



Multiple RIS-assisted federated learning system based on channel quality device selection algorithm

Yujun Cai¹ · Shufeng Li¹ · Jianbo Liu¹ · Xinruo Zhang² · Xiaowei Shen³

Received: 1 March 2025 / Accepted: 15 August 2025
© The Author(s) 2025

Abstract

Machine learning has been introduced to assist humans in handling complex computations across various domains. However, concerns about inadequate user privacy have hindered its widespread adoption in communication systems. To enhance communication efficiency and reduce the overhead associated with machine learning, reconfigurable intelligent surfaces (RIS) and over-the-air computation (AirComp) have been utilized as supportive technologies. Simultaneously, federated learning (FL) emerged as an alternative approach to preserve user privacy and mitigate communication latency issues by requiring only the exchange of gradient information. Research has demonstrated that integrating FL with RIS leads to improved performance in system models. In this context, we propose a multiple RIS-assisted over-the-air FL system that incorporates a device selection algorithm based on channel quality. This algorithm identifies optimal devices for participation in the training process according to their communication channel conditions. Simulation results confirm that our proposed model outperforms the benchmark designs in terms of test accuracy and convergence speed.

Keywords Reconfigurable intelligent surface · Federated learning · Over-the-air computation · Device selection

Introduction

The application of machine learning and artificial intelligence has covered many aspects, such as face recognition, unmanned aerial vehicles, and wearable devices. Traditional machine learning algorithms are usually deployed on cloud

data centers and need mobile networks to transmit the raw data to edge devices, which may lead to great bandwidth pressure on the networks and the risk of leaking the privacy of users. Owing to the development of the framework of mobile edge computation, it is now possible to deploy machine learning at the edge of the network. In 2016, Google researchers introduced the concept of Federated Learning (FL), a method that enables multiple decentralized edge devices to collaboratively train a shared global model. This is achieved through a Parameter Server (PS), which coordinates the updates from devices without requiring the transfer of any raw data [1]. In this way, FL can protect users' privacy. The federated learning process involves iterative model updates. Clients download the global model, compute local updates (e.g., gradients), and transmit these updates to a central PS. The server then performs model aggregation using a suitable algorithm (e.g., FedAvg) and broadcasts the updated global model [2].

The iterative exchange of model updates in FL poses limitations on scalability and efficiency. The computational burden on clients, coupled with the high communication overhead from frequent gradient transmissions, necessitates novel techniques for communication-efficient FL [2].

Many strategies have been introduced to address the communication overhead in FL. For example, Yang et al. [3]

✉ Yujun Cai
tangentchase@cuc.edu.cn

Shufeng Li
lishufeng@cuc.edu.cn

Jianbo Liu
ljb@cuc.edu.cn

Xinruo Zhang
xinruo.zhang@essex.ac.uk

Xiaowei Shen
6511393822@qq.com

¹ School of Information and Engineering Communication, Communication University of China, Chaoyang, Beijing 100024, China

² School of Computer Science and Electronic Engineering, University of Essex, Wivenhoe Park, Colchester CO4 3SQ, UK

³ Qinghai Province Medium Wave Broadcasting Management Center, Xining 810000, Qinghai, China

introduces three approaches for user scheduling: random scheduling, round-robin, and proportional fair scheduling. The work examines how different signal-to-noise ratios (SNRs) impact test accuracy by exploring the relationship between the number of communication rounds and model performance. Experiments were conducted using models like convolutional neural networks and support vector machines. In another approach, Amiri et al. in [4] propose a federated edge learning framework, which dynamically selects the number of participating devices in each communication round to maximize test accuracy. Additionally, FL via over-the-air computation (AirComp) introduces three novel scheduling methods. While the channel-based approach offers computational simplicity, it achieves the lowest accuracy. In contrast, model update-based scheduling improves accuracy but at the cost of increased complexity. Hybrid scheduling strikes a compromise between computational efficiency and accuracy [5]. Furthermore, differential privacy mechanism based on FL is also proved to better protect users' privacy as well as improve model performance [6]. The authors of [7] propose a FL scheme that uses hierarchical protection and multiple aggregations to improve security and efficiency. The scheme achieves high accuracy while simultaneously reducing training time and communication traffic.

Despite the methods above, reconfigurable intelligent surface (RIS) is introduced into the wireless communication system as a promising technology. RIS makes the transmission environment controllable and can increase both energy efficiency and spectrum efficiency. Meanwhile, the deployment of RIS has a high level of flexibility and compatibility, which makes RIS a key technology of sixth-generation (6G) communication [8]. RIS is a plain array that consists of lots of reconfigurable passive components. Each component can adjust the propagation direction and amplitude of the electronic-magnetic wave to improve the environment of wireless transmission, achieving a high-efficiency, low-cost, and low-energy-consumption wireless communication system. Numerous studies have investigated the applications of RIS in various network scenarios, such as millimeter-wave communications, physical layer security, mobile edge computing, multiple-input multiple-output (MIMO) systems, and non-orthogonal multiple access (NOMA) frameworks. These explorations consistently highlight the potential benefits of leveraging RIS technology in enhancing system performance [9–13]. Also, studies like [14] have proven that RIS can be used to help FL systems reduce communication errors. However, RIS-assisted networks still have challenges to face, such as network optimization and performance characterization.

Studies on RIS have also been done by many researchers. T. Jiang, et al proved that RIS can significantly reduce communication errors in an AirComp system [15]. The authors of [16] applied the method of [15] to FL and discovered the performance of the model was improved compared with no

RISs. Reference [17] combined RIS and Device Selection (DS) in the FL framework, which could quantitatively depict the effect of DS and communication error of FL. Delighted by the advantages and motivations of FL and RIS, it is urgent to jointly use them to minimize aggregation error and speed up the process of convergence. Most of the current studies use only one RIS array in the system network, but some consider using multiple RIS arrays at the same time. W. Ni, et al studied the problem of model aggregation for FL aided by multiple RISs, minimizing the mean-square-error (MSE) of the system [18]. In [19], the authors attempted to use multiple RIS arrays to improve the energy efficiency of a wireless network by modifying the reflection coefficients of the RIS arrays and adjusting the status of each RIS element dynamically.

In this work, we propose a RIS-assisted FL system with multiple RIS arrays and a channel quality-based device selection method based on prior works [20] and [21]. The proposed algorithm can select the optimal devices for the FL aggregation process in each iteration round, which helps to reduce the communication overhead in FL. This algorithm is an integrated framework that combines multiple RISs, FL, and device selection, allowing us to track the influence of aggregation errors and device selection during the training process. Simulation results verify that the proposed algorithm achieves a better performance of test accuracy and convergence speed compared to systems using only one RIS array and other device selection methods.

- We propose a RIS-assisted FL system that has multiple RIS arrays to assist the model aggregation. We formulate the expression of the model aggregation error and analyze the convergence of the model.
- We propose a channel quality based device selection algorithm to select the devices that have the best channel quality for the FL aggregation process in every iteration round to reduce the communication overhead.
- We combine multiple RIS, FL, and device selection into an integrated framework. Simulation results prove that our proposed algorithm can achieve better performance in test accuracy, training loss, and convergence speed compared to the system using only one RIS array and device selection method from the references.

The structure of this paper is as follows. “[System model](#)” introduces the system model of FL, which incorporates multiple RISs, device selection strategies, and over-the-air computation techniques. “[FL system performance analysis](#)” focuses on the performance analysis of the proposed FL approach. In “[Simulation results](#)”, we present simulation results to assess the effectiveness of the proposed methods. Finally, concluding remarks are provided in “[Conclusion](#)”.

Notations: In this paper, scalars, vectors, and matrices are represented by regular letters, bold lowercase letters, and bold uppercase letters, respectively. The transpose and conjugate transpose of a matrix are denoted as $(\cdot)^T$ and $(\cdot)^H$, respectively. The i -th entry of a vector $\mathbf{x}$ is represented as x_i , while x_{ij} refers to the (i, j) -th entry of a matrix $\mathbf{X}$. The notation $\text{diag}(\mathbf{x})$ represents a diagonal matrix where the diagonal entries correspond to $\mathbf{x}$. We denote the l_2 -norm as $\|\cdot\|_2$, and the circularly symmetric complex normal distribution with mean μ and variance σ^2 as $\mathcal{CN}(\mu, \sigma^2)$. Additionally, $\mathbb{E}[\cdot]$ is used to represent the expectation operator.

System model

Federated learning system

We consider an FL system as shown in Fig. 1 that consists of one PS and M user devices. During the federated learning process, each user receives the most recent global model from the server. Using its own local dataset, each user performs local training and transmits only the computed gradient data back to the server. The server collects gradient information from all the users participating in the current training round and uses it to update the model parameters. Once the update is complete, the server sends the improved model back to the users.

The goal is to minimize the empirical loss function, which can be expressed as:

$$\min_{\mathbf{w} \in \mathbb{R}^{D \times 1}} F(\mathbf{w}) = \frac{1}{K} \sum_{k=1}^K f(\mathbf{w}; \mathbf{x}_k, y_k), \quad (1)$$

where $\mathbf{w}$ represents the D -dimensional parameter vector of the model; K is the total number of training examples; $(\mathbf{x}_k, y_k)$ corresponds to the k -th sample, with $\mathbf{x}_k$ being the input feature vector and y_k the associated label. The loss

function for a given sample $(\mathbf{x}_k, y_k)$ is $f(\mathbf{w}; \mathbf{x}_k, y_k)$. Assume that device m holds a local dataset D_m , which contains K_m training samples. Thus, the total number of samples across all M device is $K = \sum_{m=1}^M K_m$, where $|D_m| = K_m$. Consequently, the loss function can be reformulated as:

$$F(\mathbf{w}) = \frac{1}{K} \sum_{m=1}^M K_m F_m(\mathbf{w}; D_m), \quad (2)$$

where

$$F_m(\mathbf{w}; D_m) = \frac{1}{K_m} \sum_{(\mathbf{x}_{mk}, y_{mk}) \in D_m} f(\mathbf{w}; \mathbf{x}_{mk}, y_{mk}).$$

During the FL process, user devices aim to optimize their respective local objective functions, F_m , by updating their local models. The local model updates are performed using batch gradient descent. The global model is refined iteratively through multiple rounds of training, which include the following steps:

Device selection: The PS selects a subset of devices, denoted as $\mathbf{M}_t \subseteq \{1, 2, \dots, M\}$, from the total M devices to take part in the current training round.

Model distribution: The PS transmits the latest global model, $\mathbf{w}_t$, to the devices included in $\mathbf{M}_t$.

Local gradient computation: The selected devices in $\mathbf{M}_t$ compute the gradients based on their local datasets. The local gradient for the m -th device is expressed as:

$$\mathbf{g}_{m,t} = \nabla F_m(\mathbf{w}_t; D_m), \quad (3)$$

where $F_m(\mathbf{w}; D_m)$ is the loss function evaluated for the local dataset D_m , and $\nabla F_m(\cdot)$ denotes the gradient computed at $\mathbf{w} = \mathbf{w}_t$.

Model aggregation: The $\mathbf{M}_t$ selected devices upload the gradients to the PS. After receiving the gradients, the server computes the sum of the gradients to update the global model. Every gradient vector has a different weight, which is related to K_m . Let $\mathbf{r}_t$ denote the weighted sum of the signal at the PS, which can be written as $\mathbf{r}_t = \sum_{m \in \mathbf{M}_t} K_m \mathbf{g}_{m,t}$. We can also define $\hat{\mathbf{r}}_t$ as the estimated value of $\mathbf{r}_t$. Because of the communication noise and channel fading, we know that most of the time $\hat{\mathbf{r}}_t \neq \mathbf{r}_t$. $\mathbf{w}_{t+1}$ can be updated by

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \frac{\lambda}{\sum_{m \in \mathbf{M}_t} K_m} \hat{\mathbf{r}}_t \quad (4)$$

after obtaining $\hat{\mathbf{r}}_t$, where λ denotes the learning rate and $\frac{1}{\sum_{m \in \mathbf{M}_t} K_m}$ denotes the scaling factor.

To reduce communication overhead, we choose a part of the devices instead of all of them to participate in each

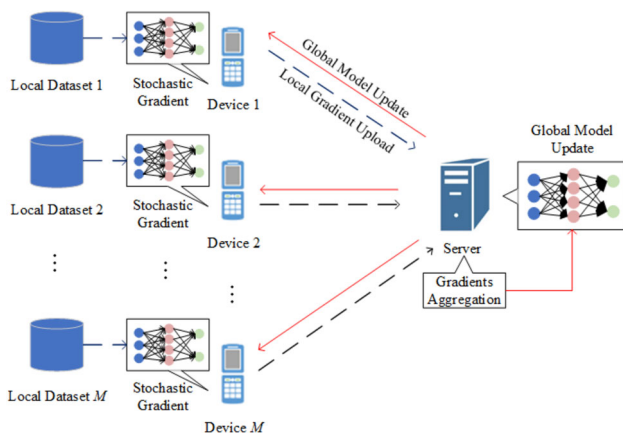


Fig. 1 Federated learning process

iteration. The details of the device selection process will be discussed in “[Device selection](#)”.

Multiple RIS system

The system network is illustrated in Fig. 2, where Z RIS arrays are used to assist the communication between user devices and the PS. Each device only has one antenna while the PS has N antennas. The number of reflecting elements is L for every RIS and the system uses a fading channel model, which has invariant channel coefficients during the training process.

Let $\mathbf{h}_{DP,m}$, $\mathbf{H}_{RP,z}$, and $\mathbf{h}_{DR,m,z}$ represent, respectively, the channel vector or matrix between the m -th device and the PS, the z -th RIS and the PS, as well as the m -th device and the z -th RIS. It is assumed that the Channel State Information (CSI) is perfectly known at the PS, the RIS controllers, and all devices. The diagonal matrix describing the z -th RIS, with dimension $L \times L$, is denoted by $\Theta_z = \text{diag}(e^{j\theta_z^1}, e^{j\theta_z^2}, \dots, e^{j\theta_z^L})$, where $\theta_z^l \in [0, 2\pi]$ corresponds to the phase shift of the l -th reflecting element on the z -th RIS.

In a multiple-RIS-assisted communication system, there will be reflected signals between RIS arrays, causing interference. It's necessary to jointly optimize the phase shifts of the RIS arrays when it comes to a real scenario, so as to increase the energy efficiency, system capacity, and the quality of the signal. However, to simplify the model, we assume the phase shift is fixed in this paper.

In the t -th training round, the signal received at the PS in time slot d can be expressed as:

$$\begin{aligned} \mathbf{y}_t[d] &= \sum_{m \in \mathbf{M}_t} \mathbf{h}_m(\Theta) x_{m,t}[d] + \mathbf{n}_t[d] \\ &= \sum_{m \in \mathbf{M}_t} (\mathbf{h}_{DP,m} + \sum_{z=1}^Z \mathbf{H}_{RP,z} \Theta_z \mathbf{h}_{DR,m,z}) x_{m,t}[d] + \mathbf{n}_t[d], \end{aligned} \quad (5)$$

where $\mathbf{n}_t[d]$ is an additive white Gaussian noise (AWGN) vector following the distribution $CN(0, \sigma_n^2)$. The channel coefficient vector between the m -th device and the PS, incorporating the RIS effects, is denoted as $\mathbf{h}_m(\Theta)$, and is defined as:

$$\mathbf{h}_m(\Theta) \triangleq \mathbf{h}_{DP,m} + \sum_{z=1}^Z \mathbf{H}_{RP,z} \Theta_z \mathbf{h}_{DR,m,z} \quad (6)$$

Meanwhile, we can obtain the power consumption of the entire system:

$$P_{\text{total}} = P_B + M \cdot P_D + Z \cdot L \cdot P_R, \quad (7)$$

where P_B , P_D , P_R are the power consumption of the BS, a device, and a RIS reflecting element.

Device selection

It is well-known that involving all devices in the iterative process results in significant communication overhead. On the other hand, selecting too few devices can slow down the convergence speed of the model. To address this trade-off, we propose a device selection method based on channel quality. This approach prioritizes devices with higher-quality channels for participation in each training round.

The channel quality for the m -th device is estimated using the l_2 norm of its corresponding channel matrix, represented as $\|\mathbf{h}_m(\Theta)\|_2$. A higher l_2 norm value indicates better channel quality. After computing the l_2 norm for all devices, the ones with the best channel quality are chosen from the total pool of M devices to take part in training. These selected devices transmit their gradient information to the PS for aggregation. The PS then updates the global model and shares it with all devices, initiating a new round of device selection.

After the device selection, we can obtain a weight coefficient for the selected devices:

$$\delta_m = \frac{\|\mathbf{h}_m(\Theta)\|_2}{\sum_{m \in \mathbf{M}_t} \|\mathbf{h}_m(\Theta)\|_2}. \quad (8)$$

Therefore, the global gradient can be expressed as:

$$\mathbf{w}_t = \frac{1}{N_t} \sum_{m \in \mathbf{M}_t} \delta_m \mathbf{w}_m \quad (9)$$

According to [22], we can define the gradient diversity as:

$$I(\mathbf{w}_t) = \frac{\sum_{m \in \mathbf{M}_t} \delta_m \|\mathbf{w}_m\|_2^2}{\|\mathbf{w}_t\|_2^2}. \quad (10)$$

It has been proved that high gradient diversity can help the SGD algorithm achieve better acceleration while maintaining the generalization performance of sequential SGD, while the lack of gradient diversity can lead to a deterioration in the performance of SGD after a certain batch size [22].

Further details about the proposed device selection algorithm are provided in “[Learning performance analysis](#)”.

Over-the-air model aggregation

AirComp is employed to facilitate the aggregation of FL models. During the t -th training round, the devices chosen from $\mathbf{M}_t$ update their respective local models. By properly tuning the transmission scaling factors, the resulting weighted sum, $\mathbf{r}_t$ can be accurately reconstructed at the PS.

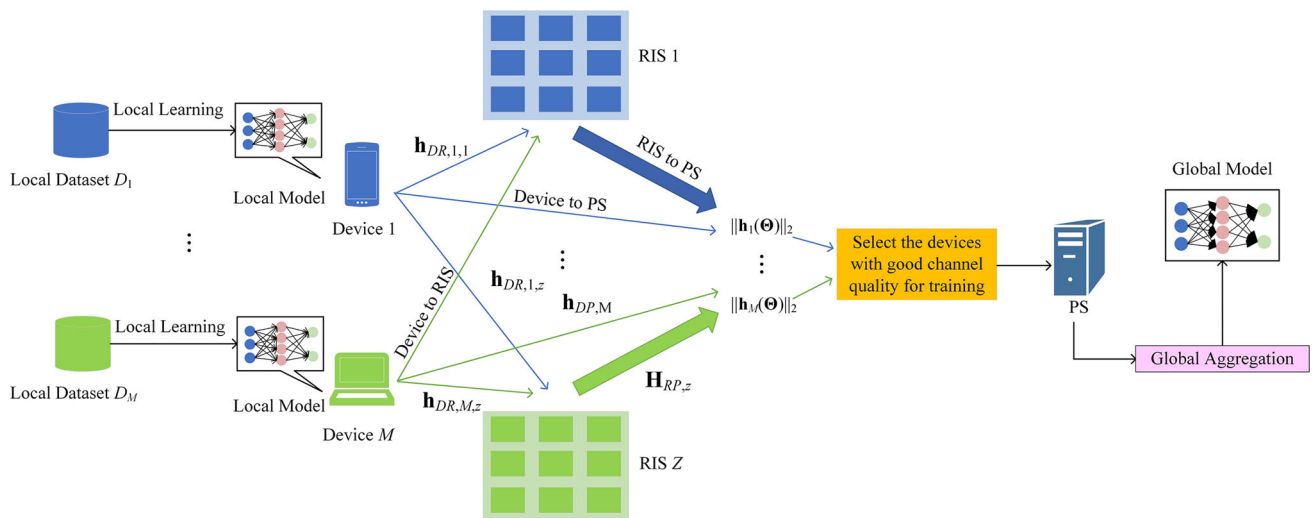


Fig. 2 Multiple RIS system

Since each round in the training process follows a similar procedure, the index t is omitted for simplicity. Let $g_m[d]$ denote the d -th entry of $\mathbf{g}_m$. Firstly, we compute the means and variances of the local gradients by

$$\bar{g}_m = \frac{1}{D} \sum_{d=1}^D g_m[d], \quad v_m^2 = \frac{1}{D} \sum_{d=1}^D (g_m[d] - \bar{g}_m)^2. \quad (11)$$

Then, the selected $\mathbf{M}_t$ devices upload the aforementioned information to the PS.

The transmit signal from the m -th device is defined as:

$$x_m = p_m s_m[d], \quad s_m[d] \triangleq \frac{g_m[d] - \bar{g}_m}{v_m}, \quad (12)$$

where p_m is the transmit equalization factor designed to compensate for channel fading and to ensure proper weighting at the PS. In the above equation, $g_m[d]$ is normalized to a zero-mean and unit-variance symbol $s_m[d]$ by leveraging $\bar{g}_m$ and v_m . This normalization guarantees that $\mathbb{E}[|x_m[d]|^2] = |p_m|^2$. Furthermore, the transmit power is constrained by:

$$\mathbb{E}[|x_m[d]|^2] = |p_m|^2 \leq P_0, \quad (13)$$

where P_0 denotes the maximum allowable transmit power.

By substituting the expression for x_m into the received signal model at the PS, we obtain:

$$\mathbf{y}[d] = \sum_{m \in \mathbf{M}_t} \mathbf{h}_m(\Theta) \frac{p_m(g_m[d] - \bar{g}_m)}{v_m} + \mathbf{n}[d]. \quad (14)$$

Next, the PS reconstructs $\hat{r}[d]$ from the received signal $\mathbf{y}[d]$ as follows:

$$\begin{aligned} \hat{r}[d] &= \frac{1}{\sqrt{\eta}} \mathbf{f}^H \mathbf{y}[d] + \sum_{m \in \mathbf{M}_t} K_m \bar{g}_m \\ &= \frac{1}{\sqrt{\eta}} \left[\sum_{m \in \mathbf{M}_t} \mathbf{f}^H \mathbf{h}_m(\Theta) \frac{p_m(g_m[d] - \bar{g}_m)}{v_m} + \mathbf{f}^H \mathbf{n}[d] \right] \\ &\quad + \sum_{m \in \mathbf{M}_t} K_m \bar{g}_m, \end{aligned} \quad (15)$$

where $\mathbf{f}$ denotes the normalized beamforming vector, satisfying $\|\mathbf{f}\|_2^2 = 1$, and η is a positive scalar for normalization. The additional term $\sum_{m \in \mathbf{M}_t} K_m \bar{g}_m$ compensates for the local mean subtraction introduced earlier in the computation of $s_m[d]$.

Integrating all the received entries $\hat{r}[1], \dots, \hat{r}[D]$ forms the vector $\hat{\mathbf{r}}$, which is then utilized to update the global model at the PS. However, the presence of channel fading and communication noise introduces unavoidable estimation errors in $\hat{\mathbf{r}}$, thereby negatively affecting the convergence of the federated learning process. The next section will specifically analyze how these communication-induced errors impact the overall convergence performance of the FL algorithm.

FL system performance analysis

Preliminaries and assumptions

The PS computes $\hat{\mathbf{r}}_t$ through model aggregation above. Accordingly, the global model is updated as:

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \frac{\lambda}{\sum_{m \in \mathbf{M}_t} K_m} \hat{\mathbf{r}}_t = \mathbf{w}_t - \lambda(\nabla F(\mathbf{w}_t) - \mathbf{e}_t), \quad (16)$$

where $\nabla F(\mathbf{w}_t) \triangleq \frac{1}{K} \sum_{k=1}^K \nabla f(\mathbf{w}_t; \mathbf{x}_k, y_k)$ represents the gradient of the global loss function $F(\mathbf{w})$ evaluated at $\mathbf{w} = \mathbf{w}_t$ and $\mathbf{e}_t$ denotes the gradient error vector, which arises from discrepancies introduced by distributed training and aggregation at the PS, which can be written as

$$\mathbf{e}_t = \underbrace{\nabla F(\mathbf{w}_t) - \frac{\mathbf{r}_t}{\sum_{m \in \mathbf{M}_t} K_m}}_{\mathbf{e}_{1,t}: \text{error from device selection}} + \underbrace{\frac{\mathbf{r}_t}{\sum_{m \in \mathbf{M}_t} K_m} - \frac{\hat{\mathbf{r}}_t}{\sum_{m \in \mathbf{M}_t} K_m}}_{\mathbf{e}_{2,t}: \text{error from fading and noise}}. \quad (17)$$

From the formula above, we can know that if $\mathbf{M}_t = \{1, 2, \dots, M\}$, $\mathbf{e}_{1,t} = 0$ and if $|\mathbf{M}_t| < M$, $\mathbf{e}_{1,t} \neq 0$. As a result, device selection introduces an error $\mathbf{e}_{1,t}$ on the gradient vector but reduces the amount of used data. These two types of errors jointly degrade the performance of FL.

Based on the findings in [17], we consider the following assumptions (A1–A4) related to the properties of the loss function $F(\cdot)$:

A1 The function F exhibits strong convexity with parameter μ . Specifically, for any $\mathbf{w}, \mathbf{w}' \in \mathbb{R}^{D \times 1}$, we have: That is, $F(\mathbf{w}) \geq F(\mathbf{w}') + (\mathbf{w} - \mathbf{w}')^T \nabla F(\mathbf{w}') + \frac{\mu}{2} \|\mathbf{w} - \mathbf{w}'\|_2^2$.

A2 The gradient $\nabla F(\cdot)$ is Lipschitz continuous with constant ω , ensuring that: $\|\nabla F(\mathbf{w}) - \nabla F(\mathbf{w}')\|_2 \leq \omega \|\mathbf{w} - \mathbf{w}'\|_2, \forall \mathbf{w}, \mathbf{w}' \in \mathbb{R}^{D \times 1}$.

A3 The function F is twice-continuously differentiable.

A4 The gradient at any training sample $\nabla f(\cdot)$, is bounded for $\mathbf{w}_t, 1 \leq t \leq T$. More specifically, $\|\nabla f(\mathbf{w}_t; \mathbf{x}_k, y_k)\|_2^2 \leq \alpha_1 + \alpha_2 \|\nabla F(\mathbf{w}_t)\|_2^2, \forall k, \forall t$, where $\alpha_1 \geq 0$ and $\alpha_2 > 0$.

These assumptions are widely accepted in stochastic optimization literature. Assumption A1 ensures that $F(\cdot)$ has a unique optimal solution $\mathbf{w}^*$. A4 provides an upper bound for the norm of the local training gradients, while the parameter α_1 in A4 permits non-zero local gradients even when the global gradient vanishes. Together, the four assumptions help establish an upper bound for the value of the loss function $F(\mathbf{w}_{t+1})$, as defined by equation (16), given that the learning rate λ is appropriately chosen.

Assuming that the loss function $F(\cdot)$ satisfies assumptions A1–A4, satisfies the conditions specified in A1–A4, the fol-

lowing inequality holds for the t -th training iteration:

$$\mathbb{E}[F(\mathbf{w}_{t+1})] \leq \mathbb{E}[F(\mathbf{w}_t)] - \frac{1}{2\omega} \|\nabla F(\mathbf{w}_t)\|_2^2 + \frac{1}{2\omega} \mathbb{E}[\|\mathbf{e}_t\|_2^2], \quad (18)$$

where ω is the Lipschitz constant defined in assumption A2.

Learning performance analysis

Given assumptions A1–A4 and the previously stated inequality, the expression for $\mathbb{E}[F(\mathbf{w}_{t+1})]$ can be determined. First, we can obtain the following inequality

$$\begin{aligned} \mathbb{E}[F(\mathbf{w}_{t+1})] &\leq \mathbb{E}[F(\mathbf{w}_t)] - \frac{1}{2\omega} \|\nabla F(\mathbf{w}_t)\|_2^2 + \frac{1}{2\omega} \mathbb{E}[\|\mathbf{e}_{1,t} + \mathbf{e}_{2,t}\|_2^2] \\ &\stackrel{(a)}{\leq} \mathbb{E}[F(\mathbf{w}_t)] - \frac{1}{2\omega} \|\nabla F(\mathbf{w}_t)\|_2^2 + \frac{1}{2\omega} \mathbb{E}[(\|\mathbf{e}_{1,t}\|_2 + \|\mathbf{e}_{2,t}\|_2)^2] \\ &\stackrel{(b)}{\leq} \mathbb{E}[F(\mathbf{w}_t)] - \frac{1}{2\omega} \|\nabla F(\mathbf{w}_t)\|_2^2 + \frac{1}{\omega} (\mathbb{E}[\|\mathbf{e}_{1,t}\|_2^2] + \mathbb{E}[\|\mathbf{e}_{2,t}\|_2^2]), \end{aligned} \quad (19)$$

where (a) is obtained using the triangle inequality, while (b) follows from the application of the AM-GM inequality. The expectations are computed with respect to the communication noise, so the expectation on $\|\mathbf{e}_{1,t}\|_2^2$ can be omitted since $\mathbf{e}_{1,t}$ is not related to the communication noise.

To further analyze $\mathbb{E}[F(\mathbf{w}_{t+1})]$, we still need expressions of $\frac{\|\mathbf{e}_{1,t}\|_2^2}{\omega}$ and $\frac{\mathbb{E}[\|\mathbf{e}_{1,t}\|_2^2]}{\omega}$. We know that $\mathbf{e}_{1,t}$ is related to the $\mathbf{M}_t$, while $\mathbf{e}_{2,t}$ is determined by both $\mathbf{M}_t$ and $\{\mathbf{f}, \Theta, \eta, p_m\}, m \in \mathbf{M}_t$. With the transmit power limit (13) and given device selection $\mathbf{M}_t$, the optimal $\{\eta, p_m, \forall m \in \mathbf{M}_t\}$ that can minimize $\mathbb{E}[\|\mathbf{e}_{2,t}\|_2^2]$ is given by

$$\eta = \min_{m \in \mathbf{M}_t} \frac{P_0 |\mathbf{f}^H \mathbf{h}_m(\Theta)|^2}{K_m^2 v_m^2}, \quad (20)$$

$$p_m = \frac{K_m \sqrt{\eta} v_m (\mathbf{f}^H \mathbf{h}_m(\Theta))^H}{|\mathbf{f}^H \mathbf{h}_m(\Theta)|^2}, \forall m \in M \quad (21)$$

Meanwhile, $\mathbb{E}[\|\mathbf{e}_{2,t}\|_2^2]$ with (20) and (21) is given by

$$\mathbb{E}[\|\mathbf{e}_{2,t}\|_2^2] = \frac{D\sigma_n^2}{P_0 \left(\sum_{m \in \mathbf{M}_t} K_m \right)^2} \max_{m \in \mathbf{M}_t} \frac{K_m^2 v_m^2}{|\mathbf{f}^H \mathbf{h}_m(\Theta)|^2}. \quad (22)$$

The above formula provides the optimal solution to $\{\eta, p_m, \forall m \in \mathbf{M}_t\}$ in (20) and (21). We define two functions

of the same variables $\{\mathbf{M}_t, \mathbf{f}, \Theta\}$ as:

$$\Phi(\mathbf{M}_t, \mathbf{f}, \Theta) = \frac{4}{K^2} \left(K - \sum_{m \in \mathbf{M}_t} K_m \right)^2 + \frac{\sigma_n^2}{P_0 \left(\sum_{m \in \mathbf{M}_t} K_m \right)^2} \max_{m \in \mathbf{M}_t} \frac{K_m^2}{|\mathbf{f}^H \mathbf{h}_m \Theta|^2}, \quad (23)$$

$$\Psi(\mathbf{M}_t, \mathbf{f}, \Theta) = 1 - \frac{\mu}{\omega} + \frac{2\mu\alpha_2 \Phi(\mathbf{M}_t, \mathbf{f}, \Theta)}{\omega}. \quad (24)$$

Based on assumptions A1–A4 and (22), we can have the bound of $\mathbb{E}[F(\mathbf{w}_{t+1}) - F(\mathbf{w}^*)]$ for any given $\{\mathbf{M}_t, \mathbf{f}, \Theta\}$ and $t = 0, 1, \dots, T-1$, that is,

$$\begin{aligned} & \mathbb{E}[F(\mathbf{w}_{t+1}) - F(\mathbf{w}^*)] \\ & \leq \frac{\alpha_1}{\omega} \Phi(\mathbf{M}_t, \mathbf{f}, \Theta) \frac{1 - (\Psi(\mathbf{M}_t, \mathbf{f}, \Theta))^t}{1 - \Psi(\mathbf{M}_t, \mathbf{f}, \Theta)} \\ & \quad + (\Psi(\mathbf{M}_t, \mathbf{f}, \Theta))^t (F(\mathbf{w}_0) - F(\mathbf{w}^*)), \end{aligned} \quad (25)$$

where $\mathbf{w}_0$ is the initial gradient of the FL model.

The formula above demonstrates that FL ensures the resulting loss remains minimized. Specifically, it shows that $\mathbb{E}[F(\mathbf{w}_{t+1}) - F(\mathbf{w}^*)]$ is bounded above by terms involving $\{\mathbf{M}_t, \mathbf{f}, \Theta\}$. However, due to factors such as communication noise, the loss introduced by device selection, and the non-strongly-convex nature of the cross-entropy function employed in the learning task, a gap persists between the converged loss and the optimal loss.

As mentioned in “Multiple RIS system”, to improve the performance of the model and decrease communication overhead, we propose a channel quality based device selection method. The steps of the proposed channel quality based DS method are shown in Algorithm 1, while the flowchart of it is shown in Fig. 3.

Algorithm 1: Channel quality based device selection

```

1: Input:  $M, \mathbf{h}(\Theta)$ 
2: Initialization: The set of the indexes of selected devices is
    $\mathbf{M}_t = \{0, 0, 0, 0, 0\}$ , the set of the  $l_2$  norms of the corresponding
   devices is  $\mathbf{N}_t = \{0, 0, 0, 0, 0\}$ 
3: for  $m = 1, 2, \dots, M$  do
4:   Compute the  $l_2$  norm of the  $m$ -th  $\mathbf{h}(\Theta)$  by  $\|\mathbf{h}_m(\Theta)\|_2$ 
5:   Compare the size of  $\|\mathbf{h}_m(\Theta)\|_2$  and elements in  $\mathbf{N}_t$ 
6:   if  $\|\mathbf{h}_m(\Theta)\|_2$  is larger than the minimum element  $n$  in  $\mathbf{N}_t$  then
7:     Replace  $n$  with  $\|\mathbf{h}_m(\Theta)\|_2$  in  $\mathbf{N}_t$ 
8:     Renew the index of  $n$  with  $m$  in  $\mathbf{M}_t$ 
9:   end if
10: end for
11: Output:  $\mathbf{M}_t$ 

```

Table 1 Simulation parameters

Parameters	Values
N (number of BS antennas)	5
L (size of RIS arrays)	10, 20, 30, 40
Z (number of RIS arrays)	1, 2, 3, 4
M (number of devices)	40
P_0 (maximum transit power)	− 10 dB
Modulation	16QAM
f_c (carrier frequency)	915MHz

Simulation results

Simulation settings

To assess the effectiveness of the proposed algorithm for image classification, we use the Fashion-MNIST dataset and the CIFAR-10 dataset. The Fashion-MNIST dataset comprises 50,000 independent and identically distributed (i.i.d.) training images and 10,000 test images, and the CIFAR-10 dataset includes the same number of images as the Fashion-MNIST dataset. We conduct experiments under both i.i.d. cases and non-i.i.d. cases. For i.i.d. cases, we use the Fashion-MNIST dataset and each device can access all the classes. For non-i.i.d. cases, we use the CIFAR-10 datasets but each device only has access to one class. The performance of FL is measured using test accuracy, calculated as the proportion of correctly classified images in the test set relative to the total number of test images. This metric produces a value within the range of 0 to 1. Additionally, the training loss is employed to examine how well various methods perform on the training data (Table 1).

For training our convolutional neural network, we design the architecture with two convolutional layers of size $5*5$, each followed by $2*2$ max pooling. These are succeeded by a batch normalization layer, a fully connected layer, a ReLU activation function, and a softmax output layer. The cross-entropy loss is used as the objective function during training. In a three-dimensional coordinate system, we place the PS at (0, 0, 0) and the RIS arrays at (20, 20, 20), (− 20, 20, 20), (− 20, − 20, 20), (20, − 20, 20), respectively. In Table 2, some of the default values of the system parameters are provided, and Fig. 4 shows the top view of the position of the RIS arrays and the BS.

Simulations on RIS array numbers

To analyze the influence of the number of RIS arrays on the model’s performance in both i.i.d. cases and non-i.i.d. cases, we present the test accuracy and training loss under varying numbers of RIS arrays in Figs. 5 and 6, respectively.

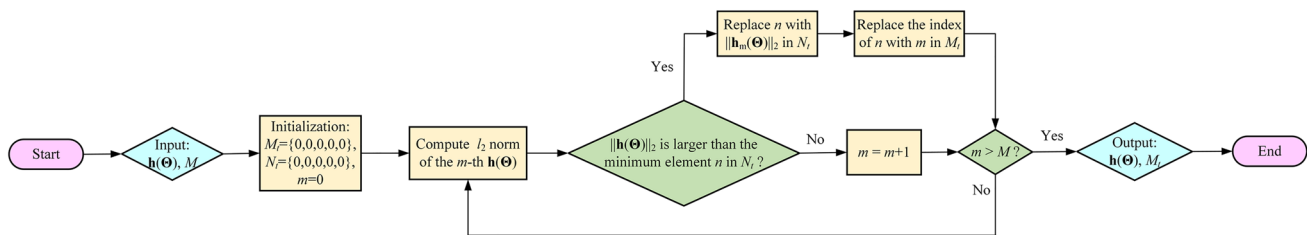


Fig. 3 Flowchart of Algorithm 1

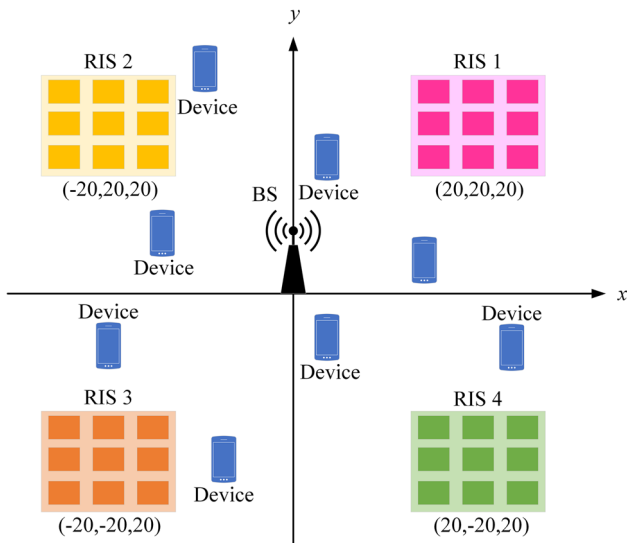


Fig. 4 Simulation setup of the system (top view)

Simulations on device selection methods

In this subsection, we examine the effects of different device selection strategies on model performance. The test accuracy and training loss for various selection methods in i.i.d. cases and non-i.i.d. cases are illustrated in Figs. 7 and 8, where we set the number of reflecting elements to 64 and select 10 devices for training in each method. Device selection allows the model to transmit less data, so the convergence process might be sped up.

As shown in Figs. 7 and 8, applying the proposed channel quality-based device selection algorithm in both cases leads to improvements in test accuracy and convergence speed compared to the method in [5] and [14], for both multiple RIS and single RIS scenarios, while also reducing the training loss. Meanwhile, in a single RIS case, the effect of the DS algorithm is more obvious than in a multiple RIS case. However, the training loss fluctuates when the training epochs increase to 400, but is still on a trend of descending.

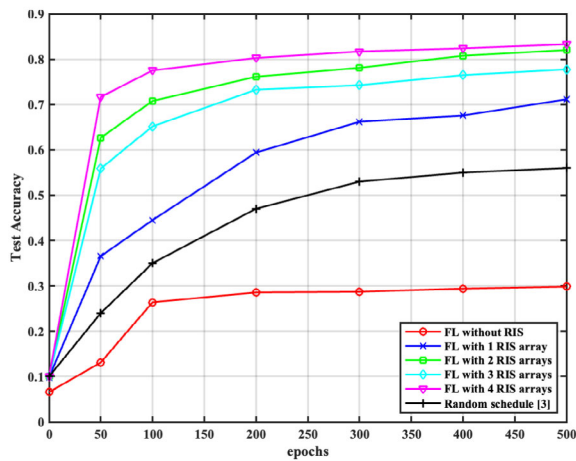
The system is configured with 64 reflecting elements, and the same device selection strategy as described in [14] is utilized for fair comparison.

As presented in Figs. 5 and 6, it can be seen that without RIS, the model cannot reach an acceptable test accuracy, whether in an i.i.d. case or not. With one RIS array added to the system, the test accuracy can increase dramatically. And with multiple RIS arrays, the test accuracy can increase more while the training loss decreases. Meanwhile, using multiple RIS arrays can also speed up the process of convergence significantly. The random schedule approach in [3], though having a better performance than the no-RIS one, failed to achieve a high level of accuracy because of the missing of RIS in i.i.d. cases, while its performance is similar to the single-RIS one in non-i.i.d. cases. The no RIS case and the random schedule approach are not shown in Fig. 5b, for the training loss becomes too large when the training epochs increase. It can be concluded that the number of RIS arrays has a great influence on the test accuracy, training loss, and convergence speed of the model.

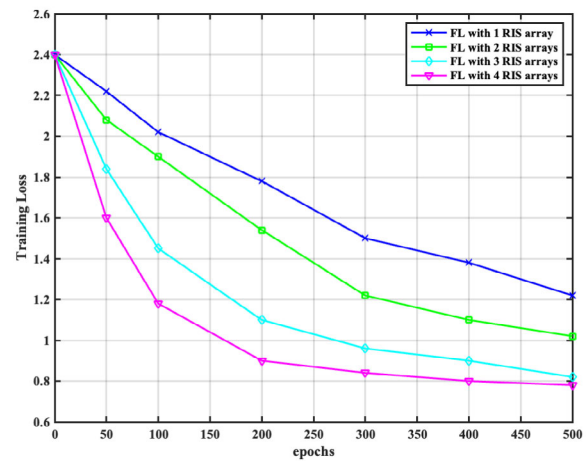
Simulations on RIS reflecting element numbers

As for the effect of RIS reflecting element numbers on the model performance, we plot the test accuracy and training loss with different RIS reflecting element numbers in Figs. 9 and 10 for i.i.d. cases and non-i.i.d. cases, respectively. We set the training epochs to 100 and the number of selected devices to 10 for each method.

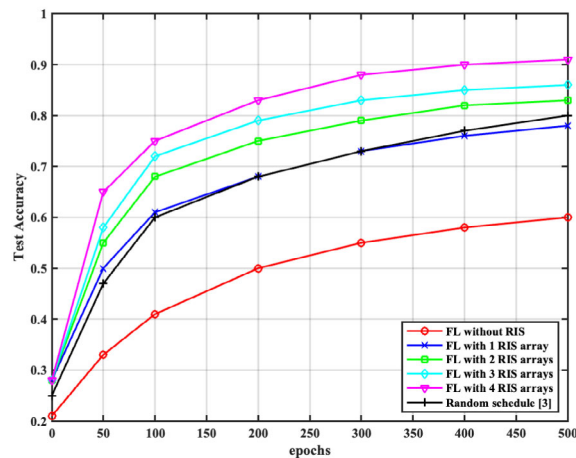
We can see from Figs. 9 and 10 that the model can reach higher test accuracy and lower training loss with more reflecting elements. In i.i.d. cases, the single RIS case with DS in [14] almost converges when the number of reflecting elements reaches 20, while the accuracy is still at an unacceptable rate of 40%. The differential privacy method [6] converges more slowly but can reach a better level of accuracy than the DS algorithm in [14]. Meanwhile, similar to the previous subsection that shows the relation between test accuracy and iteration numbers the effect of the DS algorithm in a single RIS case is more evident than in a multiple RIS case.



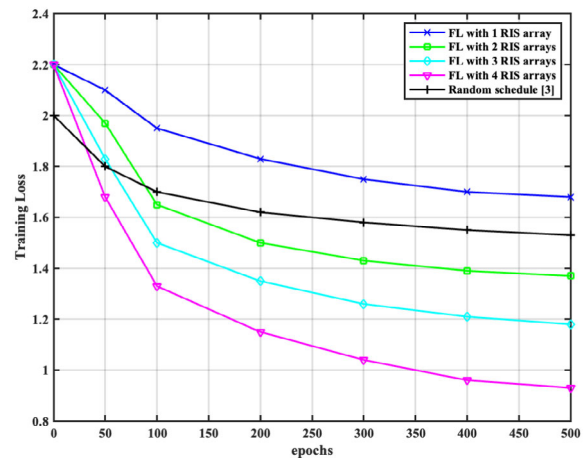
(a) Test accuracy with RIS array numbers



(b) Training loss with RIS array numbers

Fig. 5 Comparative analysis of RIS array numbers in i.i.d. cases

(a) Test accuracy with RIS array numbers



(b) Training loss with RIS array numbers

Fig. 6 Comparative analysis of RIS array numbers in non-i.i.d. cases

Simulations on number of selected devices

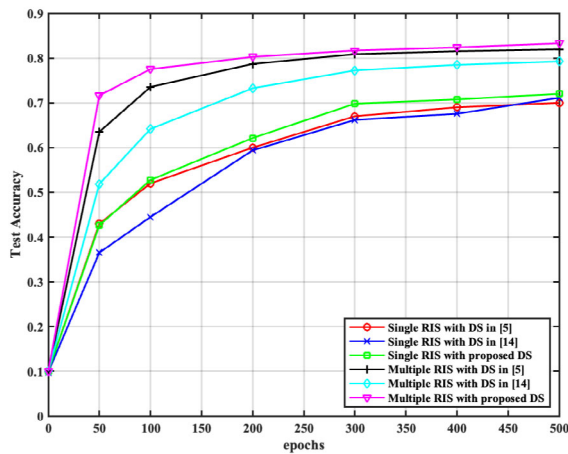
To research how the numbers of selected devices affect the model performance, we plot the test accuracy and training loss with different device selection numbers in Figs. 11 and 12 for i.i.d. cases and non-i.i.d. cases, respectively. We set the training epochs to 100 and the number of reflecting elements to 64.

As illustrated in Figs. 11a and 12a, the model achieves higher test accuracy when fewer devices are selected. This improvement stems from the enhanced channel quality associated with a smaller pool of selected devices. The FedProx method in [23] is easily affected by the number of devices, as the accuracy decreases rapidly when the number increases. In comparison, for single and multiple RIS configurations, both

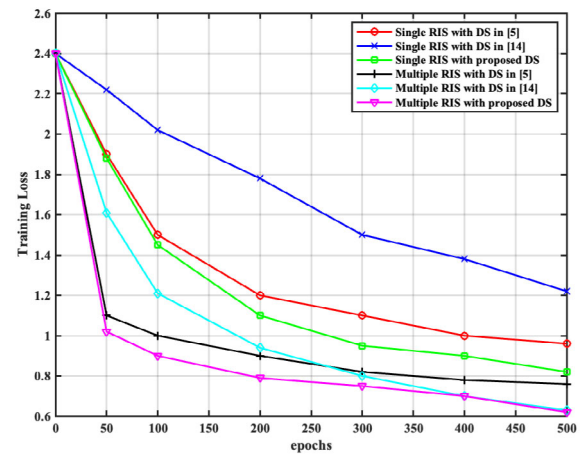
test accuracy and training loss exhibit smooth decreases and eventually stabilize as the number of selected devices grows. These simulation results further confirm that incorporating RIS arrays significantly aids in reducing the model's training loss and maintaining it within an appropriate range.

Comparison of rounds required for convergence

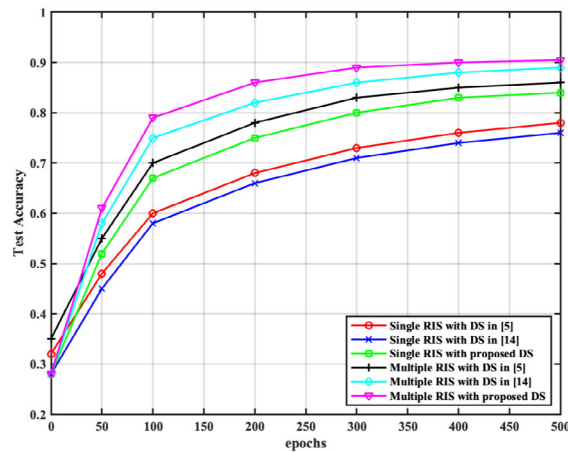
Finally, we compare the communication rounds needed for the convergence of several methods. We set the RIS array number to 4, the number of reflecting elements to 64, and the number of devices to 10. We compare the number of communication rounds required by each method to achieve 80% test accuracy of the classification task of the CIFAR-10 dataset, and Table 2 shows the results.



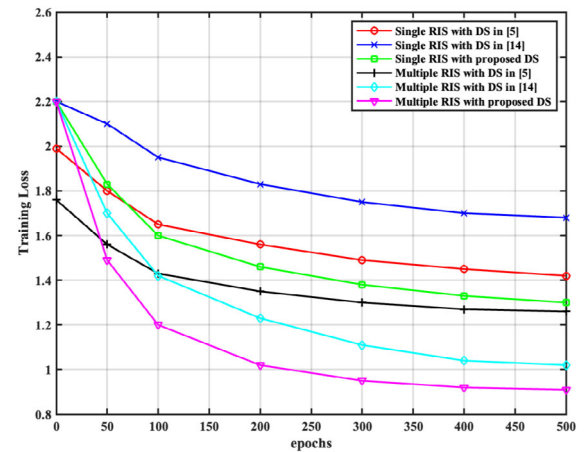
(a) Test accuracy with device selection methods



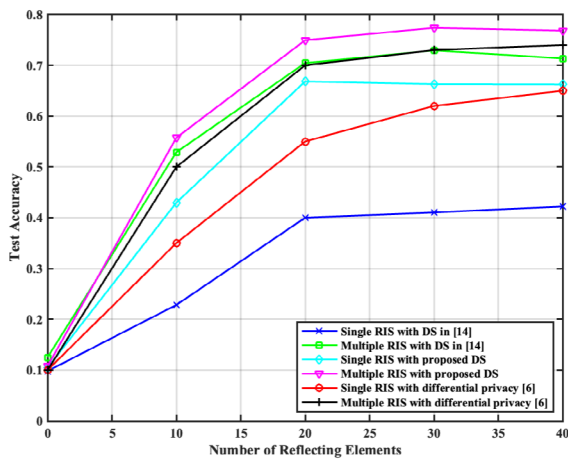
(b) Training loss with device selection methods

Fig. 7 Comparative analysis of device selection methods in i.i.d. cases

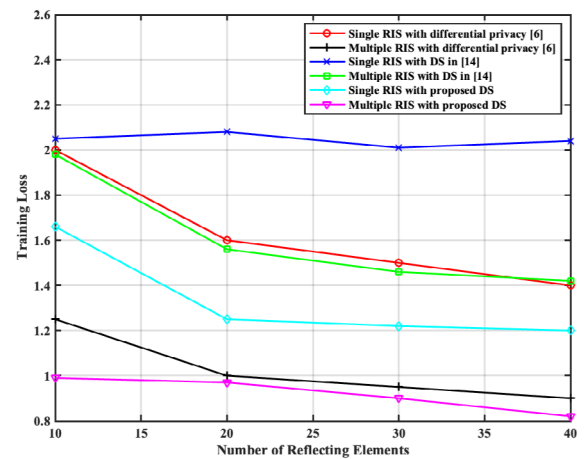
(a) Test accuracy with device selection methods



(b) Training loss with device selection methods

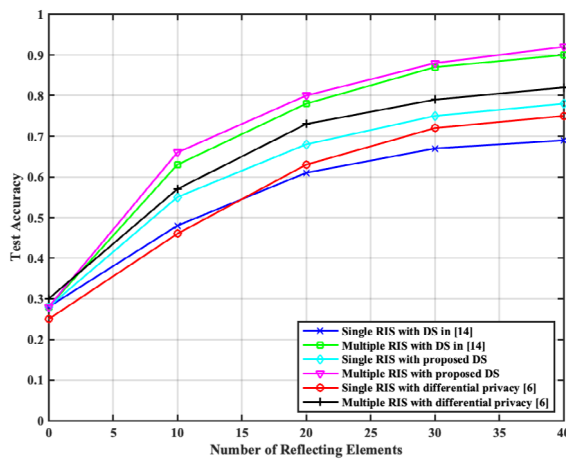
Fig. 8 Comparative analysis of device selection methods in non-i.i.d. cases

(a) Test accuracy with RIS reflecting element numbers

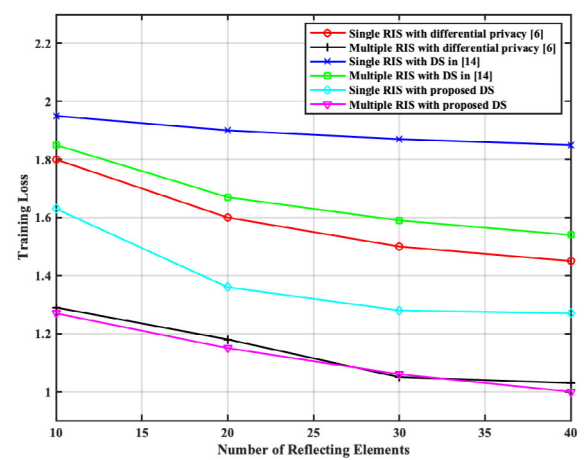


(b) Training loss with RIS reflecting element numbers

Fig. 9 Comparative analysis of RIS reflecting element numbers in i.i.d. cases

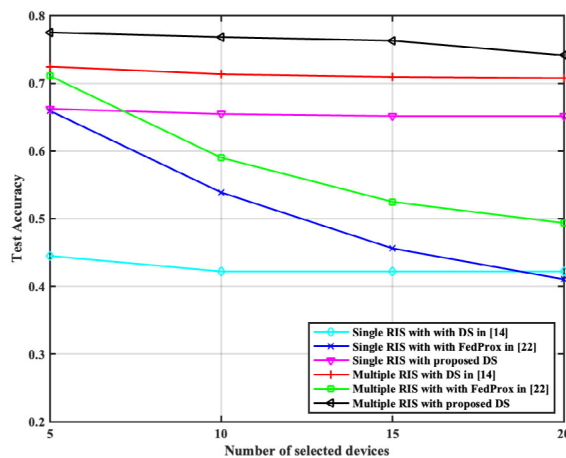


(a) Test accuracy with RIS reflecting element numbers

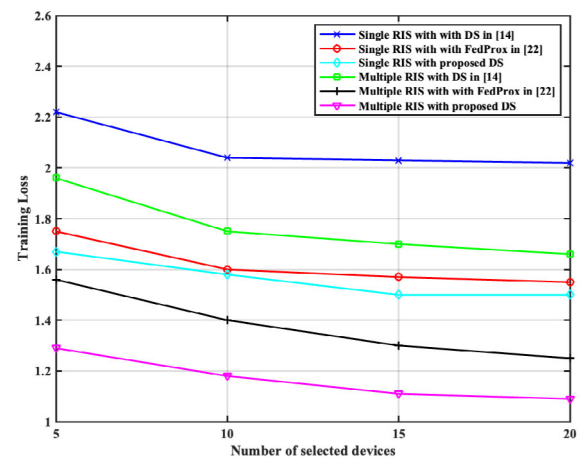


(b) Training loss with RIS reflecting element numbers

Fig. 10 Comparative analysis of RIS reflecting element numbers in non-i.i.d. cases



(a) Test accuracy with the number of selected devices



(b) Training loss with the number of selected devices

Fig. 11 Comparative analysis of the numbers of selected devices in i.i.d. cases

We can see from Table 2 that the proposed method requires a larger number of communication rounds to achieve the requested test accuracy compared to the one using differential privacy and the FedProx method, though its performance is better than random scheduling and the model-update-based method.

Limitations and future work

Although our proposed method outperforms the benchmark designs in the simulation, there are certain limitations concerning its implementation in a practical environment. Firstly, in this paper we assume that the CSI at the BS and the RIS is perfect, and ignore the reflecting signals between the RIS arrays. However, in a real case that utilizes multiple

Table 2 Communication rounds for different methods

Methods	Communication rounds
Random scheduling [3]	124
Model-update-based method [5]	110
differential privacy [6]	92
FedProx [23]	78
Proposed method	105

RIS arrays, the CSI is hardly perfect and can be influenced by estimation errors, feedback overhead, and the reflecting signals. The environment for communication will be more complicated than the one we studied in this paper. Meanwhile, we only consider 40 devices in the simulation, while

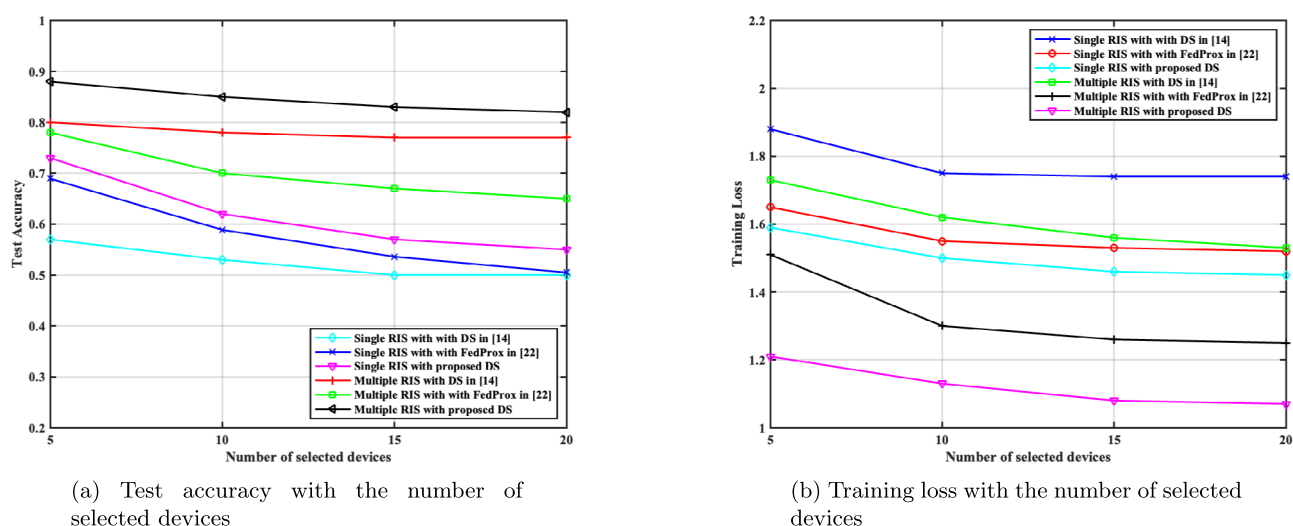


Fig. 12 Comparative analysis of the numbers of selected devices in non-i.i.d. cases

the number of communication devices can scale to thousands in practical deployments. In that case, the complexity of our proposed algorithm will not be low enough for all the devices, which may lead to a high overhead and time consumption. Besides, the proposed algorithm only considers the quality of the channel, neglecting the quality of the data. If a device has poor data quality and good channel quality, the convergence of the global model may be affected.

For future work and improvement, we plan to explore a more practical transmitting environment of a multiple-RIS-assisted communication system, taking the channel estimation errors and reflecting signals into account. The challenge of robust system design under such imperfections is widely recognized across various fields; for instance, in 3D modeling and GPS localization, researchers have developed sophisticated techniques to identify and mitigate errors caused by signal degradation and reflections, ensuring accurate global alignment and positioning in complex urban environments [24, 25]. Furthermore, we will explore the possibility of employing heuristic techniques to estimate channel quality without evaluating all devices explicitly, thereby reducing the overall complexity. Also, discussing the potential of adaptive device selection strategies that could dynamically adjust the number of devices selected based on network conditions is another future research direction. Lastly, we may consider establishing a weighted selection strategy that balances both channel quality and data quality. This would help prioritize devices with high-quality data, even in cases where their channel conditions are relatively poor.

Conclusion

In this paper, we studied the federated learning task in a RIS-assisted system. We proposed a multiple RIS-assisted FL system and a channel quality based device selection algorithm to increase the test accuracy and convergence performance of the model. Simulation results prove the advantages of the combination of multiple RIS and the proposed DS algorithm. For further studies in the future, we plan to conduct a research in a more practical environment of multiple RIS-assisted federated learning systems, which is improvement of this paper.

Author Contributions Yujun Cai: simulation and manuscript drafting. Shufeng Li and Jianbo Liu: conceptualization and supervision. Xinruo Zhang and Xiaowei Shen: manuscript revision and improvement.

Availability of data and material Enquiries about data availability should be directed to the authors. The datasets are Fashion-MNIST and CIFAR-10, and can be downloaded by torchvision.datasets in Python, and on the website of <https://www.cs.toronto.edu/~kriz/cifar.html>, respectively.

Code availability The software application is Pycharm.

Declarations

Conflict of interest There are no conflict of interest or personal relationships that could have appeared to influence the work reported in this paper.

Ethical approval Not applicable.

Open Access This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative

Commons licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

References

- McMahan B, Moore E, Ramage D, Hampson S, y Arcas BA (2017) Communication efficient learning of deep networks from decentralized data. In: Proceeding of the 20th international conference on artificial intelligence and statistics, pp 1273–1282
- Cai Y, Li S, Xia Z, Tan Y (2022) Research of federated learning communication data transmission algorithm based on user grouping policy. In: Proceeding of the 8th international conference on control science and systems engineering, pp 170–174
- Yang HH, Liu ZZ, Quek TQS, Poor HV (2020) Scheduling policies for federated learning in wireless networks. *IEEE Trans Commun* 20:317–333
- Amiri MM, Kulkarni SR, Poor HV (2021) Federated learning with downlink device selection. In: Proceeding of the IEEE 22nd international workshop on signal processing advances in wireless communications, pp 306–310
- Ma X, Sun H, Wang Q, Hu RQ (2021) User scheduling for federated learning through over-the-air computation. In: Proceeding of the IEEE 94th vehicular technology conference, pp 1–5
- Ren Y, Liu R, Wang D, Yuan L, Shen Y, Wu G, Wang Q, Yang C (2023) A study of local differential privacy mechanisms based on federated learning. *J Electr Inform Technol* 45:784–792
- Wang Z, Yu X, Wang H, Xue P (2024) A federated learning scheme for hierarchical protection and multiple aggregation. *Comput Electr Eng* 117:109240
- Yuan X, Zhang YJA, Shi Y, Yan W, Liu H (2021) Reconfigurable-intelligent-surface empowered wireless communications: challenges and opportunities. *IEEE Wirel Commun* 28:136–143
- Hu L, Wang Z, Zhu H, Zhou Y (2022) Ris-assisted over-the-air federated learning in millimeter wave mimo networks. *J Commun Inform Netw* 7:145–156
- Zhu Z, Xu J, Sun G, Wang N, Hao W (2021) Secure beamforming design for irs-assisted swipt internet of things system. *J Commun* 42:185–193
- Bai T, Pan C, Han C, Hanzo L (2021) Reconfigurable intelligent surface aided mobile edge computing. *IEEE Wirel Commun* 28:80–86
- Salen MA, Hady AA, Hussein HH, Roshdy RA (2024) Enhanced connectivity through double ris uplink mimo and advanced modulation. *Comput Electr Eng* 119:109486
- Gong C, Dai X, Cui J, Long K (2023) Performance analysis of distributed reconfigurable intelligent surface aided noma systems. *Wirel Pers Commun* 131:217–231
- Wang J, Tang W, Han YU, Jin S, Li X, Wen CK (2021) Interplay between ris and ai in wireless communications: fundamentals, architectures, applications, and open research problems. *IEEE J Sel Areas Commun* 39:2271–2288
- Jiang T, Shi Y (2019) Over-the-air computation via intelligent reflecting surfaces. In: Proceeding of IEEE global communication conference, pp 1–6
- Yang K, Shi Y, Zhou Y, Yang Z, Fu L, Chen W (2020) Federated machine learning for intelligent iot via reconfigurable intelligent surface. *IEEE Netw* 34:16–22
- Liu H, Yuan X, Zhang YJA (2021) Reconfigurable intelligent surface enabled federated learning: a unified communication-learning design approach. *IEEE Trans Wirel Commun* 20:7595–7609
- Ni W, Liu Y, Tian H (2020) Intelligent reflecting surfaces enhanced federated learning. In: Proceeding of IEEE global communication conference workshops, pp 1–6
- Le QN, Nguyen VD, Dobre OA, Zhao R (2021) Energy efficiency maximization in ris-aided cell-free network with limited backhaul. *IEEE Commun Lett* 25:1974–1978
- Wu Q, Zhang R (2019) Intelligent reflecting surface enhanced wireless network via joint active and passive beamforming. *IEEE Trans Wirel Commun* 18:5394–5409
- Huang C, Zappone A, Alexandropoulos GC, Debbah M, Yuen C (2019) Reconfigurable intelligent surfaces for energy efficiency in wireless communication. *IEEE Trans Wirel Commun* 18:4157–4170
- Yin D, Pananjady A, Lam M, Papailiopoulos D, Ramchandran K, Bartlett P (2018) Gradient diversity: a key ingredient for scalable distributed learning. In: Proceedings of the 21st international conference on artificial intelligence and statistics, vol 84. PMLR, pp 1998–2007
- Li T, Sahu AK, Zaheer M, Sanjabi M, Talwalkar A, Smith V (2020) Federated optimization in heterogeneous networks. *Proc Mach Learn Syst* 2:429–450
- Kumar A, Banno A, Ono S, Oishi T, Ikeuchi K (2013) Global coordinate adjustment of the 3d survey models under unstable gps condition. *Seisan Kenkyu* 65:91–95
- Kumar A, Sato Y, Oishi T, Ono S, Ikeuchi K (2014) Improving gps position accuracy by identification of reflected gps signals using range data for modeling of urban structures. *Seisan Kenkyu* 66:101–107

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.