
Distributed Personalized Empirical Risk Minimization

Yuyang Deng
Pennsylvania State University
yzd82@psu.edu

Mohammad Mahdi Kamani
Wyze Labs
mmkamani@alumni.psu.edu

Pouria Mahdavinia
Pennsylvania State University
pxm5426@psu.edu

Mehrdad Mahdavi
Pennsylvania State University
mzm616@psu.edu

Abstract

This paper advocates a new paradigm Personalized Empirical Risk Minimization (PERM) to facilitate learning from heterogeneous data sources without imposing stringent constraints on computational resources shared by participating devices. In PERM, we aim to learn a distinct model for each client by learning who to learn with and personalizing the aggregation of local empirical losses by effectively estimating the statistical discrepancy among data distributions, which entails optimal statistical accuracy for all local distributions and overcomes the data heterogeneity issue. To learn personalized models at scale, we propose a distributed algorithm that replaces the standard model averaging with model shuffling to simultaneously optimize PERM objectives for all devices. This also allows us to learn distinct model architectures (e.g., neural networks with different numbers of parameters) for different clients, thus confining underlying memory and compute resources of individual clients. We rigorously analyze the convergence of the proposed algorithm and conduct experiments that corroborate the effectiveness of the proposed paradigm.

1 Introduction

Recently *federated learning* (FL) has emerged as an alternative paradigm to centralized learning to encourage federated model sharing and create a framework to support edge intelligence by shifting model training and inference from data centers to potentially scattered—and perhaps self-interested—systems where data is generated [1]. While undoubtedly being a better paradigm than centralized learning, enabling the widespread adoption of FL necessitates foundational advances in the efficient use of statistical and computational resources to encourage a large pool of individuals or corporations to share their private data and resources. Specifically, due to *heterogeneity* of data and compute resources among participants, it is necessary, if not imperative, to develop distributed algorithms that are i) cognizant of statistical heterogeneity (*data-awareness*) by designing algorithms that effectively deal with highly heterogeneous data distributions across devices; and ii) confined to learning models that meet available computational resources of participant devices (*system-awareness*).

To mitigate the negative effect of data heterogeneity (non-IIDness), two common approaches are clustering and personalization. The key idea behind the clustering-based methods [2, 3, 4, 5] is to partition the devices into clusters (coalitions) of similar data distributions and then learn a single shared model for all clients within each cluster. While appealing, the partitioning methods are limited to heuristic ideas such as clustering based on the geographical distribution of devices without taking the actual data distributions into account and lack theoretical guarantees or postulate strong assumptions on initial models or data distributions [4, 5]. In personalization-based methods [6, 7, 8, 2, 9, 10, 11], the idea is to learn a distinct *personalized model* for each device alongside the global model, which

can be unified as minimizing a bi-level optimization problem [12]. Personalization aims to learn a model that has the generalization capabilities of the global model but can also perform well on the specific data distribution of each participant suffers from a few key limitations. First, as the number of clients grows, while the number of training data increases, the number of parameters to be learned increases which limits to increase in the number of clients beyond a certain point to balance data and overall model complexity tradeoff— a phenomenon known as incidental parameters problem [13]. Moreover, since the knowledge transfer among data sources happens through a single global model, it might lead to suboptimal results. To see this, consider an extreme example, where half of the users have identical data distributions, say $\mathcal{D}$, while the other half share a data distribution that is substantially different, say $-\mathcal{D}$ (e.g. two distributions with same marginal distribution on features but opposite labeling functions). In this case, the global model obtained by naively aggregating local models (e.g., fixed mixture weights) converges to a solution that suffers from low test accuracy on all local distributions which makes it preferable to learn a model for each client solely based on its local data or carefully chosen subset of data sources.

Focusing on system heterogeneity, most existing works require learning models of identical architecture to be deployed across the clients and server (model-homogeneity) [14, 15], and mostly focus on reducing number [15] or size [16, 17, 18] of communications or sampling handling chaotic availability of clients [19, 20]. The requirement of the same model makes it infeasible to train large models due to *system heterogeneity* where client devices have drastically different computational resources. A few recent studies aim to overcome this issue either by leveraging knowledge distillation methods [21, 22, 23, 24] or partial training (PT) strategies via model subsampling (either static [25, 26], random [27], or rolling [28]). However, KD-based methods require having access to a public representative dataset of all local datasets at server and ignore data heterogeneity in the distillation stage to a large extent. The focus of PT training methods is mostly on learning a single server model using heterogeneous resources of devices and does not aim at deploying a model onto each client after the global server model is trained (which is left as a future direction in [28]). The aforementioned issues lead to a fundamental question: “*What is the best strategy to learn from heterogeneous data sources to achieve optimal accuracy w.r.t. each data source, without imposing stringent constraints on computational resources shared by participating devices?*”.

We answer this question affirmatively, by proposing a new *data&system-aware* paradigm dubbed Personalized Empirical Risk Minimization (PERM), to facilitate learning from massively fragmented private data under resource constraints. Motivated by generalization bounds in multiple source domain adaptation [29, 30, 31, 32], in PERM we aim to learn a distinct model for each client by *personalizing the aggregation* of empirical losses of different data sources which enables each client to *learn who to learn with* using an effective method to empirically estimate the statistical discrepancy between their associated data distributions. We argue that PERM entails optimal statistical accuracy for all local distributions, thus overcoming the data heterogeneity issue. PERM can also be employed in other learning settings with multiple heterogeneous sources of data such as domain adaptation and multi-task learning to entail optimal statistical accuracy. While PERM overcomes the data heterogeneity issue, the number of optimization problems (i.e., distinct personalized ERM) to be solved scales linearly with the number of data sources. To simultaneously optimize all objectives in a scalable and computationally efficient manner, we propose a novel idea that replaces the standard *model averaging* in distributed learning with *model shuffling* and establish its convergence rate. This also allows us to learn distinct model architectures (e.g., neural networks with different number of parameters) for different clients, thus confining to underlying memory and compute resources of individual clients, and overcoming the system heterogeneity issue. This addresses an open question in [28] where only a single global model can be trained in a model-heterogeneous setting, while PERM allows deploying distinct models for different clients. We empirically evaluate the performance of PERM, which corroborates the statistical benefits of PERM in comparison to existing methods.

2 Personalized Empirical Risk Minimization

In this section, we formally state the problem and introduce PERM as an ideal paradigm for learning from heterogeneous data sources. We assume there are N distributed devices where each holds a distinct data shard $\mathcal{S}_i = \{(\mathbf{x}_{i,j}, y_{i,j})\}_{j=1}^{n_i}$ with n_i training samples that are realized by a local source distribution $\mathcal{D}_i$ over instance space $\Xi = \mathcal{X} \times \mathcal{Y}$. The data distributions across the devices are not independently and identically distributed (non-IID or *heterogeneous*), i.e., $\mathcal{D}_1 \neq \mathcal{D}_2 \neq$

$\dots \neq \mathcal{D}_N$, and each distribution corresponds to a local *generalization error* or *true risk* $\mathcal{L}_i(h) = \mathbb{E}_{(\mathbf{x}, y) \sim \mathcal{D}_i} [\ell(h(\mathbf{x}), y)]$, $i = 1, 2, \dots, N$ on *unseen* samples for any model $h \in \mathcal{H}$, where $\mathcal{H}$ is the hypothesis set (e.g., a linear model or a deep neural network) and $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}^+$ is a given convex or non-convex loss function. We use $\hat{\mathcal{L}}_i(h) = (1/n_i) \sum_{(\mathbf{x}, y) \in \mathcal{S}_i} \ell(h(\mathbf{x}), y)$ to denote the *local empirical risk* or training loss at i th data shard $\mathcal{S}_i$ with n_i samples.

We seek to collaboratively learn a model or personalized models that entail a good generalization on all local distributions, i.e. minimizing true risk $\mathcal{L}_i(\cdot)$, $i = 1, \dots, N$ for all data sources (all-for-all [33]). A simple non-personalized solution, particularly in FL, aims to minimize a (weighted) *empirical risk* over all data shards in a communication-efficient manner [34]:

$$\arg \min_{h \in \mathcal{H}} \sum_{i=1}^N p(i) \hat{\mathcal{L}}_i(h) \text{ with } \mathbf{p} \in \Delta_N, \quad (\text{WERM})$$

where $\Delta_N = \{\mathbf{p} \in \mathbb{R}_+^N \mid \sum_{i=1}^N p(i) = 1\}$ denotes the simplex set.

It has been shown that a *single model* learned by WERM, for example by using fixed mixing weights $p(i) = n_i/n$, where n is total number of training samples, or even agnostic to mixture of distributions [35, 36], while yielding a good performance on the *combined* datasets of all devices, can suffer from a poor generalization error on individual datasets by increasing the diversity among distributions [37, 38, 39, 40]. To overcome this issue, there has been a surge of interest in developing methods that personalize the global model to individual local distributions. These methods can be unified as the following bi-level problem (a similar unification has been made in [12]):

$$\arg \min_{h_1, h_2, \dots, h_m \in \mathcal{H}} \hat{\mathcal{F}}_i(h_i \oplus h_*) \text{ subject to } h_* = \arg \min_{h \in \mathcal{H}} \sum_{j=1}^N \alpha(j) \hat{\mathcal{L}}_j(h) \quad (\text{BERM})$$

where $\oplus$ denotes the mixing operation to combine local and global models, and $\hat{\mathcal{F}}_i$ is a modified local loss which is not necessarily same as local risk $\hat{\mathcal{L}}_i$. By carefully designing the local loss $\hat{\mathcal{F}}_i$ and mixing operation $\oplus$, we can develop different penalization schemes for FL including existing methods such as linearly interpolating global and local models [11, 2], multi-task learning [10] and meta-learning [9] as special cases. For example, BERM reduces to *zero-personalization* objective WERM when $h_i \oplus h_* = h_*$, and $\hat{\mathcal{F}}_i = \hat{\mathcal{L}}_i$. At the other end of the spectrum lies the *zero-collaboration* where the i th client trains its own model without any influence from other clients by setting $h_i \oplus h_* = h_i$, $\hat{\mathcal{F}}_i = \hat{\mathcal{L}}_i$. The personalized model with *interpolation* of global and local models can be recovered by setting $h_i \oplus h_* = \alpha h_i + (1 - \alpha) h_*$, and $\hat{\mathcal{F}}_i = \hat{\mathcal{L}}_i$. While more effective than a single global model learned via WERM, personalization methods suffer from three key issues: i) the global model is still obtained by minimizing the average empirical loss which might limit the statistical benefits of collaboration, ii) overall model complexity increases linearly with number of clients, and iii) a same model space is shared across servers and clients.

To motivate our proposal, let us consider the empirical loss $\sum_{i=1}^N \alpha(i) \hat{\mathcal{L}}_i(h)$ in WERM (or the inner level objective in BERM) with fixed mixing weights $\alpha \in \Delta_N$, and denote the optimal solution by $\hat{h}_\alpha$. The excess risk of the learned model $\hat{h}_\alpha$ on i th local distribution $\mathcal{D}_i$ w.r.t. the optimal local model $h_i^* = \arg \min_{h \in \mathcal{H}} \mathcal{L}_{\mathcal{D}_i}(h)$ (i.e. *all-for-one*) can be bounded by (informal) [31]

$$\mathcal{L}_i(\hat{h}_\alpha) \leq \mathcal{L}_i(h_i^*) + \sum_{j=1}^N \alpha(j) R_j(\mathcal{H}) + 2 \sum_{j=1}^N \alpha(j) \text{disc}_{\mathcal{H}}(\mathcal{D}_j, \mathcal{D}_i) + C \sqrt{\sum_{j=1}^N \frac{\alpha(j)^2}{n_j}} \quad (\text{GEN})$$

where $R_j(\mathcal{H})$ is the empirical Rademacher complexity $\mathcal{H}$ w.r.t. $\mathcal{S}_j$, and $\text{disc}_{\mathcal{H}}(\mathcal{D}_i, \mathcal{D}_j)$ is a pseudo-distance on the set of probability measures on Ξ to assess the discrepancy between the distributions $\mathcal{D}_i$ and $\mathcal{D}_j$ with respect to the hypothesis class $\mathcal{H}$ as defined below [29]:

Definition 1. For a model space $\mathcal{H}$ and $\mathcal{D}, \mathcal{D}'$ two probability distributions on $\Xi = \mathcal{X} \times \mathcal{Y}$,

$$\text{disc}_{\mathcal{H}}(\mathcal{D}, \mathcal{D}') = \sup_{h \in \mathcal{H}} |\mathbb{E}_{\xi \sim \mathcal{D}}(\ell(h, \xi)) - \mathbb{E}_{\xi' \sim \mathcal{D}'}(\ell(h, \xi'))|$$

Intuitively, the discrepancy between the two distributions is large, if there exists a predictor that performs well on one of them and badly on the other. On the other hand, if all functions in the hypothesis class perform similarly on both, then $\mathcal{D}$ and $\mathcal{D}'$ have low discrepancy. The above metric which is a special case of a popular family of distance measures in probability theory and mathematical

statistics known as integral probability metrics (IPMs) [41], can be estimated from finite data by replacing the expected losses with their empirical counterparts (i.e. $\mathcal{L}_i$ with $\hat{\mathcal{L}}_i$).

From GEN, it can be observed that a mismatch between pairs of distributions limits the benefits of ERM on all distributions. Indeed, the generalization risk w.r.t. $\mathcal{D}_i$ will significantly increase when the distribution divergence terms $\text{disc}_{\mathcal{H}}(\mathcal{D}_j, \mathcal{D}_i)$ are large. It leads to an ideal sample complexity $1/\sqrt{n}$ where $n = n_1 + n_2 + \dots + n_N$ is the total number of samples, which could have been obtained in the IID setting with $\alpha(j) = 1/N$ when the divergence is small as the pairwise discrepancies disappear. Also, we note that even if the global model achieves a small training error over the union of all data (e.g., over parametrized setting) and can entail a good generalization error with respect to *average distribution*, the divergence term still remains which illustrates the poor performance of the global model on all local distributions $\mathcal{D}_i, i = 1, 2, \dots, N$. This implies that even personalization of the global model as in BERM can not entail a good generalization on all local distributions as there is no effective transfer of positive knowledge among data sources in the presence of high data heterogeneity among local distributions (similar impossibility results even under seemingly generous assumptions on how distributions relate have been made in multisource domain adaptation as well [42]).

Interestingly the bound suggests that seeking optimal accuracy on *all* local distributions requires choosing a distinct mixing of local losses for each client i that minimizes the right-hand side of GEN. This indicates that in an ideal setting (i.e. *all-for-all*), we can achieve the best accuracy for each local distribution $\mathcal{D}_i$ by *personalizing* the WERM, i.e., (i) first estimating $\alpha_i, i = 1, 2, \dots, N$ for each client individually, then (ii) solving a variant of WERM for each client with obtained mixing parameters:

$$\arg \min_{h \in \mathcal{H}_i} \sum_{j=1}^N \alpha_i(j) \hat{\mathcal{L}}_j(h) \quad \text{for } i = 1, 2, \dots, N. \quad (\text{PERM})$$

By doing this each device achieves the optimal local generalization error by **learning who to learn with** based on the number of samples at each source and the mismatch between its data distribution with other clients. We also note that compared to WERM and BERM, in PERM since we solve a different aggregated empirical loss for each client, we can pick a different model space/model architecture $\mathcal{H}_i$ for each client to meet its available computational resources.

While this two-stage method is guaranteed to entail optimal test accuracy for all local distributions $\mathcal{D}_i$, however, making it scalable requires overcoming two issues. First, estimating the statistical discrepancies between each pair of data sources (i.e., $\alpha_i, i = 1, \dots, N$) is a computing burden as it requires solving $O(N^2)$ difference of (non)-convex functions in a distributed manner and requires enough samples from each source to entail good accuracy on estimating pairwise discrepancies [41]. Second, we need to solve N variants of the optimization problem in PERM, possibly each with a different model space, which is infeasible when the number of devices is huge (e.g., cross-device federated learning). In the next section, we propose a simple yet effective idea to overcome these issues in a computationally efficient manner.

3 PERM at Scale via Model Shuffling

In this section, we propose a method to efficiently estimate the empirical discrepancies among data sources followed by a model shuffling idea to simultaneously solve N versions of PERM to learn a personalized model for each client. We first start by proposing a two-stage algorithm: estimating mixing parameters followed by model shuffling. Then, we propose a single loop unified algorithm that enjoys the same computation and communication overhead as BERM (twice communication of FedAvg). For ease of exposition, we discuss the proposed algorithms by assuming all the clients share the same model architecture and later on discuss the generalization to heterogeneous model spaces. Specifically, we assume that the model space $\mathcal{H}$ is a parameterized by a convex set $\mathcal{W} \subseteq \mathbb{R}^d$ and use $f_i(\mathbf{w}) := \hat{\mathcal{L}}_i(\mathbf{w}) = \sum_{(\mathbf{x}, y) \in \mathcal{S}_i} \ell(\mathbf{w}; (\mathbf{x}, y))$ to denote the empirical loss at i th data shard.

3.1 Warmup: a two-stage algorithm

We start by proposing a two-stage method for solving N variants of PERM in parallel. In the first stage, we propose an efficient method to learn the mixing parameters for all clients. Then, in stage two, we propose a model shuffling method to solve all personalized empirical losses in parallel.

Stage 1: Mixing parameters estimation. In the first stage we aim to efficiently estimate the pairwise discrepancy among local distributions to construct mixing parameters $\alpha_i, i = 1, 2, \dots, N$. From

generalization bound GEN and Definition 1, a direct solution to estimate α_i is to solve the following convex-nonconcave minimax problem for each client:

$$\alpha_i^* = \arg \min_{\alpha \in \Delta_N} \sum_{j=1}^N \alpha(j) \max_{\mathbf{w} \in \mathcal{W}} |f_i(\mathbf{w}) - f_j(\mathbf{w})| + \sum_{j=1}^N \alpha(j)^2 / n_j \quad (1)$$

where we estimate the true risks in pairwise discrepancy terms with their empirical counterparts and drop the complexity term as it becomes identical for all sources by fixing the hypothesis space $\mathcal{H}$ and bounding it with a computable distribution-independent quantity such as VC dimension [43], or it can be controlled by choice of $\mathcal{H}$ or through data-dependent regularization. However, solving the above minimax problem itself is already challenging: the inner maximization loop is a nonconcave (or difference of convex) problem, so most of the existing minimax algorithms will fail on this problem. To our best knowledge, the only provable deterministic algorithm is [44], and it is hard to generalize it to stochastic and distributed fashion. Moreover, since we have N clients, we need to solve N variants of (1), which makes designing a scalable algorithm even harder.

To overcome aforementioned issues, we make two relaxations to estimate the per client mixing parameters. First, we optimize an upper bound of pairwise empirical discrepancies $\sup_{\mathbf{w}} |f_i(\mathbf{w}) - f_j(\mathbf{w})|$ in terms of gradient dissimilarity between local objectives $\|\nabla f_i(\mathbf{w}) - \nabla f_j(\mathbf{w})\|$ [45], which quantifies how different the local empirical losses are and widely used in analysis of learning from heterogeneous losses as in FL [46]. Second, given that the discrepancy measure based on the supremum could be excessively pessimistic in real-world scenarios, and drawing inspiration from the concept of average drift at the optimal point as a right metric to measure the effect of data heterogeneity in federated learning [47], we propose to measure discrepancy at the optimal solution obtained by solving WERM, i.e., $\mathbf{w}^* := \arg \min_{\mathbf{w} \in \mathcal{W}} (1/N) \sum_{i=1}^N f_i(\mathbf{w})$. By doing this, the problem reduces to a simple minimization for each client, given the optimal global solution. These two relaxations lead to solving the following tractable optimization problem to decide the per-client mixing parameters:

$$\arg \min_{\alpha \in \Delta_N} g_i(\mathbf{w}^*, \alpha) := \sum_{j=1}^N \alpha(j) \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + \lambda \sum_{j=1}^N \alpha(j)^2 / n_j \quad (2)$$

where we added a regularization parameter λ and used the squared of gradient dissimilarity for computational convenience. Thus, obtaining all N mixing parameters requires solving a single ERM to obtain optimal global solution and N variants of (2). To get the optimal solution in a communication-reduced manner, we adapt the Local SGD algorithm [48] (or FedAvg [14]) and find the optimal solution in intermittent communication setting [49] where the clients work in parallel and are allowed to make K stochastic updates between two communication rounds for R consecutive rounds. The detailed steps are given in Algorithm A1 in Appendix B for completeness. After obtaining the global model $\mathbf{w}^R$ we optimize over α in $g_i(\mathbf{w}^R, \alpha)$ using T_α iterations of GD to get $\hat{\alpha}_i$. Actually, we will show that as long as $\mathbf{w}^R$ converge to $\mathbf{w}^*$, $\hat{\alpha}_i, i = 1, \dots, N$ converges to solution of (2) very fast. Our proof idea is based on the following Lipschitzness observation:

$$\|\alpha_{g_i}^*(\mathbf{w}^R) - \alpha_{g_i}^*(\mathbf{w}^*)\|^2 \leq 4L^2 \kappa_g^2 \sum_{j=1}^N \left(2 \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + 4L^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) \|\mathbf{w}^R - \mathbf{w}^*\|^2$$

where $\alpha_{g_i}^*(\mathbf{w}) := \arg \min_{\alpha \in \Delta_N} g_i(\mathbf{w}, \alpha)$ and $\kappa_g := n_{\max} / (2\lambda)$ is the condition number of $g_i(\mathbf{w}, \cdot)$ where $n_{\max} = \max_{i \in [N]} n_i$. The Lipschitz constant mainly depends on *gradient dissimilarity at optimum*. As $\mathbf{w}^R$ tends to $\mathbf{w}^*$, the $\alpha_{g_i}^*(\cdot)$ becomes more Lipschitz continuous, i.e., the coefficient in front of $\|\mathbf{w}^R - \mathbf{w}^*\|^2$ getting smaller, thus leading to more accurate mixing parameters.

To establish the convergence, we make the following standard assumptions.

Assumption 1 (Smoothness and strong convexity). *We assume $f_i(\mathbf{x})$'s are L -smooth and μ -strongly convex, i.e.,*

$$\begin{aligned} \forall \mathbf{x}, \mathbf{y} : \|\nabla f_i(\mathbf{x}) - \nabla f_i(\mathbf{y})\| &\leq L \|\mathbf{x} - \mathbf{y}\|. \\ \forall \mathbf{x}, \mathbf{y} : f_i(\mathbf{y}) &\geq f_i(\mathbf{x}) + \langle \nabla f_i(\mathbf{x}), \mathbf{y} - \mathbf{x} \rangle + \frac{1}{2} \mu \|\mathbf{y} - \mathbf{x}\|^2 \end{aligned}$$

We denote the condition number by $\kappa = L/\mu$.

Assumption 2 (Bounded variance). *The variance of stochastic gradients computed at each local function is bounded, i.e., $\forall i \in [N], \forall \mathbf{w} \in \mathcal{W}, \mathbb{E}[\|\nabla f_i(\mathbf{w}; \xi) - \nabla f_i(\mathbf{w})\|^2] \leq \delta^2$.*

Assumption 3 (Bounded domain). *The domain $\mathcal{W} \subset \mathbb{R}^d$ is a bounded convex set, with diameter D under ℓ_2 metric, i.e., $\forall \mathbf{w}, \mathbf{w}' \in \mathcal{W}, \|\mathbf{w} - \mathbf{w}'\| \leq D$.*

Definition 2 (Gradient dissimilarity). *We define the following quantities to measure the gradient dissimilarity among local functions:*

$$\zeta_{i,j}(\mathbf{w}) := \|\nabla f_i(\mathbf{w}) - \nabla f_j(\mathbf{w})\|^2, \quad \bar{\zeta}_i(\mathbf{w}) := \frac{1}{N} \sum_{j=1}^N \zeta_{i,j}(\mathbf{w}),$$

$$\zeta := \sup_{\mathbf{w} \in \mathcal{W}} \max_{i \in [N]} \|\nabla f_i(\mathbf{w}) - (1/N) \sum_{j=1}^N \nabla f_j(\mathbf{w})\|^2.$$

The following theorem gives the convergence rate of estimated discrepancies to optimal counterparts.

Theorem 1. *Under Assumptions 1-3, if we run Algorithm A1 on $F(\mathbf{w}) := \frac{1}{N} \sum_{j=1}^N f_j(\mathbf{w})$ with $\gamma = \Theta\left(\frac{\log(RK)}{\mu RK}\right)$ for R rounds with synchronization gap K , for $\kappa_g = 1/(\lambda n_{\min})$, it holds that*

$$\mathbb{E}\|\alpha_i^R - \alpha_i^*\|^2 \leq \tilde{O}\left(\exp\left(-\frac{T\alpha}{\kappa_g}\right) + \kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \left(\frac{D^2}{RK} + \frac{\kappa \zeta^2}{\mu^2 R^2} + \frac{\delta^2}{\mu^2 N R K}\right)\right) \quad \forall i \in [N].$$

An immediate implication of Theorem 1 is that even we solve (2) at $\mathbf{w}^R$, the algorithm will eventually converge to optimal solution of (2) at $\mathbf{w}^*$. The core technique in the proof, as we mentioned, is to show that for a parameter within a small region centered at $\mathbf{w}^*$, the function $\alpha_{g_i}^*(\mathbf{w})$ becomes ‘more Lipschitz’. The rigorous characterization of this property is captured by Lemma 3 in appendix.

Stage 2: Scalable personalized optimization with model shuffling. After obtaining the per client mixing parameters, in the second stage we aim at solving N different personalized variants of PERM denoted by $\Phi(\hat{\alpha}_1, \mathbf{v}), \Phi(\hat{\alpha}_2, \mathbf{v}), \dots, \Phi(\hat{\alpha}_N, \mathbf{v})$ to learn local models where

$$\min_{\mathbf{v} \in \mathcal{W}} \Phi(\hat{\alpha}_i, \mathbf{v}) := \frac{1}{N} \sum_{j=1}^N \hat{\alpha}_i(j) f_j(\mathbf{v}). \quad (3)$$

Here we devise an iterative algorithm based on distributed SGD with periodic averaging (a.k.a. Local SGD [48]) to solve these N optimization problems in parallel with *no extra overhead*. The idea is to replace the model averaging in vanilla distributed (Local) SGD with *model shuffling*. Specifically, as shown in Algorithm 1 the algorithm proceeds for R epochs where each epoch runs for N communication rounds. At the beginning of each epoch r the server generates a random permutation σ_r over N clients. At each communication round j within the epoch, the server sends the model of client i to client $i_j = (i + j) \bmod N$ in the permutation σ_r along with $\alpha_i(i_j)$. After receiving a model from the server, the client updates the received model for K local steps and returns it back to the server. As it can be seen, the updates of each loss $\Phi(\hat{\alpha}_i, \mathbf{v}), i = 1, 2, \dots, N$ during an epoch is equivalent to sequentially processing individual losses in (3) which can be considered as permutation-based SGD but with the different that each component now is updated for K steps. By *interleaving the permutations*, we are able to simultaneously optimize all N objectives. We note that the computation and communication complexity of the proposed algorithm is the same as Local SGD with two differences: the model averaging is replaced with model shuffling, and the algorithms run over a permutation of devices. The convergence rate of Local SGD is well-established in literature [50, 51, 52, 53, 54], but here we establish the convergence of permutation-based variant which is interesting by its own.

Assumption 4 (Bounded Gradient). *The variance of stochastic gradients computed at each local function is bounded, i.e., $\forall i \in [N], \sup_{\mathbf{v} \in \mathcal{W}} \|\nabla f_i(\mathbf{v})\| \leq G$.*

We note that the Assumption 4 can be realized since we work with a bounded domain $\mathcal{W}$.

Theorem 2. *Let Assumptions 1-4 hold. Assume α_i^* is the solution of (2). Then if we run Algorithm 1 on the $\hat{\alpha}_i$ obtained from Algorithm A1, then Algorithm 1 with $\eta = \Theta\left(\frac{\log(NKR^3)}{\mu R}\right)$ will output the solution $\hat{\mathbf{v}}_i, \forall i \in [N]$, such that with probability at least $1 - p$, the following statement holds:*

$$\mathbb{E}[\Phi(\alpha_i^*, \hat{\mathbf{v}}_i) - \Phi(\alpha_i^*, \mathbf{v}^*(\alpha_i^*))] \leq \tilde{O}\left(\frac{D^2 L}{N K R^2}\right) + \frac{L \delta^2}{\mu^2 R} + \left(\frac{L^4 + N}{\mu^4 R^2}\right) L G^2 N \log(1/p)$$

$$+ \kappa_{\Phi}^2 L \tilde{O}\left(\exp\left(-\frac{T\alpha}{\kappa_g}\right) + \kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \left(\frac{\kappa \zeta^2}{\mu^2 R^2} + \frac{\delta^2}{\mu^2 N R K}\right)\right),$$

Algorithm 1: Shuffling Local SGD

Input: Clients $1, \dots, N$, Number of Local Steps K , Number of Epoch R , mixing parameter

$\hat{\alpha}_1, \dots, \hat{\alpha}_N$

Epoch for $r = 0, \dots, R - 1$ **do**

Server generates permutation $\sigma_r : [N] \mapsto [N]$.

parallel for $i = 1, \dots, N$ **do**

Client i sets initial model $\mathbf{v}_i^{r,0} = \mathbf{v}_i^r$.

for $j = 1, \dots, N$ **do**

Set indices $i_j = \sigma_r((i + j) \bmod N)$.

Server sends $\mathbf{v}_i^{r,j}$ to Client i_j .

$\mathbf{v}_i^{r,j+1} = \text{SGD-Update}(\mathbf{v}_i^{r,j}, \eta, i_j, K, \hat{\alpha}_i)$.

Client i does projection: $\mathbf{v}_i^{r+1} = \mathcal{P}_{\mathcal{W}}(\mathbf{v}_i^{r,N})$.

Output: $\hat{\mathbf{v}}_i = \mathcal{P}_{\mathcal{W}}(\mathbf{v}_i^R - (1/L)\nabla_{\mathbf{v}}\Phi(\hat{\alpha}_i, \mathbf{v}_i^R)), \forall i \in [N]$.

SGD-Update($\mathbf{v}, \eta, j, K, \alpha$)

Initialize $\mathbf{v}^0 = \mathbf{v}$

for $t = 0, \dots, K - 1$ **do**

$\mathbf{v}^t = \mathbf{v}^{t-1} - \eta\alpha(j)N\nabla f_j(\mathbf{v}^{t-1}; \xi_j^{t-1})$

Output: $\mathbf{v}^K$

where $\kappa = \frac{L}{\mu}$, $\kappa_g = \frac{n_{\max}}{2\lambda}$ and $\kappa_{\Phi} = \frac{\sqrt{NG}}{\mu}$, and the expectation is taken over randomness of Algorithm A1. That is, to guarantee $\mathbb{E}[\Phi(\alpha_i^*, \hat{\mathbf{v}}_i) - \Phi(\alpha_i^*, \mathbf{v}_i^*)] \leq \epsilon$, we choose $R = O\left(\max\left\{\frac{L\delta^2}{\mu\epsilon}, \frac{\kappa_{\Phi}^2\kappa_g^2L^3\zeta_i(\mathbf{w}^*)D^2}{\epsilon}\right\}\right)$ and $T_{\alpha} = O\left(\kappa_g \log\left(\frac{L\kappa_{\Phi}^2}{\epsilon}\right)\right)$.

The above theorem shows that even though we run the optimization on $\Phi(\hat{\alpha}_i, \mathbf{v})$, our obtained model $\hat{\mathbf{v}}_i$ will still converge to the optimal solution of $\Phi(\alpha_i^*, \mathbf{v})$. The convergence rate is contributed from two parts: convergence of $\hat{\alpha}_i$ (Algorithm A1) and convergence of personalized model $\hat{\mathbf{v}}_i$ (Algorithm 1). Notice that, for the convergence rate of $\hat{\mathbf{v}}_i$, we roughly recover the optimal rate of shuffling SGD [55], which is $O(1/R^2)$. However, we suffer from a $O(\delta^2/R)$ term since each client runs vanilla SGD on their local data (the SGD-Update procedure in Algorithm 1). One medication for this variance term could be deploying variance reduction or shuffling data locally at each client before applying SGD. We notice that there is a recent work [56] also considering the client-level shuffling idea, but our work differs from it in two aspects: 1) they work with local SGD type algorithm and the shuffling idea is employed for model averaging within a subset of clients, while in our algorithm, during each local update period, each client runs shuffling SGD directly on other's model 2) from a theoretical perspective, we are mostly interested in investigating whether the algorithm can converge to the true optimal solution of $\Phi(\alpha_i^*, \mathbf{v})$ if we only optimize on a surrogate function $\Phi(\hat{\alpha}_i, \mathbf{v})$.

One drawback of Algorithm 1 is that we have to wait for Algorithm A1 to finish and output $\hat{\alpha}_i$, so that we can proceed with Algorithm 1. However, if we are not satisfied with the precision of $\hat{\alpha}_i$, we may not have a chance to go back to refine it. Hence in the next subsection, we propose to interleave Algorithm 1 and Algorithm A1, and introduce a single-loop variant of PERM which will jointly optimize mixture weights and learn personalized models in an interleaving fashion.

3.2 A unified single loop algorithm

We now turn to introducing a single-stage algorithm that jointly optimizes α_i s and $\mathbf{v}_i$ s as depicted in Algorithm 2 by intertwining the two stages in Algorithm A1 and Algorithm 1 in a single unified method. The idea is to learn the global model, which is used to estimate mixing parameters, concurrent to personalized models. At each communication round, the clients compute gradients on the global model, on their data, after the server collects these gradients does a step mini-batch SGD update on the global model, and then updates the mixing parameters. Then we proceed to update the personalized models similar to Algorithm 1. We note that, unlike the two-stage method where the mixing parameters are computed at the final global model, here the mixing parameters are updated adaptively based on intermediate global models.

Theorem 3. Let Assumptions 1 to 4 to be satisfied. Assume α_i^* is the solution of (2). Then if we run Algorithm 2 with $\eta = \Theta\left(\frac{\log(NKR^3)}{\mu R}\right)$ and $\gamma = \Theta\left(\frac{\log(NKR^3)}{\mu R}\right)$, it will output the solution $\hat{\mathbf{v}}_i$,

Algorithm 2: Single Loop PERM

Input: Clients $1, \dots, N$, Number of Local Steps K , Number of Epoch R , Initial mixing parameter $\alpha_1^0 = \dots, \alpha_N^0 = \bar{\alpha} = [1/N, \dots, 1/N]$.

Epoch for $r = 0, \dots, R - 1$ **do**

Server generates permutation $\sigma_r : [N] \mapsto [N]$.

parallel for Client $i = 1, \dots, N$ **do**

Client i sets initial model $v_i^{r,0} = v_i^r$.

for $j = 1, \dots, N$ **do**

Set indices $i_j = \sigma_r((i + j) \bmod N)$.

Server sends $v_i^{r,j}$ to client i_j .

$v_i^{r,j+1} = \text{SGD-Update}(v_i^{r,j}, \eta, i_j, K, \alpha_i^r)$. // Personalized model update

Client i does projection: $v_i^{r+1} = \mathcal{P}_{\mathcal{W}}(v_i^{r,N})$.

$w^{r+1} = \mathcal{P}_{\mathcal{W}}(w^r - \gamma \frac{1}{N} \sum_{i=1}^N \frac{1}{M} \sum_{j=1}^M \nabla f_i(w^r, \xi_{i,j}^r))$ // Global model update

Compute α_i^{r+1} by running T_α steps GD on $g_i(w^{r+1}, \alpha)$ // α update

Output: $\hat{v}_i = \mathcal{P}_{\mathcal{W}}(v_i^R - (1/L) \nabla_v \Phi(\alpha_i^R, v_i^R))$, $\hat{\alpha}_i = \alpha_i^R, \forall i \in [N]$.

SGD-Update(v, η, j, K, α)

Initialize $v_j^0 = v$

for $t = 0, \dots, K - 1$ **do**

$v^t = v^{t-1} - \eta \alpha(j) N \nabla f_j(v_j^{t-1}, \xi_j^{t-1})$

Output: v^K

$\forall i \in [N]$, such that with probability at least $1 - p$, the following statement holds:

$$\mathbb{E}[\Phi(\alpha_i^*, \hat{v}_i) - \Phi(\alpha_i^*, v_i^*)] \leq O\left(\frac{LD^2}{NKR^3}\right) + \tilde{O}\left(\left(\frac{\kappa^4 L}{R^2} + \frac{NL}{\mu^2 R^2}\right) G^2 N \log(1/p) + \frac{L\delta^2}{\mu R}\right) \\ + \kappa_\Phi^2 L \tilde{O}\left(\frac{\kappa^2 \kappa_g^2 L^2 \bar{\zeta}_i(w^*) DG}{R} + R^2 \exp\left(-\frac{T_\alpha}{\kappa_g}\right) + \frac{L \kappa^2 \kappa_g^2 \bar{\zeta}_i(w^*) \delta^2}{\mu^2 M}\right),$$

where $\kappa = \frac{L}{\mu}$, $\kappa_g = \frac{n_{\max}}{2\lambda}$, $\kappa_\Phi = \frac{\sqrt{NG}}{\mu}$ and the expectation is taken over the randomness of stochastic samples in Algorithm 2. That is, to guarantee $\mathbb{E}[\Phi(\alpha_i^*, \hat{v}_i) - \Phi(\alpha_i^*, v_i^*)] \leq \epsilon$, we choose $M = O\left(\frac{L^2 \kappa^2 \kappa_g^2 \kappa_\Phi^2 \bar{\zeta}_i(w^*) \delta^2}{\mu^2 \epsilon}\right)$, $R = O\left(\max\left\{\frac{L\delta^2}{\mu \epsilon}, \frac{\kappa_\Phi^2 \kappa^2 \kappa_g^2 L^3 \bar{\zeta}_i(w^*) DG}{\epsilon}\right\}\right)$ and $T_\alpha = O\left(\kappa_g \log\left(\frac{LR^2}{\epsilon}\right)\right)$.

Compared to Theorem 2, we achieve a slightly worse rate, since we need a large mini-batch when we update global model w . However, the advantages of the single-loop algorithm are two-fold. First, as we mentioned in the previous subsection, we have the freedom to optimize $\hat{\alpha}_i$ to arbitrary accuracy, while in double loop algorithm (Algorithm A1 + Algorithm 1), once we get $\hat{\alpha}_i$, we do not have the chance to further refine it. Second, in practice, a single-loop algorithm is often easier to implement and can make better use of caches by operating on data sequentially, leading to improved performance, especially on modern processors with complex memory hierarchies.

3.3 Extension to heterogeneous model setting

In the homogeneous model setting, we assumed a shared model space $\mathcal{W}$ for clients and the server. However, in real-world FL applications, devices have diverse resources and can only train models that match their capacities. We demonstrate that the PERM paradigm can be extended to support learning in model-heterogeneous settings, where different models with varying capacities are used by the server and clients. Focusing on learning the global optimal model to estimate pairwise statistical discrepancies, we note that by utilizing partial training methods [28], where at each communication round a sub-model with a size proportional to resources of each client is sampled from the server's global model (extracted either random, static, or rolling) and is transmitted to be updated locally. Upon receiving updated sub-models, the server can simply aggregate (average) heterogeneous sub-model updates sent from the clients to update the global model. We can consider the complexity of

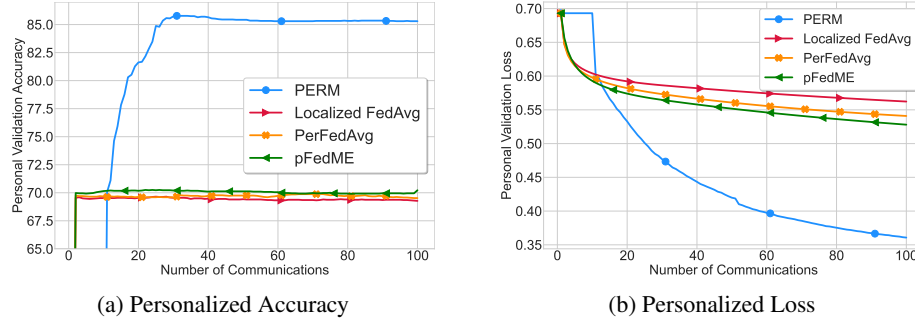


Figure 1: Comparative analysis of personalization methods, including our single-loop PERM algorithm, localized FedAvg, perFedAvg, and pFedME, with synthetic data. The disparity in personalized accuracy and loss highlights PERM’s capability to leverage relevant client correlations.

models used by clients when estimating mixing parameters by solving a modified version of (2) as:

$$\sum_{j=1}^N \alpha(j) \sqrt{\text{VC}(\mathcal{H}_j)/n_j} + \sum_{j=1}^N \alpha(j) \|\nabla f_i(\mathbf{m}_i \odot \mathbf{w}^*) - \nabla f_j(\mathbf{m}_j \odot \mathbf{w}^*)\|^2 + \lambda \sum_{j=1}^N \alpha(j)^2/n_j,$$

where we simply upper bounded the Rademacher complexity w.r.t. each data source in (GEN) with VC dimension [57]. Here $\mathbf{m}_i$ is the masking operator to extract a sub-model of the global model to compute local gradients at client i based on its available resources. By doing so, we can adjust mixing parameters based on the complexity of underlying models, as different sub-models of the global model (i.e., $\mathbf{m}_i \odot \mathbf{w}^*$ versus $\mathbf{m}_j \odot \mathbf{w}^*$) are used to compute drift between pair of gradients at the optimal solution. With regards to training personalized models with heterogeneous local models, as we solve a distinct aggregated empirical loss for each client by interleaving permutations and shuffling models, we can utilize different model spaces $\mathcal{W}_i, i = 1, \dots, N$ for different clients that meet their available resources with aforementioned partial training strategies.

4 Experimental Results

In this section we benchmark the effectiveness of PERM on synthetic data with 50 clients, where it notably outshone other renowned methods as evident in Figure 1. Our experiments concluded with the CIFAR10 dataset, employing a 2-layer convolutional neural network, where PERM, despite a warm-up phase, demonstrated unmatched convergence performance (Figure 2). Additional experiments are reported in the appendix. Across all datasets, the PERM algorithm consistently showcased its robustness and unmatched efficiency in the realm of personalized federated learning.

Experiment on synthetic data. To demonstrate the superior effectiveness of our proposed single-loop PERM algorithm compared to other existing personalization methods, we conducted an experiment using synthetic data generated according to the following specifications. We consider a scenario with a total of N clients, where we draw samples from the distribution $\mathcal{N}(\mu_1, \Sigma_i)$ for half of the clients, denoted by $i \in [1, \frac{N}{2}]$, and from $\mathcal{N}(\mu_2, \Sigma_i)$ for the remaining clients, denoted by $i \in (\frac{N}{2}, N]$. Following the approach outlined in [58], we adopt a uniform variance for all samples, with $\Sigma_{k,k} = k^{-1.2}$. Subsequently, we generate a labeling model using the distribution $\mathcal{N}(\mu_w, \Sigma_w)$.

Given a data sample $\mathbf{x} \in \mathbb{R}^d$, the labels are generated as follows: clients $1, \dots, \frac{N}{2}$ assign labels based on $y = \text{sign}(\mathbf{w}^\top \mathbf{x})$, while clients $\frac{N}{2} + 1, \dots, N$ assign labels based on $y = \text{sign}(-\mathbf{w}^\top \mathbf{x})$. For this specific experiment, we set $\mu_1 = 0.2$, $\mu_2 = -0.2$, and $\mu_w = 0.1$. The data dimension is $d = 60$, and there are 2 classes in the output. We have a total of 50 clients, each generating 500 samples following the aforementioned guidelines. We train a logistic regression model on each client’s data.

To demonstrate the superiority of our PERM algorithm, we conducted a performance comparison against other prominent personalized approaches, including the fine-tuned model of FedAvg [14] (referred to as localized FedAvg), perFedAg [9], and pFedME [7]. The results in Figure 1 highlight PERM’s efficient learning of personalized models for individual clients. In contrast, competing methods relying on globally trained models struggle to match PERM’s effectiveness in highly heterogeneous scenarios, as seen in personalized accuracy and loss. This showcases PERM’s exceptional ability to leverage relevant client learning.

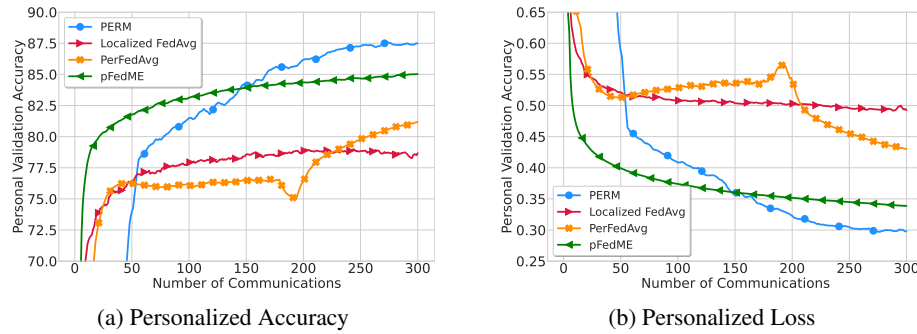


Figure 2: Comparative analysis of our single-loop PERM algorithm, localized FedAvg, pFedMe, and perFedAg, on CIFAR10 dataset and a 2-layer CNN model. Each client has access to only 2 classes of data. PERM rapidly catches up after 10 rounds of warmup without personalization involved.

Experiment on CIFAR10 dataset. We extend our experimentation to the CIFAR10 dataset using a 2-layer convolutional neural network. During this test, 50 clients participate, each limited to data from just 2 classes, resulting in a pronounced heterogeneous data distribution. We benchmark our algorithm against PerFedAvg, pFedMe, and the localized FedAvg. As illustrated in Figure 2, PERM demonstrates superior convergence performance compared to other personalized strategies. It’s noteworthy that PERM’s initial personalized validation is significantly lower than that of approaches like PerFedAvg and pFedMe. This discrepancy stems from our choice to implement 10 communication rounds as a warm-up phase before initiating personalization, whereas other models embark on personalization right from the outset.

Computational overhead. In demonstrating the computational efficiency of the proposed PERM algorithm, we present a comparison of wall-clock time of completing one round of communication of PERM and other methods. Each method undertakes 20 local steps along with their distinct computations for personalization. As depicted in Figure 3, the PERM (single loop) algorithm’s runtime is compared against personalization methods such as PerFedAvg, FedAvg, and pFedMe. Remarkably, PERM maintains a notably minimal computational overhead. The run-time is slightly worse due to overhead of estimating mixing parameters.

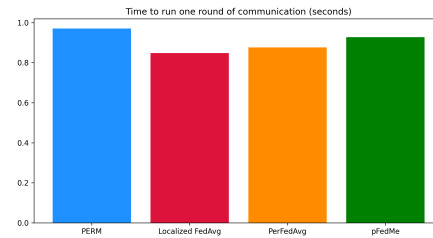


Figure 3: Runtime of different algorithms in a limited environment. We compare PERM (single loop), PerFedAvg, FedAvg, and pFedMe. PERM has a minimal overhead over FedAvg and is comparable to other personalization methods.

5 Discussion & Conclusion

This paper introduces a new *data&system-aware* paradigm for learning from multiple heterogeneous data sources to achieve optimal statistical accuracy across all data distributions without imposing stringent constraints on computational resources shared by participating devices. The proposed PERM schema, though simple, provides an efficient solution to enable each client to learn a personalized model by *learning who to learn with* via personalizing the aggregation of data sources through an efficient empirical statistical discrepancy estimation module. To efficiently solve all aggregated personalized losses, we propose a model shuffling idea to optimize all losses in parallel. PERM can also be employed in other learning settings with multiple sources of data such as domain adaptation and multi-task learning to entail optimal statistical accuracy.

We would like to embark on the scalability of PERM. The compute burden on clients and servers is roughly the same as existing methods thanks to shuffling (except for extra overhead due to estimating mixing parameters which is the same as running FedAvg in a two-stage approach and an extra communication in an interleaved approach). The only hurdle would be the required *memory at server* to maintain mixing parameters, which scales proportionally to the square of the number of clients, which can be alleviated by clustering devices which we leave as a future work.

Acknowledgement

This work was partially supported by NSF CAREER Award #2239374 NSF CNS Award #1956276.

References

- [1] Peter Kairouz, H. Brendan McMahan, Brendan Avent, Aurélien Bellet, Mehdi Bennis, Arjun Nitin Bhagoji, Keith Bonawitz, Zachary Charles, Graham Cormode, Rachel Cummings, Rafael G. L. D'Oliveira, Salim El Rouayheb, David Evans, Josh Gardner, Zachary Garrett, Adrià Gascón, Badi Ghazi, Phillip B. Gibbons, Marco Gruteser, Zaid Harchaoui, Chaoyang He, Lie He, Zhouyuan Huo, Ben Hutchinson, Justin Hsu, Martin Jaggi, Tara Javidi, Gauri Joshi, Mikhail Khodak, Jakub Konečný, Aleksandra Korolova, Farinaz Koushanfar, Sanmi Koyejo, Tancrède Lepoint, Yang Liu, Prateek Mittal, Mehryar Mohri, Richard Nock, Ayfer Özgür, Rasmus Pagh, Mariana Raykova, Hang Qi, Daniel Ramage, Ramesh Raskar, Dawn Song, Weikang Song, Sebastian U. Stich, Ziteng Sun, Ananda Theertha Suresh, Florian Tramèr, Praneeth Vepakomma, Jianyu Wang, Li Xiong, Zheng Xu, Qiang Yang, Felix X. Yu, Han Yu, and Sen Zhao. Advances and open problems in federated learning. *Foundations and Trends® in Machine Learning*, 2021.
- [2] Yishay Mansour, Mehryar Mohri, Jae Ro, and Ananda Theertha Suresh. Three approaches for personalization with applications to federated learning. *arXiv preprint arXiv:2002.10619*, 2020.
- [3] Chengxi Li, Gang Li, and Pramod K Varshney. Federated learning with soft clustering. *IEEE Internet of Things Journal*, 9(10):7773–7782, 2021.
- [4] Avishek Ghosh, Jichan Chung, Dong Yin, and Kannan Ramchandran. An efficient framework for clustered federated learning. *Advances in Neural Information Processing Systems*, 33:19586–19597, 2020.
- [5] Jie Ma, Guodong Long, Tianyi Zhou, Jing Jiang, and Chengqi Zhang. On the convergence of clustered federated learning. *arXiv preprint arXiv:2202.06187*, 2022.
- [6] Hubert Eichner, Tomer Koren, Brendan McMahan, Nathan Srebro, and Kunal Talwar. Semi-cyclic stochastic gradient descent. In *Proceedings of the 36th International Conference on Machine Learning*, PMLR, volume 97, 2019.
- [7] Canh T Dinh, Nguyen Tran, and Tuan Dung Nguyen. Personalized federated learning with moreau envelopes. *Advances in Neural Information Processing Systems*, 33, 2020.
- [8] Yutao Huang, Lingyang Chu, Zirui Zhou, Lanjun Wang, Jiangchuan Liu, Jian Pei, and Yong Zhang. Personalized federated learning: An attentive collaboration approach. *arXiv preprint arXiv:2007.03797*, 2020.
- [9] Alireza Fallah, Aryan Mokhtari, and Asuman Ozdaglar. Personalized federated learning: A meta-learning approach. *Advances in Neural Information Processing Systems*, 2020.
- [10] Virginia Smith, Chao-Kai Chiang, Maziar Sanjabi, and Ameet S Talwalkar. Federated multi-task learning. In *Advances in Neural Information Processing Systems*, pages 4424–4434, 2017.
- [11] Yuyang Deng, Mohammad Mahdi Kamani, and Mehrdad Mahdavi. Adaptive personalized federated learning. *arXiv preprint arXiv:2003.13461*, 2020.
- [12] Filip Hanzely, Boxin Zhao, and Mladen Kolar. Personalized federated learning: A unified framework and universal optimization techniques. *arXiv preprint arXiv:2102.09743*, 2021.
- [13] Tony Lancaster. The incidental parameter problem since 1948. *Journal of econometrics*, 95(2):391–413, 2000.
- [14] Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Agüera y Arcas. Communication-efficient learning of deep networks from decentralized data. In *Artificial intelligence and statistics*, pages 1273–1282. PMLR, 2017.
- [15] Sai Praneeth Karimireddy, Satyen Kale, Mehryar Mohri, Sashank Reddi, Sebastian Stich, and Ananda Theertha Suresh. Scaffold: Stochastic controlled averaging for federated learning. In *International Conference on Machine Learning*, pages 5132–5143. PMLR, 2020.
- [16] Jenny Hamer, Mehryar Mohri, and Ananda Theertha Suresh. Fedboost: A communication-efficient algorithm for federated learning. In *International Conference on Machine Learning*, pages 3973–3983. PMLR, 2020.

- [17] Farzin Haddadpour, Mohammad Mahdi Kamani, Aryan Mokhtari, and Mehrdad Mahdavi. Federated learning with compression: Unified analysis and sharp guarantees. In *International Conference on Artificial Intelligence and Statistics*, pages 2350–2358. PMLR, 2021.
- [18] Yan Sun, Li Shen, Tiansheng Huang, Liang Ding, and Dacheng Tao. Fedsspeed: Larger local interval, less communication round, and higher generalization accuracy. In *The Eleventh International Conference on Learning Representations*, 2023.
- [19] Haibo Yang, Xin Zhang, Prashant Khanduri, and Jia Liu. Anarchic federated learning. In *International Conference on Machine Learning*, pages 25331–25363. PMLR, 2022.
- [20] Shiqiang Wang and Mingyue Ji. A unified analysis of federated learning with arbitrary client participation. *Advances in Neural Information Processing Systems*, 35:19124–19137, 2022.
- [21] Chuhan Wu, Fangzhao Wu, Lingjuan Lyu, Yongfeng Huang, and Xing Xie. Communication-efficient federated learning via knowledge distillation. *Nature communications*, 13(1):1–8, 2022.
- [22] Chaoyang He, Murali Annavaram, and Salman Avestimehr. Group knowledge transfer: Federated learning of large cnns at the edge. *Advances in Neural Information Processing Systems*, 33:14068–14080, 2020.
- [23] Tao Lin, Lingjing Kong, Sebastian U Stich, and Martin Jaggi. Ensemble distillation for robust model fusion in federated learning. *Advances in Neural Information Processing Systems*, 33:2351–2363, 2020.
- [24] Sohei Itahara, Takayuki Nishio, Yusuke Koda, Masahiro Morikura, and Koji Yamamoto. Distillation-based semi-supervised federated learning for communication-efficient collaborative training with non-iid private data. *IEEE Transactions on Mobile Computing*, 22(1):191–205, 2021.
- [25] Enmao Diao, Jie Ding, and Vahid Tarokh. Heterofl: Computation and communication efficient federated learning for heterogeneous clients. *arXiv preprint arXiv:2010.01264*, 2020.
- [26] Samuel Horvath, Stefanos Laskaridis, Mario Almeida, Ilias Leontiadis, Stylianos Venieris, and Nicholas Lane. Fjord: Fair and accurate federated learning under heterogeneous targets with ordered dropout. *Advances in Neural Information Processing Systems*, 34:12876–12889, 2021.
- [27] Sebastian Caldas, Jakub Konečný, H Brendan McMahan, and Ameet Talwalkar. Expanding the reach of federated learning by reducing client resource requirements. *arXiv preprint arXiv:1812.07210*, 2018.
- [28] Samiul Alam, Luyang Liu, Ming Yan, and Mi Zhang. Fedrolex: Model-heterogeneous federated learning with rolling sub-model extraction. *Advances in Neural Information Processing Systems*, 35:29677–29690, 2022.
- [29] Shai Ben-David, John Blitzer, Koby Crammer, Alex Kulesza, Fernando Pereira, and Jennifer Wortman Vaughan. A theory of learning from different domains. *Machine learning*, 79:151–175, 2010.
- [30] Yishay Mansour and Mariano Schain. Robust domain adaptation. *Annals of Mathematics and Artificial Intelligence*, 71(4):365–380, 2014.
- [31] Nikola Konstantinov and Christoph Lampert. Robust learning from untrusted sources. In *International conference on machine learning*, pages 3488–3498. PMLR, 2019.
- [32] Koby Crammer, Michael Kearns, and Jennifer Wortman. Learning from multiple sources. *Journal of Machine Learning Research*, 9(8), 2008.
- [33] Mathieu Even, Laurent Massoulié, and Kevin Scaman. On sample optimality in personalized collaborative and federated learning. In *NeurIPS 2022-36th Conference on Neural Information Processing System*, 2022.
- [34] Jakub Konečný, H Brendan McMahan, Felix X Yu, Peter Richtárik, Ananda Theertha Suresh, and Dave Bacon. Federated learning: Strategies for improving communication efficiency. *arXiv preprint arXiv:1610.05492*, 2016.
- [35] Mehryar Mohri, Gary Sivek, and Ananda Theertha Suresh. Agnostic federated learning. In *International Conference on Machine Learning*, pages 4615–4625. PMLR, 2019.

- [36] Yuyang Deng, Mohammad Mahdi Kamani, and Mehrdad Mahdavi. Distributionally robust federated averaging. *Advances in neural information processing systems*, 33:15111–15122, 2020.
- [37] Tian Li, Anit Kumar Sahu, Manzil Zaheer, Maziar Sanjabi, Ameet Talwalkar, and Virginia Smith. Feddane: A federated newton-type method. In *2019 53rd Asilomar Conference on Signals, Systems, and Computers*, pages 1227–1231. IEEE, 2019.
- [38] Sai Praneeth Karimireddy, Satyen Kale, Mehryar Mohri, Sashank J Reddi, Sebastian U Stich, and Ananda Theertha Suresh. Scaffold: Stochastic controlled averaging for on-device federated learning. *International Conference on Machine Learning*, 119:5132–5143, 2020.
- [39] Farzin Haddadpour and Mehrdad Mahdavi. On the convergence of local descent methods in federated learning. *arXiv preprint arXiv:1910.14425*, 2019.
- [40] Tao Yu, Eugene Bagdasaryan, and Vitaly Shmatikov. Salvaging federated learning by local adaptation. *arXiv preprint arXiv:2002.04758*, 2020.
- [41] Bharath K Sriperumbudur, Kenji Fukumizu, Arthur Gretton, Bernhard Schölkopf, and Gert RG Lanckriet. On integral probability metrics, ϕ -divergences and binary classification. *arXiv preprint arXiv:0901.2698*, 2009.
- [42] Steve Hanneke and Samory Kpotufe. A no-free-lunch theorem for multitask learning. *The Annals of Statistics*, 50(6):3119–3143, 2022.
- [43] Shai Shalev-Shwartz and Shai Ben-David. *Understanding machine learning: From theory to algorithms*. Cambridge university press, 2014.
- [44] Zi Xu, Huiling Zhang, Yang Xu, and Guanghui Lan. A unified single-loop alternating gradient projection algorithm for nonconvex–concave and convex–nonconcave minimax problems. *Mathematical Programming*, pages 1–72, 2023.
- [45] Yatin Dandi, Luis Barba, and Martin Jaggi. Implicit gradient alignment in distributed and federated learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pages 6454–6462, 2022.
- [46] Yue Zhao, Meng Li, Liangzhen Lai, Naveen Suda, Damon Civin, and Vikas Chandra. Federated learning with non-iid data. *arXiv preprint arXiv:1806.00582*, 2018.
- [47] Jianyu Wang, Rudrajit Das, Gauri Joshi, Satyen Kale, Zheng Xu, and Tong Zhang. On the unreasonable effectiveness of federated averaging with heterogeneous data. *arXiv preprint arXiv:2206.04723*, 2022.
- [48] Sebastian U Stich. Local sgd converges fast and communicates little. In *International Conference on Learning Representations*, 2018.
- [49] Blake E Woodworth, Jialei Wang, Adam Smith, Brendan McMahan, and Nati Srebro. Graph oracle models, lower bounds, and gaps for parallel stochastic optimization. In *Advances in neural information processing systems*, pages 8496–8506, 2018.
- [50] Blake E Woodworth, Kumar Kshitij Patel, and Nati Srebro. Minibatch vs local sgd for heterogeneous distributed learning. *Advances in Neural Information Processing Systems*, 33:6281–6292, 2020.
- [51] Konstantin Mishchenko, Grigory Malinovsky, Sebastian Stich, and Peter Richtárik. Proxskip: Yes! local gradient steps provably lead to communication acceleration! finally! In *International Conference on Machine Learning*, pages 15750–15769. PMLR, 2022.
- [52] Honglin Yuan and Tengyu Ma. Federated accelerated stochastic gradient descent. *Advances in Neural Information Processing Systems*, 33:5332–5344, 2020.
- [53] Farzin Haddadpour, Mohammad Mahdi Kamani, Mehrdad Mahdavi, and Viveck Cadambe. Local sgd with periodic averaging: Tighter analysis and adaptive synchronization. In *Advances in Neural Information Processing Systems*, pages 11080–11092, 2019.
- [54] Eduard Gorbunov, Filip Hanzely, and Peter Richtárik. Local sgd: Unified theory and new efficient methods. In *International Conference on Artificial Intelligence and Statistics*, pages 3556–3564. PMLR, 2021.
- [55] Kwangjun Ahn, Chulhee Yun, and Suvrit Sra. Sgd with shuffling: optimal rates without component convexity and large epoch requirements. *Advances in Neural Information Processing Systems*, 33:17526–17535, 2020.

- [56] Yae Jee Cho, Pranay Sharma, Gauri Joshi, Zheng Xu, Satyen Kale, and Tong Zhang. On the convergence of federated averaging with cyclic client participation. *arXiv preprint arXiv:2302.03109*, 2023.
- [57] Mehryar Mohri, Afshin Rostamizadeh, and Ameet Talwalkar. *Foundations of machine learning*. MIT press, 2018.
- [58] Tian Li, Anit Kumar Sahu, Manzil Zaheer, Maziar Sanjabi, Ameet Talwalkar, and Virginia Smith. Federated optimization in heterogeneous networks. *arXiv preprint arXiv:1812.06127*, 2018.
- [59] Sebastian Caldas, Peter Wu, Tian Li, Jakub Konečný, H Brendan McMahan, Virginia Smith, and Ameet Talwalkar. Leaf: A benchmark for federated settings. *arXiv preprint arXiv:1812.01097*, 2018.
- [60] Tianyi Lin, Chi Jin, and Michael I Jordan. On gradient descent ascent for nonconvex-concave minimax problems. *arXiv preprint arXiv:1906.00331*, 2019.
- [61] Markus Schneider. Probability inequalities for kernel embeddings in sampling without replacement. In *Artificial Intelligence and Statistics*, pages 66–74. PMLR, 2016.

A Additional Experiments

In addition to experiments on synthetic and CIFAR10 datasets reported before, we have also conducted experiments on the EMNIST dataset, highlighting PERM’s capability to derive superior personalized models by tapping into inter-client data similarities. Additionally, further insights emerged from our tests on the MNIST dataset, revealing how PERM’s learned mixture weights adeptly respond to both homogeneous and highly heterogeneous data scenarios.

Experiment on EMNIST dataset In addition to the synthetic and CIFAR10 datasets discussed in the main body, we run experiments on the EMNIST dataset [59], which is naturally distributed in a federated setting. In this case, we chose 50 clients and use a 2-layer MLP model, each with 200 neurons. We compare the PERM algorithm with the localized model in FedAvg and perFedAvg [9]. As it can be seen in Figure 4, PERM can learn a better personalized model by attending to each client’s data according to the similarity of the data distribution between clients. The learned values of α , in Figure 5, show that the clients are learning from each others’ data, and not focused on their own data only. This signifies that the distribution of data among clients in this dataset is not highly heterogeneous. Note that, since we are using a subset of clients in the EMNIST dataset for the training (only 50 clients for 100 rounds of communication), the results would be sub-optimal. Nonetheless, the experiments are designed to show the effectiveness of different algorithms. As it can be concluded, in terms of performance, PERM consistently excels beyond its peers, demonstrating exemplary results on various benchmark datasets.

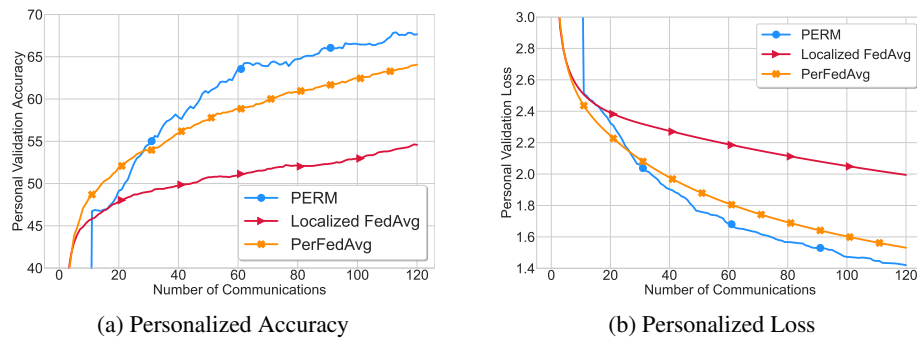


Figure 4: Comparative Analysis of Personalization methods, including our single-loop PERM algorithm, localized FedAvg, and perFedAg, with EMNIST dataset. The disparity in personalized accuracy and loss highlights PERM’s capability in leveraging relevant client correlations.

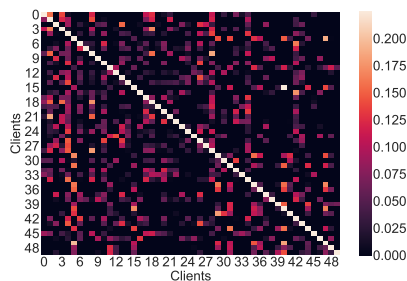


Figure 5: The heat map of the learned α values for the PERM algorithms on the EMNIST dataset with a 2-layer MLP model. The weights signify that clients mutually benefiting from one another’s data, which also highlight that the distribution of data is not significantly heterogeneous in this dataset.

The effectiveness of learned mixture weights To show the effectiveness of the two-stage PERM algorithm, as well as the effects of heterogeneity on the distribution of data among clients on the learned weights α in the algorithm, we run this algorithm on MNIST dataset. We use 50 clients, and the model is an MLP, similar to the EMNIST experiment. In this case, we consider two cases: distributing the data randomly across clients (homogeneous) and only allocating 1 class per client (highly heterogeneous). As it can be seen from Figure 6, when the data distribution is homogeneous the learned values of α as diffused across clients. However, when the data is highly heterogeneous, the learned α values will be highly sparse, indicating that each client is mostly learning from its own data and some other clients with partial distribution similarity. Notably, the matrix predominantly exhibits sparsity, indicating that each client selectively leverages information solely from a subset of other clients. This discernible pattern reinforces the inherent confidence that each client is effectively learning from a limited but strategically chosen group of clients.

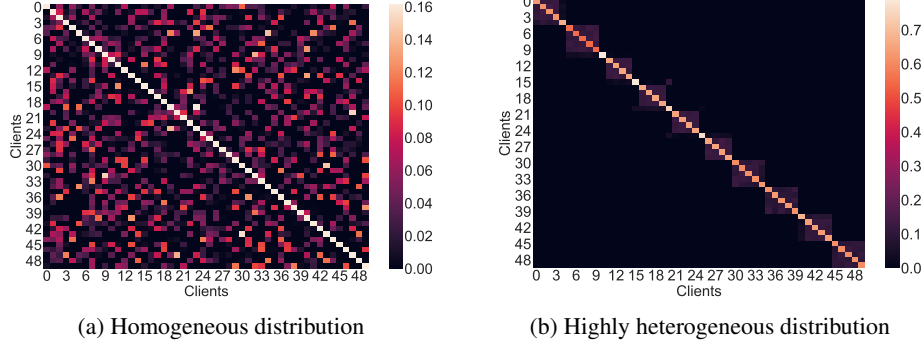


Figure 6: Comparing the performance of two-stage PERM algorithm in learning α values on heterogeneous and homogeneous data distributions. We use MNIST dataset across 50 clients with homogeneous and heterogeneous distributions.

B Proof of Two Stages Algorithm

In this section we provide the proof of convergence of two-stage implementation of PERM (computing mixing parameters followed by learning personalized models via model shuffling using permutation-based variant of distributed SGD with periodic communication).

B.1 Technical Lemmas

Lemma 1. Define $\mathbf{v}^*(\alpha) := \arg \min_{\mathbf{v} \in \mathcal{W}} \Phi(\alpha, \mathbf{v})$, and assume $\Phi(\alpha, \cdot)$ is μ -strongly convex and $\nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v})$ is L Lipschitz in α . Then, $\mathbf{v}^*(\cdot)$ is κ -Lipschitz where $\kappa = L/\mu$.

Proof. The proof is similar to Lin et al's result on minimax objective [60]. First, according to optimality conditions we have:

$$\begin{aligned} \langle \mathbf{v} - \mathbf{v}^*(\alpha), \nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v}^*(\alpha)) \rangle &\geq 0, \\ \langle \mathbf{v} - \mathbf{v}^*(\alpha'), \nabla_{\mathbf{v}} \Phi(\alpha', \mathbf{v}^*(\alpha')) \rangle &\geq 0 \end{aligned}$$

Substituting $\mathbf{v}$ with $\mathbf{v}^*(\alpha')$ and $\mathbf{v}^*(\alpha)$ in the above first and second inequalities respectively yields:

$$\begin{aligned} \langle \mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha), \nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v}^*(\alpha)) \rangle &\geq 0, \\ \langle \mathbf{v}^*(\alpha) - \mathbf{v}^*(\alpha'), \nabla_{\mathbf{v}} \Phi(\alpha', \mathbf{v}^*(\alpha')) \rangle &\geq 0 \end{aligned}$$

Adding up the above two inequalities yields:

$$\langle \mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha), \nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v}^*(\alpha)) - \nabla_{\mathbf{v}} \Phi(\alpha', \mathbf{v}^*(\alpha')) \rangle \geq 0, \quad (4)$$

Since $\Phi(\alpha, \cdot)$ is μ strongly convex, we have:

$$\langle \mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha), \nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v}^*(\alpha')) - \nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v}^*(\alpha)) \rangle \geq \mu \|\mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha)\|^2. \quad (5)$$

Adding up (4) and (5) yields:

$$\langle \mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha), \nabla_{\mathbf{v}} \Phi(\alpha, \mathbf{v}^*(\alpha')) - \nabla_{\mathbf{v}} \Phi(\alpha', \mathbf{v}^*(\alpha')) \rangle \geq \mu \|\mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha)\|^2$$

Finally, using L smoothness of Φ will conclude the proof:

$$\begin{aligned} L \|\mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha)\| \|\alpha - \alpha'\| &\geq \mu \|\mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha)\|^2 \\ \iff \kappa \|\alpha - \alpha'\| &\geq \|\mathbf{v}^*(\alpha') - \mathbf{v}^*(\alpha)\| \end{aligned}$$

□

Lemma 2 (Optimality Gap). Let $\Phi(\alpha, \mathbf{v})$ be defined in (3). Let $\hat{\mathbf{v}} = \mathcal{P}_{\mathcal{W}}(\tilde{\mathbf{v}} - \frac{1}{L} \nabla_{\mathbf{v}} \Phi(\hat{\alpha}, \tilde{\mathbf{v}}))$. If we assume each f_i is L -smooth, μ -strongly convex and with gradient bounded by $\bar{G}$, then the following statement holds true:

$$\Phi(\alpha^*, \hat{\mathbf{v}}) - \Phi(\alpha^*, \mathbf{v}^*) \leq 2L \|\tilde{\mathbf{v}} - \mathbf{v}^*(\hat{\alpha})\|^2 + \left(2\kappa_{\Phi}^2 L + \frac{4NG^2}{L} \right) \|\hat{\alpha} - \alpha^*\|^2.$$

where $\kappa_{\Phi} = \frac{\sqrt{NG}}{\mu}$, $\mathbf{v}^* = \arg \min_{\mathbf{v} \in \mathcal{W}} \Phi(\alpha^*, \mathbf{v})$.

Proof. First we show that $\nabla_v \Phi(\alpha, v)$ is $\sqrt{NG}$ Lipschitz in α . To see this:

$$\begin{aligned} \|\nabla_v \Phi(\alpha, v) - \nabla_v \Phi(\alpha', v)\| &= \left\| \sum_{j=1}^N \alpha_i(j) \nabla f_j(v) - \sum_{j=1}^N \alpha'_i(j) \nabla f_j(v) \right\| \\ &\leq \sqrt{NG} \|\alpha_i - \alpha'_i\|. \end{aligned}$$

Hence due to Lemma 1, we know $v^*(\alpha)$ is $\kappa_\Phi := \frac{\sqrt{NG}}{\mu}$ Lipschitz. According to property of projection, we have:

$$\begin{aligned} 0 &\leq \langle v - \hat{v}, L(\hat{v} - \tilde{v}) + \nabla_v \Phi(\hat{\alpha}, \tilde{v}) \rangle \\ &= \underbrace{\langle v - \hat{v}, L(\hat{v} - \tilde{v}) + \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle}_{T_1} + \underbrace{\langle v - \hat{v}, \nabla_v \Phi(\hat{\alpha}, \tilde{v}) - \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle}_{T_2}. \end{aligned}$$

For T_1 , we notice:

$$\begin{aligned} \langle v - \hat{v}, L(\hat{v} - \tilde{v}) + \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle &= L \langle v - \tilde{v}, \hat{v} - \tilde{v} \rangle + \frac{1}{\eta} \langle \tilde{v} - \hat{v}, \hat{v} - \tilde{v} \rangle + \langle v - \hat{v}, \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle \\ &= L \langle v - \tilde{v}, \hat{v} - \tilde{v} \rangle - L \|\tilde{v} - \hat{v}\|^2 + \langle v - \hat{v}, \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle \\ &\leq L(\|v - \tilde{v}\|^2 + \frac{1}{4} \|\hat{v} - \tilde{v}\|^2) - L \|\tilde{v} - \hat{v}\|^2 + \underbrace{\langle v - \hat{v}, \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle}_{\spadesuit} \end{aligned}$$

where at last step we used Young's inequality. To bound $\spadesuit$, we apply the L smoothness and μ strongly convexity of $\Phi(\alpha, \cdot)$:

$$\begin{aligned} \langle v - \hat{v}, \nabla_v \Phi(\alpha_i^*, \tilde{v}_i) \rangle &= \langle v - \tilde{v}_i, \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle + \langle \tilde{v} - \hat{v}, \nabla_v \Phi(\alpha^*, \tilde{v}) \rangle \\ &\leq \Phi(\alpha^*, v) - \Phi(\alpha^*, \tilde{v}) - \frac{\mu}{2} \|\tilde{v} - v\|^2 + \Phi(\alpha^*, \tilde{v}) - \Phi(\alpha^*, \hat{v}) + \frac{L}{2} \|\tilde{v} - \hat{v}\|^2 \\ &\leq \Phi(\alpha^*, v) - \Phi(\alpha^*, \hat{v}) - \frac{\mu}{2} \|\tilde{v} - v\|^2 + \frac{L}{2} \|\tilde{v} - \hat{v}\|^2 \end{aligned}$$

Putting above bound back yields:

$$\left\langle v - \hat{v}, \frac{1}{\eta}(\hat{v} - \tilde{v}) + \nabla_v \Phi(\alpha^*, \tilde{v}) \right\rangle \leq \Phi(\alpha^*, v) - \Phi(\alpha^*, \hat{v}) + \frac{1}{2\eta} \|\tilde{v} - v\|^2 - \left(\frac{3L}{4} - \frac{L}{2} \right) \|\tilde{v} - \hat{v}\|^2$$

Now we switch to bounding T_2 . Applying Cauchy-Schwartz yields:

$$\begin{aligned} \langle v - \hat{v}_i, \nabla_v \Phi(\hat{\alpha}_i, \tilde{v}_i) - \nabla_v \Phi(\alpha_i^*, \tilde{v}_i) \rangle &\leq \frac{L}{4} \|v - \tilde{v}_i\|^2 + \frac{L}{4} \|\tilde{v} - \hat{v}_i\|^2 + \frac{4}{L} \|\nabla_v \Phi(\hat{\alpha}_i, \tilde{v}_i) - \nabla_v \Phi(\alpha_i^*, \tilde{v}_i)\|^2 \\ &\leq \frac{L}{4} \|v - \tilde{v}_i\|^2 + \frac{L}{4} \|\tilde{v} - \hat{v}_i\|^2 + \frac{4NG^2}{L} \|\hat{\alpha}_i - \alpha_i^*\|^2 \end{aligned}$$

where at last step we apply $\sqrt{NG}$ smoothness of $\Phi(\cdot, v)$. Putting pieces together yields:

$$0 \leq \Phi(\alpha_i^*, v) - \Phi(\alpha_i^*, \hat{v}_i) + \frac{L}{2} \|\tilde{v}_i - v\|^2 + \frac{L}{2} \|v - \tilde{v}_i\|^2 + \frac{4NG^2}{L} \|\hat{\alpha}_i - \alpha_i^*\|^2$$

Re-arranging terms and setting $v = v^*(\alpha^*) = \arg \min_{v \in \mathcal{W}} \Phi(\alpha^*, v)$ yields:

$$\Phi(\alpha^*, \hat{v}) - \Phi(\alpha^*, v^*) \leq L \|\tilde{v} - v^*\|^2 + \frac{4NG^2}{L} \|\hat{\alpha} - \alpha^*\|^2.$$

At last, due to the κ_Φ -Lipschitzness property of $v^*(\cdot)$ as shown in Lemma 1, it follows that:

$$\begin{aligned} L \|\tilde{v} - v^*(\alpha^*)\|^2 &\leq L \|\tilde{v} - v^*(\hat{\alpha})\|^2 + L \|v^*(\hat{\alpha}) - v^*(\alpha^*)\|^2 \\ &\leq 2L \|\tilde{v} - v^*(\hat{\alpha})\|^2 + 2\kappa_\Phi^2 L \|\hat{\alpha} - \alpha^*\|^2, \end{aligned}$$

as desired. □

Algorithm A1: Discrepancy Estimation at Optimum

Input: Number of clients N , number of local steps K , number of communications rounds R

for $r = 0, \dots, R - 1$ **do**

parallel for client $i = 0, \dots, N - 1$ **do**

 Client i initializes model $\mathbf{w}_i^{r,0} = \mathbf{w}_i^r$.

for $t = 0, \dots, K - 1$ **do**

$\mathbf{w}_i^{r,t+1} = \mathbf{w}_i^{r,t} - \gamma \nabla f_i(\mathbf{w}_i^{r,t}; \xi_i^{r,t})$ where $\xi_i^{r,t}$ is a mini-batch sampled from $\mathcal{S}_i$.

 Client i sends $\mathbf{w}_i^{r,K}$ to Server.

 Server computes $\mathbf{w}^{r+1} = \mathcal{P}_{\mathcal{W}} \left(\frac{1}{N} \sum_{i=1}^N \mathbf{w}_i^{r,K} \right)$

 Server broadcasts $\mathbf{w}^{r+1}$ to all clients.

Server computes $\hat{\alpha}_i, i = 1, 2, \dots, N$ by running T_{α} steps of GD on $g_i(\mathbf{w}^R, \alpha)$.

Output: $\hat{\alpha}_1, \dots, \hat{\alpha}_N$.

B.2 Proof of Convergence of Theorem 1

In this section we are going to prove the result in Theorem 1. To this end, we need to show that mixing parameters we compute by first learning the global model and then solving the optimization problem in objective (2) (as depicted in Algorithm A1) converges to optimal values. Notice that in Algorithm A1 we do not solve $g_i(\mathbf{w}^*, \alpha)$ directly, but optimize $g_i(\mathbf{w}^R, \alpha)$ on α for T_{α} iterations of GD. Hence, firstly we need to show that optimizing the surrogate function will also guarantee the convergence of output of algorithm $\hat{\alpha}$ to α^* by deriving a property of the objective in (2). Formally the property is captured by the following lemma.

Lemma 3. Let $g(\mathbf{w}, \alpha) := \sum_{j=1}^N \alpha_j \|\nabla f_i(\mathbf{w}) - \nabla f_j(\mathbf{w})\|^2 + \lambda \sum_{j=1}^N \alpha_j^2 / n_j$ and $\alpha_g^*(\mathbf{w}) = \arg \min_{\alpha \in \Delta_N} g(\mathbf{w}, \alpha)$. Let $\mathbf{w}^R$ be the output of Algorithm A1. Then the following statement holds:

$$\|\alpha_g^*(\mathbf{w}^R) - \alpha_g^*(\mathbf{w}^*)\| \leq \kappa_g^2 \sum_{j=1}^N \left(2 \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + 4L^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) 4L \|\mathbf{w}^R - \mathbf{w}^*\|^2$$

where $\kappa_g := \frac{n_{\max}}{2\lambda}$.

Proof. Define function

$$W(\mathbf{z}, \alpha) = \sum_{j=1}^N \alpha_j z_j + \lambda \sum_{j=1}^N \alpha_j^2 / n_j \quad (6)$$

Apparently, $W(\mathbf{z}, \alpha)$ is linear in $\mathbf{z}$ and $2\frac{\lambda}{n_{\max}}$ strongly convex in $\mathbf{z}$. Next we show that $\nabla_{\alpha} W(\mathbf{z}, \alpha)$ is Lipschitz in $\mathbf{w}$. To see this,

$$\begin{aligned} \|\nabla_{\alpha} W(\mathbf{z}, \alpha) - \nabla_{\alpha} W(\mathbf{z}', \alpha)\| &= \|[z_1, \dots, z_N] - [z'_1, \dots, z'_N]\| \\ &\leq \|\mathbf{z} - \mathbf{z}'\|. \end{aligned}$$

Then, according to Proposition 1, $\alpha_W^*(\mathbf{z}) := \arg \min_{\alpha \in \Delta_N} W(\mathbf{z}, \alpha)$ is κ_g lipschitz in $\mathbf{z}$ where $\kappa_g = \frac{n_{\max}}{2\lambda}$, i.e., $\|\alpha_W^*(\mathbf{z}) - \alpha_W^*(\mathbf{z}')\| \leq \kappa_g \|\mathbf{z} - \mathbf{z}'\|$. Now, let us consider the objective (2):

$$g(\mathbf{w}, \alpha) := \sum_{j=1}^N \alpha_j \|\nabla f_i(\mathbf{w}) - \nabla f_j(\mathbf{w})\|^2 + \lambda \sum_{j=1}^N \alpha_j^2 / n_j$$

We define $\alpha_g^*(\mathbf{w}) = \arg \min_{\alpha \in \Delta_N} g(\mathbf{w}, \alpha)$.

We set

$$\begin{aligned} \mathbf{z}^R &= \left[\|\nabla f_i(\mathbf{w}^R) - \nabla f_1(\mathbf{w}^R)\|^2, \dots, \|\nabla f_i(\mathbf{w}^R) - \nabla f_N(\mathbf{w}^R)\|^2 \right], \\ \mathbf{z}^* &= \left[\|\nabla f_i(\mathbf{w}^*) - \nabla f_1(\mathbf{w}^*)\|^2, \dots, \|\nabla f_i(\mathbf{w}^*) - \nabla f_N(\mathbf{w}^*)\|^2 \right]. \end{aligned}$$

Then we know that

$$\|\alpha_g^*(\mathbf{w}^R) - \alpha_g^*(\mathbf{w}^*)\|^2 = \|\alpha_W^*(\mathbf{z}^R) - \alpha_W^*(\mathbf{z}^*)\|^2 \leq \kappa_g^2 \|\mathbf{z}^R - \mathbf{z}^*\|^2 \quad (7)$$

$$\begin{aligned} &\leq \kappa_g^2 \sum_{j=1}^N \left\| \nabla f_i(\mathbf{w}^R) - \nabla f_j(\mathbf{w}^R) \right\|^2 - \left\| \nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*) \right\|^2 \quad (8) \\ &\leq \kappa_g^2 \sum_{j=1}^N \left\| (\nabla f_i(\mathbf{w}^R) - \nabla f_j(\mathbf{w}^R) + \nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)) \right. \\ &\quad \times \left. (\nabla f_i(\mathbf{w}^R) - \nabla f_j(\mathbf{w}^R) - \nabla f_i(\mathbf{w}^*) + \nabla f_j(\mathbf{w}^*)) \right\|^2 \\ &\leq \kappa_g^2 \sum_{j=1}^N \left\| \nabla f_i(\mathbf{w}^R) - \nabla f_j(\mathbf{w}^R) + \nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*) \right\|^2 4L^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \end{aligned}$$

Since $\|\nabla f_i(\mathbf{w}^R) - \nabla f_j(\mathbf{w}^R)\| \leq \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\| + 2L \|\mathbf{w}^R - \mathbf{w}^*\|$, we can conclude that

$$\|\alpha_g^*(\mathbf{w}^R) - \alpha_g^*(\mathbf{w}^*)\| \leq \kappa_g^2 \sum_{j=1}^N \left(2 \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + 4L^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) 4L \|\mathbf{w}^R - \mathbf{w}^*\|^2.$$

□

With above lemma, to show the convergence of $\hat{\alpha}$ to α^* , we do the following decomposition

$$\begin{aligned} \|\hat{\alpha} - \alpha^*\|^2 &\leq 2 \|\hat{\alpha} - \alpha_g^*(\mathbf{w}^R)\|^2 + 2 \|\alpha_g^*(\mathbf{w}^R) - \alpha_g^*(\mathbf{w}^*)\|^2 \\ &\leq 2(1 - \mu\eta_\alpha)^K + 2\kappa_g^2 \sum_{j=1}^N \left(2 \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + 4L^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) 4L \|\mathbf{w}^R - \mathbf{w}^*\|^2. \end{aligned}$$

Now it remains to show the convergence of Local SGD *last iterate* $\mathbf{w}^R$ to optimal solution $\mathbf{w}^*$. By convention, we use $\mathbf{w}^t = \frac{1}{N} \sum_{i=1}^N \mathbf{w}_i^t$ to denote the virtual average iterates.

Lemma 4 (One iteration analysis of Local SGD). *Under the condition of Theorem 1, the following statement holds true for any $t \in [T]$:*

$$\begin{aligned} \mathbb{E} \|\mathbf{w}^{r,t+1} - \mathbf{w}^*\|^2 &\leq (1 - \mu\gamma) \mathbb{E} \|\mathbf{w}^{r,t} - \mathbf{w}^*\|^2 - (2\gamma - 4\gamma^2 L) \mathbb{E} (F(\mathbf{w}^*) - F(\mathbf{w}^{r,t})) \\ &\quad + (\gamma L + 2\gamma^2 L^2) \frac{1}{N} \sum_{i=1}^N \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 + \gamma^2 \frac{\delta^2}{N}. \end{aligned}$$

Proof. According to updating rule in Algorithm A1, we have the following identity:

$$\mathbb{E} \|\mathbf{w}^{r,t+1} - \mathbf{w}^*\|^2 = \mathbb{E} \|\mathbf{w}^{r,t} - \mathbf{w}^*\|^2 - 2\gamma \mathbb{E} \left\langle \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}; \mathbf{z}_i^{r,t}), \mathbf{w}^{r,t} - \mathbf{w}^* \right\rangle + \gamma^2 \mathbb{E} \left\| \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}) \right\|^2 \quad (9)$$

$$\begin{aligned} &= \mathbb{E} \|\mathbf{w}^t - \mathbf{w}^*\|^2 - 2\gamma \underbrace{\mathbb{E} \left\langle \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}), \mathbf{w}^{r,t} - \mathbf{w}^* \right\rangle}_{T_1} \\ &\quad + \gamma^2 \underbrace{\mathbb{E} \left\| \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}; \mathbf{z}_i^{r,t}) \right\|^2}_{T_2} + \gamma^2 \frac{\delta^2}{N}. \quad (10) \end{aligned}$$

For T_1 , since each f_j is L smooth and μ strongly convex, we have:

$$\begin{aligned} -2\gamma \left\langle \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}), \mathbf{w}^t - \mathbf{w}^* \right\rangle &= -2\gamma \left\langle \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^t), \mathbf{w}^t - \mathbf{w}_i^t + \mathbf{w}_i^{r,t} - \mathbf{w}^* \right\rangle \\ &\leq -2\gamma \left\langle \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^t), \mathbf{w}^{r,t} - \mathbf{w}_i^{r,t} + \mathbf{w}_i^{r,t} - \mathbf{w}^* \right\rangle \\ &\leq 2\gamma \frac{1}{N} \sum_{i=1}^N \left(f_i(\mathbf{w}^*) - f_i(\mathbf{w}^{r,t}) - \frac{\mu}{2} \|\mathbf{w}_i^{r,t} - \mathbf{w}^*\|^2 + \frac{L}{2} \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 \right). \end{aligned}$$

Due to Jensen's inequality we know: $-\frac{1}{N} \sum_{i=1}^N \frac{\mu}{2} \|\mathbf{w}_i^t - \mathbf{w}^*\|^2 \leq -\frac{\mu}{2} \|\mathbf{w}^t - \mathbf{w}^*\|^2$. Hence we know:

$$-2\gamma \left\langle \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}), \mathbf{w}^{r,t} - \mathbf{w}^* \right\rangle \leq 2\gamma \left(F(\mathbf{w}^*) - F(\mathbf{w}^{r,t}) - \frac{\mu}{2} \|\mathbf{w}^{r,t} - \mathbf{w}^*\|^2 + \frac{L}{2} \frac{1}{N} \sum_{i=1}^N \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 \right).$$

For T_2 , we have:

$$\begin{aligned} \mathbb{E} \left\| \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}) \right\|^2 &= 2\mathbb{E} \left\| \frac{1}{N} \sum_{i=1}^N \nabla f_i(\mathbf{w}_i^{r,t}) - \nabla F(\mathbf{w}^{r,t}) \right\|^2 + 2\mathbb{E} \|\nabla F(\mathbf{w}^{r,t})\|^2 \\ &\leq 2L^2 \frac{1}{N} \sum_{i=1}^N \mathbb{E} \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 + 4L (F(\mathbf{w}^{r,t}) - F(\mathbf{w}^*)). \end{aligned}$$

Now, plugging T_1 and T_2 back to (10) yields:

$$\begin{aligned} \mathbb{E} \|\mathbf{w}^{r,t+1} - \mathbf{w}^*\|^2 &\leq (1 - \mu\gamma) \mathbb{E} \|\mathbf{w}^{r,t} - \mathbf{w}^*\|^2 - (2\gamma - 4\gamma^2 L) \mathbb{E} (F(\mathbf{w}^*) - F(\mathbf{w}^{r,t})) \\ &\quad + (\gamma L + 2\gamma^2 L^2) \frac{1}{N} \sum_{i=1}^N \mathbb{E} \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 + \gamma^2 \frac{\delta^2}{N}. \end{aligned}$$

□

Lemma 5. [50, Lemma 8] For the iterates $\{\mathbf{w}_i^{r,t}\}$ generated in Algorithm A1, the following statement holds true:

$$\frac{1}{N} \sum_{i=1}^N \mathbb{E} \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 \leq 3K\gamma^2\delta^2 + 6K^2\gamma^2\zeta^2.$$

Lemma 6 (Last iterate convergence of Local SGD). Under the conditions of Theorem 1, the following statement holds true for the iterates in Algorithm A1:

$$\mathbb{E} \|\mathbf{w}^R - \mathbf{w}^*\|^2 \leq (1 - \mu\gamma)^{RK} \mathbb{E} \|\mathbf{w}^0 - \mathbf{w}^*\|^2 + \frac{1}{\mu\gamma} (\gamma L + 2\gamma^2 L^2) (3K\gamma^2\delta^2 + 6K^2\gamma^2\zeta^2) + \frac{\gamma\delta^2}{\mu N}$$

Proof. We first unroll the recursion in Lemma 4 from $t = K$ to 0, within one communication round:

$$\begin{aligned} \mathbb{E} \|\mathbf{w}^{r,K} - \mathbf{w}^*\|^2 &= (1 - \mu\gamma)^K \mathbb{E} \|\mathbf{w}^{r,0} - \mathbf{w}^*\|^2 - \sum_{t=0}^{K-1} (1 - \mu\gamma)^{K-t} (2\gamma - 4\gamma^2 L) \mathbb{E} (F(\mathbf{w}^*) - F(\mathbf{w}^{r,t})) \\ &\quad + \sum_{t=0}^{K-1} (1 - \mu\gamma)^{K-t} (\gamma L + 2\gamma^2 L^2) \frac{1}{N} \sum_{i=1}^N \mathbb{E} \|\mathbf{w}_i^{r,t} - \mathbf{w}^{r,t}\|^2 + \sum_{t=0}^{K-1} (1 - \mu\gamma)^{K-t} \gamma^2 \frac{\delta^2}{N} \end{aligned}$$

Since we choose $\gamma \leq \frac{1}{2L}$, we know $\sum_{t=0}^{K-1} (1 - \mu\gamma)^{K-t} (2\gamma - 4\gamma^2 L) \mathbb{E} (F(\mathbf{w}^*) - F(\mathbf{w}^{r,t})) \geq 0$. Plugging in Lemma 5 yields:

$$\mathbb{E} \|\mathbf{w}^R - \mathbf{w}^*\|^2 = (1 - \mu\gamma)^{RK} \mathbb{E} \|\mathbf{w}^0 - \mathbf{w}^*\|^2 + \frac{1}{\mu\gamma} (\gamma L + 2\gamma^2 L^2) (3K\gamma^2\delta^2 + 6K^2\gamma^2\zeta^2) + \frac{\gamma\delta^2}{\mu N},$$

Algorithm A2: Shuffling Local SGD (One Client)

Input: Clients $0, \dots, N-1$, Number of Local Steps K , Number of Epoch R , Mixing parameter $\hat{\alpha}$

Epoch for $r = 0, \dots, R-1$ **do**

 Server generates permutation $\sigma_r : [N] \mapsto [N]$.

 Client sets initial model $\mathbf{v}^{r,0} = \mathbf{v}^r$.

for $j = 0, \dots, N-1$ **do**

 Server sends $\mathbf{v}^{r,j}$ to Client $\sigma_r(j)$.

$\mathbf{v}^{r,j+1} = \text{SGD-Update}(\mathbf{v}^{r,j}, \eta, \sigma_r(j), K, \hat{\alpha})$.

 Client i does projection: $\mathbf{v}^{r+1} = \mathcal{P}_{\mathcal{W}}(\mathbf{v}^{r,N})$.

Output: $\hat{\mathbf{v}} = \mathbf{v}^R$.

SGD-Update($\mathbf{v}, \eta, j, K, \alpha$)

 Initialize $\mathbf{v}^0 = \mathbf{v}$

for $t = 0, \dots, K-1$ **do**

$\mathbf{v}^t = \mathbf{v}^{t-1} - \eta\alpha(j)N\nabla f_j(\mathbf{v}^{t-1}; \xi^{t-1})$

Output $\mathbf{v}^K$

Plugging in $\gamma = \frac{\log(RK)}{\mu RK}$ gives the convergence rate:

$$\mathbb{E} \|\mathbf{w}^R - \mathbf{w}^*\|^2 \leq \tilde{O} \left(\frac{\mathbb{E} \|\mathbf{w}^0 - \mathbf{w}^*\|^2}{RK} + \kappa \left(\frac{\delta^2}{\mu^2 R^2 K} + \frac{\zeta^2}{\mu^2 R^2} \right) + \frac{\delta^2}{\mu^2 N R K} \right),$$

which concludes the proof. $\square$

Equipped with above results, we are now ready to provide the convergence of main theorem.

Proof of Theorem 1. The proof simply follows from Lemma 3:

$$\begin{aligned} \|\hat{\alpha} - \alpha^*\|^2 &\leq 2 \|\hat{\alpha} - \alpha_g^*(\mathbf{w}^R)\|^2 + 2 \|\alpha_g^*(\mathbf{w}^R) - \alpha_g^*(\mathbf{w}^*)\|^2 \\ &\leq 2(1 - \mu\eta_\alpha)^{T_\alpha} + 8L\kappa_g^2 \sum_{j=1}^N \left(2 \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + 4L^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) \|\mathbf{w}^R - \mathbf{w}^*\|^2 \\ &\leq 2(1 - \mu\eta_\alpha)^{T_\alpha} + 8L\kappa_g^2 \left(2\bar{\zeta}_i(\mathbf{w}^*) + 4NL^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) \|\mathbf{w}^R - \mathbf{w}^*\|^2 \end{aligned}$$

Plugging in the convergence of $\|\mathbf{w}^R - \mathbf{w}^*\|^2$ from Lemma 6, and the stepsize $\eta_\alpha = \frac{1}{L_g}$ for α yields:

$$\mathbb{E} \|\alpha_i^R - \alpha_i^*\|^2 \leq \tilde{O} \left(\exp\left(-\frac{T_\alpha}{\kappa_g}\right) + \kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \left(\frac{D^2}{RK} + \kappa \left(\frac{\delta^2}{\mu^2 R^2 K} + \frac{\zeta^2}{\mu^2 R^2} \right) + \frac{\delta^2}{\mu^2 N R K} \right) \right).$$

$\square$

B.3 Proof of Convergence of Shuffling Local SGD

In this section, we are going to prove the convergence of proposed shuffled variant of Local SGD (Theorem 2). The whole proof framework follows the analysis of vanilla shuffling SGD, but notice that there are two differences. First, in vanilla shuffling SGD, in each epoch, algorithm only updates on each component function f_j once, while here we have to take K steps of SGD update on each component function. Second, we are considering a weighted sum objective in contrary to averaged objective in [56], which means we need to rescale the objective when we apply without-replacement concentration inequality. Even though our algorithm solves models for N clients, for the sake of simplicity, throughout the proof we only show the convergence of one client's model. The algorithm from one client point of view is described in Algorithm A2, where we drop the client index for notational convenience.

Proposition 1. Assume a sequence $\{\mathbf{w}^t\}_{t=1}^K$ is obtained by

$$\mathbf{w}^t = \mathbf{w}^{t-1} - \eta\alpha N \nabla f(\mathbf{w}^{t-1}; \xi^{t-1}), \quad t = 1, \dots, K,$$

then we have

$$\mathbf{w}^{t+1} = \mathbf{w}^0 - \left(\sum_{\tau=0}^t \prod_{t'=t}^{\tau+1} (\mathbf{I} - \alpha N \eta \mathbf{H}_{t'}) \right) \eta \alpha N \nabla f(\mathbf{w}^0) - \sum_{\tau=0}^t \prod_{t'=t}^{\tau+1} (\mathbf{I} - \alpha N \eta \mathbf{H}_{t'}) \eta \alpha N \boldsymbol{\delta}^t, \quad \forall 0 \leq t \leq K-1,$$

where $\boldsymbol{\delta}^t := \nabla f(\mathbf{w}^t; \xi^t) - \nabla f(\mathbf{w}^t)$, and by convention, we define $\prod_{j=a}^b \mathbf{A}_j = \mathbf{I}$ if $a < b$.

Proof. According to updating rule, we have:

$$\begin{aligned} \mathbf{w}^{t+1} - \mathbf{w}^0 &= \mathbf{w}^t - \mathbf{w}^0 - \eta \alpha N \nabla f(\mathbf{w}^t; \xi^t) \\ &= \mathbf{w}^t - \mathbf{w}^0 - \eta \alpha N \nabla f(\mathbf{w}^t) - \eta \alpha N \boldsymbol{\delta}^t \\ &= \mathbf{w}^t - \mathbf{w}^0 - \eta \alpha N \nabla f(\mathbf{w}^0) - \eta \alpha N (\nabla f(\mathbf{w}^t) - \nabla f(\mathbf{w}^0)) - \eta \alpha N \boldsymbol{\delta}^t. \end{aligned}$$

Since f is L smooth, and according to Mean Value Theorem, there is a matrix $\mathbf{H}_t$ satisfying $\mu \mathbf{I} \preceq \mathbf{H}_t \preceq L \mathbf{I}$, such that $\nabla f(\mathbf{w}^t) - \nabla f(\mathbf{w}^0) = \mathbf{H}_t (\mathbf{w}^t - \mathbf{w}^0)$. Hence we have:

$$\mathbf{w}^{t+1} - \mathbf{w}^0 = (\mathbf{I} - \eta \alpha N \mathbf{H}_t) (\mathbf{w}^t - \mathbf{w}^0) - \eta \alpha N \nabla f(\mathbf{w}^0) - \eta \alpha N \boldsymbol{\delta}^t.$$

Unrolling the recursion from t to 0 will conclude the proof. $\square$

The following lemma establishes the updating rule of models between epochs r and $r+1$. For notational convenience, whenever there is no confusion, we drop the superscript r in σ^r .

Lemma 7 (One epoch updating rule). *Let $\mathbf{v}^r$ and $\mathbf{v}^{r+1}$ be two iterates generated by Shuffling Local SGD (Algorithm A2), then the following updating rule holds:*

$$\mathbf{v}^{r+1} = \mathbf{v}^r - \sum_{j=1}^N \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) (\mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \boldsymbol{\delta}_j),$$

where

$$\begin{aligned} \mathbf{Q}_j &:= \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(j)) N \mathbf{H}_{t'}) \right) \eta \hat{\alpha}(\sigma(j)) N, \\ \boldsymbol{\delta}_j &:= \sum_{\tau=0}^{K-1} \prod_{t'=t}^{\tau+1} (\mathbf{I} - \hat{\alpha}(\sigma(j)) N \eta \mathbf{H}_{t'}) \eta \hat{\alpha}(\sigma(j)) N \boldsymbol{\delta}_{\sigma(j)}^t, \end{aligned}$$

by convention, we define $\prod_{j=a}^b \mathbf{A}_j = \mathbf{I}$ if $a < b$.

Proof. According to Proposition 1, we have

$$\begin{aligned} \mathbf{v}^{r,j+1} &= \mathbf{v}^{r,j} - \left(\sum_{\tau=0}^{K-1} \prod_{t'=t}^{\tau+1} (\mathbf{I} - \hat{\alpha}(\sigma(j)) N \eta \mathbf{H}_{t'}) \right) \eta \hat{\alpha}(\sigma(j)) N \nabla f_{\sigma(j)}(\mathbf{v}^{r,j}) \\ &\quad - \sum_{\tau=0}^{K-1} \prod_{t'=t}^{\tau+1} (\mathbf{I} - \hat{\alpha}(\sigma(j)) N \eta \mathbf{H}_{t'}) \eta \hat{\alpha}(\sigma(j)) N \boldsymbol{\delta}_{\sigma(j)}^t. \end{aligned}$$

Plugging our definition of $\mathbf{Q}_j$ and $\boldsymbol{\delta}_j$ yields:

$$\mathbf{v}^{r,j+1} - \mathbf{v}^r = \mathbf{v}^{r,j} - \mathbf{v}^r - \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^{r,j}) - \boldsymbol{\delta}_j.$$

Following the same reasoning in the proof of Proposition 1 will conclude the proof. $\square$

Lemma 8 (Summation by parts). *Let $\mathbf{A}_j$ and $\mathbf{B}_j$ be complex valued matrices. Then the following fact holds:*

$$\sum_{j=1}^N \mathbf{A}_j \mathbf{B}_j = \mathbf{A}_N \sum_{j=1}^N \mathbf{B}_j - \sum_{n=1}^{N-1} (\mathbf{A}_{n+1} - \mathbf{A}_n) \sum_{j=1}^n \mathbf{B}_j.$$

Proposition 2 (Spectral bound of polynomial expansion). *Given a collection of matrices $\{\mathbf{A}_t\}$ and $\{\mathbf{B}_t\}$, such that $\mathbf{A}_t \preceq L\mathbf{I}$ and $\mathbf{B}_t \preceq L\mathbf{I}$, the following bound hold:*

$$\left\| \prod_{t=l}^h (\mathbf{I} - a\mathbf{A}_t) - \mathbf{I} \right\| \leq \sum_{m=1}^{h-l} \left(\frac{e(h-l)}{m} \right)^m (aL)^m,$$

$$\left\| \prod_{t=l}^h (\mathbf{I} - a\mathbf{A}_t) - \prod_{t=l}^h (\mathbf{I} - b\mathbf{B}_t) \right\| \leq \sum_{m=1}^{h-l} \left(\frac{e(h-l)}{m} \right)^m (aL)^m + \sum_{m=1}^{h-l} \left(\frac{e(h-l)}{m} \right)^m (bL)^m.$$

Proof. We start with proving the first statement. Expanding the product yields:

$$\prod_{t=l}^h (\mathbf{I} - a\mathbf{A}_t) = \mathbf{I} + \sum_{m=1}^{h-l} (-1)^m a^m \sum_{|S|=m, |S| \subseteq \{l, \dots, h\}} \prod_{m' \in S} \mathbf{A}_{m'}.$$

Hence we have:

$$\left\| \prod_{t=l}^h (\mathbf{I} - a\mathbf{A}_t) - \mathbf{I} \right\| = \left\| \sum_{m=1}^{h-l} (-1)^m a^m \sum_{|S|=m, |S| \subseteq \{l, \dots, h\}} \prod_{m' \in S} \mathbf{A}_{m'} \right\| \leq \sum_{m=1}^{h-l} \binom{h-l}{m} (aL)^m$$

According to the upper bound for binomial coefficients: $\binom{h-l}{m} \leq \left(\frac{e(h-l)}{m} \right)^m$, we have:

$$\sum_{m=1}^{h-l} \binom{h-l}{m} (aL)^m \leq \sum_{m=1}^{h-l} \left(\frac{e(h-l)}{m} \right)^m (aL)^m.$$

Then we switch to the second one. Using the same expanding product yields:

$$\begin{aligned} & \left\| \prod_{t=l}^h (\mathbf{I} - a\mathbf{A}_t) - \prod_{t=l}^h (\mathbf{I} - b\mathbf{B}_t) \right\| \\ &= \left\| \sum_{m=1}^{h-l} (-1)^m a^m \sum_{|S|=m, |S| \subseteq \{l, \dots, h\}} \prod_{m' \in S} \mathbf{A}_{m'} - \sum_{m=1}^{h-l} (-1)^m b^m \sum_{|S|=m, |S| \subseteq \{l, \dots, h\}} \prod_{m' \in S} \mathbf{B}_{m'} \right\| \\ &\leq \sum_{m=1}^{h-l} \binom{h-l}{m} (aL)^m + \sum_{m=1}^{h-l} \binom{h-l}{m} (bL)^m \\ &\leq \sum_{m=1}^{h-l} \left(\frac{e(h-l)}{m} \right)^m (aL)^m + \sum_{m=1}^{h-l} \left(\frac{e(h-l)}{m} \right)^m (bL)^m. \end{aligned}$$

□

The following concentration result is the key to bound variance during shuffling updating. The original result holds for the average of gradients, and we will later on generalize it to an arbitrary weighted sum of gradients.

Lemma 9 ([61, Theorem 2]). *Suppose $n \geq 2$. Let $\mathbf{g}_1, \mathbf{g}_2, \dots, \mathbf{g}_n \in \mathbb{R}^d$ satisfy $\|\mathbf{g}_j\| \leq G$ for all j . Let $\bar{\mathbf{g}} = \frac{1}{n} \sum_{j=1}^n \mathbf{g}_j$. Let $\sigma \in \mathcal{S}_n$ be a uniform random permutation of n elements. Then, for $i \leq n$, with probability at least $1 - p$, we have*

$$\left\| \frac{1}{i} \sum_{j=1}^i \mathbf{g}_{\sigma(j)} - \bar{\mathbf{g}} \right\| \leq G \sqrt{\frac{8(1 - \frac{i-1}{n}) \log \frac{2}{p}}{i}}.$$

Lemma 10 (Concentration of partial sum of gradients). *Given a uniformly randomly generated permutation σ , and simplex vector α , if we assume each $\sup_{\mathbf{v} \in \mathcal{W}} \|\nabla f_j(\mathbf{v})\| \leq G$, then the following statement holds true:*

$$\left\| \sum_{j=0}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \leq G \sqrt{8n \log(1/p)} + \frac{n}{N} \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|.$$

Proof. The proof works by re-writing weighted sum of vectors to average of the these vectors:

$$\begin{aligned} \left\| \sum_{j=0}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| &= \frac{1}{N} \left\| \sum_{j=0}^n \hat{\alpha}(\sigma(j)) N \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ &= \frac{1}{N} \left(\left\| \sum_{j=0}^n \hat{\alpha}(\sigma(j)) N \nabla f_{\sigma(j)}(\mathbf{v}^r) - n \nabla \Phi(\hat{\alpha}, \mathbf{v}^r) \right\| + n \left\| \nabla \Phi(\hat{\alpha}, \mathbf{v}^r) \right\| \right) \\ &\leq G \sqrt{8n \log(1/p)} + \frac{n}{N} \left\| \nabla \Phi(\hat{\alpha}, \mathbf{v}^r) \right\|. \end{aligned}$$

□

Proposition 3 (Spectral norm bound of $\mathbf{Q}_j$). *Let $\mathbf{Q}_j$ be defined in (11). Then the following bound for the spectral norm of $\mathbf{Q}_j$ holds true for all $j \in [N]$:*

$$\|\mathbf{Q}_j\| \leq \eta \hat{\alpha}(\sigma(j)) N K (1 + \eta N L)^K$$

Proof. The proof can be completed by writing down the definition of $\mathbf{Q}_j$ and applying Cauchy-Schwartz inequality:

$$\begin{aligned} \|\mathbf{Q}_j\| &= \left\| \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(j)) N' \mathbf{H}_{t'}) \right) \eta \hat{\alpha}(\sigma(j)) N \right\| \\ &\leq \eta \hat{\alpha}(\sigma(j)) N \sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} \|(\mathbf{I} - \eta \hat{\alpha}(\sigma(j)) N' \mathbf{H}_{t'})\| \\ &\leq \eta \hat{\alpha}(\sigma(j)) N \sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (1 + \eta \hat{\alpha}(\sigma(j)) N' L) \\ &\leq \eta \hat{\alpha}(\sigma(j)) N K (1 + \eta N L)^K. \end{aligned}$$

The last step is due to we choose η such that $\eta N L \leq \frac{1}{K}$.

□

The following lemma establishes the bound regarding cumulative update between two epochs, namely, $\mathbf{v}^{r+1} - \mathbf{v}^r$. In particular, Lemma 11 below shows that: (a) in shuffling Local SGD, our update from $\mathbf{v}^r$ to $\mathbf{v}^{r+1}$ approximates performing NK times of gradient descent with $\hat{\alpha}(j) N \nabla f_{\sigma(j)}(\mathbf{v}^r)$, namely, the bias is controlled, and (b) the update itself is bounded, and can be related to the norm of full gradient.

Lemma 11. *During the dynamic of Algorithm A2, the following statements hold true with probability at least $1 - p$:*

(a)

$$\begin{aligned} \left\| \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \eta N K \sum_{j=1}^N \hat{\alpha}(j) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 &\leq 10 \eta^2 N^2 K^2 \left(\frac{e}{4R - e} \right)^2 \left\| \nabla \Phi(\hat{\alpha}, \mathbf{v}^r) \right\|^2 \\ &\quad + 128 \eta^2 N^3 K^2 \left(\frac{e}{4R - e} \right)^2 G^2 \log(1/p). \end{aligned}$$

(b) for any N' such that $0 \leq N' < N$

$$\left\| \sum_{j=1}^{N'-1} \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \leq 3e \eta N K \left(\left\| \nabla \Phi(\hat{\alpha}, \mathbf{v}^r) \right\| + G \sqrt{8N \log(1/p)} \right),$$

where

$$\mathbf{Q}_j := \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(j)) N' \mathbf{H}_{t'}) \right) \eta \hat{\alpha}(\sigma(j)) N. \quad (11)$$

Proof. We start with proving statement (a). Let $\mathbf{A}_j = \frac{\mathbf{Q}_j}{\hat{\alpha}(\sigma(j))}$ and $\mathbf{B}_j = \hat{\alpha}(\sigma(j))\nabla f_{\sigma(j)}(\mathbf{v}^r)$, applying the identity of summation by parts yields:

$$\begin{aligned} \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) &= \frac{\mathbf{Q}_{N-1}}{\hat{\alpha}(\sigma(N-1))} \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right) \sum_{j=1}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \\ &\left\| \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}) - \eta N K \sum_{j=1}^N \hat{\alpha}(j) \nabla f_j(\mathbf{v}) \right\|^2 \\ &\leq 2 \underbrace{\left\| \left(\frac{\mathbf{Q}_{N-1}}{\hat{\alpha}(\sigma(N-1))} - \eta N K \mathbf{I} \right) \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2}_{T_1} + 2 \underbrace{\left\| \sum_{n=1}^{N-1} \left(\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right) \sum_{j=1}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2}_{T_2}. \end{aligned}$$

According to Proposition 2, we have:

$$\left\| \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(j)) N \mathbf{H}_{t'}) - \mathbf{I} \right\| \leq \sum_{m=1}^{K-2-\tau} \left(\frac{e(K-2-\tau)}{m} \eta \hat{\alpha}(\sigma(j)) N L \right)^m.$$

Since we choose $\eta \leq \frac{1}{4NKR L}$, we have:

$$\left\| \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(j)) N \mathbf{H}_{t'}) - \mathbf{I} \right\| \leq \sum_{m=1}^{K-2-\tau} \left(\frac{e}{4Rm} \right)^m \leq \frac{e}{4R-e}, \quad (12)$$

where we use the fact that $\sum_{m=1}^{K-2-\tau} \left(\frac{e}{4Rm} \right)^m \leq \sum_{m=1}^{K-2-\tau} \left(\frac{e}{4R} \right)^m \leq \frac{e}{4R} \frac{1}{1-e/4R}$. Hence we know:

$$\begin{aligned} T_1 &\leq \left\| \left(\frac{\mathbf{Q}_{N-1}}{\hat{\alpha}(\sigma(N-1))} - \eta N K \mathbf{I} \right) \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\ &\leq \left\| \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(N-1)) N \mathbf{H}_{t'}) \right) \eta N - \eta N K \mathbf{I} \right\|^2 \left\| \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\ &\leq \eta^2 N^2 K \sum_{\tau=0}^{K-1} \left\| \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(N-1)) N \mathbf{H}_{t'}) - \mathbf{I} \right\|^2 \left\| \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\ &\leq \eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 \left\| \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2. \end{aligned}$$

Thus we have:

$$T_1 \leq \eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2.$$

For T_2 , we first examine the bound of $\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))}$:

$$\begin{aligned} \left\| \frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right\| &= \left\| \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(n+1)) N \mathbf{H}_{t'}) \right) \eta N - \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(n)) N \mathbf{H}_{t'}) \right) \eta N \right\| \\ &= \eta N \left\| \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(n+1)) N \mathbf{H}_{t'}) \right) - \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(n)) N \mathbf{H}_{t'}) \right) \right\| \\ &= \eta N \left\| \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(n+1)) N \mathbf{H}_{t'}) \right) - \left(\sum_{\tau=0}^{K-1} \prod_{t'=K-1}^{\tau+1} (\mathbf{I} - \eta \hat{\alpha}(\sigma(n)) N \mathbf{H}_{t'}) \right) \right\| \\ &\leq \eta N \sum_{\tau=0}^{K-1} \left(\sum_{m=1}^{K-2-\tau} \left(\frac{e(K-2-\tau)}{m} \eta \hat{\alpha}(\sigma(n)) N L \right)^m + \sum_{m=1}^{K-2-\tau} \left(\frac{e(K-2-\tau)}{m} \eta \hat{\alpha}(\sigma(n+1)) N L \right)^m \right). \end{aligned}$$

where we evoke Proposition 2 at last step. Given that $\eta \leq \frac{1}{4NKR L}$ we have:

$$\begin{aligned} \left\| \frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right\| &\leq \eta N \sum_{\tau=0}^{K-1} \left(\sum_{m=1}^{K-2-\tau} \left(\frac{e}{4Rm} \hat{\alpha}(\sigma(n)) \right)^m + \sum_{m=1}^{K-2-\tau} \left(\frac{e}{4Rm} \hat{\alpha}(\sigma(n+1)) L \right)^m \right) \\ &\leq \eta N K \left(\frac{\hat{\alpha}(\sigma(n))e}{4R-e} + \frac{\hat{\alpha}(\sigma(n+1))e}{4R-e} \right). \end{aligned}$$

where we use the reasoning in (12). Hence for $\sqrt{T_2}$:

$$\begin{aligned} \sqrt{T_2} &\leq \eta N K \sum_{n=1}^{N-1} \left(\frac{\hat{\alpha}(\sigma(n))e}{4R-e} + \frac{\hat{\alpha}(\sigma(n+1))e}{4R-e} \right) \left\| \sum_{j=1}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ &\leq \eta N K \frac{e}{4R-e} \sum_{n=1}^{N-1} (\hat{\alpha}(\sigma(n)) + \hat{\alpha}(\sigma(n+1))) \left(G \sqrt{8n \log(1/p)} + \frac{n}{N} \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\| \right) \\ &\leq \eta N K \frac{2e}{4R-e} \left(G \sqrt{8N \log(1/p)} + \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\| \right). \end{aligned}$$

where at last step we evoke Lemma 10. So we can conclude $T_2 \leq 2\eta^2 N^2 K^2 \left(\frac{2e}{4R-e} \right)^2 \left(G^2 8N \log(1/p) + \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 \right)$. Putting the bounds of T_1 and T_2 together will conclude the proof for (a).

Now we switch to proving (b). Once again by the summation of parts identity we have:

$$\sum_{j=1}^{N'} \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) = \frac{\mathbf{Q}_{N'}}{\hat{\alpha}(\sigma(N'))} \sum_{j=1}^{N'} \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N'-1} \left(\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right) \sum_{j=1}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r).$$

Taking the norm of both side yields:

$$\begin{aligned} \left\| \sum_{j=1}^{N'} \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| &= \underbrace{\left\| \frac{\mathbf{Q}_{N'}}{\hat{\alpha}(\sigma(N'))} \sum_{j=1}^{N'} \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|}_B \\ &\quad + \underbrace{\left\| \sum_{n=1}^{N'-1} \left(\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right) \sum_{j=1}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|}_C. \end{aligned}$$

Plugging our developed bound for $\|\mathbf{Q}_{N'}\|$ and $\left\| \sum_{n=1}^{N'-1} \left(\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right) \right\|$ yields:

$$\begin{aligned} B &\leq \left\| \frac{\mathbf{Q}_{N'}}{\hat{\alpha}(\sigma(N'))} \right\| \left\| \sum_{j=1}^{N'} \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ &\leq \eta N K (1 + \eta N L)^K \left(G \sqrt{8N' \log(1/p)} + \frac{N'}{N} \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\| \right). \end{aligned}$$

where at last step we evoke Lemma 10. And for C, we use the similar reasoning:

$$\begin{aligned} C &\leq \sum_{n=1}^{N'-1} \left\| \left(\frac{\mathbf{Q}_{n+1}}{\hat{\alpha}(\sigma(n+1))} - \frac{\mathbf{Q}_n}{\hat{\alpha}(\sigma(n))} \right) \right\| \left\| \sum_{j=1}^n \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ &\leq \sum_{n=1}^{N'-1} \eta N K \frac{e}{4R-e} (\hat{\alpha}(\sigma(n+1)) + \hat{\alpha}(\sigma(n))) \left(G \sqrt{8n \log(1/p)} + \frac{n}{N} \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\| \right) \\ &\leq 2\eta N K \frac{e}{4R-e} \left(G \sqrt{8N \log(1/p)} + \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\| \right). \end{aligned}$$

Putting these pieces together yields:

$$\left\| \sum_{j=1}^{N'} \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \leq 3e\eta NK \left(\|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\| + G\sqrt{8N \log(1/p)} \right).$$

□

Lemma 12. *During the dynamic of Algorithm A2, the following statements hold true with probability at least $1 - p$:*

$$\begin{aligned} & \left\| \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\ & \leq 18e^6 \eta^4 N^4 K^4 L^4 \left(\|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\|^2 + 8G^2 N \log(1/p) \right) \end{aligned}$$

Proof. We first apply Cauchy-Schwartz inequality:

$$\begin{aligned} & \left\| \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ & \leq \sum_{n=1}^{N-1} \left\| \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \right\| \left\| \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \right\| \left\| \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ & \leq (1 + \eta NK + \eta^2 N^2 KL)^{2N} \eta N L K (1 + \eta N L)^K L \sum_{n=1}^{N-1} \hat{\alpha}(\sigma(n+1)) \left\| \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ & \leq e^2 \eta N K L^2 \sum_{n=1}^{N-1} \hat{\alpha}(\sigma(n+1)) \left\| \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|. \end{aligned}$$

We proceed by applying the bound from Lemma 11 (b):

$$\left\| \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \leq 3e\eta NK \left(\|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\| + G\sqrt{8N \log(1/p)} \right).$$

Therefore, it follows that:

$$\begin{aligned} & \left\| \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\| \\ & \leq e^2 \eta N K L^2 \sum_{n=1}^{N-1} \hat{\alpha}(\sigma(n+1)) \cdot 3e\eta NK \left(\|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\| + G\sqrt{8N \log(1/p)} \right) \\ & \leq 3e^3 \eta^2 N^2 K^2 L^2 \left(\|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\| + G\sqrt{8N \log(1/p)} \right) \end{aligned}$$

□

Lemma 13 (Noise bound). *During the dynamic of Algorithm A2, the following statement for gradient noises holds true with probability at least $1 - p$:*

$$\mathbb{E} \left\| \sum_{j=1}^N \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \boldsymbol{\delta}_j \right\| \leq \eta NK e^2 \delta,$$

where

$$\boldsymbol{\delta}_j := \sum_{\tau=0}^{K-1} \prod_{t'=\tau}^{\tau+1} (\mathbf{I} - \hat{\alpha}(\sigma(j)) N \eta \mathbf{H}_{t'}) \eta \hat{\alpha}(\sigma(j)) N \boldsymbol{\delta}_{\sigma(j)}^t.$$

Proof. According to triangle and Cauchy-Schwartz inequalities we have:

$$\begin{aligned}
\left\| \sum_{j=1}^N \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \delta_j \right\| &\leq \sum_{j=1}^N \prod_{j'=N}^{j+1} \|\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}\| \|\delta_j\| \\
&\leq \sum_{j=1}^N (1 + (\eta \hat{\alpha}(\sigma(j)) NK (1 + \eta NL^K) L))^N \|\delta_j\| \\
&\leq \sum_{j=1}^N \underbrace{(1 + (\eta \hat{\alpha}(\sigma(j)) NK (1 + \eta NL^K) L))^N}_{\leq e} \cdot \eta \hat{\alpha}(\sigma(j)) NK \underbrace{(1 + \eta NL^K)}_{\leq e} \delta \\
&\leq \eta NK e^2 \delta.
\end{aligned}$$

□

B.4 Proof of Theorem 2

Proof. For notational convenience, let us define

$$\begin{aligned}
\mathbf{g}^r &:= \sum_{j=1}^N \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r), \\
\delta^r &:= \sum_{j=1}^N \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \delta_j.
\end{aligned}$$

Then we recall the updating rule of $\mathbf{v}$ (Lemma 7):

$$\mathbf{v}^{r+1} = \mathcal{P}_{\mathcal{W}}(\mathbf{v}^r - \mathbf{g}^r - \delta^r)$$

Hence we have:

$$\begin{aligned}
\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\hat{\alpha})\|^2 &= \mathbb{E} \|\mathcal{P}_{\mathcal{W}}(\mathbf{v}^r - \mathbf{g}^r - \delta^r - \mathbf{v}^*(\hat{\alpha}))\|^2 \\
&\leq \mathbb{E} \|\mathbf{v}^r - \mathbf{g}^r - \delta^r - \mathbf{v}^*(\hat{\alpha})\|^2 \\
&\leq \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\alpha})\|^2 - 2\mathbb{E} \langle \mathbf{g}^r, \mathbf{v}^r - \mathbf{v}^*(\hat{\alpha}) \rangle + \mathbb{E} \|\mathbf{g}^r\|^2 + \mathbb{E} \|\delta^r\|^2 \\
&\leq \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\alpha})\|^2 - 2\mathbb{E} \langle \eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r), \mathbf{v}^r - \mathbf{v}^*(\hat{\alpha}) \rangle - 2\mathbb{E} \langle \mathbf{g}^r - \eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r), \mathbf{v}^r - \mathbf{v}^*(\hat{\alpha}) \rangle \\
&\quad + \mathbb{E} \|\mathbf{g}^r\|^2 + \mathbb{E} \|\delta^r\|^2.
\end{aligned}$$

Now, applying strongly convexity of $\Phi(\hat{\alpha}, \cdot)$ and Cauchy-Schwartz inequality yields:

$$\begin{aligned}
\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\hat{\alpha})\|^2 &\leq (1 - \mu \eta NK) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\alpha})\|^2 - \eta NK \mathbb{E} [\Phi(\hat{\alpha}, \mathbf{v}^r) - \Phi(\hat{\alpha}, \mathbf{v}^*(\hat{\alpha}))] \\
&\quad + \frac{1}{2} \left(\frac{1}{\mu \eta NK} \mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 + \mu \eta NK \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\alpha})\|^2 \right) \\
&\quad + \mathbb{E} \|\mathbf{g}^r\|^2 + \mathbb{E} \|\delta^r\|^2 \\
&\leq (1 - \frac{1}{2} \mu \eta NK) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\alpha})\|^2 - \eta NK \mathbb{E} [\Phi(\hat{\alpha}, \mathbf{v}^r) - \Phi(\hat{\alpha}, \mathbf{v}^*(\hat{\alpha}))] \\
&\quad + \frac{1}{2\mu \eta NK} \mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 \\
&\quad + 2\mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 + 2\mathbb{E} \|\eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 + \mathbb{E} \|\delta^r\|^2.
\end{aligned}$$

Since $\Phi(\hat{\alpha}, \cdot)$ is L smooth, we have: $\mathbb{E} \|\nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 \leq 2L \mathbb{E} [\Phi(\hat{\alpha}, \mathbf{v}^r) - \Phi(\hat{\alpha}, \mathbf{v}^*(\hat{\alpha}))]$. Therefore, we have:

$$\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\hat{\alpha})\|^2 \leq (1 - \frac{1}{2} \mu \eta NK) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\alpha})\|^2 - (\eta NK - 4\eta^2 N^2 K^2 L) \mathbb{E} [\Phi(\hat{\alpha}, \mathbf{v}^r) - \Phi(\hat{\alpha}, \mathbf{v}^*(\hat{\alpha}))] \quad (13)$$

$$+ \left(\frac{1}{2\mu \eta NK} + 2 \right) \mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\alpha}, \mathbf{v}^r)\|^2 + \mathbb{E} \|\delta^r\|^2. \quad (14)$$

Now, we examine the term $\|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\|^2$. First according to summation by part (Lemma 8) by letting $\mathbf{A}_j := \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'})$ and $\mathbf{B}_j = \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r)$, we have:

$$\begin{aligned} \mathbf{g}^r &= \sum_{j=1}^N \prod_{j'=N}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \\ &= \sum_{j=1}^N \mathbf{A}_j \mathbf{B}_j = \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) - \prod_{j'=N}^{n+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \\ &= \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r). \end{aligned}$$

Hence we have:

$$\begin{aligned} &\|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\|^2 \\ &= \left\| \eta NK \sum_{j=1}^N \hat{\boldsymbol{\alpha}}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) - \left(\sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right) \right\|^2 \\ &\stackrel{(1)}{=} 2 \left\| \left(\eta NK \sum_{j=1}^N \hat{\boldsymbol{\alpha}}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right) \right\|^2 + 2 \left\| \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\ &\stackrel{(2)}{\leq} \left(20\eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 + 36e^6 \eta^4 N^4 K^4 L^4 \right) \|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\|^2 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 G^2 \log(1/p) \\ &\quad + 244e^6 \eta^4 N^4 K^4 L^4 G^2 N \log(1/p) \\ &\stackrel{(3)}{\leq} \left(20\eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 + 36e^6 \eta^4 N^4 K^4 L^4 \right) 2L(\Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r) - \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^*(\hat{\boldsymbol{\alpha}}))) \\ &\quad + \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 \right) G^2 N \log(1/p), \end{aligned}$$

where in (1) we apply Jensen's inequality, in (2) we plug in Lemma 11 (a), and Lemma 12, and in (3) we use the L -smoothness of Φ . Plugging above bound back in (19) yields:

$$\begin{aligned} &\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 \\ &\leq (1 - \frac{1}{2}\mu\eta NK) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 + \eta^2 N^2 K^2 e^4 \delta^2 \\ &\quad - \underbrace{\left(\eta NK - 4\eta^2 N^2 K^2 L - \left(\frac{1}{2\mu\eta NK} + 2 \right) \left(20\eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 - 36e^6 \eta^4 N^4 K^4 L^4 \right) \right)}_{T_1} \mathbb{E}[\Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r) - \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^*(\hat{\boldsymbol{\alpha}}))] \\ &\quad + \left(\frac{1}{2\mu\eta NK} + 2 \right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 \right) G^2 N \log(1/p). \end{aligned}$$

Since we choose $\eta = \frac{4\log(\sqrt{NKR})}{\mu NKR}$, and large enough epoch number:

$$R \geq \max \left\{ \left(\frac{40}{\mu} + 1 \right) e, 16 \log(\sqrt{NKR}), 64\kappa \log(\sqrt{NKR}) \right\},$$

we know that $T_1 \leq 0$. We thus have:

$$\begin{aligned} &\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 \\ &\leq (1 - \frac{1}{2}\mu\eta NK) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 + \eta^2 N^2 K^2 e^4 \delta^2 \\ &\quad + \left(\frac{1}{2\mu\eta NK} + 2 \right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 \right) G^2 N \log(1/p) \end{aligned}$$

Unrolling the recursion from $r = R$ to 0:

$$\begin{aligned} & \mathbb{E} \|\mathbf{v}^R - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 \\ & \leq (1 - \frac{1}{2}\mu\eta NK)^R \mathbb{E} \|\mathbf{v}^0 - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 + \frac{2}{\mu}\eta NK e^4 \delta^2 \\ & \quad + \frac{1}{\mu} \left(\frac{1}{2\mu\eta NK} + 2 \right) \left(488e^6 \eta^3 N^3 K^3 L^4 + 512\eta N^2 K \left(\frac{e}{4R - e} \right)^2 \right) G^2 N \log(1/p). \end{aligned}$$

Plugging in our choice of η will conclude the proof:

$$\mathbb{E} \|\mathbf{v}^R - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 \leq \tilde{O} \left(\frac{\mathbb{E} \|\mathbf{v}^0 - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2}{NKR^2} + \frac{\delta^2}{\mu^2 R} + \left(\frac{L^4 + N}{\mu^4 R^2} \right) G^2 N \log(1/p) \right).$$

Finally, according to Lemma 2 we can complete the proof:

$$\begin{aligned} \Phi(\boldsymbol{\alpha}_i^*, \hat{\mathbf{v}}_i) - \Phi(\boldsymbol{\alpha}_i^*, \mathbf{v}_i^*) & \leq 2L \|\mathbf{v}_i^R - \mathbf{v}^*(\hat{\boldsymbol{\alpha}}_i)\|^2 + \left(2\kappa_\Phi^2 L + \frac{4NG^2}{L} \right) \|\hat{\boldsymbol{\alpha}}_i - \boldsymbol{\alpha}^*\|^2 \\ & \leq \tilde{O} \left(\frac{LD^2}{NKR^2} + \frac{L\delta^2}{\mu^2 R} + \left(\frac{L^4 + N}{\mu^4 R^2} \right) LG^2 N \log(1/p) \right) \\ & \quad + \kappa_\Phi^2 L \tilde{O} \left(\exp \left(-\frac{T\alpha}{\kappa_g} \right) + \kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \left(\frac{D^2}{RK} + \frac{\kappa \zeta^2}{\mu^2 R^2} + \frac{\delta^2}{\mu^2 NKR} \right) \right), \end{aligned}$$

where we plug in the convergence result from Theorem 1 at last step. $\square$

C Proof of Convergence of Single Loop Algorithm

In this section, we turn to presenting the proof of single loop PERM algorithm (Algorithm 2) where the learning of mixing parameters and personalized models are coupled. Compared to Algorithm A2, here during the optimization of model, the mixing parameters are also being updated. As a result, we need to decouple the two updates which makes the analysis more involved. We begin with some technical lemmas that support the proof of main result.

C.1 Technical Lemmas

Proposition 4 (Basic Properties of SGD on Smooth Strongly Convex Function). *Let $\mathbf{w}^t$ to be the t th iterate of minibatch SGD on smooth and strongly convex function F , with minibatch size M and learning rate γ . Also assume the variance is bounded by δ . Then the following statements hold true after T iterations of SGD:*

$$\mathbb{E} \|\nabla F(\mathbf{w}^T)\|^2 \leq 2L(1 - \mu\gamma)^T (F(\mathbf{w}^0) - F(\mathbf{w}^*)) + \frac{2\gamma\kappa\delta^2}{M} \quad (15)$$

$$\mathbb{E} \|\mathbf{w}^{T+1} - \mathbf{w}^T\|^2 \leq 2\gamma^2 L(1 - \mu\gamma)^T (F(\mathbf{w}^0) - F(\mathbf{w}^*)) + \frac{2\gamma^3 \kappa \delta^2}{M} + \frac{\gamma^2 \delta^2}{M} \quad (16)$$

$$\mathbb{E} \|\mathbf{w}^T - \mathbf{w}^*\|^2 \leq \frac{2}{\mu} (1 - \mu\gamma)^T (F(\mathbf{w}^0) - F(\mathbf{w}^*)) + 2\gamma \frac{\delta^2}{\mu^2 M}. \quad (17)$$

Lemma 14 (Bounded iterates difference of $\boldsymbol{\alpha}$). *Let $\{\boldsymbol{\alpha}_i^r\}$ be iterates generated by Algorithm 2, then under conditions of Theorem 3, the following statement holds:*

$$\|\boldsymbol{\alpha}_i^r - \boldsymbol{\alpha}_i^{r-1}\|^2 \leq 6 \left(1 - \frac{1}{\kappa_g} \right)^{T\alpha} + O(\kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*)) \left(\gamma^2 L(1 - \mu\gamma)^r (F(\mathbf{w}^0) - F(\mathbf{w}^*)) + \frac{\gamma^3 \kappa \delta^2}{M} + \frac{\gamma^2 \delta^2}{M} \right)$$

Proof. Define

$$\mathbf{z}^r = \left[\|\nabla f_i(\mathbf{w}^r) - \nabla f_1(\mathbf{w}^r)\|^2, \dots, \|\nabla f_i(\mathbf{w}^r) - \nabla f_N(\mathbf{w}^r)\|^2 \right].$$

According to updating rule of α in Algorithm 2 and Lemma 3 we have:

$$\begin{aligned}
\|\alpha_i^r - \alpha_i^{r-1}\|^2 &\leq 3\|\alpha_i^r - \alpha_{g_i}^*(\mathbf{w}^r)\|^2 + 3\|\alpha_{g_i}^*(\mathbf{w}^{r-1}) - \alpha_{g_i}^*(\mathbf{w}^r)\|^2 + 3\|\alpha_{g_i}^*(\mathbf{w}^{r-1}) - \alpha_i^{r-1}\|^2 \\
&\leq 6(1 - \mu_g \eta \alpha)^{T\alpha} + 3\|\alpha_{g_i}^*(\mathbf{w}^{r-1}) - \alpha_{g_i}^*(\mathbf{w}^r)\|^2 \\
&\leq 6(1 - \mu_g \eta \alpha)^{T\alpha} + 3\kappa_g^2 \|\mathbf{z}^{r-1} - \mathbf{z}^r\|^2 \\
&\leq 6(1 - \mu_g \eta \alpha)^{T\alpha} + 3\kappa_g^2 \sum_{j=1}^N \|\nabla f_i(\mathbf{w}^r) - \nabla f_j(\mathbf{w}^r) + \nabla f_i(\mathbf{w}^{r-1}) - \nabla f_j(\mathbf{w}^{r-1})\|^2 4L^2 \|\mathbf{w}^r - \mathbf{w}^{r-1}\|^2
\end{aligned}$$

where the third inequality follows from (8). Since $\|\nabla f_i(\mathbf{w}^r) - \nabla f_j(\mathbf{w}^r)\| \leq \|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\| + 2L\|\mathbf{w}^r - \mathbf{w}^*\|$, we can conclude that

$$\begin{aligned}
\|\alpha_i^r - \alpha_i^{r-1}\|^2 &\leq 6(1 - \mu_g \eta \alpha)^{T\alpha} \\
&\quad + 12L^2 \kappa_g^2 \sum_{j=1}^N \left(8\|\nabla f_i(\mathbf{w}^*) - \nabla f_j(\mathbf{w}^*)\|^2 + 8L^2 \|\mathbf{w}^r - \mathbf{w}^*\|^2 + 8L^2 \|\mathbf{w}^{r-1} - \mathbf{w}^*\|^2 \right) \|\mathbf{w}^r - \mathbf{w}^{r-1}\|^2 \\
&\leq 6 \left(1 - \frac{1}{\kappa_g} \right)^{T\alpha} + O \left(\kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) \|\mathbf{w}^r - \mathbf{w}^{r-1}\|^2 \right) \\
&\leq 6 \left(1 - \frac{1}{\kappa_g} \right)^{T\alpha} + O \left(\kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) \right) \left(\gamma^2 L (1 - \mu\gamma)^r (F(\mathbf{w}^0) - F(\mathbf{w}^*)) + \frac{\gamma^3 \kappa \delta^2}{M} + \frac{\gamma^2 \delta^2}{M} \right)
\end{aligned}$$

where at last step we plug in Proposition 4 (16). $\square$

Lemma 15 (Convergence of α). *Let $\{\hat{\alpha}_i\}_{i=1}^N$ be the mixing parameters generated by Algorithm 2. Then under the conditions of Theorem 3, the following statement holds:*

$$\|\hat{\alpha}_i - \alpha^*\|^2 \leq 2(1 - \frac{1}{\kappa_g})^{T\alpha} + O \left(\kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \frac{2}{\mu} (1 - \mu\gamma)^T + 2\gamma \frac{\delta^2}{\mu^2 M} \right), i \in [N]$$

Proof. We notice the following decomposition:

$$\begin{aligned}
\|\hat{\alpha}_i - \alpha^*\|^2 &= \|\alpha_i^R - \alpha_{g_i}^*(\mathbf{w}^*)\|^2 \\
&\leq 2\|\alpha_i^R - \alpha_{g_i}^*(\mathbf{w}^R)\|^2 + 2\|\alpha_{g_i}^*(\mathbf{w}^R) - \alpha_{g_i}^*(\mathbf{w}^*)\|^2 \\
&\leq 2(1 - \frac{1}{\kappa_g})^{T\alpha} + O \left(\kappa_g^2 \left(\bar{\zeta}_i(\mathbf{w}^*) + NL^2 \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) 4L \|\mathbf{w}^R - \mathbf{w}^*\|^2 \right) \\
&\leq 2(1 - \frac{1}{\kappa_g})^{T\alpha} + O \left(\kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \frac{2}{\mu} (1 - \mu\gamma)^T + 2\gamma \frac{\delta^2}{\mu^2 M} \right),
\end{aligned}$$

where in the second inequality we apply Lemma 3, and in the third inequality we use Proposition 4 (17). $\square$

C.2 Proof of Theorem 3

Proof. According to Lemma 2, we have:

$$\Phi(\alpha_i^*, \hat{\mathbf{v}}_i) - \Phi(\alpha_i^*, \mathbf{v}_i^*) \leq 2L \|\mathbf{v}_i^R - \mathbf{v}^*(\hat{\alpha}_i)\|^2 + \left(2\kappa_\Phi^2 L + \frac{4NG^2}{L} \right) \|\hat{\alpha}_i - \alpha_i^*\|^2.$$

We first examine the convergence of $\|\mathbf{v}_i^R - \mathbf{v}^*(\hat{\alpha}_i)\|^2$. Applying Cauchy-Schwartz inequality yields:

$$\begin{aligned}
\|\mathbf{v}^{r+1} - \mathbf{v}^*(\alpha^{r+1})\|^2 &\leq \left(1 + \frac{1}{4a-2} \right) \|\mathbf{v}^{r+1} - \mathbf{v}^*(\alpha^r)\|^2 + (1 + 4a - 2) \|\mathbf{v}^*(\alpha^{r+1}) - \mathbf{v}^*(\alpha^r)\|^2 \\
&\leq \left(1 + \frac{1}{4a-2} \right) \|\mathbf{v}^{r+1} - \mathbf{v}^*(\alpha^r)\|^2 + (1 + 4a - 2) \kappa_\Phi^2 \|\alpha^{r+1} - \alpha^r\|^2
\end{aligned} \tag{18}$$

where $a = \frac{1}{\mu\eta NK}$, and last step is due to that $\mathbf{v}^*(\boldsymbol{\alpha})$ is $\kappa_\Phi := \frac{\sqrt{NG}}{\mu}$ Lipschitz, as proven in Lemma 2. Similar to the proof of Theorem 2, we first define

$$\begin{aligned}\mathbf{g}^r &:= \sum_{j=1}^N \prod_{j'=N-1}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r), \\ \boldsymbol{\delta}^r &:= \sum_{j=1}^N \prod_{j'=N-1}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \boldsymbol{\delta}_j.\end{aligned}$$

Then we recall the updating rule of $\mathbf{v}$:

$$\mathbf{v}^{r+1} = \mathcal{P}_{\mathcal{W}}(\mathbf{v}^r - \mathbf{g}^r - \boldsymbol{\delta}^r).$$

Hence we have:

$$\begin{aligned}\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 &= \mathbb{E} \|\mathcal{P}_{\mathcal{W}}(\mathbf{v}^r - \mathbf{g}^r - \boldsymbol{\delta}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r))\|^2 \\ &\leq \mathbb{E} \|\mathbf{v}^r - \mathbf{g}^r - \boldsymbol{\delta}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 \\ &\leq \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 - 2\mathbb{E} \langle \mathbf{g}^r, \mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r) \rangle + \mathbb{E} \|\mathbf{g}^r\|^2 + \mathbb{E} \|\boldsymbol{\delta}^r\|^2 \\ &\leq \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 - 2\mathbb{E} \langle \eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r), \mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r) \rangle \\ &\quad - 2\mathbb{E} \langle \mathbf{g}^r - \eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r), \mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r) \rangle + \mathbb{E} \|\mathbf{g}^r\|^2 + \mathbb{E} \|\boldsymbol{\delta}^r\|^2.\end{aligned}$$

Now, applying strongly convexity of $\Phi(\boldsymbol{\alpha}^r, \cdot)$ and Cauchy-Schwartz inequality yields:

$$\begin{aligned}\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 &\leq (1 - \mu\eta NK) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 - \eta NK \mathbb{E} [\Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r) - \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^*(\hat{\boldsymbol{\alpha}}))] \\ &\quad + \frac{1}{2} \left(\frac{1}{\mu\eta NK} \mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\|^2 + \mu\eta NK \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 \right) \\ &\quad + \mathbb{E} \|\mathbf{g}^r\|^2 + \mathbb{E} \|\boldsymbol{\delta}^r\|^2 \\ &\leq \left(1 - \frac{1}{2} \mu\eta NK \right) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 - \eta NK \mathbb{E} [\Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r) - \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^*(\hat{\boldsymbol{\alpha}}))] \\ &\quad + \frac{1}{2\mu\eta NK} \mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r)\|^2 \\ &\quad + 2\mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r)\|^2 + 2\mathbb{E} \|\eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r)\|^2 + \mathbb{E} \|\boldsymbol{\delta}^r\|^2.\end{aligned}$$

where in the first inequality we applied Cauchy-Schwartz inequality and strongly convexity. Since $\Phi(\boldsymbol{\alpha}^r, \cdot)$ is L smooth, we have: $\mathbb{E} \|\nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r)\|^2 \leq 2L \mathbb{E} [\Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r) - \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^*(\boldsymbol{\alpha}^r))]$. Therefore, it follows that:

$$\begin{aligned}\mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 &\leq \left(1 - \frac{1}{2} \mu\eta NK \right) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 - (\eta NK - 4\eta^2 N^2 K^2 L) \mathbb{E} [\Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r) - \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^*(\boldsymbol{\alpha}^r))] \\ &\quad + \left(\frac{1}{2\mu\eta NK} + 2 \right) \mathbb{E} \|\mathbf{g}^r - \eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r)\|^2 + \mathbb{E} \|\boldsymbol{\delta}^r\|^2\end{aligned}\quad (19)$$

Now, we examine the term $\|\mathbf{g}^r - \eta NK \nabla \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r)\|^2$ in the right hand side of above inequality. First according to summation by part (Lemma 8): we let $\mathbf{A}_j := \prod_{j'=N-1}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'})$ and $\mathbf{B}_j = \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r)$, then we have:

$$\begin{aligned}\mathbf{g}^r &= \sum_{j=1}^N \prod_{j'=N-1}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \\ &= \sum_{j=1}^N \mathbf{A}_j \mathbf{B}_j = \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) - \prod_{j'=N}^{n+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \\ &= \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r).\end{aligned}$$

Hence we have:

$$\begin{aligned}
& \| \mathbf{g}^r - \eta N K \nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r) \|^2 \\
&= \left\| \eta N K \nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r) - \sum_{j=1}^N \prod_{j'=N-1}^{j+1} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\
&= \left\| \eta N K \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) - \left(\sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=0}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right) \right\|^2 \\
&\stackrel{(1)}{\leq} 2 \left\| \left(\eta N K \sum_{j=1}^N \hat{\alpha}(\sigma(j)) \nabla f_{\sigma(j)}(\mathbf{v}^r) - \sum_{j=1}^N \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right) \right\|^2 \\
&\quad + 2 \left\| \sum_{n=1}^{N-1} \left(\prod_{j'=N}^{n+2} (\mathbf{I} - \mathbf{Q}_{j'} \mathbf{H}_{j'}) \right) \mathbf{Q}_{n+1} \mathbf{H}_{n+1} \sum_{j=1}^n \mathbf{Q}_j \nabla f_{\sigma(j)}(\mathbf{v}^r) \right\|^2 \\
&\stackrel{(2)}{\leq} \left(20\eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 + 36e^6 \eta^4 N^4 K^4 L^4 \right) \|\nabla \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r)\|^2 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 G^2 \log(1/p) \\
&\quad + 244e^6 \eta^4 N^4 K^4 L^4 G^2 N \log(1/p) \\
&\stackrel{(3)}{\leq} \left(20\eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 + 36e^6 \eta^4 N^4 K^4 L^4 \right) 2L (\Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r) - \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^*(\hat{\boldsymbol{\alpha}}))) \\
&\quad + \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 \right) G^2 N \log(1/p)
\end{aligned}$$

where in (1) we apply Jensen's inequality, in (2) we plug in Lemma 11 (a), and Lemma 12, and in (3) we use the L -smoothness of Φ . Plugging above bound back in (19) yields:

$$\begin{aligned}
& \mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 \\
&\leq (1 - \frac{1}{2} \mu \eta N K) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 + \eta^2 N^2 K^2 e^4 \delta^2 \\
&\quad - \left(\eta N K - 4\eta^2 N^2 K^2 L - \left(\frac{1}{2\mu \eta N K} + 2 \right) \left(20\eta^2 N^2 K^2 \left(\frac{e}{4R-e} \right)^2 - 36e^6 \eta^4 N^4 K^4 L^4 \right) \right) \\
&\quad \times \mathbb{E} [\Phi(\boldsymbol{\alpha}^r, \mathbf{v}^r) - \Phi(\boldsymbol{\alpha}^r, \mathbf{v}^*(\boldsymbol{\alpha}^r))] \\
&\quad + \left(\frac{1}{2\mu \eta N K} + 2 \right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e} \right)^2 \right) G^2 N \log(1/p).
\end{aligned}$$

Since we choose $\eta = \frac{4 \log(\sqrt{NKR})}{\mu NKR}$, and

$$R \geq \max \left\{ \frac{3}{8} e, \sqrt[3]{\frac{64\kappa^2 \log(\sqrt{NKR}) e^6}{9\mu}} \right\},$$

hence we have:

$$\begin{aligned}
& \mathbb{E} \|\mathbf{v}^{r+1} - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 \\
& \leq \left(1 - \frac{1}{2}\mu\eta NK\right) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 + \eta^2 N^2 K^2 e^4 \delta^2 - \underbrace{\frac{1}{2}\eta NK \mathbb{E}[\Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^r) - \Phi(\hat{\boldsymbol{\alpha}}, \mathbf{v}^*(\hat{\boldsymbol{\alpha}}))]}_{\geq 0} \\
& \quad + \left(\frac{1}{2\mu\eta NK} + 2\right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e}\right)^2\right) G^2 N \log(1/p) \\
& \leq \left(1 - \frac{1}{2}\mu\eta NK\right) \mathbb{E} \|\mathbf{v}^r - \mathbf{v}^*(\hat{\boldsymbol{\alpha}})\|^2 + \eta^2 N^2 K^2 e^4 \delta^2 \\
& \quad + \left(\frac{1}{2\mu\eta NK} + 2\right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e}\right)^2\right) G^2 N \log(1/p).
\end{aligned}$$

Putting above inequality back to (18) yields:

$$\begin{aligned}
\|\mathbf{v}^{r+1} - \mathbf{v}^*(\boldsymbol{\alpha}^{r+1})\|^2 & \leq \left(1 - \frac{1}{4a}\right) \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 + 2\eta^2 N^2 K^2 e^4 \delta^2 + (1 + 4a - 2) \kappa_{\Phi}^2 \|\boldsymbol{\alpha}^{r+1} - \boldsymbol{\alpha}^r\|^2 \\
& \quad + 2 \left(\frac{1}{2\mu\eta NK} + 2\right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e}\right)^2\right) G^2 N \log(1/p) \\
& \leq \left(1 - \frac{1}{4a}\right) \|\mathbf{v}^r - \mathbf{v}^*(\boldsymbol{\alpha}^r)\|^2 + 2\eta^2 N^2 K^2 e^4 \delta^2 \\
& \quad + 2 \left(\frac{1}{2\mu\eta NK} + 2\right) \left(244e^6 \eta^4 N^4 K^4 L^4 + 256\eta^2 N^3 K^2 \left(\frac{e}{4R-e}\right)^2\right) G^2 N \log(1/p) \\
& \quad + O \left(\frac{\kappa_{\Phi}^2}{\mu\eta NK} \left(\left(1 - \frac{1}{\kappa_g}\right)^{T_{\alpha}} + \kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) \left(\gamma^2 L (1 - \mu\gamma)^r DG + \frac{\gamma^3 \kappa \delta^2}{M} + \frac{\gamma^2 \delta^2}{M} \right) \right) \right)
\end{aligned}$$

where at second inequality we plug in Lemma 14. Unrolling the recursion from $r = R$ to 0, and plugging in $\eta = \frac{4 \log(NKR^3)}{\mu NK R}$ yields:

$$\begin{aligned}
& \|\mathbf{v}^R - \mathbf{v}^*(\boldsymbol{\alpha}^R)\|^2 \\
& \leq \left(1 - \frac{1}{4}\mu\eta NK\right)^R \|\mathbf{v}^0 - \mathbf{v}^*(\boldsymbol{\alpha}^0)\|^2 + \frac{1}{\mu} \eta NK e^4 \delta^2 \\
& \quad + 8 \frac{1}{\mu} \left(\frac{1}{2\mu\eta NK} + 2\right) \left(244e^6 \eta^3 N^3 K^3 L^4 + 256\eta^2 N^2 K \left(\frac{e}{4R-e}\right)^2\right) G^2 N \log(1/p) \\
& \quad + O \left(\frac{\kappa_{\Phi}^2}{\mu\eta NK} \sum_{r=0}^R \left(1 - \frac{1}{4a}\right)^{R-r} \left(\left(1 - \frac{1}{\kappa_g}\right)^{T_{\alpha}} + \kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) \left(\gamma^2 L (1 - \mu\gamma)^r DG + \frac{\gamma^3 \kappa \delta^2}{M} + \frac{\gamma^2 \delta^2}{M} \right) \right) \right) \\
& \leq O \left(\frac{\|\mathbf{v}^0 - \mathbf{v}^*(\boldsymbol{\alpha}^0)\|^2}{NK R^3} \right) + \tilde{O} \left(\left(\frac{\kappa^4}{R^2} + \frac{N}{\mu^2 R^2} \right) G^2 N \log(1/p) + \frac{\delta^2}{\mu R} \right) \\
& \quad + \tilde{O} \left(R \kappa_{\Phi}^2 \sum_{r=0}^R \left(1 - \frac{\log(NKR^3)}{R}\right)^{R-r} \left(\left(1 - \frac{1}{\kappa_g}\right)^{T_{\alpha}} + \kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) \left(\gamma^2 L (1 - \mu\gamma)^r DG + \frac{\gamma^3 \kappa \delta^2}{M} + \frac{\gamma^2 \delta^2}{M} \right) \right) \right)
\end{aligned}$$

Plugging in $\gamma = \frac{\log(NKR^3)}{\mu R}$ yields:

$$\begin{aligned} \|\mathbf{v}^R - \mathbf{v}^*(\boldsymbol{\alpha}^R)\|^2 &\leq O\left(\frac{\|\mathbf{v}^0 - \mathbf{v}^*(\boldsymbol{\alpha}^0)\|^2}{NKR^3}\right) + \tilde{O}\left(\left(\frac{\kappa^4}{R^2} + \frac{N}{\mu^2 R^2}\right) G^2 N \log(1/p) + \frac{\delta^2}{\mu R}\right) \\ &\quad + \tilde{O}\left(\kappa_\Phi^2 R \left(\gamma^2 L^2 \sum_{r=0}^R \left(1 - \frac{\log(NKR^3)}{R}\right)^R \kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) DG \right.\right. \\ &\quad \left.\left. + \sum_{r=0}^R \left(1 - \frac{\log(NKR^3)}{R}\right)^{R-r} \left(\left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + \frac{\kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) \gamma^2 \delta^2}{M}\right)\right)\right) \\ &\leq O\left(\frac{D^2}{NKR^3}\right) + \tilde{O}\left(\left(\frac{\kappa^4}{R^2} + \frac{N}{\mu^2 R^2}\right) G^2 N \log(1/p) + \frac{\delta^2}{\mu R}\right) \\ &\quad + \tilde{O}\left(\frac{\kappa_\Phi^2 \kappa^2 \kappa_g^2 L^2 \bar{\zeta}_i(\mathbf{w}^*) DG}{R} + \kappa_\Phi^2 R^2 \left(\left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + \frac{\kappa_g^2 \kappa^2 \bar{\zeta}_i(\mathbf{w}^*) \delta^2}{\mu^2 MR^2}\right)\right) \end{aligned}$$

Since $\hat{\mathbf{v}}_i = \mathbf{v}^R$ and $\hat{\boldsymbol{\alpha}}_i = \boldsymbol{\alpha}^R$, we have the convergence of $\|\hat{\mathbf{v}}_i - \mathbf{v}^*(\hat{\boldsymbol{\alpha}}_i)\|^2$. Plugging this convergence rate together with the convergence of $\|\hat{\boldsymbol{\alpha}}_i - \boldsymbol{\alpha}^*\|^2$ from Lemma 15:

$$\begin{aligned} \|\hat{\boldsymbol{\alpha}}_i - \boldsymbol{\alpha}^*\|^2 &\leq O\left(2\left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + O\left(\kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \frac{2}{\mu} (1 - \mu\gamma)^R + 2\gamma \frac{\delta^2}{\mu^2 M}\right)\right) \\ &\leq \tilde{O}\left(\left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + O\left(\kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) L^2 \frac{2}{\mu} \frac{1}{R} + \frac{\delta^2}{\mu^3 RM}\right)\right) \end{aligned}$$

together with applying Lemma 2 leads to:

$$\begin{aligned} \Phi(\boldsymbol{\alpha}_i^*, \hat{\mathbf{v}}_i) - \Phi(\boldsymbol{\alpha}_i^*, \mathbf{v}_i^*) &\leq 2L \|\hat{\mathbf{v}}_i - \mathbf{v}^*(\hat{\boldsymbol{\alpha}}_i)\|^2 + \left(2\kappa_\Phi^2 L + \frac{4NG^2}{L}\right) \|\hat{\boldsymbol{\alpha}}_i - \boldsymbol{\alpha}_i^*\|^2 \\ &\leq O\left(\frac{LD^2}{NKR^3}\right) + \tilde{O}\left(\left(\frac{\kappa^4 L}{R^2} + \frac{NL}{\mu^2 R^2}\right) G^2 N \log(1/p) + \frac{L\delta^2}{\mu R}\right) \\ &\quad + \tilde{O}\left(\frac{\kappa_\Phi^2 \kappa^2 \kappa_g^2 L^3 \bar{\zeta}_i(\mathbf{w}^*) DG}{R} + \kappa_\Phi^2 LR^2 \left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + \frac{L\kappa_\Phi^2 \kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) \delta^2}{\mu^2 M}\right) \\ &\quad + \left(2\kappa_\Phi^2 L + \frac{4NG^2}{L}\right) \tilde{O}\left(\left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + \frac{\kappa_g^2 \bar{\zeta}_i(\mathbf{w}^*) \kappa L}{R} + \frac{\delta^2}{\mu^3 RM}\right) \\ &\leq O\left(\frac{LD^2}{NKR^3}\right) + \tilde{O}\left(\left(\frac{\kappa^4 L}{R^2} + \frac{NL}{\mu^2 R^2}\right) G^2 N \log(1/p) + \frac{L\delta^2}{\mu R}\right) \\ &\quad + \tilde{O}\left(\frac{\kappa_\Phi^2 \kappa^2 \kappa_g^2 L^3 \bar{\zeta}_i(\mathbf{w}^*) DG}{R} + \kappa_\Phi^2 LR^2 \left(1 - \frac{1}{\kappa_g}\right)^{T_\alpha} + \frac{L^2 \kappa^2 \kappa_g^2 \kappa_\Phi^2 \bar{\zeta}_i(\mathbf{w}^*) \delta^2}{\mu^2 M}\right). \end{aligned}$$

thus completing the proof. $\square$