

FED-CBS: A HETEROGENEITY-AWARE CLIENT SAMPLING MECHANISM FOR FEDERATED LEARNING VIA CLASS-IMBALANCE REDUCTION

Anonymous authors

Paper under double-blind review

ABSTRACT

Due to often limited communication bandwidth of edge devices, most existing federated learning (FL) methods randomly select only a subset of devices to participate in training in each communication round. Compared with engaging all the available clients, such random-selection mechanism could lead to significant performance degradation on non-IID (independent and identically distributed) data. In this paper, we present our key observation that the essential reason resulting in such performance degradation is the class-imbalance of the grouped data from randomly selected clients. Based on this observation, we design an efficient heterogeneity-aware client sampling mechanism, namely, Federated Class-balanced Sampling (Fed-CBS), which can effectively reduce class-imbalance of the grouped dataset from the intentionally selected clients. We first propose a measure of class-imbalance and then employ homomorphic encryption to derive this measure in a privacy-preserving way. Based on this measure, we design a computation-efficient client sampling strategy such that the actively selected clients will generate a more class-balanced grouped dataset with theoretical guarantees. Experimental results show that Fed-CBS outperforms the status quo approaches in terms of test accuracy and the rate of convergence, while achieving comparable or even better performance than the ideal setting where all the available clients participate in the FL training.

1 INTRODUCTION

With the booming of IoT devices, a huge amount of data is generated at the network edge, providing valuable resources for learning insightful information and enabling intelligent applications such as self-driving, video analytics, anomaly detection, etc. The traditional wisdom is to train machine learning models by collecting data from edge devices and performing centralized training. The data migration between the edge and the center server, however, raises serious privacy concerns. Federated learning (FL) (Konečný et al., 2016; McMahan et al., 2017a; Kairouz et al., 2019) is a popular technique to mitigate such privacy concerns which allows a large number of clients to collaboratively learn a shared model. The learning process is orchestrated by a central server: the participating clients download a global model from the central server and then compute local model updates using their local data. The clients then transmit the local updates to the central server where the local updates are aggregated and the global model is updated accordingly.

In practice, due to limited communication and computing capabilities, one usually can not engage all the available clients in FL training to fully utilize all the local data. Therefore, most FL methods only randomly select a subset of the available clients to participate in the training in each communication round. However, in practice, the data held by different clients are often typically non-IID (independent and identically distributed) due to various user preferences and usage patterns. This leads to a serious problem that the random client selection strategy often fails to learn a global model that can generalize well for most of the participating clients under non-IID settings. This has been identified by many recent studies (Cho et al., 2020; Goetz et al., 2019; Nishio & Yonetani, 2019; Yang et al., 2020). Several heuristic client selection mechanisms have been proposed to tackle this challenge. For example, the method in Goetz et al. (2019) makes the clients with larger local loss have a higher probability to be selected in the training. Power-of-Choice (Cho et al., 2020) selects several clients

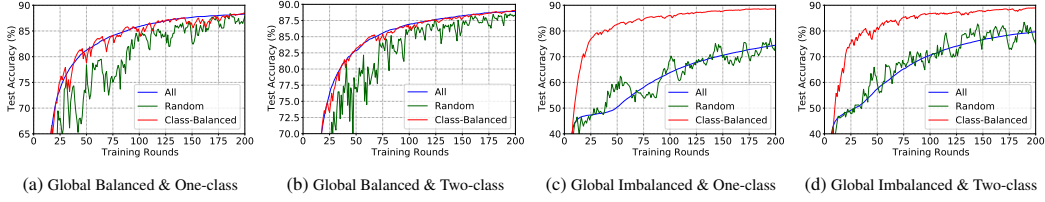


Figure 1: Three different FL client selection strategies on MNIST. In Figure (a) and (b), the global dataset of all the clients’ training data is class-balanced. In Figure (c) and (d), the global dataset is class-imbalanced. Each client has only one class of data in (a) and (c) and each client has two classes of data in (b) and (d). The results show significant performance degradation with imbalanced data from random client selection.

with the largest loss from a randomly sampled subset of all the available clients. However, selecting the clients with the larger local loss may not guarantee that the final model can have a smaller global loss. Another limitation of previous research on client selection is the comparison between their strategy and the ideal settings, where all the available clients participate in the training, is ignored. In general, existing works not only miss an important criterion that can measure their performance, but also fail to investigate the essential reason why random client selection leads to performance degradation on non-IID data compared with engaging all the available clients.

In this paper, we focus on the image classification tasks. First, we demonstrate our key observation for the essential reason why random client selection results in performance degradation on non-IID data, which is the *class-imbalance* of the grouped dataset from randomly selected clients. Based on our observation, we design an efficient heterogeneity-aware client sampling mechanism, *i.e.*, Federated Class-Balanced Sampling (Fed-CBS), which effectively reduces the class-imbalance in FL. Fed-CBS is orthogonal to numerous existing techniques to improve the performance of FL (Wang et al., 2020a;b; Chen et al., 2020; Karimireddy et al., 2019) on non-IID data, meaning Fed-CBS can be integrated with these methods to further improve their performance. Our major contributions are summarized as follows:

- We reveal that the class-imbalance is the essential reason why random client selection leads to performance degradation on non-IID data in Section 2.
- To effectively reduce the class-imbalance, we design an efficient heterogeneity-aware client sampling mechanism, *i.e.*, Fed-CBS, based on our proposed class-imbalance metric in Section 3. Besides, we provide theoretical analysis on its convergence in Section 4.
- We empirically evaluate Fed-CBS on FL benchmark non-IID datasets in Section 5. The results demonstrate that Fed-CBS can improve the accuracy of FL models on CIFAR-10 by 2% $\sim$ 7% and accelerate the convergence time by $1.3\times \sim 2.8\times$, compared with the state-of-the-art method (Yang et al., 2020) that also aims to reduce class-imbalance via client selection. Furthermore, our Fed-CBS achieves comparable or even better performance than the ideal setting where all the available devices are involved in the training.

2 PRELIMINARY AND RELATED WORK

2.1 PITFALL OF CLASS-IMBALANCE IN CLIENT SELECTION

We first clarify three definitions. **Local dataset** is the client’s own dataset, which is not accessible to other clients and the server. Due to the heterogeneity of local data distribution, the phenomenon of class-imbalance happens frequently in most of the clients’ local dataset. **Global dataset** is the union of all the available client local datasets in FL, which can be class-balanced or class-imbalanced. **Grouped dataset** is the union of several clients’ local datasets, which have been selected to participate in the training in one communication round of FL. Grouped dataset is a subset of the global dataset.

Some recent works (Duan et al., 2019; Wang et al., 2020b; Yang et al., 2020) have identified the issue of class-imbalance in the grouped dataset by random selection under non-IID settings. Since class-imbalance degrades the classification accuracy on minority classes (Huang et al., 2016) and leads to low training efficiency, this motivates us to verify whether the class-imbalance of the grouped dataset, caused by random selection, is the essential reason accounting for the performance degradation. We conduct experiments on MNIST (LeCun & Cortes, 2010) to verify our presumption¹. As Figure 1a

¹Detailed experiment settings are listed in the Appendix D.1

	Exact Local Label Distribution	Estimated Local Label Distribution
#1	[0.00, 0.00, 0.00, 0.28 , 0.00, 0.00, 0.36 , 0.00, 0.00, 0.36]	[0.02, 0.01, 0.02, 0.13 , 0.02, 0.02, 0.20 , 0.02, 0.02, 0.54]
#2	[0.00, 0.00, 0.22 , 0.16 , 0.00, 0.00, 0.20 , 0.00, 0.22 , 0.20]	[0.05, 0.03, 0.14 , 0.10 , 0.04, 0.03, 0.17 , 0.03, 0.16 , 0.25]
#3	[0.00, 0.00, 0.00, 1.00 , 0.00, 0.00, 0.00, 0.00, 0.00, 0.00]	[0.00, 0.00, 0.00, 1.00 , 0.00, 0.00, 0.00, 0.00, 0.00, 0.00]
#4	[0.00, 0.00, 0.00, 0.20 , 0.00, 0.00, 0.27 , 0.00, 0.28 , 0.25]	[0.03, 0.02, 0.03, 0.16 , 0.03, 0.02, 0.30 , 0.02, 0.20 , 0.19]
#5	[0.00, 0.00, 0.00, 0.29 , 0.00, 0.00, 0.36 , 0.00, 0.00, 0.35]	[0.02, 0.01, 0.02, 0.14 , 0.02, 0.01, 0.15 , 0.01, 0.02, 0.60]

Table 1: The exact local label distributions and the estimated ones of the first 5 clients in the experiment of Yang et al. (2020).

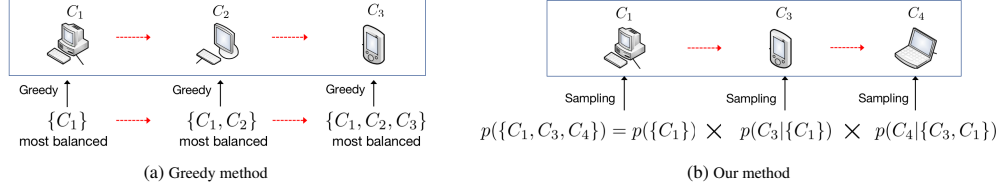


Figure 2: An example demonstrating the weakness of greedy method to deal with class imbalance. Supposing we work on a 6-class classification task and aim to select 3 clients from 4 available clients C_1, C_2, C_3, C_4 . Each of them has 30 images. The compositions of their local datasets are $[5, 5, 5, 5, 5, 5]$, $[6, 6, 6, 6, 6, 0]$, $[0, 0, 0, 10, 10, 0]$ and $[10, 10, 10, 0, 0, 0]$ respectively. The greedy method in Yang et al. (2020) is deterministic. It can only derive one result $\{C_1, C_2, C_3\}$ instead of the optimal solution $\{C_1, C_3, C_4\}$ (see the text description). But our method is based on probability modeling, which directly models the distribution of the optimal solution $\{C_1, C_3, C_4\}$. Thus when sampling from it, the optimal solution can be returned with high probability.

and Figure 1b show, the random client selection gets the worst performance when the global label distribution is class-balanced. If we keep the grouped dataset class-balanced by manually selecting the clients based on their local label distributions, we can achieve the accuracy comparable to the case of fully engaging all the clients in training.

Another corollary which comes naturally is that when the global dataset is inherently class-imbalanced, engaging all clients in the training may lead to worse performance than manually keeping class balanced in the grouped dataset. The results in Figure 1c and Figure 1d validate our hypothesis and verify the importance of class-imbalance reduction. This also indicates that only keeping diversity in the data and fairness for clients is not enough, which was ignored in the previous literature (Balakrishnan et al., 2021; Huang et al., 2021; Wang et al., 2020b; Yang et al., 2020).

2.2 RELATED WORK

Some effort has been made to improve client selection for FL in previous literature. Cho et al. (2020); Goetz et al. (2019) select clients with larger local loss, which cannot guarantee that the final global model has a smaller global loss. Focusing on the diversity in client selection, the authors of Balakrishnan et al. (2021) select clients by maximizing a submodular facility location function defined over gradient space. A fairness guaranteed algorithm termed RBCS-F was proposed in Huang et al. (2021), which models the fairness guaranteed client selection as a Lyapunov optimization problem. Although diversity and fairness are important, the experimental results in Section 2.1 demonstrate that they are not enough for client selection if class-imbalance issue is not considered. The authors in Ribero & Vikalo (2020) model the progression of model weights by an Ornstein-Uhlenbeck process and design a sampling strategy for selecting clients with significant weight updates. However, the work only considers identical data distribution. Following the existing works (Cho et al., 2020; Goetz et al., 2019), we only focus on the data heterogeneity caused by non-IID data across clients. The system heterogeneity considered in (Lai et al., 2021; Nishio & Yonetani, 2019; Reisizadeh et al., 2020) is beyond the scope of this paper and we leave it as future work.

To the best of our knowledge, Duan et al. (2019) and Yang et al. (2020) are the first two attempts to improve client selection by reducing the class-imbalance. An extra virtual component called mediator is introduced in *Astraea* of Duan et al. (2019), which has access to the local label distributions of the clients. With these distributions, *Astraea* will conduct client selection in a greedy way. The method of Yang et al. (2020) first estimates the local label distribution of each client based on the gradient of model parameters and adopts the same greedy way to select clients as *Astraea*. Since directly knowing the exact value of local label distributions of clients in *Astraea* will cause severe concerns on privacy leakage, we consider the method in Yang et al. (2020) as the state-of-the-art work aiming to improve client selection through class-imbalance reduction.

The weakness of the attempt in Yang et al. (2020) is threefold. First, their method requires a class-balanced auxiliary dataset that consists of all classes of data at server. However, that is not always available in some large scale FL systems since it requires the server to collect raw data from clients, which breaches privacy. Second, their estimations of the clients' local label distribution are not accurate as shown in Table 1. Theorem 1 in Yang et al. (2020), which supports their estimations, can not be generalized to multi-class classification tasks since it has only been proved in the original paper (Anand et al., 1993) for two-class classification problems. Finally, the performance of greedily conducting the client selection is not guaranteed due to the nature of greedy algorithm. We provide an example in Figure 2 to show its weakness. Their method will select C_1 as the first client since it is the most class-balanced one. Then C_2 will be selected because the grouped dataset of $C_1 \cup C_2$ is the most class-balanced among the ones of $C_1 \cup C_2$, $C_1 \cup C_3$ and $C_1 \cup C_4$. Similarly, it will choose C_3 since the grouped dataset of $C_1 \cup C_2 \cup C_3$ is more class-balanced than $C_1 \cup C_2 \cup C_4$. Their method is deterministic and thus only one combination $\{C_1, C_2, C_3\}$ is obtained. However, this is clearly not the optimal solution since $\{C_1, C_3, C_4\}$ is more class-balanced than $\{C_1, C_2, C_3\}$. The above weaknesses motivate us to design a more efficient sampling method to deal with this task.

3 METHODOLOGY

We first propose a metric to measure class-imbalance in Section 3.1. Then we present a method to calculate it in a privacy-preserving way in Section 3.2. Based on this measurement, we then design our client sampling mechanism and show its superiority in Section 3.3.

3.1 CLASS-IMBALANCE MEASUREMENT

Assume there are B classes of data in an image classification task, where $B \geq 2$. In the k -th communication round, we assume there are N_k available clients and we select M clients from them. To make the notation concise, we ignore the index "k" and assume the set of indices for the available clients is $\{1, 2, 3, \dots, N\}$ and the n -th available client has its own training dataset $\mathcal{D}_n$. We adopt the following vector, $\alpha_n = [\alpha_{(n,1)}, \alpha_{(n,2)}, \dots, \alpha_{(n,b)}, \dots, \alpha_{(n,B)}]$ to represent the local label distribution of $\mathcal{D}_n$, where $\alpha_{(n,b)} \geq 0$ and $\sum_{b=1}^B \alpha_{(n,b)} = 1$. We aim to find a subset $\mathcal{M}$ of $\{1, 2, 3, \dots, N\}$ of size M , such that the grouped dataset $\mathcal{D}_{\mathcal{M}}^g = \bigcup_{n \in \mathcal{M}} \mathcal{D}_n$ is class-balanced. Assuming the n -th client's local dataset has q_n training data in total, one can derive the following vector $\alpha_{\mathcal{M}}^g$ which represents the label distribution of the grouped dataset $\mathcal{D}_{\mathcal{M}}^g$,

$$\alpha_{\mathcal{M}}^g = \frac{\sum_{n \in \mathcal{M}} q_n \alpha_n}{\sum_{n \in \mathcal{M}} q_n} = \left[\frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,1)}}{\sum_{n \in \mathcal{M}} q_n}, \dots, \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n}, \dots, \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,B)}}{\sum_{n \in \mathcal{M}} q_n} \right].$$

We propose the following function to measure the magnitude of class-imbalance of $\mathcal{M}$, which we call *Quadratic Class-Imbalance Degree (QCID)*:

$$QCID(\mathcal{M}) \triangleq \sum_{b=1}^B \left(\frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} - \frac{1}{B} \right)^2.$$

3.2 PRIVACY-PRESERVING QCID DERIVATION

Expanding the the expression of *QCID* allows us to explore how the pairwise relationships of the clients' local label distributions $\{\alpha_m\}$ affects the class-imbalance degree of $\mathcal{M}$, where $m \in \mathcal{M}$. In the meantime, we must conduct this while keeping clients' local distributions hidden from the server for privacy reasons. Based on Theorem 3.1, there are several ways to achieve the privacy goal.

Theorem 3.1 *The QCID value is decided by the sum of inner products between each two vectors $\alpha_m, \alpha_{m'} \in \{\alpha_m\}$ with $m \in \mathcal{M}$, i.e.,*

$$QCID(\mathcal{M}) = \frac{\sum_{n \in \mathcal{M}, n' \in \mathcal{M}} q_n q_{n'} \alpha_n \alpha_{n'}^T}{(\sum_{n \in \mathcal{M}} q_n)^2} - \frac{1}{B}$$

Theorem 3.1 reveals the fact that there is no need to know the local label distribution of each client to calculate the *QCID*, as long as we have access to the inner products between each other. To derive the *QCID* for any subset $\mathcal{M} \subseteq \{1, 2, 3, \dots, N\}$, we only need to know the $N \times N$ matrix S with element $s_{n,n'}$ being $\alpha_n \alpha_{n'}^T$, which is the inner product between the label distributions of

clients n and n' . For our privacy goal, Fully Homomorphic Encryption (FHE) (Brakerski et al., 2014; Chen et al., 2017; Fan & Vercauteren, 2012) is one option since it enables an untrusted party to perform computation (addition and multiplication) on encrypted data. Alternatively, server-side trusted execution environments (TEEs), e.g., Intel SGX (Anati et al., 2013), can be used.

In Figure 3, we provide an example of the framework to illustrate how to derive S without knowing the local label distributions $\{\alpha_i\}$ using FHE. Our framework can be implemented using off-the-shelf FHE libraries such as Chen et al. (2017); Halevi & Shoup (2014). There is a selector in our framework. It is usually from third party and keeps a unique private key, denoted by K_1^{-1} . The corresponding public key is denoted as K_1 . In the confidential transmission between server and clients, each client first uses K_1 to encrypt their label distribution vector α_k as $K_1(\alpha_k)$, and then transmits it to server. Since only the server has access to $K_1(\alpha_k)$, no one else including the selector can decrypt α_k . When the server gets all $K_1(\alpha_k)$, it will conduct FHE computation to get the matrix $K_1(S) = K_1(\{\alpha_i^T \alpha_j\}_{ij}) = \{K_1(\alpha_i)^T K_1(\alpha_j)\}_{ij}$. Then the server transmits the $K_1(S)$ to selector where K_1^{-1} is used to access the final result S . Since only the selector has K_1^{-1} , only it knows S . After that, the selector will conduct client selection following some strategy to derive the result $\mathcal{M}$ and transmit it back to the server. At last, the server will conduct FL aggregation.

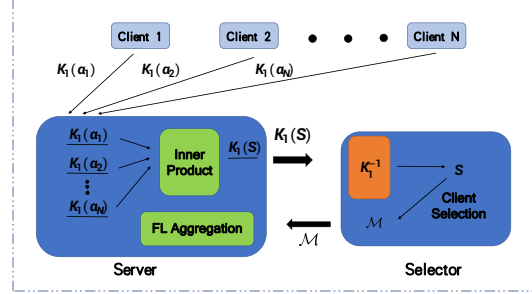


Figure 3: An example of FHE to securely transmit S .

The server, selector or any other clients except client n can not get α_n , which protects the privacy of the clients. Furthermore, only the clients have no access to the inner product results S , which prevents malicious clients or server inferring the label distributions of the other clients. We also prove that it is impossible even for the selector to derive $\{\alpha_i\}$ from S in Section B.

3.3 A CLIENT SAMPLING MECHANISM

To select the most class-balanced grouped dataset $\mathcal{D}_{\mathcal{M}}^g$, we need to find the optimal subset $\mathcal{M}^*$ that has the lowest $QCID$ value, which is defined as follows:

$$\mathcal{M}^* \triangleq \arg \min_{\mathcal{M} \subseteq \{1,2,3,\dots,N\}} \frac{\sum_{n \in \mathcal{M}, n' \in \mathcal{M}} q_n q_{n'} \alpha_n \alpha_{n'}^T}{(\sum_{n \in \mathcal{M}} q_n)^2} - \frac{1}{B}.$$

The main challenge is the computational complexity. To find the exact optimal $\mathcal{M}^*$, we need to loop through all the possible cases and find the lowest $QCID$ value. The computational complexity thereafter will be $\mathcal{O}\left(\binom{N}{M} \times M^2\right)$, which is unacceptable when N is extremely large.

A probability approach To overcome the computational bottleneck, instead of treating $\mathcal{M}$ as a determined set, we consider it as a sequence of random variables, i.e. $\mathcal{M} = \{C_1, C_2, \dots, C_m, \dots, C_M\}$ and assign it with some probability. Our expectation is that $\mathcal{M}$ should have higher probability to be sampled with if it is more class-balanced. This means $P(C_1 = c_1, C_2 = c_2, \dots, C_m = c_m, \dots, C_M = c_M)$ should be larger if $\mathcal{M} = \{c_1, c_2, \dots, c_M\}$ has a lower $QCID$ value.

Our sampling strategy generates the elements in $\mathcal{M}$ in a sequential manner, i.e., we first sample $\mathcal{M}_1 = \{c_1\}$ according to the probability of $P(C_1 = c_1)$, then sample c_2 to form $\mathcal{M}_2 = \{c_1, c_2\}$ according to the conditional probability $P(C_2 = c_2 | C_1 = c_1)$. The same procedure applies for the following clients until we finally obtain $\mathcal{M} = \{c_1, c_2, \dots, c_M\}$. In the following, we will design proper conditional probabilities such that the joint distribution of client selection satisfies our expectation.

Let T_n denote the number of times that client n has been selected. Once client n has been selected in a communication round, $T_n \rightarrow T_n + 1$, otherwise, $T_n \rightarrow T_n$. Inspired by combinatorial upper confidence bounds (CUCB) algorithm (Chen et al., 2013) and previous work in Yang et al. (2020), in the k -th communication round, the first element is designed to be sampled with the following probability:

$$P(C_1 = c_1) \propto \frac{1}{[QCID(\mathcal{M}_1)]^{\beta_1}} + \lambda \sqrt{\frac{3 \ln k}{2T_{c_1}}}, \quad \beta_1 > 0,$$

where λ above is the exploration factor to balance the trade-off between exploitation and exploration. The second term will add higher probability to the clients that have never been sampled before in the following communication rounds. After C_1 , the second client C_2 will be sampled with probability

$$P(C_2 = c_2 | C_1 = c_1) \propto \frac{\frac{1}{[QCID(\mathcal{M}_2)]^{\beta_2}}}{\frac{1}{[QCID(\mathcal{M}_1)]^{\beta_1}} + \alpha \sqrt{\frac{3 \ln k}{2T_{c_1}}}}, \quad \beta_2 > 0.$$

For the m -th client, where $2 < m \leq M$, we define

$$P(C_m = c_m | C_{m-1} = c_{m-1}, \dots, C_2 = c_2, C_1 = c_1) \propto \frac{[QCID(\mathcal{M}_{m-1})]^{\beta_{m-1}}}{[QCID(\mathcal{M}_m)]^{\beta_m}}, \quad \beta_{m-1}, \beta_m > 0.$$

With the above sampling process, the final probability to sample $\mathcal{M}$ is $P(C_1 = c_1, C_2 = c_2, \dots, C_M = c_M) = P(C_1 = c_1) \times P(C_2 = c_2 | C_1 = c_1) \cdots \times P(C_M = c_M | C_{M-1} = c_{M-1}, \dots, C_2 = c_2, C_1 = c_1) \propto 1/[QCID(\mathcal{M})]^{\beta_M}$. Since $\beta_M > 0$, this matches our goal that the $\mathcal{M}$ with lower $QCID$ value should have higher probability to be sampled with. Our mechanism, Fed-CBS, is summarized in the Algorithm 1 in Appendix (Section A).

Details and analysis For any $1 < m < M$, we have

$$P(C_1 = c_1, C_2 = c_2, \dots, C_m = c_m) \propto \frac{1}{[QCID(\mathcal{M}_m)]^{\beta_m}}.$$

This means when we generate the first m elements of $\mathcal{M}$, we expect the $\mathcal{M}_m$ could be more class-balanced since the $\mathcal{M}_m$ with lower $QCID$ value has higher probability to be sampled. This is different from the greedy algorithm in Yang et al. (2020), which chooses the c_m from $\{1, 2, \dots, N\} \setminus \mathcal{M}_{m-1}$ that can make the $\mathcal{M}_m$ the most class-balanced one. Instead, our method can generate the global optimal set of clients in the sense of probability. An example has been provided in Figure 2 to show our superiority. After selecting the first two clients, $\{C_1, C_3\}$ of our method is less class-balanced than $\{C_1, C_2\}$ of greedy method. However, after making the last choice, our method has the chance to derive a perfectly class-balanced set $\{C_1, C_3, C_4\}$. In contrast, greedy method can only get one result $\{C_1, C_2, C_3\}$, which is less class-balanced.

We require the distribution of $P(C_1 = c_1, C_2 = c_2, \dots, C_m = c_m)$ to be more dispersed when m is small. This is because we expect our sampling strategy to explore more possible cases of client composition at the beginning. We require the distribution of $P(C_1 = c_1, C_2 = c_2, \dots, C_m = c_m)$ to be less dispersed when m is large. This is because as we approach the end of our sampling process, we expect our sampling strategy can find the $\mathcal{M}_m$ that is more class-balanced. We can set $0 < \beta_1 < \beta_2 < \dots < \beta_M$ to satisfy the above requirements.

Remark: We set a lower bound for $QCID(\mathcal{M}_m)$ as L_b since $QCID(\mathcal{M}_m) = 0$ in some special cases will cause $P(C_m = c_m | C_{m-1} = c_{m-1}, \dots, C_1 = c_1) \rightarrow \infty$.

Below we present two theorems to show the superiority of our proposed sampling strategy.

Theorem 3.2 (Class-Imbalance Reduction) We denote the probability of selecting $\mathcal{M}$ with our strategy with β_M as P_{β_M} and probability of selecting $\mathcal{M}$ with the random selection as P_{rand} . Our method can reduce the expectation of $QCID$ compared to the randomly selection mechanism. In other words, we have

$$\mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M}) < \mathbb{E}_{\mathcal{M} \sim P_{rand}} QCID(\mathcal{M}).$$

Furthermore, if increasing the value β_M , the expectation of $QCID$ can be further reduced, i.e., for $\beta'_M > \beta_M$, we have

$$\mathbb{E}_{\mathcal{M} \sim P_{\beta'_M}} QCID(\mathcal{M}) < \mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M}).$$

Theorem 3.3 (Computation Complexity Reduction) The computation complexity of our method is $\mathcal{O}(N \times M^2)$, which is much smaller than the exhaustive search of $\mathcal{O}\left(\binom{N}{M} \times M^2\right)$.

Theorem 3.3 shows that the computation complexity of our method is independent of the number of classes. Since the dimension of neural networks is typically much larger than the class distribution vector α_n , the additional communication cost is almost negligible.

4 CONVERGENCE ANALYSIS

To analyze the convergence of our method, as well as quantify the relationship between the class-imbalance and the convergence, we first define our objective functions and adopt some general assumptions. Our global objective function $\tilde{F} > 0$ can be decomposed as $\tilde{F} = \frac{1}{B} \sum_{b=1}^B \tilde{F}_b$, where $\tilde{F}_b$ is the averaged loss function with respect to all the data of the b -th class in global dataset. Similarly, the n -th client's local objective function F_n can be decomposed as $F_n = \sum_{b=1}^B \alpha_{(n,b)} F_{n,b}$, where $F_{n,b}$ is the averaged loss function with respect to all the data of the b -th class in the n -th client's local dataset and $\alpha_{(n,b)}$ is defined in Equation 3.1. Besides, let $\mathbf{w}^{(k)}$ denote the global model parameter at the k -th communication round and $\mathbf{w}^{(0)}$ denote the initial global model parameter.

Assumption 4.1 (Smoothness) *The global objective function $\tilde{F}$ and each client's averaged loss function $F_{n,b}$ are Lipschitz smooth, i.e. $\|\nabla \tilde{F}(\mathbf{w}) - \nabla \tilde{F}(\mathbf{w}')\| \leq L_{\tilde{F}} \|\mathbf{w} - \mathbf{w}'\|$ and $\|\nabla F_{n,b}(\mathbf{w}) - \nabla F_{n,b}(\mathbf{w}')\| \leq L_{n,b} \|\mathbf{w} - \mathbf{w}'\|, \forall n, b, \mathbf{w}, \mathbf{w}'$.*

Assumption 4.2 (Unbiased Gradient and Bounded Variance) *The stochastic gradient g_n at each client is an unbiased estimator of the local gradient: $\mathbb{E}_{\xi} [g_n(\mathbf{w} | \xi)] = \nabla F_n(\mathbf{w})$, with bounded variance $\mathbb{E}_{\xi} [\|g_n(\mathbf{w} | \xi) - \nabla F_n(\mathbf{w})\|^2] \leq \sigma^2, \forall \mathbf{w}$, where $\sigma^2 \geq 0$.*

Assumption 4.3 (Bounded Dissimilarity) *There exists two non-negative constants $\delta \geq 1, \gamma^2 \geq 0$ such that $\sum_{b=1}^B \frac{1}{B} \|\nabla \tilde{F}_b(\mathbf{w})\|^2 \leq \delta \left\| \sum_{b=1}^B \frac{1}{B} \nabla \tilde{F}_b(\mathbf{w}) \right\|^2 + \gamma^2, \forall \mathbf{w}$.*

Assumption 4.4 (Class-wise Similarity) *For each class b , the discrepancy between the gradient of global averaged loss function and the local one has been bounded by some constant in l^2 norm. That means, for every n and b , we have $\|\nabla \tilde{F}_b(\mathbf{w}) - \nabla F_{n,b}(\mathbf{w})\|^2 \leq \kappa_{n,b}^2, \forall \mathbf{w}$.*

Assumption 4.1, 4.2 and 4.3 have been widely adopted in the previous literature on the theoretical analysis of FL such as Cho et al. (2020); Li et al. (2019); Wang et al. (2020a). Assumption 4.4 is based on the similarity among the data from the same class. Similar to the standard setting (Wang et al., 2020a), the convergence is measured by the norm of the gradients, stated in Theorem 4.5.

Theorem 4.5 *Under Assumption 4.1 to 4.4, if the total communication rounds K is pre-determined and the learning rate is set as $\eta = \frac{s}{10L\sqrt{\tau(\tau-1)K}}$, where $s < 1$, $L = \max_{\{n,b\}} L_{n,b}$ and τ is the number of local update iterations, the minimal gradient norm of $\tilde{F}$ is bounded as:*

$$\min_{k \leq K} \mathbb{E} \left\| \nabla \tilde{F}(\mathbf{w}^{(k)}) \right\|^2 \leq \frac{1}{V} \left[\frac{\sigma^2 s^2}{25\tau K} + \frac{sL_{\tilde{F}}\sigma^2}{10L\sqrt{\tau(\tau-1)K}} + \frac{10L\sqrt{\tau(\tau-1)}\tilde{F}(\mathbf{w}^{(0)})}{s\sqrt{K}} + \rho \right],$$

where $V = \frac{1}{3} - \delta B \mathbb{E}[QCID]$, $\rho = 5\kappa^2 + \gamma^2 \mathbb{E}[QCID]$ and $\kappa = \max_{\{n,b\}} \kappa_{n,b}$.

If the class-imbalance in client selection is reduced, $\mathbb{E}[QCID]$ will decrease, and consequently, $\frac{1}{V}$ and $\frac{\mathbb{E}[QCID]}{V}$ will also decrease, making the convergence bound on the right side tighter. Theorem 4.5 not only provides a convergence guarantee for Fed-CBS, but also proves the class-imbalance reduction in client selection could benefit FL, i.e., more class-balance leads to faster convergence.

5 EXPERIMENTS

We conduct thorough experiments on three public benchmark datasets, CIFAR-10 (Krizhevsky et al.), Fashion-MNIST (Xiao et al., 2017) and FEMNIST in the Leaf Benchmark (Caldas et al., 2018). In all the experiments, we conduct the simulation of the cross-device federated learning (CDFL), where the system runs on a large amount of client with only a fraction of them available in each communication round and we make client selection on those available clients. The results show that our method can achieve faster and more stable convergence, compared with four baselines: random selection (rand), Power-of-choice Selection Strategy (pow-d) (Cho et al., 2020), the method in Yang et al. (2020) (Fed-cub), the ideal setting where we select all the available clients (all). To compare them efficiently, we present the results of the Cifar-10 where the whole dataset is divided to 200 (or 120) clients in the main text, since we need to engage all the clients for the ideal setting. To simulate

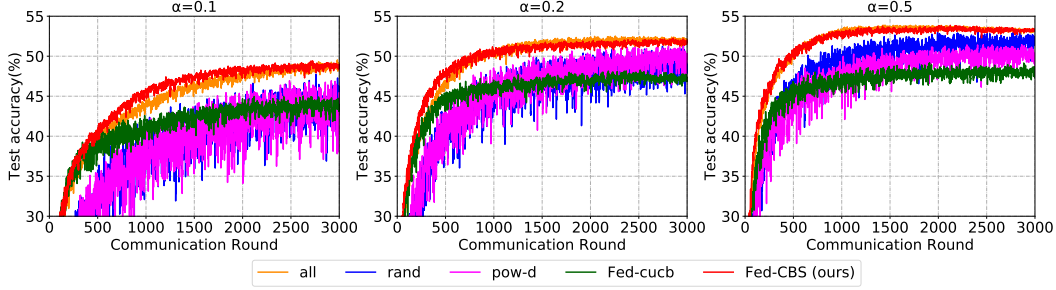


Figure 4: Test accuracy on Cifar-10 under three heterogeneous settings.

		all	rand	pow-d	Fed-cucb	Fed-CBS
Communication Rounds	$\alpha=0.1$	757 ± 155	951 ± 202	1147 ± 130	861 ± 328	654 ± 96
	$\alpha=0.2$	746 ± 95	762 ± 105	741 ± 111	803 ± 220	475 ± 110
	$\alpha=0.5$	426 ± 67	537 ± 115	579 ± 140	1080 ± 309	384 ± 74
$\mathbb{E}[QCID](10^{-2})$	$\alpha=0.1$	1.01 ± 0.01	8.20 ± 0.21	12.36 ± 0.26	7.09 ± 2.27	0.62 ± 0.20
	$\alpha=0.2$	0.93 ± 0.03	7.54 ± 0.27	10.6 ± 0.48	5.93 ± 1.01	0.51 ± 0.12
	$\alpha=0.5$	0.72 ± 0.03	5.87 ± 0.24	7.36 ± 0.57	6.47 ± 0.77	0.36 ± 0.04

Table 2: The communication rounds required for targeted test accuracy and the averaged QCID values. The targeted test accuracy is 45% for $\alpha = 0.1$, 47% for $\alpha = 0.2$ and 50% for $\alpha = 0.5$. The results are the mean and the standard deviation over 4 different random seeds.

more realistic settings where there are thousands of clients, we conduct our method on the FEMNIST in the Leaf Benchmark with more than 3000 clients. Since in practice it is impossible to engage all the clients in the training, we compare our method with randomly selecting more clients. Due to space limit, we move the results of FEMNIST, Fashion-MNIST and ablation studies to Section D.5 & E in Appendix. For Fashion-MNIST, we adopt FedNova (Wang et al., 2020a) to show that our method can be organically integrated to the existing orthogonal works which aim at improving FL.

Experiment Setup We target cross-device settings where the devices are resource-constrained, i.e., most of the devices do not have sufficient computational power and memory to support the training of large models. Therefore, we adopt a compact model with two convolutional layers followed with three fully-connected layers and FedAvg (McMahan et al., 2017b) as the FL optimizer. The batch size is 50 for each client. In each communication round, all of them conduct the same number of local updates, which allows the client with the largest local dataset to conduct 5 local training epochs. In our method, we set the $\beta_m = m$, $\gamma = 10$ and $L_b = 10^{-20}$. The local optimizer is SGD with a weight decay of 0.0005. The learning rate is 0.01 initially and the decay factor is 0.9992. We terminate the FL training after 3000 communication rounds and then evaluate the model’s performance on the test dataset of CIFAR-10. More details of the experiment setup are listed in Section D.2.

5.1 RESULTS FOR CLASS-BALANCED GLOBAL DATASETS

In this experiment, we set 200 clients in total with a class-balanced global dataset. The non-IID data partition among clients is based on Dirichlet distribution parameterized by the concentration parameter α in Hsu et al. (2019). Roughly speaking, as α decreases, the data distribution will become more non-iid. In each communication round, we uniformly and randomly set 30% of them (i.e., 60 clients) available and select 10 clients from those 60 available ones to participate in the training.

As shown in Table 5, our method can achieve the lowest *QCID* value compared with other client selection strategies. As a benefit of successfully reducing the class-imbalance, our method outperforms the other three baseline methods and achieves comparable performance to the ideal setting where all the available clients are engaged into the training. As shown in Table 5 and Figure 4, our method can achieve faster and more stable convergence. It is worth noting that due to the inaccurate estimation of distribution and the weakness of greedy method discussed in Section 2.2, the performance of Fed-cucb is much worse than ours.

5.2 RESULTS FOR CLASS-IMBALANCED GLOBAL DATASETS

In the real-world settings, the global dataset of all the clients is not always class-balanced. Hence, we investigate two different cases to show the superiority of our method and provide more details of their

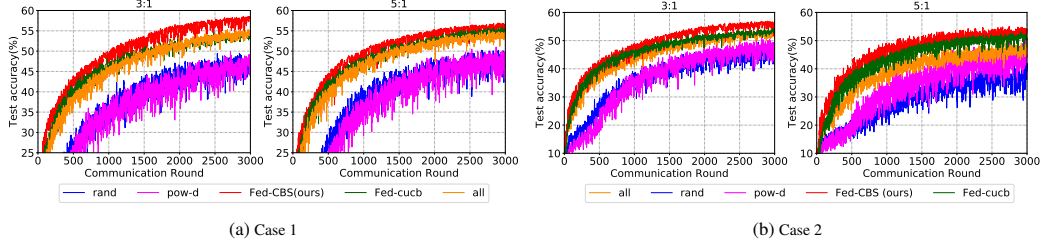


Figure 5: Test accuracy on Cifar-10 with class-imbalanced global dataset in Case 1&2.

		all	rand	pow-d	Fed-cub	Fed-CBS
Case 1	3:1	55.17 $\pm$ 0.94	50.99 $\pm$ 0.97	53.51 $\pm$ 0.34	55.11 $\pm$ 0.26	56.86$\pm$0.34
	5:1	50.93 $\pm$ 1.64	47.36 $\pm$ 2.34	52.73 $\pm$ 1.85	53.75 $\pm$ 0.58	54.94$\pm$0.73
Case 2	3:1	54.01 $\pm$ 0.60	50.81 $\pm$ 2.03	53.98 $\pm$ 1.87	54.48 $\pm$ 1.31	57.71$\pm$0.50
	5:1	50.42 $\pm$ 1.27	48.33 $\pm$ 3.03	53.54 $\pm$ 1.18	53.38 $\pm$ 1.48	57.99$\pm$0.46

Table 3: Best test accuracy for our method and other four baselines.

settings in Section D.3. For the convenience of constructing a class-imbalanced global dataset, each client only has one class of data with the same quantity. We report the best test accuracy in Table 3 and present the corresponding *QCID* values in Section D.4.

5.2.1 CASE 1: UNIFORM AVAILABILITY

Settings. There are 120 clients in total and the global dataset of these 120 clients is class-imbalanced. To measure the degree of class imbalance, we let the global dataset have the same amount of n_1 data samples for five classes and the same amount of n_2 data samples for the other five classes. The ratio r between n_1 and n_2 is respectively set to 3 : 1 and 5 : 1 in the experiments. In each communication round, we uniformly set 30% of them (*i.e.*, 36 clients) available with replacement and select 10 clients to participate in the training.

As shown in Table 3 and Figure 5a, our method can achieve faster and more stable convergence, and even slightly better performance than the ideal setting where all the available clients are engaged. The performance of Fed-cub (Yang et al., 2020) is better than the results on class-balanced global dataset, which is partly due to the simplicity of each client’s local dataset composition in our experiments. The third line in Figure 1 indicates Fed-cub can accurately estimate this simple type of label distributions.

5.2.2 CASE 2: NON-UNIFORM AVAILABILITY

Settings. There are 200 clients in total. In each communication round, 30% of them (*i.e.*, 60 clients) are set available uniformly in each training round with replacement. By non-uniformly setting the availability, the global dataset of those 60 available clients is always class-imbalanced. To measure the degree of class imbalance, we make the global dataset have the same amount of n_1 data samples for the five classes and have the same amount of n_2 data samples for the other five classes. The ratio r between n_1 and n_2 is set to 3 : 1 and 5 : 1. We select 10 clients to participate in the training.

As shown in Table 3 and Figure 5a, our method consistently achieves higher test accuracy and more stable convergence, which also outperforms the ideal setting where all the available clients are engaged. Since the global dataset of the available 60 clients in each communication round is always class-imbalanced, engaging all of them is not the optimal selection strategy in terms of test accuracy.

6 CONCLUSION

We unveil the essential reason of performance degradation on non-IID data with random client selection strategy in FL training, *i.e.*, the class-imbalance. Motivated by this insight, we propose an efficient heterogeneity-aware client sampling mechanism named Fed-CBS. Extensive experiments demonstrate that Fed-CBS significantly outperforms the status quo approaches while yielding comparable or even better performance than the ideal setting where all the available clients participate in the training. We also provide the theoretical convergence guarantee for Fed-CBS. Our mechanism is very general and can be easily applied to various application domains (e.g., NLP and medical tasks). As for future work, we will try to incorporate the system heterogeneity into our framework.

REFERENCES

- Rangachari Anand, Kishan G Mehrotra, Chilukuri K Mohan, and Sanjay Ranka. An improved algorithm for neural network classification of imbalanced training sets. *IEEE Transactions on Neural Networks*, 4(6):962–969, 1993.
- Ittai Anati, Shay Gueron, Simon Johnson, and Vincent Scarlata. Innovative technology for cpu based attestation and sealing. In *Proceedings of the 2nd international workshop on hardware and architectural support for security and privacy*, volume 13, pp. 7. Citeseer, 2013.
- Ravikumar Balakrishnan, Tian Li, Tianyi Zhou, Nageen Himayat, Virginia Smith, and Jeff Bilmes. Diverse client selection for federated learning: Submodularity and convergence analysis. In *ICML 2021 International Workshop on Federated Learning for User Privacy and Data Confidentiality*, Virtual, July 2021.
- Zvika Brakerski, Craig Gentry, and Vinod Vaikuntanathan. (leveled) fully homomorphic encryption without bootstrapping. *ACM Transactions on Computation Theory (TOCT)*, 6(3):1–36, 2014.
- Sebastian Caldas, Peter Wu, Tian Li, Jakub Konečný, H Brendan McMahan, Virginia Smith, and Ameet Talwalkar. Leaf: A benchmark for federated settings. *arXiv preprint arXiv:1812.01097*, 2018. URL <https://arxiv.org/abs/1812.01097>.
- Hao Chen, Kim Laine, and Rachel Player. Simple encrypted arithmetic library-seal v2.1. In *International Conference on Financial Cryptography and Data Security*, pp. 3–18. Springer, 2017.
- Wei Chen, Yajun Wang, and Yang Yuan. Combinatorial multi-armed bandit: General framework and applications. In Sanjoy Dasgupta and David McAllester (eds.), *Proceedings of the 30th International Conference on Machine Learning*, volume 28 of *Proceedings of Machine Learning Research*, pp. 151–159, Atlanta, Georgia, USA, 17–19 Jun 2013. PMLR. URL <https://proceedings.mlr.press/v28/chen13a.html>.
- Wei Chen, Kartikeya Bhardwaj, and Radu Marculescu. Fedmax: mitigating activation divergence for accurate and communication-efficient federated learning. *arXiv preprint arXiv:2004.03657*, 2020.
- Yae Jee Cho, Jianyu Wang, and Gauri Joshi. Client selection in federated learning: Convergence analysis and power-of-choice selection strategies. *ArXiv*, abs/2010.01243, 2020.
- Moming Duan, Duo Liu, Xianzhang Chen, Yujuan Tan, Jinting Ren, Lei Qiao, and Liang Liang. Astraea: Self-balancing federated learning for improving classification accuracy of mobile deep learning applications. In *2019 IEEE 37th international conference on computer design (ICCD)*, pp. 246–254. IEEE, 2019.
- Junfeng Fan and Frederik Vercauteren. Somewhat practical fully homomorphic encryption. *IACR Cryptol. ePrint Arch.*, 2012:144, 2012.
- Jack Goetz, Kshitiz Malik, Duc Bui, Seungwhan Moon, Honglei Liu, and Anuj Kumar. Active federated learning. *arXiv preprint arXiv:1909.12641*, 2019.
- Shai Halevi and Victor Shoup. Algorithms in HElib. In *Annual Cryptology Conference*, pp. 554–571. Springer, 2014.
- Tzu-Ming Harry Hsu, Hang Qi, and Matthew Brown. Measuring the effects of non-identical data distribution for federated visual classification. *ArXiv*, abs/1909.06335, 2019.
- Chen Huang, Yining Li, Chen Change Loy, and Xiaoou Tang. Learning deep representation for imbalanced classification. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 5375–5384, 2016.
- T. Huang, W. Lin, W. Wu, L. He, K. Li, and A. Y. Zomaya. An efficiency-boosting client selection scheme for federated learning with fairness guarantee. *IEEE Transactions on Parallel Distributed Systems*, 32(07):1552–1564, jul 2021. ISSN 1558-2183. doi: 10.1109/TPDS.2020.3040887.

- 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, Badih Ghazi, Phillip B. Gibbons, Marco Gruteser, Zaid Harchaoui, Chaoyang He, Lie He, Zhouyuan Huo, Ben Hutchinson, Justin Hsu, Martin Jaggi, Tara Javidi, Gauri Joshi, Mikhail Khodak, Jakub Konečný, Aleksandra Korolova, Farinaz Koushanfar, Sanmi Koyejo, Tancrède Lepoint, Yang Liu, Prateek Mittal, Mehryar Mohri, Richard Nock, Ayfer Özgür, Rasmus Pagh, Mariana Raykova, Hang Qi, Daniel Ramage, Ramesh Raskar, Dawn Song, Weikang Song, Sebastian U. Stich, Ziteng Sun, Ananda Theertha Suresh, Florian Tramèr, Praneeth Vepakomma, Jianyu Wang, Li Xiong, Zheng Xu, Qiang Yang, Felix X. Yu, Han Yu, and Sen Zhao. Advances and open problems in federated learning, 2019. URL <https://arxiv.org/abs/1912.04977>.
- Sai Praneeth Karimireddy, Satyen Kale, Mehryar Mohri, Sashank J Reddi, Sebastian U Stich, and Ananda Theertha Suresh. Scaffold: Stochastic controlled averaging for on-device federated learning. 2019.
- Jakub Konečný, H Brendan McMahan, Daniel Ramage, and Peter Richtárik. Federated optimization: Distributed machine learning for on-device intelligence. *arXiv preprint arXiv:1610.02527*, 2016.
- Alex Krizhevsky, Vinod Nair, and Geoffrey Hinton. Cifar-10 (canadian institute for advanced research). URL <http://www.cs.toronto.edu/~kriz/cifar.html>.
- Fan Lai, Xiangfeng Zhu, Harsha V. Madhyastha, and Mosharaf Chowdhury. Oort: Efficient federated learning via guided participant selection. In *15th USENIX Symposium on Operating Systems Design and Implementation (OSDI 21)*, pp. 19–35. USENIX Association, July 2021. ISBN 978-1-939133-22-9. URL <https://www.usenix.org/conference/osdi21/presentation/lai>.
- Yann LeCun and Corinna Cortes. MNIST handwritten digit database. 2010. URL <http://yann.lecun.com/exdb/mnist/>.
- Xiang Li, Kaixuan Huang, Wenhao Yang, Shusen Wang, and Zhihua Zhang. On the convergence of fedavg on non-iid data. *arXiv preprint arXiv:1907.02189*, 2019.
- Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Aguera y Arcas. Communication-efficient learning of deep networks from decentralized data. In *Artificial intelligence and statistics*, pp. 1273–1282. PMLR, 2017a.
- H Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, et al. Communication-efficient Learning of Deep Networks from Decentralized Data. *Artificial Intelligence and Statistics*, 2017b.
- Takayuki Nishio and Ryo Yonetani. Client selection for federated learning with heterogeneous resources in mobile edge. *ICC 2019 - 2019 IEEE International Conference on Communications (ICC)*, pp. 1–7, 2019.
- Amirhossein Reisizadeh, Isidoros Tziotis, Hamed Hassani, Aryan Mokhtari, and Ramtin Pedarsani. Straggler-resilient federated learning: Leveraging the interplay between statistical accuracy and system heterogeneity. *arXiv preprint arXiv:2012.14453*, 2020.
- Monica Ribero and Haris Vikalo. Communication-efficient federated learning via optimal client sampling. *arXiv preprint arXiv:2007.15197*, 2020.
- H. Schneider and G.P. Barker. *Matrices and Linear Algebra*. Dover Books on Mathematics Series. Dover Publications, 1989. ISBN 9780486660141. URL <https://books.google.com/books?id=HJIT3CSb0wIC>.
- Jianyu Wang, Qinghua Liu, Hao Liang, Gauri Joshi, and H Vincent Poor. Tackling the objective inconsistency problem in heterogeneous federated optimization. *arXiv preprint arXiv:2007.07481*, 2020a.
- Lixu Wang, Shichao Xu, Xiao Wang, and Qi Zhu. Addressing class imbalance in federated learning. *arXiv preprint arXiv:2008.06217*, 2020b.

Han Xiao, Kashif Rasul, and Roland Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms, 2017. URL <http://arxiv.org/abs/1708.07747>. cite arxiv:1708.07747Comment: Dataset is freely available at <https://github.com/zalandoresearch/fashion-mnist> Benchmark is available at <http://fashion-mnist.s3-website.eu-central-1.amazonaws.com/>.

Miao Yang, Akitanoshou Wong, Hongbin Zhu, Haifeng Wang, and Hua Qian. Federated learning with class imbalance reduction, 2020.

A FED-CBS ALGORITHM

Algorithm 1 Fed-CBS

Initialization: initial local model $\mathbf{w}^{(0)}$, client index subset $\mathcal{M} = \emptyset$, K communication rounds,
 $k = 0, T_n = 1$
while $k < K$ **do**
 Client Selection:
 for n **in** $\{1, 2, \dots, N\}$ **do**
 if $n \in \mathcal{M}$ **then**
 $T_n \rightarrow T_n + 1$
 else
 $T_n \rightarrow T_n$.
 end if
 end for
 Update $\mathcal{M}$ using our proposed sampling strategy in Section 3.3
 Local Updates:
 for $n \in \mathcal{M}$ **do**
 $\mathbf{w}_n^{(k+1)} \leftarrow \text{Update}(\mathbf{w}^{(k)})$.
 end for
 Global Aggregation:
 $\mathbf{w}^{(k+1)} \leftarrow \text{Aggregate}(\mathbf{w}_n^{(k+1)})$ for $n \in \mathcal{M}$
end while

B PRIVACY PROTECTION IN THE FRAMEWORK IN FIGURE 3

Theorem B.1 *The selector can not derive the exact values of $\{\alpha_i\}$ from $\mathbf{S}$.*

Proof B.1 *By the definitions of $\{\alpha_i\}$, we define the following matrix $\mathbf{A}_\alpha$*

$$\mathbf{A}_\alpha \triangleq \begin{bmatrix} \alpha_1 \\ \dots \\ \alpha_n \\ \dots \\ \alpha_N \end{bmatrix} = \begin{bmatrix} \alpha_{(1,1)} & \alpha_{(1,2)} & \dots & \alpha_{(1,b)} & \dots & \alpha_{(1,B)} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \alpha_{(n,1)} & \alpha_{(n,2)} & \dots & \alpha_{(n,b)} & \dots & \alpha_{(n,B)} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \alpha_{(1,1)} & \alpha_{(1,2)} & \dots & \alpha_{(1,b)} & \dots & \alpha_{(1,B)} \end{bmatrix}$$

By the definitions of $\mathbf{S}$, we have

$$\mathbf{S} = \mathbf{A}_\alpha \cdot \mathbf{A}_\alpha^\top \quad (1)$$

To derive the exact values of $\{\alpha_i\}$ based on $\mathbf{S}$, we need to solve the problem 1. However, given $\mathbf{S}$, the $\mathbf{A}_\alpha$ which satisfies $\mathbf{S} = \mathbf{A}_\alpha \cdot \mathbf{A}_\alpha^\top$ is not unique. If $\bar{\mathbf{A}}_\alpha$ is a solution to the problem 1, then for any orthogonal matrix $\mathbf{Q}$ i.e. $\mathbf{Q} \cdot \mathbf{Q}^\top = \mathbf{I}$ where the $\mathbf{I}$ is the identity matrix, the new matrix $\bar{\mathbf{A}}_\alpha \cdot \mathbf{Q}$ is also solution to the problem 1. This is because

$$\bar{\mathbf{A}}_\alpha \cdot \mathbf{Q} \cdot (\bar{\mathbf{A}}_\alpha \cdot \mathbf{Q})^\top = \bar{\mathbf{A}}_\alpha \cdot \mathbf{Q} \cdot \mathbf{Q}^\top \cdot \bar{\mathbf{A}}_\alpha^\top = \bar{\mathbf{A}}_\alpha \cdot \bar{\mathbf{A}}_\alpha^\top = \mathbf{S}$$

Hence, the $\mathbf{A}_\alpha$ which satisfies $\mathbf{S} = \mathbf{A}_\alpha \cdot \mathbf{A}_\alpha^\top$ is not unique and we finish our proof.

To understand the Theorem B.1, we provide the following example. If we conduct the following permutation on the columns of $\mathbf{A}_\alpha$, we can derive a new matrix $\bar{\mathbf{A}}_\alpha$.

$$\bar{\mathbf{A}}_\alpha \triangleq \begin{bmatrix} \alpha_{(1,2)} & \dots & \alpha_{(1,b)} & \dots & \alpha_{(1,1)} & \alpha_{(1,B)} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \alpha_{(n,2)} & \dots & \alpha_{(n,b)} & \dots & \alpha_{(n,1)} & \alpha_{(n,B)} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \alpha_{(1,2)} & \dots & \alpha_{(1,b)} & \dots & \alpha_{(1,1)} & \alpha_{(1,B)} \end{bmatrix}$$

We can find that $\bar{A}_\alpha$ also satisfies $S = \bar{A}_\alpha \cdot \bar{A}_\alpha^\top$. Actually, there are also many other permutations which can derive the solutions to the problem 1. Hence, in our framework shown in 3, the selector can not estimate the exact label distribution of the clients.

C PROOF OF THEOREM 3.1, 3.2, 3.3 AND 4.5

C.1 PROOF OF THEOREM 3.1

Proof C.1

$$\begin{aligned}
QCID(\mathcal{M}) &= \sum_{b=1}^B \left(\frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} - \frac{1}{B} \right)^2 \\
&= \sum_{b=1}^B \left(\frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} \right)^2 + \frac{1}{B} - 2 * \frac{1}{\sum_{n \in \mathcal{M}} q_n} * \frac{1}{B} * \sum_{n \in \mathcal{M}} q_n \\
&= \sum_{b=1}^B \left(\frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} \right)^2 - \frac{1}{B} \\
&= \frac{1}{(\sum_{n \in \mathcal{M}} q_n)^2} \sum_{n \in \mathcal{M}, n' \in \mathcal{M}} q_n q_{n'} \left(\sum_{b=1}^B \alpha_{(n,b)} \alpha_{(n',b)} \right) - \frac{1}{B} \\
&= \frac{1}{(\sum_{n \in \mathcal{M}} q_n)^2} \sum_{n \in \mathcal{M}, n' \in \mathcal{M}} q_n q_{n'} \alpha_n \alpha_{n'}^T - \frac{1}{B}
\end{aligned}$$

C.2 PROOF OF THEOREM 3.2

Proof C.2 To select M clients from N available clients, there are $\binom{N}{M}$ different choices to construct $\mathcal{M}$, denoted by $\mathcal{M}^{(1)}, \mathcal{M}^{(2)}, \dots, \mathcal{M}^{(\binom{N}{M})}$, respectively. Let $x_i \triangleq QCID(\mathcal{M}^{(i)})$ and $\bar{N} \triangleq \binom{N}{M}$. Then we have

$$\begin{aligned}
\mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M}) &= x_1 \frac{\frac{1}{x_1^{\beta_M}}}{\frac{1}{x_1^{\beta_M}} + \frac{1}{x_2^{\beta_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M}}} + x_2 \frac{\frac{1}{x_2^{\beta_M}}}{\frac{1}{x_1^{\beta_M}} + \frac{1}{x_2^{\beta_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M}}} \\
&\quad + \dots + x_{\bar{N}} \frac{\frac{1}{x_{\bar{N}}^{\beta_M}}}{\frac{1}{x_1^{\beta_M}} + \frac{1}{x_2^{\beta_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M}}}
\end{aligned}$$

$$\text{And } \mathbb{E}_{\mathcal{M} \sim P_{rand}} QCID(\mathcal{M}) = \frac{1}{\bar{N}} (x_1 + x_2 + \dots + x_{\bar{N}})$$

Without loss of generality, we assume $x_1 \leq x_2 \leq \dots \leq x_{\bar{N}}$ and define the following y_i for the notation simplicity:

$$y_i = \begin{cases} \frac{1}{x_i^{\beta_M}} & \text{if } 0 \leq i \leq \bar{N} \\ \frac{1}{x_{i-\bar{N}}^{\beta_M}} & \text{if } \bar{N} < i \leq 2\bar{N} - 1 \end{cases} \quad (2)$$

Now we calculate the following ratio:

$$\begin{aligned}
\frac{\mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M})}{\mathbb{E}_{\mathcal{M} \sim P_{rand}} QCID(\mathcal{M})} &= \frac{\bar{N} (x_1 \frac{1}{x_1^{\beta_M}} + x_2 \frac{1}{x_2^{\beta_M}} + \dots + x_{\bar{N}} \frac{1}{x_{\bar{N}}^{\beta_M}})}{(x_1 + x_2 + \dots + x_{\bar{N}}) (\frac{1}{x_1^{\beta_M}} + \frac{1}{x_2^{\beta_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M}})} \\
&= \frac{\sum_{j=1}^{\bar{N}} (\sum_{i=1}^{\bar{N}} x_i \frac{1}{x_i^{\beta_M}})}{\sum_{j=1}^{\bar{N}} (\sum_{i=1}^{\bar{N}} x_i y_{j+i-1})}
\end{aligned}$$

Since we assume that $x_1 \leq x_2 \leq \dots \leq x_{\bar{N}}$, we have $\frac{1}{x_1^{\beta_M}} \geq \frac{1}{x_2^{\beta_M}} \geq \dots \geq \frac{1}{x_{\bar{N}}^{\beta_M}}$. Besides, it is easy to find x_i and $x_{i'}$ satisfying $x_i \neq x_{i'}$. Then for each $1 \leq j \leq \bar{N}$, according to the rearrangement inequality, we have

$$\begin{aligned} \sum_{i=1}^{\bar{N}} x_i \frac{1}{x_i^{\beta_M}} &< \sum_{i=1}^{\bar{N}} x_i y_{j+i-1} \Rightarrow \sum_{j=1}^{\bar{N}} \sum_{i=1}^{\bar{N}} x_i \frac{1}{x_i^{\beta_M}} < \sum_{j=1}^{\bar{N}} \sum_{i=1}^{\bar{N}} x_i y_{j+i-1} \Rightarrow \\ \frac{\mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M})}{\mathbb{E}_{\mathcal{M} \sim P_{rand}} QCID(\mathcal{M})} &< 1 \Rightarrow \mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M}) < \mathbb{E}_{\mathcal{M} \sim P_{rand}} QCID(\mathcal{M}) \end{aligned}$$

Similarly, for β'_M such that $\beta'_M \geq \beta_M$, denote $\beta'_M = \beta_M + \Delta\beta$. We have

$$\begin{aligned} \mathbb{E}_{\mathcal{M} \sim P_{\beta'_M}} QCID(\mathcal{M}) &= x_1 \frac{\frac{1}{x_1^{\beta'_M}}}{\frac{1}{x_1^{\beta'_M}} + \frac{1}{x_2^{\beta'_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta'_M}}} + x_2 \frac{\frac{1}{x_2^{\beta'_M}}}{\frac{1}{x_1^{\beta'_M}} + \frac{1}{x_2^{\beta'_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta'_M}}} + \dots + x_{\bar{N}} \frac{\frac{1}{x_{\bar{N}}^{\beta'_M}}}{\frac{1}{x_1^{\beta'_M}} + \frac{1}{x_2^{\beta'_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta'_M}}} \\ &= x_1 \frac{\frac{1}{x_1^{\beta_M + \Delta\beta}}}{\frac{1}{x_1^{\beta_M + \Delta\beta}} + \frac{1}{x_2^{\beta_M + \Delta\beta}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M + \Delta\beta}}} + x_2 \frac{\frac{1}{x_2^{\beta_M + \Delta\beta}}}{\frac{1}{x_1^{\beta_M + \Delta\beta}} + \frac{1}{x_2^{\beta_M + \Delta\beta}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M + \Delta\beta}}} \\ &\quad + \dots + x_{\bar{N}} \frac{\frac{1}{x_{\bar{N}}^{\beta_M + \Delta\beta}}}{\frac{1}{x_1^{\beta_M + \Delta\beta}} + \frac{1}{x_2^{\beta_M + \Delta\beta}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M + \Delta\beta}}} \end{aligned}$$

Now we calculate the following ratio:

$$\begin{aligned} \frac{\mathbb{E}_{\mathcal{M} \sim P_{\beta'_M}} QCID(\mathcal{M})}{\mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M})} &= \frac{\left(\frac{x_1}{x_1^{\beta_M + \Delta\beta}} + \frac{x_2}{x_2^{\beta_M + \Delta\beta}} + \dots + \frac{x_{\bar{N}}}{x_{\bar{N}}^{\beta_M + \Delta\beta}} \right) \left(\frac{1}{x_1^{\beta_M}} + \frac{1}{x_2^{\beta_M}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M}} \right)}{\left(\frac{1}{x_1^{\beta_M + \Delta\beta}} + \frac{1}{x_2^{\beta_M + \Delta\beta}} + \dots + \frac{1}{x_{\bar{N}}^{\beta_M + \Delta\beta}} \right) \left(\frac{x_1}{x_1^{\beta_M}} + \frac{x_2}{x_2^{\beta_M}} + \dots + \frac{x_{\bar{N}}}{x_{\bar{N}}^{\beta_M}} \right)} \\ &= \frac{\sum_{1 \leq i \leq j \leq \bar{N}} \frac{1}{x_i^{\beta_M} x_j^{\beta_M}} \left(\frac{x_i}{x_i^{\Delta\beta}} + \frac{x_j}{x_j^{\Delta\beta}} \right)}{\sum_{1 \leq i \leq j \leq \bar{N}} \frac{1}{x_i^{\beta_M} x_j^{\beta_M}} \left(\frac{x_j}{x_j^{\Delta\beta}} + \frac{x_i}{x_i^{\Delta\beta}} \right)} \end{aligned}$$

Since we assume that $x_1 \leq x_2 \leq \dots \leq x_{\bar{N}}$, we have $\frac{1}{x_1^{\Delta\beta}} \geq \frac{1}{x_2^{\Delta\beta}} \geq \dots \geq \frac{1}{x_{\bar{N}}^{\Delta\beta}}$. Then for each $1 \leq i \leq j \leq \bar{N}$, according to the rearrangement inequality, we have

$$\frac{x_i}{x_i^{\Delta\beta}} + \frac{x_j}{x_j^{\Delta\beta}} \leq \frac{x_i}{x_j^{\Delta\beta}} + \frac{x_j}{x_i^{\Delta\beta}}$$

Furthermore, among all the (x_i, x_j) pairs, it is easy to find one $(x_i, x_{i'})$ such that it satisfies $x_i \neq x_{i'}$. Thus we have

$$\frac{x_i}{x_i^{\Delta\beta}} + \frac{x_{i'}}{x_{i'}^{\Delta\beta}} < \frac{x_i}{x_{i'}^{\Delta\beta}} + \frac{x_{i'}}{x_i^{\Delta\beta}}$$

Consequently, we have

$$\begin{aligned} \sum_{1 \leq i \leq j \leq \bar{N}} \frac{1}{x_i^{\beta_M} x_j^{\beta_M}} \left(\frac{x_i}{x_i^{\Delta\beta}} + \frac{x_j}{x_j^{\Delta\beta}} \right) &< \sum_{1 \leq i \leq j \leq \bar{N}} \frac{1}{x_i^{\beta_M} x_j^{\beta_M}} \left(\frac{x_i}{x_j^{\Delta\beta}} + \frac{x_j}{x_i^{\Delta\beta}} \right) \Rightarrow \\ \frac{\mathbb{E}_{\mathcal{M} \sim P_{\beta'_M}} QCID(\mathcal{M})}{\mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M})} &< 1 \Rightarrow \mathbb{E}_{\mathcal{M} \sim P_{\beta'_M}} QCID(\mathcal{M}) < \mathbb{E}_{\mathcal{M} \sim P_{\beta_M}} QCID(\mathcal{M}) \end{aligned}$$

C.3 PROOF OF THEOREM 3.3

Proof C.3 According to Schneider & Barker (1989), we first define the principle submatrix, which is a submatrix where the set of remaining row indices is the same as the remaining set of column indices.

Before selecting the first client, we need to calculate the following value for all clients $c_1 \in \{1, 2, 3, \dots, N\}$,

$$P(C_1 = c_1) \propto \frac{1}{[QCID(\mathcal{M}_1)]^{\beta_1}} + \lambda \sqrt{\frac{3 \ln k}{2T_{c_1}}}, \quad \beta_1 > 0.$$

To derive the $QCID(\mathcal{M}_1)$ for each $c_1 \in \{1, 2, 3, \dots, N\}$, according to Theorem 3.1, we need to find the principle submatrix of S , denoted by S_1 , in which the set of column indices is $\mathcal{M}_1$. Then we need to calculate the sum of all the elements in S_1 . Since there are N different values for c_1 and the dimension of S_1 is 1×1 , we need to conduct the computation for N times.

After selecting $\mathcal{M}_1 = \{c_1\}$, we need to select $c_2 \in \{1, 2, 3, \dots, N\} / \mathcal{M}_1$ to form $\mathcal{M}_2 = \mathcal{M}_1 \cup \{c_2\}$.

Before selecting the second client, we need to calculate the following value for all the $\mathcal{M}_2 = \{c_1, c_2\}$ where $c_2 \in \{1, 2, 3, \dots, N\} / \mathcal{M}_1$,

$$P(C_2 = c_2 | C_1 = c_1) \propto \frac{\frac{1}{[QCID(\mathcal{M}_2)]^{\beta_2}}}{\frac{1}{[QCID(\mathcal{M}_1)]^{\beta_1}} + \alpha \sqrt{\frac{3 \ln k}{2T_{c_1}}}}$$

To derive the $QCID(\mathcal{M}_2)$ for each $c_2 \in \{1, 2, 3, \dots, N\} / \{c_1\}$, according to Theorem 3.1, we need to find the principle submatrix of S , denoted by S_2 , in which the set of column indices is $\mathcal{M}_2$. Then we need to calculate the sum of all the elements in S_2 . Since there are $N - 1$ different values for c_2 , there will be $N - 1$ different S_2 . Also, because we have already calculate the sum of all the elements in S_1 , which is a submatrix of S_2 , in our first step, we now only need to sum over all the other elements in S_2 . Since the dimension of S_2 is 2×2 , we need to do the computation for $(N - 1) \times (2^2 - 1)$ times.

This procedure goes on. After selecting $\mathcal{M}_{m-1} = \{c_1, c_2, \dots, c_{m-1}\}$, where $3 \leq m \leq M$, we need to select $c_m \in \{1, 2, 3, \dots, N\} / \mathcal{M}_{m-1}$ to form $\mathcal{M}_m = \mathcal{M}_{m-1} \cup \{c_m\}$. Before selecting the m -th client, we need to calculate the following value for all the $\mathcal{M}_m = \{c_1, c_2, \dots, c_m\}$ where $c_m \in \{1, 2, 3, \dots, N\} / \mathcal{M}_{m-1}$,

$$P(C_m = c_m | C_{m-1} = c_{m-1}, \dots, C_2 = c_2, C_1 = c_1) \propto \frac{[QCID(\mathcal{M}_{m-1})]^{\beta_{m-1}}}{[QCID(\mathcal{M}_m)]^{\beta_m}}$$

To derive the $QCID(\mathcal{M}_m)$ for each $c_m \in \{1, 2, 3, \dots, N\} / \{c_1, c_2, \dots, c_{m-1}\}$, according to Theorem 3.1, we need to find the principle submatrix of S , denoted by S_m , in which the set of column indices is $\mathcal{M}_m$. Then we need to calculate the sum of all the elements in S_m . Since there are $N - (m - 1)$ different values for c_m , there will be $N - (m - 1)$ different S_m . Since we have already calculate the sum of all the elements in S_{m-1} , which is a submatrix of S_m , in our previous step, now we only need to sum all the other elements in S_m . Since the dimension of S_m is $m \times m$, we need to conduct the computation for $(N - (m - 1)) \times (m^2 - (m - 1)^2)$ times.

In summary, in our strategy, the total times of computations we need to conduct are

$$\begin{aligned} & N + (N - 1) \times (2^2 - 1) + \dots + (N - (m - 1)) \times (m^2 - (m - 1)^2) + \dots + (N - M) \times (M^2 - (M - 1)^2) \\ & \leq N + N \times (2^2 - 1) + \dots + N \times (m^2 - (m - 1)^2) + \dots + N \times (M^2 - (M - 1)^2) \\ & = N \times M^2, \end{aligned}$$

which finishes the proof that the computation complexity for our method is $\mathcal{O}(N \times M^2)$.

C.4 PROOF OF THEOREM 4.5

Proof C.4 Suppose there are N available clients and their indices are denoted by $\{1, 2, 3, \dots, N\}$. Our goal is to get a subset $\mathcal{M}$ of $\{1, 2, 3, \dots, N\}$ following the probability law S of some client selection strategy. Let $\mathbf{w}_n^{(k,t)}$ denote the model parameter of client n after t local updates in the k -th communication round and $\mathbf{w}^{(k,0)}$ denote the global model parameter at the beginning of the k -th communication round. According to the proof of Theorem 1 in Wang et al. (2020a), we can define the following auxiliary variables for the setting where we adopt FedAvg as the FL optimizer and all the client conduct τ local updates in each communication round k :

$$\text{Normalized Stochastic Gradient: } \mathbf{d}_n^{(k)} = \sum_{t=0}^{\tau-1} g_n(\mathbf{w}_n^{(k,t)})$$

$$\text{Normalized Gradient: } \mathbf{h}_n^{(k)} = \sum_{t=0}^{\tau-1} \nabla F_n(\mathbf{w}_n^{(k,t)}).$$

$$\text{Normalized Class-wise Gradient: } \mathbf{h}_{(n,b)}^{(k)} = \sum_{t=0}^{\tau-1} \nabla F_{(n,b)}(\mathbf{w}_n^{(k,t)}).$$

$$\text{It is easy to verify that } \mathbf{h}_n^{(k)} = \sum_{b=1}^B \alpha_{(n,b)} \mathbf{h}_{(n,b)}^{(k)}.$$

According to the proof of Theorem 1 in Wang et al. (2020a), one can show that $\mathbb{E}[\mathbf{d}_n^{(k)} - \mathbf{h}_n^{(k)}] = 0$.

Besides, since clients are independent to each other, we have $\mathbb{E}\langle \mathbf{d}_n^{(k)} - \mathbf{h}_n^{(k)}, \mathbf{d}_{n'}^{(k)} - \mathbf{h}_{n'}^{(k)} \rangle = 0, \forall n \neq n'$. Recall that the update rule of the global model can be written as follows:

$$\mathbf{w}^{(k+1,0)} - \mathbf{w}^{(k,0)} = -\eta \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{d}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n},$$

where η is the learning rate. According to the Lipschitz-smooth assumption for the global objective function $\tilde{F}$ (Assumption 4.1), it follows that

$$\begin{aligned} & \mathbb{E}[\tilde{F}(\mathbf{w}^{(k+1,0)})] - \tilde{F}(\mathbf{w}^{(k,0)}) \\ & \leq -\eta \underbrace{\mathbb{E}\left[\left\langle \nabla \tilde{F}(\mathbf{w}^{(k,0)}), \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{d}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\rangle\right]}_{T_1} + \underbrace{\frac{\eta^2 L_{\tilde{F}}}{2} \mathbb{E}\left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{d}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2\right]}_{T_2} \end{aligned} \quad (3)$$

where the expectation is taken over randomly selected indices set $\mathcal{M}$ as well as mini-batches $\xi_i^{(k,t)}, \forall n \in \{1, 2, \dots, m\}, t \in \{0, 1, \dots, \tau - 1\}$

Similar to the proof in Wang et al. (2020a), to bound the T_1 in (3), we should notice that

$$\begin{aligned} T_1 &= \mathbb{E}\left[\left\langle \nabla \tilde{F}(\mathbf{w}^{(k,0)}), \frac{\sum_{n \in \mathcal{M}} q_n (\mathbf{d}_n^{(k)} - \mathbf{h}_n^{(k)})}{\sum_{n \in \mathcal{M}} q_n} \right\rangle\right] + \mathbb{E}\left[\left\langle \nabla \tilde{F}(\mathbf{w}^{(k,0)}), \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\rangle\right] \\ &= \mathbb{E}\left[\left\langle \nabla \tilde{F}(\mathbf{w}^{(k,0)}), \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\rangle\right] \\ &= \frac{1}{2} \|\nabla \tilde{F}(\mathbf{w}^{(k,0)})\|^2 + \frac{1}{2} \mathbb{E}\left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2\right] - \frac{1}{2} \mathbb{E}\left[\left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) - \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2\right] \end{aligned} \quad (4)$$

where the last equation uses the fact: $2\langle a, b \rangle = \|a\|^2 + \|b\|^2 - \|a - b\|^2$.

T_2 is similar as the one in Wang et al. (2020a). According to the proof in Section C.2 of Wang et al. (2020a), we have the following bound for T_2 ,

$$\begin{aligned} T_2 &\leq 2\sigma^2 \mathbb{E}\left[\frac{\sum_{n \in \mathcal{M}} q_n^2}{(\sum_{n \in \mathcal{M}} q_n)^2}\right] + 2\mathbb{E}\left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2\right] \\ &\leq 2\sigma^2 + 2\mathbb{E}\left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2\right] \end{aligned} \quad (5)$$

Plugging (4) and (5) back into (3), we have

$$\begin{aligned}
& \mathbb{E} \left[\tilde{F} \left(\mathbf{w}^{(k+1,0)} \right) \right] - \tilde{F} \left(\mathbf{w}^{(k,0)} \right) \\
& \leq - \underbrace{\eta \mathbb{E} \left[\left\langle \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right), \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{d}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\rangle \right]}_{T_1} + \underbrace{\frac{\eta^2 L_{\tilde{F}}}{2} \mathbb{E} \left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{d}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \right]}_{T_2} \\
& \leq - \frac{1}{2} \eta \left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) \right\|^2 - \frac{1}{2} \eta \mathbb{E} \left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \right] + \frac{1}{2} \eta \mathbb{E} \left[\left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) - \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \right] \\
& \quad + \eta^2 L_{\tilde{F}} \sigma^2 + \eta^2 L_{\tilde{F}} \mathbb{E} \left[\left\| \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \right] \tag{6}
\end{aligned}$$

If we set $\eta \leq \frac{1}{2L}$, we have

$$\begin{aligned}
& \mathbb{E} \left[\tilde{F} \left(\mathbf{w}^{(k+1,0)} \right) \right] - \tilde{F} \left(\mathbf{w}^{(k,0)} \right) \\
& \leq - \frac{1}{2} \eta \left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) \right\|^2 + \frac{1}{2} \eta \mathbb{E} \left[\left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) - \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \right] + \eta^2 L_{\tilde{F}} \sigma^2. \tag{7}
\end{aligned}$$

Now we focus on the $\mathbb{E} \left[\left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) - \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \right]$ in the following:

$$\begin{aligned}
& \mathbb{E} \left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) - \frac{\sum_{n \in \mathcal{M}} q_n \mathbf{h}_n^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 = \mathbb{E} \left\| \frac{1}{B} \sum_{b=1}^B \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) - \frac{\sum_{n \in \mathcal{M}} q_n \left(\sum_{b=1}^B \alpha_{(n,b)} \mathbf{h}_{(n,b)}^{(k)} \right)}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
& = \mathbb{E} \left\| \frac{1}{B} \sum_{b=1}^B \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) - \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} \mathbf{h}_{(n,b)}^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
& \leq 2 \mathbb{E} \left\| \frac{1}{B} \sum_{b=1}^B \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) - \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right)}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
& \quad + 2 \mathbb{E} \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right)}{\sum_{n \in \mathcal{M}} q_n} - \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} \mathbf{h}_{(n,b)}^{(k)}}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
& = 2 \mathbb{E} \underbrace{\left\| \sum_{b=1}^B \left(\frac{1}{B} - \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} \right) \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) \right\|^2}_{T_3} \\
& \quad + 2 \mathbb{E} \underbrace{\left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) - \mathbf{h}_{(n,b)}^{(k)}]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2}_{T_4} \tag{8}
\end{aligned}$$

For T_3 , according to the Cauchy-Schwarz inequality and Assumption 4.3, we have

$$\begin{aligned}
& \mathbb{E} \left\| \sum_{b=1}^B \left(\frac{1}{B} - \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} \right) \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) \right\|^2 \\
& \leq B \mathbb{E} \left[\sum_{b=1}^B \left(\frac{1}{B} - \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} \right)^2 \left\| \frac{1}{B} \sum_{b=1}^B \nabla \tilde{F}_b \left(\mathbf{w}^{(k,0)} \right) \right\|^2 \right] \\
& = B \left\| \nabla \tilde{F} \left(\mathbf{w}^{(k,0)} \right) \right\|^2 \mathbb{E}[QCID(\mathcal{M})] + \gamma^2 \mathbb{E}[QCID(\mathcal{M})] \tag{9}
\end{aligned}$$

For T_4 , we have

$$\begin{aligned}
& \mathbb{E} \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla \tilde{F}_b(\mathbf{w}^{(k,0)}) - \mathbf{h}_{(n,b)}^{(k)}]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
&= \mathbb{E} \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla \tilde{F}_b(\mathbf{w}^{(k,0)}) - \nabla F_{(n,b)}(\mathbf{w}^{(k,0)}) + \nabla F_{(n,b)}(\mathbf{w}^{(k,0)}) - \mathbf{h}_{(n,b)}^{(k)}]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
&\leq 2 \mathbb{E} \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla \tilde{F}_b(\mathbf{w}^{(k,0)}) - \nabla F_{(n,b)}(\mathbf{w}^{(k,0)})]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
&\quad + 2 \mathbb{E} \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla F_{(n,b)}(\mathbf{w}^{(k,0)}) - \mathbf{h}_{(n,b)}^{(k)}]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
&\leq 2\kappa^2 \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)}}{\sum_{n \in \mathcal{M}} q_n} + 2 \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla F_{(n,b)}(\mathbf{w}^{(k,0)}) - \mathbf{h}_{(n,b)}^{(k)}]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2 \\
&\leq 2\kappa^2 + 2 \left\| \sum_{b=1}^B \frac{\sum_{n \in \mathcal{M}} q_n \alpha_{(n,b)} [\nabla F_{(n,b)}(\mathbf{w}^{(k,0)}) - \mathbf{h}_{(n,b)}^{(k)}]}{\sum_{n \in \mathcal{M}} q_n} \right\|^2, \tag{10}
\end{aligned}$$

where $\kappa = \max_{\{n,b\}} \kappa_{\{n,b\}}$. According to the results from the proof in C.5 in Wang et al. (2020a), we have

$$\begin{aligned}
& \mathbb{E} \left\| \nabla F_{(n,b)}(\mathbf{w}^{(k,0)}) - \mathbf{h}_{(n,b)}^{(k)} \right\|^2 \leq \frac{L_{n,b}^2}{\tau} \sum_{k=0}^{\tau-1} \mathbb{E} \left[\left\| \mathbf{w}^{(k,0)} - \mathbf{w}_n^{(k,t)} \right\|^2 \right] \\
&\leq \frac{L^2}{\tau} \sum_{k=0}^{\tau-1} \mathbb{E} \left[\left\| \mathbf{w}^{(k,0)} - \mathbf{w}_n^{(k,t)} \right\|^2 \right] \\
&\leq \frac{2\eta^2 L^2 \sigma^2}{1-D} (\tau-1) + \frac{D}{1-D} \mathbb{E} \left[\left\| \nabla F_i(\mathbf{w}^{(k,0)}) \right\|^2 \right] \\
&\leq \frac{2\eta^2 L^2 \sigma^2}{1-D} (\tau-1) + \frac{2D}{1-D} \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 + \frac{2D}{1-D} \mathbb{E} \left[\left\| \nabla F_i(\mathbf{w}^{(k,0)}) - \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 \right] \\
&\leq \frac{2\eta^2 L^2 \sigma^2}{1-D} (\tau-1) + \frac{2D}{1-D} \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 + \frac{2D}{1-D} \mathbb{E} \left[\left\| \nabla F_i(\mathbf{w}^{(k,0)}) - \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 \right] \\
&\leq \frac{2\eta^2 L^2 \sigma^2}{1-D} (\tau-1) + \frac{2D}{1-D} \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 + \frac{1}{B} \frac{2D}{1-D} \kappa^2 \tag{11}
\end{aligned}$$

where $L = \max_{\{n,b\}} L_{n,b}$ and $D = 4\eta^2 L^2 \tau (\tau-1)$.

Combining the results in (7), (8), (9), (10) and (11), it is easy to derive that

$$\begin{aligned}
& \mathbb{E} \left[\tilde{F}(\mathbf{w}^{(k+1,0)}) \right] - \tilde{F}(\mathbf{w}^{(k,0)}) \\
&\leq -\left(\frac{1}{2} - B\delta\mathbb{E}[QCID]\right) - \frac{4D}{1-D} \eta \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 + 4\eta\kappa^2 + \frac{4\eta^3 L^2 \sigma^2}{1-D} (\tau-1) \\
&\quad + \frac{1}{B} \frac{4D}{1-D} \eta\kappa^2 + \eta^2 L_{\tilde{F}}^2 \sigma^2 + \gamma^2 \eta \mathbb{E}[QCID] \tag{12}
\end{aligned}$$

Now we have

$$\begin{aligned}
& \frac{\mathbb{E} [\tilde{F}(\mathbf{w}^{(k+1,0)})] - \tilde{F}(\mathbf{w}^{(k,0)})}{\eta} \\
& \leq -\left(\frac{1}{2} - B\delta\mathbb{E}[QCID] - \frac{4D}{1-D}\right) \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 + 4\kappa^2 \\
& \quad + \frac{4\eta^2 L^2 \sigma^2}{1-D} (\tau - 1) + \frac{1}{B} \frac{4D}{1-D} \kappa^2 + \gamma^2 \mathbb{E}[QCID] + \eta L_{\tilde{F}} \sigma^2
\end{aligned} \tag{13}$$

Taking the total expectation and averaging over all rounds, one can obtain

$$\begin{aligned}
\frac{\mathbb{E} [\tilde{F}(\mathbf{w}^{(K,0)})] - \tilde{F}(\mathbf{x}^{(0,0)})}{\eta K} & \leq -\left(\frac{1}{2} - B\delta\mathbb{E}[QCID] - \frac{4D}{1-D}\right) \frac{1}{K} \sum_{t=1}^{K-1} \mathbb{E} \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 \\
& \quad + \gamma^2 \mathbb{E}[QCID] + \left(4 + \frac{1}{B} \frac{4D}{1-D}\right) \kappa^2 + \frac{4\eta^2 L^2 \sigma^2}{1-D} (\tau - 1) + \eta L \sigma^2
\end{aligned} \tag{14}$$

Finally, we have

$$\begin{aligned}
\min_{k \leq K} \mathbb{E} \left\| \nabla F(\mathbf{w}^{(k,0)}) \right\|^2 & \leq \frac{1}{K} \sum_{k=1}^{K-1} \mathbb{E} \left\| \nabla \tilde{F}(\mathbf{w}^{(k,0)}) \right\|^2 \\
& \leq \frac{1}{\left(\frac{1}{2} - B\delta\mathbb{E}[QCID(\mathcal{M})] - \frac{4D}{1-D}\right)} \left[\frac{\tilde{F}(\mathbf{w}^{(0,0)}) - \tilde{F}_{min}}{\eta K} \right. \\
& \quad \left. + \left(4 + \frac{1}{B} \frac{4D}{1-D}\right) \kappa^2 + \frac{4\eta^2 L^2 \sigma^2}{1-D} (\tau - 1) + \eta L_{\tilde{F}} \sigma^2 + \gamma^2 \mathbb{E}[QCID] \right]
\end{aligned} \tag{15}$$

If setting $\eta = \frac{s}{10L\sqrt{\tau(\tau-1)K}}$ with $s < 1$, we have

$$\begin{aligned}
\min_{k \leq K} \mathbb{E} \left\| \nabla F(\mathbf{w}^{(k,0)}) \right\|^2 & \leq \frac{1}{\frac{1}{3} - B\delta\mathbb{E}[QCID(\mathcal{M})]} \left[\frac{\tilde{F}(\mathbf{w}^{(0,0)}) - \tilde{F}_{min}}{s\sqrt{K}/(10L\sqrt{\tau(\tau-1)})} \right. \\
& \quad \left. + 5\kappa^2 + \frac{\sigma^2 s^2}{25\tau K} + \frac{sL_{\tilde{F}}\sigma^2}{10L\sqrt{\tau(\tau-1)K}} + \gamma^2 \mathbb{E}[QCID] \right]
\end{aligned} \tag{17}$$

Since $\tilde{F}$ is larger than 0, $\tilde{F}_{min} > 0$. Now we let $\mathbf{w}^{(k)}$ denote the global model parameter at the k -th communication round and $\mathbf{w}^{(0)}$ denote the initial parameter. After changing the notations, we can finish our proof by the following:

$$\begin{aligned}
\min_{k \leq K} \mathbb{E} \left\| \nabla \tilde{F}(\mathbf{w}^{(k)}) \right\|^2 & \leq \frac{1}{\frac{1}{3} - B\delta\mathbb{E}[QCID(\mathcal{M})]} \left[\frac{\sigma^2 s^2}{25\tau K} + \frac{sL_{\tilde{F}}\sigma^2}{10L\sqrt{\tau(\tau-1)K}} \right. \\
& \quad \left. + 5\kappa^2 + \frac{10L\sqrt{\tau(\tau-1)}\tilde{F}(\mathbf{w}^{(0)})}{s\sqrt{K}} + \gamma^2 \mathbb{E}[QCID] \right],
\end{aligned}$$

D SUPPLEMENTAL EXPERIMENT SETTINGS AND RESULTS

We run all the other experiments on a sever with a single NVIDIA TITAN RTX GPU and Intel Xeon Gold 6254 CPU.

D.1 THE EXPERIMENTAL SETTINGS IN SECTION 2.1

We adopt an MLP model with one hidden layer of 64 units and FedAvg (McMahan et al., 2017b) as the FL optimizer. In Figure 1a, we allocate the MNIST data to $N = 100$ clients with each client only

accessing to the same amount of data from one class. In Figure 1b, each client is associated with the same amount of data from two classes. In Figure 1c and 1d, we first allocate the whole MNIST dataset to $N = 200$ clients and pick 100 to construct a class-imbalanced global dataset. The global dataset with the 100 clients has the same amount of n_1 data samples for five classes and has the same amount of n_2 data samples for the other five classes. The ration r between n_1 and n_2 is set to $3 : 1$.

In each training round (communication round), all of the clients conduct 5 local training epochs. The batch size is 50 for each client. The local optimizer is SGD with a weight decay of 0.0005. The learning rate is 0.01 initially and the decay factor is 0.9992. We terminate the FL training after 200 training rounds (communication rounds) and then evaluate the model’s performance on the test dataset of MNIST.

D.2 ADDITIONAL EXPERIMENTAL SETTINGS IN SECTION 5

The model we adopt has two convolutional layers with the number of kernel being 6 and 16, respectively. And all convolution kernels are of size 5×5 . The outputs of convolutional layers are fed into two hidden layers with 120 and 84 units.

In our implementation of Power-of-choice selection strategy (pow-d)(Cho et al., 2020), we first sample a candidate set $\mathcal{A}$ of 20 clients without replacement such that client n is chosen with probability proportional to the size of their local dataset q_n . Then the server sends the current global model to the clients in set $\mathcal{A}$, and these clients compute and send back to the server their local loss. To derive $\mathcal{M}$, we select M clients which have the highest loss from $\mathcal{A}$.

In our implementation of the method in (Yang et al., 2020) (Fed-cucb), the exploration factor to balance the trade-off between exploitation and exploration is set as 0.2 and forgetting factor as 0.99, which are the same as the settings in Yang et al. (2020).

With help with FHE, we can derive the matrix of inner products S accurately. Hence, in the simulation of our method, Fed-CBS, we ignore the process of deriving S and focus on our sampling strategy.

D.3 ADDITIONAL DETAILS FOR THE EXPERIMENTAL SETTINGS IN CASE 1 AND CASE 2

Case 1 In this setting, we have 120 clients in total, and each client has only one class of data.

When $n_1 : n_2 = 3 : 1$, there are 18 clients having the data from the 1st class, 18 clients having the data from the 2nd class, 18 clients having the data from the 3rd class, 18 clients having the data from the 4th class and 18 clients having the data from the 5th class. There are 6 clients having the data from the 6th class, 6 clients having the data from the 7th class, 6 clients having the data from the 8th class, 6 clients having the data from the 9th class and 6 clients having the data from the 10th class.

When $n_1 : n_2 = 5 : 1$, there are 20 clients having the data from the 1st class, 20 clients having the data from the 2nd class, 20 clients having the data from the 3rd class, 20 clients having the data from the 4th class and 20 clients having the data from the 5th class. There are 4 clients having the data from the 6th class, 4 clients having the data from the 7th class, 4 clients having the data from the 8th class, 4 clients having the data from the 9th class and 4 clients having the data from the 10th class.

Then we uniformly set 30% (36 clients) of them available. Since there are more clients which contain the data from the first 5 classes among the above 120 clients. The global dataset of these 36 clients is often class-imbalanced.

Case 2 In this setting, we have 200 clients in total and each client has only one class of data. For all the $i \in \{1, 2, \dots, 10\}$, there are 20 clients having the data from the i -th class.

When $n_1 : n_2 = 3 : 1$, we randomly pick 9 clients from the 20 clients which have the data from the 1st class and set them available. We randomly pick 9 clients from the 20 clients which have the data from the 2nd class and set them available. Similarly, for the k -th class ($2 < k \leq 5$), we randomly pick 9 clients from the 20 clients which have the data from the k -th class and set them available. On the contrary, we randomly pick 3 clients from the 20 clients which have the data from the 6th class and set them available. We randomly pick 3 clients from the 20 clients which have the data from the 7th class and set them available. Similarly, for $7 < k \leq 10$, we randomly pick 3 clients from the 20 clients which have the data from the k -th class and set them available. There are 60 clients in total.

When $n_1 : n_2 = 5 : 1$, we randomly pick 10 clients from the 20 clients which have the data from the 1st class and set them available. We randomly pick 10 clients from the 20 clients which have the data from the 2nd class and set them available. For the k -th class ($2 < k \leq 5$), we randomly pick 10 clients from the 20 clients which have the data from the k -th class and set them available. On the contrary, we randomly pick 2 clients from the 20 clients which have the data from the 6th class and set them available. We randomly pick 2 clients from the 20 clients which have the data from the 7th class and set them available. And for the other k -th class ($7 < k \leq 10$), we randomly pick 2 clients from the 20 clients which have the data from the k -th class and set them available. There are 60 clients in total.

Since there are more clients that contain the data from the first 5 classes among the above 60 clients, the global dataset of these 60 clients is always class-imbalanced.

The difference between the settings of Case 1 and Case 2 is that we uniformly set 30% clients available in Case 1 but non-uniformly set 30% clients available in Case 2. Nevertheless, the global datasets of the available clients are both class-imbalanced in both cases.

D.4 THE AVERAGED QCID VALUES FOR CASE 1 AND CASE 2 IN SECTION 5.2

$\mathbb{E}[QCID](10^{-2})$		all	rand	pow-d	Fed-cucb	Fed-CBS
Case 1	3:1	2.90±0.02	9.33±0.17	13.70±0.39	1.39±0.37	0.57±0.04
	5:1	6.17±0.04	12.36±0.20	16.63±0.74	3.43±0.76	2.41±0.07
Case 2	3:1	2.50±0.00	9.91±0.16	13.68±0.72	1.89±1.72	0.001±0.001
	5:1	4.44±0.00	11.70±0.20	15.68±0.96	2.63±2.40	0.002±0.001

Table 4: The averaged $QCID$ values for four baselines and our method. Our method, Fed-CBS, has successfully reduced the class-imbalance. Since the global dataset of all the 60 available clients is always class-imbalanced and the ratio is always fixed in case 2, the $QCID$ value is fixed and the derivation of it is always zero.

D.5 EXPERIMENT RESULTS OF FASHION-MNIST DATASET

Experiment Setup We adopt an MLP model with one hidden layer of 64 units and FedNova (Wang et al., 2020a) as the FL optimizer. Similar to the setup in the experiment of CIFAR-10, the batch size is 50 for each client. In each communication round, all of them conduct the same number of local updates, which allows the client with the largest local dataset to conduct 5 local training epochs. In our method, we set the $\beta_m = m$, $\gamma = 10$ and $L_b = 10^{-20}$. The local optimizer is SGD with a weight decay of 0.0005. The learning rate is 0.01 initially and the decay factor is 0.9992. We terminate the FL training after 3000 communication rounds and then evaluate the model’s performance on the test dataset of Fashion-MNIST.

		all	rand	pow-d	Fed-cucb	Fed-CBS
Communication Rounds	$\alpha=0.1$	115±17	185±27	135±22	124±37	92±6
	$\alpha=0.2$	173±45	284±54	218±55	216±24	166±36
	$\alpha=0.5$	258±44	331±55	281±54	284±51	218±36
$\mathbb{E}[QCID](10^{-2})$	$\alpha=0.1$	1.40±0.11	8.20±0.19	11.72±0.33	4.24±0.59	0.15±0.02
	$\alpha=0.2$	1.39±0.22	7.67±0.26	10.31±0.24	4.43±0.38	0.21±0.01
	$\alpha=0.5$	0.94±0.07	5.93±0.26	7.68±0.28	4.34±0.85	0.22±0.01

Table 5: The communication rounds required for targeted test accuracy and the averaged $QCID$ values on Fashion-MNIST dataset. The targeted test accuracy is 78% for $\alpha = 0.1$, 80% for $\alpha = 0.2$ and 82% for $\alpha = 0.5$. The results are the mean and the standard deviation over 4 different random seeds.

D.5.1 RESULTS FOR CLASS-BALANCED GLOBAL DATASET

Similar to the experiment settings, in this experiment, we set 200 clients in total with a class-balanced global dataset. The non-IID data partition among clients is based on the settings of Dirichlet distribution parameterized by the concentration parameter α in Hsu et al. (2019). In each communication round, we uniformly and randomly set 30% of them (i.e., 60 clients) available and select 10 clients from those 60 available ones to participate in the training.

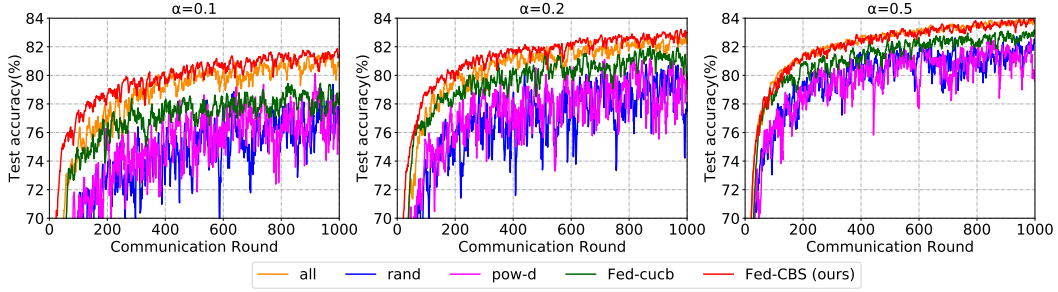


Figure 6: Test accuracy on Fashion-MNIST dataset under three heterogeneous settings.

As shown in Table 6, our method successfully reduces the class-imbalance, since it achieves the lowest *QCID* value compared with other client selection strategies. Our method outperforms the other three baseline methods and achieves comparable performance to the ideal setting where all the available clients are engaged in the training. As shown in Table 4 and Figure 6, our method can achieve faster and more stable convergence. It is worth noting that due to the inaccurate estimation of distribution and the weakness of greedy method discussed in Section 2.2, the performance of Fed-cucb is much worse than ours.

D.5.2 RESULTS FOR CLASS-IMBALANCED GLOBAL DATASET: CASE 1

Similar to the settings for Cifar-10, there are 120 clients in total and each client only has one class of data with the same quantity. The global dataset of these 120 clients is always class-imbalanced. To measure the degree of class imbalance, we let the global dataset have the same amount n_1 of data samples for five classes and have the same amount n_2 of data samples for the other five classes. The ratio r between n_1 and n_2 is set to 3 : 1 and 5 : 1 respectively in the experiments. In each communication round, we randomly set 30% of them (*i.e.*, 36 clients) available and select 10 clients to participate in the training.

As shown in the Table 6 and Figure 7a, our method can achieve faster and more stable convergence, and even better performance than the ideal setting where all the available clients are engaged. The performance of Fed-cucb (Yang et al., 2020) is better than the results on class-balanced global dataset, which is partly due to the simplicity of each client’s local dataset composition in our experiments as discussed in the experiments of Cifar-10.

D.5.3 RESULTS FOR CLASS-IMBALANCED GLOBAL DATASET: CASE 2

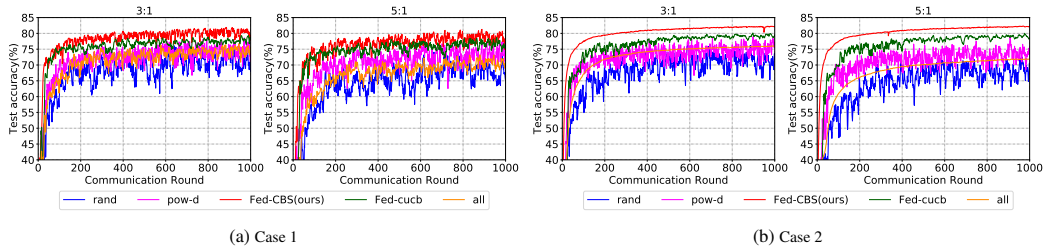


Figure 7: Test accuracy on Fashion-MNIST with class-imbalanced global dataset in Case 1 and Case 2.

		all	rand	pow-d	Fed-cucb	Fed-CBS
Case 1	3:1	78.42±0.79	78.46±0.90	81.08±0.21	80.83±0.91	81.75±0.34
	5:1	72.42±2.22	75.49±2.56	80.15±0.41	80.50±0.95	81.42±0.50
Case 2	3:1	74.64±1.87	78.80±0.55	81.13±0.41	79.94±0.31	81.95±0.57
	5:1	67.16±4.13	74.17±2.01	80.05±0.39	80.00±0.58	81.92±0.57

Table 6: Best test accuracy for our method and other four baselines on Fashion-MNIST dataset.

Similar to the settings of Cifar-10, we assume that there are 200 clients in total. In each communication round, 30% of them (*i.e.*, 60 clients) are set available in each training round. The global dataset of those 60 available clients is always class-imbalanced. To measure the degree of class imbalance, we make the global dataset have the same amount n_1 of data for the five classes and have the same amount n_2 of data for the other five classes. The ratio r between n_1 and n_2 is set to 3 : 1 and 5 : 1. We select 10 clients from these 60 clients to participate in the training.

As shown in the Table 6 and Figure 7b, our method can achieve higher test accuracy and more stable convergence, which outperforms the ideal setting where all the available clients are engaged. Since the global dataset of the available 60 clients in each communication round is always class-imbalanced, the performance of engaging all of them is not good.

E ABLATION STUDIES AND DISCUSSION

E.1 ACCURATE ESTIMATION VS INACCURATE ESTIMATION FOR FED-CUCB

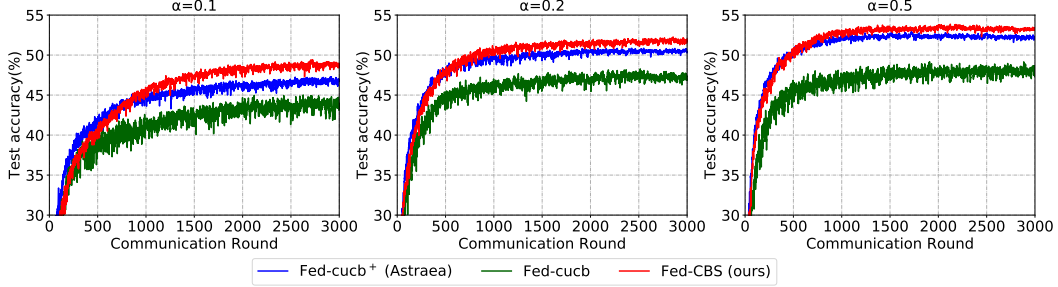


Figure 8: Test accuracy on Cifar-10 for Fed-cucb, Fed-cucb⁺ and Fed-CBS.

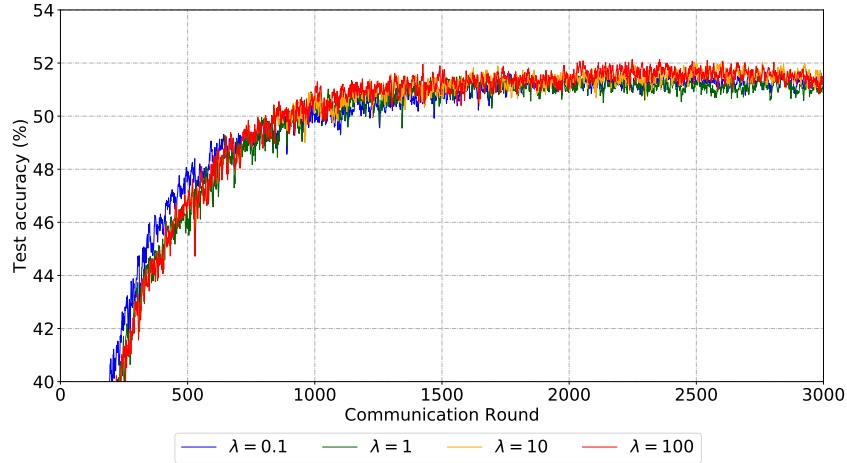
		Fed-cucb ⁺ (Astraea)	Fed-cucb	Fed-CBS
Best Accuracy (%)	$\alpha=0.1$	49.10 $\pm$ 0.70	46.84 $\pm$ 0.73	50.36$\pm$0.58
	$\alpha=0.2$	50.61 $\pm$ 0.77	48.80 $\pm$ 1.05	51.95$\pm$0.57
	$\alpha=0.5$	52.71 $\pm$ 0.27	50.98 $\pm$ 0.56	54.21$\pm$0.34
$\mathbb{E}(QCID) (10^{-2})$	$\alpha=0.1$	0.83 $\pm$ 0.18	7.09 $\pm$ 2.27	0.62$\pm$0.20
	$\alpha=0.2$	0.68 $\pm$ 0.05	5.93 $\pm$ 1.01	0.51$\pm$0.12
	$\alpha=0.5$	0.43 $\pm$ 0.04	6.47 $\pm$ 0.77	0.36$\pm$0.04

Table 7: Best accuracy and the averaged $QCID$ values.

As discussed in Sections 2.2 and 5.1, the estimation of the label distribution in Fed-cucb (Yang et al., 2020) is not accurate, which leads to performance degradation. Hence there comes a natural question, would the performance of Fed-cucb get improved if it got an exact estimation of the local label distribution? In our simulation, we manually let the Fed-cucb know the exact value of each client’s local label distribution and name it as Fed-cucb⁺. Actually, Fed-cucb⁺ is the core part of Astraea (Duan et al., 2019) without data augmentation. Hence, comparing our method with Fed-cucb⁺ can show the superiority of our sampling strategy over the greedy method in Fed-cucb (Yang et al., 2020) and Astraea (Duan et al., 2019).

E.2 THE EFFECT OF EXPLORATION FACTOR λ

As shown in the Figure 8 and Table 7, Fed-cucb⁺ does improve the performance of Fed-cucb, which verifies the importance of accurate estimation. However, our Fed-CBS still outperforms Fed-cucb⁺. Although, it seems that the accuracy of Fed-cucb⁺ increases a little faster than Fed-CBS at the beginning of the training, our method will achieve higher accuracy as the training proceeds further. As discussed in the Remark 3.3 in Section 3.3 and the Figure 2 of Section 5.1, this is due to the pitfall of greedy method, where one will miss the optimal solution. This has been verified by the averaged $QCID$ value in Table 7 and Fed-CBS achieves lower $\mathbb{E}(QCID)$ than Fed-cucb⁺ (Astraea).

Figure 9: Test accuracy with different exploration factor λ .

Another potential weakness of greedy method is the diversity of client composition. Following their selection process, once the first choice of client has been made, the following choices are fixed successively. Hence there are only limited kinds of client composition. It is interesting to investigate the relationship between the training performance and the diversity of client composition and we leave it as future work.

In our sampling strategy, when we sample the first client, we introduce the exploration factor λ to balance the tradeoff between exploitation and exploration. When the λ is small, our method will tend to exploit the class-balanced clients since their *QCID* values are smaller. For fairness, we hope every client can get the chance to be selected. Hence, we can increase the λ and then our method will tend to explore the clients which have seldom been selected before. However, it might cost many communication rounds for exploration and lead to slower convergence.

We conduct some experiments to verify the effect of exploration factor λ . The setting are the same as the ones in Section 5.1 when $\alpha = 0.2$. As shown in the Figure 9, as the λ becomes larger, the increase of accuracy will become a little slower at the start of the training. This because the it might cost more communication rounds for exploration. As the training proceeds, the accuracy with larger λ becomes a little higher than the ones with smaller λ . Overall, the improvement on the convergence speed and best accuracy is very slight, which means the performance of FL training is not very sensitive to the values of exploration factor λ . Generally, if we want to slightly fasten the convergence, we can decrease the value of λ . If we want to improve the best accuracy, we can increase the value of λ .

E.3 THE PERFORMANCE WITH DIFFERENT AMOUNTS OF SELECTED CLIENTS

In this section, we want to investigate how the amount of selected clients will affect the FL training performance. Generally, we think as the amount of selected clients increases, the FL training process can achieve better performance. However, once that amount reaches some threshold ϵ , the improvement will become slighter. This is because we find that select only a subset of all the available can achieve comparable results with engaging all the available clients into the training. As for how to decide the threshold ϵ , we provide the following two principles based on *QCID* and our experience.

- First, if we work on a classification task with B classes, we can select at least B clients. This is because in some special cases, each client will only have one class of data in their local dataset, such as the settings in Section 5.2. Hence, if less than B clients are selected, the grouped dataset of the selected clients will miss some classes of data.
- Second, to avoid missing some classes of data, we increase the threshold ϵ such that the averaged *QCID* value could smaller than $\frac{1}{B^2}$. This is because if the grouped dataset misses at least one class of data, the *QCID* will be larger than $\frac{1}{B^2}$.

We conduct some experiments to verify our prediction on the effect of the amount of selected clients. The settings are the same as the ones in Section 5.1 when $\alpha = 0.2$. As shown in the left figure of Figure 10, as the amount of selected clients increases, the FL training process can achieve better performance. However, when the amount M is larger than 10, the improvement is slighter. In the right figure of Figure 10, we can find that the averaged $QCID$ value of selecting 5 clients is larger than $(\frac{1}{10})^2 = 0.01$ and its performance is obviously worse than the others. These results verify the effectiveness of our principles on how to set the threshold ϵ . It is worth noting that due to limitation of communication bandwidth, we cannot select as many clients as possible. Hence, how to identify the appropriate threshold ϵ is critical to the FL training.

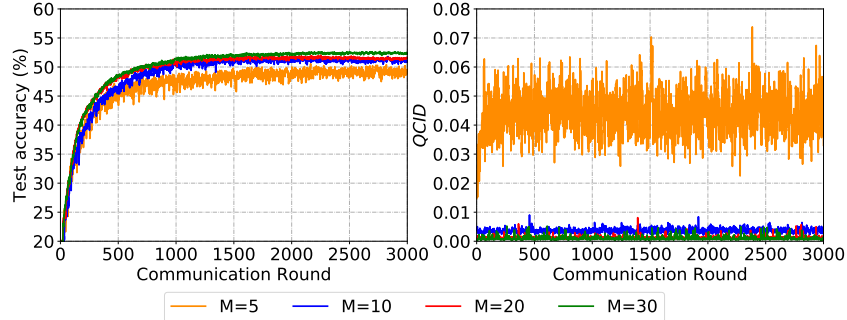


Figure 10: Left: The performance with different amounts of selected clients. Right: The $QCID$ with different amounts of selected clients.

Besides, we notice that if we increase the number of clients selected, the performance of random selection might improve. However, it is worth noting that our method can achieve comparable results with less selected clients, which reduces the computation and communication costs. Indeed, this has been reflected in our comparison between the baseline "all" and our method "Fed-CBS" in Section 5.1, which shows our method's strength. We **randomly** set 30% clients (60 clients) available in each round. Then our method will select 10 clients from them. The baseline "all" is the ideal case where we select all the 60 clients. In fact, it is a special case of "randomly selecting more clients". Although it is much better than randomly selecting 10 clients, our method is still comparable to it.

E.4 ADDITIONAL EXPERIMENTAL RESULTS ON FEMNIST DATASET

We also conduct some experiment on the FEMNIST Dataset. There are 3500 (> 1000) clients in total and we randomly set 10% ($< 30\%$) of them available in each round. Then our method tries to select 30 clients from them. That is less than 1% of all the 3550 clients and also less than the number of classes (64). Besides, we also run three baselines, randomly selecting 30 clients, randomly selecting 120 (> 100) clients, selecting 30 clients with fedcub. We present the results in Table 8 and Figure 11. Our performance is still the best. Due to the global imbalance, rand-120 is even worse than rand-30.

	rand-30	rand-120	Fed-cub	Fed-CBS
Communication Rounds	1106 $\pm$ 24	1394 $\pm$ 11	1124 $\pm$ 31	980 $\pm$ 17

Table 8: The communication rounds required for targeted test accuracy (75%). The results are the mean and the standard deviation over 3 different random seeds.

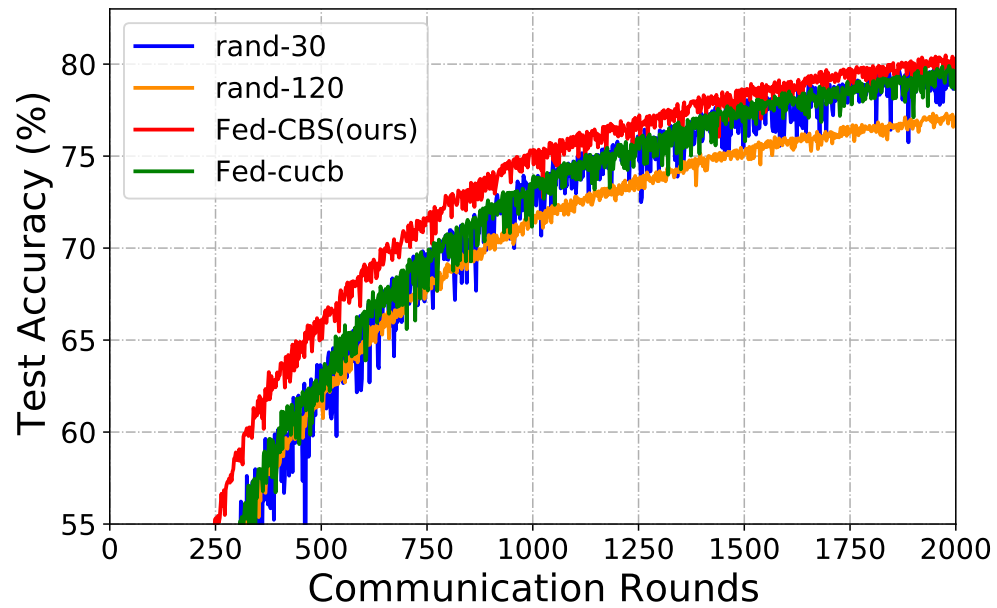


Figure 11: Test accuracy for FEMNIST