
FEDNEST: Federated Bilevel, Minimax, and Compositional Optimization

Davoud Ataee Tarzanagh¹ Mingchen Li² Christos Thrampoulidis³ Samet Oymak²

Abstract

Standard federated optimization methods successfully apply to stochastic problems with *single-level* structure. However, many contemporary ML problems – including adversarial robustness, hyperparameter tuning, actor-critic – fall under nested bilevel programming that subsumes minimax and compositional optimization. In this work, we propose FEDNEST: A federated alternating stochastic gradient method to address general nested problems. We establish provable convergence rates for FEDNEST in the presence of heterogeneous data and introduce variations for bilevel, minimax, and compositional optimization. FEDNEST introduces multiple innovations including federated hypergradient computation and variance reduction to address inner-level heterogeneity. We complement our theory with experiments on hyperparameter & hyper-representation learning and minimax optimization that demonstrate the benefits of our method in practice.

1. Introduction

In the federated learning (FL) paradigm, multiple clients cooperate to learn a model under the orchestration of a central server (McMahan et al., 2017) without directly exchanging local client data with the server or other clients. The locality of data distinguishes FL from traditional distributed optimization and also motivates new methodologies to address heterogeneous data across clients. Additionally, cross-device FL across many edge devices presents additional challenges since only a small fraction of clients participate in each round, and clients cannot maintain *state* across rounds (Kairouz et al., 2019).

Traditional distributed SGD methods are often unsuitable in

¹Email: tarzanaq@umich.edu, University of Michigan

²Emails: {mlil76@, oymak@ece.}ucr.edu, University of California, Riverside ³Email: cthrampo@ece.ubc.ca, University of British Columbia. Correspondence to: Davoud Ataee Tarzanagh <tarzanaq@umich.edu>.

FL and incur high communication costs. To overcome this issue, popular FL methods, such as FEDAVG (McMahan et al., 2017), use local client updates, i.e. clients update their models multiple times before communicating with the server (aka, local SGD). Although FEDAVG has seen great success, recent works have exposed convergence issues in certain settings (Karimireddy et al., 2020; Hsu et al., 2019). This is due to a variety of factors, including *client drift*, where local models move away from globally optimal models due to objective and/or systems heterogeneity.

Existing FL methods, such as FEDAVG, are widely applied to stochastic problems with *single-level* structure. Instead, many machine learning tasks – such as adversarial learning (Madry et al., 2017), meta learning (Bertinetto et al., 2018), hyperparameter optimization (Franceschi et al., 2018), reinforcement/imitation learning (Wu et al., 2020; Arora et al., 2020), and neural architecture search (Liu et al., 2018) – admit *nested* formulations that go beyond the standard single-level structure. Towards addressing such nested problems, bilevel optimization has received significant attention in the recent literature (Ghadimi & Wang, 2018; Hong et al., 2020; Ji et al., 2021); albeit in non-FL settings. On the other hand, federated versions have been elusive perhaps due to the additional challenges surrounding heterogeneity, communication, and inverse Hessian approximation.

Contributions: This paper addresses these challenges and develops **FEDNEST**: A federated machinery for nested problems with provable convergence and lightweight communication. FEDNEST is composed of **FEDINN**: a federated stochastic variance reduction algorithm (FEDSVRG) to solve the inner problem while avoiding client drift, and **FEDOUT**: a communication-efficient federated hypergradient algorithm for solving the outer problem. Importantly, we allow both inner & outer objectives to be finite sums over heterogeneous client functions. FEDNEST runs a variant of FEDSVRG on inner & outer variables in an alternating fashion as outlined in Algorithm 1. We make multiple algorithmic and theoretical contributions summarized below.

- **The variance reduction** of FEDINN enables robustness in the sense that local models converge to the globally optimal inner model despite client drift/heterogeneity unlike FEDAVG. While FEDINN is similar to FEDSVRG (Konečný et al., 2018) and FEDLIN (Mittra et al., 2021), we make two key contributions: (i) We

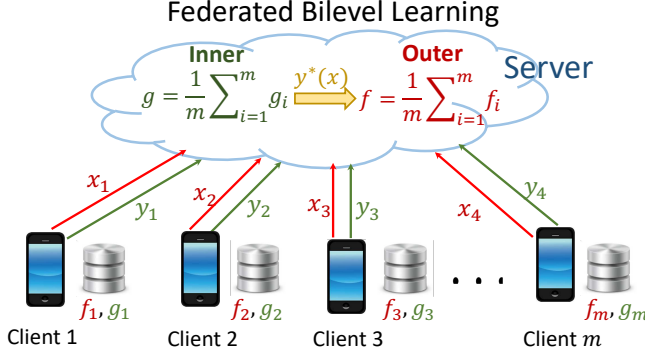


Figure 1: Depiction of federated bilevel nested optimization and high-level summary of FEDNEST (Algorithm 1). At outer loop, FEDIHGP uses multiple rounds of matrix-vector products to facilitate hypergradient computation while only communicating vectors. At inner loop, FEDINN uses FEDSVRG to avoid client drift and find the unique global minima. Both are crucial for establishing provable convergence of FEDNEST.

leverage the global convergence of FEDINN to ensure accurate hypergradient computation which is crucial for our bilevel proof. (ii) We establish new convergence guarantees for single-level stochastic non-convex FEDSVRG, which are then integrated within our FEDOUT.

- **Communication efficient bilevel optimization:** Within FEDOUT, we develop an efficient federated method for hypergradient estimation that bypass Hessian computation. Our approach approximates the global Inverse Hessian-Gradient-Product (IHGP) via computation of matrix-vector products over few communication rounds.
- **LFEDNEST:** To further improve communication efficiency, we additionally propose a Light-FEDNEST algorithm, which computes hypergradients locally and only needs a single communication round for the outer update. Experiments reveal that LFEDNEST becomes very competitive as client functions become more homogeneous.
- **Unified federated nested theory:** We specialize our bilevel results to *minimax* and *compositional* optimization with emphasis on the former. For these, FEDNEST significantly simplifies and leads to faster convergence. Importantly, our results are on par with the state-of-the-art non-federated guarantees for nested optimization literature without additional assumptions (Table 1).
- We provide extensive **numerical experiments**¹ on bilevel and minimax optimization problems. These demonstrate the benefits of FEDNEST, efficiency of LFEDNEST, and shed light on tradeoffs surrounding communication, computation, and heterogeneity.

2. Federated Nested Problems & FEDNEST

We will first provide the background on bilevel nested problems and then introduce our general federated method.

¹FEDNEST code is available at <https://github.com/ucr-optml/FEDNEST>.

Communication-efficiency:

- ✓ FEDIHGP avoids explicit Hessian
- ✓ LFEDNEST for local hypergradients

Client heterogeneity:

- ✓ FEDINN avoids client drift

Finite sample bilevel theory:

- ✓ Stochastic inner & outer analysis

Specific nested optimization problems:

- ✓ Bilevel
- ✓ Minimax
- ✓ Compositional

Stochastic Bilevel Optimization

	FEDNEST	Non-Federated		
		ALSET	BSA	TTSA
batch size		$\mathcal{O}(1)$		
samples in ξ	$\mathcal{O}(\kappa_g^5 \epsilon^{-2})$	$\mathcal{O}(\kappa_g^5 \epsilon^{-2})$	$\mathcal{O}(\kappa_g^6 \epsilon^{-2})$	$\mathcal{O}(\kappa_g^p \epsilon^{-2.5})$
samples in ζ	$\mathcal{O}(\kappa_g^9 \epsilon^{-2})$	$\mathcal{O}(\kappa_g^9 \epsilon^{-2})$	$\mathcal{O}(\kappa_g^9 \epsilon^{-3})$	$\mathcal{O}(\kappa_g^p \epsilon^{-2.5})$

Stochastic Minimax Optimization

	FEDNEST	Non-Federated		
		ALSET	SGDA	SMD
batch size	$\mathcal{O}(1)$	$\mathcal{O}(1)$	$\mathcal{O}(\epsilon^{-1})$	N.A.
samples	$\mathcal{O}(\kappa_f^3 \epsilon^{-2})$			

Stochastic Compositional Optimization

	FEDNEST	Non-Federated		
		ALSET	SCGD	NASA
batch size		$\mathcal{O}(1)$		
samples	$\mathcal{O}(\epsilon^{-2})$	$\mathcal{O}(\epsilon^{-2})$	$\mathcal{O}(\epsilon^{-4})$	$\mathcal{O}(\epsilon^{-2})$

Table 1: Sample complexity of FEDNEST and comparable non-FL methods to find an ϵ -stationary point of f : $\kappa_g := \ell_{g,1}/\mu_g$ and $\kappa_f := \ell_{f,1}/\mu_f$. κ_g^p denotes a polynomial function of κ_g . ALSET (Chen et al., 2021a), BSA (Ghadimi & Wang, 2018), TTSA (Hong et al., 2020), SGDA (Lin et al., 2020), SMD (Rafique et al., 2021), SCGD (Wang et al., 2017), and NASA (Ghadimi et al., 2020).

Notation. For a differentiable function $h(\mathbf{x}, \mathbf{y}) : \mathbb{R}^{d_1} \times \mathbb{R}^{d_2} \rightarrow \mathbb{R}$ in which $\mathbf{y} = \mathbf{y}(\mathbf{x}) : \mathbb{R}^{d_1} \rightarrow \mathbb{R}^{d_2}$, we denote $\nabla h \in \mathbb{R}^{d_1}$ the gradient of h as a function of $\mathbf{x}$ and $\nabla_{\mathbf{x}} h$, $\nabla_{\mathbf{y}} h$ the partial derivatives of h with respect to $\mathbf{x}$ and $\mathbf{y}$, respectively. We let $\nabla_{\mathbf{x}\mathbf{y}}^2 h$ and $\nabla_{\mathbf{y}}^2 h$ denote the Jacobian and Hessian of h , respectively. We consider FL optimization over m clients and we denote $\mathcal{S} = \{1, \dots, m\}$. For vectors $\mathbf{v} \in \mathbb{R}^d$ and matrix $\mathbf{M} \in \mathbb{R}^{d \times d}$, we denote $\|\mathbf{v}\|$ and $\|\mathbf{M}\|$ the respective Euclidean and spectral norms.

2.1. Preliminaries on Federated Nested Optimization

In **federated bilevel learning**, we consider the following nested optimization problem as depicted in Figure 1:

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}^{d_1}} \quad & f(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) \\ \text{subj. to} \quad & \mathbf{y}^*(\mathbf{x}) \in \operatorname{argmin}_{\mathbf{y} \in \mathbb{R}^{d_2}} \frac{1}{m} \sum_{i=1}^m g_i(\mathbf{x}, \mathbf{y}). \end{aligned} \quad (1a)$$

Algorithm 1 FEDNEST

```

1: Inputs:  $K, T \in \mathbb{N}$ ;  $(\mathbf{x}^0, \mathbf{y}^0) \in \mathbb{R}^{d_1+d_2}$ ; FEDINN,
2:   FEDOUT with stepsizes  $\{(\alpha^k, \beta^k)\}_{k=0}^{K-1}$ 
3: for  $k = 0, \dots, K-1$  do
4:    $\mathbf{y}^{k,0} = \mathbf{y}^k$ 
5:   for  $t = 0, \dots, T-1$  do
6:      $\mathbf{y}^{k,t+1} = \text{FEDINN}(\mathbf{x}^k, \mathbf{y}^{k,t}, \beta^k)$ 
7:   end for
8:    $\mathbf{y}^{k+1} = \mathbf{y}^{k,T}$ 
9:    $\mathbf{x}^{k+1} = \text{FEDOUT}(\mathbf{x}^k, \mathbf{y}^{k+1}, \alpha^k)$ 
10: end for
    
```

Recall that m is the number of clients. Here, to model objective heterogeneity, each client i is allowed to have its own individual outer & inner functions (f_i, g_i) . Moreover, we consider a general stochastic oracle model, access to local functions (f_i, g_i) is via stochastic sampling as follows:

$$\begin{aligned} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) &:= \mathbb{E}_{\xi \sim \mathcal{C}_i} [f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x}); \xi)], \\ g_i(\mathbf{x}, \mathbf{y}) &:= \mathbb{E}_{\zeta \sim \mathcal{D}_i} [g_i(\mathbf{x}, \mathbf{y}; \zeta)], \end{aligned} \quad (1b)$$

where $(\xi, \zeta) \sim (\mathcal{C}_i, \mathcal{D}_i)$ are outer/inner sampling distributions for the i^{th} client. We emphasize that for $i \neq j$, the tuples $(f_i, g_i, \mathcal{C}_i, \mathcal{D}_i)$ and $(f_j, g_j, \mathcal{C}_j, \mathcal{D}_j)$ can be different.

Example 1 (Hyperparameter tuning). *Each client has local validation and training datasets associated with objectives $(f_i, g_i)_{i=1}^m$ corresponding to validation and training losses, respectively. The goal is finding hyper-parameters $\mathbf{x}$ that lead to learning model parameters $\mathbf{y}$ that minimize the (global) validation loss.*

The stochastic bilevel problem (1) subsumes two popular problem classes with the nested structure: *Stochastic Mini-Max & Stochastic Compositional*. Therefore, results on the general nested problem (1) also imply the results in these special cases. Below, we briefly describe them.

Minimax optimization. If $g_i(\mathbf{x}, \mathbf{y}; \zeta) := -f_i(\mathbf{x}, \mathbf{y}; \xi)$ for all $i \in \mathcal{S}$, the stochastic bilevel problem (1) reduces to the stochastic minimax problem

$$\min_{\mathbf{x} \in \mathbb{R}^{d_1}} f(\mathbf{x}) := \frac{1}{m} \max_{\mathbf{y} \in \mathbb{R}^{d_2}} \sum_{i=1}^m \mathbb{E}[f_i(\mathbf{x}, \mathbf{y}; \xi)]. \quad (2)$$

Motivated by applications in fair beamforming, training generative-adversarial networks (GANs) and robust machine learning, significant efforts have been made for solving (2) including (Daskalakis & Panageas, 2018; Gidel et al., 2018; Mokhtari et al., 2020; Thekumparampil et al., 2019).

Example 2 (GANs). *We train a generative model $g_{\mathbf{x}}(\cdot)$ and an adversarial model $a_{\mathbf{y}}(\cdot)$ using client datasets $\mathcal{C}_i$. The local functions may for example take the form $f_i(\mathbf{x}, \mathbf{y}) = \mathbb{E}_{s \sim \mathcal{C}_i} \{\log a_{\mathbf{y}}(s)\} + \mathbb{E}_{z \sim \mathcal{D}_{\text{noise}}} \{\log[1 - a_{\mathbf{y}}(g_{\mathbf{x}}(z))]\}$.*

Compositional optimization. Suppose $f_i(\mathbf{x}, \mathbf{y}; \xi) := f_i(\mathbf{y}; \xi)$ and g_i is quadratic in $\mathbf{y}$ given as $g_i(\mathbf{x}, \mathbf{y}; \zeta) := \|\mathbf{y} - \mathbf{r}_i(\mathbf{x}; \zeta)\|^2$. Then, the bilevel problem (1) reduces to

$$\begin{aligned} \min_{\mathbf{x} \in \mathbb{R}^{d_1}} \quad & f(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m f_i(\mathbf{y}^*(\mathbf{x})) \\ \text{subj. to} \quad & \mathbf{y}^*(\mathbf{x}) = \operatorname{argmin}_{\mathbf{y} \in \mathbb{R}^{d_2}} \frac{1}{m} \sum_{i=1}^m g_i(\mathbf{x}, \mathbf{y}) \end{aligned} \quad (3)$$

with $f_i(\mathbf{y}^*(\mathbf{x})) := \mathbb{E}_{\xi \sim \mathcal{C}_i} [f_i(\mathbf{y}^*(\mathbf{x}); \xi)]$ and $g_i(\mathbf{x}, \mathbf{y}) := \mathbb{E}_{\zeta \sim \mathcal{D}_i} [g_i(\mathbf{x}, \mathbf{y}; \zeta)]$. Optimization problems in the form of (3) occur for example in model agnostic meta-learning and policy evaluation in reinforcement learning (Finn et al., 2017; Ji et al., 2020b; Dai et al., 2017; Wang et al., 2017).

Assumptions. Let $\mathbf{z} = (\mathbf{x}, \mathbf{y}) \in \mathbb{R}^{d_1+d_2}$. Throughout, we make the following assumptions on inner/outer objectives.

Assumption A (Well-behaved objectives). *For all $i \in [m]$:*

- (A1) $f_i(\mathbf{z}), \nabla f_i(\mathbf{z}), \nabla g_i(\mathbf{z}), \nabla^2 g_i(\mathbf{z})$ are $\ell_{f,0}, \ell_{f,1}, \ell_{g,1}, \ell_{g,2}$ -Lipschitz continuous, respectively; and
- (A2) $g_i(\mathbf{x}, \mathbf{y})$ is μ_g -strongly convex in $\mathbf{y}$ for all $\mathbf{x} \in \mathbb{R}^{d_1}$.

Throughout, we use $\kappa_g = \ell_{g,1}/\mu_g$ to denote the condition number of the inner function g .

Assumption B (Stochastic samples). *For all $i \in [m]$:*

- (B1) $\nabla f_i(\mathbf{z}; \xi), \nabla g_i(\mathbf{z}; \zeta), \nabla^2 g_i(\mathbf{z}; \zeta)$ are unbiased estimators of $\nabla f_i(\mathbf{z}), \nabla g_i(\mathbf{z}), \nabla^2 g_i(\mathbf{z})$, respectively; and
- (B2) Their variances are bounded, i.e., $\mathbb{E}_{\xi} [\|\nabla f_i(\mathbf{z}; \xi) - \nabla f_i(\mathbf{z})\|^2] \leq \sigma_f^2$, $\mathbb{E}_{\zeta} [\|\nabla^2 g_i(\mathbf{z}; \zeta) - \nabla^2 g_i(\mathbf{z})\|^2] \leq \sigma_{g,1}^2$, and $\mathbb{E}_{\zeta} [\|\nabla g_i(\mathbf{z}; \zeta) - \nabla g_i(\mathbf{z})\|^2] \leq \sigma_{g,2}^2$ for some $\sigma_f^2, \sigma_{g,1}^2$, and $\sigma_{g,2}^2$.

These assumptions are common in the bilevel optimization literature (Ghadimi & Wang, 2018; Chen et al., 2021a; Ji et al., 2021). Assumption A requires that the inner and outer functions are well-behaved. Specifically, strong-convexity of the inner objective is a recurring assumption in bilevel optimization theory implying a unique solution to the inner minimization in (1).

2.2. Proposed Algorithm: FEDNEST

In this section, we develop FEDNEST, which is formally presented in Algorithm 1. The algorithm operates in two nested loops. The outer loop operates in rounds $k \in \{1, \dots, K\}$. Within each round, an inner loop operating for T iterations is executed. Given estimates $\mathbf{x}^k$ and $\mathbf{y}^k$, each iteration $t \in \{1, \dots, T\}$ of the inner loop produces a new *global* model $\mathbf{y}^{k,t+1}$ of the inner optimization variable $\mathbf{y}^*(\mathbf{x}^k)$ as the output of an optimizer FEDINN. The final estimate $\mathbf{y}^{k+1} = \mathbf{y}^{k,T}$ of the inner variable is then used by an optimizer FEDOUT to update the outer *global* model $\mathbf{x}^{k+1}$.

The subroutines FEDINN and FEDOUT are gradient-based optimizers. Each subroutine involves a certain number of

local training steps indexed by $\nu \in \{0, \dots, \tau_i - 1\}$ that are performed at the i^{th} client. The local steps of FEDINN iterate over *local* models $\mathbf{y}_{i,\nu}$ of the inner variable. Accordingly, FEDOUT iterates over *local* models $\mathbf{x}_{i,\nu}$ of the global variable. A critical component of FEDOUT is a communication-efficient federated hypergradient estimation routine, which we call FEDIHGP. The implementation of FEDINN, FEDOUT and FEDIHGP is critical to circumvent the algorithmic challenges of federated bilevel optimization. In the remaining of this section, we detail the challenges and motivate our proposed implementations. Later, in Section 3, we provide a formal convergence analysis of FEDNEST.

2.3. Key Challenge: Federated Hypergradient Estimation

FEDOUT is a gradient-based optimizer for the outer minimization in (1); thus each iteration involves computing $\nabla f(\mathbf{x}) = (1/m) \sum_{i=1}^m \nabla f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))$. Unlike single-level FL, the fact that the outer objective f depends explicitly on the inner minimizer $\mathbf{y}^*(\mathbf{x})$ introduces a new challenge. A good starting point to understand the challenge is the following evaluation of $\nabla f(\mathbf{x})$ in terms of partial derivatives. The result is well-known from properties of implicit functions.

Lemma 2.1. *Under Assumption A, for all $i \in [m]$:*

$$\nabla f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) = \nabla^D f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) + \nabla^I f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})),$$

where the direct and indirect gradient components are:

$$\nabla^D f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) := \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})), \quad (4a)$$

$$\begin{aligned} \nabla^I f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) &:= -\nabla_{\mathbf{x}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) \\ &\quad \cdot [\nabla_{\mathbf{y}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))]^{-1} \nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})). \end{aligned} \quad (4b)$$

We now use the above formula to describe the two core challenges of bilevel FL optimization.

First, evaluation of any of the terms in (4) requires access to the minimizer $\mathbf{y}^*(\mathbf{x})$ of the inner problem. On the other hand, one may at best hope for a good *approximation* to $\mathbf{y}^*(\mathbf{x})$ produced by the inner optimization subroutine. Of course, this challenge is inherent in any bilevel optimization setting, but is exacerbated in the FL setting because of *client drift*. Specifically, when clients optimize their individual (possibly different) local inner objectives, the global estimate of the inner variable produced by SGD-type methods may drift far from (a good approximation to) $\mathbf{y}^*(\mathbf{x})$. We explain in Section 2.5 how FEDINN solves that issue.

The second challenge comes from the stochastic nature of the problem. Observe that the indirect component in (4b) is nonlinear in the Hessian $\nabla_{\mathbf{x}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))$, complicating an unbiased stochastic approximation of $\nabla f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))$. As we expose here, solutions to this complication developed in the non-federated bilevel optimization literature, are

not directly applicable in the FL setting. Indeed, existing stochastic bilevel algorithms, e.g. (Ghadimi & Wang, 2018), define $\bar{\nabla} f(\mathbf{x}, \mathbf{y}) := \bar{\nabla}^D f(\mathbf{x}, \mathbf{y}) + \bar{\nabla}^I f(\mathbf{x}, \mathbf{y})$ as a surrogate of $\nabla f(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))$ by replacing $\mathbf{y}^*(\mathbf{x})$ in definition (4) with an approximation $\mathbf{y}$ and using the following stochastic approximations:

$$\bar{\nabla}^D f(\mathbf{x}, \mathbf{y}) \approx \nabla_{\mathbf{x}} f(\mathbf{x}, \mathbf{y}; \dot{\xi}), \quad (5a)$$

$$\bar{\nabla}^I f(\mathbf{x}, \mathbf{y}) \approx -\nabla_{\mathbf{x}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}; \zeta_{N'+1})$$

$$\left[\frac{N}{\ell_{g,1}} \prod_{n=1}^{N'} \left(\mathbf{I} - \frac{1}{\ell_{g,1}} \nabla_{\mathbf{y}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}; \zeta_n) \right) \right] \nabla_{\mathbf{y}} f(\mathbf{x}, \mathbf{y}; \dot{\xi}). \quad (5b)$$

Here, N' is drawn from $\{0, \dots, N-1\}$ uniformly at random (UAR) and $\{\dot{\xi}, \zeta_1, \dots, \zeta_{N'+1}\}$ are i.i.d. samples. Ghadimi & Wang (2018); Hong et al. (2020) have shown that using (5), the inverse Hessian estimation bias exponentially decreases with the number of samples N .

One might hope to directly leverage the above approach in a local computation fashion by replacing the global outer function f with the individual function f_i . However, note from (4b) and (5b) that the proposed stochastic approximation of the indirect gradient involves in a nonlinear way the *global* Hessian, which is *not* available at the client². Communication efficiency is one of the core objectives of FL making the idea of communicating Hessians between clients and server prohibitive. *Is it then possible, in a FL setting, to obtain an accurate stochastic estimate of the indirect gradient while retaining communication efficiency?* In Section 2.4, we show how FEDOUT and its subroutine FEDIHGP, a matrix-vector products-based (thus, communication efficient) federated hypergradient estimator, answer this question affirmatively.

2.4. Outer Optimizer: FEDOUT

This section presents the outer optimizer FEDOUT, formally described in Algorithm 2. As a subroutine of FEDNEST (see Line 9, Algorithm 1), at each round $k = 0, \dots, K-1$, FEDOUT takes the most recent global outer model $\mathbf{x}^k$ together the updated (by FEDINN) global inner model $\mathbf{y}^{k+1}$ and produces an update $\mathbf{x}^{k+1}$. To lighten notation, for a round k , denote the function's input as $(\mathbf{x}, \mathbf{y}^+)$ (instead of $(\mathbf{x}^k, \mathbf{y}^{k+1})$) and the output as $\mathbf{x}^+$ (instead of $\mathbf{x}^{k+1}$). For each client $i \in \mathcal{S}$, FEDOUT uses stochastic approximations of $\bar{\nabla}^I f_i(\mathbf{x}, \mathbf{y}^+)$ and $\bar{\nabla}^D f_i(\mathbf{x}, \mathbf{y}^+)$, which we call $\mathbf{h}_i^I(\mathbf{x}, \mathbf{y}^+)$ and $\mathbf{h}_i^D(\mathbf{x}, \mathbf{y}^+)$, respectively. The specific choice of these approximations (see Line 5) is critical and is discussed in detail later in this section. Before that, we explain

²We note that the approximation in (5) is not the only construction, and bilevel optimization can accommodate other forms of gradient surrogates (Ji et al., 2021). Yet, all these approximations require access (in a nonlinear fashion) to the *global* Hessian; thus, they suffer from the same challenge in FL setting.

Algorithm 2 $\mathbf{x}^+ = \text{FEDOUT}(\mathbf{x}, \mathbf{y}^+, \alpha)$ for stochastic bilevel and minimax problems

```

1:  $\mathbb{F}_i(\cdot) \leftarrow \nabla_{\mathbf{x}} f_i(\cdot, \mathbf{y}^+; \cdot)$ 
2:  $\mathbf{x}_{i,0} = \mathbf{x}$  and  $\alpha_i \in (0, \alpha]$ 
3: Choose  $N \geq 1$  and set  $\mathbf{p}_{N'} = \text{FEDIHGP}(\mathbf{x}, \mathbf{y}^+, N)$ 
4: for  $i \in \mathcal{S}$  in parallel do
5:    $\mathbf{h}_i = \mathbb{F}_i(\mathbf{x}; \xi_i) - \nabla_{\mathbf{x}\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}^+; \zeta_i) \mathbf{p}_{N'}$ 
6:    $\mathbf{h}_i = \mathbb{F}_i(\mathbf{x}; \xi_i)$ 
7: end for
8:  $\mathbf{h} = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{h}_i$ 
9: for  $i \in \mathcal{S}$  in parallel do
10:   for  $\nu = 0, \dots, \tau_i - 1$  do
11:      $\mathbf{h}_{i,\nu} = \mathbb{F}_i(\mathbf{x}_{i,\nu}; \xi_{i,\nu}) - \mathbb{F}_i(\mathbf{x}; \xi_{i,\nu}) + \mathbf{h}$ 
12:      $\mathbf{x}_{i,\nu+1} = \mathbf{x}_{i,\nu} - \alpha_i \mathbf{h}_{i,\nu}$ 
13:   end for
14: end for
15:  $\mathbf{x}^+ = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{x}_{i,\tau_i}$ 
    
```

how each client uses these proxies to form local updates of the outer variable. In each round, starting from a common global model $\mathbf{x}_{i,0} = \mathbf{x}$, each client i performs τ_i local steps (in parallel):

$$\mathbf{x}_{i,\nu+1} = \mathbf{x}_{i,\nu} - \alpha_i \mathbf{h}_{i,\nu}, \quad (6)$$

and then the server aggregates local models via $\mathbf{x}^+ = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{x}_{i,\tau_i}$. Here, $\alpha_i \in (0, \alpha]$ is the local stepsize,

$$\begin{aligned} \mathbf{h}_{i,\nu} := & \mathbf{h}^\top(\mathbf{x}, \mathbf{y}^+) + \mathbf{h}^\mathbf{D}(\mathbf{x}, \mathbf{y}^+) \\ & - \mathbf{h}_i^\mathbf{D}(\mathbf{x}, \mathbf{y}^+) + \mathbf{h}_i^\mathbf{D}(\mathbf{x}_{i,\nu}, \mathbf{y}^+), \end{aligned} \quad (7)$$

and, $\mathbf{h}(\mathbf{x}, \mathbf{y}) := |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{h}_i(\mathbf{x}, \mathbf{y}) = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} (\mathbf{h}_i^\mathbf{D}(\mathbf{x}, \mathbf{y}^+) - \mathbf{h}_i^\top(\mathbf{x}, \mathbf{y}^+))$.

The key features of updates (6)–(7) are exploiting past gradients (variance reduction) to account for objective heterogeneity. Indeed, the ideal update in FEDOUT would perform the update $\mathbf{x}_{i,\nu+1} = \mathbf{x}_{i,\nu} - \alpha_i (\mathbf{h}^\top(\mathbf{x}_{i,\nu}, \mathbf{y}^+) + \mathbf{h}^\mathbf{D}(\mathbf{x}_{i,\nu}, \mathbf{y}^+))$ using the global gradient estimates. But this requires each client i to have access to both direct and indirect gradients of all other clients—which it does not, since clients do not communicate between rounds. To overcome this issue, each client i uses global gradient estimates, i.e., $\mathbf{h}^\top(\mathbf{x}, \mathbf{y}^+) + \mathbf{h}^\mathbf{D}(\mathbf{x}, \mathbf{y}^+)$ from the beginning of each round as a guiding direction in its local update rule. However, since both $\mathbf{h}^\mathbf{D}$ and $\mathbf{h}^\top$ are computed at a previous $(\mathbf{x}, \mathbf{y}^+)$, client i makes a correction by subtracting off the stale direct gradient estimate $\mathbf{h}_i^\mathbf{D}(\mathbf{x}, \mathbf{y}^+)$ and adding its own local estimate $\mathbf{h}_i^\mathbf{D}(\mathbf{x}_{i,\nu}, \mathbf{y}^+)$. Our local update rule in Step 11 of Algorithm 2 is precisely of this form, i.e., $\mathbf{h}_{i,\nu}$ approximates $\mathbf{h}^\top(\mathbf{x}_{i,\nu}, \mathbf{y}^+) + \mathbf{h}^\mathbf{D}(\mathbf{x}_{i,\nu}, \mathbf{y}^+)$ via (7). Note here that the described local correction of FEDOUT only applies

Algorithm 3 $\mathbf{p}_{N'} = \text{FEDIHGP}(\mathbf{x}, \mathbf{y}^+, N)$: Federated approximation of inverse-Hessian-gradient product

```

1: Select  $N' \in \{0, \dots, N-1\}$  UAR.
2: Select  $\mathcal{S}_0 \in \mathcal{S}$  UAR.
3: for  $i \in \mathcal{S}_0$  in parallel do
4:    $\mathbf{p}_{i,0} = \nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}^+; \xi_{i,0})$ 
5: end for
6:  $\mathbf{p}_0 = \frac{N}{\ell_{g,1}} |\mathcal{S}_0|^{-1} \sum_{i \in \mathcal{S}_0} \mathbf{p}_{i,0}$ 
7: if  $N' = 0$  then
8:   Return  $\mathbf{p}_{N'}$ 
9: end if
10: Select  $\mathcal{S}_1, \dots, \mathcal{S}_{N'} \in \mathcal{S}$  UAR.
11: for  $n = 1, \dots, N'$  do
12:   for  $i \in \mathcal{S}_n$  in parallel do
13:      $\mathbf{p}_{i,n} = \left( \mathbf{I} - \frac{1}{\ell_{g,1}} \nabla_{\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}^+; \zeta_{i,n}) \right) \mathbf{p}_{i,n-1}$ 
14:   end for
15:    $\mathbf{p}_n = |\mathcal{S}_n|^{-1} \sum_{i \in \mathcal{S}_n} \mathbf{p}_{i,n}$ 
16: end for
    
```

to the direct gradient component (the indirect component would require global Hessian information). An alternative approach leading to LFEDNEST is discussed in Section 2.6.

FEDOUT applied to special nested problems. Algorithm 2 naturally allows the use of other optimizers for minimax & compositional optimization. For example, in the minimax problem (2), the bilevel gradient components are $\nabla^\mathbf{D} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) = \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))$ and $\nabla^\top f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) = 0$ for all $i \in \mathcal{S}$. Hence, the hyper-gradient estimate (7) reduces to

$$\mathbf{h}_{i,\nu} = \mathbf{h}^\mathbf{D}(\mathbf{x}, \mathbf{y}^+) - \mathbf{h}_i^\mathbf{D}(\mathbf{x}, \mathbf{y}^+) + \mathbf{h}_i^\mathbf{D}(\mathbf{x}_{i,\nu}, \mathbf{y}^+). \quad (8)$$

For the compositional problem (3), Hessian becomes the identity matrix, the direct gradient is the zero vector, and $\nabla_{\mathbf{x}\mathbf{y}} g(\mathbf{x}, \mathbf{y}) = -(1/m) \sum_{i=1}^m \nabla \mathbf{r}_i(\mathbf{x})^\top$. Hence, $\mathbf{h}_i = \ell_{g,1} \nabla \mathbf{r}_i(\mathbf{x})^\top \mathbf{p}_0$ for all $i \in \mathcal{S}$.

More details on these special cases are provided in Appendices D and E.

Indirect gradient estimation & FEDIHGP. Here, we aim to address one of the key challenges in nested FL: inverse Hessian gradient product. Note from (5b) that the proposed stochastic approximation of the indirect gradient involves in a nonlinear way the *global* Hessian, which is *not* available at the client. To get around this, we use a client sampling strategy and recursive reformulation of (5b) so that $\bar{\nabla}^\top f_i(\mathbf{x}, \mathbf{y})$ can be estimated in an efficient federated manner. In particular, given $N \in \mathbb{N}$, we select $N' \in \{0, \dots, N-1\}$ and $\mathcal{S}_0, \dots, \mathcal{S}_{N'} \in \mathcal{S}$ UAR. For all $i \in \mathcal{S}$, we then define

$$\mathbf{h}_i^\top(\mathbf{x}, \mathbf{y}) = -\nabla_{\mathbf{x}\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}; \zeta_i) \mathbf{p}_{N'}, \quad (9a)$$

where $\mathbf{p}_{N'} = |\mathcal{S}_0|^{-1} \widehat{\mathbf{H}}_{\mathbf{y}} \sum_{i \in \mathcal{S}_0} \nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}; \xi_{i,0})$ and $\widehat{\mathbf{H}}_{\mathbf{y}}$

is the approximate inverse Hessian:

$$\frac{N}{\ell_{g,1}} \prod_{n=1}^{N'} \left(\mathbf{I} - \frac{1}{\ell_{g,1} |\mathcal{S}_n|} \sum_{i=1}^{|\mathcal{S}_n|} \nabla_{\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}; \zeta_{i,n}) \right). \quad (9b)$$

The subroutine FEDIHGP provides a recursive strategy to compute $\mathbf{p}_{N'}$ and FEDOUT multiplies $\mathbf{p}_{N'}$ with the global Jacobian to drive an indirect gradient estimate. Importantly, these approximations require only matrix-vector products and vector communications.

Lemma 2.2. *Under Assumptions A and B, the approximate inverse Hessian $\widehat{\mathbf{H}}_{\mathbf{y}}$ defined in (9b) satisfies the following for any $\mathbf{x}$ and $\mathbf{y}$:*

$$\begin{aligned} \left\| [\nabla_{\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y})]^{-1} - \mathbb{E}_{\mathcal{W}}[\widehat{\mathbf{H}}_{\mathbf{y}}] \right\| &\leq \frac{1}{\mu_g} \left(\frac{\kappa_g - 1}{\kappa_g} \right)^N, \\ \mathbb{E}_{\mathcal{W}} \left[\left\| [\nabla_{\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y})]^{-1} - \widehat{\mathbf{H}}_{\mathbf{y}} \right\| \right] &\leq \frac{2}{\mu_g}. \end{aligned} \quad (10)$$

Here, $\mathcal{W} := \{\mathcal{S}_n, \xi_i, \zeta_i, \xi_{i,0}, \zeta_{i,n} \mid i \in \mathcal{S}_n, 0 \leq n \leq N'\}$. Further, for all $i \in \mathcal{S}$, $\mathbf{h}_i^T(\mathbf{x}, \mathbf{y})$ defined in (9a) satisfies

$$\left\| \mathbb{E}_{\mathcal{W}}[\mathbf{h}_i^T(\mathbf{x}, \mathbf{y})] - \bar{\nabla}^T f_i(\mathbf{x}, \mathbf{y}) \right\| \leq b, \quad (11)$$

where $b := \kappa_g \ell_{f,1} ((\kappa_g - 1)/\kappa_g)^N$.

2.5. Inner Optimizer: FEDINN

In FL, each client performs multiple local training steps in isolation on its own data (using for example SGD) before communicating with the server. Due to such local steps, FEDAVG suffers from a *client-drift* effect under objective heterogeneity; that is, the local iterates of each client drift-off towards the minimum of their own local function. In turn, this can lead to convergence to a point different from the global optimum $\mathbf{y}^*(\mathbf{x})$ of the inner problem; e.g., see (Mitra et al., 2021). This behavior is particularly undesirable in a nested optimization setting since it directly affects the outer optimization; see, e.g. (Liu et al., 2021, Section 7).

In light of this observation, we build on the recently proposed FEDLIN (Mitra et al., 2021) which improves FEDSVRG (Konečný et al., 2018) to solve the inner problem; see Algorithm 4. For each $i \in \mathcal{S}$, let $\mathbf{q}_i(\mathbf{x}, \mathbf{y})$ denote an unbiased estimate of the gradient $\nabla_{\mathbf{y}} g_i(\mathbf{x}, \mathbf{y})$. In each round, starting from a common global model $\mathbf{y}$, each client i performs τ_i local SVRG-type training steps in parallel: $\mathbf{y}_{i,\nu+1} = \mathbf{y}_{i,\nu} - \beta_i \mathbf{q}_{i,\nu}$, where $\mathbf{q}_{i,\nu} := \mathbf{q}_i(\mathbf{x}, \mathbf{y}_{i,\nu}) - \mathbf{q}_i(\mathbf{x}, \mathbf{y}) + \mathbf{q}(\mathbf{x}, \mathbf{y})$, $\beta_i \in (0, \beta]$ is the local inner step-size, and $\mathbf{q}(\mathbf{x}, \mathbf{y}) := |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{q}_i(\mathbf{x}, \mathbf{y})$. We note that for the optimization problems (1), (2), and (3), $\mathbf{q}_i(\mathbf{x}, \mathbf{y}_{i,\nu})$ is equal to $\nabla_{\mathbf{y}} g_i(\mathbf{x}, \mathbf{y}_{i,\nu}; \zeta_{i,\nu})$, $-\nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}_{i,\nu}; \xi_{i,\nu})$, and $\mathbf{y}_{i,\nu} - \mathbf{r}_i(\mathbf{x}; \zeta_{i,\nu})$, respectively; see Appendices C–E.

Algorithm 4 $\mathbf{y}^+ = \text{FEDINN}(\mathbf{x}, \mathbf{y}, \beta)$

```

1:  $\mathbb{G}_i(\cdot) \leftarrow \nabla_{\mathbf{y}} g_i(\mathbf{x}, \cdot)$  (bilevel),  $-\nabla_{\mathbf{y}} f_i(\mathbf{x}, \cdot)$  (minimax)
2:  $\mathbf{y}_{i,0} = \mathbf{y}$  and  $\beta_i \in (0, \beta]$ 
3: for  $i \in \mathcal{S}$  in parallel do
4:    $\mathbf{q}_i = \mathbb{G}_i(\mathbf{y}; \zeta_i)$ 
5: end for
6:  $\mathbf{q} = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{q}_i$ 
7: for  $i \in \mathcal{S}$  in parallel do
8:   for  $\nu = 0, \dots, \tau_i - 1$  do
9:      $\mathbf{q}_{i,\nu} = \mathbb{G}_i(\mathbf{y}_{i,\nu}; \zeta_{i,\nu}) - \mathbb{G}_i(\mathbf{y}; \zeta_{i,\nu}) + \mathbf{q}$ 
10:     $\mathbf{y}_{i,\nu+1} = \mathbf{y}_{i,\nu} - \beta_i \mathbf{q}_{i,\nu}$ 
11:   end for
12: end for
13:  $\mathbf{y}^+ = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} \mathbf{y}_{i,\tau_i}$ 

```

2.6. Light-FEDNEST: Communication Efficiency via Local Hypergradients

Each FEDNEST epoch k requires $2T + N + 3$ communication rounds as follows: $2T$ rounds for SVRG of FEDINN, N iterations for inverse Hessian approximation within FEDIHGP and 3 additional aggregations. Note that, these are vector communications and we fully avoid Hessian communication. In Appendix A, we also propose simplified variants of FEDOUT and FEDIHGP, which are tailored to homogeneous or high-dimensional FL settings. These algorithms can then either use *local* Jacobian / inverse Hessian or their approximation, and can use either SVRG or SGD.

Light-FEDNEST: Specifically, we propose LFEDNEST where each client runs IHGP locally. This reduces the number of rounds to $T + 1$, saving $T + N + 2$ rounds (see experiments in Section 4 for performance comparison and Appendix A for further discussion.)

3. Convergence Analysis for FEDNEST

In this section, we present convergence results for FEDNEST. All proofs are relegated to Appendices C–E.

Theorem 3.1. *Suppose Assumptions A and B hold. Further, assume $\alpha_i^k = \alpha_k / \tau_i$ and $\beta_i^k = \beta_k / \tau_i$ for all $i \in \mathcal{S}$, where*

$$\beta_k = \frac{\bar{\beta} \alpha_k}{T}, \quad \alpha_k = \min \left\{ \bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3, \frac{\bar{\alpha}}{\sqrt{K}} \right\} \quad (12)$$

for some positive constants $\bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3, \bar{\alpha}$, and $\bar{\beta}$ independent of K . Then, for any $T \geq 1$, the iterates $\{(\mathbf{x}^k, \mathbf{y}^k)\}_{k \geq 0}$ generated by FEDNEST satisfy

$$\begin{aligned} \frac{1}{K} \sum_{k=1}^K \mathbb{E} \left[\|\nabla f(\mathbf{x}^k)\|^2 \right] &= \mathcal{O} \left(\frac{\bar{\alpha} \max(\sigma_{g,1}^2, \sigma_{g,2}^2, \sigma_f^2)}{\sqrt{K}} \right. \\ &\quad \left. + \frac{1}{\min(\bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3) K} + b^2 \right), \end{aligned}$$

where $b = \kappa_g \ell_{f,1} ((\kappa_g - 1)/\kappa_g)^N$ and N is the input parameter to FEDIHGP.

Corollary 3.1 (Bilevel). *Under the same conditions as in Theorem 3.1, if $N = \mathcal{O}(\kappa_g \log K)$ and $T = \mathcal{O}(\kappa_g^4)$, then*

$$\frac{1}{K} \sum_{k=1}^K \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O} \left(\frac{\kappa_g^4}{K} + \frac{\kappa_g^{2.5}}{\sqrt{K}} \right).$$

For ϵ -accurate stationary point, we need $K = \mathcal{O}(\kappa_g^5 \epsilon^{-2})$.

Above, we choose $N \propto \kappa_g \log K$ to guarantee $b^2 \lesssim 1/\sqrt{K}$. In contrast, we use $T \gtrsim \kappa_g^4$ inner SVRG epochs. From Section 2.6, this would imply the communication cost is dominated by SVRG epochs N and $\mathcal{O}(\kappa_g^4)$ rounds.

From Corollary 3.1, we remark that FEDNEST matches the guarantees of centralized alternating SGD methods, such as ALSET (Chen et al., 2021a) and BSA (Ghadimi & Wang, 2018), despite federated setting, i.e. communication challenge, heterogeneity in the client objectives, and device heterogeneity.

3.1. Minimax Federated Learning

We focus on special features of federated minimax problems and customize the general results to yield improved convergence results for this special case. Recall from (2) that $g_i(\mathbf{x}, \mathbf{y}) = -f_i(\mathbf{x}, \mathbf{y})$ which implies that $b = 0$ and following Assumption A, $f_i(\mathbf{x}, \mathbf{y})$ is μ_f -strongly concave in $\mathbf{y}$ for all $\mathbf{x}$.

Corollary 3.2 (Minimax). *Denote $\kappa_f = \ell_{f,1}/\mu_f$. Assume same conditions as in Theorem 3.1 and $T = \mathcal{O}(\kappa_f)$. Then,*

$$\frac{1}{K} \sum_{k=1}^K \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O} \left(\frac{\kappa_f^2}{K} + \frac{\kappa_f}{\sqrt{K}} \right).$$

Corollary 3.2 implies that for the minimax problem, the convergence rate of FEDNEST to the stationary point of f is $\mathcal{O}(1/\sqrt{K})$. Again, we note this matches the convergence rate of non-FL algorithms (see also Table 1) such as SGDA (Lin et al., 2020) and SMD (Rafique et al., 2021).

3.2. Compositional Federated Learning

Observe that in the compositional problem (3), the outer function is $f_i(\mathbf{x}, \mathbf{y}; \xi) = f_i(\mathbf{y}; \xi)$ and the inner function is $g_i(\mathbf{x}, \mathbf{y}; \zeta) = \frac{1}{2} \|\mathbf{y} - \mathbf{r}_i(\mathbf{x}; \zeta)\|^2$, for all $i \in \mathcal{S}$. Hence, $b = 0$ and $\kappa_g = 1$.

Corollary 3.3 (Compositional). *Under the same conditions as in Theorem 3.1, if we select $T = 1$ in (12). Then,*

$$\frac{1}{K} \sum_{k=1}^K \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O} \left(\frac{1}{\sqrt{K}} \right).$$

Corollary 3.3 implies that for the compositional problem (3), the convergence rate of FEDNEST to the stationary point of f is $\mathcal{O}(1/\sqrt{K})$. This matches the convergence rate of non-federated stochastic algorithms such as SCGD (Wang et al., 2017) and NASA (Ghadimi et al., 2020) (Table 1).

3.3. Single-Level Federated Learning

Building upon the general results for stochastic nonconvex nested problems, we establish new convergence guarantees for *single-level* stochastic non-convex federated SVRG which is integrated within our FEDOUT. Note that in the single-level setting, the optimization problem (1) reduces to

$$\min_{\mathbf{x} \in \mathbb{R}^{d_1}} f(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m f_i(\mathbf{x}) \quad (13)$$

with $f_i(\mathbf{x}) := \mathbb{E}_{\xi \sim \mathcal{C}_i} [f_i(\mathbf{x}; \xi)]$, where $\xi \sim \mathcal{C}_i$ is sampling distribution for the i^{th} client.

We make the following assumptions on (13) that are counterparts of Assumptions A and B.

Assumption C (Lipschitz continuity). *For all $i \in [m]$, $\nabla f_i(\mathbf{x})$ is L_f -Lipschitz continuous.*

Assumption D (Stochastic samples). *For all $i \in [m]$, $\nabla f_i(\mathbf{x}; \xi)$ is an unbiased estimator of $\nabla f_i(\mathbf{x})$ and its variance is bounded, i.e., $\mathbb{E}_{\xi} [\|\nabla f_i(\mathbf{x}; \xi) - \nabla f_i(\mathbf{x})\|^2] \leq \sigma_f^2$.*

Theorem 3.2 (Single-Level). *Suppose Assumptions C and D hold. Further, assume $\alpha_i^k = \alpha_k/\tau_i$ for all $i \in \mathcal{S}$, where*

$$\alpha_k = \min \left\{ \bar{\alpha}_1, \frac{\bar{\alpha}}{\sqrt{K}} \right\} \quad (14)$$

for some $\bar{\alpha}_1, \bar{\alpha} > 0$. Then,

$$\frac{1}{K} \sum_{k=1}^K \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O} \left(\frac{\Delta_f}{\bar{\alpha}_1 K} + \frac{\frac{\Delta_f}{\bar{\alpha}} + \bar{\alpha} \sigma_f^2}{\sqrt{K}} \right),$$

where $\Delta_f := f(\mathbf{x}^0) - \mathbb{E}[f(\mathbf{x}^K)]$.

Theorem 3.2 extends recent results by (Mitra et al., 2021) from the stochastic strongly convex to the stochastic nonconvex setting. The above rate is also consistent with existing single-level non-FL guarantees (Ghadimi & Lan, 2013).

4. Numerical Experiments

In this section, we numerically investigate the impact of several attributes of our algorithms on a hyper-representation problem (Franceschi et al., 2018), a hyper-parameter optimization problem for loss function tuning (Li et al., 2021), and a federated minimax optimization problem.

4.1. Hyper-Representation Learning

Modern approaches in meta learning such as MAML (Finn et al., 2017) and reptile (Nichol & Schulman, 2018) learn

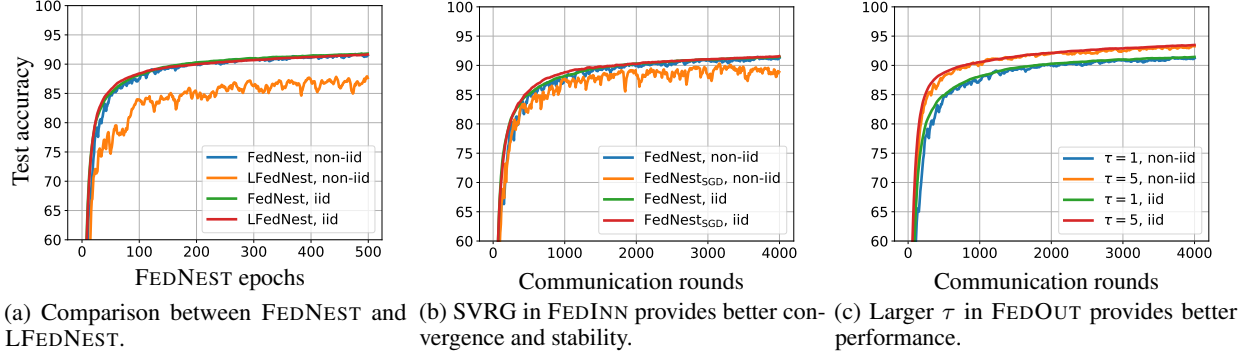


Figure 2: Hyper-representation experiments on a 2-layer MLP and MNIST dataset.

representations (that are shared across all tasks) in a bilevel manner. Similarly, the hyper-representation problem optimizes a classification model in a two-phased process. The outer objective optimizes the model backbone to obtain better feature representation on validation data. The inner problem optimizes a header for downstream classification tasks on training data. In this experiment, we use a 2-layer multilayer perceptron (MLP) with 200 hidden units. The outer problem optimizes the hidden layer with 157,000 parameters, and the inner problem optimizes the output layer with 2,010 parameters. We study both i.i.d. and non-i.i.d. ways of partitioning the MNIST data exactly following FEDAVG (McMahan et al., 2017), and split each client’s data evenly to train and validation datasets. Thus, each client has 300 train and 300 validation samples.

Figure 2 demonstrates the impact on test accuracy of several important components of FEDNEST. Figure 2a compares FEDNEST and LFEDNEST. Both algorithms perform well on the i.i.d. setup, while on the non-i.i.d. setup, FEDNEST achieves i.i.d. performance, significantly outperforming LFEDNEST. These findings are in line with our discussions in Section 2.6. LFEDNEST saves on communication rounds compared to FEDNEST and performs well on homogeneous clients. However, for heterogeneous clients, the isolation of local Hessian in LFEDNEST (see Algorithm 5 in Appendix A) degrades the test performance. Next, Figure 2b demonstrates the importance of SVRG in FEDINN algorithm for heterogeneous data (as predicted by our theoretical considerations in Section 2.5). To further clarify the algorithm difference in Figures 2b and 3a, we use FEDNEST_{SGD} to denote the FEDNEST algorithm where SGD is used in FEDINN. Finally, Figure 2c elucidates the role of local epoch τ in FEDOUT: larger τ saves on communication and improves test performance by enabling faster convergence.

4.2. Loss Function Tuning on Imbalanced Dataset

We use bilevel optimization to tune a loss function for learning an imbalanced MNIST dataset. We aim to maximize the class-balanced validation accuracy (which helps mi-

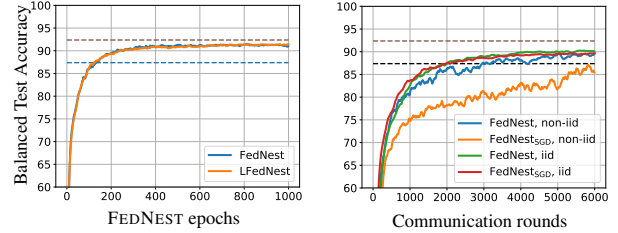

 Figure 3: Loss function tuning on a 3-layer MLP and imbalanced MNIST dataset to maximize *class-balanced test accuracy*. The *brown dashed* line is the accuracy on non-federated bilevel optimization (Li et al., 2021), and the *black dashed* line is the accuracy without tuning the loss function.

Figure 3: Loss function tuning on a 3-layer MLP and imbalanced MNIST dataset to maximize *class-balanced test accuracy*. The *brown dashed* line is the accuracy on non-federated bilevel optimization (Li et al., 2021), and the *black dashed* line is the accuracy without tuning the loss function.

nority/tail classes). Following the problem formulation in (Li et al., 2021) we tune the so-called VS-loss function (Kini et al., 2021) in a federated setting. In particular, we first create a long-tail imbalanced MNIST dataset by exponentially decreasing the number of examples per class (e.g. class 0 has 6,000 samples, class 1 has 3,597 samples and finally, class 9 has only 60 samples). We partition the dataset to 100 clients following again FEDAVG (McMahan et al., 2017) on both i.i.d. and non-i.i.d. setups. Different from the hyper-representation experiment, we employ 80%-20% train-validation on each client and use a 3-layer MLP model with 200, 100 hidden units, respectively. It is worth noting that, in this problem, the outer objective f (aka validation cost) only depends on the hyperparameter $\mathbf{x}$ through the optimal model parameters $\mathbf{y}^*(\mathbf{x})$; thus, the direct gradient $\nabla^D f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))$ is zero for all $i \in \mathcal{S}$.

Figure 3 displays test accuracy vs epochs/rounds for our federated bilevel algorithms. The horizontal dashed lines serve as centralized baselines: brown depicts accuracy reached by bilevel optimization in non-FL setting, and, black depicts accuracy without any loss tuning. Compared to these, Figure 3a shows that FEDNEST achieves near non-federated performance. In Figure 3b, we investigate the key role of SVRG in FEDINN by comparing it with possible alternative

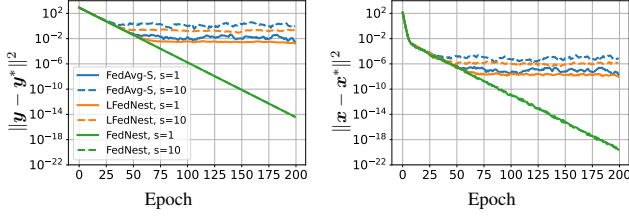


Figure 4: FEDNEST converges linearly despite heterogeneity. LFEDNEST slightly outperforms FEDAVG-S.

implementation that uses SGD-type updates. The figure confirms our discussion in Section 2.5: SVRG offers significant performance gains that are pronounced by client heterogeneity.

4.3. Federated Minimax Problem

We conduct experiments on the minimax problem (2) with

$$f_i(x, y) := -\left[\frac{1}{2}\|y\|^2 - b_i^\top y + y^\top A_i x\right] + \frac{\lambda}{2}\|x\|^2,$$

to compare standard FEDAVG saddle-point (FEDAVG-S) method updating (x, y) simultaneously (Hou et al., 2021) and our alternative approaches (LFEDNEST and FEDNEST). This is a saddle-point formulation of $\min_{x \in \mathbb{R}^{d_1}} \frac{1}{2} \left\| \frac{1}{m} \sum_{i=1}^m A_i x - b_i \right\|^2$. We set $\lambda = 10$, $b_i = b'_i - \frac{1}{m} \sum_{i=1}^m b'_i$ and $A_i = t_i I$, where $b'_i \sim \mathcal{N}(0, s^2 I_d)$, and t_i is drawn UAR over $(0, 0.1)$. Figure 4 shows that LFEDNEST and FEDNEST outperform FEDAVG-S thanks to their alternating nature. FEDNEST significantly improves the convergence of LFEDNEST due to controlling client-drift. To our knowledge, FEDNEST is the only *alternating* federated SVRG for minimax problems.

5. Related Work

Federated learning. FEDAVG was first introduced by McMahan et al. (2017), who showed it can dramatically reduce communication costs. For identical clients, FEDAVG coincides with local SGD (Zinkevich et al., 2010) which has been analyzed by many works (Stich, 2019; Yu et al., 2019; Wang & Joshi, 2018). Recently, many variants of FEDAVG have been proposed to tackle issues such as convergence and client drift. Examples include FEDPROX (Li et al., 2020b), SCAFFOLD (Karimireddy et al., 2020), FEDSPLIT (Pathak & Wainwright, 2020), FEDNOVA (Wang et al., 2020), and, the most closely relevant to us FEDLIN (Mittra et al., 2021). A few recent studies are also devoted to the extension of FEDAVG to the minimax optimization (Rasouli et al., 2020; Deng et al., 2020) and compositional optimization (Huang et al., 2021). In contrast to these methods, FEDNEST makes alternating SVRG updates between the global variables x and y , and yields sample complexity bounds and batch size choices that are on par with the non-FL guarantees (Table 1). Evaluations in the Appendix H.1 reveal that both alternating

updates and SVRG provides a performance boost over these prior approaches.

Bilevel optimization. This class of problems was first introduced by (Bracken & McGill, 1973), and since then, different types of approaches have been proposed. See (Sinha et al., 2017; Liu et al., 2021) for surveys. Earlier works in (Aiyoshi & Shimizu, 1984; Lv et al., 2007) reduced the bilevel problem to a single-level optimization problem. However, the reduced problem is still difficult to solve due to for example a large number of constraints. Recently, more efficient gradient-based algorithms have been proposed by estimating the hypergradient of $\nabla f(x)$ through iterative updates (Maclaurin et al., 2015; Franceschi et al., 2017; Domke, 2012; Pedregosa, 2016). The asymptotic and non-asymptotic analysis of bilevel optimization has been provided in (Franceschi et al., 2018; Shaban et al., 2019; Liu et al., 2020) and (Ghadimi & Wang, 2018; Hong et al., 2020), respectively. There is also a line of work focusing on minimax optimization (Nemirovski, 2004; Daskalakis & Panageas, 2018) and compositional optimization (Wang et al., 2017). Closely related to our work are (Lin et al., 2020; Rafique et al., 2021; Chen et al., 2021a) and (Ghadimi et al., 2020; Chen et al., 2021a) which provide non-asymptotic analysis of SGD-type methods for minimax and compositional problems with outer nonconvex objective, respectively.

A more in-depth discussion of related work is given in Appendix B. We summarize the complexities of different methods for FL/non-FL bilevel optimization in Table 1.

6. Conclusions

We presented a new class of federated algorithms for solving general nested stochastic optimization spanning bilevel and minimax problems. FEDNEST runs a variant of federated SVRG on inner & outer variables in an alternating fashion. We established provable convergence rates for FEDNEST under arbitrary client heterogeneity and introduced variations for min-max and compositional problems and for improved communication efficiency (LFEDNEST). We showed that, to achieve an ϵ -stationary point of the nested problem, FEDNEST requires $O(\epsilon^{-2})$ samples in total, which matches the complexity of the non-federated nested algorithms in the literature.

Acknowledgements

Davoud Ataee Tarzanagh was supported by ARO YIP award W911NF1910027 and NSF CAREER award CCF-1845076. Christos Thrampoulidis was supported by NSF Grant Numbers CCF-2009030 and HDR-1934641, and an NSERC Discovery Grant. Mingchen Li and Samet Oymak were supported by the NSF CAREER award CCF-2046816, Google Research Scholar award, and ARO grant W911NF2110312.

References

- Aiyoshi, E. and Shimizu, K. A solution method for the static constrained stackelberg problem via penalty method. *IEEE Transactions on Automatic Control*, 29(12):1111–1114, 1984.
- Al-Khayyal, F. A., Horst, R., and Pardalos, P. M. Global optimization of concave functions subject to quadratic constraints: an application in nonlinear bilevel programming. *Annals of Operations Research*, 34(1):125–147, 1992.
- Arora, S., Du, S., Kakade, S., Luo, Y., and Saunshi, N. Provable representation learning for imitation learning via bi-level optimization. In *International Conference on Machine Learning*, pp. 367–376. PMLR, 2020.
- Barazandeh, B., Huang, T., and Michailidis, G. A decentralized adaptive momentum method for solving a class of min-max optimization problems. *Signal Processing*, 189: 108245, 2021a.
- Barazandeh, B., Tarzanagh, D. A., and Michailidis, G. Solving a class of non-convex min-max games using adaptive momentum methods. In *ICASSP 2021-2021 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pp. 3625–3629. IEEE, 2021b.
- Basu, D., Data, D., Karakus, C., and Diggavi, S. Qsparse-local-SGD: Distributed SGD with quantization, sparsification and local computations. In *Advances in Neural Information Processing Systems*, pp. 14668–14679, 2019.
- Bertinetto, L., Henriques, J. F., Torr, P. H., and Vedaldi, A. Meta-learning with differentiable closed-form solvers. *arXiv preprint arXiv:1805.08136*, 2018.
- Bracken, J. and McGill, J. T. Mathematical programs with optimization problems in the constraints. *Operations Research*, 21(1):37–44, 1973.
- Brown, G. W. Iterative solution of games by fictitious play. *Activity analysis of production and allocation*, 13(1):374–376, 1951.
- Chen, T., Sun, Y., and Yin, W. Closing the gap: Tighter analysis of alternating stochastic gradient methods for bilevel problems. *Advances in Neural Information Processing Systems*, 34, 2021a.
- Chen, T., Sun, Y., and Yin, W. A single-timescale stochastic bilevel optimization method. *arXiv preprint arXiv:2102.04671*, 2021b.
- Dagr  ou, M., Ablin, P., Vaiter, S., and Moreau, T. A framework for bilevel optimization that enables stochastic and global variance reduction algorithms. *arXiv preprint arXiv:2201.13409*, 2022.
- Dai, B., He, N., Pan, Y., Boots, B., and Song, L. Learning from conditional distributions via dual embeddings. In *Artificial Intelligence and Statistics*, pp. 1458–1467. PMLR, 2017.
- Daskalakis, C. and Panageas, I. The limit points of (optimistic) gradient descent in min-max optimization. *arXiv preprint arXiv:1807.03907*, 2018.
- Deng, Y. and Mahdavi, M. Local stochastic gradient descent ascent: Convergence analysis and communication efficiency. In *International Conference on Artificial Intelligence and Statistics*, pp. 1387–1395. PMLR, 2021.
- Deng, Y., Kamani, M. M., and Mahdavi, M. Distributionally robust federated averaging. *Advances in Neural Information Processing Systems*, 33:15111–15122, 2020.
- Diakonikolas, J., Daskalakis, C., and Jordan, M. Efficient methods for structured nonconvex-nonconcave min-max optimization. In *International Conference on Artificial Intelligence and Statistics*, pp. 2746–2754. PMLR, 2021.
- Domke, J. Generic methods for optimization-based modeling. In *Artificial Intelligence and Statistics*, pp. 318–326. PMLR, 2012.
- Edmunds, T. A. and Bard, J. F. Algorithms for nonlinear bilevel mathematical programs. *IEEE transactions on Systems, Man, and Cybernetics*, 21(1):83–89, 1991.
- Finn, C., Abbeel, P., and Levine, S. Model-agnostic meta-learning for fast adaptation of deep networks. In *International Conference on Machine Learning*, pp. 1126–1135. PMLR, 2017.
- Franceschi, L., Donini, M., Frasconi, P., and Pontil, M. Forward and reverse gradient-based hyperparameter optimization. In *International Conference on Machine Learning*, pp. 1165–1173. PMLR, 2017.
- Franceschi, L., Frasconi, P., Salzo, S., Grazzi, R., and Pontil, M. Bilevel programming for hyperparameter optimization and meta-learning. In *International Conference on Machine Learning*, pp. 1568–1577. PMLR, 2018.
- Ghadimi, S. and Lan, G. Stochastic first-and zeroth-order methods for nonconvex stochastic programming. *SIAM Journal on Optimization*, 23(4):2341–2368, 2013.
- Ghadimi, S. and Wang, M. Approximation methods for bilevel programming. *arXiv preprint arXiv:1802.02246*, 2018.
- Ghadimi, S., Ruszczynski, A., and Wang, M. A single timescale stochastic approximation method for nested stochastic optimization. *SIAM Journal on Optimization*, 30(1):960–979, 2020.

- Gidel, G., Berard, H., Vignoud, G., Vincent, P., and Lacoste-Julien, S. A variational inequality perspective on generative adversarial networks. *arXiv preprint arXiv:1802.10551*, 2018.
- Grazzi, R., Franceschi, L., Pontil, M., and Salzo, S. On the iteration complexity of hypergradient computation. In *International Conference on Machine Learning*, pp. 3748–3758. PMLR, 2020.
- Guo, Z., Xu, Y., Yin, W., Jin, R., and Yang, T. On stochastic moving-average estimators for non-convex optimization. *arXiv preprint arXiv:2104.14840*, 2021.
- Hansen, P., Jaumard, B., and Savard, G. New branch-and-bound rules for linear bilevel programming. *SIAM Journal on scientific and Statistical Computing*, 13(5):1194–1217, 1992.
- Hong, M., Wai, H.-T., Wang, Z., and Yang, Z. A two-timescale framework for bilevel optimization: Complexity analysis and application to actor-critic. *arXiv preprint arXiv:2007.05170*, 2020.
- Hou, C., Thekumparampil, K. K., Fanti, G., and Oh, S. Efficient algorithms for federated saddle point optimization. *arXiv preprint arXiv:2102.06333*, 2021.
- Hsu, T.-M. H., Qi, H., and Brown, M. Measuring the effects of non-identical data distribution for federated visual classification. *arXiv preprint arXiv:1909.06335*, 2019.
- Huang, F. and Huang, H. Biadam: Fast adaptive bilevel optimization methods. *arXiv preprint arXiv:2106.11396*, 2021.
- Huang, F., Li, J., and Huang, H. Compositional federated learning: Applications in distributionally robust averaging and meta learning. *arXiv preprint arXiv:2106.11264*, 2021.
- Ji, K. and Liang, Y. Lower bounds and accelerated algorithms for bilevel optimization. *ArXiv*, abs/2102.03926, 2021.
- Ji, K., Yang, J., and Liang, Y. Provably faster algorithms for bilevel optimization and applications to meta-learning. *ArXiv*, abs/2010.07962, 2020a.
- Ji, K., Yang, J., and Liang, Y. Theoretical convergence of multi-step model-agnostic meta-learning. *arXiv preprint arXiv:2002.07836*, 2020b.
- Ji, K., Yang, J., and Liang, Y. Bilevel optimization: Convergence analysis and enhanced design. In *International Conference on Machine Learning*, pp. 4882–4892. PMLR, 2021.
- Ji, S. A pytorch implementation of federated learning. Mar 2018. doi: 10.5281/zenodo.4321561.
- Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., Bonawitz, K., Charles, Z., Cormode, G., Cummings, R., et al. Advances and open problems in federated learning. *arXiv preprint arXiv:1912.04977*, 2019.
- Karimireddy, S. P., Kale, S., Mohri, M., Reddi, S., Stich, S., and Suresh, A. T. Scaffold: Stochastic controlled averaging for federated learning. In *International Conference on Machine Learning*, pp. 5132–5143. PMLR, 2020.
- Khaled, A., Mishchenko, K., and Richtárik, P. First analysis of local GD on heterogeneous data. *arXiv preprint arXiv:1909.04715*, 2019.
- Khanduri, P., Zeng, S., Hong, M., Wai, H.-T., Wang, Z., and Yang, Z. A near-optimal algorithm for stochastic bilevel optimization via double-momentum. *arXiv preprint arXiv:2102.07367*, 2021.
- Khodak, M., Tu, R., Li, T., Li, L., Balcan, M.-F. F., Smith, V., and Talwalkar, A. Federated hyperparameter tuning: Challenges, baselines, and connections to weight-sharing. *Advances in Neural Information Processing Systems*, 34, 2021.
- Kini, G. R., Paraskevas, O., Oymak, S., and Thrampoulidis, C. Label-imbalanced and group-sensitive classification under overparameterization. *arXiv preprint arXiv:2103.01550*, 2021.
- Konečný, J., McMahan, H. B., Ramage, D., and Richtárik, P. Federated optimization: Distributed machine learning for on-device intelligence. *International Conference on Learning Representations*, 2018.
- Li, J., Gu, B., and Huang, H. Improved bilevel model: Fast and optimal algorithm with theoretical guarantee. *arXiv preprint arXiv:2009.00690*, 2020a.
- Li, M., Zhang, X., Thrampoulidis, C., Chen, J., and Oymak, S. Autobalance: Optimized loss functions for imbalanced data. *Advances in Neural Information Processing Systems*, 34, 2021.
- Li, T., Sahu, A. K., Zaheer, M., Sanjabi, M., Talwalkar, A., and Smith, V. Federated optimization in heterogeneous networks. *Proceedings of Machine Learning and Systems*, 2:429–450, 2020b.
- Li, X., Huang, K., Yang, W., Wang, S., and Zhang, Z. On the convergence of FedAvg on non-IID data. *arXiv preprint arXiv:1907.02189*, 2019.

- Lin, T., Jin, C., and Jordan, M. On gradient descent ascent for nonconvex-concave minimax problems. In *International Conference on Machine Learning*, pp. 6083–6093. PMLR, 2020.
- Liu, H., Simonyan, K., and Yang, Y. Darts: Differentiable architecture search. *arXiv preprint arXiv:1806.09055*, 2018.
- Liu, M., Mroueh, Y., Ross, J., Zhang, W., Cui, X., Das, P., and Yang, T. Towards better understanding of adaptive gradient algorithms in generative adversarial nets. *arXiv preprint arXiv:1912.11940*, 2019.
- Liu, R., Mu, P., Yuan, X., Zeng, S., and Zhang, J. A generic first-order algorithmic framework for bi-level programming beyond lower-level singleton. In *International Conference on Machine Learning*, pp. 6305–6315. PMLR, 2020.
- Liu, R., Gao, J., Zhang, J., Meng, D., and Lin, Z. Investigating bi-level optimization for learning and vision from a unified perspective: A survey and beyond. *arXiv preprint arXiv:2101.11517*, 2021.
- Luo, L., Ye, H., Huang, Z., and Zhang, T. Stochastic recursive gradient descent ascent for stochastic nonconvex-strongly-concave minimax problems. *Advances in Neural Information Processing Systems*, 33:20566–20577, 2020.
- Lv, Y., Hu, T., Wang, G., and Wan, Z. A penalty function method based on kuhn–tucker condition for solving linear bilevel programming. *Applied Mathematics and Computation*, 188(1):808–813, 2007.
- Maclaurin, D., Duvenaud, D., and Adams, R. Gradient-based hyperparameter optimization through reversible learning. In *International conference on machine learning*, pp. 2113–2122. PMLR, 2015.
- Madry, A., Makelov, A., Schmidt, L., Tsipras, D., and Vladu, A. Towards deep learning models resistant to adversarial attacks. *arXiv preprint arXiv:1706.06083*, 2017.
- McMahan, B., Moore, E., Ramage, D., Hampson, S., and y Arcas, B. A. Communication-efficient learning of deep networks from decentralized data. In *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics, AISTATS 2017, 20-22 April 2017, Fort Lauderdale, FL, USA*, pp. 1273–1282, 2017. URL <http://proceedings.mlr.press/v54/mcmahan17a.html>.
- Mitra, A., Jaafar, R., Pappas, G., and Hassani, H. Linear convergence in federated learning: Tackling client heterogeneity and sparse gradients. *Advances in Neural Information Processing Systems*, 34, 2021.
- Mohri, M., Sivek, G., and Suresh, A. T. Agnostic federated learning. In *International Conference on Machine Learning*, pp. 4615–4625. PMLR, 2019.
- Mokhtari, A., Ozdaglar, A., and Pattathil, S. A unified analysis of extra-gradient and optimistic gradient methods for saddle point problems: Proximal point approach. In *International Conference on Artificial Intelligence and Statistics*, pp. 1497–1507. PMLR, 2020.
- Moore, G. M. *Bilevel programming algorithms for machine learning model selection*. Rensselaer Polytechnic Institute, 2010.
- Nazari, P., Tarzanagh, D. A., and Michailidis, G. Dadam: A consensus-based distributed adaptive gradient method for online optimization. *arXiv preprint arXiv:1901.09109*, 2019.
- Nedić, A. and Ozdaglar, A. Subgradient methods for saddle-point problems. *Journal of optimization theory and applications*, 142(1):205–228, 2009.
- Nemirovski, A. Prox-method with rate of convergence $o(1/t)$ for variational inequalities with lipschitz continuous monotone operators and smooth convex-concave saddle point problems. *SIAM Journal on Optimization*, 15(1):229–251, 2004.
- Nichol, A. and Schulman, J. Reptile: a scalable metalearning algorithm. *arXiv preprint arXiv:1803.02999*, 2(3):4, 2018.
- Nouiehed, M., Sanjabi, M., Huang, T., Lee, J. D., and Razaviyayn, M. Solving a class of non-convex min-max games using iterative first order methods. *Advances in Neural Information Processing Systems*, 32, 2019.
- Pathak, R. and Wainwright, M. J. Fedsplit: an algorithmic framework for fast federated optimization. *Advances in Neural Information Processing Systems*, 33:7057–7066, 2020.
- Pedregosa, F. Hyperparameter optimization with approximate gradient. In *International conference on machine learning*, pp. 737–746. PMLR, 2016.
- Rafique, H., Liu, M., Lin, Q., and Yang, T. Weakly-convex-concave min-max optimization: provable algorithms and applications in machine learning. *Optimization Methods and Software*, pp. 1–35, 2021.
- Rasouli, M., Sun, T., and Rajagopal, R. Fedgan: Federated generative adversarial networks for distributed data. *arXiv preprint arXiv:2006.07228*, 2020.
- Reddi, S. J., Charles, Z., Zaheer, M., Garrett, Z., Rush, K., Konečný, J., Kumar, S., and McMahan, H. B. Adaptive

- federated optimization. In *International Conference on Learning Representations*, 2020.
- Reisizadeh, A., Farnia, F., Pedarsani, R., and Jadbabaie, A. Robust federated learning: The case of affine distribution shifts. *Advances in Neural Information Processing Systems*, 33:21554–21565, 2020.
- Shaban, A., Cheng, C.-A., Hatch, N., and Boots, B. Truncated back-propagation for bilevel optimization. In *The 22nd International Conference on Artificial Intelligence and Statistics*, pp. 1723–1732. PMLR, 2019.
- Shen, Y., Du, J., Zhao, H., Zhang, B., Ji, Z., and Gao, M. Fedmm: Saddle point optimization for federated adversarial domain adaptation. *arXiv preprint arXiv:2110.08477*, 2021.
- Shi, C., Lu, J., and Zhang, G. An extended kuhn–tucker approach for linear bilevel programming. *Applied Mathematics and Computation*, 162(1):51–63, 2005.
- Sinha, A., Malo, P., and Deb, K. A review on bilevel optimization: from classical to evolutionary approaches and applications. *IEEE Transactions on Evolutionary Computation*, 22(2):276–295, 2017.
- Stich, S. U. Local SGD converges fast and communicates little. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=Slg2JnRcFX>.
- Stich, S. U. and Karimireddy, S. P. The error-feedback framework: Better rates for SGD with delayed gradients and compressed communication. *arXiv preprint arXiv:1909.05350*, 2019.
- Thekumparampil, K. K., Jain, P., Netrapalli, P., and Oh, S. Efficient algorithms for smooth minimax optimization. *arXiv preprint arXiv:1907.01543*, 2019.
- Tran Dinh, Q., Liu, D., and Nguyen, L. Hybrid variance-reduced sgd algorithms for minimax problems with nonconvex-linear function. *Advances in Neural Information Processing Systems*, 33:11096–11107, 2020.
- Wang, J. and Joshi, G. Cooperative SGD: A unified framework for the design and analysis of communication-efficient SGD algorithms. *arXiv preprint arXiv:1808.07576*, 2018.
- Wang, J., Liu, Q., Liang, H., Joshi, G., and Poor, H. V. Tackling the objective inconsistency problem in heterogeneous federated optimization. *Advances in neural information processing systems*, 2020.
- Wang, M., Liu, J., and Fang, E. Accelerating stochastic composition optimization. *Advances in Neural Information Processing Systems*, 29, 2016.
- Wang, M., Fang, E. X., and Liu, H. Stochastic compositional gradient descent: algorithms for minimizing compositions of expected-value functions. *Mathematical Programming*, 161(1-2):419–449, 2017.
- Wang, S., Tuor, T., Salonidis, T., Leung, K. K., Makaya, C., He, T., and Chan, K. Adaptive federated learning in resource constrained edge computing systems. *IEEE Journal on Selected Areas in Communications*, 37(6):1205–1221, 2019.
- Wu, Y. F., Zhang, W., Xu, P., and Gu, Q. A finite-time analysis of two time-scale actor-critic methods. *Advances in Neural Information Processing Systems*, 33:17617–17628, 2020.
- Xie, J., Zhang, C., Zhang, Y., Shen, Z., and Qian, H. A federated learning framework for nonconvex-pl minimax problems. *arXiv preprint arXiv:2105.14216*, 2021.
- Yan, Y., Xu, Y., Lin, Q., Liu, W., and Yang, T. Optimal epoch stochastic gradient descent ascent methods for minimax optimization. *Advances in Neural Information Processing Systems*, 33:5789–5800, 2020.
- Yang, J., Kiyavash, N., and He, N. Global convergence and variance reduction for a class of nonconvex-nonconcave minimax problems. *Advances in Neural Information Processing Systems*, 33:1153–1165, 2020.
- Yoon, T. and Ryu, E. K. Accelerated algorithms for smooth convex-concave minimax problems with $\mathcal{O}(1/k^2)$ rate on squared gradient norm. In *International Conference on Machine Learning*, pp. 12098–12109. PMLR, 2021.
- Yu, H., Yang, S., and Zhu, S. Parallel restarted SGD with faster convergence and less communication: Demystifying why model averaging works for deep learning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pp. 5693–5700, 2019.
- Zinkevich, M., Weimer, M., Li, L., and Smola, A. J. Parallelized stochastic gradient descent. In *Advances in neural information processing systems*, pp. 2595–2603, 2010.

APPENDIX

FEDNEST: Federated Bilevel, Minimax, and Compositional Optimization

The appendix is organized as follows: Section A introduces the LFEDNEST algorithm. Section B discusses the related work. We provide all details for the proof of the main theorems in Sections C, D, E, and F for federated bilevel, minimax, compositional, and single-level optimization, respectively. In Section G, we state a few auxiliary technical lemmas. Finally, in Section H, we provide the detailed parameters of our numerical experiments (Section 4) and then introduce further experiments.

A. LFEDNEST

Implementing FEDINN and FEDOUT naively by using the global direct and indirect gradients and sending the local information to the server that would then calculate the global gradients leads to a communication and space complexity of which can be prohibitive for large-sized d_1 and d_2 . One can consider possible local variants of FEDINN and FEDOUT tailored to such scenarios. Each of the possible algorithms (See Table 2) can then either use the global gradient or only the local gradient, either use a SVRG or SGD.

Algorithm 5 $x^+ = \text{LFEDOUT}(x, y, \alpha)$ for stochastic bilevel, minimax, and compositional problems

```

1:  $x_{i,0} = x$  and  $\alpha_i \in (0, \alpha]$  for each  $i \in \mathcal{S}$ .
2: Choose  $N \in \{1, 2, \dots\}$  (the number of terms of Neumann series).
3: for  $i \in \mathcal{S}$  in parallel do
4:   for  $\nu = 0, \dots, \tau_i - 1$  do
5:     Select  $N' \in \{0, \dots, N - 1\}$  UAR.
6:      $h_{i,\nu} = \nabla_x f_i(x_{i,\nu}, y; \xi_{i,\nu}) - \frac{N}{\ell_{g,1}} \nabla_{xy}^2 g_i(x_{i,\nu}, y; \zeta_{i,\nu}) \prod_{n=1}^{N'} (I - \frac{1}{\ell_{g,1}} \nabla_y^2 g_i(x_{i,\nu}, y; \zeta_{i,n})) \nabla_x f_i(y_{i,\nu}, y; \xi_{i,\nu})$ 
7:      $h_{i,\nu} = \nabla_x f_i(x_{i,\nu}, y; \xi_{i,\nu})$ 
8:      $h_{i,\nu} = \nabla r_i(x_{i,\nu}; \zeta_{i,\nu})^\top \nabla f_i(y_{i,\nu}; \xi_{i,\nu})$ 
9:      $x_{i,\nu+1} = x_{i,\nu} - \alpha_i h_{i,\nu}$ 
10:   end for
11: end for
12:  $x^+ = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} x_{i,\tau_i}$ 

```

Algorithm 6 $y^+ = \text{LFEDINN}(x, y, \beta)$ for stochastic bilevel, minimax, and compositional problems

```

1:  $y_{i,0} = y$  and  $\beta_i \in (0, \beta]$  for each  $i \in \mathcal{S}$ .
2: for  $i \in \mathcal{S}$  in parallel do
3:   for  $\nu = 0, \dots, \tau_i - 1$  do
4:      $q_{i,\nu} = \nabla_y g_i(x, y_{i,\nu}; \zeta_{i,\nu})$   $q_{i,\nu} = -\nabla_y f_i(x, y_{i,\nu}; \xi_{i,\nu})$   $q_{i,\nu} = y_{i,\nu} - r_i(x; \zeta_{i,\nu})$ 
5:      $y_{i,\nu+1} = y_{i,\nu} - \beta_i q_{i,\nu}$ 
6:   end for
7: end for
8:  $y^+ = |\mathcal{S}|^{-1} \sum_{i \in \mathcal{S}} y_{i,\tau_i}$ 

```

	definition		properties			
	outer optimization	inner optimization	global outer gradient	global IHGP	global inner gradient	# communication rounds
FEDNEST	Algorithm 2 (SVRG on x)	Algorithm 4 (SVRG on y)	yes	yes	yes	$2T + N + 3$
LFEDNEST	Algorithm 5 (SGD on x)	Algorithm 6 (SGD on y)	no	no	no	$T + 1$
FEDNEST_{SGD}	Algorithm 2 (SVRG on x)	Algorithm 6 (SGD on y)	yes	yes	no	$T + N + 3$
LFEDNEST_{SVRG}	Algorithm 5 (SGD on x)	Algorithm 4 (SVRG on y)	no	no	yes	$2T + 1$

Table 2: Definition of studied algorithms by using inner/outer optimization algorithms and server updates and resulting properties of these algorithms. T and N denote the number of inner iterations and terms of Neumann series, respectively.

B. Related Work

We provide an overview of the current literature on non-federated nested (bilevel, minimax, and compositional) optimization and federated learning.

B.1. Bilevel Optimization

A broad collection of algorithms have been proposed to solve bilevel nonlinear programming problems. Aiyoshi & Shimizu (1984); Edmunds & Bard (1991); Al-Khayyal et al. (1992); Hansen et al. (1992); Shi et al. (2005); Lv et al. (2007); Moore (2010) reduce the bilevel problem to a single-level optimization problem using for example the Karush-Kuhn-Tucker (KKT) conditions or penalty function methods. A similar idea was also explored in Khodak et al. (2021) where the authors provide a reformulation of the hyperparameter optimization (bilevel objective) into a single-level objective and develop a federated online method to solve it. However, the reduced single-level problem is usually difficult to solve (Sinha et al., 2017).

In comparison, alternating gradient-based approaches designed for the bilevel problems are more attractive due to their simplicity and effectiveness. This type of approaches estimate the hypergradient $\nabla f(x)$ for iterative updates, and are generally divided to approximate implicit differentiation (AID) and iterative differentiation (ITD) categories. ITD-based approaches (Maclaurin et al., 2015; Franceschi et al., 2017; Finn et al., 2017; Grazi et al., 2020) estimate the hypergradient $\nabla f(x)$ in either a reverse (automatic differentiation) or forward manner. AID-based approaches (Pedregosa, 2016; Grazi et al., 2020; Ghadimi & Wang, 2018) estimate the hypergradient via implicit differentiation which involves solving a linear system. Our algorithms follow the latter approach.

Theoretically, bilevel optimization has been studied via both asymptotic and non-asymptotic analysis (Franceschi et al., 2018; Liu et al., 2020; Li et al., 2020a; Shaban et al., 2019; Ghadimi & Wang, 2018; Ji et al., 2021; Hong et al., 2020). In particular, (Franceschi et al., 2018) provided the asymptotic convergence of a backpropagation-based approach as one of ITD-based algorithms by assuming the inner problem is strongly convex. (Shaban et al., 2019) gave a similar analysis for a *truncated* backpropagation approach. Non-asymptotic complexity analysis for bilevel optimization has also been explored. Ghadimi & Wang (2018) provided a finite-time convergence analysis for an AID-based algorithm under three different loss geometries, where $f(\cdot)$ is either strongly convex, convex or nonconvex, and $g(x, \cdot)$ is strongly convex. (Ji et al., 2021) provided an improved non-asymptotic analysis for AID- and ITD-based algorithms under the nonconvex-strongly-convex geometry. (Ji & Liang, 2021) provided the first-known lower bounds on complexity as well as tighter upper bounds. When the objective functions can be expressed in an expected or finite-time form, (Ghadimi & Wang, 2018; Ji et al., 2021; Hong et al., 2020) developed stochastic bilevel algorithms and provided the non-asymptotic analysis. (Chen et al., 2021a) provided a tighter analysis of SGD for stochastic bilevel problems. (Chen et al., 2021b; Guo et al., 2021; Khanduri et al., 2021; Ji et al., 2020a; Huang & Huang, 2021; Dagréou et al., 2022) studied accelerated SGD, SAGA, momentum, and adaptive-type bilevel optimization methods. More results can be found in the recent review paper (Liu et al., 2021) and references therein.

B.1.1. MINIMAX OPTIMIZATION

Minimax optimization has a long history dating back to (Brown, 1951). Earlier works focused on the deterministic convex-concave regime (Nemirovski, 2004; Nedić & Ozdaglar, 2009). Recently, there has emerged a surge of studies of stochastic minimax problems. The alternating version of the gradient descent ascent (SGDA) has been studied by incorporating the idea of optimism (Daskalakis & Panageas, 2018; Gidel et al., 2018; Mokhtari et al., 2020; Yoon & Ryu, 2021). (Rafique et al., 2021; Thekumparampil et al., 2019; Nouiehed et al., 2019; Lin et al., 2020) studied SGDA in the nonconvex-strongly concave setting. Specifically, the $\mathcal{O}(\epsilon^{-2})$ sample complexity has been established in (Lin et al., 2020) under an increasing batch size $\mathcal{O}(\epsilon^{-1})$. Chen et al. (2021a) provided the $\mathcal{O}(\epsilon^{-2})$ sample complexity under an $\mathcal{O}(1)$ constant batch size. In the same setting, accelerated GDA algorithms have been developed in (Luo et al., 2020; Yan et al., 2020; Tran Dinh et al., 2020). Going beyond the one-side concave settings, algorithms and their convergence analysis have been studied for nonconvex-nonconcave minimax problems with certain benign structure; see e.g., (Gidel et al., 2018; Liu et al., 2019; Yang et al., 2020; Diakonikolas et al., 2021; Barazandeh et al., 2021b). A comparison of our results with prior work can be found in Table 1.

B.1.2. COMPOSITIONAL OPTIMIZATION

Stochastic compositional gradient algorithms (Wang et al., 2017; 2016) can be viewed as an alternating SGD for the special compositional problem. However, to ensure convergence, the algorithms in (Wang et al., 2017; 2016) use two sequences of variables being updated in two different time scales, and thus the iteration complexity of (Wang et al., 2017) and (Wang et al., 2016) is worse than $\mathcal{O}(\epsilon^{-2})$ of the standard SGD. Our work is closely related to ALSET (Chen et al., 2021a), where an $\mathcal{O}(\epsilon^{-2})$ sample complexity has been established in a non-FL setting.

B.2. Federated Learning

FL involves learning a centralized model from distributed client data. Although this centralized model benefits from all client data, it raises several types of issues such as generalization, fairness, communication efficiency, and privacy (Mohri et al., 2019; Stich, 2019; Yu et al., 2019; Wang & Joshi, 2018; Stich & Karimireddy, 2019; Basu et al., 2019; Nazari et al., 2019; Barazandeh et al., 2021a). FEDAVG (McMahan et al., 2017) can tackle some of these issues such as high communication costs. Many variants of FEDAVG have been proposed to tackle other emerging issues such as convergence and *client drift*. Examples include adding a regularization term in the client objectives towards the broadcast model (Li et al., 2020b), proximal splitting (Pathak & Wainwright, 2020; Mitra et al., 2021), variance reduction (Karimireddy et al., 2020; Mitra et al., 2021) and adaptive updates (Reddi et al., 2020). When clients are homogeneous, FEDAVG is closely related to local SGD (Zinkevich et al., 2010), which has been analyzed by many works (Stich, 2019; Yu et al., 2019; Wang & Joshi, 2018; Stich & Karimireddy, 2019; Basu et al., 2019).

In order to analyze FEDAVG in heterogeneous settings, (Li et al., 2020b; Wang et al., 2019; Khaled et al., 2019; Li et al., 2019) derive convergence rates depending on the amount of heterogeneity. They showed that the convergence rate of FEDAVG gets worse with client heterogeneity. By using control variates to reduce client drift, the SCAFFOLD method (Karimireddy et al., 2020) achieves convergence rates that are independent of the amount of heterogeneity. Relatedly, FEDNOVA (Wang et al., 2020) and FEDLIN (Mitra et al., 2021) provided the convergence of their methods despite arbitrary local objective and systems heterogeneity. In particular, (Mitra et al., 2021) showed that FEDLIN guarantees linear convergence to the global minimum of deterministic objective, despite arbitrary objective and systems heterogeneity. As explained in the main body, our algorithms critically leverage these ideas after identifying the additional challenges that client drift brings to federated bilevel settings.

B.2.1. FEDERATED MINIMAX LEARNING

A few recent studies are devoted to federated minimax optimization (Rasouli et al., 2020; Reisizadeh et al., 2020; Deng et al., 2020; Hou et al., 2021). In particular, (Reisizadeh et al., 2020) consider minimax problem with inner problem satisfying PL condition and the outer one being either nonconvex or satisfying PL. However, the proposed algorithm only communicates $\mathbf{x}$ to the server. Xie et al. (2021) consider a general class of nonconvex-PL minimax problems in the cross-device federated learning setting. Their algorithm performs multiple local update steps on a subset of active clients in each round and leverages global gradient estimates to correct the bias in local update directions. Deng & Mahdavi (2021) studied federated optimization for a family of smooth nonconvex minimax functions. Shen et al. (2021) proposed a distributed minimax optimizer called FEDMM, designed specifically for the federated adversary domain adaptation problem.

Hou et al. (2021) proposed a SCAFFOLD saddle point algorithm (SCAFFOLD-S) for solving strongly convex-concave minimax problems in the federated setting. To the best of our knowledge, all the aforementioned developments require a bound on the *heterogeneity* of the local functions, and do not account for the effects of systems heterogeneity which is also a key challenge in FL. In addition, our work proposes the first *alternating* federated SVRG-type algorithm for minimax problems with iteration complexity that matches to the non-federated setting (see, Table 1).

C. Proof for Federated Bilevel Optimization

Throughout the proof, we will use $\mathcal{F}_{i,\nu}^{k,t}$ to denote the filtration that captures all the randomness up to the ν -th local step of client i in inner round t and outer round k . With a slight abuse of notation, $\mathcal{F}_{i,-1}^{k,t}$ is to be interpreted as $\mathcal{F}^{k,t}, \forall i \in \mathcal{S}$. For simplicity, we remove subscripts k and t from the definition of stepsize and model parameters. For example, $\mathbf{x}$ and $\mathbf{x}^+$ denote $\mathbf{x}^k$ and $\mathbf{x}^{k+1}$, respectively. We further set

$$\bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) := \mathbb{E}[\mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) | \mathcal{F}_{i,\nu-1}]. \quad (15)$$

Proof of Lemma 2.1

Proof. Given $\mathbf{x} \in \mathbb{R}^{d_1}$, the optimality condition of the inner problem in (1) is $\nabla_{\mathbf{y}} g(\mathbf{x}, \mathbf{y}) = 0$. Now, since $\nabla_{\mathbf{x}} (\nabla_{\mathbf{y}} g(\mathbf{x}, \mathbf{y})) = 0$, we obtain

$$0 = \sum_{j=1}^m (\nabla_{\mathbf{x}\mathbf{y}}^2 g_j(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) + \nabla \mathbf{y}^*(\mathbf{x}) \nabla_{\mathbf{y}}^2 g_j(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))),$$

which implies

$$\nabla \mathbf{y}^*(\mathbf{x}) = - \left(\sum_{i=1}^m \nabla_{\mathbf{x}\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) \right) \left(\sum_{i=1}^m \nabla_{\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) \right)^{-1}.$$

The results follows from a simple application of the chain rule to f as follows:

$$\nabla f(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) = \nabla_{\mathbf{x}} f(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) + \nabla \mathbf{y}^*(\mathbf{x}) \nabla_{\mathbf{y}} f(\mathbf{x}, \mathbf{y}^*(\mathbf{x})).$$

□

Proof of Lemma 2.2

Proof. By independency of N' , $\zeta_{i,n}$, and $\mathcal{S}_n$, and under Assumption B, we have

$$\begin{aligned} \mathbb{E}_{\mathcal{W}} [\widehat{\mathbf{H}}_{\mathbf{y}}] &= \mathbb{E}_{\mathcal{W}} \left[\frac{N}{\ell_{g,1}} \prod_{n=1}^{N'} \left(\mathbf{I} - \frac{1}{\ell_{g,1} |\mathcal{S}_n|} \sum_{i=1}^{|\mathcal{S}_n|} \nabla_{\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}; \zeta_{i,n}) \right) \right] \\ &= \mathbb{E}_{N'} \left[\mathbb{E}_{\mathcal{S}_{1:N'}} \left[\mathbb{E}_{\zeta} \left[\frac{N}{\ell_{g,1}} \prod_{n=1}^{N'} \left(\mathbf{I} - \frac{1}{\ell_{g,1} |\mathcal{S}_n|} \sum_{i=1}^{|\mathcal{S}_n|} \nabla_{\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}; \zeta_{i,n}) \right) \right] \right] \right] \\ &= \frac{1}{\ell_{g,1}} \sum_{n=0}^{N-1} \left[\mathbf{I} - \frac{1}{\ell_{g,1}} \nabla_{\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}) \right]^n, \end{aligned} \quad (16)$$

where the last equality follows from the uniform distribution of N' .

Note that since $\mathbf{I} \succeq \frac{1}{\ell_{g,1}} \nabla_{\mathbf{y}}^2 g_i \succeq \frac{\mu_g}{\ell_{g,1}}$ for all $i \in [m]$ due to Assumption A, we have

$$\begin{aligned} \mathbb{E}_{\mathcal{W}} [\|\widehat{\mathbf{H}}_{\mathbf{y}}\|] &\leq \frac{N}{\ell_{g,1}} \mathbb{E}_{\mathcal{W}} \left[\prod_{n=1}^{N'} \left\| \mathbf{I} - \frac{1}{\ell_{g,1} |\mathcal{S}_n|} \sum_{i=1}^{|\mathcal{S}_n|} \nabla_{\mathbf{y}}^2 g_i(\mathbf{x}, \mathbf{y}; \zeta_{i,n}) \right\| \right] \\ &\leq \frac{N}{\ell_{g,1}} \mathbb{E}_{N'} \left[1 - \frac{\mu_g}{\ell_{g,1}} \right]^{N'} = \frac{1}{\ell_{g,1}} \sum_{n=0}^{N-1} \left[1 - \frac{\mu_g}{\ell_{g,1}} \right]^n \leq \frac{1}{\mu_g}. \end{aligned}$$

The reminder of the proof is similar to (Ghadimi & Wang, 2018).

□

The following lemma extends (Ghadimi & Wang, 2018, Lemma 2.2) and (Chen et al., 2021a, Lemma 2) to the finite-sum problem (1). Proofs follow similarly by applying their analysis to the inner & outer functions (f_i, g_i) , $\forall i \in \mathcal{S}$.

Lemma C.1. *Under Assumptions A and B, for all $\mathbf{x}_1, \mathbf{x}_2$:*

$$\|\nabla f(\mathbf{x}_1) - \nabla f(\mathbf{x}_2)\| \leq L_f \|\mathbf{x}_1 - \mathbf{x}_2\|, \quad (17a)$$

$$\|\mathbf{y}^*(\mathbf{x}_1) - \mathbf{y}^*(\mathbf{x}_2)\| \leq L_y \|\mathbf{x}_1 - \mathbf{x}_2\|, \quad (17b)$$

$$\|\nabla \mathbf{y}^*(\mathbf{x}_1) - \nabla \mathbf{y}^*(\mathbf{x}_2)\| \leq L_{yx} \|\mathbf{x}_1 - \mathbf{x}_2\|, \quad (17c)$$

Also, for all $i \in \mathcal{S}$, $\nu \in \{0, \dots, \tau_i - 1\}$, $\mathbf{x}_1, \mathbf{x}_2$, and $\mathbf{y}$, we have:

$$\|\bar{\nabla} f_i(\mathbf{x}_1, \mathbf{y}) - \bar{\nabla} f_i(\mathbf{x}_1, \mathbf{y}^*(\mathbf{x}_1))\| \leq M_f \|\mathbf{y}^*(\mathbf{x}_1) - \mathbf{y}\|, \quad (17d)$$

$$\|\bar{\nabla} f_i(\mathbf{x}_2, \mathbf{y}) - \bar{\nabla} f_i(\mathbf{x}_1, \mathbf{y})\| \leq M_f \|\mathbf{x}_2 - \mathbf{x}_1\|, \quad (17e)$$

$$\mathbb{E} [\|\bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}) - \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y})\|^2] \leq \tilde{\sigma}_f^2, \quad (17f)$$

$$\mathbb{E} [\|\mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+)\|^2 | \mathcal{F}_{i,\nu-1}] \leq \tilde{D}_f^2. \quad (17g)$$

Here,

$$\begin{aligned} L_y &:= \frac{\ell_{g,1}}{\mu_g} = \mathcal{O}(\kappa_g), \\ L_{yx} &:= \frac{\ell_{g,2} + \ell_{g,2} L_y}{\mu_g} + \frac{\ell_{g,1}}{\mu_g^2} (\ell_{g,2} + \ell_{g,2} L_y) = \mathcal{O}(\kappa_g^3), \\ M_f &:= \ell_{f,1} + \frac{\ell_{g,1} \ell_{f,1}}{\mu_g} + \frac{\ell_{f,0}}{\mu_g} \left(\ell_{g,2} + \frac{\ell_{g,1} \ell_{g,2}}{\mu_g} \right) = \mathcal{O}(\kappa_g^2), \\ L_f &:= \ell_{f,1} + \frac{\ell_{g,1} (\ell_{f,1} + M_f)}{\mu_g} + \frac{\ell_{f,0}}{\mu_g} \left(\ell_{g,2} + \frac{\ell_{g,1} \ell_{g,2}}{\mu_g} \right) = \mathcal{O}(\kappa_g^3), \\ \tilde{\sigma}_f^2 &:= \sigma_f^2 + \frac{3}{\mu_g^2} \left((\sigma_f^2 + \ell_{f,0}^2) (\sigma_{g,2}^2 + 2\ell_{g,1}^2) + \sigma_f^2 \ell_{g,1}^2 \right), \\ \tilde{D}_f^2 &:= \left(\ell_{f,0} + \frac{\ell_{g,1}}{\mu_g} \ell_{f,1} + \ell_{g,1} \ell_{f,1} \frac{1}{\mu_g} \right)^2 + \tilde{\sigma}_f^2 = \mathcal{O}(\kappa_g^2), \end{aligned} \quad (18)$$

where the other constants are provided in Assumptions A and B.

C.1. Descent of Outer Objective

The following lemma characterizes the descent of the outer objective.

Lemma C.2 (Descent Lemma). *Suppose Assumptions A and B hold. Further, assume $\tau_i \geq 1$ and $\alpha_i = \alpha/\tau_i, \forall i \in \mathcal{S}$ for some positive constant α . Then, FEDOUT guarantees:*

$$\begin{aligned} \mathbb{E} [f(\mathbf{x}^+)] - \mathbb{E} [f(\mathbf{x})] &\leq -\frac{\alpha}{2} \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + \frac{\alpha^2 L_f}{2} \tilde{\sigma}_f^2 \\ &\quad - \frac{\alpha}{2} (1 - \alpha L_f) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] \\ &\quad + \frac{3\alpha}{2} \left(b^2 + M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \frac{M_f^2}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] \right). \end{aligned} \quad (19)$$

Proof. It follows from Algorithm 2 that $\mathbf{x}_{i,0} = \mathbf{x}, \forall i \in \mathcal{S}$, and

$$\mathbf{x}^+ = \mathbf{x} - \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+). \quad (20)$$

Now, using the Lipschitz property of ∇f in Lemma C.1, we have

$$\begin{aligned}
 \mathbb{E}[f(\mathbf{x}^+)] - \mathbb{E}[f(\mathbf{x})] &\leq \mathbb{E}[\langle \mathbf{x}^+ - \mathbf{x}, \nabla f(\mathbf{x}) \rangle] + \frac{L_f}{2} \mathbb{E}[\|\mathbf{x}^+ - \mathbf{x}\|^2] \\
 &= -\mathbb{E}\left[\left\langle \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+), \nabla f(\mathbf{x}) \right\rangle\right] \\
 &\quad + \frac{L_f}{2} \mathbb{E}\left[\left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2\right].
 \end{aligned} \tag{21}$$

In the following, we bound each term on the right hand side (RHS) of (21). For the first term, we have

$$\begin{aligned}
 -\mathbb{E}\left[\left\langle \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+), \nabla f(\mathbf{x}) \right\rangle\right] &= -\mathbb{E}\left[\frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbb{E}[\langle \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+), \nabla f(\mathbf{x}) \rangle \mid \mathcal{F}_{i,\nu-1}]\right] \\
 &= -\mathbb{E}\left[\left\langle \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+), \nabla f(\mathbf{x}) \right\rangle\right] \\
 &= -\frac{\alpha}{2} \mathbb{E}\left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2\right] - \frac{\alpha}{2} \mathbb{E}[\|\nabla f(\mathbf{x})\|^2] \\
 &\quad + \frac{\alpha}{2} \mathbb{E}\left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \nabla f(\mathbf{x}) \right\|^2\right],
 \end{aligned} \tag{22}$$

where the first equality follows from the law of total expectation; the second equality uses the fact that $\bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) = \mathbb{E}[\mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \mid \mathcal{F}_{i,\nu-1}]$; and the last equality is obtained from our assumption $\alpha_i = \alpha/\tau_i, \forall i \in \mathcal{S}$.

Next, we bound the last term in (22). Note that

$$\begin{aligned}
 \left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \nabla f(\mathbf{x}) \right\|^2 &= \left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} (\bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+)) \right\|^2 \\
 &\quad + \left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+) - \nabla f(\mathbf{x}) \right\|^2 \\
 &\leq 3 \left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} (\bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+)) \right\|^2 \\
 &\quad + 3 \left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} (\bar{\nabla} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+)) \right\|^2 \\
 &\quad + 3 \|\bar{\nabla} f(\mathbf{x}, \mathbf{y}^+) - \nabla f(\mathbf{x})\|^2,
 \end{aligned}$$

where the inequality uses Lemma G.1.

Hence,

$$\begin{aligned}
 &\mathbb{E}\left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \nabla f(\mathbf{x}) \right\|^2\right] \\
 &\leq 3b^2 + \frac{3M_f^2}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E}[\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 3M_f^2 \mathbb{E}[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2],
 \end{aligned} \tag{23}$$

where the inequality uses Lemmas 2.2 and C.1.

Substituting (23) into (22) yields

$$\begin{aligned}
 & -\mathbb{E} \left[\left\langle \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+), \nabla f(\mathbf{x}) \right\rangle \right] \\
 & \leq -\frac{\alpha}{2} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] - \frac{\alpha}{2} \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] \\
 & + \frac{3\alpha}{2} \left(b^2 + M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \frac{M_f^2}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] \right).
 \end{aligned} \tag{24}$$

Next, we bound the second term on the RHS of (21). Observe that

$$\begin{aligned}
 & \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] \\
 & = \alpha^2 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} (\mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) + \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+)) \right\|^2 \right] \\
 & \leq \alpha^2 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] + \alpha^2 \tilde{\sigma}_f^2,
 \end{aligned} \tag{25}$$

where the inequality follows from Lemmas G.3 and C.1.

Plugging (25) and (24) into (21) completes the proof. $\square$

C.2. Error of FEDINN

The following lemma establishes the progress of FEDINN. It should be mentioned that the assumption on $\beta_i, \forall i \in \mathcal{S}$ is identical to the one listed in (Mitra et al., 2021, Theorem 4).

Lemma C.3 (Error of FEDINN). *Suppose Assumptions A and B hold. Further, assume*

$$\tau_i \geq 1, \quad \alpha_i = \frac{\alpha}{\tau_i}, \quad \beta_i = \frac{\beta}{\tau_i}, \quad \forall i \in \mathcal{S},$$

where $0 < \beta < \min(1/(6\ell_{g,1}), 1)$ and α is some positive constant. Then, FEDINN guarantees:

$$\mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] \leq \left(1 - \frac{\beta\mu_g}{2}\right)^T \mathbb{E} [\|\mathbf{y} - \mathbf{y}^*(\mathbf{x})\|^2] + 25T\beta^2\sigma_{g,1}^2, \quad \text{and} \tag{26a}$$

$$\begin{aligned}
 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x}^+)\|^2] & \leq a_1(\alpha) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] \\
 & + a_2(\alpha) \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + a_3(\alpha) \tilde{\sigma}_f^2.
 \end{aligned} \tag{26b}$$

Here,

$$\begin{aligned}
 a_1(\alpha) & := L_{\mathbf{y}}^2 \alpha^2 + \frac{L_{\mathbf{y}} \alpha}{4M_f} + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{2\eta}, \\
 a_2(\alpha) & := 1 + 4M_f L_{\mathbf{y}} \alpha + \frac{\eta L_{\mathbf{y}\mathbf{x}} \tilde{D}_f^2 \alpha^2}{2}, \\
 a_3(\alpha) & := \alpha^2 L_{\mathbf{y}}^2 + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{2\eta},
 \end{aligned} \tag{27}$$

for any $\eta > 0$.

Proof. Note that

$$\begin{aligned} \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(x^+)\|^2] &= \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(x)\|^2] + \mathbb{E} [\|\mathbf{y}^*(x^+) - \mathbf{y}^*(x)\|^2] \\ &\quad + 2\mathbb{E} [\langle \mathbf{y}^+ - \mathbf{y}^*(x), \mathbf{y}^*(x) - \mathbf{y}^*(x^+) \rangle]. \end{aligned} \quad (28)$$

Next, we upper bound each term on the RHS of (28).

Bounding the first term in (28):

From (Mitra et al., 2021, Theorem 4), for all $t \in \{0, \dots, T-1\}$, we obtain

$$\mathbb{E} [\|\mathbf{y}^{t+1} - \mathbf{y}^*(x)\|^2] \leq \left(1 - \frac{\beta\mu_g}{2}\right) \mathbb{E} [\|\mathbf{y}^t - \mathbf{y}^*(x)\|^2] + 25\beta^2\sigma_{g,1}^2,$$

which together with our setting $\mathbf{y}^+ = \mathbf{y}^T$ implies

$$\mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(x)\|^2] \leq \left(1 - \frac{\beta\mu_g}{2}\right)^T \mathbb{E} [\|\mathbf{y} - \mathbf{y}^*(x^*)\|^2] + 25T\beta^2\sigma_{g,1}^2. \quad (29)$$

Bounding the second term in (28):

By similar steps as in (25), we have

$$\begin{aligned} \mathbb{E} [\|\mathbf{y}^*(x^+) - \mathbf{y}^*(x)\|^2] &\leq L_{\mathbf{y}}^2 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(x, \mathbf{y}^+) \right\|^2 \right] \\ &\leq L_{\mathbf{y}}^2 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(x_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] + \alpha^2 L_{\mathbf{y}}^2 \tilde{\sigma}_f^2, \end{aligned} \quad (30)$$

where the inequalities are obtained from Lemmas C.1 and G.3.

Bounding the third term in (28):

Observe that

$$\begin{aligned} \mathbb{E} [\langle \mathbf{y}^+ - \mathbf{y}^*(x), \mathbf{y}^*(x) - \mathbf{y}^*(x^+) \rangle] &= -\mathbb{E} [\langle \mathbf{y}^+ - \mathbf{y}^*(x), \nabla \mathbf{y}^*(x)(x^+ - x) \rangle] \\ &\quad - \mathbb{E} [\langle \mathbf{y}^+ - \mathbf{y}^*(x), \mathbf{y}^*(x^+) - \mathbf{y}^*(x) - \nabla \mathbf{y}^*(x)(x^+ - x) \rangle]. \end{aligned} \quad (31)$$

For the first term on the R.H.S. of the above equality, we have

$$\begin{aligned} -\mathbb{E} [\langle \mathbf{y}^+ - \mathbf{y}^*(x), \nabla \mathbf{y}^*(x)(x^+ - x) \rangle] &= -\mathbb{E} \left[\langle \mathbf{y}^+ - \mathbf{y}^*(x), \frac{1}{m} \nabla \mathbf{y}^*(x) \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(x_{i,\nu}, \mathbf{y}^+) \rangle \right] \\ &\leq \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(x)\| \left\| \frac{1}{m} \nabla \mathbf{y}^*(x) \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(x_{i,\nu}, \mathbf{y}^+) \right\| \right] \\ &\leq L_{\mathbf{y}} \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(x)\| \left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(x_{i,\nu}, \mathbf{y}^+) \right\| \right] \\ &\leq 2\gamma \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(x)\|^2] + \frac{L_{\mathbf{y}}^2 \alpha^2}{8\gamma} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(x_{i,\nu}, \mathbf{y}^+) \right\|^2 \right], \end{aligned} \quad (32)$$

where the first equality uses the fact that $\bar{\mathbf{h}}_i(x_{i,\nu}, \mathbf{y}^+) = \mathbb{E} [\mathbf{h}_i(x_{i,\nu}, \mathbf{y}^+) | \mathcal{F}_{i,\nu-1}]$; the second inequality follows from Lemma C.1; and the last inequality is obtained from the Young's inequality such that $ab \leq 2\gamma a^2 + \frac{b^2}{8\gamma}$.

Further, using Lemma C.1, we have

$$\begin{aligned} &-\mathbb{E} [\langle \mathbf{y}^+ - \mathbf{y}^*(x), \mathbf{y}^*(x^+) - \mathbf{y}^*(x) - \nabla \mathbf{y}^*(x)(x^+ - x) \rangle] \\ &\leq \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(x)\| \|\mathbf{y}^*(x^+) - \mathbf{y}^*(x) - \nabla \mathbf{y}^*(x)(x^+ - x)\|] \\ &\leq \frac{L_{\mathbf{y}x}}{2} \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(x)\| \|x^+ - x\|^2], \end{aligned} \quad (33)$$

where the inequality follows from Lemma C.1.

From Algorithm 2, we have

$$\mathbf{x}_{i,0} = \mathbf{x}, \forall i \in \mathcal{S}, \quad \text{and} \quad \mathbf{x}^+ = \mathbf{x} - \frac{1}{m} \sum_{i=1}^m \frac{\alpha}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+). \quad (34)$$

Note that $\mathcal{F}_0 = \mathcal{F}_{i,0}$ for all $i \in \mathcal{S}$. Hence,

$$\begin{aligned} & \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2 \|\mathbf{x}^+ - \mathbf{x}\|^2 \right] \\ & \leq \frac{1}{m} \sum_{i=1}^m \frac{\alpha^2}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2 \mathbb{E} \left[\|\mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+)\|^2 \mid \mathcal{F}_{i,\tau_i-1} \right] \right] \\ & \leq \alpha^2 \tilde{D}_f^2 \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2 \right], \end{aligned} \quad (35)$$

where the last inequality uses Lemma C.1.

Note also that for any $\eta > 0$, we have $1 \leq \frac{\eta}{2} + \frac{1}{2\eta}$. Combining this inequality with (33) and using (35) give

$$\begin{aligned} & -\mathbb{E}[\langle \mathbf{y}^+ - \mathbf{y}^*(\mathbf{x}), \mathbf{y}^*(\mathbf{x}^+) - \mathbf{y}^*(\mathbf{x}) - \nabla \mathbf{y}^*(\mathbf{x})(\mathbf{x}^+ - \mathbf{x}) \rangle] \\ & \leq \frac{L_{\mathbf{y}\mathbf{x}}}{2} \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\| \|\mathbf{x}^+ - \mathbf{x}\|^2 \right] \\ & \leq \frac{\eta L_{\mathbf{y}\mathbf{x}}}{4} \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\| \|\mathbf{x}^+ - \mathbf{x}\|^2 \right] + \frac{L_{\mathbf{y}\mathbf{x}}}{4\eta} \mathbb{E} \left[\|\mathbf{x}^+ - \mathbf{x}\|^2 \right] \\ & \leq \frac{\eta L_{\mathbf{y}\mathbf{x}} \tilde{D}_f^2 \alpha^2}{4} \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2 \right] + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{4\eta} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{4\eta} \tilde{\sigma}_f^2, \end{aligned} \quad (36)$$

where the last inequality uses (35) and Lemma 2.2.

Let $\gamma = M_f L_{\mathbf{y}} \alpha$. Plugging (36) and (32) into (31), we have

$$\begin{aligned} \mathbb{E}[\langle \mathbf{y}^+ - \mathbf{y}^*(\mathbf{x}), \mathbf{y}^*(\mathbf{x}) - \mathbf{y}^*(\mathbf{x}^+) \rangle] & \leq \left(2\gamma + \frac{\eta L_{\mathbf{y}\mathbf{x}} \tilde{D}_f^2}{4} \alpha^2 \right) \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2 \right] \\ & \quad + \left(\frac{L_{\mathbf{y}}^2 \alpha^2}{8\gamma} + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{4\eta} \right) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{4\eta} \tilde{\sigma}_f^2 \\ & = \left(2M_f L_{\mathbf{y}} \alpha + \frac{\eta L_{\mathbf{y}\mathbf{x}} \tilde{D}_f^2}{4} \alpha^2 \right) \mathbb{E} \left[\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2 \right] \\ & \quad + \left(\frac{L_{\mathbf{y}} \alpha}{8M_f} + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{4\eta} \right) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] + \frac{L_{\mathbf{y}\mathbf{x}} \alpha^2}{4\eta} \tilde{\sigma}_f^2. \end{aligned} \quad (37)$$

Substituting (37), (30), and (29) into (28) completes the proof. $\square$

C.3. Drifting Errors of FEDOUT

The following lemma provides a bound on the *drift* of each $\mathbf{x}_{i,\nu}$ from $\mathbf{x}$ for stochastic nonconvex bilevel problems. It should be mentioned that similar drifting bounds for *single-level* problems are provided under either strong convexity (Mitra et al., 2021) and/or bounded dissimilarity assumptions (Wang et al., 2020; Reddi et al., 2020; Li et al., 2020b).

Lemma C.4 (Drifting Error of FEDOUT). *Suppose Assumptions A and B hold. Further, assume $\tau_i \geq 1$ and $\alpha_i \leq 1/(5M_f \tau_i)$, $\forall i \in \mathcal{S}$. Then, for each $i \in \mathcal{S}$ and $\forall \nu \in \{0, \dots, \tau_i - 1\}$, FEDOUT guarantees:*

$$\begin{aligned} \mathbb{E} \left[\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2 \right] & \leq 36\tau_i^2 \alpha_i^2 (M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 3b^2) \\ & \quad + 27\tau_i \alpha_i^2 \tilde{\sigma}_f^2. \end{aligned} \quad (38)$$

Proof. The result trivially holds for $\tau_i = 1$. Let $\tau_i > 1$ and define

$$\begin{aligned} \mathbf{v}_{i,\nu} &:= \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\mathbf{h}}_i(\mathbf{x}, \mathbf{y}^+) \\ &\quad + \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+) + \bar{\mathbf{h}}(\mathbf{x}, \mathbf{y}^+) - \bar{\nabla} f(\mathbf{x}, \mathbf{y}^+), \\ \mathbf{w}_{i,\nu} &:= \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) + \bar{\mathbf{h}}_i(\mathbf{x}, \mathbf{y}^+) \\ &\quad - \mathbf{h}_i(\mathbf{x}, \mathbf{y}^+) + \mathbf{h}(\mathbf{x}, \mathbf{y}^+) - \bar{\mathbf{h}}(\mathbf{x}, \mathbf{y}^+), \\ \mathbf{z}_{i,\nu} &:= \bar{\nabla} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+) + \bar{\nabla} f(\mathbf{x}, \mathbf{y}^+) - \nabla f(\mathbf{x}) + \nabla f(\mathbf{x}). \end{aligned} \quad (39)$$

One will notice that

$$\mathbf{v}_{i,\nu} + \mathbf{w}_{i,\nu} + \mathbf{z}_{i,\nu} = \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \mathbf{h}_i(\mathbf{x}, \mathbf{y}^+) + \mathbf{h}(\mathbf{x}, \mathbf{y}^+).$$

Hence, from Algorithm 2, for each $i \in \mathcal{S}$, and $\forall \nu \in \{0, \dots, \tau_i - 1\}$, we have

$$\mathbf{x}_{i,\nu+1} - \mathbf{x} = \mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i(\mathbf{v}_{i,\nu} + \mathbf{w}_{i,\nu} + \mathbf{z}_{i,\nu}), \quad (40)$$

which implies that

$$\begin{aligned} \mathbb{E} [\|\mathbf{x}_{i,\nu+1} - \mathbf{x}\|^2] &= \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i(\mathbf{v}_{i,\nu} + \mathbf{z}_{i,\nu})\|^2] + \alpha_i^2 \mathbb{E} [\|\mathbf{w}_{i,\nu}\|^2] \\ &\quad - 2\mathbb{E} [\mathbb{E} [\langle \mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i(\mathbf{v}_{i,\nu} + \mathbf{z}_{i,\nu}), \alpha_i \mathbf{w}_{i,\nu} \rangle \mid \mathcal{F}_{i,\nu-1}]] \\ &= \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i(\mathbf{v}_{i,\nu} + \mathbf{z}_{i,\nu})\|^2] + \alpha_i^2 \mathbb{E} [\|\mathbf{w}_{i,\nu}\|^2]. \end{aligned} \quad (41)$$

Here, the last equality uses Lemma G.3 since $\mathbb{E}[\mathbf{w}_{i,\nu} \mid \mathcal{F}_{i,\nu-1}] = 0$, by definition.

From Lemmas G.1, 2.2, and C.1, for $\mathbf{v}_{i,\nu}$, $\mathbf{w}_{i,\nu}$, and $\mathbf{z}_{i,\nu}$ defined in (39), we have

$$\begin{aligned} \mathbb{E} [\|\mathbf{v}_{i,\nu}\|^2] &\leq 3\mathbb{E} [\|\bar{\nabla} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+)\|^2 \\ &\quad + \|\bar{\mathbf{h}}_i(\mathbf{x}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+)\|^2 + \|\bar{\nabla} f(\mathbf{x}, \mathbf{y}^+) - \bar{\mathbf{h}}(\mathbf{x}, \mathbf{y}^+)\|^2] \\ &\leq 9b^2, \end{aligned} \quad (42a)$$

$$\begin{aligned} \mathbb{E} [\|\mathbf{w}_{i,\nu}\|^2] &\leq 3\mathbb{E} [\|\mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+)\|^2 \\ &\quad + \|\bar{\mathbf{h}}_i(\mathbf{x}, \mathbf{y}^+) - \mathbf{h}_i(\mathbf{x}, \mathbf{y}^+)\|^2 + \|\mathbf{h}(\mathbf{x}, \mathbf{y}^+) - \bar{\mathbf{h}}(\mathbf{x}, \mathbf{y}^+)\|^2] \\ &\leq 9\tilde{\sigma}_f^2, \end{aligned} \quad (42b)$$

and

$$\begin{aligned} \mathbb{E} [\|\mathbf{z}_{i,\nu}\|^2] &\leq 3\mathbb{E} [\|\bar{\nabla} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) - \bar{\nabla} f_i(\mathbf{x}, \mathbf{y}^+)\|^2 \\ &\quad + \|\bar{\nabla} f(\mathbf{x}, \mathbf{y}^+) - \nabla f(\mathbf{x})\|^2 + \|\nabla f(\mathbf{x})\|^2] \\ &\leq 3(M_f^2 \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2]). \end{aligned} \quad (42c)$$

Now, the first term in the RHS of (41) can be bounded as follows:

$$\begin{aligned} \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i(\mathbf{v}_{i,\nu} + \mathbf{z}_{i,\nu})\|^2] &\leq \left(1 + \frac{1}{2\tau_i - 1}\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 2\tau_i \mathbb{E} [\|\alpha_i(\mathbf{v}_{i,\nu} + \mathbf{z}_{i,\nu})\|^2] \\ &\leq \left(1 + \frac{1}{2\tau_i - 1}\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 4\tau_i \alpha_i^2 (\mathbb{E} [\|\mathbf{z}_{i,\nu}\|^2] + \mathbb{E} [\|\mathbf{v}_{i,\nu}\|^2]) \\ &\leq \left(1 + \frac{1}{2\tau_i - 1}\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 4\tau_i \alpha_i^2 (\mathbb{E} [\|\mathbf{z}_{i,\nu}\|^2] + 9b^2) \\ &\leq \left(1 + \frac{1}{2\tau_i - 1} + 12\tau_i \alpha_i^2 M_f^2\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] \\ &\quad + 12\tau_i \alpha_i^2 (M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 3b^2). \end{aligned} \quad (43)$$

Here, the first inequality follows from Lemma G.2; the second inequality uses Lemma G.1; and the third and last inequalities follow from (42a) and (42c).

Substituting (43) into (41) gives

$$\begin{aligned}
 \mathbb{E} [\|\mathbf{x}_{i,\nu+1} - \mathbf{x}\|^2] &\leq \left(1 + \frac{1}{2\tau_i - 1} + 12\tau_i\alpha_i^2 M_f^2\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] \\
 &\quad + 12\tau_i\alpha_i^2 (M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 3b^2) + 9\alpha_i^2 \tilde{\sigma}_f^2 \\
 &\leq \left(1 + \frac{1}{\tau_i - 1}\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] \\
 &\quad + 12\tau_i\alpha_i^2 (M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 3b^2) + 9\alpha_i^2 \tilde{\sigma}_f^2.
 \end{aligned} \tag{44}$$

Here, the first inequality uses (42b) and the last inequality follows by noting $\alpha_i \leq 1/(5M_f\tau_i)$.

For all $\tau_i > 1$, we have

$$\begin{aligned}
 \sum_{j=0}^{\nu-1} \left(1 + \frac{1}{\tau_i - 1}\right)^j &= \frac{\left(1 + \frac{1}{\tau_i - 1}\right)^\nu - 1}{\left(1 + \frac{1}{\tau_i - 1}\right) - 1} \\
 &\leq \tau_i \left(1 + \frac{1}{\tau_i}\right)^\nu \leq \tau_i \left(1 + \frac{1}{\tau_i}\right)^{\tau_i} \leq \exp(1)\tau_i < 3\tau_i.
 \end{aligned} \tag{45}$$

Now, iterating equation (44) and using $\mathbf{x}_{i,0} = \mathbf{x}, \forall i \in \mathcal{S}$, we obtain

$$\begin{aligned}
 \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] &\leq (12\tau_i\alpha_i^2 (M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 3b^2) + 9\alpha_i^2 \tilde{\sigma}_f^2) \sum_{j=0}^{\nu-1} \left(1 + \frac{1}{\tau_i - 1}\right)^j \\
 &\leq 3\tau_i (12\tau_i\alpha_i^2 (M_f^2 \mathbb{E} [\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2] + \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 3b^2) + 9\alpha_i^2 \tilde{\sigma}_f^2),
 \end{aligned}$$

where the second inequality uses (45). This completes the proof. $\square$

Remark C.1. Lemma C.4 shows that the bound on the client-drift scales linearly with τ_i and the inner error $\|\mathbf{y}^+ - \mathbf{y}^*(\mathbf{x})\|^2$ in general nested FL. We aim to control such a drift by selecting $\alpha_i = \mathcal{O}(1/\tau_i)$ for all $i \in \mathcal{S}$ and using the inner error bound provided in Lemma C.3.

Next, we provide the proof of our main result which can be adapted to general nested problems (bilevel, min-max, compositional).

C.4. Proof of Theorem 3.1

Proof. We define the following Lyapunov function

$$\mathbb{W}^k := f(\mathbf{x}^k) + \frac{M_f}{L_y} \|\mathbf{y}^k - \mathbf{y}^*(\mathbf{x}^k)\|^2. \tag{46}$$

Motivated by (Chen et al., 2021a), we bound the difference between two Lyapunov functions. That is,

$$\mathbb{W}^{k+1} - \mathbb{W}^k = f(\mathbf{x}^{k+1}) - f(\mathbf{x}^k) + \frac{M_f}{L_y} (\|\mathbf{y}^{k+1} - \mathbf{y}^*(\mathbf{x}^{k+1})\|^2 - \|\mathbf{y}^k - \mathbf{y}^*(\mathbf{x}^k)\|^2). \tag{47}$$

The first two terms on the RHS of (47) quantifies the descent of outer objective f and the reminding terms measure the descent of the inner errors.

From our assumption, we have $\alpha_i^k = \alpha_k/\tau_i, \beta_i^k = \beta_k/\tau_i, \forall i \in \mathcal{S}$. Substituting these stepsizes into the bounds provided in

Lemmas C.2 and C.3, and using (47), we get

$$\begin{aligned} \mathbb{E}[\mathbb{W}^{k+1}] - \mathbb{E}[\mathbb{W}^k] &\leq \left(\frac{L_f \alpha_k^2}{2} + \frac{M_f}{L_y} a_3(\alpha_k) \right) \tilde{\sigma}_f^2 + \frac{3\alpha_k}{2} b^2, \\ &\quad - \frac{\alpha_k}{2} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] + \frac{3M_f^2 \alpha_k}{2m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E}[\|\mathbf{x}_{i,\nu}^k - \mathbf{x}^k\|^2] \end{aligned} \quad (48a)$$

$$- \left(\frac{\alpha_k}{2} - \frac{L_f \alpha_k^2}{2} - \frac{M_f}{L_y} a_1(\alpha_k) \right) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \bar{\mathbf{h}}_i(\mathbf{x}_{i,\nu}^k, \mathbf{y}^{k+1}) \right\|^2 \right] \quad (48b)$$

$$+ \frac{M_f}{L_y} \left(\frac{3M_f L_y \alpha_k}{2} + a_2(\alpha_k) \right) \mathbb{E}[\|\mathbf{y}^{k+1} - \mathbf{y}^*(\mathbf{x}^k)\|^2] - \frac{M_f}{L_y} \mathbb{E}[\|\mathbf{y}^k - \mathbf{y}^*(\mathbf{x}^k)\|^2], \quad (48c)$$

where $a_1(\alpha) - a_3(\alpha)$ are defined in (27).

Let

$$\alpha_i^k = \frac{\alpha_k}{\tau_i}, \quad \forall i \in \mathcal{S} \quad \text{where} \quad \alpha_k \leq \frac{1}{216M_f^2 + 5M_f}. \quad (49)$$

The above choice of α_k satisfies the condition of Lemma C.4 and we have $54M_f^2 \alpha_k^3 \leq \alpha_k^2/4$. Hence, from Lemma C.4, we get

$$\begin{aligned} (48a) &\leq -\frac{\alpha_k}{2} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] + \frac{81}{2} \alpha_k^3 M_f^2 \tilde{\sigma}_f^2 \\ &\quad + 54M_f^2 \alpha_k^3 (M_f^2 \mathbb{E}[\|\mathbf{y}^{k+1} - \mathbf{y}^*(\mathbf{x}^k)\|^2] + \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] + 3b^2) \\ &\leq -\frac{\alpha_k}{4} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] + \frac{\alpha_k^2}{4} \tilde{\sigma}_f^2 + \frac{\alpha_k^2}{4} (M_f^2 \mathbb{E}[\|\mathbf{y}^{k+1} - \mathbf{y}^*(\mathbf{x}^k)\|^2] + 3b^2) \\ &\leq -\frac{\alpha_k}{4} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] + \frac{\alpha_k^2}{4} \tilde{\sigma}_f^2 + \frac{3\alpha_k^2}{4} b^2 \\ &\quad + \frac{25M_f}{L_y} \left(\frac{M_f L_y}{4} \alpha_k^2 \right) T \beta_k^2 \sigma_{g,1}^2 + \frac{M_f}{L_y} \left(\frac{M_f L_y}{4} \alpha_k^2 \right) \left(1 - \frac{\beta_k \mu_g}{2} \right)^T, \end{aligned} \quad (50)$$

where the first inequality uses (49) and the last inequality follows from (26a).

To guarantee the descent of $\mathbb{W}^k$, the following constraints need to be satisfied

$$\begin{aligned} (48b) &\leq 0, \\ \implies \frac{\alpha_k}{2} - \frac{L_f \alpha_k^2}{2} - \frac{M_f}{L_y} \left(L_y^2 \alpha_k^2 + \frac{L_y \alpha_k}{4M_f} + \frac{L_y \alpha_k^2}{2\eta} \right) &\geq 0, \\ \implies \alpha_k &\leq \frac{1}{2L_f + 4M_f L_y + \frac{2M_f L_y \alpha_k}{L_y \eta}}, \end{aligned} \quad (51)$$

where the second line uses (27).

Further, substituting (26) in (48c) gives

$$\begin{aligned} (48c) &\leq \frac{25M_f}{L_y} \left(\frac{3M_f L_y \alpha_k}{2} + a_2(\alpha_k) \right) T \beta_k^2 \sigma_{g,1}^2 \\ &\quad + \frac{M_f}{L_y} \left(\left(\frac{3M_f L_y \alpha_k}{2} + a_2(\alpha_k) \right) \left(1 - \frac{\beta_k \mu_g}{2} \right)^T - 1 \right) \mathbb{E}[\|\mathbf{y}^k - \mathbf{y}^*(\mathbf{x}^k)\|^2]. \end{aligned} \quad (52)$$

Substituting (50)–(52) into (48) gives

$$\begin{aligned}
 \mathbb{E}[\mathbb{W}^{k+1}] - \mathbb{E}[\mathbb{W}^k] &\leq -\frac{\alpha_k}{4} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] \\
 &\quad + \left(\frac{3}{2}\alpha_k + \frac{3}{4}\alpha_k^2\right)b^2 \\
 &\quad + \left(\left(\frac{L_f}{2} + \frac{1}{4}\right)\alpha_k^2 + \frac{M_f}{L_y}a_3(\alpha_k)\right)\tilde{\sigma}_f^2 \\
 &\quad + \frac{25M_f}{L_y} \left(\frac{M_f L_y}{4}\alpha_k^2 + \frac{3M_f L_y \alpha_k}{2} + a_2(\alpha_k)\right) T\beta_k^2 \sigma_{g,1}^2 \\
 &\quad + \frac{M_f}{L_y} \left(\left(\frac{M_f L_y}{4}\alpha_k^2 + \frac{3M_f L_y \alpha_k}{2} + a_2(\alpha_k)\right) \left(1 - \frac{\beta_k \mu_g}{2}\right)^T - 1\right) \mathbb{E}[\|\mathbf{y}^k - \mathbf{y}^*(\mathbf{x}^k)\|^2]. \quad (53a)
 \end{aligned}$$

Let $\beta_k < \min(1/(6\ell_{g,1}), 1)$. Then, we have $\beta_k \mu_g/2 < 1$. This together with (27) implies that for any $\alpha_k > 0$

$$\begin{aligned}
 (53a) &\leq 0, \\
 &\Rightarrow \left(1 + \frac{M_f L_y}{4}\alpha_k^2 + \frac{11M_f L_y \alpha_k}{2} + \frac{\eta L_{yx} \tilde{D}_f^2 \alpha_k^2}{2}\right) \left(1 - \frac{\beta_k \mu_g}{2}\right)^T - 1 \leq 0, \\
 &\Rightarrow \exp\left(\frac{M_f L_y}{4}\alpha_k^2 + \frac{11M_f L_y \alpha_k}{2} + \frac{\eta L_{yx} \tilde{D}_f^2 \alpha_k^2}{2}\right) \exp\left(-\frac{T\beta_k \mu_g}{2}\right) - 1 \leq 0, \\
 &\Rightarrow \beta_k \geq \frac{11M_f L_y + \eta L_{yx} \tilde{D}_f^2 \alpha_k + \frac{M_f L_y \alpha_k}{2}}{\mu_g} \cdot \frac{\alpha_k}{T}. \quad (54)
 \end{aligned}$$

From (49), (51) and (54), we select

$$\alpha_k = \min\{\bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3, \frac{\bar{\alpha}}{\sqrt{K}}\}, \quad \beta_k = \frac{\bar{\beta} \alpha_k}{T}, \quad (55)$$

where

$$\begin{aligned}
 \bar{\beta} &:= \frac{1}{\mu_g} \left(11M_f L_y + \eta L_{yx} \tilde{D}_f^2 \bar{\alpha}_1 + \frac{M_f L_y \bar{\alpha}_1}{2}\right), \\
 \bar{\alpha}_1 &:= \frac{1}{2L_f + 4M_f L_y + \frac{2M_f L_{yx}}{L_y \eta}}, \quad \bar{\alpha}_2 := \frac{T}{8\ell_{g,1} \bar{\beta}}, \quad \bar{\alpha}_3 := \frac{1}{216M_f^2 + 5M_f}, \quad (56)
 \end{aligned}$$

With the above choice of stepsizes, (53) can be simplified as

$$\begin{aligned}
 \mathbb{E}[\mathbb{W}^{k+1}] - \mathbb{E}[\mathbb{W}^k] &\leq -\frac{\alpha_k}{4} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] \\
 &\quad + \left(\frac{3}{2}\alpha_k + \frac{3}{4}\alpha_k^2\right)b^2 \\
 &\quad + \left(\frac{L_f + \frac{1}{2}}{2}\alpha_k^2 + \frac{M_f}{L_y}a_3(\alpha_k)\right)\tilde{\sigma}_f^2 \\
 &\quad + \frac{25M_f}{L_y} \left(\frac{M_f L_y}{4}\alpha_k^2 + \frac{3M_f L_y \alpha_k}{2} + a_2(\alpha_k)\right) T\beta_k^2 \sigma_{g,1}^2 \\
 &\leq -\frac{\alpha_k}{4} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] + c_1 \alpha_k^2 \sigma_{g,1}^2 + \left(\frac{3}{2}\alpha_k + \frac{3}{4}\alpha_k^2\right)b^2 + c_2 \alpha_k^2 \tilde{\sigma}_f^2, \quad (57)
 \end{aligned}$$

where the constants c_1 and c_2 are defined as

$$\begin{aligned}
 c_1 &= \frac{25M_f}{L_y} \left(1 + \frac{11M_f L_y}{2} \bar{\alpha}_1 + \left(\frac{M_f L_y + 2\eta L_{yx} \tilde{D}_f^2}{4}\right) \bar{\alpha}_1^2\right) \bar{\beta}^2 \frac{1}{T}, \\
 c_2 &= \frac{L_f + \frac{1}{2}}{2} + M_f L_y + \frac{L_{yx} M_f}{4\eta L_y}. \quad (58)
 \end{aligned}$$

Then telescoping gives

$$\begin{aligned}
 \frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] &\leq \frac{4}{\sum_{k=0}^{K-1} \alpha_k} \left(\Delta_{\mathbb{W}} + \sum_{k=0}^{K-1} \frac{3}{2} \left(\alpha_k + \frac{\alpha_k^2}{2} \right) b^2 + c_1 \alpha_k^2 \sigma_{g,1}^2 + c_2 \alpha_k^2 \tilde{\sigma}_f^2 \right) \\
 &\leq \frac{4\Delta_{\mathbb{W}}}{\min\{\bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3\}K} + \frac{4\Delta_{\mathbb{W}}}{\bar{\alpha}\sqrt{K}} + 6 \left(1 + \frac{\bar{\alpha}}{2\sqrt{K}} \right) b^2 + \frac{4c_1\bar{\alpha}}{\sqrt{K}} \sigma_{g,1}^2 + \frac{4c_2\bar{\alpha}}{\sqrt{K}} \tilde{\sigma}_f^2 \\
 &= \mathcal{O} \left(\frac{1}{\min\{\bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3\}K} + \frac{\bar{\alpha} \max(\sigma_{g,1}^2, \sigma_{g,2}^2, \sigma_f^2)}{\sqrt{K}} + b^2 \right),
 \end{aligned} \tag{59}$$

where $\Delta_{\mathbb{W}} := \mathbb{W}^0 - \mathbb{E}[\mathbb{W}^K]$. □

C.5. Proof of Corollary 3.1

Proof. Let $\eta = \frac{M_f}{L_y} = \mathcal{O}(\kappa_g)$ in (58). It follows from (18), (56), and (58) that

$$\bar{\alpha}_1 = \mathcal{O}(\kappa_g^{-3}), \quad \bar{\alpha}_2 = \mathcal{O}(T\kappa_g^{-3}), \quad \bar{\alpha}_3 = \mathcal{O}(\kappa_g^{-4}), \quad c_1 = \mathcal{O}(\kappa_g^9/T), \quad c_2 = \mathcal{O}(\kappa_g^3). \tag{60}$$

Further, $N = \mathcal{O}(\kappa_g \log K)$ gives $b = \frac{1}{K^{1/4}}$. Now, if we select $\bar{\alpha} = \mathcal{O}(\kappa_g^{-2.5})$ and $T = \mathcal{O}(\kappa_g^4)$, Eq. (59) gives

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O} \left(\frac{\kappa_g^4}{K} + \frac{\kappa_g^{2.5}}{\sqrt{K}} \right).$$

To achieve ε -optimal solution, we need $K = \mathcal{O}(\kappa_g^5 \varepsilon^{-2})$, and the samples in ξ and ζ are $\mathcal{O}(\kappa_g^5 \varepsilon^{-2})$ and $\mathcal{O}(\kappa_g^9 \varepsilon^{-2})$, respectively. □

D. Proof for Federated Minimax Optimization

Note that the minimax optimization problem (2) has the following bilevel form

$$\begin{aligned}
 \min_{\mathbf{x} \in \mathbb{R}^{d_1}} \quad & f(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) \\
 \text{subj. to} \quad & \mathbf{y}^*(\mathbf{x}) = \operatorname{argmin}_{\mathbf{y} \in \mathbb{R}^{d_2}} - \frac{1}{m} \sum_{i=1}^m f_i(\mathbf{x}, \mathbf{y}).
 \end{aligned} \tag{61a}$$

Here,

$$f_i(\mathbf{x}, \mathbf{y}) = \mathbb{E}_{\xi \sim \mathcal{C}_i} [f_i(\mathbf{x}, \mathbf{y}; \xi)] \tag{61b}$$

is the loss functions of the i^{th} client.

In this case, the hypergradient of (61) is

$$\nabla f_i(\mathbf{x}) = \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) + \nabla_{\mathbf{x}} \mathbf{y}^*(\mathbf{x})^\top \nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) = \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y}^*(\mathbf{x})), \tag{62}$$

where the second equality follows from the optimality condition of the inner problem, i.e., $\nabla_{\mathbf{y}} f(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) = 0$.

For each $i \in \mathcal{S}$, we can approximate $\nabla f_i(\mathbf{x})$ on a vector $\mathbf{y}$ in place of $\mathbf{y}^*(\mathbf{x})$, denoted as $\bar{\nabla} f_i(\mathbf{x}, \mathbf{y}) := \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y})$. We also note that in the minimax case $\mathbf{h}_i$ is an unbiased estimator of $\bar{\nabla} f_i(\mathbf{x}, \mathbf{y})$. Thus, $b = 0$. Therefore, we can apply FEDNEST using

$$\begin{aligned}
 \mathbf{q}_{i,\nu} &= -\nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}_{i,\nu}; \xi_{i,\nu}) + \nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}; \xi_{i,\nu}) - \frac{1}{m} \sum_{i=1}^m \nabla_{\mathbf{y}} f_i(\mathbf{x}, \mathbf{y}; \xi_i), \\
 \mathbf{h}_{i,\nu} &= \nabla_{\mathbf{x}} f_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+; \xi_{i,\nu}) - \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y}^+; \xi_{i,\nu}) + \frac{1}{m} \sum_{i=1}^m \nabla_{\mathbf{x}} f_i(\mathbf{x}, \mathbf{y}^+; \xi_i).
 \end{aligned} \tag{63}$$

D.1. Supporting Lemmas

Let $\mathbf{z} = (\mathbf{x}, \mathbf{y}) \in \mathbb{R}^{d_1+d_2}$. We make the following assumptions that are counterparts of Assumptions A and B.

Assumption E. For all $i \in [m]$:

(E1) $f_i(\mathbf{z}), \nabla f_i(\mathbf{z}), \nabla^2 f_i(\mathbf{z})$ are respectively $\ell_{f,0}, \ell_{f,1}, \ell_{f,2}$ -Lipschitz continuous; and

(E2) $f_i(\mathbf{x}, \mathbf{y})$ is μ_f -strongly convex in $\mathbf{y}$ for any fixed $\mathbf{x} \in \mathbb{R}^{d_1}$.

We use $\kappa_f = \ell_{f,1}/\mu_f$ to denote the condition number of the inner objective with respect to $\mathbf{y}$.

Assumption F. For all $i \in [m]$:

(F1) $\nabla f_i(\mathbf{z}; \xi)$ is unbiased estimators of $\nabla f_i(\mathbf{z})$; and

(F2) Its variance is bounded, i.e., $\mathbb{E}_\xi[\|\nabla f_i(\mathbf{z}; \xi) - \nabla f_i(\mathbf{z})\|^2] \leq \sigma_f^2$, for some σ_f^2 .

In the following, we re-derive Lemma C.1 for the finite-sum minimax problem (61).

Lemma D.1. Under Assumptions E and F, we have $\bar{\mathbf{h}}_i(\mathbf{x}, \mathbf{y}) = \bar{\nabla} f_i(\mathbf{x}, \mathbf{y})$ for all $i \in \mathcal{S}$ and (17a)–(17g) hold with

$$\begin{aligned} L_{yx} &= \frac{\ell_{f,2} + \ell_{f,2}L_{\mathbf{y}}}{\mu_f} + \frac{\ell_{f,1}(\ell_{f,2} + \ell_{f,2}L_{\mathbf{y}})}{\mu_f^2} = \mathcal{O}(\kappa_f^3), \\ M_f &= \ell_{f,1} = \mathcal{O}(1), \quad L_f = (\ell_{f,1} + \frac{\ell_{f,1}^2}{\mu_f}) = \mathcal{O}(\kappa_f), \\ L_{\mathbf{y}} &= \frac{\ell_{f,1}}{\mu_f} = \mathcal{O}(\kappa_f), \quad \tilde{\sigma}_f^2 = \sigma_f^2, \quad \tilde{D}_f^2 = \ell_{f,0}^2 + \sigma_f^2, \end{aligned} \tag{64}$$

where $\ell_{f,0}, \ell_{f,1}, \ell_{f,2}, \mu_f$, and σ_f are given in Assumptions E and F.

D.2. Proof of Corollary 3.2

Proof. Let $\eta = 1$. From (55) and (56), we have

$$\alpha_k = \min \left\{ \bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3, \frac{\bar{\alpha}}{\sqrt{K}} \right\}, \quad \beta_k = \frac{\bar{\beta}\alpha_k}{T}, \tag{65a}$$

where

$$\begin{aligned} \bar{\beta} &= \frac{1}{\mu_g} \left(11\ell_{f,1}L_{\mathbf{y}} + L_{yx}\tilde{D}_f^2\bar{\alpha}_1 + \frac{\ell_{f,1}L_{\mathbf{y}}\bar{\alpha}_1}{2} \right), \\ \bar{\alpha}_1 &= \frac{1}{2L_f + 4\ell_{f,1}L_{\mathbf{y}} + \frac{2\ell_{f,1}L_{yx}}{L_{\mathbf{y}}}}, \quad \bar{\alpha}_2 = \frac{T}{8\ell_{g,1}\bar{\beta}}, \quad \bar{\alpha}_3 = \frac{1}{216\ell_{f,1}^2 + 5\ell_{f,1}}. \end{aligned} \tag{65b}$$

Using the above choice of stepsizes, (59) reduces to

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] \leq \frac{4\Delta_{\mathbb{W}}}{K \min\{\bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3\}} + \frac{4\Delta_{\mathbb{W}}}{\bar{\alpha}\sqrt{K}} + \frac{4(c_1 + c_2)\bar{\alpha}}{\sqrt{K}}\sigma_f^2, \tag{66}$$

where $\Delta_{\mathbb{W}} = \mathbb{W}^0 - \mathbb{E}[\mathbb{W}^K]$,

$$\begin{aligned} c_1 &= \frac{25\ell_{f,1}}{L_{\mathbf{y}}} \left(1 + \frac{11\ell_{f,1}L_{\mathbf{y}}}{2}\bar{\alpha}_1 + \left(\frac{\ell_{f,1}L_{\mathbf{y}} + 2L_{yx}\tilde{D}_f^2}{4} \right) \bar{\alpha}_1^2 \right) \bar{\beta}^2 \frac{1}{T}, \\ c_2 &= \frac{L_f + \frac{1}{2}}{2} + \ell_{f,1}L_{\mathbf{y}} + \frac{L_{yx}\ell_{f,1}}{4L_{\mathbf{y}}}. \end{aligned} \tag{67}$$

Let $\bar{\alpha} = \mathcal{O}(\kappa_f^{-1})$. Since by our assumption, $T = \mathcal{O}(\kappa_f)$, it follows from (64) and (73) that

$$\bar{\alpha}_1 = \mathcal{O}(\kappa_f^{-2}), \quad \bar{\alpha}_2 = \mathcal{O}(\kappa_f^{-1}), \quad \bar{\alpha}_3 = \mathcal{O}(1), \quad c_1 = \mathcal{O}(\kappa_f^2), \quad c_2 = \mathcal{O}(\kappa_f^2). \quad (68)$$

Substituting (68) in (66) and (67) gives

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E}[\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O}\left(\frac{\kappa_f^2}{K} + \frac{\kappa_f}{\sqrt{K}}\right). \quad (69)$$

To achieve ε -accuracy, we need $K = \mathcal{O}(\kappa_f^2 \varepsilon^{-2})$. $\square$

E. Proof for Federated Compositional Optimization

Note that in the stochastic compositional problem (3), the inner function $f_i(\mathbf{x}, \mathbf{y}; \xi) = f_i(\mathbf{y}; \xi)$ for all $i \in \mathcal{S}$, and the outer function is $g_i(\mathbf{x}, \mathbf{y}; \zeta) = \frac{1}{2} \|\mathbf{y} - \mathbf{r}_i(\mathbf{x}; \zeta)\|^2$, for all $i \in \mathcal{S}$. In this case, we have

$$\nabla_{\mathbf{y}} g_i(\mathbf{x}, \mathbf{y}) = \mathbf{y}_i - \mathbf{r}_i(\mathbf{x}; \zeta), \quad \nabla_{\mathbf{y}\mathbf{y}} g(\mathbf{x}, \mathbf{y}; \zeta) = \mathbf{I}_{d_2 \times d_2}, \quad \text{and} \quad \nabla_{\mathbf{x}\mathbf{y}} g(\mathbf{x}, \mathbf{y}; \zeta) = -\frac{1}{m} \sum_{i=1}^m \nabla \mathbf{r}_i(\mathbf{x}; \zeta)^\top. \quad (70)$$

Hence, the hypergradient of (3) has the following form

$$\begin{aligned} \nabla f_i(\mathbf{x}) &= \nabla_{\mathbf{x}} f_i(\mathbf{y}^*(\mathbf{x})) \\ &\quad - \nabla_{\mathbf{x}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}^*(\mathbf{x})) [\nabla_{\mathbf{y}\mathbf{y}}^2 g(\mathbf{x}, \mathbf{y}^*(\mathbf{x}))]^{-1} \nabla_{\mathbf{y}} f_i(\mathbf{y}^*(\mathbf{x})) \\ &= \left(\frac{1}{m} \sum_{i=1}^m \nabla \mathbf{r}_i(\mathbf{x}) \right)^\top \nabla_{\mathbf{y}} f_i(\mathbf{y}^*(\mathbf{x})). \end{aligned} \quad (71)$$

We can obtain an approximate gradient $\nabla f_i(\mathbf{x})$ by replacing $\mathbf{y}^*(\mathbf{x})$ with $\mathbf{y}$; that is $\bar{\nabla} f_i(\mathbf{x}, \mathbf{y}) = (\frac{1}{m} \sum_{i=1}^m \nabla \mathbf{r}_i(\mathbf{x}))^\top \nabla_{\mathbf{y}} f_i(\mathbf{y})$. It should be mentioned that in the compositional case $b = 0$. Thus, we can apply FEDNEST and LFEDNEST using the above gradient approximations.

E.1. Supporting Lemmas

Let $\mathbf{z} = (\mathbf{x}, \mathbf{y}) \in \mathbb{R}^{d_1+d_2}$. We make the following assumptions that are counterparts of Assumptions A and B.

Assumption G. For all $i \in [m]$, $f_i(\mathbf{z})$, $\nabla f_i(\mathbf{z})$, $\mathbf{r}_i(\mathbf{z})$, $\nabla \mathbf{r}_i(\mathbf{z})$ are respectively $\ell_{f,0}$, $\ell_{f,1}$, $\ell_{\mathbf{r},0}$, $\ell_{\mathbf{r},1}$ -Lipschitz continuous.

Assumption H. For all $i \in [m]$:

- (H1) $\nabla f_i(\mathbf{z}; \xi)$, $\mathbf{r}_i(\mathbf{x}; \zeta)$, $\nabla \mathbf{r}_i(\mathbf{x}; \zeta)$ are unbiased estimators of $\nabla f_i(\mathbf{z})$, $\mathbf{r}_i(\mathbf{x})$, and $\nabla \mathbf{r}_i(\mathbf{x})$.
- (H2) Their variances are bounded, i.e., $\mathbb{E}_\xi[\|\nabla f_i(\mathbf{z}; \xi) - \nabla f_i(\mathbf{z})\|^2] \leq \sigma_f^2$, $\mathbb{E}_\zeta[\|\mathbf{r}_i(\mathbf{z}; \zeta) - \mathbf{r}_i(\mathbf{z})\|^2] \leq \sigma_{\mathbf{r},0}^2$, and $\mathbb{E}_\zeta[\|\nabla \mathbf{r}_i(\mathbf{z}; \zeta) - \nabla \mathbf{r}_i(\mathbf{z})\|^2] \leq \sigma_{\mathbf{r},1}^2$ for some $\sigma_f^2, \sigma_{\mathbf{r},0}^2$, and $\sigma_{\mathbf{r},1}^2$.

The following lemma is the counterpart of Lemma C.1. The proof is similar to (Chen et al., 2021a, Lemma 7).

Lemma E.1. Under Assumptions G and H, we have $\bar{\mathbf{h}}_i(\mathbf{x}, \mathbf{y}) = \bar{\nabla} f_i(\mathbf{x}, \mathbf{y})$ for all $i \in \mathcal{S}$, and (17a)–(17g) hold with

$$\begin{aligned} M_f &= \ell_{\mathbf{r},0} \ell_{f,1}, \quad L_{\mathbf{y}} = \ell_{\mathbf{r},0}, \quad L_f = \ell_{\mathbf{r},0}^2 \ell_{f,1} + \ell_{f,0} \ell_{\mathbf{r},1}, \quad L_{\mathbf{y}\mathbf{x}} = \ell_{\mathbf{r},1}, \\ \tilde{\sigma}_f^2 &= \ell_{\mathbf{r},0}^2 \sigma_f^2 + (\ell_{f,0}^2 + \sigma_f^2) \sigma_{\mathbf{r},1}^2, \quad \tilde{D}_f^2 = (\ell_{f,0}^2 + \sigma_f^2) (\ell_{\mathbf{r},0}^2 + \sigma_{\mathbf{r},1}^2). \end{aligned} \quad (72)$$

E.2. Proof of Corollary 3.3

Proof. By our assumption $T = 1$. Let $\bar{\alpha} = 1$ and $\eta = 1/L_{\mathbf{y}\mathbf{x}}$. From (55) and (56), we obtain

$$\alpha_k = \min \left\{ \bar{\alpha}_1, \bar{\alpha}_2, \bar{\alpha}_3, \frac{1}{\sqrt{K}} \right\}, \quad \beta_k = \bar{\beta} \alpha_k, \quad (73a)$$

where

$$\begin{aligned}\bar{\beta} &= \frac{1}{\mu_g} \left(11\ell_{f,1}\ell_{\mathbf{r},0}^2 + \tilde{D}_f^2\bar{\alpha}_1 + \frac{\ell_{f,1}\ell_{\mathbf{r},0}^2\bar{\alpha}_1}{2} \right), \\ \bar{\alpha}_1 &= \frac{1}{2\ell_{f,0}\ell_{\mathbf{r},1} + 6\ell_{f,1}\ell_{\mathbf{r},0}^2 + 2\ell_{f,1}\ell_{\mathbf{r},1}^2}, \quad \bar{\alpha}_2 = \frac{1}{8\ell_{\mathbf{r},0}\bar{\beta}}, \quad \bar{\alpha}_3 = \frac{1}{216(\ell_{\mathbf{r},0}\ell_{f,1})^2 + 5\ell_{\mathbf{r},0}\ell_{f,1}}.\end{aligned}\tag{73b}$$

Then, using (59), we obtain

$$\frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] = \mathcal{O} \left(\frac{1}{\sqrt{K}} \right).\tag{74}$$

This completes the proof. $\square$

F. Proof for Federated Single-Level Optimization

Next, we re-derive Lemmas C.2 and C.4 for single-level nonconvex FL under Assumptions C and D.

Lemma F.1 (Counterpart of Lemma C.2). *Suppose Assumptions C and D hold. Further, assume $\tau_i \geq 1$ and $\alpha_i = \alpha/\tau_i, \forall i \in \mathcal{S}$ for some positive constant α . Then, FEDOUT guarantees:*

$$\begin{aligned}\mathbb{E} [f(\mathbf{x}^+)] - \mathbb{E} [f(\mathbf{x})] &\leq -\frac{\alpha}{2}(1 - \alpha L_f) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \nabla f_i(\mathbf{x}_{i,\nu}) \right\|^2 \right] \\ &\quad - \frac{\alpha}{2} \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + \frac{\alpha L_f^2}{2m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + \frac{\alpha^2 L_f}{2} \sigma_f^2.\end{aligned}\tag{75}$$

Proof. By applying Algorithm 2 to the single-level optimization problem (13), we have

$$\mathbf{x}_{i,0} = \mathbf{x} \quad \forall i \in \mathcal{S}, \quad \mathbf{x}^+ = \mathbf{x} - \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}).$$

where

$$\mathbf{h}_i(\mathbf{x}_{i,\nu}) := \nabla_{\mathbf{x}} f_i(\mathbf{x}_{i,\nu}; \xi_{i,\nu}) - \nabla_{\mathbf{x}} f_i(\mathbf{x}; \xi_{i,\nu}) + \frac{1}{m} \sum_{i=1}^m \nabla_{\mathbf{x}} f_i(\mathbf{x}; \xi_i).$$

This together with Assumption C implies that

$$\begin{aligned}\mathbb{E} [f(\mathbf{x}^+)] - \mathbb{E} [f(\mathbf{x})] &\leq \mathbb{E} [\langle \mathbf{x}^+ - \mathbf{x}, \nabla f(\mathbf{x}) \rangle] + \frac{L_f}{2} \mathbb{E} [\|\mathbf{x}^+ - \mathbf{x}\|^2] \\ &= -\mathbb{E} \left[\left\langle \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}), \nabla f(\mathbf{x}) \right\rangle \right] \\ &\quad + \frac{L_f}{2} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}) \right\|^2 \right].\end{aligned}\tag{76}$$

For the first term on the RHS of (76), we obtain

$$\begin{aligned}
 -\mathbb{E} \left[\left\langle \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}), \nabla f(\mathbf{x}) \right\rangle \right] &= -\mathbb{E} \left[\frac{1}{m} \sum_{i=1}^m \frac{\alpha}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\langle \mathbf{h}_i(\mathbf{x}_{i,\nu}), \nabla f(\mathbf{x}) \rangle \mid \mathcal{F}_{i,\nu-1}] \right] \\
 &= -\frac{\alpha}{2} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \nabla f_i(\mathbf{x}_{i,\nu}) \right\|^2 \right] - \frac{\alpha}{2} \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] \\
 &\quad + \frac{\alpha}{2} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \nabla f_i(\mathbf{x}_{i,\nu}) - \nabla f(\mathbf{x}) \right\|^2 \right], \tag{77} \\
 &= -\frac{\alpha}{2} \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \nabla f_i(\mathbf{x}_{i,\nu}) \right\|^2 \right] - \frac{\alpha}{2} \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] \\
 &\quad + \frac{\alpha L_f^2}{2m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2],
 \end{aligned}$$

where the first equality follows from the law of total expectation and the last inequality is obtained from Assumption C.

For the second term on the RHS of (76), Assumption D gives

$$\begin{aligned}
 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \alpha_i \sum_{\nu=0}^{\tau_i-1} \mathbf{h}_i(\mathbf{x}_{i,\nu}, \mathbf{y}^+) \right\|^2 \right] &= \alpha^2 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} (\mathbf{h}_i(\mathbf{x}_{i,\nu}) - \nabla f_i(\mathbf{x}_{i,\nu}) + \nabla f_i(\mathbf{x}_{i,\nu})) \right\|^2 \right] \\
 &\leq \alpha^2 \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \nabla f_i(\mathbf{x}_{i,\nu}) \right\|^2 \right] + \alpha^2 \sigma_f^2. \tag{78}
 \end{aligned}$$

Plugging (78) and (77) into (76) gives the desired result. $\square$

Lemma F.2 (Counterpart of Lemma C.4). *Suppose Assumptions C and D hold. Further, assume $\tau_i \geq 1$ and $\alpha_i = \alpha/\tau_i, \forall i \in \mathcal{S}$, where $\alpha \leq 1/(3L_f)$. Then, for all $\nu \in \{0, \dots, \tau_i - 1\}$, FEDOUT gives*

$$\mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] \leq 12\tau_i^2 \alpha_i^2 \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 27\tau_i \alpha_i^2 \sigma_f^2. \tag{79}$$

Proof. The result trivially holds for $\tau_i = 1$. Similar to what is done in the proof of Lemma C.4, let $\tau_i > 1$ and define

$$\begin{aligned}
 \mathbf{v}_{i,\nu} &:= \nabla f_i(\mathbf{x}_{i,\nu}) - \nabla f_i(\mathbf{x}) + \nabla f(\mathbf{x}), \\
 \mathbf{w}_{i,\nu} &:= \mathbf{h}_i(\mathbf{x}_{i,\nu}) - \nabla f_i(\mathbf{x}_{i,\nu}) + \nabla f_i(\mathbf{x}) - \mathbf{h}_i(\mathbf{x}) + \mathbf{h}(\mathbf{x}) - \nabla f(\mathbf{x}).
 \end{aligned} \tag{80}$$

where $\mathbf{h}_i(\mathbf{x}) = \nabla_{\mathbf{x}} f_i(\mathbf{x}; \xi_{i,\nu})$ and $\mathbf{h}(\mathbf{x}) = 1/m \sum_{i=1}^m \nabla_{\mathbf{x}} f_i(\mathbf{x}; \xi_i)$.

From Algorithm 2, for each $i \in \mathcal{S}$, and $\forall \nu \in \{0, \dots, \tau_i - 1\}$, we obtain

$$\begin{aligned}
 \mathbf{x}_{i,\nu+1} - \mathbf{x} &= \mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i (\mathbf{h}_i(\mathbf{x}_{i,\nu}) - \mathbf{h}_i(\mathbf{x}) + \mathbf{h}(\mathbf{x})) \\
 &= \mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i (\mathbf{v}_{i,\nu} + \mathbf{w}_{i,\nu}),
 \end{aligned}$$

which implies that

$$\begin{aligned}
 \mathbb{E} [\|\mathbf{x}_{i,\nu+1} - \mathbf{x}\|^2] &= \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i \mathbf{v}_{i,\nu}\|^2] + \alpha_i^2 \mathbb{E} [\|\mathbf{w}_{i,\nu}\|^2] \\
 &\quad - 2\mathbb{E} [\mathbb{E} [\langle \mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i \mathbf{v}_{i,\nu}, \alpha_i \mathbf{w}_{i,\nu} \rangle \mid \mathcal{F}_{i,\nu-1}]] \\
 &= \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i \mathbf{v}_{i,\nu}\|^2] + \alpha_i^2 \mathbb{E} [\|\mathbf{w}_{i,\nu}\|^2].
 \end{aligned} \tag{81}$$

Here, the last equality uses Lemma G.3 since $\mathbb{E}[\mathbf{w}_{i,\nu} \mid \mathcal{F}_{i,\nu-1}] = 0$.

From Assumption D and Lemma G.1, for $\mathbf{w}_{i,\nu}$ defined in (80), we have

$$\begin{aligned}\mathbb{E} [\|\mathbf{w}_{i,\nu}\|^2] &\leq 3\mathbb{E} [\|\mathbf{h}_i(\mathbf{x}_{i,\nu}) - \nabla f_i(\mathbf{x}_{i,\nu})\|^2 + \|\nabla f_i(\mathbf{x}) - \mathbf{h}_i(\mathbf{x})\|^2 + \|\mathbf{h}(\mathbf{x}) - \nabla f(\mathbf{x})\|^2] \\ &\leq 9\sigma_f^2.\end{aligned}\quad (82)$$

Substituting (82) into (81), we get

$$\begin{aligned}\mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x} - \alpha_i \mathbf{v}_{i,\nu}\|^2] &\leq \left(1 + \frac{1}{2\tau_i - 1}\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 2\tau_i \alpha_i^2 \mathbb{E} [\|\mathbf{v}_{i,\nu}\|^2] + 9\alpha_i^2 \sigma_f^2 \\ &\leq \left(1 + \frac{1}{2\tau_i - 1} + 4\tau_i \alpha_i^2 L_f^2\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 4\tau_i \alpha_i^2 \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 9\alpha_i^2 \sigma_f^2 \\ &\leq \left(1 + \frac{1}{\tau_i - 1}\right) \mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] + 4\tau_i \alpha_i^2 \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 9\alpha_i^2 \sigma_f^2.\end{aligned}\quad (83)$$

Here, the first inequality follows from Lemma G.2; the second inequality uses Assumption C and Lemma G.1; and the last inequality follows by noting $\alpha_i = \alpha/\tau_i, \forall i \in \mathcal{S}$ and $\alpha \leq 1/(3L_f)$.

Now, iterating equation (83) and using $\mathbf{x}_{i,0} = \mathbf{x}, \forall i \in \mathcal{S}$, we obtain

$$\begin{aligned}\mathbb{E} [\|\mathbf{x}_{i,\nu} - \mathbf{x}\|^2] &\leq (4\tau_i \alpha_i^2 \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 9\alpha_i^2 \sigma_f^2) \sum_{j=0}^{\nu-1} \left(1 + \frac{1}{\tau_i - 1}\right)^j \\ &\leq 12\tau_i^2 \alpha_i^2 \mathbb{E} [\|\nabla f(\mathbf{x})\|^2] + 27\tau_i \alpha_i^2 \sigma_f^2,\end{aligned}\quad (84)$$

where the second inequality uses (45). This completes the proof. $\square$

F.1. Proof of Theorem 3.2

Proof. Let $\bar{\alpha}_1 := 1/(3L_f(1 + 8L_f))$. Note that by our assumption $\alpha_k \leq \bar{\alpha}_1$. Hence, the stepsize α_k satisfies the condition of Lemma F.2, and we have $6L_f^2 \alpha_k^3 \leq \alpha_k^2/4$. This together with Lemmas F.1 and F.2 gives

$$\begin{aligned}\mathbb{E} [f(\mathbf{x}^{k+1})] - \mathbb{E} [f(\mathbf{x}^k)] &\leq -\frac{\alpha_k}{2} \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] + \frac{L_f^2 \alpha_k}{2m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\|\mathbf{x}_{i,\nu}^k - \mathbf{x}^k\|^2] \\ &\quad - \frac{\alpha_k}{2} (1 - \alpha_k L_f) \mathbb{E} \left[\left\| \frac{1}{m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \nabla f_i(\mathbf{x}_{i,\nu}^k) \right\|^2 \right] + \frac{\alpha_k^2 L_f}{2} \sigma_f^2 \\ &\leq -\frac{\alpha_k}{2} \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] + \frac{L_f^2 \alpha_k}{2m} \sum_{i=1}^m \frac{1}{\tau_i} \sum_{\nu=0}^{\tau_i-1} \mathbb{E} [\|\mathbf{x}_{i,\nu}^k - \mathbf{x}^k\|^2] + \frac{\alpha_k^2 L_f}{2} \sigma_f^2 \\ &\leq -\frac{\alpha_k}{2} \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] + 6L_f^2 \alpha_k^3 \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] + \left(\frac{27}{2} \alpha_k^3 L_f^2 + \frac{\alpha_k^2 L_f}{2} \right) \sigma_f^2 \\ &\leq -\frac{\alpha_k}{4} \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] + (4 + L_f) \alpha_k^2 \sigma_f^2,\end{aligned}\quad (85)$$

where the second and last inequalities follow from (14).

Summing (85) over k and using our choice of stepsize in (14), we obtain

$$\begin{aligned}\frac{1}{K} \sum_{k=0}^{K-1} \mathbb{E} [\|\nabla f(\mathbf{x}^k)\|^2] &\leq \frac{4\Delta_f}{K} \cdot \min \left\{ \frac{1}{\bar{\alpha}_1}, \frac{\sqrt{K}}{\bar{\alpha}} \right\} + 4(4 + L_f) \sigma_f^2 \cdot \frac{\bar{\alpha}}{\sqrt{K}} \\ &\leq \frac{4\Delta_f}{\bar{\alpha}_1 K} + \left(\frac{4\Delta_f}{\bar{\alpha}} + 4(4 + L_f) \bar{\alpha} \sigma_f^2 \right) \frac{1}{\sqrt{K}},\end{aligned}\quad (86)$$

where $\Delta_f = f(\mathbf{x}^0) - \mathbb{E}[f(\mathbf{x}^K)]$. $\square$

G. Other Technical Lemmas

We collect additional technical lemmas in this section.

Lemma G.1. For any set of vectors $\{\mathbf{x}_i\}_{i=1}^m$ with $\mathbf{x}_i \in \mathbb{R}^d$, we have

$$\left\| \sum_{i=1}^m \mathbf{x}_i \right\|^2 \leq m \sum_{i=1}^m \|\mathbf{x}_i\|^2. \quad (87)$$

Lemma G.2. For any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^d$, the following holds for any $c > 0$:

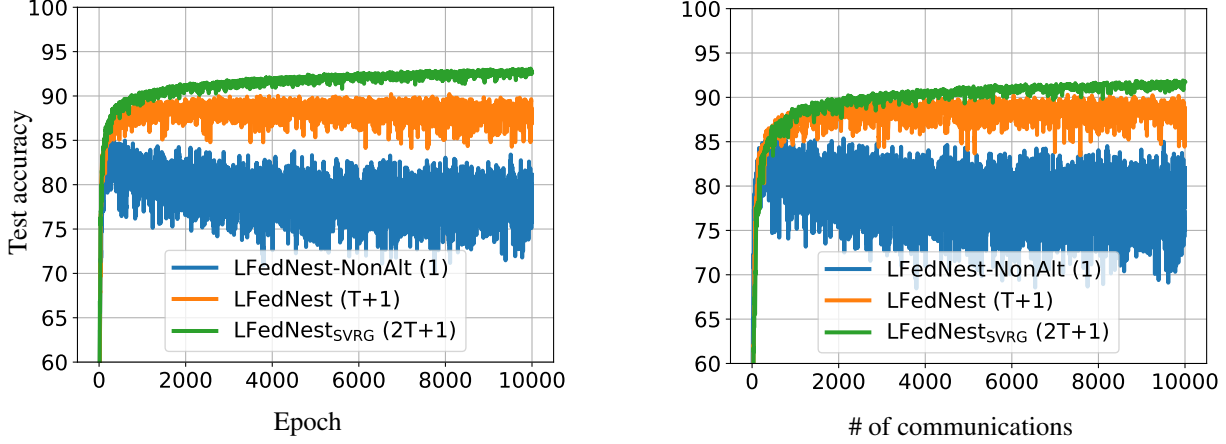
$$\|\mathbf{x} + \mathbf{y}\|^2 \leq (1 + c)\|\mathbf{x}\|^2 + \left(1 + \frac{1}{c}\right)\|\mathbf{y}\|^2. \quad (88)$$

Lemma G.3. For any set of independent, mean zero random variables $\{\mathbf{x}_i\}_{i=1}^m$ with $\mathbf{x}_i \in \mathbb{R}^d$, we have

$$\mathbb{E} \left[\left\| \sum_{i=1}^m \mathbf{x}_i \right\|^2 \right] = \sum_{i=1}^m \mathbb{E} [\|\mathbf{x}_i\|^2]. \quad (89)$$

H. Additional Experimental Results

In this section, we first provide the detailed parameters in Section 4 and then discuss more experiments. In Section 4, our federated algorithm implementation is based on (Ji, 2018), both hyper-representation and loss function tuning use batch size 64 and Neumann series parameter $N = 5$. We conduct 5 SGD/SVRG epoch of local updates in FEDINN and $\tau = 1$ in FEDOUT. In FEDNEST, we use $T = 1$, have 100 clients in total, and 10 clients are selected in each FEDNEST epoch.



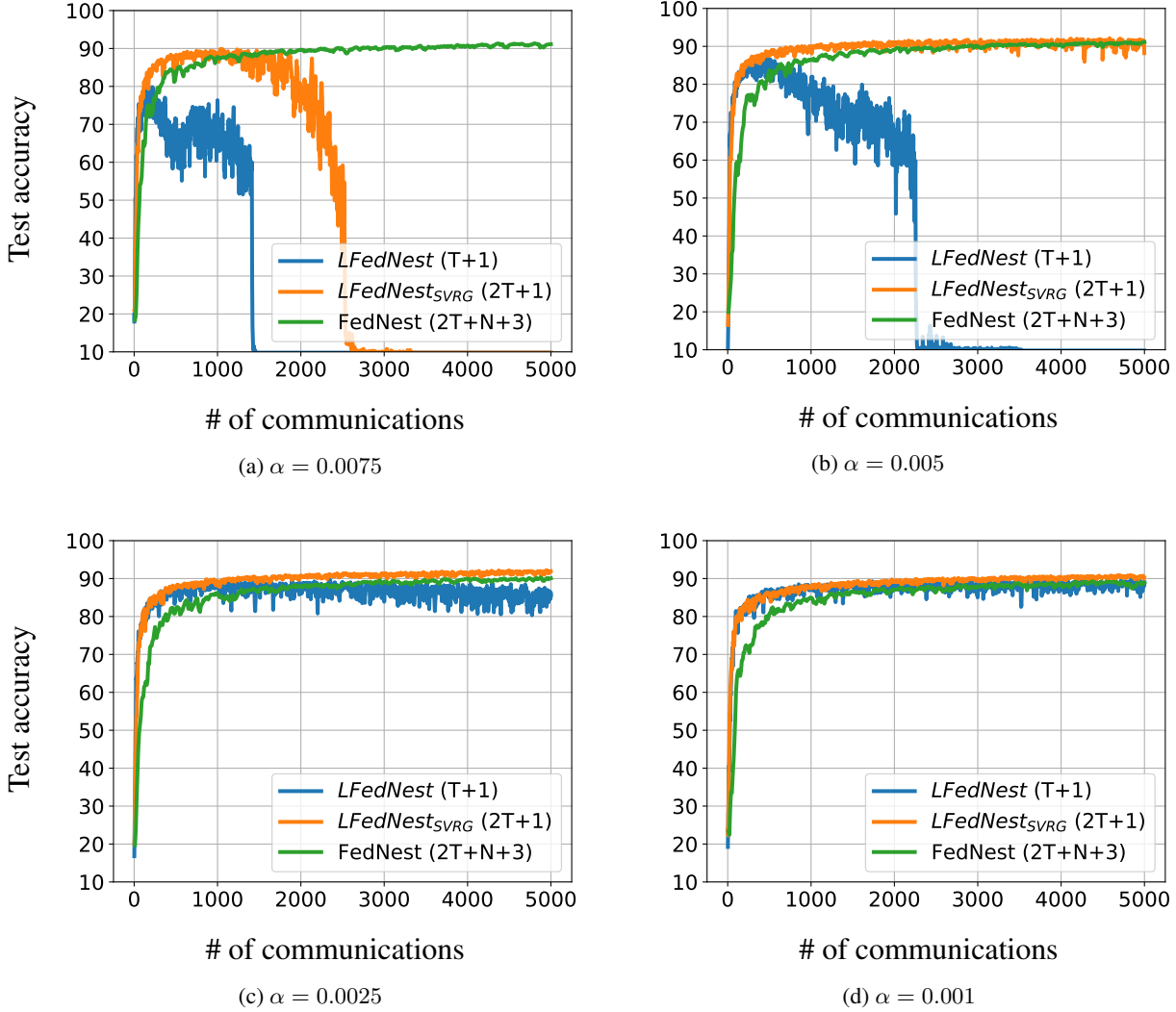
(a) The test accuracy w.r.t to the algorithm epochs.

(b) The test accuracy w.r.t to the number of communications.

Figure 5: Hyper representation experiment comparing LFEDNEST, LFEDNEST_{SVRG} and LFEDNEST-NONALT on **non-i.i.d** dataset. The number in parentheses corresponds to communication rounds shown in Table 2.

H.1. The effect of the alternating between inner and outer global variables

In our addition experiments, we investigate the effect of the alternating between inner and outer global variables $\mathbf{x}$ and $\mathbf{y}$. We use LFEDNEST-NONALT to denote the training where each client updates their local $\mathbf{y}_i$ and then update local $\mathbf{x}_i$ w.r.t. local $\mathbf{y}_i$ for all $i \in \mathcal{S}$. Hence, the nested optimization is performed locally (within the clients) and the joint variable $[\mathbf{x}_i, \mathbf{y}_i]$ is communicated with the server. One can notice that only one communication is conducted when server update global $\mathbf{x}$ and $\mathbf{y}$ by aggregating all $\mathbf{x}_i$ and $\mathbf{y}_i$.


 Figure 6: Learning rate analysis on **non-i.i.d.** data with respect to **# of communications**.

As illustrated in Figure 5, the test accuracy of LFEDNEST-NONALT remains around 80%, but both standard LFEDNEST and LFEDNEST_{SVRG} achieves better performance. Here, the number of inner iterations is set to $T = 1$. The performance boost reveals the necessity of both averaging and SVRG in FEDINN, where the extra communication makes clients more consistent.

H.2. The effect of the learning rate and the global inverse Hessian

Figure 6 shows that on non-i.i.d. dataset, both SVRG and FEDOUT have the effect of stabilizing the training. Here, we set $T = 1$ and $N = 5$. As we observe in (a)-(d), where the learning rate decreases, the algorithms with more communications are easier to achieve convergence. We note that LFEDNEST successfully converges in (d) with a very small learning rate. In contrast, in (a), FEDNEST (using the global inverse Hessian) achieves better test accuracy in the same communication round with a larger learning rate.