t] &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta \cdot \frac{n}{3} \cdot \mathbb{E}[A_n^{(t-1)}] + (1-\beta) \sum_{i=0}^{n-1} \mathbb{E}[A_i^{(t)}] \right] \\
 &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta \cdot \frac{n}{3} \cdot 2n \cdot \mathbb{E}[I_{t-1} + N_{t-1}] + (1-\beta) n^2 \mathbb{E}[I_t + N_t] \right] \\
 &\leq L^2 \xi_t^2 \left[\frac{2}{3} \beta \cdot \mathbb{E}[I_{t-1} + N_{t-1}] + (1-\beta) \mathbb{E}[I_t + N_t] \right].
 \end{aligned}$$

For $t = 1$, by applying Lemma 1 and 7 consecutively to (44) we have

$$I_t \leq L^2 \sum_{i=0}^{n-1} \left\| w_i^{(t)} - w_0^{(t)} \right\|^2 \stackrel{(21)}{\leq} \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} \left((1-\beta)^2 A_i^{(t)} \right) \leq \frac{L^2 \xi_t^2 (1-\beta)^2}{n^2} \sum_{i=0}^{n-1} A_i^{(t)}.$$

Similarly we get

$$\mathbb{E}[I_t] \stackrel{(40)}{\leq} \frac{L^2 \xi_t^2 (1-\beta)^2}{n^2} n^2 \cdot \mathbb{E}[I_t + N_t] = L^2 \xi_t^2 (1-\beta)^2 \mathbb{E}[I_t + N_t]. \tag{45}$$

(b) Using the similar argument as above we can derive that

$$\begin{aligned}
 J_t &= \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t-1)}(i+1)) - g_i^{(t-1)} \right\|^2 \\
 &\stackrel{\text{Step 1}}{=} \sum_{i=0}^{n-1} \left\| \nabla f(w_0^{(t)}; \pi^{(t-1)}(i+1)) - \nabla f(w_i^{(t-1)}; \pi^{(t-1)}(i+1)) \right\|^2 \\
 &\stackrel{(3)}{\leq} L^2 \sum_{i=0}^{n-1} \left\| w_i^{(t-1)} - w_0^{(t)} \right\|^2.
 \end{aligned}$$

Applying Lemma 1 for $t > 2$ to the above estimate, we get

$$J_t \leq \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} \left[\beta \frac{(n-i)^2}{n^2} A_n^{(t-2)} + (1-\beta) B_i^{(t-1)} \right]$$

$$\begin{aligned}
 &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta \frac{A_n^{(t-2)}}{n^2} \sum_{i=0}^{n-1} (n-i)^2 + (1-\beta) \sum_{i=0}^{n-1} B_i^{(t-1)} \right] \\
 &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta n \cdot A_n^{(t-2)} + (1-\beta) \sum_{i=0}^{n-1} B_i^{(t-1)} \right],
 \end{aligned}$$

where the last line follows from the fact that $\sum_{i=0}^{n-1} (n-i)^2 \leq n^3$. Similar to the previous part, taking total expectation we get

$$\begin{aligned}
 \mathbb{E}[J_t] &\leq \frac{L^2 \xi_t^2}{n^2} \left[\beta n \cdot \mathbb{E}[A_n^{(t-2)}] + (1-\beta) \sum_{i=0}^{n-1} \mathbb{E}[B_i^{(t-1)}] \right] \\
 &\stackrel{(41)}{\leq} \frac{L^2 \xi_t^2}{n^2} \left[\beta n \cdot \mathbb{E}[A_n^{(t-2)}] + (1-\beta) 2n^2 \cdot \mathbb{E}[I_{t-1} + N_{t-1}] \right] \\
 &\stackrel{(40)}{\leq} \frac{L^2 \xi_t^2}{n^2} \left[2\beta n^2 \cdot \mathbb{E}[I_{t-2} + N_{t-2}] + 2(1-\beta)n^2 \cdot \mathbb{E}[I_{t-1} + N_{t-1}] \right] \\
 &\leq 2L^2 \xi_t^2 \left[\beta \cdot \mathbb{E}[I_{t-2} + N_{t-2}] + (1-\beta) \cdot \mathbb{E}[I_{t-1} + N_{t-1}] \right].
 \end{aligned}$$

Now applying Lemma 1 for the special case $t = 2$, we obtain

$$J_t \leq \frac{L^2 \xi_t^2}{n^2} \sum_{i=0}^{n-1} (1-\beta)^2 B_i^{(t-1)}.$$

Therefore, we have

$$\mathbb{E}[J_t] \leq \frac{L^2 \xi_t^2}{n^2} (1-\beta)^2 \sum_{i=0}^{n-1} \mathbb{E}[B_i^{(t-1)}] \leq \frac{L^2 \xi_t^2}{n^2} (1-\beta)^2 \cdot 2n^2 \mathbb{E}[I_{t-1} + N_{t-1}] \leq 2L^2 \xi_t^2 \cdot \mathbb{E}[I_{t-1} + N_{t-1}],$$

where we have used the result (41) from Lemma 7 and the fact that $1-\beta \leq 1$.

(c) First, from the result of Part (a), for $t > 1$, we have

$$\mathbb{E}[I_t] \leq L^2 \xi_t^2 \left[\frac{2}{3} \beta \mathbb{E}[I_{t-1} + N_{t-1}] + (1-\beta) \mathbb{E}[N_t] \right] + (1-\beta) L^2 \xi_t^2 \mathbb{E}[I_t].$$

Since $L^2 \xi_t^2 \leq \frac{3}{5}$ (due to the choice of our learning rate) and $1-\beta \leq 1$, we further get

$$\mathbb{E}[I_t] \leq \frac{2}{5} \left[\beta \mathbb{E}[I_{t-1} + N_{t-1}] + \frac{3}{2} (1-\beta) \mathbb{E}[N_t] \right] + \frac{3}{5} \mathbb{E}[I_t],$$

which is equivalent to

$$\mathbb{E}[I_t] \leq \beta \mathbb{E}[I_{t-1} + N_{t-1}] + \frac{3}{2} (1-\beta) \mathbb{E}[N_t].$$

Adding $\mathbb{E}[N_t]$ to both sides of this inequality, for $t > 1$, we get

$$\mathbb{E}[I_t + N_t] \leq \beta \mathbb{E}[I_{t-1} + N_{t-1}] + \frac{5-3\beta}{2} \mathbb{E}[N_t]. \quad (46)$$

Finally, for $t = 1$, since $L^2 \xi_t^2 \leq \frac{3}{5}$ and $1-\beta \leq 1$, we have

$$\mathbb{E}[I_t] \leq L^2 \xi_t^2 (1-\beta)^2 (\mathbb{E}[I_t] + \mathbb{E}[N_t]) \leq \frac{3}{5} (1-\beta)^2 \mathbb{E}[I_t] + \frac{3}{5} (1-\beta)^2 \mathbb{E}[N_t] \leq \frac{3}{5} \mathbb{E}[I_t] + \frac{3}{5} (1-\beta) \mathbb{E}[N_t],$$

which is equivalent to

$$\frac{2}{5} \mathbb{E}[I_t] \leq \frac{3}{5} (1-\beta) \mathbb{E}[N_t]. \quad (47)$$

This leads to our desired result for $t = 1$ as

$$\mathbb{E}[I_1 + N_1] \leq \frac{3}{2} (1-\beta) \mathbb{E}[N_1] + \mathbb{E}[N_1] = \frac{5-3\beta}{2} \mathbb{E}[N_1]. \quad (48)$$

Combining two cases, we obtain the desired result in Part (c). $\square$

8.5.3. PROOF OF LEMMA 9: UPPER BOUNDING THE KEY QUANTITY $\mathbb{E}[\beta J_t + (1 - \beta)I_t]$

First, we analyze the case $t > 2$ using the results of Lemma 8 as follows:

$$\begin{aligned}
 \beta \mathbb{E}[J_t] + (1 - \beta) \mathbb{E}[I_t] &\leq 2\beta L^2 \xi_t^2 [\beta \mathbb{E}[I_{t-2} + N_{t-2}] + (1 - \beta) \mathbb{E}[I_{t-1} + N_{t-1}]] \\
 &\quad + (1 - \beta) L^2 \xi_t^2 \left[\frac{2}{3} \beta \mathbb{E}[I_{t-1} + N_{t-1}] + (1 - \beta) \mathbb{E}[I_t + N_t] \right] \\
 &\leq \frac{2}{3} L^2 \xi_t^2 [2\mathbb{E}[I_t + N_t] + 4\beta \mathbb{E}[I_{t-1} + N_{t-1}] + 3\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}]] \\
 &=: \frac{2}{3} L^2 \xi_t^2 S_t,
 \end{aligned} \tag{49}$$

where the last line follows since $1 - \beta \leq 1$ and $S_t := 2\mathbb{E}[I_t + N_t] + 4\beta \mathbb{E}[I_{t-1} + N_{t-1}] + 3\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}]$ for $t > 2$.

Next, we bound the term S_t in (49) as

$$\begin{aligned}
 S_t &= 2\mathbb{E}[I_t + N_t] + 4\beta \mathbb{E}[I_{t-1} + N_{t-1}] + 3\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] \\
 &\leq 2 \left[\beta \mathbb{E}[I_{t-1} + N_{t-1}] + \frac{5 - 3\beta}{2} \mathbb{E}[N_t] \right] + 4\beta \mathbb{E}[I_{t-1} + N_{t-1}] + 3\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] \quad \text{apply (46) for } t \\
 &= (5 - 3\beta) \mathbb{E}[N_t] + 6\beta \mathbb{E}[I_{t-1} + N_{t-1}] + 3\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] \\
 &\leq (5 - 3\beta) \mathbb{E}[N_t] + 6\beta \left[\beta \mathbb{E}[I_{t-2} + N_{t-2}] + \frac{5 - 3\beta}{2} \mathbb{E}[N_{t-1}] \right] + 3\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] \quad \text{apply (46) for } t - 1 \\
 &= (5 - 3\beta) \mathbb{E}[N_t] + 3\beta(5 - 3\beta) \mathbb{E}[N_{t-1}] + 9\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}].
 \end{aligned}$$

We consider the last term, which can be bounded as

$$\begin{aligned}
 \beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] &= \beta^3 \mathbb{E}[I_{t-3} + N_{t-3}] + \frac{5 - 3\beta}{2} \beta^2 \mathbb{E}[N_{t-2}] \quad \text{apply (46) recursively for } t - 2 > 1 \\
 &= \beta^{t-1} \mathbb{E}[I_1 + N_1] + \frac{5 - 3\beta}{2} \sum_{j=2}^{t-2} \beta^{t-j} \mathbb{E}[N_j] \\
 &\leq \beta^{t-1} \frac{5 - 3\beta}{2} \mathbb{E}[N_1] + \frac{5 - 3\beta}{2} \sum_{j=2}^{t-2} \beta^{t-j} \mathbb{E}[N_j] \quad \text{apply (48)} \\
 &\leq \frac{5 - 3\beta}{2} \sum_{j=1}^{t-2} \beta^{t-j} \mathbb{E}[N_j].
 \end{aligned}$$

Note that this bound is also true for the case $t - 2 = 1$:

$$\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] = \beta^2 \mathbb{E}[I_1 + N_1] \leq \beta^2 \frac{5 - 3\beta}{2} \mathbb{E}[N_j] \quad \text{apply (48)}$$

Substituting this inequality into S_t , for $t > 2$, we get

$$\begin{aligned}
 S_t &\leq (5 - 3\beta) \mathbb{E}[N_t] + 3\beta(5 - 3\beta) \mathbb{E}[N_{t-1}] + 9\beta^2 \mathbb{E}[I_{t-2} + N_{t-2}] \\
 &\leq \frac{9}{2} (5 - 3\beta) \mathbb{E}[N_t] + \frac{9}{2} (5 - 3\beta) \beta \mathbb{E}[N_{t-1}] + \frac{9}{2} (5 - 3\beta) \sum_{j=1}^{t-2} \beta^{t-j} \mathbb{E}[N_j] \\
 &\leq \frac{9}{2} (5 - 3\beta) \sum_{j=1}^t \beta^{t-j} \mathbb{E}[N_j].
 \end{aligned}$$

Now we analyze similarly for the case $t = 2$ as follows:

$$\beta \mathbb{E}[J_2] + (1 - \beta) \mathbb{E}[I_2] \leq 2\beta L^2 \xi_2^2 \mathbb{E}[I_1 + N_1] + (1 - \beta) L^2 \xi_2^2 \left[\frac{2}{3} \beta \mathbb{E}[I_1 + N_1] + (1 - \beta) \mathbb{E}[I_2 + N_2] \right]$$

$$\begin{aligned}
 &\leq \frac{2}{3}L^2\xi_2^2[2\mathbb{E}[I_2 + N_2] + 4\beta\mathbb{E}[I_1 + N_1]] \\
 &=: \frac{2}{3}L^2\xi_2^2S_2,
 \end{aligned}$$

where the last line follows since $1 - \beta \leq 1$ and $S_2 := 2\mathbb{E}[I_2 + N_2] + 4\beta\mathbb{E}[I_1 + N_1]$.

Next, we bound the term S_2 as follows:

$$\begin{aligned}
 S_2 &= 2\mathbb{E}[I_2 + N_2] + 4\beta\mathbb{E}[I_1 + N_1] \\
 &\leq 2\left(\beta\mathbb{E}[I_1 + N_1] + \frac{5-3\beta}{2}\mathbb{E}[N_2]\right) + 4\beta\mathbb{E}[I_1 + N_1] && \text{apply (46) for } t=2 \\
 &= (5-3\beta)\mathbb{E}[N_2] + 6\beta\mathbb{E}[I_1 + N_1] \\
 &\leq (5-3\beta)\mathbb{E}[N_2] + 6\beta\frac{5-3\beta}{2}\mathbb{E}[N_1] && \text{apply (48)} \\
 &= (5-3\beta)\mathbb{E}[N_2] + 3\beta(5-3\beta)\mathbb{E}[N_1] \\
 &\leq \frac{9}{2}(5-3\beta)\sum_{j=1}^2\beta^{2-j}\mathbb{E}[N_j].
 \end{aligned}$$

Hence, the statements $\mathbb{E}[\beta J_t + (1-\beta)I_t] \leq \frac{2}{3}L^2\xi_t^2S_t$ and $S_t \leq \frac{9}{2}(5-3\beta)\sum_{j=1}^t\beta^{t-j}\mathbb{E}[N_j]$ are true for every $t \geq 2$.

Combining these two cases, we have the following estimate

$$\begin{aligned}
 \mathbb{E}[\beta J_t + (1-\beta)I_t] &\leq \frac{2}{3}L^2\xi_t^2S_t \leq L^2\xi_t^2 \cdot \frac{2}{3} \cdot \frac{9}{2}(5-3\beta)\sum_{j=1}^t\beta^{t-j}\mathbb{E}[N_j] \\
 &\stackrel{(16)}{\leq} 3L^2\xi_t^2(5-3\beta)\sum_{j=1}^t\beta^{t-j}\mathbb{E}\left[(\Theta+n)\|\nabla F(w_0^{(j)})\|^2 + \sigma^2\right] \\
 &\leq 3L^2\xi_t^2(5-3\beta)\left[(\Theta+n)\sum_{j=1}^t\beta^{t-j}\mathbb{E}\left[\|\nabla F(w_0^{(j)})\|^2\right] + \sum_{j=1}^t\beta^{t-j}\sigma^2\right] \\
 &\leq 3L^2\xi_t^2(5-3\beta)\left[(\Theta+n)\sum_{j=1}^t\beta^{t-j}\mathbb{E}\left[\|\nabla F(w_0^{(j)})\|^2\right] + \frac{\sigma^2}{1-\beta}\right],
 \end{aligned}$$

where the last line follows since $\sum_{j=1}^t\beta^{t-j} \leq \frac{1}{1-\beta}$ for every $t \geq 2$. Hence, we have proved (42). $\square$

8.5.4. THE PROOF OF LEMMA 10

First, from the assumption $0 < \eta_t \leq \frac{1}{L\sqrt{D}}$, $t \geq 1$, we have $0 < \eta_t^2 \leq \frac{1}{DL^2}$. Next, from (17), we have $\xi_t = \max(\eta_t; \eta_{t-1})$ for $t > 1$ and $\xi_1 = \eta_1$, which lead to $0 < \xi_t^2 \leq \frac{1}{DL^2}$ for $t \geq 1$. Moreover, from the definition of $D = \max\left(\frac{5}{3}, \frac{6(5-3\beta)(\Theta+n)}{n(1-\beta)}\right)$ in Theorem 2, we have $L^2\xi_t^2 \leq \frac{3}{5}$ and $6L^2\xi_t^2(5-3\beta)(\Theta+n) \leq n(1-\beta)$ for $t \geq 1$.

Similar to the estimate (39) in Theorem 2, we get the following result:

$$F(w_0^{(2)}) \leq F(w_0^{(1)}) - \frac{\eta_1}{2}(1-\beta)\|\nabla F(w_0^{(1)})\|^2 + \frac{\eta_1}{2}(1-\beta)\frac{1}{n}I_1.$$

Taking total expectation and applying Lemma 8 we further have:

$$\begin{aligned}
 \mathbb{E}\left[F(w_0^{(2)})\right] &\leq \mathbb{E}\left[F(w_0^{(1)})\right] - \frac{\eta_1}{2}(1-\beta)\mathbb{E}\left[\|\nabla F(w_0^{(1)})\|^2\right] + \frac{\eta_1}{2n}(1-\beta)\mathbb{E}[I_1] \\
 &\stackrel{(45)}{\leq} \mathbb{E}\left[F(w_0^{(1)})\right] - \frac{\eta_1}{2}(1-\beta)\mathbb{E}\left[\|\nabla F(w_0^{(1)})\|^2\right] + \frac{\eta_1}{2n}(1-\beta)L^2\xi_1^2(1-\beta)^2 \cdot \mathbb{E}[I_1 + N_1] \\
 &\stackrel{(48)}{\leq} \mathbb{E}\left[F(w_0^{(1)})\right] - \frac{\eta_1}{2}(1-\beta)\mathbb{E}\left[\|\nabla F(w_0^{(1)})\|^2\right] + \frac{\eta_1}{2n}(1-\beta)^3L^2\xi_1^2 \cdot \frac{5-3\beta}{2}\mathbb{E}[N_1].
 \end{aligned}$$

Noting that $\tilde{w}_{t-1} = w_0^{(t)}$, $1 - \beta < 1$ and $N_t = (\Theta + n)\|\nabla F(w_0^{(t)})\|^2 + \sigma^2$, we get

$$\mathbb{E}[F(\tilde{w}_1)] \leq \mathbb{E}[F(\tilde{w}_0)] - \frac{\eta_1}{2}(1 - \beta)\mathbb{E}[\|\nabla F(\tilde{w}_0)\|^2] + \frac{\eta_1 L^2 \xi_1^2 (1 - \beta)(5 - 3\beta)}{4n} \left[(\Theta + n)\mathbb{E}[\|\nabla F(\tilde{w}_0)\|^2] + \sigma^2 \right].$$

Rearranging the terms and dividing both sides by $(1 - \beta)$, we have

$$\frac{\eta_1}{2}\mathbb{E}[\|\nabla F(\tilde{w}_0)\|^2] \leq \frac{\mathbb{E}[F(\tilde{w}_0) - F(\tilde{w}_1)]}{(1 - \beta)} + \frac{\eta_1 L^2 \xi_1^2 \cdot (5 - 3\beta)}{4n} (\Theta + n)\mathbb{E}[\|\nabla F(\tilde{w}_0)\|^2] + \frac{\eta_1 \sigma^2 L^2 \xi_1^2 (5 - 3\beta)}{4n}.$$

Since ξ_t satisfies $6L^2 \xi_t^2 (5 - 3\beta)(\Theta + n) \leq n(1 - \beta)$ as above, we can deduce from the last estimate that

$$\frac{\eta_1}{2}\mathbb{E}[\|\nabla F(\tilde{w}_0)\|^2] \leq \frac{\mathbb{E}[F(\tilde{w}_0) - F(\tilde{w}_1)]}{(1 - \beta)} + \frac{(1 - \beta)\eta_1}{4}\mathbb{E}[\|\nabla F(\tilde{w}_0)\|^2] + \frac{3\sigma^2(5 - 3\beta)L^2}{2n(1 - \beta)} \cdot \xi_1^3,$$

which proves (43). $\square$

9. The Proof of Technical Results in Section 4

This section provides the full proofs of the results in Section 4. Before proving Theorem 3, we introduce some common quantities and provide four technical lemmas.

9.1. Technical Lemmas

For the sake of our analysis, we introduce the following quantities:

$$h_i^{(t)} := g_i^{(t)} - \nabla f(w_0^{(t)}; \pi(i + 1)), \quad t \geq 1, i = 0, \dots, n - 1; \quad (50)$$

$$k_i^{(t-1)} := g_i^{(t-1)} - \nabla f(w_0^{(t)}; \pi(i + 1)), \quad t \geq 2, i = 0, \dots, n - 1; \quad (51)$$

$$p_t := \frac{1}{n} \sum_{i=0}^{n-1} m_{i+1}^{(t-1)}, \quad t \geq 2; \quad (52)$$

$$q_t := \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left[\left(\sum_{j=n-i}^{n-1} \beta^j \right) g_i^{(t-1)} + \left(\sum_{j=0}^{n-i-1} \beta^j \right) g_i^{(t)} \right], \quad t \geq 2. \quad (53)$$

The proof of Theorem 3 requires the following lemmas as intermediate steps of its proof.

Lemma 12. *Let $\{w_i^{(t)}\}$ be generated by Algorithm 2 with $0 \leq \beta < 1$ and $\eta_i^{(t)} := \frac{\eta_t}{n}$ for every $t \geq 1$. Suppose that Assumption 2 holds. Then we have*

$$(a) \quad \|m_{i+1}^{(t)}\|^2 \leq G^2, \quad \text{for } t \geq 1 \text{ and } i = 0, \dots, n - 1; \quad (54)$$

$$(b) \quad \|p_t\|^2 \leq G^2, \text{ for } t \geq 2 \quad \text{and} \quad \|\nabla F(w_0^{(t)})\|^2 \leq G^2, \text{ for } t \geq 1. \quad (55)$$

Lemma 13. *Under the same setting as of Lemma 12 and Assumption 1(b), if ξ_t is defined by (17), then for $i = 0, 1, \dots, n - 1$ and $t \geq 2$, it holds that*

$$\max \left(\|h_i^{(t)}\|^2, \|k_i^{(t-1)}\|^2 \right) \leq L^2 G^2 \xi_t^2. \quad (56)$$

Lemma 14. *Under the same setting as of Lemma 12, for $i = 0, 1, \dots, n - 1$ and $t \geq 2$, we have*

$$m_{i+1}^{(t)} = \beta^n m_{i+1}^{(t-1)} + (1 - \beta) \left(\beta^{n-1} g_{i+1}^{(t-1)} + \dots + \beta^{i+1} g_{n-1}^{(t-1)} + \beta^i g_0^{(t)} + \dots + \beta g_{i-1}^{(t)} + g_i^{(t)} \right).$$

From this expression we obtain the following sum:

$$\frac{1}{n} \sum_{i=0}^{n-1} m_{i+1}^{(t)} = \beta^n p_t + q_t, \text{ for } t \geq 2; \quad (57)$$

where p_t and q_t are defined in (52) and (53), respectively.

Lemma 15. *Under the same setting as in Theorem 3, the initial objective value $F(\tilde{w}_1)$ is upper bounded by*

$$F(\tilde{w}_1) \leq F(\tilde{w}_0) + \frac{1}{2L} \|\nabla F(\tilde{w}_0)\|^2 + L\eta_1^2 G^2. \quad (58)$$

9.2. The Proof of Theorem 3: Key Bound for Algorithm 2

From the update $w_{i+1}^{(t)} := w_i^{(t)} - \eta_i^{(t)} m_{i+1}^{(t)}$ at Step 2 and Step 2 of Algorithm 2 with $\eta_i^{(t)} := \frac{\eta_t}{n}$, for $i = 1, \dots, n-1$, we have

$$w_i^{(t)} = w_{i-1}^{(t)} - \frac{\eta_t}{n} m_i^{(t)} = w_0^{(t)} - \frac{\eta_t}{n} \sum_{j=0}^{i-1} m_{j+1}^{(t)}. \quad (59)$$

Now, letting $i = n$ in the estimate and noting that $w_0^{(t+1)} = w_n^{(t)}$ for all $t \geq 1$, we obtain

$$w_0^{(t+1)} - w_0^{(t)} = w_n^{(t)} - w_0^{(t)} = -\frac{\eta_t}{n} \sum_{j=0}^{n-1} m_{j+1}^{(t)} \quad (60)$$

From this update and the L -smoothness of F from Assumption 1(c), for $t \geq 2$, we can derive

$$\begin{aligned} F(w_0^{(t+1)}) &\stackrel{(5)}{\leq} F(w_0^{(t)}) + \nabla F(w_0^{(t)})^\top (w_0^{(t+1)} - w_0^{(t)}) + \frac{L}{2} \|w_0^{(t+1)} - w_0^{(t)}\|^2 \\ &\stackrel{(60)}{=} F(w_0^{(t)}) - \eta_t \nabla F(w_0^{(t)})^\top \left(\frac{1}{n} \sum_{j=0}^{n-1} m_{j+1}^{(t)} \right) + \frac{L\eta_t^2}{2} \left\| \frac{1}{n} \sum_{j=0}^{n-1} m_{j+1}^{(t)} \right\|^2 \\ &\stackrel{(57)}{=} F(w_0^{(t)}) - \eta_t \nabla F(w_0^{(t)})^\top (\beta^n p_t + q_t) + \frac{L\eta_t^2}{2} \|\beta^n p_t + q_t\|^2 \\ &\stackrel{(a)}{\leq} F(w_0^{(t)}) - \eta_t \beta^n \nabla F(w_0^{(t)})^\top p_t - \eta_t \nabla F(w_0^{(t)})^\top q_t + \frac{L\eta_t^2}{2} \beta^n \|p_t\|^2 + \frac{L\eta_t^2}{2(1-\beta^n)} \|q_t\|^2 \\ &\stackrel{(b)}{\leq} F(w_0^{(t)}) + \eta_t \beta^n G^2 - \eta_t \nabla F(w_0^{(t)})^\top q_t + \frac{L\eta_t^2}{2} \beta^n G^2 + \frac{L\eta_t^2}{2(1-\beta^n)} \|q_t\|^2 \\ &\stackrel{(c)}{\leq} F(w_0^{(t)}) + 2\eta_t \beta^n G^2 - \eta_t \nabla F(w_0^{(t)})^\top q_t + \frac{L\eta_t^2}{2(1-\beta^n)} \|q_t\|^2 \end{aligned}$$

where (a) comes from the convexity of squared norm: $\|\beta^n p_t + q_t\|^2 \leq \beta^n \|p_t\|^2 + (1-\beta^n) \|q_t\|^2$ for $0 \leq \beta^n < 1$, (b) is from the inequalities $\|F(w_0^{(t)})\| \leq G$ and $\|p_t\| \leq G$ noted in Lemma 12, and (c) follows from the fact that $L\eta_t \leq 1$.

Next, we focus on processing the term $\nabla F(w_0^{(t)})^\top q_t$. We can further upper bound the above expression as

$$\begin{aligned} F(w_0^{(t+1)}) &\leq F(w_0^{(t)}) + 2\eta_t \beta^n G^2 - \eta_t \nabla F(w_0^{(t)})^\top q_t + \frac{L\eta_t^2}{2(1-\beta^n)} \|q_t\|^2 \\ &= F(w_0^{(t)}) + 2\eta_t \beta^n G^2 - \frac{\eta_t}{1-\beta^n} (1-\beta^n) \nabla F(w_0^{(t)})^\top q_t + \frac{L\eta_t^2}{2(1-\beta^n)} \|q_t\|^2 \\ &\stackrel{(d)}{=} F(w_0^{(t)}) + 2\eta_t \beta^n G^2 + \frac{\eta_t}{2(1-\beta^n)} \|(1-\beta^n) \nabla F(w_0^{(t)}) - q_t\|^2 \\ &\quad - \frac{\eta_t}{2(1-\beta^n)} \|(1-\beta^n) \nabla F(w_0^{(t)})\|^2 - \frac{\eta_t}{2(1-\beta^n)} \|q_t\|^2 + \frac{L\eta_t^2}{2(1-\beta^n)} \|q_t\|^2 \\ &\stackrel{(e)}{\leq} F(w_0^{(t)}) + 2\eta_t \beta^n G^2 + \frac{\eta_t}{2(1-\beta^n)} M_t - \frac{\eta_t(1-\beta^n)}{2} \|\nabla F(w_0^{(t)})\|^2, \end{aligned} \quad (61)$$

where $M_t := \|(1-\beta^n) \nabla F(w_0^{(t)}) - q_t\|^2$. Here, (d) follows from the equality $u^\top v = \frac{1}{2} (\|u\|^2 + \|v\|^2 - \|u-v\|^2)$, and (e) comes from the fact that $\eta_t \leq \frac{1}{L}$.

Note that $1 - \beta^n = (1 - \beta) \sum_{j=0}^{n-1} \beta^j$ and $\nabla F(w_0^{(t)}) = \frac{1}{n} \sum_{i=0}^{n-1} \nabla f(w_0^{(t)}; \pi(i+1))$, we can rewrite

$$(1 - \beta^n) \nabla F(w_0^{(t)}) = \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left[\left(\sum_{j=0}^{n-1} \beta^j \right) \nabla f(w_0^{(t)}; \pi(i+1)) \right]. \quad (62)$$

Recall the definition of q_t , $h_i^{(t)}$, and $k_i^{(t-1)}$ from (53), (50), and (51), respectively, we can bound M_t as

$$\begin{aligned} M_t &:= \|q_t - (1 - \beta^n) \nabla F(w_0^{(t)})\|^2 \\ &\stackrel{(53), (62)}{=} \left\| \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left[\left(\sum_{j=n-i}^{n-1} \beta^j \right) g_i^{(t-1)} + \left(\sum_{j=0}^{n-i-1} \beta^j \right) g_i^{(t)} - \left(\sum_{j=0}^{n-1} \beta^j \right) \nabla f(w_0^{(t)}; \pi(i+1)) \right] \right\|^2 \\ &\stackrel{(50), (51)}{=} \left\| \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left[\left(\sum_{j=n-i}^{n-1} \beta^j \right) k_i^{(t-1)} + \left(\sum_{j=0}^{n-i-1} \beta^j \right) h_i^{(t)} \right] \right\|^2 \\ &\leq \frac{(1 - \beta)^2}{n} \sum_{i=0}^{n-1} \left\| \left(\sum_{j=n-i}^{n-1} \beta^j \right) k_i^{(t-1)} + \left(\sum_{j=0}^{n-i-1} \beta^j \right) h_i^{(t)} \right\|^2, \end{aligned}$$

where the last line comes from the Cauchy-Schwarz inequality. We normalize the last squared norms as follows:

$$\begin{aligned} M_t &\leq \frac{(1 - \beta)^2}{n} \left(\sum_{j=0}^{n-1} \beta^j \right)^2 \sum_{i=0}^{n-1} \left\| \frac{\sum_{j=n-i}^{n-1} \beta^j}{\sum_{j=0}^{n-1} \beta^j} k_i^{(t-1)} + \frac{\sum_{j=0}^{n-i-1} \beta^j}{\sum_{j=0}^{n-1} \beta^j} h_i^{(t)} \right\|^2 \\ &\leq \frac{(1 - \beta)^2}{n} \left(\sum_{j=0}^{n-1} \beta^j \right)^2 \sum_{i=0}^{n-1} \left[\frac{\sum_{j=n-i}^{n-1} \beta^j}{\sum_{j=0}^{n-1} \beta^j} \|k_i^{(t-1)}\|^2 + \frac{\sum_{j=0}^{n-i-1} \beta^j}{\sum_{j=0}^{n-1} \beta^j} \|h_i^{(t)}\|^2 \right] \quad \text{by convexity of } \|\cdot\|^2 \\ &\leq \frac{(1 - \beta)^2}{n} \left(\sum_{j=0}^{n-1} \beta^j \right)^2 \sum_{i=0}^{n-1} \max \left(\|k_i^{(t-1)}\|^2, \|h_i^{(t)}\|^2 \right) \\ &= \frac{(1 - \beta^n)^2}{n} \sum_{i=0}^{n-1} \max \left(\|k_i^{(t-1)}\|^2, \|h_i^{(t)}\|^2 \right) \quad \text{since } (1 - \beta) \sum_{j=0}^{n-1} \beta^j = 1 - \beta^n \\ &\leq (1 - \beta^n)^2 L^2 G^2 \xi_t^2, \end{aligned}$$

where the last estimate follows from Lemma 13 with $\xi_t^2 = \max(\eta_t^2, \eta_{t-1}^2)$ for $t \geq 2$. Substituting the last estimate into (61) for $t \geq 2$, we obtain

$$F(w_0^{(t+1)}) \leq F(w_0^{(t)}) + 2\eta_t \beta^n G^2 + \frac{\eta_t(1 - \beta^n)}{2} L^2 G^2 \xi_t^2 - \frac{\eta_t(1 - \beta^n)}{2} \|\nabla F(w_0^{(t)})\|^2.$$

Since $\tilde{w}_{t-1} = w_0^{(t)}$ and $\tilde{w}_t = w_0^{(t+1)}$, this inequality leads to

$$\eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 \leq \frac{2[F(\tilde{w}_{t-1}) - F(\tilde{w}_t)]}{(1 - \beta^n)} + \frac{4\beta^n G^2}{1 - \beta^n} \eta_t + L^2 G^2 \xi_t^3, \quad t \geq 2.$$

Now, using the fact that $\xi_1 = \eta_1$, we can derive the following statement for $t = 1$:

$$\eta_1 \|\nabla F(\tilde{w}_0)\|^2 \leq \frac{\eta_1 \|\nabla F(\tilde{w}_0)\|^2}{1 - \beta^n} + \frac{4\beta^n G^2}{1 - \beta^n} \eta_1 + L^2 G^2 \xi_1^3.$$

Summing the previous statement for $t = 2, 3, \dots, T$, and using the last one, we can deduce that

$$\begin{aligned} \sum_{t=1}^T \left(\eta_t \|\nabla F(\tilde{w}_{t-1})\|^2 \right) &\leq \frac{2[F(\tilde{w}_1) - F_*] + \eta_1 \|\nabla F(\tilde{w}_0)\|^2}{(1 - \beta^n)} + \frac{4\beta^n G^2}{1 - \beta^n} \sum_{t=1}^T \eta_t + L^2 G^2 \sum_{t=1}^T \xi_t^3 \\ &\stackrel{(58)}{\leq} \frac{2[F(\tilde{w}_0) - F_*] + \left(\frac{1}{L} + \eta_1\right) \|\nabla F(\tilde{w}_0)\|^2 + 2L\eta_1^2 G^2}{(1 - \beta^n)} + \frac{4\beta^n G^2}{1 - \beta^n} \sum_{t=1}^T \eta_t + L^2 G^2 \sum_{t=1}^T \xi_t^3. \end{aligned}$$

Finally, dividing both sides of this estimate by $\sum_{t=1}^T \eta_t$, we obtain (10). $\square$

9.3. The Proof of Corollaries 7 and 8: Constant and Diminishing Learning Rates

The proof of Corollary 7. First, from $\beta = \left(\frac{\nu}{T^{2/3}}\right)^{1/n}$ we have $\beta^n = \frac{\nu}{T^{2/3}}$. Since $\beta \leq \left(\frac{R-1}{R}\right)^{1/n}$ for some $R \geq 1$, we have $\frac{1}{1-\beta^n} \leq R$. We also have $\eta_t := \frac{\gamma}{T^{1/3}}$ and $\eta_t \leq \frac{1}{L}$ for all $t \geq 1$. Moreover, $\sum_{t=1}^T \eta_t = \gamma T^{2/3}$ and $\sum_{t=1}^T \xi_t^3 = \gamma^3$. Substituting these expressions into (10), we obtain

$$\begin{aligned} \mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] &\leq \frac{\Delta_1}{\gamma(1-\beta^n)T^{2/3}} + \frac{L^2 G^2 \gamma^2}{T^{2/3}} + \frac{4\beta^n G^2}{1-\beta^n} \\ &\leq \frac{\Delta_1 R}{\gamma T^{2/3}} + \frac{L^2 G^2 \gamma^2}{T^{2/3}} + 4\beta^n G^2 R \\ &= \frac{\Delta_1 R}{\gamma T^{2/3}} + \frac{L^2 G^2 \gamma^2}{T^{2/3}} + \frac{4\nu G^2 R}{T^{2/3}} \\ &\leq \frac{\frac{R}{\gamma} \Delta_1 + L^2 G^2 \gamma^2 + 4\nu G^2 R}{T^{2/3}}. \end{aligned}$$

Here, we have $\Delta_1 = 2[F(\tilde{w}_0) - F_*] + \left(\frac{1}{L} + \frac{\gamma}{T^{1/3}}\right) \|\nabla F(\tilde{w}_0)\|^2 + \frac{2LG^2\gamma^2}{T^{2/3}}$ as defined in Theorem 3. Substituting this expression into the last inequality, we obtain the main inequality in Corollary 7. $\square$

The proof of Corollary 8. We have $\eta_t = \frac{\gamma}{(t+\lambda)^{1/3}}$ and $\eta_t \leq \frac{1}{L}$ for all $t \geq 1$. Moreover, we have

$$\begin{aligned} \sum_{t=1}^T \eta_t &= \gamma \sum_{t=1}^T \frac{1}{(t+\lambda)^{1/3}} \geq \gamma \int_1^T \frac{d\tau}{(\tau+\lambda)^{1/3}} \geq \gamma [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}], \\ \sum_{t=1}^T \xi_t^3 &= 2\eta_1^3 + \gamma^3 \sum_{t=3}^T \frac{1}{t-1+\lambda} \leq \frac{2\gamma^3}{1+\lambda} + \gamma^3 \int_{t=2}^T \frac{d\tau}{\tau-1+\lambda} \leq \gamma^3 \left[\frac{2}{1+\lambda} + \log(T-1+\lambda) \right]. \end{aligned}$$

Substituting these expressions with $\beta = \left(\frac{\nu}{T^{2/3}}\right)^{1/n}$ into (10), we obtain

$$\begin{aligned} \mathbb{E}[\|\nabla F(\hat{w}_T)\|^2] &\leq \frac{\Delta_1}{\gamma [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}] (1-\beta^n)} + \frac{L^2 G^2 \gamma^2 \left(\frac{2}{1+\lambda} + \log(T-1+\lambda) \right)}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}} + \frac{4\beta^n G^2}{1-\beta^n} \\ &\leq \frac{\Delta_1 R}{\gamma [(T+\lambda)^{2/3} - (1+\lambda)^{2/3}]} + \frac{L^2 G^2 \gamma^2 \left(\frac{2}{1+\lambda} + \log(T-1+\lambda) \right)}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}} + 4\beta^n G^2 R \\ &= \frac{\frac{R}{\gamma} \Delta_1}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}} + \frac{L^2 G^2 \gamma^2 \left(\frac{2}{1+\lambda} + \log(T-1+\lambda) \right)}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}} + \frac{4\nu G^2 R}{T^{2/3}} \\ &= \frac{\frac{R}{\gamma} \Delta_1 + \frac{2}{1+\lambda} L^2 G^2 \gamma^2}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}} + \frac{L^2 G^2 \gamma^2 \log(T-1+\lambda)}{(T+\lambda)^{2/3} - (1+\lambda)^{2/3}} + \frac{4\nu G^2 R}{T^{2/3}}. \end{aligned}$$

Substituting Δ_1 from Theorem 3 and $\eta_1 = \frac{\gamma}{(1+\lambda)^{1/3}}$ into the last estimate, we obtain the inequality in Corollary 8. $\square$

9.4. The Proof of Technical Lemmas

We now provide the full proof of lemmas that serve for the proof of Theorem 3 above.

9.4.1. THE PROOF OF LEMMA 12

We prove the first estimate of Lemma 12 by induction.

- For $t = 1$, it is obvious that $\|m_0^{(1)}\|^2 = \|\tilde{m}_0\|^2 = 0 \leq G^2$.
- For $t = 1$, assume that (54) holds for $i - 1$, that is, $\|m_i^{(1)}\|^2 \leq G^2$. From the update $m_{i+1}^{(1)} = \beta m_i^{(1)} + (1-\beta)g_k^{(1)}$, by our induction hypothesis, convexity of $\|\cdot\|^2$, $0 \leq \beta < 1$, and Assumption 2, we have

$$\|m_{i+1}^{(1)}\|^2 \leq \beta \|m_i^{(1)}\|^2 + (1-\beta) \|g_i^{(1)}\|^2 \leq \beta G^2 + (1-\beta) G^2 = G^2.$$

Consequently, for $t = 1$, we have $\|m_i^{(1)}\|^2 \leq G^2$ for all $i = 0, \dots, n-1$.

- Now, we prove (54) for $t > 1$. Assume that (54) holds for epoch $t-1$, that is $\|m_{i+1}^{(t-1)}\|^2 \leq G^2$ for every $i = 0, \dots, n-1$. Since $m_0^{(t)} := m_n^{(t-1)}$, we have $\|m_0^{(t)}\|^2 = \|\tilde{m}_{t-1}\|^2 = \|m_n^{(t-1)}\|^2 \leq G^2$. Using the previous argument from the case $t = 1$, we also get $\|m_{i+1}^{(t)}\|^2 \leq G^2$ for all $i = 0, \dots, n-1$.

Combining all the cases above, we have shown that (54) holds for all $i = 0, \dots, n-1$ and $t = 1, \dots, T$.

Now using the Cauchy-Schwarz inequality, we get

$$\|p_t\|^2 = \left\| \frac{1}{n} \sum_{i=0}^{n-1} m_{i+1}^{(t-1)} \right\|^2 \leq \frac{1}{n} \sum_{i=0}^{n-1} \|m_{i+1}^{(t-1)}\|^2 \leq G^2, \text{ for } t \geq 2.$$

Similarly, for $t \geq 1$, we also have

$$\|\nabla F(w_0^{(t)})\|^2 = \left\| \frac{1}{n} \sum_{i=0}^{n-1} \nabla f(w_0^{(t)}; \pi(i+1)) \right\|^2 \leq \frac{1}{n} \sum_{i=0}^{n-1} \|\nabla f(w_0^{(t)}; \pi(i+1))\|^2 \leq G^2,$$

which proves the second estimate (55) of Lemma 12. $\square$

9.4.2. THE PROOF OF LEMMA 13

Using the L -smoothness assumption of $f(\cdot; i)$, we derive the following for $t \geq 2$:

$$\begin{aligned} \|h_i^{(t)}\|^2 &\stackrel{(50)}{=} \|\nabla f(w_i^{(t)}; \pi(i+1)) - \nabla f(w_0^{(t)}; \pi(i+1))\|^2 \leq L^2 \|w_i^{(t)} - w_0^{(t)}\|^2, \\ \|k_i^{(t-1)}\|^2 &\stackrel{(51)}{=} \|\nabla f(w_i^{(t-1)}; \pi(i+1)) - \nabla f(w_0^{(t)}; \pi(i+1))\|^2 \leq L^2 \|w_i^{(t-1)} - w_0^{(t)}\|^2. \end{aligned}$$

Now from the update $w_{i+1}^{(t)} := w_i^{(t)} - \eta_i^{(t)} m_{i+1}^{(t)}$ in Algorithm 2 with $\eta_i^{(t)} := \frac{\eta_t}{n}$, for $i = 1, \dots, n-1$ and $t \geq 2$, we have

$$\begin{aligned} w_i^{(t)} &= w_{i-1}^{(t)} - \frac{\eta_t}{n} m_i^{(t)} = w_0^{(t)} - \frac{\eta_t}{n} \sum_{j=0}^{i-1} m_{j+1}^{(t)}, \\ w_0^{(t)} &= w_n^{(t-1)} = w_{n-1}^{(t-1)} - \frac{\eta_{t-1}}{n} m_n^{(t-1)} = w_i^{(t-1)} - \frac{\eta_{t-1}}{n} \sum_{j=i}^{n-1} m_{j+1}^{(t-1)}. \end{aligned}$$

Therefore, for $i = 0, \dots, n-1$ and $t \geq 2$, using these expressions and (54), we can bound

$$\begin{aligned} \|w_i^{(t)} - w_0^{(t)}\|^2 &= \eta_t^2 \left\| \frac{1}{n} \sum_{j=0}^{i-1} m_{j+1}^{(t)} \right\|^2 \stackrel{(54)}{\leq} \eta_t^2 G^2, \\ \|w_i^{(t-1)} - w_0^{(t)}\|^2 &= \eta_{t-1}^2 \left\| \frac{1}{n} \sum_{j=i}^{n-1} m_{j+1}^{(t-1)} \right\|^2 \stackrel{(54)}{\leq} \eta_{t-1}^2 G^2. \end{aligned}$$

Substituting these inequalities into the first two expressions, for $i = 0, \dots, n-1$ and $t \geq 2$, we get

$$\max \left(\|h_i^{(t)}\|^2, \|k_i^{(t-1)}\|^2 \right) \leq L^2 \max \left(\eta_t^2 G^2; \eta_{t-1}^2 G^2 \right) \leq L^2 G^2 \max(\eta_t^2; \eta_{t-1}^2) = L^2 G^2 \max(\eta_t; \eta_{t-1})^2 = L^2 G^2 \xi_t^2,$$

which proves (56), where the last line follows from the facts that $\eta_t > 0$ and $\eta_{t-1} > 0$. $\square$

9.4.3. THE PROOF OF LEMMA 14

The first statement follows directly from the momentum update rule, i.e.:

$$\begin{aligned} m_{i+1}^{(t)} &= \beta m_i^{(t)} + (1-\beta) g_i^{(t)} = \beta^2 m_{i-1}^{(t)} + (1-\beta) \beta g_{i-1}^{(t)} + (1-\beta) g_i^{(t)} \\ &= \beta^{i+1} m_0^{(t)} + (1-\beta) \left(\beta^i g_0^{(t)} + \dots + \beta g_{i-1}^{(t)} + g_i^{(t)} \right) \end{aligned}$$

$$\begin{aligned}
 &\stackrel{(a)}{=} \beta^{i+1} m_n^{(t-1)} + (1 - \beta) \left(\beta^i g_0^{(t)} + \dots + \beta g_{i-1}^{(t)} + g_i^{(t)} \right) \\
 &= \beta^n m_{i+1}^{(t-1)} + (1 - \beta) \left(\beta^{n-1} g_{i+1}^{(t-1)} + \dots + \beta^{i+1} g_{n-1}^{(t-1)} + \beta^i g_0^{(t)} + \dots + \beta g_{i-1}^{(t)} + g_i^{(t)} \right),
 \end{aligned}$$

where (a) follows from the rule $m_0^{(t)} = \tilde{m}_{t-1} = m_n^{(t-1)}$.

Next, summing up this expression from $i := 0$ to $i := n - 1$ and averaging, we get

$$\frac{1}{n} \sum_{i=0}^{n-1} m_{i+1}^{(t)} = \beta^n p_t + q_t, \tag{63}$$

where

$$p_t := \frac{1}{n} \sum_{i=0}^{n-1} m_{i+1}^{(t-1)} \quad \text{and} \quad q_t := \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left(\beta^{n-1} g_{i+1}^{(t-1)} + \dots + \beta^{i+1} g_{n-1}^{(t-1)} + \beta^i g_0^{(t)} + \dots + g_i^{(t)} \right).$$

Now we consider the term q_t in the last expression:

$$\begin{aligned}
 q_t &= \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left(\beta^{n-1} g_{i+1}^{(t-1)} + \dots + \beta^{i+1} g_{n-1}^{(t-1)} + \beta^i g_0^{(t)} + \dots + \beta g_{i-1}^{(t)} + g_i^{(t)} \right) \\
 &= \frac{1 - \beta}{n} \left(\beta^{n-1} g_1^{(t-1)} + \beta^{n-2} g_2^{(t-1)} \dots + \beta^1 g_{n-1}^{(t-1)} + g_0^{(t)} \right) + \dots \\
 &\quad + \frac{1 - \beta}{n} \left(\beta^{n-1} g_{i+1}^{(t-1)} + \beta^{n-2} g_{i+2}^{(t-1)} + \dots + \beta^{i+1} g_{n-1}^{(t-1)} + \beta^i g_0^{(t)} \dots + \beta g_{i-1}^{(t)} + g_i^{(t)} \right) + \dots \\
 &\quad + \frac{1 - \beta}{n} \left(\beta^{n-1} g_0^{(t)} + \beta^{n-2} g_1^{(t)} \dots + \beta g_{n-2}^{(t)} + g_{n-1}^{(t)} \right).
 \end{aligned}$$

Reordering the terms $g_i^{(t)}$ and $g_i^{(t-1)}$ and noting that $\sum_{j=n}^{n-1} \beta^j = 0$ by convention, we get

$$\begin{aligned}
 q_t &= \frac{1 - \beta}{n} \sum_{i=1}^{n-1} \left[\left(\sum_{j=n-i}^{n-1} \beta^j \right) g_i^{(t-1)} \right] + \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left[\left(\sum_{j=0}^{n-i-1} \beta^j \right) g_i^{(t)} \right] \\
 &= \frac{1 - \beta}{n} \sum_{i=0}^{n-1} \left[\left(\sum_{j=n-i}^{n-1} \beta^j \right) g_i^{(t-1)} + \left(\sum_{j=0}^{n-i-1} \beta^j \right) g_i^{(t)} \right].
 \end{aligned}$$

Substituting this expression of q_t into (63), we obtain Equation (57) of Lemma 14. $\square$

9.4.4. THE PROOF OF LEMMA 15: UPPER BOUNDING THE INITIAL OBJECTIVE VALUE

Using (54), we have

$$\|w_0^{(2)} - w_0^{(1)}\|^2 = \|w_n^{(1)} - w_0^{(1)}\|^2 = \eta_1^2 \left\| \frac{1}{n} \sum_{i=0}^{n-1} m_{i+1}^{(1)} \right\|^2 \stackrel{(54)}{\leq} \eta_1^2 G^2. \tag{64}$$

Since F is L -smooth, we can derive

$$\begin{aligned}
 F(w_0^{(2)}) &\stackrel{(5)}{\leq} F(w_0^{(1)}) + \nabla F(w_0^{(1)})^\top (w_0^{(2)} - w_0^{(1)}) + \frac{L}{2} \|w_0^{(2)} - w_0^{(1)}\|^2 \\
 &\stackrel{(a)}{\leq} F(w_0^{(1)}) + \frac{1}{2L} \|\nabla F(w_0^{(1)})\|^2 + L \|w_0^{(2)} - w_0^{(1)}\|^2, \\
 &\stackrel{(64)}{\leq} F(w_0^{(1)}) + \frac{1}{2L} \|\nabla F(w_0^{(1)})\|^2 + L \eta_1^2 G^2,
 \end{aligned}$$

where (a) follows since $u^\top v \leq \frac{\|u\|^2}{2} + \frac{\|v\|^2}{2}$.

Substituting $w_0^{(2)} := \tilde{w}_1$ and $w_0^{(1)} := \tilde{w}_0$ into this estimate, we obtain (58). $\square$

10. Detailed Implementation and Additional Experiments

In this Supp. Doc., we explain the detailed hyper-parameter tuning strategy in Section 5 and provide additional experiments for our proposed methods.

10.1. Comparing SMG with Other Methods

We compare our SMG algorithm with Stochastic Gradient Descent (SGD) and two other methods: SGD with Momentum (SGD-M) (Polyak, 1964) and Adam (Kingma & Ba, 2014). To have a fair comparison, a random reshuffling strategy is applied to all methods. We tune each algorithm with a constant learning rate using grid search and select the hyper-parameters that perform best according to their results. We report the additional results on the squared norm of gradients of this experiment in Figure 7. The hyper-parameters tuning strategy for each method is given below:

- For SGD, the first (coarse) searching grid is $\{0.1, 0.01, 0.001\}$ and the fine grid is like $\{0.5, 0.4, 0.2, 0.1, 0.08, 0.06, 0.05\}$ if 0.1 is the best value we get after the first stage.
- For our new SMG algorithm, we fixed the parameter $\beta := 0.5$ since it usually performs the best in our experiments. We let the coarse searching grid be $\{1, 0.1, 0.01\}$ and the fine grid be $\{0.5, 0.4, 0.2, 0.1, 0.08, 0.06, 0.05\}$ if 0.1 is the best value of the first stage. Note that our SMG algorithm may work with a bigger learning rate than the traditional SGD algorithm.
- For SGD-M, we update the weights using the following rule:

$$\begin{aligned} m_{i+1}^{(t)} &:= \beta m_i^{(t)} + g_i^{(t)} \\ w_{i+1}^{(t)} &:= w_i^{(t)} - \eta_i^{(t)} m_{i+1}^{(t)}, \end{aligned}$$

where $g_i^{(t)}$ is the $(i + 1)$ -th gradient at epoch t . Note that this momentum update is implemented in PyTorch with the default value $\beta = 0.9$. Hence, we choose this setting for SGD-M, and we tune the learning rate using the two searching grids $\{0.1, 0.01, 0.001\}$ and $\{0.5, 0.4, 0.2, 0.1, 0.08, 0.06, 0.05\}$ as in the SGD algorithm.

- For Adam, we fixed two hyper-parameters $\beta_1 := 0.9$, $\beta_2 := 0.999$ as in the original paper. Since the default learning rate for Adam is 0.001, we let our coarse searching grid be $\{0.01, 0.001, 0.0001\}$, and the fine grid be $\{0.002, 0.001, 0.0005\}$ if 0.001 performs the best in the first stage. We note that since the best learning rate for Adam is usually 0.001, its hyper-parameter tuning process requires little effort than other algorithms in our experiments.

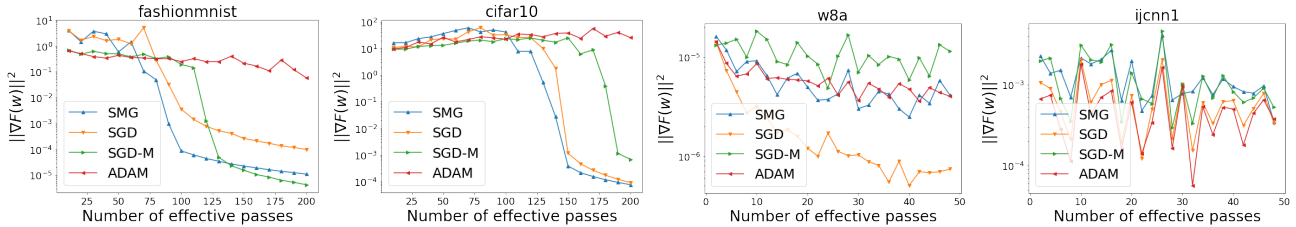


Figure 7. The squared norm of gradient produced by SMG, SGD, SGD-M, and Adam for different datasets.

10.2. The Choice of Hyper-parameter β

As presented in Section 5, in this experiment, we investigate the sensitivity of our proposed SMG algorithm to the hyper-parameter β . We choose a constant learning rate for each dataset and run the algorithm for different values of β in the linear grid $\{0.1, 0.2, 0.3, 0.4, 0.5, 0.6, 0.7, 0.8, 0.9\}$. However, the choice of $\beta \geq 0.6$ does not lead to a good performance, and, therefore, we omit to report them in our results. Figure 8 shows the comparison of $\|\nabla F(w)\|^2$ on different values of β for FashionMNIST, CIFAR-10, w8a, and ijcnn1 datasets. Our empirical study here shows that the choice of β in $[0.1, 0.5]$ works reasonably well, and $\beta := 0.5$ seems to be the best in our test.

10.3. Different Learning Rate Schemes

As presented in Section 5, in the last experiment, we examine the performance of our SMG method with the effect of four different learning rate schemes. We choose the hyper-parameter $\beta = 0.5$ since it tends to give the best results in our test.

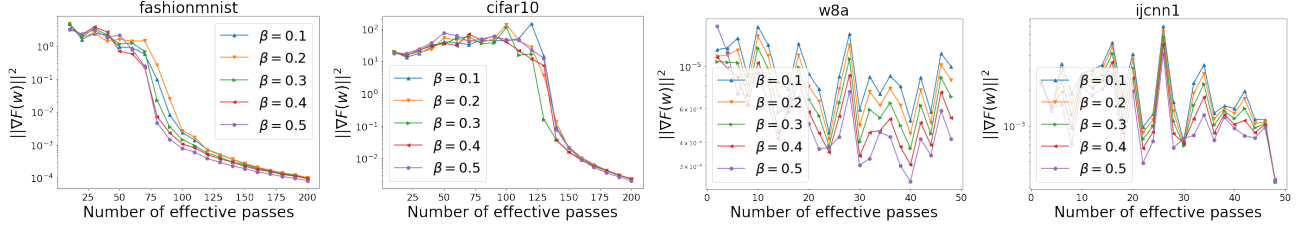


Figure 8. The squared norm of gradient reported by SMG with different β on FashionMNIST, CIFAR-10, w8a, and ijcnn1 datasets

The additional result on the squared norm gradient is reported in Figure 9.

- *Constant learning rate:* Similar to the first section, we set the coarse searching grid as $\{1, 0.1, 0.01\}$ and let the fine grid be $\{0.5, 0.4, 0.2, 0.1, 0.08, 0.06, 0.05\}$ if 0.1 is the best value of the first stage.
- *Cosine annealing learning rate:* For a fixed number of epoch T , we need one hyper-parameter η for the cosine learning rate scheme: $\eta_t = \eta(1 + \cos(t\pi/T))$. We choose a coarse grid $\{1, 0.1, 0.01, 0.001\}$ and a fine grid $\{0.5, 0.4, 0.2, 0.1, 0.08, 0.06, 0.05\}$ if 0.1 is the best value for the parameter η in the first stage.
- *Diminishing learning rate:* In order to apply the diminishing scheme $\eta_t = \frac{\gamma}{(t+\lambda)^{1/3}}$, we need two hyper-parameters γ and λ . At first, we let the searching grid for λ be $\{1, 2, 4, 8\}$. For the γ value, we set its searching grid so that the initial value η_1 lies in the first grid of $\{1, 0.1, 0.01\}$, and then lies in a fine grid centered at the best one of the coarse grid.
- *Exponential learning rate:* We choose two hyper-parameters for the exponential scheme $\eta_t = \eta\alpha^t$. First, let the searching grid for decay rate α be $\{0.99, 0.995, 0.999\}$. Then we set the searching grid for η similarly to the diminishing case (such that initial learning rate η_1 is in the coarse grid $\{1, 0.1, 0.01\}$, and then a second grid centered at the best value of the first).

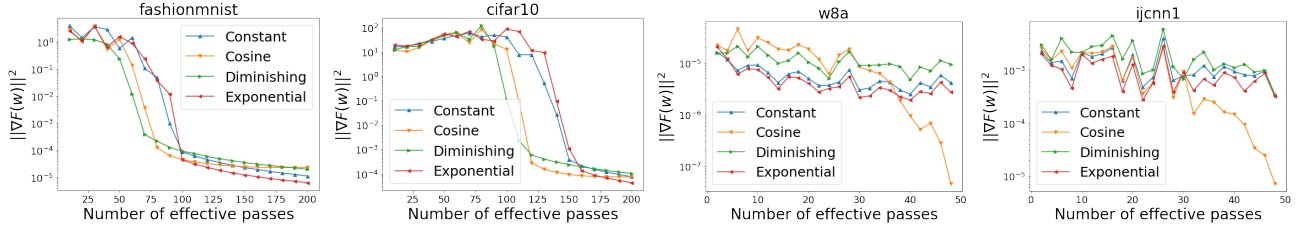


Figure 9. The squared norm of gradient produced by SMG under four different learning rate schemes on FashionMNIST, CIFAR-10, w8a, and ijcnn1 datasets.

10.4. Additional Experiment with Single Shuffling Scheme

Algorithm 2 (Single Shuffling Momentum Gradient - SSMG) uses the single shuffling strategy which covers incremental gradient as a special case. Therefore in this experiment we compare our proposed methods (SMG and SSMG) with SGD, SGD with Momentum (SGD-M) and Adam using the single shuffling scheme.

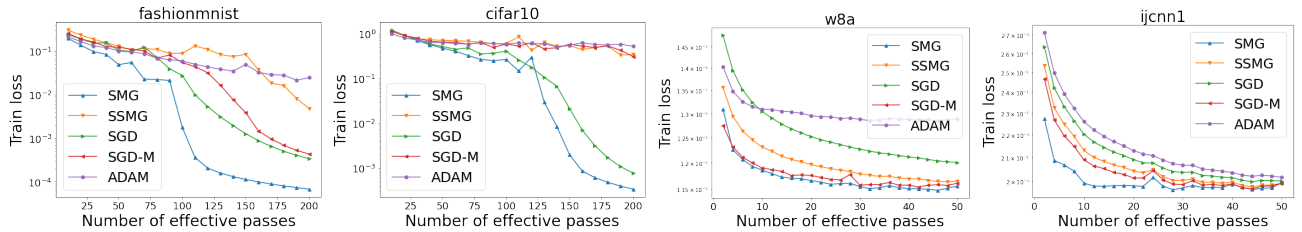


Figure 10. The train loss produced by SMG, SSMG, SGD, SGD-M, and Adam on the four different datasets: Fashion-MNIST, CIFAR-10, w8a, and ijcnn1.

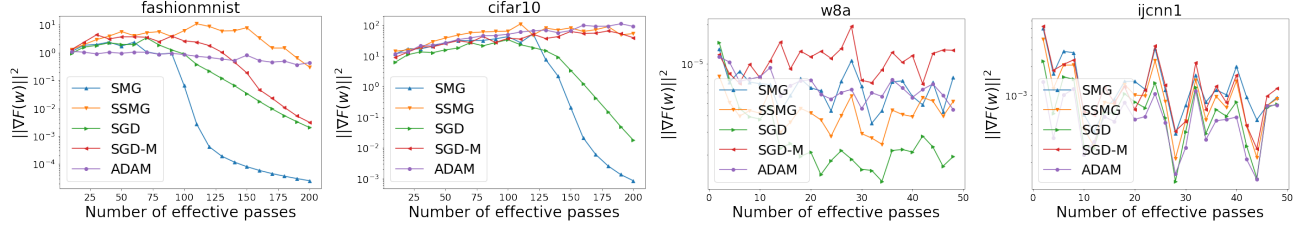


Figure 11. The squared norm of gradient produced by SMG, SSMG, SGD, SGD-M, and Adam on the four different datasets: Fashion-MNIST, CIFAR-10, w8a, and ijcnn1.

For the SSMG algorithm, the hyper-parameter β is chosen in the grid $\{0.1, 0.5, 0.9\}$. For the learning rate, we let the coarse searching grid be $\{0.1, 0.01, 0.001\}$ and the fine grid be $\{0.5, 0.4, 0.2, 0.1, 0.08, 0.06, 0.05\}$ if 0.1 is the best value of the first stage. For other methods, the hyper-parameters tuning strategy is similar to the settings in subsection 10.1.

The train loss $F(w)$ and the squared norm of gradient $\|\nabla F(w)\|^2$ of each methods are reported in Figures 10 and 11. We observe that SSMG does not work well in the first two datasets compared to SMG, SGD, and SGD-M, but it performs reasonably well in the last two datasets on the binary classification problem.