

Strategic Data Revocation in Federated Unlearning

Ningning Ding

Northwestern University, USA
ningning.ding@northwestern.edu

Ermin Wei

Northwestern University, USA
ermin.wei@northwestern.edu

Randall Berry

Northwestern University, USA
rberry@northwestern.edu

Abstract—By allowing users to erase their data’s impact on federated learning models, federated unlearning protects users’ right to be forgotten and data privacy. Despite a burgeoning body of research on federated unlearning’s technical feasibility, there is a paucity of literature investigating the considerations behind users’ requests for data revocation. This paper proposes a non-cooperative game framework to study users’ data revocation strategies in federated unlearning. We prove the existence of a Nash equilibrium. However, users’ best response strategies are coupled via model performance and unlearning costs, which makes the equilibrium computation challenging. We obtain the Nash equilibrium by establishing its equivalence with a much simpler auxiliary optimization problem. We also summarize users’ multi-dimensional attributes into a single-dimensional metric and derive the closed-form characterization of an equilibrium, when users’ unlearning costs are negligible. Moreover, we compare the cases of allowing and forbidding partial data revocation in federated unlearning. Interestingly, the results reveal that allowing partial revocation does not necessarily increase users’ data contributions or payoffs due to the game structure. Additionally, we demonstrate that positive externalities may exist between users’ data revocation decisions when users incur unlearning costs, while this is not the case when their unlearning costs are negligible.

Index Terms—Federated unlearning, data revocation strategies, game theory, allowing or forbidding partial revocation

I. INTRODUCTION

A. Background and Motivations

Federated unlearning is an emerging technique designed to erase the influence of certain users’ data from a trained federated learning model. This is motivated by the success of federated learning over large numbers of edge users and the recent regulations guaranteeing data owners’ rights to be forgotten.

Federated learning is a distributed machine learning paradigm, where many users collaboratively train a shared learning model under a server’s coordination. Although federated learning helps protect privacy by allowing distributed training without sharing users’ raw data, the trained model can still leak users’ information (e.g., through a backdoor attack) [1], [2]. This motivates recent regulations to protect users’ data privacy and their rights to revoke their data from a trained model. Examples include the European Union General Data Protection Regulation (GDPR) [3] and the California Consumer Privacy Act (CCPA) [4]. A naive strategy to accomplish this is to retrain the model from scratch using the remaining users’ data, which is time-consuming and computationally

expensive. Therefore, federated unlearning methods focus on removing the influence of revoked data without retraining the whole model from scratch. Typically, this process needs the staying users to perform additional calculations based on the learned model and remaining data to obtain an unlearned model [5].

Given the right to be forgotten, users may seek to revoke their data from the learned model for a variety of reasons; a few examples follow. First, if users perceive their benefits from the learned model as insufficient to offset their costs, they may revoke their data. Second, users typically lack full knowledge of the associated benefits (e.g., model performance) and costs (e.g., privacy costs tied to data uniqueness) until they participate in federated learning [6]. Consequently, if the actual outcomes deviate from their initial estimations, users may request their data to be unlearned. Third, users can mitigate or even eliminate certain costs (e.g., data privacy costs) by revoking their data after participating in federated learning [7].

It is challenging for users to optimally decide whether to revoke data and how much data to revoke. Users need to make complex assessments of the benefits and costs associated with their actions. These considerations include the local performance of the trained model, their potential privacy leakage from participation, and federated unlearning costs if they stay in the system (i.e., additional training for obtaining an unlearned model). Moreover, users’ strategies for data revocation indirectly impact each other’s payoffs. For example, each user’s data revocation action will affect the performance of the unlearned global model and the required efforts for federated unlearning (e.g., required training rounds for users). It is hard for each user to precisely anticipate such impacts and make optimal decisions, especially given a large number of heterogeneous users with multi-dimensional attributes on data and payoffs. This motivates our first key question:

Question 1. *What are heterogeneous users’ optimal data revocation strategies in federated unlearning?*

Furthermore, existing federated unlearning literature (e.g., [5], [8], [9]) usually operated under the assumption that users either completely leave the system (i.e., revoke all their data) or remain entirely committed (i.e., do not revoke any data). Partial data revocation remains relatively uncharted territory. However, this may harbor the potential to increase the total data size remaining in the system and users’ payoffs, compared to the scenario where partial data revocation is forbidden. This lead to the second key question in this paper:

This work is supported by NSF ECCS-2030251 and NSF ECCS-2216970.

Question 2. Compared with forbidding partial data revocation, is allowing partial revocation more beneficial in terms of users' payoffs and total remaining data size?

B. Contributions

We summarize our key contributions below.

- *Data revocation in federated unlearning:* We propose a game-theoretic model to characterize the interaction among users in federated unlearning. To the best of our knowledge, this is the first analytical work to study users' strategic data revocation in federated unlearning, with partial data revocation allowed.
- *Users' data revocation strategies at equilibrium:* Given the complicated mutual influence of users' decisions through model performance and unlearning costs, the problem is challenging to analyze (e.g., with non-convex, non-monotonic, and piece-wise best responses). We calculate the equilibrium through a problem transformation that applies to general payoff functions. Moreover, when users' unlearning costs are negligible, we summarize users' multi-dimensional attributes into a single-dimensional metric that enables us to give a closed-form equilibrium.
- *Comparison between allowing and forbidding partial data revocation:* We investigate whether allowing users' partial revocation is better than when it is forbidden. The results counter-intuitively show that allowing partial revocation may not increase users' remaining data amount or payoffs because of the strategic interactions among users.
- *Insights:* Our results reveal trade-offs among model performance, data privacy concerns, and federated unlearning costs, highlighting how these are influenced by factors such as model dependency, data size, and data uniqueness. We also show that users may follow the others' data revocation decisions (i.e., a positive externality) when they incur unlearning costs but will not (i.e., a negative externality) when unlearning costs are negligible. This is because the model performance increases in users' remaining data sizes while unlearning costs increase in users' revoked data sizes.

C. Related Work

Machine unlearning was first introduced in 2015 [10]. It focuses on how to efficiently erase the effects of certain data on a *centrally* trained machine learning model. The current machine unlearning methods often involve properly modifying training data or altering the trained model parameters to reduce the model's reliance on the deleted information (e.g., [11]). However, the server in a federated network has no access to the local datasets of a large number of users, which makes data modification and manipulation impossible. Moreover, users in a federated system contribute to the final model through iterative training processes, which makes it not straightforward to partition data impact and alter model parameters.

Some pioneering works have recently proposed federated unlearning mechanisms using methods such as gradient subtraction (e.g., [5], [8]), gradient scaling (e.g., [7], [12]), or knowledge distillation (e.g., [13]). These papers did not study

which user(s) will revoke data and how much data will be revoked. This is of great importance in understanding users' behaviors and developing related mechanisms or regulations. To fill this gap, we focus on users' data revocation strategies in this paper.

Furthermore, there is a wide spectrum of literature on studying (or incentivizing) users' strategic data contributions in federated learning, utilizing game theory or other economic methodologies (e.g., [14]–[18]). However, these studies did not incorporate the unique aspects of federated unlearning (e.g., unlearning costs), so it is hard to apply their results to our context. Ding et al. in [9] proposed an incentive mechanism for federated unlearning, aimed at retaining strategic users intending to leave the system. Nevertheless, their assumption of all-or-nothing user engagement fails to account for scenarios where users may wish to partially revoke their data. They also assumed that users do not care about the global model's performance. To the best of our knowledge, this paper is the first to study users' strategic data revocation in federated unlearning, considering a more general scenario with users' partial data revocation and interests in the model performance.

The rest of the paper is organized as follows. We first introduce the system model in Section II. In Section III, we analyze users' equilibrium strategies on data revocation. To give more insights, we study a special scenario in Section IV, where users' unlearning costs are negligible. In Section V, we investigate whether allowing partial revocation is beneficial compared with when it is forbidden. We present simulation results in Section VI and conclude this paper in Section VII.

II. SYSTEM MODEL

We consider a setting in which a set of users have already participated in federated learning using their local data and are faced with a federated unlearning decision. The heterogeneous users now simultaneously decide the amount of data to revoke from the learned model and collaboratively accomplish the unlearning process. We first introduce the objectives and process of federated learning and unlearning, and then we specify users' strategies and payoffs, respectively. Finally, we frame a non-cooperative game model, highlighting the users' participation within the system.

A. Federated Learning and Federated Unlearning

1) *Federated Learning:* Consider a set $\mathcal{I} = \{1, 2, \dots, I\}$ of users. The purpose of federated learning is to compute a model parameter w by using all users' local data. The optimal model parameter w^* minimizes the global loss function, which is an average of all users' loss functions $\{F_i(w)\}_{i \in \mathcal{I}}$ [19], [20]:

$$w^* = \arg \min_w \frac{1}{I} \sum_{i \in \mathcal{I}} F_i(w). \quad (1)$$

During the federated learning process, each user computes a local model update based on local data and sends the model update to a central server. The server aggregates the local updates and sends the global shared model back to the users for the next round of training. The process is then repeated until the global model converges [21].

2) *Federated Unlearning*: A federated learning process maps users' data into a model space, while a federated unlearning process maps a learned model, users' data set, and the data set that is required to be forgotten into an unlearned model space. The goal of federated unlearning is to make the unlearned model have the same distribution as a retrained model (i.e., retrained from scratch using the remaining data).¹

A natural method for federated unlearning is to let the staying users (excluding leaving users) use the remaining data to continue training from the learned model w^* . The training continues until it converges to a new optimal model parameter $\tilde{w}^*$, which minimizes the global loss function of staying users:

$$\tilde{w}^* = \arg \min_w \frac{1}{I - I_{\text{leave}}} \sum_{i \in \mathcal{I} \setminus \mathcal{I}_{\text{leave}}} \tilde{F}_i(w), \quad (2)$$

where $\mathcal{I}_{\text{leave}}$ is the set of users who leave the system through federated unlearning. This method is typically more efficient than training from scratch, as the minimum point may not change much after some users leave. As in (2), the unlearned model greatly depends on users' data revocation strategies.

B. Users' Strategies

We use $d_i^{\max} > 0$ to denote the size of data that a user $i \in \mathcal{I}$ has previously utilized in training a federated learning model. Thus, $d_i^{\max}$ is the maximum data size that user i can revoke through the unlearning phase.

Each user i decides the data size $d_i \in [0, d_i^{\max}]$ to remain in the system after revocation. In other words, the data size that user i revokes through federated unlearning is $d_i^{\max} - d_i$.² We assume that the revoked data is uniformly sampled at random from $d_i^{\max}$, so that the remaining data distribution does not change. Users' different data revocation strategies will inevitably result in diverse payoffs.

C. Users' Payoffs

Each user i 's payoff from data revocation consists of three parts.

1) *Global Model's Performance*: To capture the global model's performance on user i 's local data, we employ a widely-adopted model that provides a good approximation of the experimental statistics (e.g., [14], [22]–[24]):

$$P_i(d_i, \mathbf{d}_{-i}) = \ln \left(\sum_{j \in \mathcal{I}} d_j + \epsilon_i \right), \quad (3)$$

where $\mathbf{d}_{-i} \triangleq \{d_j\}_{j \in \mathcal{I}, j \neq i}$ denotes the strategies of other users except for user i . The logarithmic concave model captures the diminishing marginal utility of additional data. Intuitively, the model performance increases in the total data size remaining in the system. However, when the remaining data size is already substantial, any further increase in the data size does not

significantly improve the model performance. The factor ϵ_i captures the heterogeneity in model performance and bounds user i 's cost when all of the users' data is revoked, i.e. $\sum_{j \in \mathcal{I}} d_j = 0$. We refer to ϵ_i as the *independence index*, where a small value (e.g., $\epsilon_i \rightarrow 0$) implies that the user has a high dependency on the global model.

2) *Privacy Cost*: A user i 's perceived data privacy cost increases in its remaining data size and the uniqueness of its data (e.g., [25], [26]):

$$C_i^p(d_i) = \xi_i d_i \ell_i, \quad (4)$$

where ξ_i is user i 's *marginal privacy cost* and ℓ_i is user i 's *data uniqueness level*. The data uniqueness level can be characterized by the computed gradient $\|\nabla F_i(w^*)\|$, as the gradient reflects the distance of user i 's data from the average of other users' distributions. For simplicity, we define $\ell_i = \|\nabla F_i(w^*)\|$.

3) *Unlearning Cost*: According to our previous analysis based on Scaffold and FedAvg algorithms [9], the number of unlearning communication rounds increases in the uniqueness levels of the users' revoked data:

$$\sum_{j \in \mathcal{I}} \left(1 - \frac{d_j}{d_j^{\max}} \right) \ell_j^2, \quad (5)$$

where $(1 - d_j/d_j^{\max})$ represents the proportion of user j 's revoked data size to its total data size $d_j^{\max}$. When the distance between the learned and unlearned models is large (as implied by a high proportion and uniqueness of the revoked data), it necessitates an increased number of communication rounds for effective federated unlearning.

User i 's unlearning cost increases in the unlearning rounds and its remained data size (used in each round's unlearning):³

$$C_i^u(d_i, \mathbf{d}_{-i}) = \theta_i d_i \sum_{j \in \mathcal{I}, j \neq i} \left(1 - \frac{d_j}{d_j^{\max}} \right) \ell_j^2, \quad (6)$$

where θ_i is user i 's marginal unlearning cost.

Combining the three parts in (3), (4), and (6), each user i 's payoff is

$$U_i(d_i, \mathbf{d}_{-i}) = P_i(d_i, \mathbf{d}_{-i}) - C_i^p(d_i) - C_i^u(d_i, \mathbf{d}_{-i}). \quad (7)$$

As users' payoffs are affected by each other's strategies, the users are engaged in a game.

D. Game Formulation

We formally define the resulting game as follows.

Game 1 (Users' Data Revocation Game). *This game consists of*

- *Players*: I users in set $\mathcal{I}$.

³For tractability, we assume that the influence of each user's revocation decision on the unlearning communication round is negligible, due to the significant volume of participating users. As a result, each user dismisses its individual impact on the unlearning rounds (indicated as $j \neq i$ in (6)). This assumption is not restrictive in applications like Gboard (Google's keyboard) which encompasses billions of users.

¹The distribution is due to the randomness in the training process (e.g. randomly sampled data and random ordering of batches).

²This paper focuses on analyzing the case of allowing partial data revocation. We will compare it with the case of forbidding partial revocation (i.e., binary decision $d_i \in \{0, d_i^{\max}\}$) in Section V.

- *Strategy space:* each user $i \in \mathcal{I}$ decides the data size to remain in the system $d_i \in [0, d_i^{\max}]$.
- *Payoff function:* each user $i \in \mathcal{I}$ maximizes its payoff $U_i(d_i, \mathbf{d}_{-i})$ in (7).

In this game, each user needs to trade off the model performance, privacy costs, and unlearning costs, when considering how much data to revoke. For instance, the contribution of additional data beyond a certain threshold may not significantly enhance the model's performance but lead to substantial privacy and unlearning costs.

As we will see in the next section, each user's revocation strategy affects other users through the global model performance and unlearning costs in a complex way. This makes it challenging to analyze the Nash equilibrium of this game.

III. USERS' DATA REVOCATION EQUILIBRIUM

In this section, we calculate the Nash equilibrium of Game 1, which is defined as follows.

Definition 1 (Equilibrium of Game 1). *The Nash equilibrium of Game 1 is a decision profile $\{d_i^*\}_{i \in \mathcal{I}}$, such that each user $i \in \mathcal{I}$ achieves its maximum payoff assuming other users are following the equilibrium strategies, i.e.,*

$$U_i(d_i^*, \mathbf{d}_{-i}^*) \geq U_i(d_i, \mathbf{d}_{-i}^*), \forall d_i \in [0, d_i^{\max}]. \quad (8)$$

A. Best Responses

Based on Definition 1, we first specify each user's best response to the other users' revocation decisions in Lemma 1.

Lemma 1. *The best response of user $i \in \mathcal{I}$ is*

$$d_i^*(\mathbf{d}_{-i}) = \left[\tilde{d}_i \right]_0^{d_i^{\max}} = \min \left\{ d_i^{\max}, \max \left\{ \tilde{d}_i, 0 \right\} \right\}, \quad (9)$$

where

$$\tilde{d}_i \triangleq \left(\xi_i \ell_i + \theta_i \sum_{j \in \mathcal{I}, j \neq i} \left(1 - \frac{d_j}{d_j^{\max}} \right) \ell_j^2 \right)^{-1} - \sum_{j \in \mathcal{I}, j \neq i} d_j - \epsilon_i. \quad (10)$$

Lemma 1 shows that users with larger costs (i.e., ξ_i and θ_i), data uniqueness level ℓ_i , and independency index ϵ_i tend to revoke more data (i.e., smaller $d_i^*(\mathbf{d}_{-i})$). However, users' data revocation decisions $\mathbf{d}$ influence each other in a complex way. To provide more insights into this interdependency, we consider two extreme cases of Lemma 1.

- If all other users do not revoke any data, i.e., $d_j = d_j^{\max}, \forall j \in \mathcal{I}, j \neq i$, then user i 's best response is

$$d_i^*(\mathbf{d}_{-i}) = \left[(\xi_i \ell_i)^{-1} - \sum_{j \in \mathcal{I}, j \neq i} d_j^{\max} - \epsilon_i \right]_0^{d_i^{\max}}. \quad (11)$$

When the marginal privacy cost ξ_i and data uniqueness level ℓ_i are large, user i will revoke all its data (i.e., $d_i^* = 0$) and be a free rider. However, this may not be an equilibrium, as other users may have the incentive to deviate.

- If all other users fully revoke data, i.e., $d_j = 0, \forall j \in \mathcal{I}, j \neq i$, then user i 's best response is

$$d_i^*(\mathbf{d}_{-i}) = \left[\left(\xi_i \ell_i + \theta_i \sum_{j \in \mathcal{I}, j \neq i} \ell_j^2 \right)^{-1} - \epsilon_i \right]_0^{d_i^{\max}}. \quad (12)$$

When user i has a high dependency on the global model (i.e., small ϵ_i), it will achieve the unlearning by itself.⁴ Specifically, if user i incurs high costs, it will use partial data; if its costs are small, it will not revoke any data. Conversely, when user i does not really value the global model (i.e., large ϵ_i), it will also revoke all its data (i.e., $d_i^* = 0$) like others. Thus, when users' independence indexes $\{\epsilon_i\}_{i \in \mathcal{I}}$ are sufficiently high, all users revoking all their data is an equilibrium.

Based on Lemma 1, we can also analyze the mutual influence of users' data (i.e., the externality) as follows:

Corollary 1 (Externality). *When $d_i^*(\mathbf{d}_{-i}) \in (0, d_i^{\max})$, if*

$$\theta_i \frac{d_j \ell_j^2}{d_j^{\max}} \geq \xi_i \ell_i + \theta_i \sum_{k \in \mathcal{I}, k \neq i, j} \left(1 - \frac{d_k}{d_k^{\max}} \right) \ell_k^2, \quad (13)$$

then we have

$$\frac{\partial d_i^*(\mathbf{d}_{-i})}{\partial d_j} \geq 0; \quad (14)$$

otherwise, we have

$$\frac{\partial d_i^*(\mathbf{d}_{-i})}{\partial d_j} < 0. \quad (15)$$

Corollary 1 demonstrates that if user j 's data uniqueness level ℓ_j and remaining data size d_j are sufficiently large (as indicated by (13)), user i 's partial contribution will increase in user j 's remaining data size (i.e., (14)). This suggests that users may follow others' data revocation decisions, potentially leading to cascaded revocations among users. We will provide illustrative numerical examples of this phenomenon in Section VI-B.

As shown in Lemma 1, the best responses of users are complex piece-wise functions that incorporate multi-dimensional parameters for user attributes. It is hard to directly calculate the equilibrium based on these best responses, as there will be many possible non-convex and non-monotonic cases depending on the parameters' values. We will illustrate the challenge through simulations in Section VI-A. Next, we first present the Nash equilibrium of homogeneous users in Section III-B and then for heterogeneous users in Section III-C.

⁴Note that when only one user remains within the system, the consequences of leaving or staying are different. By choosing to stay, the user can benefit from a warm starting point provided by the platform, facilitating the training process. However, in this situation, the user still incurs privacy costs. This is because the platform (or server) still can access the model, thus presenting a potential privacy risk.

B. Nash Equilibrium of Homogeneous Dependent Users

In the homogeneous dependent case, all I users have the same θ , ℓ , ξ , ϵ , and $d^{\max}$, and they depend on the global model, i.e., $\epsilon \rightarrow 0$.

The equilibrium in this case is in Proposition 1.

Proposition 1. *The Nash equilibrium of homogeneous dependent users is*

- If $d^{\max} > \frac{1}{I\xi\ell}$, then $\forall i \in \mathcal{I}$,

$$d_i^* = \frac{1}{2\theta\ell(I-1)} \left(d^{\max}(\xi + \theta\ell(I-1)) - \sqrt{d^{\max^2}(\xi + \theta\ell(I-1))^2 - 4\theta(I-1)d^{\max}/I} \right) \in (0, d^{\max}). \quad (16)$$

- If $d^{\max} \leq \frac{1}{I\xi\ell}$, then

$$d_i^* = d^{\max}, \forall i \in \mathcal{I}. \quad (17)$$

Proposition 1 shows that if the data size used in federated learning $d^{\max}$ and the costs $\xi\ell$ are excessively large, users will partially revoke their data (as in (16)). This is because contributing more data will not significantly improve the global model's local performance but will lead to high costs on privacy and unlearning. Conversely, when the amount of data utilized in federated learning and the costs are relatively small, users will not revoke any data to ensure the good performance of the global model (i.e., (17)).

C. Nash Equilibrium of Heterogeneous Users

In this subsection, we focus on the more complex case with heterogeneous users. We first prove the existence and non-uniqueness of the equilibrium. We next present some special equilibria and finally obtain the equilibrium under general conditions.

1) Existence and Non-Uniqueness:

Lemma 2. *Pure strategy Nash equilibrium exists but may not be unique in Game 1.*

Game 1 is a concave game, so a pure strategy Nash equilibrium exists [27]. To illustrate the non-uniqueness, we next present multiple special equilibria that can coexist in Proposition 2.

2) *Special Nash Equilibrium:* We first specify the conditions for two special Nash equilibria, where all users have the same strategies: fully revoke or not revoke their data.

Proposition 2 (Same strategies at equilibrium). *If and only if*

$$\epsilon_i \geq \left(\xi_i \ell_i + \theta_i \sum_{j \in \mathcal{I}, j \neq i} \ell_j^2 \right)^{-1}, \forall i \in \mathcal{I}, \quad (18)$$

there exists an equilibrium where all users fully revoke their data, i.e.,

$$d_i^* = 0, \forall i \in \mathcal{I}. \quad (19)$$

If and only if

$$\epsilon_i \leq (\xi_i \ell_i)^{-1} - \sum_{j \in \mathcal{I}} d_j^{\max}, \forall i \in \mathcal{I}, \quad (20)$$

there exists an equilibrium where no user revokes any data:

$$d_i^* = d_i^{\max}, \forall i \in \mathcal{I}. \quad (21)$$

Note that it is possible for the two conditions (18) and (20) to simultaneously hold, so the two equilibria in (19) and (21) may coexist. This happens when users have highly heterogeneous data (i.e., large data uniqueness ℓ), large unlearning costs θ , small data sizes d , and small privacy concerns ξ .

Proposition 2 shows that when the trained model is not very valuable to users (i.e., large independence indexes $\{\epsilon_i\}_{i \in \mathcal{I}}$ and users' costs are significant (i.e., small right-hand side of (18)), users will revoke all their data. On the other hand, if users really depend on the global model (i.e., small $\{\epsilon_i\}_{i \in \mathcal{I}}$ in (20)), privacy costs $\xi_i \ell_i$ are small, and there is still room for improving model performance (i.e., contributed data sizes $\sum_{j \in \mathcal{I}} d_j^{\max}$ are small), then all user fully contributing their data without any revocation is an equilibrium. This finding is consistent with the homogeneous dependent case where $\epsilon \rightarrow 0$, as described in Proposition 1.

There also exists Nash equilibrium where some users fully revoke their data, some users partially revoke data, and others do not revoke any data. Proposition 3 is an example.

Proposition 3 (Different strategies at equilibrium). *If conditions (22), (23), and (24) on the next page hold, there exists an equilibrium:*

$$\begin{aligned} d_i^* &= 0, \forall i \in \mathcal{I}_1, \\ d_k^* &= \left(\xi_k \ell_k + \theta_k \sum_{i \in \mathcal{I}_1} \ell_i^2 \right)^{-1} - \sum_{j \in \mathcal{I}_2} d_j^{\max} - \epsilon_k \in (0, d_k^{\max}), \quad (25) \\ d_j^* &= d_j^{\max}, \forall j \in \mathcal{I}_2. \end{aligned}$$

Proposition 3 shows that heterogeneous users may employ different strategies at equilibrium, including full, partial, and no data revocations. Each user's strategy depends on both its and others' attributes. Specifically, conditions (22) and (23) share similar insights with (18) and (20), where users need to trade off model independence indexes, costs, and model performance. For the right-hand side of (22) and (23), the first term represents user i 's costs (including privacy cost and unlearning cost), and the remaining two terms reflect the model performance. Condition (24) represents a middle case between (22) and (23) (i.e., user k has medium model independence, costs, and model performance), so user k chooses partial revocation. Note that there can be multiple users who partially revoke data (as in Section VI-B), but we refrain from including the complex conditions when this occurs for brevity.

We next present a general approach to computing Nash equilibria.

3) *General Approach for Nash Equilibrium Computation:* Many methods in the literature for calculating Nash equilibria are not applicable to our problem.

- A typical method involves best response updates simultaneously or sequentially (e.g., [28]). However, such a method does not necessarily converge for our problem, as shown by the example in Fig. 1.

$$\epsilon_i \geq \left(\xi_i \ell_i + \theta_i \sum_{i' \in \mathcal{I}_1 \setminus \{i\}} \ell_{i'}^2 + \theta_i \left(1 - \left(\left(\xi_k \ell_k + \theta_k \sum_{i' \in \mathcal{I}_1} \ell_{i'}^2 \right)^{-1} - \sum_{j \in \mathcal{I}_2} d_j^{\max} - \epsilon_k \right) / d_k^{\max} \right) \ell_k^2 \right)^{-1} - \left(\xi_k \ell_k + \theta_k \sum_{i' \in \mathcal{I}_1} \ell_{i'}^2 \right)^{-1} + \epsilon_k, \forall i \in \mathcal{I}_1, \quad (22)$$

$$\epsilon_j \leq \left(\xi_j \ell_j + \theta_j \sum_{i \in \mathcal{I}_1} \ell_i^2 + \theta_j \left(1 - \left(\left(\xi_k \ell_k + \theta_k \sum_{i \in \mathcal{I}_1} \ell_i^2 \right)^{-1} - \sum_{j' \in \mathcal{I}_2} d_{j'}^{\max} - \epsilon_k \right) / d_k^{\max} \right) \ell_k^2 \right)^{-1} - \left(\xi_k \ell_k + \theta_k \sum_{i \in \mathcal{I}_1} \ell_i^2 \right)^{-1} + \epsilon_k, \forall j \in \mathcal{I}_2, \quad (23)$$

$$\left(\xi_k \ell_k + \theta_k \sum_{i \in \mathcal{I}_1} \ell_i^2 \right)^{-1} - \sum_{j \in \mathcal{I}_2 \cup \{k\}} d_j^{\max} < \epsilon_k < \left(\xi_k \ell_k + \theta_k \sum_{i \in \mathcal{I}_1} \ell_i^2 \right)^{-1} - \sum_{j \in \mathcal{I}_2} d_j^{\max}, \quad (24)$$

where $\mathcal{I}_1 \cup \mathcal{I}_2 \cup \{k\} = \mathcal{I}$.

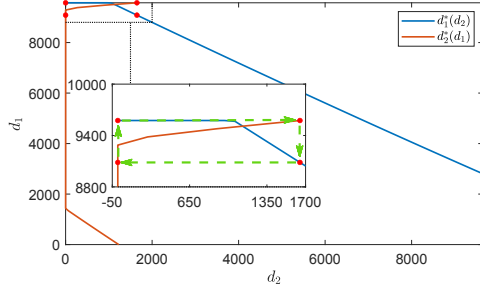


Fig. 1. Divergence of sequential best response update: an example of two users. The two users' best responses keep circulating among the four red points along the green trajectory, unable to converge to a Nash equilibrium where no user has an incentive to alter its strategy.

- Another widely-adopted method is gradient (better) response updates (e.g., [29]). However, the identification of appropriate stepsizes under different parameter settings for this problem remains unclear. A small stepsize can cause very slow convergence while a large stepsize may fail to converge at all. Moreover, the literature usually requires some technical assumptions for convergence (such as diagonally strictly concave and asymptotically stable), which do not apply in our context.

To ensure efficient convergence without requiring any assumption, we propose a general approach for obtaining the Nash equilibrium next. We first define the following function for the convenience of presentation:

$$W_i(\mathbf{d}_{-i}) \triangleq \ln \left(d_i^*(\mathbf{d}_{-i}) + \sum_{j \in \mathcal{I}, j \neq i} d_j + \epsilon_i \right) - \xi_i d_i^*(\mathbf{d}_{-i}) \ell_i - \theta_i d_i^*(\mathbf{d}_{-i}) \sum_{j \in \mathcal{I}, j \neq i} \left(1 - \frac{d_j}{d_j^{\max}} \right) \ell_j^2. \quad (26)$$

This function is obtained by substituting each user i 's best response $d_i^*(\mathbf{d}_{-i})$ in (9) into its payoff function $U_i(d_i, \mathbf{d}_{-i})$ in (7).

Inspired by [30], we transform calculating the equilibrium of Game 1 into solving the following optimization problem:

Problem 1.

$$\begin{aligned} \min \quad & \sum_{i \in \mathcal{I}} (W_i(\mathbf{d}_{-i}) - U_i(d_i, \mathbf{d}_{-i})) \\ \text{s.t.} \quad & 0 \leq d_i \leq d_i^{\max}, \quad i \in \mathcal{I} \\ \text{var.} \quad & \{d_i\}_{i \in \mathcal{I}} \end{aligned}$$

Theorem 1. *Problem 1 has a global solution and the corresponding objective value (i.e., minimum value) is 0. Moreover, each global minimum point of Problem 1 is a Nash equilibrium of Game 1.*

Consequently, it suffices to describe a method to solve Problem 1. The minimum point can be easily found through numerous global optimization methods (e.g., [31]–[33]) and optimization toolboxes (e.g., surrogateopt of MATLAB [34]).

There are several advantages of this approach for finding the equilibrium. First, this approach applies without requiring any additional assumptions on the payoff functions or decision variables. Second, it gives simple stopping criteria for iterative algorithms, as we seek to make the objective value equal to zero. Third, it also enables the use of local optimization methods. By checking whether a local optimum has an objective value of zero, one can verify if it is also a global optimum. Lastly, instead of tackling I separate problems and searching for a fixed point, we obtain the equilibrium by solving a single optimization problem.

We will show the corresponding numerical results in Section VI. Next, to present more insights, we consider a special scenario with negligible unlearning costs.

IV. SPECIAL SCENARIO: NEGLIGIBLE UNLEARNING COST

In this section, we consider a special scenario where users' unlearning costs are negligible (e.g., when users have abundant or even surplus computation resources). We formally state this in Assumption 1.

Assumption 1. *Users' unlearning costs are negligible, i.e.,*

$$\theta_i d_i \sum_{j \in \mathcal{I}, j \neq i} \left(1 - \frac{d_j}{d_j^{\max}} \right) \ell_j^2 \rightarrow 0, \forall d_i \in [0, d_i^{\max}], \forall i \in \mathcal{I}. \quad (27)$$

In this case, the payoff function of user $i \in \mathcal{I}$ is

$$U_i(d_i, \mathbf{d}_{-i}) = \ln \left(\sum_{j \in \mathcal{I}} d_j + \epsilon_i \right) - \xi_i d_i \ell_i, \quad (28)$$

i.e., users only care about the model performance and privacy concerns when deciding how much data to revoke.

The results and analysis in Section III still hold after setting $\theta = 0$. Next, we present some additional results and insights.

Lemma 3. Under Assumption 1, the best response of user $i \in \mathcal{I}$ is

$$d_i^*(\mathbf{d}_{-i}) = \left[(\xi_i \ell_i)^{-1} - \sum_{j \in \mathcal{I}, j \neq i} d_j - \epsilon_i \right]_0^{d_i^{\max}}. \quad (29)$$

Lemma 3 shows that when unlearning costs are negligible, a user's remaining data size d_i^* tends to decrease in other users' remaining data sizes $\sum_{j \in \mathcal{I}, j \neq i} d_j$. That is, a user will revoke more data if the others revoke less. This is different from the scenario with considerable unlearning costs where users may follow others' revocation decisions (Corollary 1). This is because when unlearning costs are not negligible, users affect each other through both the model performance and the unlearning costs. The former increases in users' remaining data sizes while the latter increases in users' revoked data sizes.

In this context, the users' multi-dimensional heterogeneity still poses challenges to computing the equilibrium. To overcome this, we propose to summarize it into a one-dimensional metric, which can greatly streamline the equilibrium calculation process. We define $(\xi \ell)^{-1} - \epsilon$ as the *remaining metric* and proceed to rank users based on this metric, as per the guidelines set in Assumption 2.

Assumption 2. Users follow a strict ascending order of the remaining metric:

$$(\xi_1 \ell_1)^{-1} - \epsilon_1 < \dots < (\xi_I \ell_I)^{-1} - \epsilon_I \quad (30)$$

and have the same maximum data size:

$$d_i^{\max} = d^{\max}, \forall i \in \mathcal{I}. \quad (31)$$

Given Assumption 2, we have the following relationship among the users' data revocation strategies.

Proposition 4. Under Assumptions 1 and 2, users' remaining data sizes in equilibrium satisfy:

$$d_1^* \leq d_2^* \leq \dots \leq d_I^*. \quad (32)$$

Proposition 4 shows that a user possessing a smaller remaining metric (i.e., a smaller index i) will revoke more data in equilibrium (i.e., a smaller d_i^*). This is due to its larger privacy cost ξ , data uniqueness ℓ , and model independence ϵ .

Based on Proposition 4, Theorem 2 shows that there is a unique equilibrium and that the remaining metric $(\xi \ell)^{-1} - \epsilon$ determines the users' data revocation decisions at this equilibrium.⁵

Theorem 2. Under Assumptions 1 and 2, there exists a unique Nash equilibrium and it falls into one of the following two cases:

(i) If there exists a user j satisfying

$$(I - j)d^{\max} \leq (\xi_j \ell_j)^{-1} - \epsilon_j \leq (I - j + 1)d^{\max}, \quad (33)$$

⁵Note that as we consider a strict order of users in (30), Theorem 2 shows at most one user who partially revokes data at the equilibrium. Otherwise, users with the same remaining metric can simultaneously choose partial revocation in equilibrium.

then the Nash equilibrium is

$$d_i^* = \begin{cases} 0, & \text{if } i < j, \\ (\xi_i \ell_i)^{-1} - (I - j)d^{\max} - \epsilon_i, & \text{if } i = j, \\ d^{\max}, & \text{if } i > j. \end{cases} \quad (34)$$

(ii) If no user satisfies (33), there must exist a user j satisfying

$$(\xi_j \ell_j)^{-1} - \epsilon_j < (I - j)d^{\max} < (\xi_{j+1} \ell_{j+1})^{-1} - \epsilon_{j+1}, \quad (35)$$

and the Nash equilibrium is

$$d_i^* = \begin{cases} 0, & \text{if } i \leq j, \\ d^{\max}, & \text{if } i > j. \end{cases} \quad (36)$$

In the next section, we investigate whether allowing partial revocation is beneficial to users.

V. IS ALLOWING PARTIAL REVOCATION BENEFICIAL?

In this section, we compare users' payoffs and total remaining data size in the two cases of allowing and forbidding partial data revocation.⁶ Specifically, we first focus on the general scenario where users have unlearning costs in Section V-A. Then, we use the special scenario with negligible unlearning costs to further provide insights in Section V-B.

A. General Scenario: Considerable Unlearning Cost

One might presume that allowing partial revocation will lead to more contributed data. This is because users who might have chosen to fully revoke their data without this option, might now opt for partial revocation instead. One might also presume that allowing partial revocation will increase users' payoffs as users have more choices. However, as shown in Proposition 5, the intuition may not hold.

Proposition 5. Compared with forbidding partial revocation, allowing partial revocation can increase, decrease, or not change the total amount of remaining data and users' payoffs.

This counter-intuitive phenomenon arises due to the strategic interactions of the users in the underlying game. Specifically, allowing partial revocation can also make some users who do not revoke data in the forbidding partial revocation case revoke some data, leading to possibly reduced total remaining data size. The resulting potential loss in payoffs from increasing the users' action space is similar to *Braess's paradox* in transportation networks [35].

B. Special Scenario: Negligible Unlearning Cost

In the following corollary, we give examples of the possible outcomes in Proposition 5 under the assumption of negligible unlearning costs. These examples show that users' payoffs and the remaining data sizes can either increase, decrease, or remain unchanged when partial revocation is allowed.

Corollary 2. When users have negligible unlearning costs and are forbidden to partially revoke data,

⁶For the forbidding partial revocation case, we just need to change each user i 's strategy space from continuous actions $[0, d_i^{\max}]$ to discrete actions $\{0, d_i^{\max}\}$ and conduct a similar analysis as in previous sections.

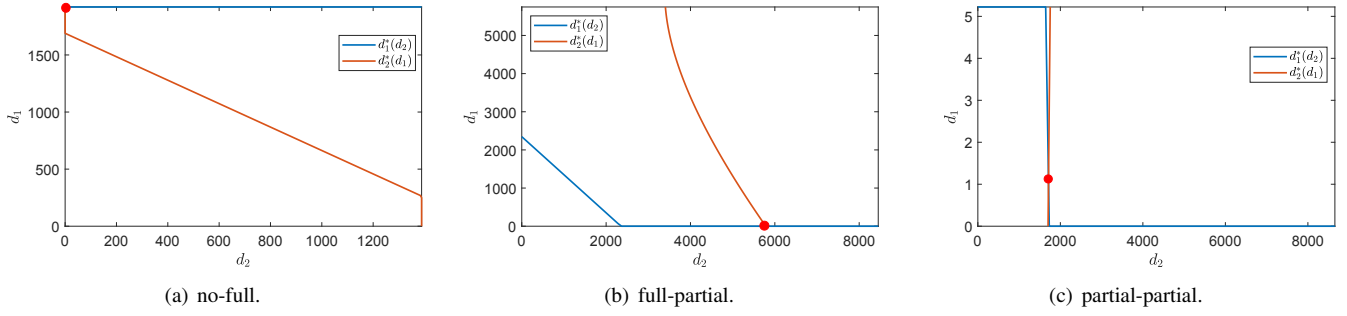


Fig. 2. Three examples of users' best responses to other users' data revocation strategies when there are two users ($I = 2$). The intersection of the two best response functions (i.e., red dot) in each figure is the Nash equilibrium. The caption of each sub-figure indicates the equilibrium choices of user 1 and user 2, e.g. "no-full" means that user 1 chooses no revocation while user 2 chooses full revocation.

- if the following conditions hold:

$$\ln \frac{(I-j+1)d^{\max} + \epsilon_i}{(I-j)d^{\max} + \epsilon_i} \geq \xi_i \ell_i d^{\max}, \forall i \in \{j, j+1, \dots, I\}, \quad (37)$$

user j in Theorem 2 case (i) will not revoke any data at the equilibrium and other users' equilibrium strategies remain unchanged, i.e.,

$$d_{i, \text{forbid}}^* = \begin{cases} 0, & \text{if } i < j, \\ d^{\max}, & \text{if } i \geq j. \end{cases} \quad (38)$$

Moreover, compared to allowing partial revocation, user j 's payoff decreases while the other users' payoffs increase, and the total remaining data size also increases.

- if the following conditions hold:

$$\ln \frac{(I-j+1)d^{\max} + \epsilon_i}{(I-j)d^{\max} + \epsilon_i} \leq \xi_i \ell_i d^{\max}, \forall i \in \{1, 2, \dots, j\}, \quad (39)$$

user j in Theorem 2 case (i) will fully revoke its data at the equilibrium and the other users' equilibrium strategies remain unchanged, i.e.,

$$d_{i, \text{forbid}}^* = \begin{cases} 0, & \text{if } i \leq j, \\ d^{\max}, & \text{if } i > j. \end{cases} \quad (40)$$

Compared to allowing partial revocation, all users' payoffs will decrease, and the total remaining data size will also decrease.

- for Theorem 2 case (ii), the equilibrium strategies of users do not change. Moreover, all users' payoffs and the total remaining data size also remain unchanged.

The condition (37) indicates that users with indexes larger than or equal to j will not deviate from the equilibrium in (38). The conditions for users with indexes smaller than j are naturally satisfied. In this case, user j ' strategy changes from partial revocation (when allowing partial revocation) to no revocation (when forbidding partial revocation), while the strategies of other users remain constant. This results in more data remaining in the system. As a consequence of the improved model performance, other users experience an increase in their payoffs. Nevertheless, due to the constraints imposed on the strategic space, the payoff for user j sees a reduction. The insights are similar for the remaining two cases.

VI. SIMULATIONS

In this section, we perform numerical experiments to validate our analytical results and reveal more insights. Specifically, we first illustrate the best responses of users and the complexity of computing Nash equilibria in Section VI-A, then we calculate users' revocation strategies at the equilibrium in Section VI-B, and finally we compare the allowing and forbidding partial revocation in Section VI-C.

A. Best Response

To illustrate the complexity of calculating Nash equilibria based on users' best responses, we first consider $I = 2$ users in the system with randomly generated attribute parameters (i.e., ℓ_i , ξ_i , θ_i , ϵ_i , and $d_i^{\max}$).⁷

Each user may adopt one of three types of strategies at the equilibrium: full revocation, partial revocation, and no revocation, each with different expressions as in Lemma 1. Thus, even for only two users, we can generally divide the equilibrium calculation into nine categories. Examples in three of these categories are shown in Fig. 2. Moreover, when their attribute parameters take different values, users' best response expressions can have different properties, such as non-convex (e.g., Fig. 2(a) and Fig. 2(c)), non-linear (e.g., Fig. 2(b)), or non-monotonic (e.g., Fig. 1), complicating analysis of the equilibrium.

B. Users' Data Revocation Strategies

We use the approach in Theorem 1 to calculate the equilibria with $I = 10$ users. We randomly generate users' different data uniqueness levels ℓ keeping the other parameters fixed.⁸

Fig. 3 shows users' data revocation strategies at the equilibrium when they have or do not have unlearning costs. When unlearning costs are negligible (green bar), users with smaller data uniqueness ℓ_i (i.e., users 7-10) keep more data in the system, and all users either fully revoke or do not revoke data at equilibrium. This validates the results in Proposition 4 and

⁷The results and insights under different groups of randomly generated parameters are similar, so we only present representative groups of them in this paper. Moreover, unless otherwise specified, we focus on the general scenario where users have unlearning costs in simulations.

⁸Besides the data uniqueness level, we have similarly explored varying the other parameters. The results are similar and are omitted due to space limitations.

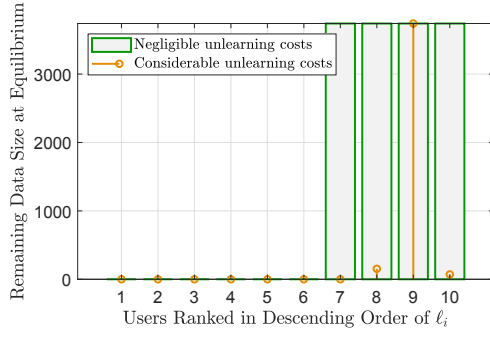


Fig. 3. Users' data revocation strategies at the equilibrium with or without unlearning costs.

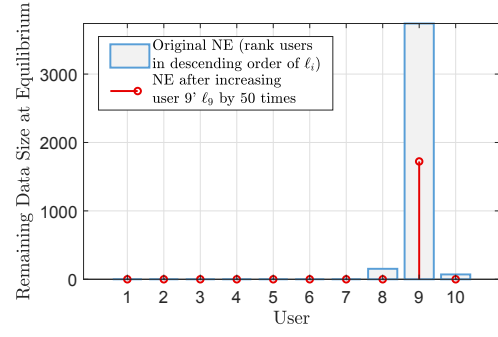
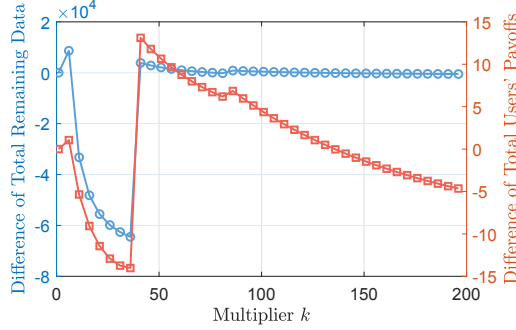
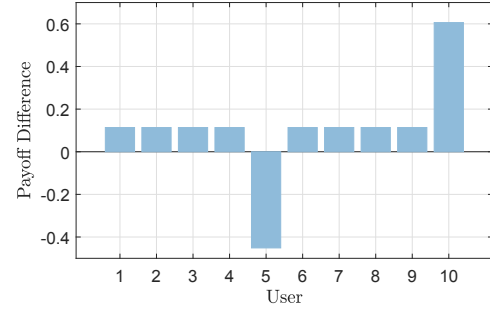


Fig. 4. Users' data revocation strategies at the Nash equilibrium (NE) before and after user 9's data uniqueness level changes, when users have unlearning costs.



(a) Difference of total remaining data sizes in two cases (allow—forbid) and difference of users' total payoffs in two cases (allow—forbid), when users' data uniqueness levels are multiplied by k .



(b) Difference of each user's payoffs in two cases (allow—forbid) when multiplier $k = 6$.

Fig. 5. Comparison between allowing and forbidding partial data revocation in terms of users' payoffs and total remaining data size.

Theorem 2. When users have unlearning costs (orange dot), users' remaining data sizes do not follow a strict order but roughly show that users with larger data uniqueness revoke more data, as indicated in Lemma 1. The lack of ordering in terms of data sizes is because, as shown in (10), users' data sizes and data uniqueness levels affect each other in a non-monotonic way.

Fig. 4 shows how a user's data uniqueness level affects the other users' data revocation strategies when users have unlearning costs. After we increase user 9's data uniqueness level by 50 times, user 9's equilibrium strategy changes from no revocation (blue bar) to partial revocation (redpoint). Users 8 and 10 also revoke more data, i.e., they change their strategies from partial revocation to full revocation. Other users still revoke all their data. This validates the insights of Corollary 1 that when users have unlearning costs, users may follow the data revocation decisions of others who have large data uniqueness levels and remaining data sizes.

C. Compare Allowing and Forbidding Partial Revocation

We numerically compare users' payoffs and remaining data sizes in the two cases of allowing and forbidding partial revocation. Fig. 5(a) shows the differences in payoffs and remaining data sizes between the two cases (allowing case minus forbidding case), when users' randomly generated data uniqueness levels are multiplied by k . The differences can take

positive, zero, or negative values. This confirms that allowing partial data revocation can either enhance, leave unchanged, or diminish users' total payoff and remaining data size compared with when it is forbidden.

Fig. 5(b) further shows that 10 users may have different preferences on the two revocation cases. User 5 prefers forbidding partial revocation while others prefer allowing partial revocation. The comparison insights are consistent with Propositions 5 and Corollary 2.

VII. CONCLUSION

This paper focuses on data revocation in federated unlearning. To the best of our knowledge, this is the first study of users' strategic data revocation in federated unlearning with partial revocation allowed. We propose a game theoretic framework to capture users' decision-making trade-offs among the model performance, data privacy concerns, and federated unlearning costs. Despite inherent challenges, we obtain the Nash equilibrium of the game through both analytical and numerical methods. We use this equilibrium to deliver some interesting insights, such as that allowing partial data revocation may not always benefit users compared to when it is not allowed, and users' unlearning costs can cause their cascaded data revocation. Possible future work includes the design of incentive mechanisms to improve the users' payoffs.

REFERENCES

- [1] H. Wang, K. Sreenivasan, S. Rajput, H. Vishwakarma, S. Agarwal, J.-y. Sohn, K. Lee, and D. Papailiopoulos, "Attack of the tails: Yes, you really can backdoor federated learning," *Advances in Neural Information Processing Systems*, vol. 33, pp. 16 070–16 084, 2020.
- [2] L. Lyu, H. Yu, and Q. Yang, "Threats to federated learning: A survey," *arXiv preprint arXiv:2003.02133*, 2020.
- [3] P. Voigt and A. Von dem Bussche, "The eu general data protection regulation (gdpr): A Practical Guide, 1st Ed., Cham: Springer International Publishing, vol. 10, no. 3152676, pp. 10–5555, 2017.
- [4] E. L. Harding, J. J. Vanto, R. Clark, L. Hannah Ji, and S. C. Ainsworth, "Understanding the scope and impact of the california consumer privacy act of 2018," *Journal of Data Protection & Privacy*, vol. 2, no. 3, pp. 234–253, 2019.
- [5] Y. Liu, Z. Ma, X. Liu, and J. Ma, "Learn to forget: User-level memorization elimination in federated learning," *arXiv preprint arXiv:2003.10933*, 2020.
- [6] Y. Jiang, J. Konečný, K. Rush, and S. Kannan, "Improving federated learning personalization via model agnostic meta learning," *arXiv preprint arXiv:1909.12488*, 2019.
- [7] X. Gao, X. Ma, J. Wang, Y. Sun, B. Li, S. Ji, P. Cheng, and J. Chen, "Verifi: Towards verifiable federated unlearning," *arXiv preprint arXiv:2205.12709*, 2022.
- [8] G. Liu, X. Ma, Y. Yang, C. Wang, and J. Liu, "Federaser: Enabling efficient client-level data removal from federated learning models," in *2021 IEEE/ACM 29th International Symposium on Quality of Service*, 2021.
- [9] N. Ding, Z. Sun, E. Wei, and R. Berry, "Incentive mechanism design for federated learning and unlearning," in *ACM International Symposium on Theory, Algorithmic Foundations, and Protocol Design for Mobile Networks and Mobile Computing (MobiHoc)*, 2023.
- [10] Y. Cao and J. Yang, "Towards making systems forget with machine unlearning," in *IEEE Symposium on Security and Privacy*, 2015.
- [11] T. T. Nguyen, T. T. Huynh, P. L. Nguyen, A. W.-C. Liew, H. Yin, and Q. V. H. Nguyen, "A survey of machine unlearning," *arXiv preprint arXiv:2209.02299*, 2022.
- [12] Y. Liu, L. Xu, X. Yuan, C. Wang, and B. Li, "The right to be forgotten in federated learning: An efficient realization with rapid retraining," in *IEEE Conference on Computer Communications (INFOCOM)*, 2022.
- [13] C. Wu, S. Zhu, and P. Mitra, "Federated unlearning with knowledge distillation," *arXiv preprint arXiv:2201.09441*, 2022.
- [14] Y. Zhan, P. Li, Z. Qu, D. Zeng, and S. Guo, "A learning-based incentive mechanism for federated learning," *IEEE Internet of Things Journal*, vol. 7, no. 7, pp. 6360–6368, 2020.
- [15] N. Ding, Z. Fang, and J. Huang, "Optimal contract design for efficient federated learning with multi-dimensional private information," *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 1, pp. 186–200, 2020.
- [16] M. Zhang, E. Wei, and R. Berry, "Faithful edge federated learning: Scalability and privacy," *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 12, pp. 3790–3804, 2021.
- [17] Z. Wang, L. Gao, and J. Huang, "Socially-optimal mechanism design for incentivized online learning," in *IEEE Conference on Computer Communications (INFOCOM)*, 2022.
- [18] N. Zhang, Q. Ma, and X. Chen, "Enabling long-term cooperation in cross-silo federated learning: A repeated game perspective," *IEEE Transactions on Mobile Computing*, 2022.
- [19] S. P. Karimireddy, S. Kale, M. Mohri, S. Stich, and A. T. Suresh, "Scaffold: Stochastic controlled averaging for federated learning," in *International Conference on Machine Learning*, 2020.
- [20] R. Pathak and M. J. Wainwright, "Fedsplit: An algorithmic framework for fast federated optimization," *Advances in Neural Information Processing Systems*, vol. 33, pp. 7057–7066, 2020.
- [21] H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-efficient learning of deep networks from decentralized data," in *Artificial Intelligence and Statistics*, 2017.
- [22] Y. Zhan, P. Li, K. Wang, S. Guo, and Y. Xia, "Big data analytics by crowdlearning: Architecture and mechanism design," *IEEE Network*, vol. 34, no. 3, pp. 143–147, 2020.
- [23] T. Hastie, R. Tibshirani, J. H. Friedman, and J. H. Friedman, *The elements of statistical learning: data mining, inference, and prediction*. Springer, 2009, vol. 2.
- [24] L. Lü, C.-H. Jin, and T. Zhou, "Similarity index based on local paths for link prediction of complex networks," *Physical Review E*, vol. 80, no. 4, p. 046122, 2009.
- [25] Y.-A. De Montjoye, C. A. Hidalgo, M. Verleysen, and V. D. Blondel, "Unique in the crowd: The privacy bounds of human mobility," *Scientific reports*, vol. 3, no. 1, pp. 1–5, 2013.
- [26] D. Romanini, S. Lehmann, and M. Kivelä, "Privacy and uniqueness of neighborhoods in social networks," *Scientific reports*, vol. 11, no. 1, p. 20104, 2021.
- [27] J. B. Rosen, "Existence and uniqueness of equilibrium points for concave n-person games," *Econometrica: Journal of the Econometric Society*, pp. 520–534, 1965.
- [28] P. Ramazi and M. Cao, "Asynchronous decision-making dynamics under best-response update rule in finite heterogeneous populations," *IEEE Transactions on Automatic Control*, vol. 63, no. 3, pp. 742–751, 2017.
- [29] J. Lei and U. V. Shanbhag, "Distributed variable sample-size gradient-response and best-response schemes for stochastic nash equilibrium problems," *SIAM Journal on Optimization*, vol. 32, no. 2, pp. 573–603, 2022.
- [30] A. Chernov, "On some approaches to find nash equilibrium in concave games," *Automation and Remote Control*, vol. 80, pp. 964–988, 2019.
- [31] S. M. Elsakov and V. I. Shiryaev, "Homogeneous algorithms of multiextremal optimization for time-consuming objective functions," *Numerical Methods and Programming (Vychislitel'nye Metody i Programirovanie)*, vol. 12, pp. 48–69, 2011.
- [32] J. Wang, S. C. Clark, E. Liu, and P. I. Frazier, "Parallel bayesian global optimization of expensive functions," *Operations Research*, vol. 68, no. 6, pp. 1850–1865, 2020.
- [33] S. Piyavskii, "An algorithm for finding the absolute extremum of a function," *USSR Computational Mathematics and Mathematical Physics*, vol. 12, no. 4, pp. 57–67, 1972.
- [34] "surrogateopt," <https://www.mathworks.com/help/gads/surrogateopt.html#d124e70078>.
- [35] D. Braess, A. Nagurney, and T. Wakolbinger, "On a paradox of traffic planning," *Transportation science*, vol. 39, no. 4, pp. 446–450, 2005.