

Multi-Server FL with Overlapping Clients: A Latency-Aware Relay Framework

Yun Ji¹, Zeyu Chen¹, Xiaoxiong Zhong², Yanan Ma³, Sheng Zhang^{1,*}, and Yuguang Fang³

¹Tsinghua Shenzhen International Graduate School, Tsinghua University, Shenzhen 518055, China

²Department of New Networks, Peng Cheng Laboratory, Shenzhen 518000, China

³Department of Computer Science, City University of Hong Kong, Kowloon, Hong Kong

**Corresponding author: Sheng Zhang, email: zhangsh@sz.tsinghua.edu.cn*

Abstract—Multi-server Federated Learning (FL) has emerged as a promising solution to mitigate communication bottlenecks of single-server FL. In a typical multi-server FL architecture, the regions covered by different edge servers (ESs) may overlap. Under this architecture, clients located in the overlapping areas can access edge models from multiple ESs. Building on this observation, we propose a cloud-free multi-server FL framework that leverages Overlapping Clients (OCs) as relays for inter-server model exchange while uploading the local updated model to ESs. This enables ES models to be relayed across multiple hops through neighboring ESs by OCs without introducing new communication links. We derive a new convergence upper bound for non-convex objectives under non-IID data and an arbitrary number of cells, which explicitly quantifies the impact of inter-server propagation depth on convergence error. Guided by this theoretical result, we formulate an optimization problem that aims to maximize dissemination range of each ES model among all ESs within a limited latency. To solve this problem, we develop a conflict-graph-based local search algorithm optimizing the routing strategy and scheduling the transmission times of individual ESs to its neighboring ESs. This enables ES models to be relayed across multiple hops through neighboring ESs by OCs, achieving the widest possible transmission coverage for each model without introducing new communication links. Extensive experimental results show remarkable performance gains of our scheme compared to existing methods.

Index Terms—Multi-server federated learning, Edge computing, Overlapping clients.

I. INTRODUCTION

Federated Learning (FL) has gained considerable attention due to its ability to enable collaborative model training while preserving data privacy [1], [2]. A typical FL system operates in rounds, where clients perform local training based on their own datasets and then upload model updates to a cloud server for aggregation and broadcasting in the next round of training. However, the limitation of this single-server architecture is the long-range communication latency between clients and the cloud server (CS), especially in large-scale FL systems. To address this issue, multi-server FL architectures have been proposed [3], [4]. A prominent example is hierarchical federated learning (HFL) [3], where ESs are responsible for aggregating local model updates from their associated clients and subsequently uploading their aggregated models to the CS. However, frequent communications between ESs and the CS can still lead to high latency and resource consumption.

In [5] and [6], a new FL architecture utilizing multiple servers is studied, which exploits the realistic deployment of 5G-and-beyond networks where a client can be located in the overlapping coverage areas of multiple servers. The key idea is that clients download multiple models from all the edge servers they can access and train their local models based on the average of these models. However, OCs must wait for all relevant ESs to deliver their models before training, leading to considerable delay when inter-edge heterogeneity is pronounced. Moreover, when the number of OCs is limited, their contributions to global model convergence is marginal.

To solve these challenges, we have proposed FedOC [7], where OCs can act as relay nodes that forward edge models between adjacent cells upon reception, thereby enabling distributed aggregation across all ESs. However, this work proved convergence only for three cells, and each ES model in the current round could reach only its neighboring ESs. In this paper, we extend the convergence proof to an arbitrary number of cells and optimize ES transmission timing to maximize the propagation range of each ES model among all ESs, thereby accelerating convergence within a finite time horizon.

In summary, we highlight our main contributions as follows:

- We propose a distributed aggregated multi-server FL framework, where we exploit the OCs for model exchange between neighboring cells and facilitate cross-server collaboration. We derive the upper bound of the gap between the expected and optimal global loss values with respect to the propagation range of each ES model among all ESs. Based on this, we merge model relaying with local update uploading into a single transmission, and further optimize the model relay strategy and transmission scheduling to neighboring ESs for each ES model, so as to maximize the propagation range of each ES model within a finite time without introducing new communication links.
- To solve the optimization problem, we design a conflict-graph-based local search algorithm that jointly considers path dependency and communication constraints. For each ES l , the algorithm identifies both the farthest reachable path and the farthest received model path, constructing a conflict graph to capture mutual transmission conflicts. By combining greedy initialization with

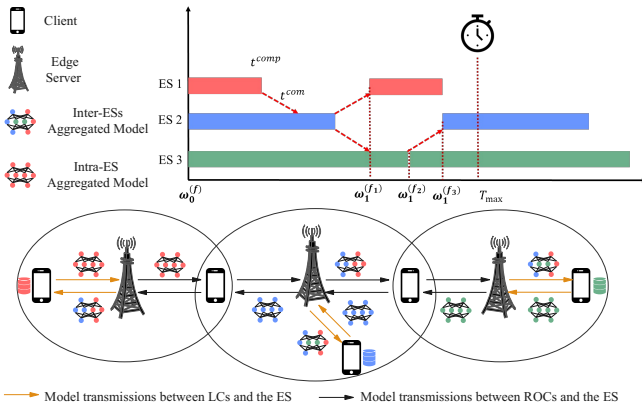


Fig. 1: Illustration of our proposed algorithm.

iterative local search, it maximizes the path weight, i.e., the total aggregated data volume across all ESs, thereby enhancing overall model propagation depth.

- We conduct extensive experiments on multiple datasets under varying numbers of cells and heterogeneous data distribution. The results demonstrate that our algorithm consistently outperforms the compared benchmarks in terms of both convergence accuracy and speed.

II. SYSTEM MODEL

A. Network Architecture

The system consists of a set of L ESs indexed by $\mathcal{L} = \{1, 2, \dots, L\}$, where each ES serves a group of clients within its coverage area, and some clients fall within the overlapping regions of multiple ESs. We refer to the local coverage area of each ES as a *cell*. Clients are categorized into two types: local clients (LCs), which are located in non-overlapping regions and communicate with only one ES, and overlapping clients (OCs), which are located in overlapping regions and can communicate with multiple ESs.

As shown in Fig. 1, we model overlaps using a chain topology, where each ES l overlaps with its immediate neighbors, ES $l-1$ and $l+1$, except for boundary ESs (ES 1 and ES L), which overlap with only one neighbor. While not necessarily implying a physical line, this structure captures typical cases of localized overlaps, e.g., due to limited signal range or linear deployment patterns like roads or urban blocks. It simplifies analysis while preserving the key challenges posed by overlapping coverage and multi-server coordination [7]. A single client in each overlap region, denoted $b_{l,l+1}$, serves as the Relay Overlapping Client (ROC) to forward models between ESs, while the rest are Normal OCs (NOCs). The commonly used notations are summarized in Table I.

B. Algorithm Description

Now we describe our FL algorithm tailored to the above setup. As shown in Fig. 1, the training process can be summarized as following stages, i.e., 1) local model update and upload. 2) intra-ES model aggregation. 3) inter-cell model

TABLE I: Frequently Used Nomenclature and Notations

Notation	Description
$n^{(k)}$	The number of data samples of client k
$b_{l,l+1}$	The relay overlapping client (ROC) in the overlapped area between cell l and $l+1$
R	The total number of rounds
$\mathcal{S}_l$	The set of clients transmitting local models to ES l
$\tilde{N}_r^{(f_l)}$	The total number of data volume clients in $\mathcal{S}_l$
N	The total number of data samples of all clients
$\tilde{\mathbf{w}}_{r,E}^{(f_l)}$	The initial ES model of cell l is obtained by aggregating the models of clients in $\mathcal{S}_l$.
$\mathbf{w}_{r,E}^{(l+1,l)}$	The model transmitted from ES $l+1$ to ES l by ROC $b_{l+1,l}$
$\mathbf{w}_{r+1}^{(f_l)}$	The ES model of cell l for broadcasting to clients at the start of round $r+1$

exchange through ROC relays. 4) ES model aggregation and broadcasting.

At the beginning of each round, clients download the edge models from ESs and update their local models, with operations varying depending on client roles:

For each LC k , local training is performed for E iterations using stochastic gradient descent (SGD) on the received edge model. The update in the e -th iteration is computed as

$$\mathbf{w}_{r,e+1}^{(k)} = \mathbf{w}_{r,e}^{(k)} - \eta_{r,e} \nabla \ell_k(\mathbf{w}_{r,e}^{(k)}), \quad (1)$$

where $e = 0, 1, \dots, E-1$, and $\nabla \ell_k(\mathbf{w}_{r,e}^{(k)})$ is the gradient of the local loss function. After local updating, each LC k transmits the local model $\mathbf{w}_{r,E}^{(k)}$ to corresponding ES.

For OCs, an overlapping client $k \in \mathcal{V}_{l,l+1}$ could receive two models: $\mathbf{w}_r^{(f_l)}$ and $\mathbf{w}_r^{(f_{l+1})}$. To reduce the wait latency, we assign each OC to the nearest ES as its participating server for training. Once it receives the model disseminated by that ES, it performs local training as the same as LCs based on the received model. Finally, each NOC transmits $\mathbf{w}_{r,E}^{(k)}$ to selected ES i , ROCs keep their updated models locally and upload them during the subsequent phase. Each ES l aggregates the received local models as

$$\tilde{\mathbf{w}}_{r,E}^{(f_l)} = \frac{\sum_{k \in \mathcal{S}_r^{(l)}} n^{(k)} \mathbf{w}_{r,E}^{(k)}}{\tilde{N}_r^{(f_l)}}, \quad (2)$$

where $\tilde{N}_r^{(f_l)} = \sum_{k \in \mathcal{S}_r^{(l)}} n^{(k)}$, and $\mathcal{S}_r^{(l)}$ consists of LCs and the NOCs that upload local models to ES l .

In the stage of inter-cell model exchange, since no additional communication links are introduced, each ES l can communicate with its ROCs at most once in each direction to relay the model to other ESs. Each ES first decides whether to transmit its model to neighboring ESs, and then determines whether to wait for the models received from its neighbors to be aggregated before forwarding them. In this way, the transmission timing to different neighbors is optimized to achieve the widest possible model propagation within a finite time horizon. The above routing strategy will be elaborated in the section IV. The relay rule of each ROC is: once ROC $b_{l,l+1}$ receives an edge model from either ES l or ES $l+1$, it will first aggregate this model with its own local model $\mathbf{w}_{r,E}^{(b_{l,l+1})}$

trained during the model update stage, and then forward it to the other ES, and we denote the model that ROC $b_{l,l+1}$ transmits to ES l as $\mathbf{w}_{r,E}^{(l+1,l)}$.

For each ES l , after receiving the models and corresponding data volume from ROCs $b_{l,l+1}$ and $b_{l-1,l}$, ES l will aggregate these models with its intra-cell model to update edge model $\mathbf{w}_{r+1}^{(f_l)}$ for broadcast in $r+1$ -th iteration, which can be written as

$$\mathbf{w}_{r+1}^{(f_l)} = \frac{N_r^{(l+1,l)} \mathbf{w}_{r,E}^{(l+1,l)} + N_r^{(l-1,l)} \mathbf{w}_{r,E}^{(l-1,l)} + \tilde{N}_r^{(f_l)} \tilde{\mathbf{w}}_{r,E}^{(f_l)}}{N_r^{(l+1,l)} + N_r^{(l-1,l)} + \tilde{N}_r^{(f_l)}}. \quad (3)$$

For boundary cells, the missing neighbor's model is taken as zero. so (3) holds for all $l \in \mathcal{L}$.

Noted that $\mathbf{w}_{r,E}^{(l-1,l)}$ may contain not only the locally aggregated model of cell $l-1$ but also models relayed from other cells through multi-hop transmissions. For instance, the model of cell i can be forwarded along the path $\mathcal{P}_{i \rightarrow (l-1)} = \{\text{ES } i, \text{ROC } b_{i,i+1}, \text{ES } (i+1), \dots, \text{ROC } b_{l-2,l-1}, \text{ES } (l-1)\}$, where each relay node aggregates the received model with its own before forwarding, provided that the timing constraint is satisfied. We denote $\mathbf{p}_r^{(f_l)} = [p_r^{(1,l)}, \dots, p_r^{(L,l)}]$, where $p_r^{(i,l)} \in \{0,1\}$ indicates whether cell i 's model participates in ES l 's aggregation. Based on (3), we have:

$$\begin{aligned} \mathbf{w}_{r+1}^{(f_l)} &= \frac{1}{N_r^{(f_l)}} \left[\sum_{i=1}^{l-1} p_r^{(i,l)} \left(\tilde{N}_r^{(f_l)} \tilde{\mathbf{w}}_{r,E}^{(f_l)} + n^{(b_{i,i+1})} \mathbf{w}_{r,E}^{(b_{i,i+1})} \right) \right. \\ &\quad \left. + p_r^{(l,l)} \tilde{N}_r^{(f_l)} \tilde{\mathbf{w}}_{r,E}^{(f_l)} \right. \\ &\quad \left. + \sum_{i=l+1}^L p_r^{(i,l)} \left(\tilde{N}_r^{(f_l)} \tilde{\mathbf{w}}_{r,E}^{(f_l)} + n^{(b_{i-1,i})} \mathbf{w}_{r,E}^{(b_{i-1,i})} \right) \right] \\ &= \sum_{i=1}^L \frac{p_r^{(i,l)} \tilde{N}_r^{(f_l)} \tilde{\mathbf{w}}_{r,E}^{(f_l)}}{\sum_{i=1}^L p_r^{(i,l)} \tilde{N}_r^{(f_l)}}, \end{aligned} \quad (4)$$

where $\tilde{\mathbf{w}}_{r,E}^{(f_l)}$ is denoted as

$$\tilde{\mathbf{w}}_{r,E}^{(f_l)} = \sum_{k \in \hat{\mathcal{K}}_r^{(f_l)}} \frac{n^{(k)}}{\hat{N}_r^{(f_l)}} \mathbf{w}_{r,E}^{(k)}, \quad (5)$$

where $\hat{N}_r^{(f_l)} = \sum_{k \in \hat{\mathcal{K}}_r^{(f_l)}} n^{(k)}$, and the client set $\hat{\mathcal{K}}_r^{(f_l)}$ satisfies:

$$\hat{\mathcal{K}}_r^{(f_l)} = \begin{cases} \mathcal{S}_i \cup b_{i,i+1}, & 1 \leq i < l \\ \mathcal{S}_i, & i = l \\ \mathcal{S}_i \cup b_{i-1,i}, & l+1 \leq i \leq L. \end{cases} \quad (6)$$

C. Latency Analysis

Several key event timings are critical to our algorithm:

- **Broadcast time** (t_l^{cast}): Time for ES l to broadcast its model to its covered clients.

- **Cell update time** (t_l^{comp}): Time for all clients in a cell to complete local computations and uploading model updates for aggregation.
- **Relay start time** ($t_{l,j}^{\text{start}}$): Time when ES l begins transmitting its local aggregated ES model to a neighboring cell $j \in \{l-1, l+1\}$ through the ROC $b_{l,j}$.
- **Communication transmission time** ($t_{l,j}^{\text{com}}$): Time taken for ES l to transmit its model to a neighboring cell j through the ROC $b_{l,j}$.
- **Model aggregation time** (t_l^{agg}): Time when ES l final aggregates all the received ES models.

Assume a total spectrum budget of B for model transmission. Adjacent cells operate on disjoint sub-bands, each allocated $B/2$ bandwidth. Each ES broadcasts the model to its clients and uniformly allocates the available bandwidth among them for local model uploads. After aggregation, the ES reclaims the bandwidth and forwards the model to neighboring ESs via ROCs. The communication latency can be given by

$$t_{l,l+1}^{\text{com}} = \frac{M}{\frac{B}{4} \left[\log \left(1 + \frac{4\delta_{l,l+1}P}{BN_0} \right) + \log \left(1 + \frac{4\delta_{l,l+1}P}{BN_0} \right) \right]}, \quad (7)$$

where M is the model size, and $\delta_{l,l+1}$ represents the channel gain from ES l to $l+1$, with P and p being the transmission powers of the ES and client, respectively. N_0 is the noise power spectral density. The relay start time is constrained by

$$t_{l,j}^{\text{start}} \geq t_l^{\text{cast}} + t_l^{\text{comp}}, \quad \forall j \in \{l-1, l+1\}. \quad (8)$$

The aggregation time t_l^{agg} is determined by:

$$t_l^{\text{agg}} = \max \left\{ p^{(j,l)} (t_{j,l}^{\text{start}} + t_{j,l}^{\text{com}}) \mid j \in \{l-1, l, l+1\} \right\}. \quad (9)$$

III. THEORETICAL RESULTS

In this paper, we consider a C -class classification problem over a compact feature space $\mathcal{X}$ with label space $\mathcal{Y} = \{1, 2, \dots, C\}$. Following [8], we define the population loss under the cross-entropy as $\ell(\mathbf{w}) = \sum_{i=1}^C P_{y=i} \mathbb{E}_{\mathbf{x}|y=i} [\log f_i(\mathbf{x}, \mathbf{w})]$, where $P_{y=i}$ denotes the probability of class i , and f_i is the predicted probability of class i parameterized by the model weights [8].

Now we introduce the widely used assumptions in FL:

Assumption 1. $\nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \log f_i(\mathbf{x}, \mathbf{w})$ is $\lambda_{\mathbf{x}|y=i}$ -Lipschitz for each class $i \in [C]$, i.e., for all $\mathbf{v}$ and $\mathbf{w}$, $\|\nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \log f_i(\mathbf{x}, \mathbf{w}) - \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \log f_i(\mathbf{x}, \mathbf{v})\| \leq \lambda_{\mathbf{x}|y=i} \|\mathbf{w} - \mathbf{v}\|$.

Theorem 1. Let Assumption 1 holds, assume $\sum_{i=1}^C P_{y=i} \lambda_{\mathbf{x}|y=i} = \lambda$, and set $\mathbf{w}^*$ to the optimal global model. If the learning rate satisfies $\eta_{r,e} = \frac{1}{(r+1)(E-1)} \forall r, e$, then we have

$$\begin{aligned} &\ell(\mathbf{w}_R^{(f_l)}) - \ell(\mathbf{w}^*) \\ &\leq \frac{\lambda}{2} \left[\sum_{j=1}^L D_{R-1}^{(j)} \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + \underbrace{\frac{\beta_{R-1}}{R}}_{\epsilon^{\text{intra}}} \right] \end{aligned}$$

$$+ \underbrace{\frac{\sum_{j=1}^3 \hat{N}_{R-1}^{(f_j)} H^{(j)} \left(\sum_{i=1}^C \left\| P_{y=i}^{(c_j)} - P_{y=i}^{(c)} \right\| \right)}{NR(E-1)}}_{\epsilon^{inter}} + F_{R-1}^{(l)} \Big], \quad (10)$$

$$\text{where } F_{R-1}^{(l)} = \sum_{j=1}^L \left\| \frac{p_{R-1}^{(j,q)} \hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,q)} \hat{N}_{R-1}^{(f_j)}} - \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \right\| \left\| \hat{\mathbf{w}}_{R-1,E}^{(f_j)} \right\|,$$

$$\text{and } \beta_{R-1} = \frac{E-1}{N} \sum_{j=1}^L \sum_{k \in \hat{\mathcal{K}}_r^{(f_j)}} n^{(k)} \left(\sum_{i=1}^C |P_{y=i}^{(k)} - P_{y=i}^{(c_j)}| \right) \left(\prod_{e=1}^{E-1} a_{R-1,e}^{(k)} \right) g_{\max}(\mathbf{w}_{R-1}^{(c_j)}).$$

Proof. Please refer to the Appendix. $\square$

ϵ^{intra} quantifies the intra-cell heterogeneity, mainly reflected by the divergence between each client's local data distribution and that of the entire cell, i.e., $\sum_{i=1}^C |P_{y=i}^{(k)} - P_{y=i}^{(c_j)}|$ for $k \in \hat{\mathcal{K}}_r^{(f_j)}$, $\forall j$. This term vanishes as $R \rightarrow \infty$. ϵ^{inter} captures the inter-cell heterogeneity across different cells and consists of two parts. The first part corresponds to the divergence between each cell's distribution and that of the entire system, i.e., $\sum_{i=1}^C \left\| P_{y=i}^{(c_j)} - P_{y=i}^{(c)} \right\|$, which converges to zero as $R \rightarrow \infty$. The second part, $F_{R-1}^{(l)}$, arises from the model aggregation procedure in our algorithm. As the number of external cell models incorporated into $\mathbf{w}_r^{(f_i)}$ increases, i.e., more elements in $\mathbf{p}_r^{(f_i)} = [p_r^{(1,l)}, \dots, p_r^{(L,l)}]$ equals to 1, $F_r^{(l)}$ approaches zero. In the extreme case where every ES model $\mathbf{w}_r^{(f_i)}$ aggregates all ES models at each round, i.e., $\mathbf{p}_r^{(f_i)}$, $\forall l$ is an all-one vector, the process degenerates into centralized FL with full cloud aggregation, and thus $F_r^{(l)} = 0$. To sum up, the loss error is significantly influenced by the probability of each cell participating in the aggregation of other cell models in each round.

IV. PROBLEM FORMULATION AND SOLUTION

Based on **Theorem 1**, we need to minimize $F_r^{(l)}$ for all l , which is equivalent to maximizing the propagation range of each ES model within a limited time. This is achieved by optimizing the routing decisions and the relay start times of each ES, enabling broader model sharing among cells.

The transmission path from cell j to l is defined as path = {ES j , ES $j+1, \dots$, ES l }. Successful transmission requires: (1) a valid link between adjacent nodes, i.e., $p^{(q,q+1)} = 1$, (2) each model arriving before the next forwarding step, i.e., $t_{q,q+1}^{start} + t_{q,q+1}^{com} \leq t_{q+1,q+2}^{start}$.

We define the indicator variable:

$$s^{(q,q+1)} = \mathbb{1} \left\{ t_{q,q+1}^{start} + t_{q,q+1}^{com} \leq t_{q+1,q+2}^{start} \right\}. \quad (11)$$

Then, we have

$$p^{(j,l)} = p^{(l-1,l)} \prod_{q=j}^{l-2} (s^{(q,q+1)} p^{(q,q+1)}), \quad j \leq l-1, \quad (12)$$

$$p^{(j,l)} = p^{(l+1,l)} \prod_{q=j}^{l+2} (s^{(q,q-1)} p^{(q,q-1)}), \quad l+1 \leq j \leq L. \quad (13)$$

We denote the optimization problem as

$$\begin{aligned} \mathbf{P1}: \quad & \min_{\mathbf{t}^{start}, \mathbf{p}} \sum_{l=1}^L \sum_{j=1}^L \left| \frac{p^{(j,l)} \hat{N}^{(f_j)}}{\sum_{j=1}^L p^{(j,l)} \hat{N}^{(f_j)}} - \frac{\hat{N}^{(f_j)}}{\sum_{j=1}^L \hat{N}^{(f_j)}} \right| \\ \text{s.t.} \quad & (8), (9), (11) - (13), \\ & t_l^{agg} \leq T_{\max}, \quad \forall l \in \mathcal{L}. \end{aligned} \quad (14)$$

where $\mathbf{t}^{start} = \{t_{l,j}^{start}\}$, $\forall l, j$, $\mathbf{p} = \{\mathbf{p}^{(f_i)}\}$, $\forall l$ and $T_{\max}$ denotes the maximum duration of each round.

Since the transmission processes in opposite directions (toward left and right ESs) are independent in the constraints, the variables for each direction can be optimized separately. After removing the absolute value, **P1** can be reformulated to

$$\begin{aligned} \mathbf{P2}: \quad & \max_{\mathbf{t}_{right}^{start}, \mathbf{p}_{right}} \sum_{l=1}^L \sum_{j=1}^l (p^{(j,l)} - \mathbb{1}\{p^{(j,l)} = 0\}) \hat{N}^{(f_j)} \\ & + \max_{\mathbf{t}_{left}^{start}, \mathbf{p}_{left}} \sum_{l=1}^L \sum_{j=l+1}^L (p^{(j,l)} - \mathbb{1}\{p^{(j,l)} = 0\}) \hat{N}^{(f_j)} \\ \text{s.t.} \quad & (8), (9), (11) - (13), \\ & t_l^{agg} \leq T_{\max}, \quad \forall l \in \mathcal{L}. \end{aligned} \quad (15)$$

where $\mathbf{t}_{right}^{start} = \{t_{l,j}^{start}\}$, $\mathbf{p}_{right} = \{p^{(l,j)}\}$, $\forall l \leq j$, $\mathbf{t}_{left}^{start} = \{t_{l,j}^{start}\}$, $\mathbf{p}_{left} = \{p^{(l,j)}\}$, $\forall l > j$. We illustrate the algorithm using the first problem as an example. Firstly $\{p^{(j,j+1)}\}$ for $1 \leq j \leq L-1$ can be determined straightforwardly. If $t_{j,j+1}^{local} + t_{j,j+1}^{com} \leq T_{\max}$, then $p^{(j,j+1)} = 1$. This ensures that ES $j+1$ at least can receive the model from ES j . If $p^{(j,l)} = 1$, all intermediate relays from ES j to ES l must also be active. Thus, maximizing the objective is equivalent to finding the longest feasible relay paths. To achieve this, each ES j transmits its cell model to ES $j+1$ once intra-cell aggregation is completed. After receiving the model, ES $j+1$ forwards the aggregated model at $\max\{t_{j,j+1}^{start} + t_{j,j+1}^{com}, t_{j+1}^{start} + t_{j+1}^{com}\}$, ensuring $s^{(j,j+1)} = 1$ as defined in (11). This process continues until $s_{q,q+1} = 0$, where no feasible $t_{q,q+1}^{start}$ satisfies $t_{q,q+1}^{start} + t_{q,q+1}^{com} \leq T_{\max}$. The resulting longest path from ES j is $\{j, j+1, \dots, q\}$. Accordingly, the optimal transmission path for each ES l is $\mathcal{P}_{right}^{(q,l)} = \{q, q+1, \dots, l\}$, where q is the earliest ES that can successfully transmit its model to ES l . Conflicts arise when two paths share intermediate nodes. We model this as a conflict graph $\mathcal{G}(\mathcal{V}, \mathcal{E})$, where vertices represent paths weighted by the total data volume of cells along this path, which we denote as $D^{(q,l)}$ for path $\mathcal{P}^{(q,l)}$, and edges denote conflicts. The problem is then equivalent to selecting a maximum-weight independent set. For small networks, exhaustive search provides the optimum. For larger networks, we propose a Conflict-Graph-Based Local Search Algorithm. The algorithm consists of three steps as follows.

Step 1. Generate an initial independent set via a greedy selection of non-conflicting high-weight vertices.

Step 2. Sort all paths by their end node numbers in ascending order, and find the sub-optimal and non-conflicted path for each node between the end node of one path and

Algorithm 1 Conflict-Graph-Based Local Search

Input: Routing decision p_{right} (or p_{left}), paths $\mathcal{P}^{(q,l)}$ with weights $D^{(q,l)}$.

Output: Relay start times $t_{\text{right}}^{\text{start}}$ (or $t_{\text{left}}^{\text{start}}$).

```

1: Initialize independent set  $\mathcal{I}^{\text{ini}}$  greedily by using Step 1.
2: Compute objective  $U(\mathcal{I}^{\text{ini}})$ .
3: for each  $i \in \mathcal{I}^{\text{ini}}$  do
4:   Replace  $i$  by feasible  $j$  to form  $\mathcal{I}^{\text{new}}$ , s.t.  $\{j, \mathcal{I} \setminus i\}$  still
   remain independent.
5:   Calculate  $\mathcal{C}(\mathcal{I}^{\text{new}})$  and  $U(\mathcal{I}^{\text{new}})$  by using step 2.
6:   if  $U(\mathcal{I}^{\text{new}}) > U(\mathcal{I}^{\text{final}})$  then
7:     Update  $\mathcal{I}^{\text{final}} = \mathcal{I}^{\text{new}}$ .
8:   end if
9: end for
10: Calculate  $\mathcal{C}(\mathcal{I}^{\text{final}})$  and  $U(\mathcal{I}^{\text{final}})$  by using step 2.
11: for each node in path  $\mathcal{P}^{(q,l)} \in \mathcal{C}(\mathcal{I}^{\text{final}})$  do
12:    $t_{j,j+1}^{\text{start}} = \max(t_{j-1,j}^{\text{start}} + t_{j-1,j}^{\text{com}}, t_j^{\text{cast}} + t_j^{\text{comp}})$ .
13: end for
14: return  $t_{\text{right}}^{\text{start}}$  (or  $t_{\text{left}}^{\text{start}}$ ).
```

the start node of the next path. Then adding these paths into $\mathcal{C}(\mathcal{I}^{\text{ini}})$ which we named as the full paths set of $\mathcal{I}^{\text{ini}}$, and calculate the objective of **P2**, which we denote as $U(\mathcal{I}^{\text{ini}})$.

Step 3. Iteratively refine the set by local search update higher $U(\mathcal{I}^{\text{final}})$.

The detailed algorithm is summarized in **Algorithm 1**.

V. PERFORMANCE EVALUATION

A. Simulation Setup

We evaluate a multi-server network with L ESs and $K = 60$ clients, each ES covering a circular area of radius 600m, where clients are randomly distributed. The wireless channel model includes Rayleigh small-scale fading and large-scale path-loss ($128.1 + 37.6 \log_{10}(d(\text{km}))$). Experiments are conducted on MNIST and CIFAR-10 with heterogeneous data distributions: each client has samples from two classes, and each ES is restricted to five classes, creating strong imbalance. The maximum per-round time T_{max} is aligned with FedOC, i.e., the maximum time for fair comparison. For the MNIST dataset, we adopt a lightweight Convolutional Neural Network (CNN) with 21840 model parameters [3]. For CIFAR-10, we employ a deeper six-layer CNN with 1.14 million parameters [4]. Key parameters are summarized in Table II.

We compare our method with the following baselines:

- **HFL [3]:** Hierarchical FL without overlapping regions.
- **FedMES [5]:** OCs train local models based on the aggregated multiple ES models.
- **FL-EOCD [9]:** After caching and training on aggregation of received ES models, each OC uploads a model that combines its local update with the cached ES models.
- **FedOC [2]:** ESs directly forward aggregated models to neighbors via ROCs after local aggregation, without transmission time optimization.

TABLE II: Simulation Parameters

Parameters	Values
Number of clients, K	60
Number of iterations, R	500
Number of local training epochs, $U_k(r)$	5
Client's transmit power, p (W)	1
ES's transmit power, P (W)	5
Total channel bandwidth, B (MHz)	50
Client's one-epoch update time (seconds)	MNIST: [0.1, 0.2], CIFAR-10: [1, 2]
Noise power spectral density, N_0 (dBm/Hz)	-174
The number of model parameters	MNIST: 21840, CIFAR-10: 1.14 million
Batch size	20
Initial learning rate, η	MNIST: 0.01, CIFAR-10: 0.1
Exponential decay factor of η in SGD	MNIST: 0.995, CIFAR-10: 0.992

TABLE III: The average number of clients aggregated per cell.

Dataset	Cells	FedOC	Ours
MNIST	3 cells	47.67	53.00
	5 cells	31.40	50.00
	6 cells	26.76	50.00
CIFAR-10	3 cells	46.67	55.33
	5 cells	31.40	53.40
	6 cells	27.79	34.35

As shown in Table III, we compare the average number of clients aggregated per cell between FedOC and our method. Under the same round time T_{max} , our approach enables each cell to receive significantly more ES models, including those from distant cells, without extra communication links. This confirms the effectiveness of our algorithm.

Fig. 2 compares the accuracy of different benchmarks across various numbers of cells. HFL and FedMES with one OC only perform intra-cell aggregation in each round, thus their models plateau at a lower accuracies. When the OCs increase to $\lfloor K/(2L) \rfloor$ clients per region, the accuracy of FedMES increases significantly, benefiting from its reliance on multiple OCs to aggregate ES models, which indirectly enhances the generalization of each ES model. However, relying solely on the limited aggregation within OCs provides only constrained gains in model generalization, especially when the number of cells increases. FL-EOCD fares better: by integrating other cells' models, albeit with a delay, it achieves higher accuracy than HFL/FedMES in this scenario. In some MNIST cases (e.g., 5 cells) it can even outperform FedOC. However, its advantage diminishes on the more challenging CIFAR-10 task due to model staleness. In contrast, FedOC leverages real-time transmission between neighboring ESs via ROCs, achieving performance second only to our algorithm in most cases. Our algorithm builds upon FedOC by incorporating optimizations for the routing and transmission times of model exchanges, enabling models to be transmitted to more distant ESs without

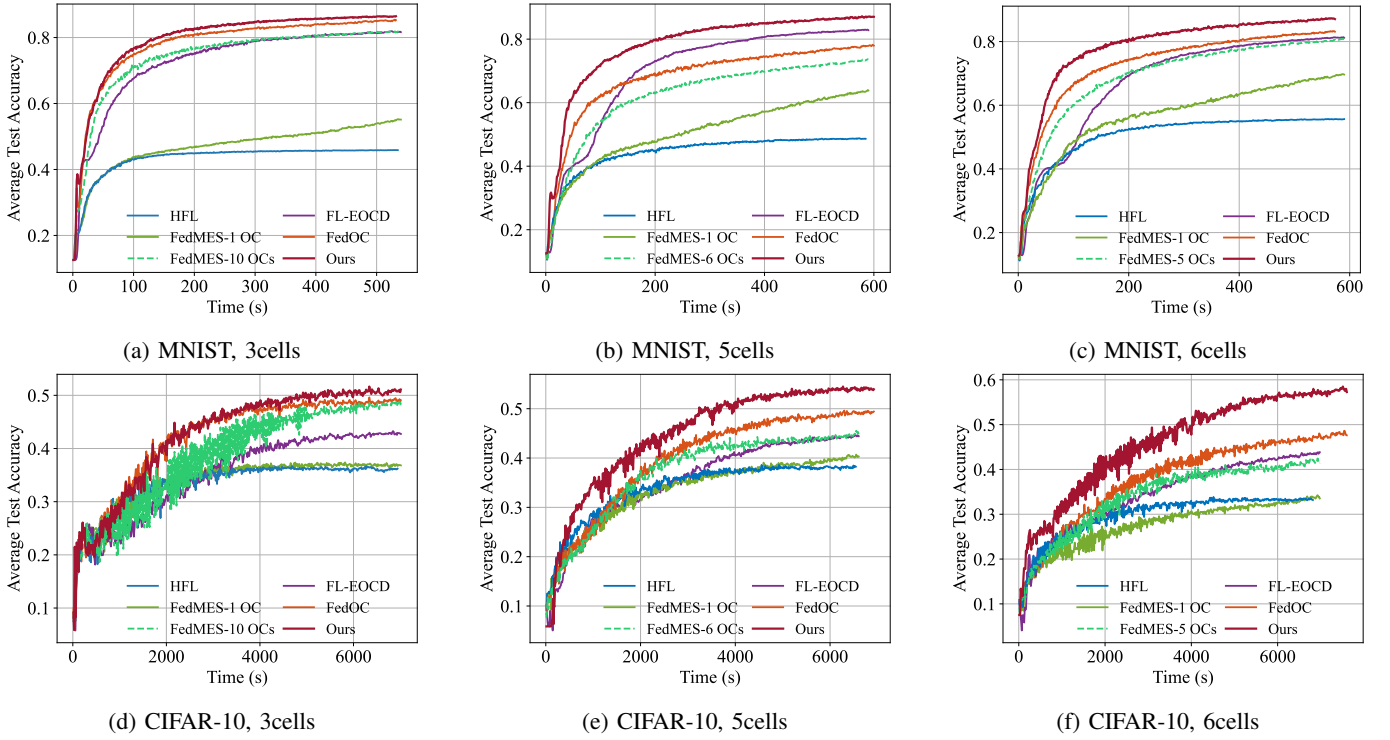


Fig. 2: Average test accuracy versus training time

adding extra communication links or time. This optimization accelerates model fusion, leading to faster convergence and higher model accuracy.

VI. CONCLUSION

In this paper, we presented a cloud-free multi-server FL framework leveraging overlapping clients as relays to enable multi-hop model dissemination among ESs without introducing additional communication links. We derived a new convergence bound of loss error of our algorithm, which characterizes the relationship between each ES model propagation depth among all ESs and the convergence error. Building upon this, we formulate an optimization problem maximizing the propagation range of edge models under latency constraints by jointly scheduling inter-server route decisions and transmission times. To solve this problem, we proposed a conflict-graph-based local search algorithm that models path dependencies and transmission conflicts as a maximum-weight independent set problem. By combining greedy initialization and iterative local refinement, it efficiently derives near-optimal relay schedules with low computational complexity. Experimental results demonstrated that the proposed framework consistently outperforms existing schemes, achieving faster convergence and higher accuracy. The gains become more pronounced as the number of cells increases, confirming that optimized relay scheduling enables deeper dissemination of models across multiple ESs and thereby accelerates convergence.

APPENDIX

PROOF OF THE THEOREM 1

Proof. From (6), the set $\hat{\mathcal{K}}_r^{(f_i)}$ is not only associated with cell i but also depends on cell l , which complicates the analysis. Since this difference only arises from the assignment of a single ROC and does not affect the overall client data of each cell, we approximate by assigning each ROC to the left cell in the mathematical formulation. That is,

$$\hat{\mathcal{K}}_r^{(f_i)} = \mathcal{S}_i \cup b_{i,i+1}, \quad 1 \leq i < L - 1. \quad (16)$$

Under these definitions, each ROC is mathematically associated with a unique cell. As shown in (4) the edge model of each ES l is the weighted sum of all participating ES models.

Since the population loss $\ell(\mathbf{w})$ satisfies $\ell(\mathbf{w}) = \mathbb{E}_{\mathbf{x}, y \sim P} \left[\sum_{i=1}^C \mathbb{1}_{y=i} \log f_i(\mathbf{x}, \mathbf{w}) \right] = \sum_{i=1}^C P_{y=i} \mathbb{E}_{\mathbf{x}|y=i} [\log f_i(\mathbf{x}, \mathbf{w})]$, the local update in (1) can be written as

$$\mathbf{w}_{r,e+1}^{(k)} = \mathbf{w}_{r,e}^{(k)} - \eta_{r,e} \sum_{i=1}^C P_{y=i}^{(k)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{r,e}^{(k)}) \right], \quad (17)$$

where $P_{y=i}^{(k)}$ represents the data proportion of class i on clients k .

To derive the convergence bound of the loss function, we introduce a cell-centralized SGD algorithm, where each ES collects the datasets of its associated clients, performs local model training, and then uploads the updates to the CS for

aggregation. Let $\mathbf{w}_r^{(c)}$ denote the global model of r -th round, and $\mathbf{w}_{r,e+1}^{(c_l)}$ represent the e -th update in the r -th round under the centralized setting. Then cell-centralized SGD performs the following update

$$\mathbf{w}_{r,e+1}^{(c_l)} = \mathbf{w}_{r,e}^{(c_l)} - \eta_{r,e} \sum_{i=1}^C P_{y=i}^{(c_l)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{r,e}^{(c_l)}) \right], \quad (18)$$

where $e = 0, 1, \dots, E-1$, and $P_{y=i}^{(c_l)}$ represents the data distribution of participating clients in cell l . After E times update, the aggregated model for next round $r+1$ in centralized settings can be written as

$$\mathbf{w}_{r+1}^{(c)} = \sum_{i=1}^L \frac{\hat{N}_r^{(f_i)} \mathbf{w}_{r,E}^{(c_i)}}{\sum_{i=1}^L \hat{N}_r^{(f_i)}}. \quad (19)$$

Combining **Assumption 1** with $\sum_{i=1}^C P_{y=i} \lambda_{x|y=i} = \lambda$, we obtain that $\ell(\mathbf{w})$ is λ -smooth, i.e., $\|\nabla \ell(\mathbf{w}) - \nabla \ell(\mathbf{v})\| \leq \lambda \|\mathbf{w} - \mathbf{v}\|$, $\forall \mathbf{v}, \mathbf{w}$. Then we have

$$\begin{aligned} \ell(\mathbf{w}_R^{(f_i)}) - \ell(\mathbf{w}^*) &\leq \frac{\lambda}{2} \|\mathbf{w}_R^{(f_i)} - \mathbf{w}^*\| \\ &= \frac{\lambda}{2} \|\mathbf{w}_R^{(f_i)} - \mathbf{w}_R^{(c)} + \mathbf{w}_R^{(c)} - \mathbf{w}^*\| \\ &\leq \frac{\lambda}{2} \underbrace{\|\mathbf{w}_R^{(f_i)} - \mathbf{w}_R^{(c)}\|}_{A_1} + \frac{\lambda}{2} \underbrace{\|\mathbf{w}_R^{(c)} - \mathbf{w}^*\|}_{A_2}, \end{aligned} \quad (20)$$

We next focus on bounding A_1 . Based on the definitions of $\mathbf{w}_R^{(f_i)}$ and $\mathbf{w}_R^{(c)}$ in (4) and (19), we have

$$\begin{aligned} A_1 &= \left\| \mathbf{w}_R^{(f_i)} - \mathbf{w}_R^{(c)} \right\| \\ &= \left\| \sum_{j=1}^L \frac{p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)} \hat{\mathbf{w}}_{R-1,E}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}} - \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)} \mathbf{w}_{R-1,E}^{(c_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \right\| \\ &= \left\| \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \left(\frac{p_{R-1}^{(j,l)} \sum_{j=1}^L \hat{N}_{R-1}^{(f_j)} \hat{\mathbf{w}}_{R-1,E}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}} - \mathbf{w}_{R-1,E}^{(c_j)} \right) \right\| \\ &= \left\| \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \left[\left(\frac{p_{R-1}^{(j,l)} \sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}} - 1 \right) \hat{\mathbf{w}}_{R-1,E}^{(f_j)} + \left(\hat{\mathbf{w}}_{R-1,E}^{(f_j)} - \mathbf{w}_{R-1,E}^{(c_j)} \right) \right] \right\| \\ &\leq \sum_{j=1}^L \left\| \frac{p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}} - \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \right\| \|\hat{\mathbf{w}}_{R-1,E}^{(f_j)}\| \\ &\quad + \underbrace{\left\| \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \left(\hat{\mathbf{w}}_{R-1,E}^{(f_j)} - \mathbf{w}_{R-1,E}^{(c_j)} \right) \right\|}_{B_1}. \end{aligned} \quad (21)$$

We substitute (5), (18) into B_1 , and we have

$$\begin{aligned} B_1 &= \left\| \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \left(\hat{\mathbf{w}}_{R-1,E}^{(f_j)} - \mathbf{w}_{R-1,E}^{(c_j)} \right) \right\| \\ &\quad \cdot \sum_{i=1}^C P_{y=i}^{(k)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-1}^{(k)}) \right] - \mathbf{w}_{R-1,E-1}^{(c_j)} \\ &\quad + \eta_{R-1,E-1} \sum_{i=1}^C P_{y=i}^{(c_j)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-1}^{(c_j)}) \right] \Big\| \\ &\stackrel{1}{=} \left\| \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} \left[\left(\mathbf{w}_{R-1,E-1}^{(k)} - \mathbf{w}_{R-1,E-1}^{(c_j)} \right) \right. \right. \\ &\quad \left. \left. + \eta_{R-1,E-1} \sum_{i=1}^C P_{y=i}^{(k)} \left(\nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-1}^{(k)}) \right] - \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-1}^{(c_j)}) \right] \right) \right] \right\| \\ &\stackrel{2}{\leq} \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} \cdot \left(1 + \eta_{R-1,E-1} \sum_{i=1}^C P_{y=i}^{(k)} \lambda_{x|y=i} \right) \cdot \left\| \mathbf{w}_{R-1,E-1}^{(k)} - \mathbf{w}_{R-1,E-1}^{(c_j)} \right\| \end{aligned} \quad (22)$$

where equality 1 holds because for each class $i \in \{1, 2, \dots, C\}$, $P_{y=1}^{(c_j)} = \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} P_{y=1}^{(k)}$, i.e., the data distribution over all the clients in cell j is the same as the data distribution over the whole cell population. Inequality 2 holds because of the **Assumption 1**. Inspired by [8], we give the bounding of $\left\| \mathbf{w}_{R-1,E-1}^{(k)} - \mathbf{w}_{R-1,E-1}^{(c_j)} \right\|$ for client $k \in \hat{\mathcal{K}}_r^{(c_j)}$ in cell j :

$$\begin{aligned} &\left\| \mathbf{w}_{R-1,E-1}^{(k)} - \mathbf{w}_{R-1,E-1}^{(c_j)} \right\| \\ &\stackrel{3}{=} \left\| \left(\mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right) \right. \\ &\quad \left. + \left(\eta_{R-1,E-2} \sum_{i=1}^C P_{y=i}^{(k)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(c_j)}) \right] - \eta_{R-1,E-2} \sum_{i=1}^C P_{y=i}^{(k)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(k)}) \right] \right) \right\| \\ &\leq \left\| \mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right\| \\ &\quad + \eta_{R-1,E-2} \left\| \sum_{i=1}^C P_{y=i}^{(k)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(k)}) \right] \right. \\ &\quad \left. - \sum_{i=1}^C P_{y=i}^{(k)} \nabla_{\mathbf{w}} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(c_j)}) \right] \right\| \end{aligned}$$

$$\begin{aligned}
& + \sum_{i=1}^C P_{y=i}^{(k)} \nabla \mathbf{w} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(c_j)}) \right] \\
& - \sum_{i=1}^C P_{y=i}^{(c_j)} \nabla \mathbf{w} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(c_j)}) \right] \Big\| \\
& \leq \left\| \mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right\| \\
& + \eta_{R-1,E-2} \left\| \sum_{i=1}^C P_{y=i}^{(k)} \left(\nabla \mathbf{w} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(k)}) \right] \right. \right. \\
& \quad \left. \left. - \nabla \mathbf{w} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(c_j)}) \right] \right) \right\| + \\
& \left\| \sum_{i=1}^C \left(P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right) \nabla \mathbf{w} \mathbb{E}_{\mathbf{x}|y=i} \left[\log f_i(\mathbf{x}, \mathbf{w}_{R-1,E-2}^{(c_j)}) \right] \right\| \\
& \stackrel{4}{\leq} \left(1 + \eta_{R-1,E-2} \sum_{i=1}^C P_{y=i}^{(k)} \lambda_{\mathbf{x}|y=i} \right) \left\| \mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right\| \\
& + \eta_{R-1,E-2} \sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| g_{\max}(\mathbf{w}_{R-1,E-2}^{(c_j)}) \\
& \stackrel{5}{=} a_{R-1,E-2}^{(k)} \left\| \mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right\| \\
& + \eta_{R-1,E-2} \sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| g_{\max}(\mathbf{w}_{R-1,E-2}^{(c_j)}) \quad (23)
\end{aligned}$$

where equality 3 holds because (17) and (18). Inequality 4 holds because the **Assumption 1** and we denote $g_{\max}(\mathbf{w}_{r,q}^{(c_j)}) = \max_{i=1}^C \left\| \nabla \mathbf{w} \mathbb{E}_{\mathbf{x}|y=i} \log f_i(\mathbf{x}, \mathbf{w}_{r,q}^{(c_j)}) \right\|$ for $r \in \{1, 2, \dots, R\}$ and $q \in \{0, 1, \dots, E\}$. equality 5 holds because we denote $a_{r,q}^{(k)} = 1 + \eta_{r,q} \sum_{i=1}^C P_{y=i}^{(k)} \lambda_{\mathbf{x}|y=i}$ for $r \in \{1, 2, \dots, R\}$ and $q \in \{0, 1, \dots, E\}$. From (23) we deduce the relation between $\left\| \mathbf{w}_{R-1,E-1}^{(k)} - \mathbf{w}_{R-1,E-1}^{(c_j)} \right\|$ and $\left\| \mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right\|$.

By repeatedly using this relation, we can recursively obtain:

$$\begin{aligned}
& \left\| \mathbf{w}_{R-1,E-1}^{(k)} - \mathbf{w}_{R-1,E-1}^{(c_j)} \right\| \\
& \leq a_{R-1,E-2}^{(k)} \left\| \mathbf{w}_{R-1,E-2}^{(k)} - \mathbf{w}_{R-1,E-2}^{(c_j)} \right\| \\
& + \sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| \eta_{R-1,E-2} g_{\max}(\mathbf{w}_{R-1,E-2}^{(c_j)}) \\
& \leq \left(\prod_{q=0}^{E-2} a_{R-1,q}^{(k)} \right) \left\| \mathbf{w}_{R-1,0}^{(k)} - \mathbf{w}_{R-1,0}^{(c_j)} \right\| + \sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| \\
& \quad \cdot \sum_{q=0}^{E-2} \eta_{R-1,q} \left(\prod_{t=q+1}^{E-2} a_{R-1,t}^{(k)} \right) g_{\max}(\mathbf{w}_{R-1,q}^{(c_j)}) \\
& \leq \left(\prod_{q=0}^{E-2} a_{R-1,q}^{(k)} \right) \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + \sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| \\
& \quad \cdot \sum_{q=0}^{E-2} \eta_{R-1,q} \left(\prod_{t=q+1}^{E-2} a_{R-1,t}^{(k)} \right) g_{\max}(\mathbf{w}_{R-1,q}^{(c_j)}), \quad (24)
\end{aligned}$$

where we assume $\prod_a^b a_{r,e}^{(k)} = 1$, $\forall b < a$, r and e . Institute (24) into (22), we have

$$\begin{aligned}
B_1 & \leq \sum_{j=1}^L \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} \left[\left(\prod_{e=0}^{E-1} a_{R-1,e}^{(k)} \right) \right. \\
& \quad \cdot \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + \sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| \\
& \quad \cdot \sum_{d=0}^{E-2} \eta_{R-1,d} \left(\prod_{e=d+1}^{E-1} a_{R-1,e}^{(k)} \right) g_{\max}(\mathbf{w}_{R-1,d}^{(c_j)}) \Big] \\
& = \sum_{j=1}^L \left(D_{R-1}^{(j)} \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + G_{R-1}^{(j)} \right), \quad (25)
\end{aligned}$$

where $D_{R-1}^{(j)} = \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} \left(\prod_{e=0}^{E-1} a_{R-1,e}^{(k)} \right)$, and $G_{R-1}^{(j)} = \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} \left(\sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| \right) \sum_{e=0}^{E-2} \eta_{R-1,d} \left(\prod_{d=e+1}^{E-1} a_{R-1,d}^{(k)} \right) g_{\max}(\mathbf{w}_{R-1,e}^{(c_j)})$. Plug (25) into (21), we have

$$\begin{aligned}
A_1 & = \left\| \mathbf{w}_R^{(f_l)} - \mathbf{w}_R^{(c)} \right\| \\
& \leq \sum_{j=1}^L \left| \frac{p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}} - \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \right| \left\| \hat{\mathbf{w}}_{R-1}^{(f_j)} \right\| \\
& + \sum_{j=1}^L \left(D_{R-1}^{(j)} \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + G_{R-1}^{(j)} \right) \\
& \stackrel{6}{=} \sum_{j=1}^L D_{R-1}^{(j)} \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + F_{R-1}^{(l)} + G_{R-1} \quad (26)
\end{aligned}$$

where equality 6 holds because we denote

$$F_{R-1}^{(l)} = \sum_{j=1}^L \left| \frac{p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L p_{R-1}^{(j,l)} \hat{N}_{R-1}^{(f_j)}} - \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \right| \left\| \hat{\mathbf{w}}_{R-1,E}^{(f_j)} \right\|, \quad (27)$$

and $G_{R-1} = \sum_{j=1}^L G_{R-1}^{(j)}$.

In the following, we first derive a theoretical upper bound for G_{R-1} . We denote

$$G_{R-1}^{(j)} = \sum_{e=0}^{E-2} \eta_{R-1,e} \beta_{R-1,e}^{(j)}, \quad (28)$$

and

$$\beta_{R-1,e}^{(j)} = \frac{\hat{N}_{R-1}^{(f_j)}}{\sum_{j=1}^L \hat{N}_{R-1}^{(f_j)}} \sum_{k=1}^{\hat{K}_r^{(f_j)}} \frac{n^{(k)}}{\hat{N}_{R-1}^{(f_j)}} \left(\sum_{i=1}^C \left| P_{y=i}^{(k)} - P_{y=i}^{(c_j)} \right| \right)$$

$$\cdot \left(\prod_{d=e+1}^{E-1} a_{R-1,d}^{(k)} \right) g_{\max}(\mathbf{w}_{R-1,e}^{(c_j)}). \quad (29)$$

We assume that there exist constants $g_{\max}$ and $\beta_r^{(j)}$ such that $g_{\max}(\mathbf{w}_{r,e}^{(c_j)}) \leq g_{\max}$ and $\beta_{r,e}^{(j)} \leq \beta_r^{(j)}$, $\forall r, e, j$. Moreover, we denote $\sum_{j=1}^L \beta_r^{(j)} = \beta_r$. Since $a_{r,q}^{(k)} = 1 + \eta_{r,q} \sum_{i=1}^C P_{y=i}^{(k)} \lambda_{\mathbf{x}|y=i} > 1$, $\forall r, q$, based on (22), we further assume $\beta_{\min} \leq \beta_r \leq \beta_{\max}$. Then we have

$$G_{R-1} = \sum_{e=0}^{E-2} \eta_{R-1,e} \sum_{j=1}^L \beta_{R-1,e}^{(j)} \leq \sum_{e=0}^{E-2} \eta_{R-1,e} \beta_{R-1} \leq \frac{\beta_{R-1}}{R}, \quad (30)$$

Finally, we will give the upper bound of A_2 . A_2 is the model divergence between the cell-centralized SGD algorithm and $\mathbf{w}^*$. If we see each cell as a client, the cell-centralized SGD is the standard FedAvg. Under the assumption that $\mathbf{w}^*$ is the final model of the global centralized SGD, where data in all clients are collected in the CS to train a global model. Under the assumptions of **Theorem 1**, [8] gives the bound of A_2 as

$$A_2 = \left\| \mathbf{w}_R^{(c)} - \mathbf{w}^* \right\| \leq \frac{\sum_{j=1}^3 \hat{N}_{R-1}^{(f_j)} H^{(j)} \left(\sum_{i=1}^C \left\| P_{y=i}^{(c_j)} - P_{y=i}^{(c)} \right\| \right)}{NR(E-1)}. \quad (31)$$

where N is the total data volume of all clients and $H^{(j)} = \sum_{e=0}^{E-1} \left(\prod_{d=e+1}^{E-1} a_{R-1,d}^{(c_j)} \right) g_{\max}(\mathbf{w}_{R-1,e}^{(c)})$.

Substituting (26), (30) and (31) into (20), we have

$$\begin{aligned} & \ell(\mathbf{w}_R^{(f_1)}) - \ell(\mathbf{w}^*) \\ & \leq \frac{\lambda}{2} \left[\sum_{j=1}^L D_{R-1}^{(j)} \left\| \mathbf{w}_{R-1}^{(f_j)} - \mathbf{w}_{R-1}^{(c)} \right\| + \underbrace{\frac{\beta_{R-1}}{R}}_{\epsilon^{\text{inter}}} \right. \\ & \quad \left. + \underbrace{\frac{\sum_{j=1}^3 \hat{N}_{R-1}^{(f_j)} H^{(j)} \left(\sum_{i=1}^C \left\| P_{y=i}^{(c_j)} - P_{y=i}^{(c)} \right\| \right)}{NR(E-1)}}_{\epsilon^{\text{inter}}} + F_{R-1}^{(l)} \right] \end{aligned} \quad (32)$$

which completes the proof. $\square$

REFERENCES

- [1] B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y. Arcas, "Communication-Efficient Learning of Deep Networks from Decentralized Data," in *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*. PMLR, Apr. 2017, pp. 1273–1282, iSSN: 2640-3498.
- [2] G. Zhu, Y. Deng, X. Chen, H. Zhang, Y. Fang, and T. F. Wong, "ESFL: Efficient Split Federated Learning over Resource-constrained Heterogeneous Wireless Devices," *IEEE Internet of Things Journal*, vol. 11, no. 16, pp. 27 153–27 166, 2024.
- [3] L. Liu, J. Zhang, S. Song, and K. B. Letaief, "Client-Edge-Cloud Hierarchical Federated Learning," in *ICC 2020 - 2020 IEEE International Conference on Communications (ICC)*. Dublin, Ireland: IEEE, Jun. 2020, pp. 1–6.
- [4] X. Chen, G. Zhu, Y. Deng, and Y. Fang, "Federated Learning Over Multihop Wireless Networks With In-Network Aggregation," *IEEE Trans. Wireless Commun.*, vol. 21, no. 6, pp. 4622–4634, Jun. 2022.
- [5] D.-J. Han, M. Choi, J. Park, and J. Moon, "FedMes: Speeding Up Federated Learning With Multiple Edge Servers," *IEEE J. Select. Areas Commun.*, vol. 39, no. 12, pp. 3870–3885, Dec. 2021.
- [6] Z. Qu, X. Li, J. Xu, B. Tang, Z. Lu, and Y. Liu, "On the Convergence of Multi-Server Federated Learning With Overlapping Area," *IEEE Trans. on Mobile Comput.*, vol. 22, no. 11, pp. 6647–6662, Nov. 2023.
- [7] Y. Ji, Z. Chen, X. Zhong, Y. Ma, S. Zhang, and Y. Fang, "Fedoc: Multi-server fl with overlapping client relays in wireless edge networks," 2025. [Online]. Available: <https://arxiv.org/abs/2509.19398>
- [8] Y. Zhao, M. Li, L. Lai, N. Suda, D. Civin, and V. Chandra, "Federated Learning with Non-IID Data," Jul. 2022, arXiv:1806.00582 [cs].
- [9] M. S. Al-Abiad, M. Obeed, M. J. Hossain, and A. Chaaban, "Decentralized Aggregation for Energy-Efficient Federated Learning via D2D Communications," *IEEE Trans. Commun.*, vol. 71, no. 6, pp. 3333–3351, Jun. 2023.