

Cross-Silo Federated Learning for Multi-Tier Networks with Vertical and Horizontal Data Partitioning

ANIRBAN DAS, Rensselaer Polytechnic Institute, USA

TIMOTHY CASTIGLIA, Rensselaer Polytechnic Institute, USA

SHIQIANG WANG, IBM Research, USA

STACY PATTERSON, Rensselaer Polytechnic Institute, USA

We consider federated learning in tiered communication networks. Our network model consists of a set of silos, each holding a vertical partition of the data. Each silo contains a hub and a set of clients, with the silo's vertical data shard partitioned horizontally across its clients. We propose Tiered Decentralized Coordinate Descent (TDCD), a communication-efficient decentralized training algorithm for such two-tiered networks. The clients in each silo perform multiple local gradient steps before sharing updates with their hub to reduce communication overhead. Each hub adjusts its coordinates by averaging its workers' updates, and then hubs exchange intermediate updates with one another. We present a theoretical analysis of our algorithm and show the dependence of the convergence rate on the number of vertical partitions and the number of local updates. We further validate our approach empirically via simulation-based experiments using a variety of datasets and objectives.

CCS Concepts: • **Computing methodologies** → **Machine learning algorithms**; **Distributed algorithms**; **Neural networks**; • **Mathematics of computing** → **Nonconvex optimization**.

Additional Key Words and Phrases: coordinate descent, federated learning, machine learning, stochastic gradient descent

1 INTRODUCTION

In recent times, we have seen an exponential increase in the amount of data produced at the edge of communication networks. In many settings, it is infeasible to transfer an entire dataset to a centralized cloud for downstream analysis, either due to practical constraints such as high communication cost or latency, or to maintain user privacy and security [11]. This has led to the deployment of distributed machine learning and deep-learning techniques where computation is performed collaboratively by a set of clients, each close to its own data source. Federated learning has emerged as a popular technique in this space, which performs iterative collaborative training of a global machine learning model over data distributed over a large number of clients without sending raw data over the network.

The more commonly studied approach of federated learning is *horizontal federated learning*. In horizontal federated learning, the clients' datasets share the same set of features, but each client holds only a subset of the sample space, i.e., the data is horizontally partitioned among clients [11, 15, 21]. In this setting, the clients train a copy of a model on their local datasets for a few iterations and then communicate their updates in the form of model weights or gradients directly to a centralized parameter server. The parameter server then creates the centralized model by aggregating the individual client updates, and the process is repeated until the desired convergence criterion is met.

Another scenario that arises in federated learning is when clients have different sets of features, but there is a sizable overlap in the sample ID space among their datasets [33]. For example, the training dataset may be distributed across silos in a multi-organizational context, such as healthcare,

Authors' addresses: Anirban Das, Rensselaer Polytechnic Institute, Troy, NY, USA, dasa2@rpi.edu; Timothy Castiglia, Rensselaer Polytechnic Institute, Troy, NY, USA, castit@rpi.edu; Shiqiang Wang, IBM Research, Yorktown Heights, NY, USA, wangshiq@us.ibm.com; Stacy Patterson, Rensselaer Polytechnic Institute, Troy, NY, USA, sep@cs.rpi.edu.

banking, finance, retail, etc. [24, 33]. Each silo holds a distinct set of features (e.g., customer/patient features); the data within each silo may even be of a different modality; for example, one silo may have audio features, whereas another silo has image data. At the same time, there exists an overlap in the sample ID space. More specifically, the silos have a large number of customers in common. Further, the silos may not want to share data with one another directly due to different policy restrictions, regulations, privacy/security concerns, or even bandwidth limitations in the communication network. The paradigm of training a global model over such feature-partitioned data is called *vertical federated learning* [19, 34].

A common approach for model training over vertically partitioned data builds on parallel coordinate gradient descent-type algorithms [12, 20, 22]. Here, each silo executes independent local training iterations to optimize a global model along its subset of the coordinates. The silos exchange intermediate information to ensure that the parallel updates lead to convergence of the global model. In early vertical federated learning works [4, 6, 34], the silos exchanged intermediate information after every local training iteration. However, such frequent information exchanges can lead to high training latency, especially when the communication latency is high and the intermediate information is large [19]. The authors in a more recent work [19] proposed an algorithm that addresses the communication bottleneck by performing multiple local training iterations in each silo before the silos exchange intermediate information for the next round of training. While this approach is similar to the multiple local iterations in horizontal federated learning, an important distinction is that in vertical federated learning, each silo updates only on its own subset of coordinates in contrast to updating all the coordinates in the case of horizontal federated learning.

Previous works on vertical federated learning have assumed that a silo’s entire dataset is contained in a single client [4, 6, 34]. However, in silos with a large organizational structure, it is natural that the data will be distributed horizontally across multiple clients. Previous works thus fail to capture the case where the dataset within each silo is further horizontally partitioned across multiple clients. As a motivating example, we consider two smartphone application providers who wish to train a global model over the datasets stored on the smartphones of their respective customer bases. Here, the two application companies do not want to share the customer data directly with each other. At the same time, the users do not wish to share raw personal data with the companies. In this case, the data is vertically partitioned across different applications of the different companies. Further, the data of a single company is horizontally partitioned across different application users. The silos can be thought of as the top *tier* of our system architecture. The clients inside the silos can be thought of as the bottom *tier*.

We can think of another use case when there is a significant overlap in the customer ID space between a banking silo $\mathcal{B}$, e.g., a banking holding company with multiple independent subsidiary banks, and an insurance silo $\mathcal{I}$, e.g., an insurance holding company with multiple independent subsidiaries. The two silos want to build a global model collaboratively to predict credit scores, for example. Here, the entire data is vertically partitioned between the bank and the insurance holding companies in the top tier. Each subsidiary of $\mathcal{B}$ acts as a client and may have a separate customer base, operating in a different geographical location in a country. Thus the entire customer data of the bank silo is horizontally partitioned across its subsidiaries in the bottom tier. Silo $\mathcal{I}$ for the insurance holding company also has a similar structure. Silos $\mathcal{B}$ and $\mathcal{I}$ do not want to share information directly with each other for reasons such as privacy and security concerns and regulations.

The settings above have aspects of both the horizontal federated learning and vertical federated learning paradigms. However, it is not possible to directly apply any one paradigm in this setting. We cannot directly apply horizontal federated learning since horizontal federated learning requires each client to have the same set of features for training the global model. In contrast, clients from

different silos hold disjoint sets of features in our setting. These disjoint feature sets may not be sufficient to train an accurate model independently. For example, a pair of labels can be only distinguishable by features in silo A, while another pair of labels can be only distinguishable by features in silo B. Further, we cannot use existing vertical federated algorithms directly since the data in each silo is also horizontally distributed across clients. Therefore, such a system setting calls for a new algorithm that combines aspects of both horizontal and vertical federated learning.

We, therefore, propose Tiered Decentralized Coordinate Descent (TDCD), a federated learning algorithm for training over data that is vertically partitioned across silos and further horizontally partitioned within silos. Each silo internally consists of multiple clients (shown as the grey circles in Fig. 1a) connected to a hub (shown as the orange circles in Fig. 1a), which is a server in the silo. The hub can be a cloud server maintained by a mobile application company. The hub can also be a server maintained by the company in charge of orchestrating the IT infrastructure of a particular silo. A hub itself does not contain data but facilitates training by coordinating clients' information. The goal is to jointly train a model on the features of the data contained across silos, without explicitly sharing raw data between the clients and the hubs and between clients across different silos.

TDCD works by performing a non-trivial combination of parallel coordinate descent on the top tier between silos and distributed stochastic gradient descent, with multiple local steps per communication round, in the bottom tier of clients inside each silo. More specifically, each hub maintains a block of the trainable parameters. Clients within a silo train these parameters on their local data sets for multiple iterations. The hub then aggregates the client models to update its parameter block. To execute the local iterations, clients need embeddings from the data from clients in other silos. The hubs orchestrate this information exchange every time they update their parameter blocks. By waiting for several training iterations to exchange this information, the communication cost of the algorithm is reduced. Our approach is thus a communication efficient combination of learning with both vertically and horizontally partitioned data in a multi-tiered network. TDCD is novel since it interleaves both the horizontal and vertical federated learning paradigms and thus has to account for both the perturbed gradients from the horizontal federated learning component and the stale information from the vertical federated learning component. This combination leads to a different convergence analysis.

Specifically, our contributions are the following:

- (1) We present a system model for decentralized learning in a two-tier network, where data is both vertically and horizontally partitioned.
- (2) We develop TDCD, a communication-efficient decentralized learning algorithm using principles from coordinated descent and stochastic gradient descent.
- (3) We analyze the convergence of TDCD for non-convex objectives and show how it depends on the number of silos, the number of clients, and the number of local training steps. With appropriate assumptions, TDCD achieves a convergence rate of $O\left(\frac{1}{\sqrt{R}}\right)$, where R is the total number of communication rounds between hubs and clients.
- (4) We validate our analysis via experiments using different objectives and different datasets. We further explore the effect of communication latency on the convergence rate.

We observe from our experiments that when communication latency is high with respect to the computation latency at the clients, a larger number of local iterations leads to faster convergence. Further, our experimental results indicate that TDCD shows little performance degradation as the number of vertical and horizontal partitions increases.

1.1 Related Work

Machine learning model training with vertically-partitioned features has been studied in the context of coordinate descent algorithms. A coordinate descent algorithm with gradient steps was first proposed in [25], where a fully centralized system is considered. In the decentralized context, parallel coordinate descent methods have been proposed in [12, 20, 22]. These algorithms typically require a shared memory architecture or raw data sharing and shuffling of the vertical partitions between the workers. Since TDCD operates in a federated learning setting, such access to raw data and shared memory is not possible, and therefore, these algorithms are not applicable.

In recent years, in distributed learning, where data is horizontally partitioned across multiple clients periodic averaging-based methods like federated averaging and its variations have been extensively studied. Federated averaging methods have been studied in the context of convex objectives [14, 23, 28] and non-convex objectives [7, 17, 26, 35]. The common approach is to take multiple (stochastic) gradient steps in parallel locally to achieve better communication efficiency as well preserve privacy by not sharing the raw data. Other works have analyzed versions of horizontal federated learning for hierarchical or multi-tier networks where data is partitioned horizontally across all clients participating in the process [1, 3, 18, 27]. In horizontal hierarchical federated learning, all clients update all the weights of the global model. In contrast, in TDCD, the clients only contribute towards updating their own silo's subset of the global model weights.

Several works have focused on distributed training with vertical partitions in a federated setting. The authors in works [4, 6, 8, 13, 31, 34] propose vertical federated learning algorithms for single-tier communication networks, but they do not use local iterations in the parties during training. In all of these works, each party communicates in each iteration of training, which is communication-wise expensive. A more recent work on vertical federated learning [19] solved this problem by allowing clients to take multiple local steps of training before communicating to the hub. Another recent work [29] considers the problem of training a latent Dirichlet allocation model in a cross-silo federated learning setting. However, all these works only addressed the case where all data in a silo is contained within a single client. In this work, we consider vertical and horizontal partitions of the dataset simultaneously in the two tiers. TDCD performs model training in such a multi-tiered system architecture by fusing horizontal and vertical learning approaches in a novel manner.

Recent work by Wang et al. [30] also considers the cross-silo setting. In their method, the silos first exchange information to form augmented local datasets. They then perform horizontal federated learning to train a global model. This is in contrast to our system, where silos only store a subset of all features and must combine horizontal and vertical federated learning to effectively train a model.

1.2 Paper Outline

The rest of the paper is organized as follows. We describe our system model, data partitioning, and loss function in Section 2. In Section 3, we present TDCD, followed by its convergence analysis in Section 4. The proofs of the theorem and supporting lemmas are deferred to Appendix A. We provide an experimental evaluation of TDCD in Section 5 and finally conclude in Section 6.

2 SYSTEM MODEL AND PROBLEM FORMULATION

This section describes the system architecture, the allocation of the training data, and the loss function we seek to minimize.

2.1 System Architecture and Training Data

We consider a decentralized system consisting of N silos, shown Fig. 1a. Each silo consists of a hub and multiple clients connected to it in a hub-and-spoke fashion. Each silo j has K_j clients. The number of clients per silo may be unequal. Our network model thus has two tiers as shown in Fig. 1a: the top tier of hubs, shown in orange and the bottom tier of clients in each silo, shown in gray. The hub of each silo can communicate with the hubs of all other silos. Thus, the hub communication network forms a complete graph. The clients within a silo only communicate with that silo's hub. If we map this system model on the bank and insurance example in Section 1, there are two silos, a banking silo corresponding to a banking holding company and an insurance silo corresponding to an insurance holding company. The clients in the banking silo are subsidiary banks of the banking holding company, and similarly, for the insurance silo, the clients are different divisions of the insurance holding company.

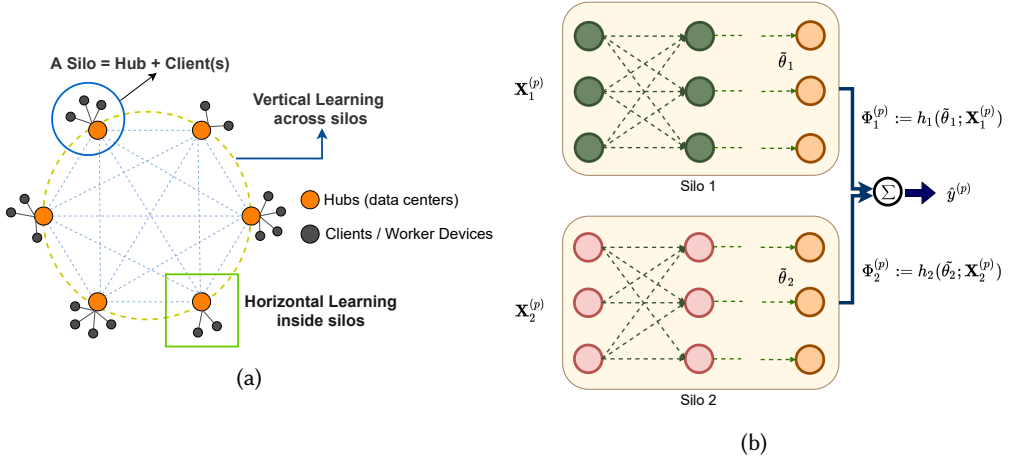


Fig. 1. (a) System architecture (b) An example of client models and the inputs and outputs to the system. The figure shows a client k at each of the two silos and the concatenation of their output embedding to find the target $\hat{y}$ corresponding to a single data sample p .

The training data consists of M sample IDs that are common across all silos. Each sample has D features. The data is partitioned vertically across the N silos so that each silo owns a disjoint subset of the features. Within each silo, its data is partitioned horizontally across its clients so that each client holds data for the vertical partition of a subset of IDs. This data partitioning is illustrated in Figure 2. More formally, we can express the entire training dataset by a matrix $\mathbf{X} \in \mathbb{R}^{M \times D}$. We denote the set of data, i.e., the columns of $\mathbf{X}$, owned by silo j by $\mathbf{X}_j$. The vertical slice of data in silo j , $\mathbf{X}_j$, is, therefore, a matrix of dimension $M \times D_j$, where D_j denotes the number of features held by silo j . Within silo j , each client holds some rows of $\mathbf{X}_j$. We denote the horizontal shard of $\mathbf{X}_j$ that is held by client k in silo j by the matrix $\mathbf{X}_{k,j}$. Lastly, we denote a sample p of the dataset (single row of $\mathbf{X}$) as $\mathbf{X}^{(p)}$, and $\mathbf{X}_j^{(p)}$ denotes the features of the p th sample corresponding to silo j . We note that for each sample p , $\mathbf{X}_j^{(p)}$ is only held in a single client in silo j . A sample ID in the banking and insurance example corresponds to an individual who is a common customer to both the banking and the insurance silos. The features of the samples in the banking silo consist of balance, installments, debit history, credit history, etc., and the banking features for each customer are stored by the bank subsidiary that serves this customer. Similarly, the insurance silo has insurance-related features

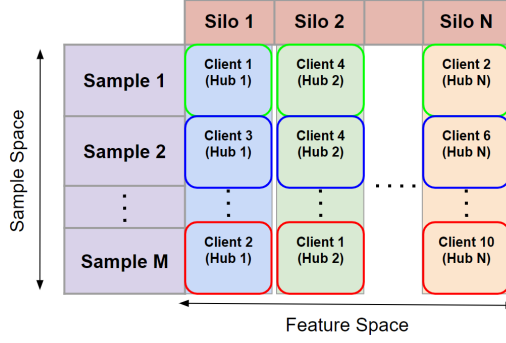


Fig. 2. Illustration of data partitioning among hubs and clients in silos. Each silo j owns a vertical partition of the full dataset X . The features of each sample i are distributed among clients of different silos. Each client in each hub owns a subset of features of some sample IDs.

such as policy details, premium, age, dependent information, past claims, etc., and each customer’s insurance features are stored by the insurance subsidiary that serves that customer. We assume that the datasets do not change for the duration of the training.

We denote the label for a sample p by $y^{(p)}$. We assume that each client stores a copy of the sample labels for each of its sample IDs. We denote the copy of sample labels for $X_{k,j}$ by $y_{k,j}$. In many low-risk scenarios, such as credit history, labels can be shared across silos. Labels can also be shared when the sample IDs themselves are anonymized. In scenarios where the labels are sensitive, for example, the case when the labels are only present in a single silo, it is still possible to extend TDCD to work without label sharing, albeit with extra communication steps, in an approach similar to [19]. We elaborate on this further in Section 3, Remark 2.

2.2 Model Formulation and Loss Function

The objective is to train a global model $\tilde{\theta}$ that can be decomposed into N coordinate block partitions, each corresponding to the model parameters assigned to a silo:

$$\tilde{\theta} = [\tilde{\theta}_1^\top, \dots, \tilde{\theta}_N^\top]^\top \quad (1)$$

where $\tilde{\theta}_j$ is the block of coordinate partition for silo j . $\tilde{\theta}$ could be weights of a linear model such as linear or logistic regression, for example, or could be trainable weights of a neural network. Let us represent the function computed by a silo j as h_j , where h_j is a map from the higher dimensional input space X_j to the lower-dimensional embedding space Φ_j parameterized by $\tilde{\theta}_j$. We denote the embedding that is output by silo j for a sample p by $\Phi_j^{(p)} \stackrel{\text{def}}{=} h_j(\tilde{\theta}_j; X_j^{(p)})$.

In Fig. 1b we show an example of a global model expressed by Equation 1. In this example, the global model is a neural network (e.g., a CNN or an LSTM) partitioned across Silo 1 and Silo 2. Each partition of the global model is also a neural network and each client in a silo j holds a copy of the neural network architecture represented by h_j . Each silo j is responsible for updating its partition of $\tilde{\theta}$, i.e., $\tilde{\theta}_j$ in the training algorithm. A client can independently compute the embedding $\Phi_j^{(p)}$ if p is a sample ID contained in that client.

Table 1. Description of all the variables used in Algorithm 1.

Notation	Definition
N	Number of silos / number of vertical partitions.
K_j	Number of clients in silo j .
$\mathbf{X}_j^{(p)}, y^{(p)}$	Features of p th sample held by in silo j and corresponding target label.
$\mathcal{L}$	Objective function.
l	Loss function.
h_j	Function computed by the silo j . Can be represented by a neural network.
$\tilde{\boldsymbol{\theta}}$	Global model parameters.
$\tilde{\boldsymbol{\theta}}_j$	Block of global model parameters corresponding to silo j , maintained by hub j .
$\boldsymbol{\theta}_{k,j}^t$	Local version of the weights $\tilde{\boldsymbol{\theta}}_j^t$ each client in hub j updates in iteration t . $\tilde{\boldsymbol{\theta}}_j^t = \frac{1}{K_j} \sum_{k=1}^{K_j} [\boldsymbol{\theta}_{k,j}^t]$.
$\Phi_j^{(p)}$	Intermediate information for the j th coordinate block for a single sample p .
ζ	Set of sample IDs forming a mini-batch. ζ^{t_0} denotes the set at iteration t_0 .
$\zeta_{k,j}^{t_0}$	Subset of sample IDs from a mini-batch ζ^{t_0} held at k th client of silo j .
η	Learning rate.
$\Phi_j^{\zeta^{t_0}}$	Intermediate information of mini-batch ζ^{t_0} formed at hub j .
$\Phi_{k,j}^{\zeta^{t_0}}$	Subset of information from $\Phi_j^{\zeta^{t_0}}$ that is relevant to samples from ζ^{t_0} held by client k of silo j .
$\Phi_{-j}^{\zeta^{t_0}}$	Intermediate information that is received by silo j from all other silos for the mini-batch ζ^{t_0} .
$\Phi_{-k,j}^{\zeta^{t_0}}$	Subset of information from $\Phi_{-j}^{\zeta^{t_0}}$ that is relevant to samples from ζ^{t_0} held by client k of silo j .
$g_{k,j}$	Local stochastic partial derivative of the loss function in client k of hub j .
$\pi_{k,j}$	A projection function such that $\pi_{k,j}(\Phi_{-j}^{\zeta}) = \Phi_{-k,j}^{\zeta}$.
Q	Number of local iterations between each global communication round.

The goal of the training algorithm is to minimize an objective function with the following structure:

$$\min_{\tilde{\boldsymbol{\theta}}} \mathcal{L}(\tilde{\boldsymbol{\theta}}; \mathbf{X}; \mathbf{y}) \stackrel{\text{def}}{=} \frac{1}{M} \sum_{p=1}^M \ell(\Phi_1^{(p)}; \Phi_2^{(p)}; \dots; \Phi_N^{(p)}; y^{(p)}) \quad (2)$$

where ℓ is the loss corresponding to a data sample p with features $\mathbf{X}^{(p)}$. Similar to the work in [19], our system architecture considers the following form for ℓ :

$$\ell(\Phi_1^{(p)}; \Phi_2^{(p)}; \dots; \Phi_N^{(p)}; y^{(p)}) \stackrel{\text{def}}{=} \ell \left(\sum_{j=1}^N h_j(\tilde{\boldsymbol{\theta}}_j; \mathbf{X}_j^{(p)}); y^{(p)} \right). \quad (3)$$

The objective function $\mathcal{L}$ can be non-convex. Since ℓ is a function of $\tilde{\boldsymbol{\theta}}$, consequently, $\mathcal{L}$ is a function of $\tilde{\boldsymbol{\theta}}$. When the context is clear, we drop $\mathbf{X}, \mathbf{y}$ from $\mathcal{L}(\cdot)$ to simplify the notation.

3 PROPOSED ALGORITHM

In this section, we present our Tiered Decentralized Coordinate Descent algorithm. The pseudocode is given in Algorithm 1. The notation used in this section is summarized in Table 1.

In Algorithm 1, in iteration $t = 0$ (line 4), the hub in each silo initializes its model weights $\tilde{\boldsymbol{\theta}}_j^0$. Then, in iteration $t = 0$, and every Q th iteration thereafter, the hubs first select a set of sample IDs ζ^{t_0} , samples, randomly drawn from the global dataset $\mathbf{X}$, to form a mini-batch. This selection can

Algorithm 1 Tiered Decentralized Coordinate Descent (TDCD)

```

1: for  $t = 0, \dots, T$  do
2:   if  $t = 0$  then
3:     for  $j = 1, \dots, N$  silos in parallel do
4:       Initialize hub  $j$  model:  $\tilde{\theta}_j^0$ 
5:     end for
6:   end if
7:   if  $t \pmod{Q} = 0$  then
8:     for  $j = 1, \dots, N$  silos in parallel do
9:       Randomly select a set of sample IDs:  $\zeta^{t_0}$ .
10:      for  $k = 1, \dots, K_j$  clients in parallel do
11:        Hub  $j$  sends copy of  $\tilde{\theta}_j^t, \zeta^{t_0}$  to client
12:        Client  $k$  initializes local model:  $\theta_{k,j}^t \leftarrow \tilde{\theta}_j^t$ 
13:        Client  $k$  calculates intermediate information  $\Phi_{k,j}^{\zeta^{t_0}}$  and sends it to hub  $j$ 
14:      end for
15:    end for
16:    for  $j = 1, \dots, N$  silos in parallel do
17:      Hub  $j$  forms  $\Phi_j^{\zeta^{t_0}} \leftarrow \bigcup_{k=1}^{K_j} \{\Phi_{k,j}^{\zeta^{t_0}}\}$ 
18:      All hubs exchange  $\Phi_j^{\zeta^{t_0}}, j = 1, \dots, N$ 
19:      Hub  $j$  calculates  $\Phi_{-j}^{\zeta^{t_0}} \leftarrow \sum_{l=1, l \neq j}^N \Phi_l^{\zeta^{t_0}}$ 
20:    end for
21:    for  $j = 1, \dots, N$  silos in parallel do
22:      For each client  $k$ , hub  $j$  calculates  $\Phi_{-k,j}^{\zeta^{t_0}} = \pi_{k,j}(\Phi_{-j}^{\zeta^{t_0}})$  and sends it to client  $k$ 
23:    end for
24:  end if
25:  for  $j = 1, \dots, N$  silos in parallel do
26:    for  $k = 1, \dots, K_j$  clients in parallel do
27:       $\theta_{k,j}^{t+1} \leftarrow \theta_{k,j}^t - \eta g_{k,j}(\Phi_{-k,j}^{\zeta^{t_0}}, \Phi_{k,j}^{\zeta^{t_0}}, \zeta_{k,j}^{t_0})$  (See Eqn. 4)
28:    end for
29:  end for
30:  if  $(t+1) \pmod{Q} = 0$  then
31:    for  $j = 1, \dots, N$  silos in parallel do
32:      Hub  $j$  computes  $\tilde{\theta}_j^{t+1} \leftarrow \frac{1}{K_j} \sum_{k=1}^{K_j} \theta_{k,j}^t$ 
33:    end for
34:  end if
35: end for

```

be achieved, for example, via pseudo-random number generators at each silo that are initialized with the same seed. This step is described in line 9. In lines 10-14, each hub j first sends a copy of its weights $\tilde{\theta}_j^t$ and the sample IDs ζ^{t_0} to the clients in its silo. The clients replace their local model weights with the updated weights received from the hub, in line 12. We define $\theta_{k,j}^t$ as the local version of the weights $\tilde{\theta}_j^t$ that client k updates. Next, TDCD computes the information needed for each client to calculate its local stochastic partial derivatives. We call the embedding $\Phi_j^{(p)}$ the *intermediate information* for the j th coordinate block for a single sample p . In TDCD, to obtain the set of intermediate information for a single mini-batch ζ , client k in silo j computes the set of information $\Phi_{k,j}^{\zeta} = \{\Phi_j^{(p)}\}_{p \in \zeta}$, for example, by executing forward propagation through its local copy of the neural network h_j . Let $\zeta_{k,j}^{t_0}$ denote the subset of sample IDs from ζ^{t_0} held at client k of

silo j . In line 13, each client j computes the intermediate information corresponding to the sample IDs in $\zeta_{k,j}^{t_0}$. We denote this by $\Phi_{k,j}^{\zeta_{k,j}^{t_0}}$. The client then sends this information to its hub.

In the next step between lines 16-20, each hub gathers the intermediate information from its clients. Each hub computes the union over the sets of updates $\{\Phi_{k,j}^{\zeta_{k,j}^{t_0}}\}$ from each of its clients to obtain $\Phi_j^{\zeta_{k,j}^{t_0}}$. This information must be propagated to the clients in the other silos so that they can compute their stochastic partial derivatives for the mini-batch ζ^{t_0} . Each hub j broadcasts $\Phi_j^{\zeta_{k,j}^{t_0}}$ to all other hubs in line 18. For hub j , we denote the intermediate information obtained from other hubs for sample p by $\Phi_{-j}^{(p)} \stackrel{\text{def}}{=} \sum_{l=1, l \neq j}^N \Phi_j^{(p)}$. Let $\Phi_{-j}^{\zeta_{k,j}^{t_0}}$ denote the entire intermediate information that is received by silo j from all other silos for the set of sample IDs in ζ^{t_0} . We similarly denote $\Phi_{-k,j}^{\zeta_{k,j}^{t_0}}$ as the subset of information from $\Phi_{-j}^{\zeta_{k,j}^{t_0}}$ corresponding to $\zeta_{k,j}^{t_0}$. We define projection functions $\pi_{k,j}$ such that $\pi_{k,j}(\Phi_{-j}^{\zeta_{k,j}^{t_0}}) = \Phi_{-k,j}^{\zeta_{k,j}^{t_0}}$. Hub j sends $\Phi_{-k,j}^{\zeta_{k,j}^{t_0}}$ to each clients k in line 22. Alternatively, a hub can send the entire $\Phi_{-j}^{\zeta_{k,j}^{t_0}}$ to the client, and the client can extract the rows corresponding to its own samples.

Finally, after receiving this intermediate information, each client k of silo j can calculate its local stochastic partial derivative for its subset of ζ^{t_0} denoted by $\zeta_{k,j}^{t_0}$. Let $g_{k,j}$ denote the local stochastic partial derivative of the loss function with respect to the block of coordinates j in client k :

$$g_{k,j}(\Phi_{-k,j}^{\zeta_{k,j}^{t_0}}, \Phi_{k,j}^{\zeta_{k,j}^{t_0}}; \zeta_{k,j}^{t_0}) \stackrel{\text{def}}{=} \frac{1}{|\zeta_{k,j}^{t_0}|} \sum_{p \in \zeta_{k,j}^{t_0}} \nabla_{\theta_{k,j}} \ell(\Phi_1^{(p)}; \Phi_2^{(p)}; \dots; \Phi_N^{(p)}; y^{(p)}). \quad (4)$$

To train its local model, each client executes Q local stochastic gradient steps on the feature block corresponding to its silo (lines 25-29):

$$\theta_{k,j}^{t+1} = \theta_{k,j}^t - \eta g_{k,j}(\Phi_{-k,j}^{\zeta_{k,j}^{t_0}}, \Phi_{k,j}^{\zeta_{k,j}^{t_0}}; \zeta_{k,j}^{t_0}). \quad (5)$$

Here, η is the learning rate. Note that t_0 represents the most recent iteration $t_0 < t$ in which the client received intermediate information from the other silos. A client uses the same subset of sample IDs $\zeta_{k,j}^{t_0}$ for all Q iterations. During local training, when a client executes (5), it recalculates its own local intermediate information $\Phi_{k,j}^{\zeta_{k,j}^{t_0}}$ using the most up to date value of its set of coordinates $\theta_{k,j}^t$. However, it reuses the intermediate information $\Phi_{-k,j}^{\zeta_{k,j}^{t_0}}$ that it received from the other silos throughout the Q iterations of local training. Therefore the information about the rest of the coordinates after iteration $t_0 + 1$ is effectively stale.

Finally, after Q local iterations, in line 32, the hubs average the model weights from the clients,

where hub j updates the j th block coordinates of global model weights as $\tilde{\theta}_j^{t+1} = \frac{1}{K_j} \sum_{k=1}^{K_j} [\theta_{k,j}^t]$. This step is similar to horizontal federated learning. The entire process from start to the completion of Q iterations of parallel training inside silos comprises a *communication round*, and this communication round is repeated until convergence.

In TDCD, clients only communicate their local model weights and intermediate information every Q iterations. This contrasts to distributed SGD algorithms, where the clients need to synchronize with a coordinating hub in each iteration. This allows TDCD to save bandwidth by increasing Q , especially when the size of the model or intermediate information is large. We explore how Q impacts the convergence of TDCD in the next section.

Remark 1. Throughout the training process, each hub maintains the block of model parameters for its silo. To form the full global model, each hub's model definition and parameters can be copied

to a central server for downstream inference purposes. This can be done periodically for model checkpointing or at the end of the training procedure.

Remark 2. We assume that every client has the labels for all of its samples. However, in cases when the labels are sensitive and sharing the labels for a sample ID across silos is not feasible, the label information for a sample ID may only be present in a client in one silo. In this case, we could modify our algorithm in the following way, similar to [19]: the clients in all silos send the intermediate information for a sample to the client that has the label for the sample. The client with the label information calculates the loss and the partial derivatives, which can then be propagated back to the other clients for use in the local gradient steps. This modification would significantly increase the communication cost of the algorithm. We note that the modified algorithm is mathematically equivalent to TDCD, albeit with a higher communication cost. Hence, the convergence analysis given in Section 4 can be trivially extended to this case.

4 CONVERGENCE ANALYSIS

In this section, we provide the convergence analysis of the TDCD algorithm. We first make the following set of common assumptions [2, 19, 35] about the loss function $\mathcal{L}$ and the local stochastic partial derivatives $g_{k,j}$ at each client.

ASSUMPTION 1. *The gradient of the loss function is Lipschitz continuous with constant L ; further, the local stochastic partial derivative of $\mathcal{L}$, with respect to each coordinate block j , denoted by $g_{k,j}(\cdot)$, is Lipschitz continuous with constant $L_{k,j}$, i.e., for all $\theta_1, \theta_2 \in \mathbb{R}^D$*

$$\|\nabla \mathcal{L}(\theta_1) - \nabla \mathcal{L}(\theta_2)\| \leq L \|\theta_1 - \theta_2\| \quad (6)$$

$$\|g_{k,j}(\theta_1) - g_{k,j}(\theta_2)\| \leq L_{k,j} \|\theta_1 - \theta_2\|. \quad (7)$$

ASSUMPTION 2. *The function $\mathcal{L}$ is lower bounded so that for all $\theta \in \mathbb{R}^D$, $\mathcal{L}(\theta) \geq \mathcal{L}_{\inf}$.*

ASSUMPTION 3. *Let ζ be a mini-batch drawn uniformly at random from all samples. Then, for all $\theta \in \mathbb{R}^D$,*

$$\mathbb{E}_{\zeta} \left[g_{k,j}(\Phi_{-j}^{\zeta}, \Phi_j^{\zeta}; \zeta) \right] = \nabla_{(j)} \mathcal{L}(\theta) \quad (8)$$

$$\mathbb{E}_{\zeta} \left[\|g_{k,j}(\Phi_{-j}^{\zeta}, \Phi_j^{\zeta}; \zeta) - \nabla_{(j)} \mathcal{L}(\theta)\|^2 \right] \leq \sigma_j^2 \quad (9)$$

where the expectation is over the choice of mini-batch ζ , and $\nabla_{(j)} \mathcal{L}(\cdot)$ denotes the partial derivative of $\mathcal{L}$, with respect to coordinate block j . Recall $\Phi_j^{\zeta} = h_j(\theta_j; \mathbf{X}_j^{\zeta})$ for a given θ and hence the local stochastic partial derivatives $g_{k,j}$ are functions of θ .

We also use the following definition:

$$L_{\max} \stackrel{\text{def}}{=} \max_k L_{k,j}. \quad (10)$$

Our analysis is based on the evolution of the global model $\tilde{\theta}$ following Algorithm 1. It is noted that the components of $\tilde{\theta}$, $\tilde{\theta}_j$ are realized every Q iterations, but for theoretical analysis we study the evolution of a virtual $\tilde{\theta}_j$ at each iteration with $\tilde{\theta}_j^t = \frac{1}{K_j} \sum_{k=1}^{K_j} \theta_{k,j}^t$.

We further define the following two quantities :

$$G_{(j)}^t = \frac{1}{K_j} \sum_{k=1}^{K_j} g_{k,j}(\Phi_{-k,j}^{\zeta^t}, \Phi_{k,j}^{\zeta^t}; \zeta_{k,j}^t) \quad (11)$$

$$G^t = [(G_{(1)}^t)^T, \dots, (G_{(N)}^t)^T]^T. \quad (12)$$

We can then write the evolution of the global model as follows,

$$\tilde{\theta}^{t+1} = \tilde{\theta}^t - \eta G^t. \quad (13)$$

We now provide the main theoretical result of this paper. We observe that since each mini-batch is used for Q local iterations, this reuse leads to a bias in some of the stochastic partial derivatives. Nevertheless, with some care, it is possible to prove the algorithm convergence. The proof of the theorem is deferred to Appendix A

THEOREM 1. *Under Assumptions 1, 2, 3, when the learning rate η satisfies the following condition:*

$$0 \leq \eta \leq \frac{1}{8Q \max(L, L_{\max})} \quad (14)$$

the expected averaged squared gradients of $\mathcal{L}$ over $T = QR > 0$ local iterations satisfies the following bound:

$$\frac{1}{R} \sum_{t_0=0}^{R-1} \mathbb{E} \left[\|\nabla \mathcal{L}(\tilde{\theta}^{t_0})\|^2 \right] \leq \frac{4(\mathcal{L}(\tilde{\theta}^0) - \mathbb{E}[\mathcal{L}(\tilde{\theta}^T)])}{\eta QR} + 4(\eta LQ + 4\eta^2 Q^2 L_{\max}^2 + 8\eta^3 Q^3 L L_{\max}^2) \sum_{j=1}^N \sigma_j^2. \quad (15)$$

The convergence error in Theorem 1 results from the parallel updates on the coordinate blocks, which depend on the number of vertical partitions N via the summation over σ_j , as well the staleness of the intermediate information due to multiple local iterations, i.e., the value of Q . With an increase in the number of vertical partitions N , the error term increases linearly. The error also depends on Q ; however, all the terms containing Q also contains the learning rate η with the same exponent. Hence, in practice, if Q is offset by a suitable learning rate η , then we can leverage multiple local iterations to achieve faster convergence in terms of wall clock time as we show in Section 5. Choosing a very small η will decrease the convergence error, but it will but increase the first term on the right-hand side of (15), leading to slower convergence.

Remark 3. When $N = 1$, i.e., there is only one silo, then the algorithm reduces to local SGD with horizontally partitioned data, and the convergence rate is similar to that obtained in [35] with respect to the orders of Q, K_j, L .

Remark 4. When $K_j = 1$, for $j = 1, 2, \dots, N$, i.e., there exists only one client per silo, the algorithm reduces to vertical federated learning with a single tier, and the convergence rate is similar with respect to the orders of Q, N, L to that obtained in [19].

Remark 5. If $\eta \propto \frac{1}{Q\sqrt{R}}$, then the (average) expected gradient squared norm converges with a rate of $O\left(\frac{1}{\sqrt{R}}\right)$, where R is the total number of communication rounds.

5 EXPERIMENTAL RESULTS

We verify the convergence properties of TDCD with respect to the different algorithm parameters of the system.

5.1 Datasets and Models

CIFAR-10: We train a CNN model on the CIFAR-10 dataset [16]. CIFAR-10 is a set of $3 \times 32 \times 32$ pixels color images with 50,000 images in the training set and 10,000 images in the test set. The goal is multi class classification with ten classes of images. We use $N = 2$ for all the experiments and divide each CIFAR-10 image vertically into two parts in all three channels ($3 \times 32 \times 16$). Each silo trains a local ResNet18 model with a penultimate linear layer, and finally, we use a cross-entropy

loss function that is parameterized by the outputs from all the silos. We assign the left ($3 \times 32 \times 16$) vertical partition of all images to one silo and the right vertical partition of all images to the other silo. Comparing with the system model in Fig. 1b, each h_j of a silo is represented by the ResNet18 model described above.

MIMIC-III: We train an LSTM model using the MIMIC-III (Medical Information Mart for Intensive Care) dataset, [10], which consists of anonymized information of patients admitted to critical care units in a hospital. We follow the data processing steps from [9] to obtain 14,681 training samples and 3,236 test samples. Each sample consists of 48 time steps corresponding to 48 hours, and each time step has 76 features. The goal is to predict in-hospital mortality as a binary classification task. For training, we split the data vertically along the feature axis. Each silo trains an LSTM model with a linear layer, and finally we use binary cross-entropy as class prediction output. Comparing with the system model in Fig. 1b, each h_j of a silo is represented by the LSTM model described above.

ModelNet40: We train a CNN model on the ModelNet40 dataset [32]. ModelNet40 is a set of images of CAD models of different objects. For each CAD model, there are 12 images from different camera views. The dataset consists of 9,843 CAD models in the training set and 2,468 CAD models in the test set. The goal is multi-class classification with ten classes of objects. We use $N = 12$ for all the experiments, assigning a single view of each CAD model to each silo. Each silo trains a local ResNet18 model with a penultimate linear layer, and finally, we use a cross-entropy loss function that is parameterized by the outputs from all the silos. Comparing with the system model in Fig. 1b, each h_j of a silo is represented by the ResNet18 model described above.

5.2 Experimental Setup

In our experiments, we simulate a multi-tier communication network and train TDCD on the CIFAR-10, MIMIC-III, and ModelNet40 datasets¹. In all figures, N represents the number of silos (vertical partitions of the dataset), and K_j represents the number of clients (horizontal partitions) in each silo. For simplicity in our experiments, we consider that all the silos have the same number of clients $K_j = K$ for all j . Thus, there are $N \times K$ clients, in total, participating in the TDCD algorithm. To distribute the training data among the clients in a silo, we generate a random permutation of the M sample IDs. Each client in silo j is assigned a contiguous block of $\frac{M}{K}$ sample IDs from this permutation. Across all experiments, we use a mini-batch size of 2000 for CIFAR-10 and MIMIC-III, and we use a mini-batch size of 320 for ModelNet40. In expectation, each client holds data for an equal fraction of each mini-batch.

The training loss is calculated using the global model $\tilde{\theta}$ and the full training data matrix. In CIFAR-10, we use the test accuracy metric calculated on the entire test dataset as generalization performance measure. Further, since MIMIC-III is an imbalanced class dataset, with only 16% positive sample, we use the F1 score as generalization performance measure on the test dataset. F1 score is the harmonic mean of the precision and recall performance of the global model $\tilde{\theta}$ calculated on the entire test dataset. For ModelNet40, we use top-5 test accuracy as our performance measure. When using top-5 accuracy, a prediction is considered correct if any of the 5 highest probabilities of the model's output matches the correct class label.

In each experiment, for each value of Q , we choose the learning rate using a grid search. For CIFAR-10 we search for a learning rate in the range $[0.0001, 0.00001]$, for MIMIC-III we search in the range $[0.1, 0.001]$, and for ModelNet40 we search in the range $[0.001, 0.00005]$. For each Q and K , we let TDCD train for 5,000 iterations for CIFAR-10, 10,000 iterations for MIMIC-III, and 4,000 iterations for ModelNet40, and pick the learning rate with the lowest training loss.

¹The source code is available at <https://github.com/rpi-nsf/TDCD>.

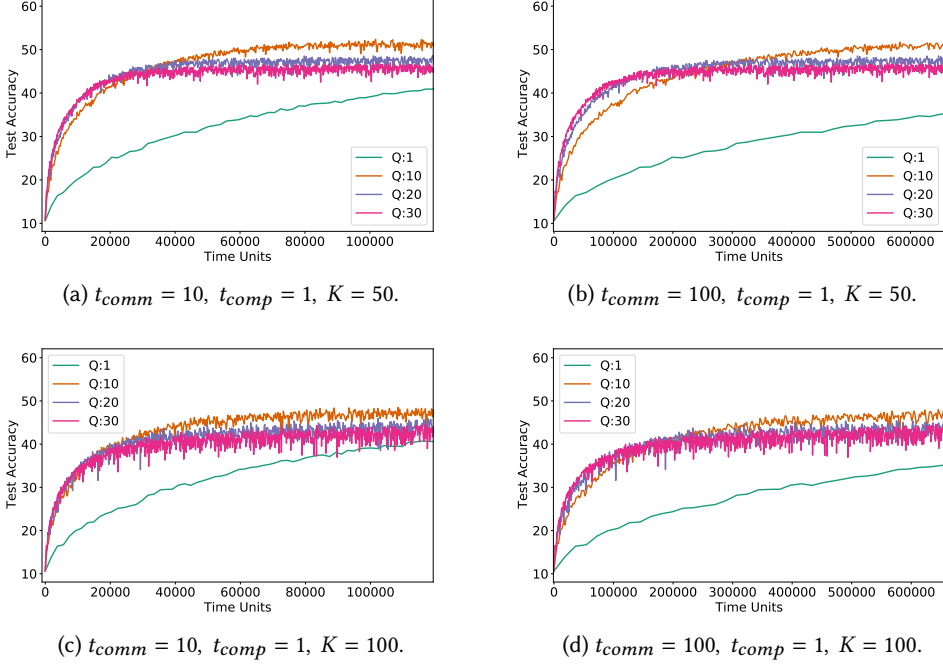


Fig. 3. Performance of TDCD on CIFAR-10. We show the test accuracy score vs. training time units t_Q . $N = 2$ in all four figures. $K = 50$ in Figures (a) and (b), and $K = 100$ in Figures (c) and (d). In Figures (a) and (c), we use a lower ratio of communication vs. computation latency with $t_{comm} = 10$. In Figures (b) and (d), we use a higher ratio of communication vs. computation latency with $t_{comm} = 100$. $t_{comp} = 1$ in all figures.

Lastly, we study the performance of TDCD in terms of both the number of training iterations and simulated latency time units. We denote the round trip communication latency between the hub and clients, and between hubs and one another, as t_{comm} . We denote the local computation latency of each step of training by $t_{comp} = 1$ time unit. Since for each communication round, i.e., Q local iterations, clients and hubs communicate twice, and hubs communicate with each other once, the total latency of one communication round can be simplified as $t_Q = 3 \times t_{comm} + Q \times t_{comp}$. In our experiments, we simulate different ratios between the computation time t_{comp} and communication time t_{comm} to explore the impact of lower and higher ratios of computation speed to communication latency.

5.3 Results

Here, we describe our experimental results.

5.3.1 Communication Efficiency. In the first set of experiments, we study the training latency of TDCD with simulated latency time units. Since in many communication networks, the communication latency between nodes dominates the computation latency inside a node, we study higher values of t_{comm} compared to t_{comp} . We study two cases, one with a low t_{comm} of 10 units and one with a high t_{comm} latency of 100 units, with $t_{comp} = 1$ unit in both cases.

In Fig. 3a and Fig. 3b, we show results for CIFAR-10 with $N = 2$ silos (vertical partitions) and $K = 50$ clients per silo. The Y-axis gives test accuracy, and the X-axis gives the training latency in

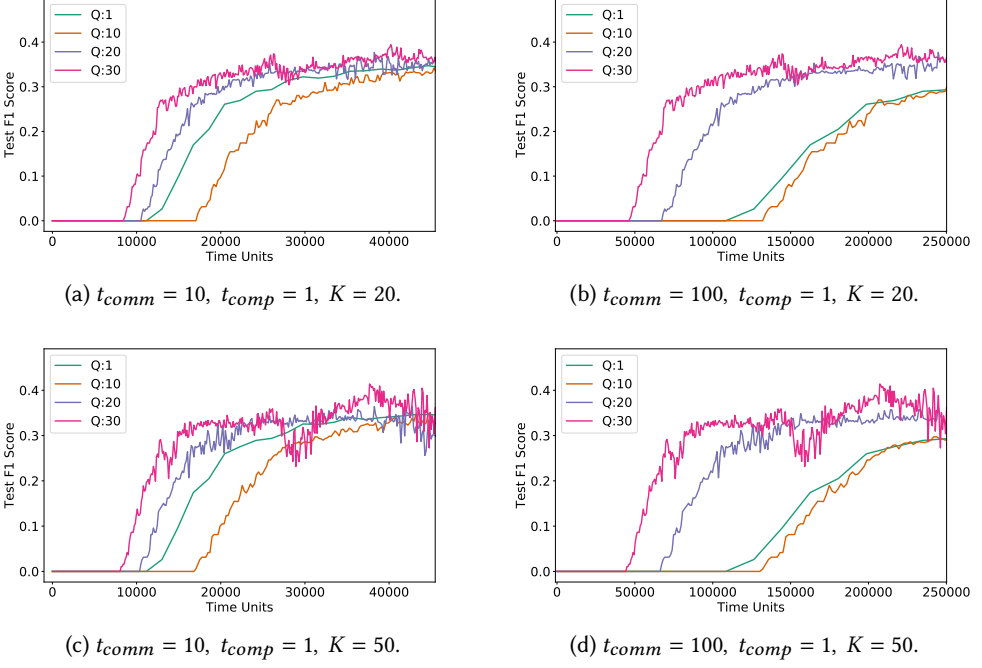


Fig. 4. Performance of TDCD on MIMIC-III. We show the test F1 score vs. training time units t_Q . $N = 4$ in all four figures. $K = 20$ in Figures (a) and (b) and $K = 50$ in Figures (c) and (d). In Figures (a) and (c), we use a lower ratio of communication vs. computation latency with $t_{comm} = 10$. In Figures (b) and (d), we use a higher ratio of communication vs. computation latency with $t_{comm} = 100$. $t_{comp} = 1$ in all figures.

time units. When $t_{comm} = 10$, we observe in Fig. 3a that the experiments with $Q > 1$ converge faster than for $Q = 1$. Similar trends can be seen for the higher value of $t_{comm} = 100$ in Fig. 3b. Moreover, comparing the results for $t_{comm} = 10$ and $t_{comm} = 100$, we observe that the gap in test accuracy between $Q > 1$ and $Q = 1$ widens as the ratio between t_{comm} and t_{comp} increases. We thus see more benefits from larger a Q when the communication latency is larger.

Results for MIMIC-III are shown in Fig. 4a and Fig. 4b for $N = 4$ silos and $K = 20$ clients per silo. The Y-axis gives the test F1 Score, and the X-axis gives the training latency in time units. The results follow similar trends to those in the previous set of experiments. Comparing the case with $t_{comm} = 10$ in Fig. 4a with the case with $t_{comm} = 100$ in Fig. 4b, we again observe that the gap in the test F1 score widens as the ratio between t_{comm} and t_{comp} increases, resulting in similar benefits from a larger Q as in CIFAR-10.

Results for ModelNet40 are shown in Fig. 5a and Fig. 5b for $N = 12$ silos and $K = 10$ clients per silo. The Y-axis gives the top-5 test accuracy, and the X-axis gives the training latency in time units. As with previous results, we see that $Q > 1$ converges faster than $Q = 1$. We again observe that increasing the ratio between t_{comm} and t_{comp} results in larger Q performing better than smaller Q .

For completeness, we also include figures for the same experiments that plot the training loss and test performance with respect to the number of local iterations. These results are shown in Fig. 6a and Fig. 6b for CIFAR-10, Fig. 7a and Fig. 7b for MIMIC-III, and Fig. 8a and Fig. 8b for ModelNet40. We observe for that CIFAR-10 with $K = 50$ clients, in Fig. 6a with increasing Q , the convergence rate in terms of iterations is worse. The convergence error is higher as well. The results of model

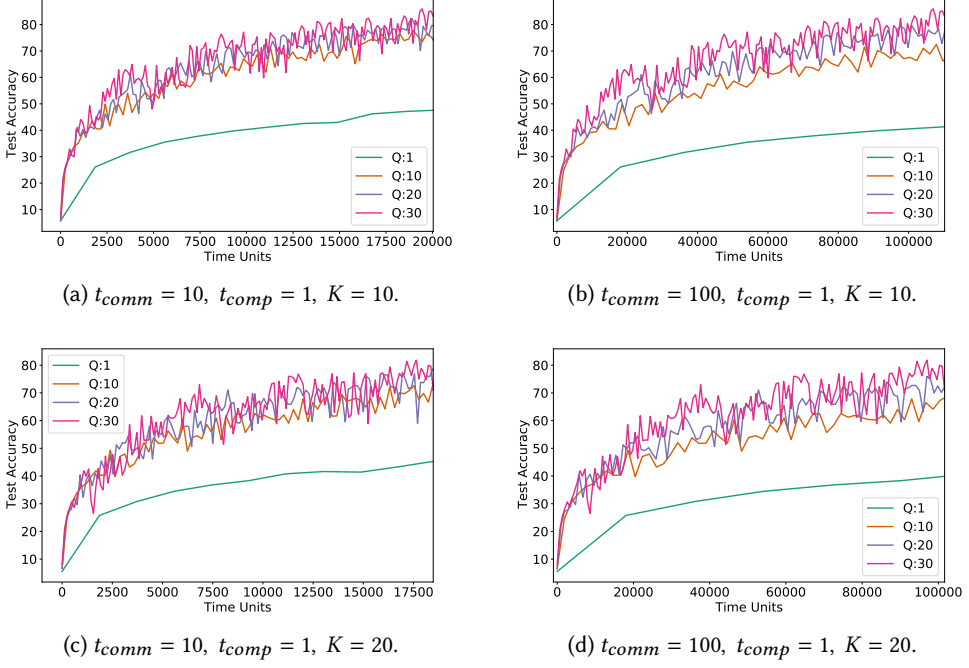


Fig. 5. Performance of TDCD on ModelNet40. We show the top-5 test accuracy score vs. training time units t_Q . $N = 12$ in all four figures. $K = 10$ in Figures (a) and (b), and $K = 20$ in Figures (c) and (d). In Figures (a) and (c), we use a lower ratio of communication vs. computation latency with $t_{comm} = 10$. In Figures (b) and (d), we use a higher ratio of communication vs. computation latency with $t_{comm} = 100$. $t_{comp} = 1$ in all figures.

test performance in Fig. 6b also reflect a performance decay as Q increases. We observe similar trends for the convergence rate and test performance for MIMIC-III with $K = 20$ clients in Fig. 7a and Fig. 7b and ModelNet40 with 10 clients in Fig. 8a and Fig. 8b.

These observations are in accordance with Theorem 1. However, as shown in Fig. 3 and Fig. 4, there is definite benefit to larger Q in terms of training latency.

5.3.2 Performance in Larger Networks. Next, we repeat the experiments in the previous section using a larger number of clients per silo. This also results in larger number of horizontal data partitions in each silo.

For CIFAR-10, we increase the number of clients in each silo from $K = 50$ to $K = 100$, resulting in total 200 participants with $N = 2$ vertical partitions. We show the test performance results vs. the training latency in Fig. 3c, Fig. 3d. As in the previous section, the training latency improves as Q increases. Comparing the results for $K = 100$ with $K = 50$, we observe that the convergence rate is similar, but the convergence error and corresponding test accuracy are slightly worse in the larger network. It is our intuition that this is because each client holds a smaller number of samples, and thus the variance of its stochastic partial derivatives is higher.

For MIMIC-III, we increase from $K = 20$ to $K = 50$ clients, resulting in 200 total clients, with $N = 4$ vertical partitions. We show the test performance in terms of the F1 score vs. the training latency in Fig. 4c and Fig. 4d for MIMIC-III. The results are similar to those in Sec. 5.3.1, though

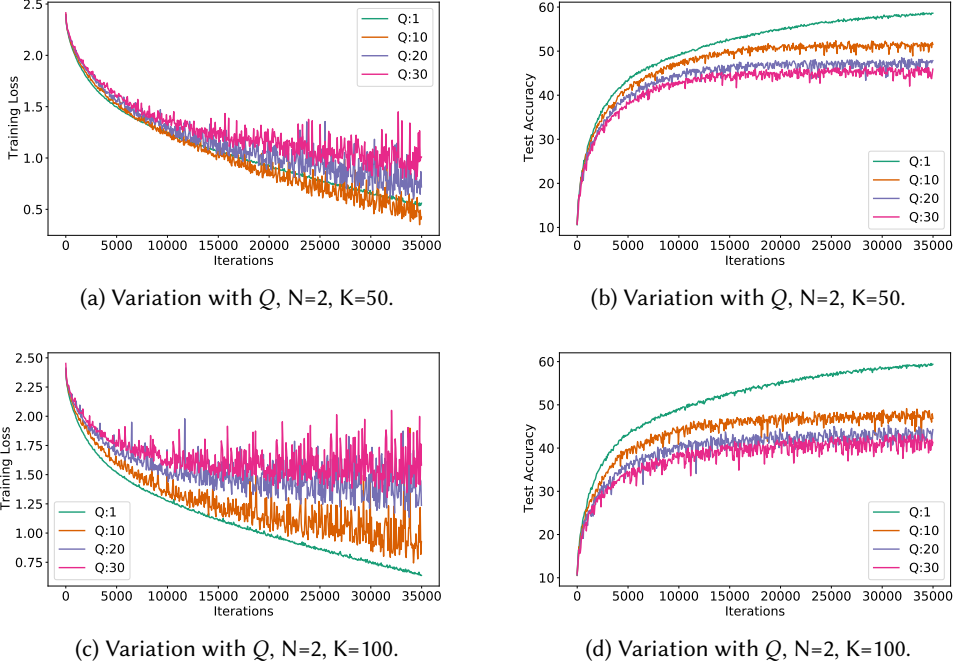


Fig. 6. Performance of TDCD on CIFAR-10 in terms of the number of training iterations for various values of Q . $N = 2$ in all figures. $K = 50$ in Figures (a) and (b), and $K = 100$ in Figures (c) and (d).

we observe a smaller performance degradation with the larger K than in the experiments with CIFAR-10.

For ModelNet40, we increase from $K = 10$ to $K = 20$ clients, resulting in 240 total clients, with $N = 12$ vertical partitions. We show the test performance in terms of the top-5 accuracy vs. the training latency in Fig. 5a and Fig. 5b for ModelNet40. As with the other datasets, test accuracy is slightly worse in the larger network.

In Fig. 6c and Fig. 6d, Fig. 7c and Fig. 7d, and Fig. 8a and Fig. 8b, we show the training loss and test performance as a function of the number of training iterations in these larger networks. In all cases, TDCD converges for all of the tested values of Q . However, we observe that for the same value of Q across the figures, larger values of K result in a larger convergence error in the training loss curves. We again believe this is due to the increased variance of the stochastic partial derivatives that the clients compute during training.

5.3.3 Effect of the Number of Vertical Partitions N . Finally, we study the impact of the number of vertical partitions N on the convergence rate. We perform this experiment using the MIMIC-III dataset. We show the results in Fig. 9a and Fig. 9b for $Q = 10$ and $K = 50$. We plot the number of iterations on the X-axis. Since we fix the value of Q , the same trend will be observed for both the number of iterations and the training latency. It is to be noted that varying N introduces other forms of dissimilarity in this experiment as the neural network architectures are different. We observe that the convergence error is worse for higher values of N . We also observe that lower values of N reach comparable F1 scores faster than higher values, in terms of local iterations. This is as expected from Theorem 1.

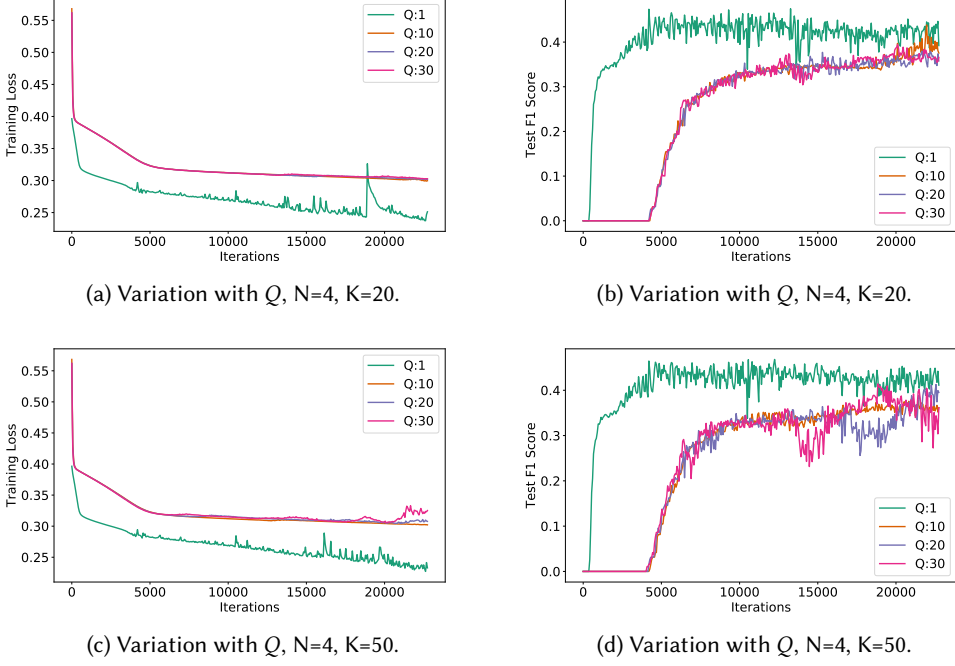


Fig. 7. Performance of TDCD on MIMIC-III in terms of the number of training iterations for various values of Q . $N = 4$ in all figures. $K = 20$ in Figures (a) and (b), and $K = 50$ in Figures (c) and (d).

6 CONCLUSION

We have introduced TDCD, a communication efficient decentralized algorithm for a multi-tier network model with both horizontally and vertically partitioned data. We have provided a theoretical analysis of the algorithm convergence and its dependence on the number of vertical partitions and the number of local iterations. Finally, we have presented experimental results that validate our theory in practice. In future work, we plan to explore the possibility of hubs communicating with each other asynchronously to share information.

A PROOF OF THE THEOREM AND SUPPORTING LEMMAS

In this section we provide the proofs of our theorem and the associated helping lemmas.

To facilitate the analysis, we first define an auxiliary local vector that represents the local view of the global model at each client. Let $y_{k,j}^t$ denote the auxiliary weight vector used to calculate the partial derivative $g_{k,j}$,

$$y_{k,j}^t = [(\theta_{-j}^{t_0})^\top, (\theta_{k,j}^t)^\top]^\top \quad (16)$$

where, $\theta_{-j}^{t_0}$ denotes the vector of all coordinates of $\tilde{\theta}$ excluding block j at iteration $t_0 \leq t$. t_0 is the most recent iteration when the client k last updated the value of $\theta_{-j}^{t_0}$ from its hub. $(\cdot)^\top$ represents transpose operation. The partial derivative $g_{k,j}$ at iteration t is a function of $y_{k,j}^t$ and the minibatch selected at t_0 . With slight abuse of notation, we let:

$$g_{k,j}(y_{k,j}^t) \stackrel{\text{def}}{=} g_{k,j}(\Phi_{-k,j}^{\zeta_{t_0}}, \Phi_{k,j}^{\zeta_{t_0}}; \zeta_{k,j}^{t_0}). \quad (17)$$

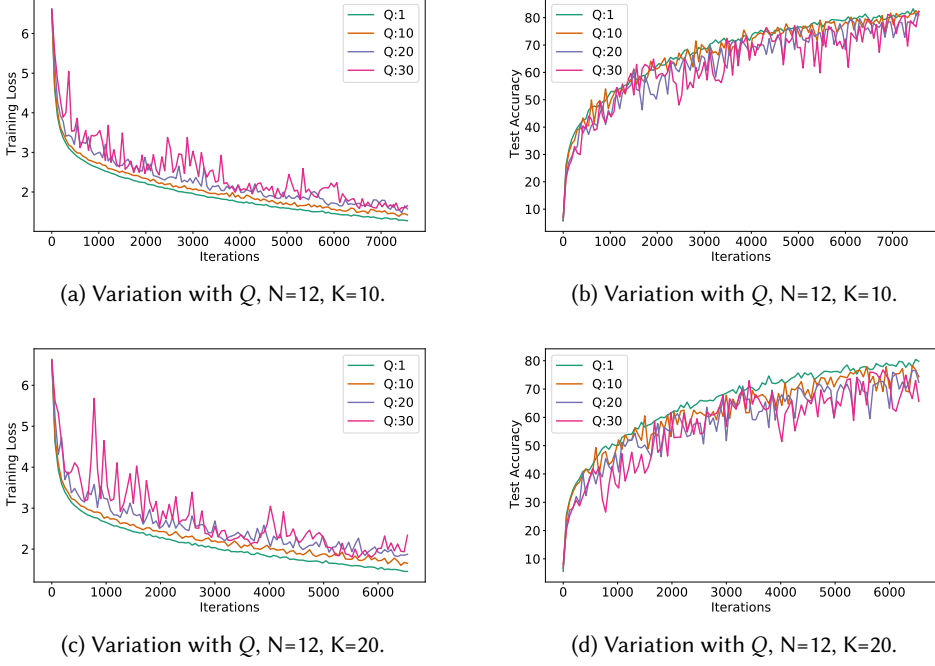


Fig. 8. Training loss and top-5 accuracy of TDCD on ModelNet40 in terms of the number of training iterations for various values of Q . $N = 12$ in all figures. $K = 10$ in Figures (a) and (b), and $K = 20$ in Figures (c) and (d).

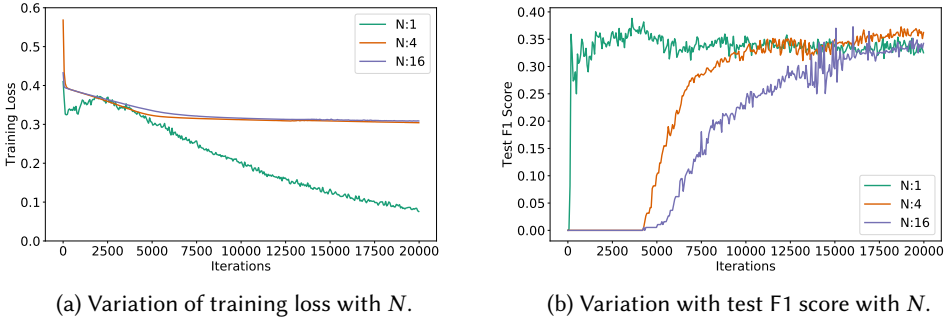


Fig. 9. Performance of TDCD on MIMIC III. Figure (a) shows training loss vs iterations t and Figure (b) shows F1 Score vs. iterations t for various values of N . We have $K = 50$, $Q = 10$.

We recall the definition of the local stochastic partial derivative $g_{k,j}^t$ from (4). $g_{k,j}^t$ at iteration t is a function of $\mathbf{y}_{k,j}^t$ and the minibatch selected at t_0 . Our analysis looks at the evolution of the $\mathbf{y}_{k,j}^t$, $t_0 \leq t \leq t_0 + Q$, i.e. the interval between the communication round t_0 and communication round $t_0 + Q$.

We note that TDCD reuses each mini-batch for Q local iterations between each communication round. This reuse implies that the local stochastic partial derivatives are not unbiased during the

local iterations $t_0 + 1 \leq t \leq t_0 + Q - 1$. Here, t_0 denotes the last iteration before t when clients synced with their hubs. Using the conditional expectation on Assumption 3 for the gradients calculated at iteration t_0 , the following is satisfied at t_0 :

$$\mathbb{E}_{\zeta^{t_0}} \left[g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \mid \{\mathbf{y}_{k,j}^r\}_{r=0}^{t_0} \right] = \nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0}) \quad (18)$$

$$\mathbb{E}_{\zeta^{t_0}} \left[\|g_{k,j}(\mathbf{y}_{k,j}^{t_0}) - \nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0})\|^2 \mid \{\mathbf{y}_{k,j}^r\}_{r=0}^{t_0} \right] \leq \sigma_j^2 \quad (19)$$

$$\mathbb{E}_{\zeta^{t_0}} [\|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \mid \{\mathbf{y}_{k,j}^r\}_{r=0}^{t_0}] \leq \sigma_j^2 + \mathbb{E}_{\zeta^{t_0}} \|\nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0})\|^2. \quad (20)$$

where, the expectation is over the mini-batch selected at iteration t_0 , conditioned on the algorithm iterates $\mathbf{y}_{k,j}^r$ up until t_0 .

For simplicity, we define $\mathbf{Y}^{t_0} \stackrel{\text{def}}{=} \{\mathbf{y}_{k,j}^r\}_{r=0}^{t_0}$ and ease the notation of the conditional expectation as the following:

$$\mathbb{E}_{\zeta^{t_0}} \left[\cdot \mid \mathbf{Y}^{t_0} \right] \stackrel{\text{def}}{=} \mathbb{E}^{t_0}. \quad (21)$$

Lastly, we denote the next local iteration when the clients will exchange intermediate information with their respective hub by

$$t_0^+ = t_0 + Q - 1.$$

A.1 Proof of Supporting Lemmas

We recall that we can write the evolution of the global model as:

$$\tilde{\boldsymbol{\theta}}^{t+1} = \tilde{\boldsymbol{\theta}}^t - \eta \mathbf{G}^t. \quad (22)$$

Here, we study the evolution of $\tilde{\boldsymbol{\theta}}$ at each iteration.

We next give several lemmas to simplify our analysis.

We first define a lemma to relate the expected local and expected global gradient.

LEMMA 1.

$$\mathbb{E} \left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 \leq \sigma_j^2 + \mathbb{E} \left\| \nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 \quad (23)$$

where $\mathbb{E}$ is the total expectation.

PROOF. We observe that at iteration t , where the latest communication between clients and hubs took place at iteration $t_0 \leq t$:

$$\sigma_j^2 \geq \mathbb{E} \left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) - \nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 \quad (24)$$

$$= \mathbb{E} \left[\mathbb{E}^{t_0} \left[\left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) - \nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 \right] \right] \quad (25)$$

$$= \mathbb{E} \left[\mathbb{E}^{t_0} \left[\left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) - \mathbb{E}^{t_0} [g_{k,j}(\mathbf{y}_{k,j}^{t_0})] \right\|^2 \right] \right] \quad (26)$$

$$= \mathbb{E} \left[\mathbb{E}^{t_0} \left[\left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 - \left\| \mathbb{E}^{t_0} [g_{k,j}(\mathbf{y}_{k,j}^{t_0})] \right\|^2 \right] \right] \quad (27)$$

$$\stackrel{(a)}{=} \mathbb{E} \left[\mathbb{E}^{t_0} \left[\left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 \right] \right] - \mathbb{E} \left[\left\| \mathbb{E}^{t_0} [g_{k,j}(\mathbf{y}_{k,j}^{t_0})] \right\|^2 \right] \quad (28)$$

$$= \mathbb{E} \left\| g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\|^2 - \mathbb{E} \left\| \nabla_{(j)} \mathcal{L}(\mathbf{y}_{k,j}^{t_0}) \right\|^2. \quad (29)$$

Here the equality in (a) follows from (8) in Assumption 3. We can obtain Lemma 1 by rearranging (29). $\square$

LEMMA 2. Under assumptions on the learning rate $\eta \leq \frac{1}{2QL_{\max}}$,

$$\sum_{t=t_0}^{t_0^*} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \leq 4Q^2(1+Q)L_{\max}^2\eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2. \quad (30)$$

PROOF. We start with simplifying the L.H.S. expression:

$$\sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (31)$$

$$= \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) + g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (32)$$

$$\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left[(1+n) \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t-1})\|^2 + \left(1 + \frac{1}{n}\right) \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \quad (33)$$

$$\stackrel{(a)}{\leq} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left[(1+n) L_{\max}^2 \|\mathbf{y}_{k,j}^t - \mathbf{y}_{k,j}^{t-1}\|^2 + \left(1 + \frac{1}{n}\right) \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \quad (34)$$

$$\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left[(1+n) L_{\max}^2 \eta^2 \|g_{k,j}(\mathbf{y}_{k,j}^{t-1})\|^2 + \left(1 + \frac{1}{n}\right) \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \quad (35)$$

$$\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left[(1+n) L_{\max}^2 \eta^2 \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0}) + g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \\ + \leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left(1 + \frac{1}{n}\right) \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (36)$$

$$\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left[2(1+n) L_{\max}^2 \eta^2 \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 + \left(1 + \frac{1}{n}\right) \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \\ + 2(1+n) L_{\max}^2 \eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (37)$$

$$= \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left(2(1+n) L_{\max}^2 \eta^2 + \left(1 + \frac{1}{n}\right) \right) \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ + 2(1+n) L_{\max}^2 \eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2. \quad (38)$$

where (a) follows from Assumption 1.

We are now interested in forming a recurrence relation.

On setting $n = Q$ and $\eta \leq \frac{1}{2QL_{\max}}$ in the first term of (38), we get,

$$\sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (39)$$

$$\begin{aligned} &\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left(2(1+Q)L_{\max}^2\eta^2 + \left(1 + \frac{1}{Q}\right) \right) \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ &\quad + 2(1+Q)L_{\max}^2\eta^2 \sum_{j=1}^N \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \end{aligned} \quad (40)$$

$$\begin{aligned} &\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left(\frac{1+Q}{2Q^2} + 1 + \frac{1}{Q} \right) \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ &\quad + 2(1+Q)L_{\max}^2\eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \end{aligned} \quad (41)$$

$$\begin{aligned} &\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left(\frac{2Q}{2Q^2} + 1 + \frac{1}{Q} \right) \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ &\quad + 2(1+Q)L_{\max}^2\eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \end{aligned} \quad (42)$$

$$\begin{aligned} &\leq \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left(1 + \frac{2}{Q} \right) \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ &\quad + 2(1+Q)L_{\max}^2\eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2. \end{aligned} \quad (43)$$

We now have a recurrence relation:

$$\begin{aligned} &\underbrace{\sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2}_{\Psi^t} \\ &\leq \underbrace{\left(1 + \frac{2}{Q}\right)}_A \underbrace{\sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t-1}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2}_{\Psi^{t-1}} \\ &\quad + \underbrace{2(1+Q)L_{\max}^2\eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2}_{\Xi}. \end{aligned} \quad (44)$$

More specifically, we have :

$$\Psi^t \leq A\Psi^{t-1} + \Xi. \quad (45)$$

We then expand the recurrence relation as follows from $t = t_0 + 1, t_0 + 2, \dots$:

$$\Psi^t \leq A^{t-t_0} \Psi^{t_0} + \Xi \sum_{p=0}^{t-t_0-1} A^p. \quad (46)$$

We note that $\Psi_0 = \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0}) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 = 0$. Therefore we have that

$$\Psi^t \leq \Xi \sum_{p=0}^{t-t_0-1} A^p \quad (47)$$

$$\leq \Xi \frac{A^{t-t_0} - 1}{A - 1}. \quad (48)$$

Now summing Ψ^t over the set of local iterations between t_0 and t_0^+ , where $t_0^+ \stackrel{\text{def}}{=} t_0 + Q - 1$, we get:

$$\sum_{t=t_0}^{t_0^+} \Psi^t \leq \Xi \sum_{t=t_0}^{t_0^+} \frac{A^{t-t_0} - 1}{A - 1} \quad (49)$$

$$\leq \frac{\Xi}{A - 1} \left(\left(\sum_{t=t_0}^{t_0^+} A^{t-t_0} \right) - Q \right) \quad (50)$$

$$\leq \frac{\Xi}{A - 1} \left(\frac{A^Q - 1}{A - 1} - Q \right) \quad (51)$$

$$\leq \frac{\Xi}{(1 + \frac{2}{Q}) - 1} \left(\frac{(1 + \frac{2}{Q})^Q - 1}{(1 + \frac{2}{Q}) - 1} - Q \right) \quad (52)$$

$$= \frac{Q\Xi}{2} \left(\frac{Q(1 + \frac{2}{Q})^Q - 1}{2} - Q \right) \quad (53)$$

$$= \frac{Q^2\Xi}{2} \left(\frac{(1 + \frac{2}{Q})^Q - 1}{2} - 1 \right) \quad (54)$$

$$\stackrel{(a)}{\leq} \frac{Q^2\Xi}{2} \left(\frac{e^2 - 1}{2} - 1 \right) \quad (55)$$

$$\leq 2Q^2\Xi \quad (56)$$

where we upper bound the terms inside the parenthesis of (a) by 4.

Plugging in the value of Ξ from (44) into (56) we get:

$$\begin{aligned} & \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ & \leq 4Q^2(1 + Q)L_{\max}^2\eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2. \end{aligned} \quad (57)$$

□

A.2 Proof of the Theorem 1

We now prove our main result.

PROOF. We have that $t_0^+ = t_0 + Q - 1$. We return to the expression in (22), on which we use the Lipschitz condition and first order taylor series expansion to obtain the following:

$$\begin{aligned} & \mathcal{L}(\tilde{\theta}^{t_0^+}) - \mathcal{L}(\tilde{\theta}^{t_0}) \\ & \leq \left\langle \nabla \mathcal{L}(\tilde{\theta}^{t_0}), \tilde{\theta}^{t_0^+} - \tilde{\theta}^{t_0} \right\rangle + \frac{L}{2} \|\tilde{\theta}^{t_0^+} - \tilde{\theta}^{t_0}\|^2 \end{aligned} \quad (58)$$

$$\stackrel{(a)}{\leq} - \left\langle \nabla \mathcal{L}(\tilde{\theta}^{t_0}), \eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t \right\rangle + \frac{L}{2} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \quad (59)$$

$$= -\eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \left\langle \nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), \mathbf{G}_{(j)}^t \right\rangle + \frac{L}{2} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \quad (60)$$

$$= -\eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left\langle \nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), g_{k,j}(\mathbf{y}_{k,j}^t) \right\rangle + \frac{L}{2} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \quad (61)$$

$$= \eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left\langle -\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\rangle \quad (62)$$

$$- \eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left\langle \nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\rangle + \frac{L}{2} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \quad (63)$$

$$\stackrel{(b)}{\leq} \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left[\|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \quad (64)$$

$$- \eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left\langle \nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\rangle + \frac{L}{2} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \quad (65)$$

where (a) follows from (22) and (b) follows since $A \cdot B = \frac{1}{2}(A^2 + B^2 - (A + B)^2) \leq \frac{1}{2}(A^2 + B^2)$.

Taking conditional expectation at t_0 on both sides of (65), we get:

$$\begin{aligned} & \mathbb{E}^{t_0} [\mathcal{L}(\tilde{\theta}^{t_0^+})] - \mathcal{L}(\tilde{\theta}^{t_0}) \\ & \leq \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \left[\mathbb{E}^{t_0} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \right] \\ & \quad - \eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left\langle \nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\rangle + \frac{L}{2} \mathbb{E}^{t_0} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \end{aligned} \quad (66)$$

$$\begin{aligned} & \leq \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\ & \quad - \eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \left\langle \nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0}), g_{k,j}(\mathbf{y}_{k,j}^{t_0}) \right\rangle + \frac{L}{2} \mathbb{E}^{t_0} \|\eta \sum_{t=t_0}^{t_0^+} \mathbf{G}^t\|^2 \end{aligned} \quad (67)$$

$$\begin{aligned}
&= \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad - \eta \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{L}{2} \mathbb{E}^{t_0} \|\eta \sum_{t=t_0}^{t_0^+} G^t\|^2
\end{aligned} \tag{68}$$

$$\begin{aligned}
&= -\frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad + \frac{L}{2} \mathbb{E}^{t_0} \|\eta \sum_{t=t_0}^{t_0^+} G^t\|^2
\end{aligned} \tag{69}$$

$$\begin{aligned}
&\stackrel{(a)}{\leq} -\frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad + \eta^2 LQ \mathbb{E}^{t_0} \sum_{t=t_0}^{t_0^+} \|G^t - G^{t_0}\|^2 + \eta^2 LQ \mathbb{E}^{t_0} \sum_{t=t_0}^{t_0^+} \|G^{t_0}\|^2
\end{aligned} \tag{70}$$

$$\begin{aligned}
&\leq -\frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad + \eta^2 LQ \mathbb{E}^{t_0} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 + \eta^2 LQ \mathbb{E}^{t_0} \sum_{t=t_0}^{t_0^+} \|G^{t_0}\|^2
\end{aligned} \tag{71}$$

$$\begin{aligned}
&= -\frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} (1 + 2\eta LQ) \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad + \eta^2 LQ \mathbb{E}^{t_0} \sum_{t=t_0}^{t_0^+} \|G^{t_0}\|^2
\end{aligned} \tag{72}$$

$$\begin{aligned}
&\leq -\frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + \frac{\eta}{2} (1 + 2\eta LQ) \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^t) - g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad + \eta^2 LQ \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2
\end{aligned} \tag{73}$$

$$\begin{aligned}
&\stackrel{(b)}{\leq} -\frac{\eta}{2} \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + 2\eta Q^2 (1 + 2\eta LQ) (1 + Q) L_{\max}^2 \eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \\
&\quad + \eta^2 LQ \sum_{t=t_0}^{t_0^+} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2
\end{aligned} \tag{74}$$

$$= -\frac{\eta Q}{2} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\theta}^{t_0})\|^2 + 2\eta Q^2 (1 + 2\eta LQ) (1 + Q) L_{\max}^2 \eta^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2$$

$$+ \eta^2 L Q^2 \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (75)$$

$$= -\frac{\eta Q}{2} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 2\eta Q^2(1 + 2\eta L Q)(1 + Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \mathbb{E}^{t_0} \|g_{k,j}(\mathbf{y}_{k,j}^{t_0})\|^2 \quad (76)$$

$$\stackrel{(c)}{\leq} -\frac{\eta Q}{2} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 2\eta Q^2(1 + 2\eta L Q)(1 + Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} (\|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2 + \sigma_j^2) \quad (77)$$

$$\leq -\left(\frac{\eta Q}{2} - \eta^2 L Q^2 - 2\eta Q^2(1 + 2\eta L Q)(1 + Q)L_{\max}^2 \eta^2\right) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 2\eta Q^2(1 + 2\eta L Q)(1 + Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \quad (78)$$

$$= -\frac{\eta Q}{2} (1 - 2\eta L Q - 4Q(1 + 2\eta L Q)(1 + Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 2\eta Q^2(1 + 2\eta L Q)(1 + Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \quad (79)$$

$$\leq -\frac{\eta Q}{2} (1 - 2\eta L Q - 8Q^2 L_{\max}^2 \eta^2 - 16Q^3 L L_{\max}^2 \eta^3) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 4\eta Q^3(1 + 2\eta L Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \quad (80)$$

where in (a) the extra Q in the last term arises since we pull the summation over t from inside the norm to outside and break the norm into two parts. We obtain (b) from Lemma 2 and (c) follows from (20).

Let $\eta \leq \frac{1}{8Q \max(L, L_{\max})}$. We can then bound (80) as follows:

$$\mathbb{E}^{t_0} [\mathcal{L}(\tilde{\boldsymbol{\theta}}^{t_0})] - \mathcal{L}(\tilde{\boldsymbol{\theta}}^{t_0}) \leq -\frac{\eta Q}{2} \left(1 - \frac{1}{4} - \frac{1}{8} - \frac{1}{32}\right) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 4\eta Q^3(1 + 2\eta L Q)L_{\max}^2 \eta^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \quad (81)$$

$$\leq -\frac{\eta Q}{4} \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \|\nabla \mathcal{L}_{(j)}(\tilde{\boldsymbol{\theta}}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 4\eta^3 Q^3 L_{\max}^2 + 8\eta^4 L Q^4 L_{\max}^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \quad (82)$$

$$\stackrel{(a)}{\leq} -\frac{\eta Q}{4} \|\nabla \mathcal{L}(\tilde{\theta}^{t_0})\|^2$$

$$+ (\eta^2 L Q^2 + 4\eta^3 Q^3 L_{\max}^2 + 8\eta^4 L Q^4 L_{\max}^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \quad (83)$$

where (a) follows from using the definition of $\mathcal{L}$ in the first term.

Rearranging the terms in (83), we get:

$$\begin{aligned} \|\nabla \mathcal{L}(\tilde{\theta}^{t_0})\|^2 &\leq 4 \frac{\mathcal{L}(\tilde{\theta}^{t_0}) - \mathbb{E}^{t_0}[\mathcal{L}(\tilde{\theta}^{t_0^+})]}{\eta Q} \\ &\quad + 4(\eta L Q + 4\eta^2 Q^2 L_{\max}^2 + 8\eta^3 Q^3 L L_{\max}^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2. \end{aligned} \quad (84)$$

Now averaging over all the communication rounds, i.e. all the intervals of Q local iterations over $t_0 = 0, 1, \dots, R-1$, such that $T = QR$, and taking total expectation:

$$\begin{aligned} &\frac{1}{R} \sum_{t_0=0}^{R-1} \mathbb{E} \left[\|\nabla \mathcal{L}(\tilde{\theta}^{t_0})\|^2 \right] \\ &\leq 4 \frac{1}{R} \mathbb{E} \sum_{t_0}^{R-1} \frac{\mathcal{L}(\tilde{\theta}^{t_0}) - \mathbb{E}^{t_0}[\mathcal{L}(\tilde{\theta}^{t_0^+})]}{\eta Q} \\ &\quad + \frac{1}{R} \mathbb{E} \sum_{t_0}^{R-1} 4(\eta L Q + 4\eta^2 Q^2 L_{\max}^2 + 8\eta^3 Q^3 L L_{\max}^2) \sum_{j=1}^N \frac{1}{K_j} \sum_{k=1}^{K_j} \sigma_j^2 \\ &\leq \frac{4(\mathcal{L}(\tilde{\theta}^0) - \mathbb{E}[\mathcal{L}(\tilde{\theta}^T)])}{\eta QR} + 4(\eta L Q + 4\eta^2 Q^2 L_{\max}^2 + 8\eta^3 Q^3 L L_{\max}^2) \sum_{j=1}^N \sigma_j^2. \end{aligned} \quad (85)$$

□

ACKNOWLEDGMENTS

This work is supported by the Rensselaer-IBM AI Research Collaboration (<http://airc.rpi.edu>), part of the IBM AI Horizons Network (<http://ibm.biz/AIHorizons>), and by the National Science Foundation under grants CNS 1553340 and CNS 1816307. A preliminary conference version of this work has been published in [5].

REFERENCES

- [1] Mehdi Salehi Heydar Abad, Emre Ozfatura, Deniz Gunduz, and Ozgur Ercetin. 2020. Hierarchical federated learning across heterogeneous cellular networks. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 8866–8870.
- [2] Léon Bottou, Frank E Curtis, and Jorge Nocedal. 2018. Optimization methods for large-scale machine learning. *Siam Review* 60, 2 (2018), 223–311.
- [3] Timothy Castiglia, Anirban Das, and Stacy Patterson. 2021. Multi-Level Local SGD: Distributed SGD for Heterogeneous Hierarchical Networks. In *International Conference on Learning Representations*.
- [4] Tianyi Chen, Xiao Jin, Yuejiao Sun, and Wotao Yin. 2020. VAFL: a Method of Vertical Asynchronous Federated Learning. [arXiv:2007.06081](https://arxiv.org/abs/2007.06081) [cs.LG]

- [5] Anirban Das and Stacy Patterson. 2021. Multi-tier federated learning for vertically partitioned data. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 3100–3104.
- [6] Siwei Feng and Han Yu. 2020. Multi-Participant Multi-Class Vertical Federated Learning. arXiv:2001.11154 [cs.LG]
- [7] Farzin Haddadpour and Mehrdad Mahdavi. 2019. On the Convergence of Local Descent Methods in Federated Learning. arXiv:1910.14425 [cs.LG]
- [8] Stephen Hardy, Wilko Henecka, Hamish Ivey-Law, Richard Nock, Giorgio Patrini, Guillaume Smith, and Brian Thorne. 2017. Private federated learning on vertically partitioned data via entity resolution and additively homomorphic encryption. arXiv:1711.10677 [cs.LG]
- [9] Hrayr Harutyunyan, Hrant Khachatrian, David C Kale, Greg Ver Steeg, and Aram Galstyan. 2019. Multitask learning and benchmarking with clinical time series data. *Scientific data* 6, 1 (2019), 1–18.
- [10] Alistair EW Johnson, Tom J Pollard, Lu Shen, H Lehman Li-Wei, Mengling Feng, Mohammad Ghassemi, Benjamin Moody, Peter Szolovits, Leo Anthony Celi, and Roger G Mark. 2016. MIMIC-III, a freely accessible critical care database. *Scientific data* 3, 1 (2016), 1–9.
- [11] Peter Kairouz, H. Brendan McMahan, Brendan Avent, Aurélien Bellet, Mehdi Bennis, Arjun Nitin Bhagoji, Kallista Bonawitz, Zachary Charles, Graham Cormode, Rachel Cummings, Rafael G. L. D'Oliveira, Hubert Eichner, Salim El Rouayheb, David Evans, Josh Gardner, Zachary Garrett, Adrià Gascón, Badi Ghazi, Phillip B. Gibbons, Marco Gruteser, Zaid Harchaoui, Chaoyang He, Lie He, Zhouyuan Huo, Ben Hutchinson, Justin Hsu, Martin Jaggi, Tara Javidi, Gauri Joshi, Mikhail Khodak, Jakub Konečný, Aleksandra Korolova, Farinaz Koushanfar, Sanmi Koyejo, Tancrède Lepoint, Yang Liu, Prateek Mittal, Mehryar Mohri, Richard Nock, Ayfer Özgür, Rasmus Pagh, Hang Qi, Daniel Ramage, Ramesh Raskar, Mariana Raykova, Dawn Song, Weikang Song, Sebastian U. Stich, Ziteng Sun, Ananda Theertha Suresh, Florian Tramèr, Praneeth Vepakomma, Jianyu Wang, Li Xiong, Zheng Xu, Qiang Yang, Felix X. Yu, Han Yu, and Sen Zhao. 2021. Advances and Open Problems in Federated Learning. *Foundations and Trends® in Machine Learning* 14, 1–2 (2021), 1–210.
- [12] Dongyeop Kang, Woosang Lim, Kijung Shin, Lee Sael, and U Kang. 2014. Data/feature distributed stochastic coordinate descent for logistic regression. In *Proceedings of the 23rd ACM International Conference on Conference on Information and Knowledge Management*. 1269–1278.
- [13] Yan Kang, Yang Liu, and Tianjian Chen. 2020. Fedmvt: Semi-supervised vertical federated learning with multiview training. (2020). arXiv:2008.10838.
- [14] Ahmed Khaled, Konstantin Mishchenko, and Peter Richtárik. 2020. Tighter theory for local SGD on identical and heterogeneous data. In *International Conference on Artificial Intelligence and Statistics*. PMLR, 4519–4529.
- [15] Jakub Konečný, H. Brendan McMahan, Daniel Ramage, and Peter Richtárik. 2016. Federated Optimization: Distributed Machine Learning for On-Device Intelligence. arXiv:1610.02527 [cs.LG]
- [16] Alex Krizhevsky, Geoffrey Hinton, et al. 2009. Learning multiple layers of features from tiny images.
- [17] Tian Li, Anit Kumar Sahu, Manzil Zaheer, Maziar Sanjabi, Ameet Talwalkar, and Virginia Smith. 2020. Federated optimization in heterogeneous networks. *Proceedings of Machine Learning and Systems* 2 (2020), 429–450.
- [18] Lumin Liu, Jun Zhang, SH Song, and Khaled B Letaief. 2020. Client-Edge-Cloud Hierarchical Federated Learning. In *Proceedings of the IEEE International Conference on Communications (ICC)*. IEEE, 1–6.
- [19] Yang Liu, Yan Kang, Xinwei Zhang, Liping Li, Yong Cheng, Tianjian Chen, Mingyi Hong, and Qiang Yang. 2020. A Communication Efficient Collaborative Learning Framework for Distributed Features. arXiv:1912.11187 [cs.LG] Presented in Workshop on Federated Learning for Data Privacy and Confidentiality, NeurIPS 2019.
- [20] Dhruv Mahajan, S Sathiya Keerthi, and S Sundararajan. 2017. A distributed block coordinate descent method for training l1 regularized linear classifiers. *The Journal of Machine Learning Research* 18, 1 (2017), 3167–3201.
- [21] Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Agüera y Arcas. 2017. Communication-efficient learning of deep networks from decentralized data. In *Artificial intelligence and statistics*. PMLR, 1273–1282.
- [22] Peter Richtárik and Martin Takáč. 2016. Distributed coordinate descent method for learning with big data. *The Journal of Machine Learning Research* 17, 1 (2016), 2657–2681.
- [23] Sebastian U. Stich. 2019. Local SGD Converges Fast and Communicates Little. In *International Conference on Learning Representations*. OpenReview.net. <https://openreview.net/forum?id=S1g2JnRCFX>
- [24] Chang Sun, Lianne Ippel, Johan Van Soest, Birgit Wouters, Alexander Malic, Onaopepo Adekunle, Bob van den Berg, Ole Müssmann, Annemarie Koster, Carla van der Kallen, et al. 2019. A Privacy-Preserving Infrastructure for Analyzing Personal Health Data in a Vertically Partitioned Scenario. *Studies in health technology and informatics* 264 (2019), 373–377.
- [25] Paul Tseng and Sangwoon Yun. 2009. A coordinate gradient descent method for nonsmooth separable minimization. *Mathematical Programming* 117, 1-2 (Aug. 2009), 387–423.
- [26] Jianyu Wang and Gauri Joshi. 2021. Cooperative sgd: A unified framework for the design and analysis of local-update sgd algorithms. *Journal of Machine Learning Research* 22, 213 (2021), 1–50.

- [27] Jiayi Wang, Shiqiang Wang, Rong-Rong Chen, and Mingyue Ji. 2020. Local Averaging Helps: Hierarchical Federated Learning and Convergence Analysis. (2020). [arXiv:2010.12998](#).
- [28] Shiqiang Wang, Tiffany Tuor, Theodoros Salonidis, Kin K Leung, Christian Makaya, Ting He, and Kevin Chan. 2019. Adaptive federated learning in resource constrained edge computing systems. *IEEE Journal on Selected Areas in Communications* 37, 6 (June 2019), 1205–1221.
- [29] Yansheng Wang, Yongxin Tong, and Dingyuan Shi. 2020. Federated latent Dirichlet allocation: A local differential privacy based framework. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 34. 6283–6290.
- [30] Yansheng Wang, Yongxin Tong, Dingyuan Shi, and Ke Xu. 2021. An Efficient Approach for Cross-Silo Federated Learning to Rank. In *Proceedings of the IEEE International Conference on Data Engineering*. IEEE, 1128–1139.
- [31] Yuncheng Wu, Shaofeng Cai, Xiaokui Xiao, Gang Chen, and Beng Chin Ooi. 2020. Privacy preserving vertical federated learning for tree-based models. *Proceedings of the VLDB Endowment* 13, 12 (Aug. 2020), 2090–2103.
- [32] Zhirong Wu, Shuran Song, Aditya Khosla, Fisher Yu, Linguang Zhang, Xiaoou Tang, and Jianxiong Xiao. 2015. 3d shapenets: A deep representation for volumetric shapes. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 1912–1920.
- [33] Qiang Yang, Yang Liu, Tianjian Chen, and Yongxin Tong. 2019. Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology (TIST)* 10, 2 (2019), 1–19.
- [34] Shengwen Yang, Bing Ren, Xuhui Zhou, and Liping Liu. 2019. Parallel Distributed Logistic Regression for Vertical Federated Learning without Third-Party Coordinator. [arXiv:1911.09824](#) [cs.LG]
- [35] Hao Yu, Sen Yang, and Shenghuo Zhu. 2019. 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*, Vol. 33. 5693–5700.