

Differentially-Private Collaborative Online Personalized Mean Estimation

Yauhen Yakimenka*, Chung-Wei Weng[†], Hsuan-Yin Lin[†], Eirik Rosnes[†], and Jörg Kliewer*

*Helen and John C. Hartmann Department of Electrical and Computer Engineering,
New Jersey Institute of Technology, Newark, New Jersey 07102, USA

[†]Simula UiB, N-5006 Bergen, Norway

Abstract—We consider the problem of collaborative personalized mean estimation under a privacy constraint in an environment of several agents continuously receiving data according to arbitrary unknown agent-specific distributions. In particular, we provide a method based on hypothesis testing coupled with differential privacy. Two privacy mechanisms are proposed and we provide a theoretical convergence analysis of the proposed algorithm for any bounded unknown distributions on the agents' data. Numerical results show that for a considered scenario the proposed approach converges much faster than a fully local approach where agents do not share data, and performs comparably to ideal performance where all data is public. This illustrates the benefit of private collaboration in an online setting.

I. INTRODUCTION

Collaborative learning has attracted significant attention lately through popular frameworks such as federated learning (FL) [1]–[3] (partially decentralized) and fully decentralized approaches like swarm learning [4]. However, different agents in the learning environment may have different objectives and hence the individually collected data may be heterogeneous and specific for each personalized learning task. Despite this, collaboration can significantly accelerate learning among a set of agents sharing a limited set of common objectives. A crucial part of any collaborative algorithm for personalized learning is the identification of agents with data from similar distributions, in particular in an online setting in which data becomes available continuously over time.

Personalized approaches for distributed learning [5]–[7] have attracted significant interest recently. In [7], Hanzely and Richtárik introduced the concept of personalized FL. In contrast to conventional FL, personalized FL looks for a trade-off between a global model and local models learned by each agent from its own dataset, as formulated in terms of a correction term to the traditional empirical risk minimization objective function. As shown in [7], personalization in general yields reduced communication complexity.

For the online setting, previous work on collaborative learning has mostly focused on the multi-armed bandit (MAB) model, mostly considering a *single* MAB instance (the arm means do not vary across the agents), see, e.g., [8]–[10] and references therein, while some recent works also consider the case where the arm means vary across agents [11] and with some amount of personalization by optimizing a mixture between a global and local objectives [12].

In this paper, we consider the problem of collaborative online personalized mean estimation, first introduced in [13],

This work was in part supported by US NSF grants 1815322, 1908756, and 2107370.

in which each agent continuously receives data according to an unknown agent-specific distribution. The aim of each agent is to calculate an accurate estimate of the mean of its underlying distribution as quickly as possible. As in [13], we assume an unknown underlying class structure where agents in the same class receive data from distributions with the same mean.

A major limitation of the algorithm proposed in [13] is that data is directly shared with other agents in the learning environment without any protection, which is in contrast to FL where there is no sharing of data amongst the agents. This may leak sensitive user information to other agents in the environment. In order to provide some level of user data privacy, we propose to add random noise to the data before it is released to other agents in the environment according to the principle of differential privacy (DP) [14], [15]. Two (online) privacy mechanisms inspired by those in [16], [17] are proposed. Moreover, as opposed to the initial work in [13], we consider an approach based on hypothesis testing, and provide a theoretical convergence analysis of our proposed method for any bounded distributions on the data (see Theorem 1). Numerical results show that our proposed approach converges faster than a fully local setting where agents do not share data, and the best scheme performs comparably to ideal performance where all data is public. This illustrates the benefit of collaboration in an online setting while preserving users' data privacy. Due to lack of space, all proofs are omitted.

II. PRELIMINARIES

A. Notation

In general, but with some exceptions, we use uppercase and lowercase letters for random variables (RVs) and their realization, respectively, and italics for sets, e.g., X , x , and $\mathcal{X}$ represent a RV, its realization, and a set, respectively. The expectation of a RV X is denoted by $\mathbb{E}[X]$, while its variance is denoted by $\text{Var}[X]$. We define $[n] \triangleq \{1, 2, \dots, n\}$ and $[i : j] \triangleq \{i, i + 1, \dots, j\}$, while $\mathbb{R}$ denotes the real numbers. $\mathcal{N}(\mu, \sigma^2)$ denotes the Gaussian distribution with mean μ and variance σ^2 . $X \sim \mathcal{P}$ denotes that X is distributed according to the distribution $\mathcal{P}$. Standard order notations $O(\cdot)$ and $o(\cdot)$ are used for asymptotic results, while $\Phi(\cdot)$ is the cumulative distribution function of the standard Gaussian distribution. $w_H(n)$ denotes the Hamming weight of the binary representation of the nonnegative integer n .

B. System Model (Problem Formulation)

There are M independent agents. Each agent $a \in [M]$ wants to estimate the mean of its sample $X_a^{(1)}, X_a^{(2)}, \dots \in \mathcal{X}_a \subset \mathbb{R}$,

where the sample follows an arbitrary unknown distribution $\mathcal{D}_a$ over the bounded set $\mathcal{X}_a$ with an (unknown) mean μ_a and known standard deviation $\sigma_a < \infty$, where $X_a^{(t)}$ arrives to the agent a at time t . We assume the time is discrete and synchronized between the agents. For simplicity, we assume that $\mathcal{X}_a$ is the same for all agents a . The agents have limited memory and thus decide to keep only the current mean of the sample: $\bar{X}_a^{(t)} = \frac{1}{t} \sum_{i=1}^t X_a^{(i)}$.

It is known that some agents $a \neq b$ might have samples from the same distribution $\mathcal{D}_a = \mathcal{D}_b$, and they want to exploit this. However, there is no preliminary information on the agents' distributions and also, the agents would like to keep their particular sample values private.

The collaborative algorithm consists of the agents exchanging their current sample means. In order to preserve privacy, a DP mechanism [14], [15] is applied before releasing the current sample mean to other agents. More precisely, at each time step t , the agent a receives $X_a^{(t)}$, updates its sample mean $\bar{X}_a^{(t)}$, and also chooses another agent $b \in [M] \setminus \{a\}$ to query. The agent b then sends its current sample mean to a , but privatized: $\bar{X}_b^{(t)} + Z_{b \rightarrow a}^{(t)}$, where $Z_{b \rightarrow a}^{(t)}$ is a Gaussian RV ("noise") with zero mean and variance depending on the DP mechanism employed (see Section IV). While $Z_{b \rightarrow a}^{(t)}$ is independent from the sample, it in general depends on the noise generated by agent b at different times. We call a particular construction of the noise $Z_{b \rightarrow a}^{(t)}$ a *private release mechanism*.

Based on the content of its memory at time t , agent a calculates its current estimate of μ_a , which we denote by $\mu_a^{(t)}$. The goal is to construct a collaborative protocol that allows for faster convergence of $\mu_a^{(t)}$ to μ_a . In this work, we measure the speed of convergence in terms of the average expected squared deviation from the mean, i.e., by $\sum_{a \in [M]} \mathbb{E}[(\mu_a^{(t)} - \mu_a)^2] / M$ as a function of t .

Also, the agents want regular updates so that at every moment t they have a good estimate of μ_a . Hence, we do not consider the approach where one waits until $t \approx t_{\max}$ and then query every agent's last sample mean.

C. Differential Privacy

We start by defining the concept of DP. Then, we provide a key lemma based on the Gaussian mechanism.

Definition 1: A randomized function $F: \mathcal{X}^n \rightarrow \mathcal{Y}$ is (ϵ, δ) -differentially private if for all subsets $\mathcal{S} \subseteq \mathcal{Y}$ and for all $(x_1, \dots, x_n) \in \mathcal{X}^n$ and $(x'_1, \dots, x'_n) \in \mathcal{X}^n$ which differ in a single component, i.e., $x_i \neq x'_i$ for exactly one $i \in [n]$,

$$\Pr[F(x_1, \dots, x_n) \in \mathcal{S}] \leq e^\epsilon \Pr[F(x'_1, \dots, x'_n) \in \mathcal{S}] + \delta.$$

Lemma 1: Let $(x_1, \dots, x_n) \in \mathcal{X}^n$ where $\mathcal{X} = [\mu - L, \mu + L]$ for some finite values μ and L . Then, the noise-corrupted sample mean $(x_1 + \dots + x_n)/n + Z/n$, where $Z \sim \mathcal{N}(0, \sigma_{\text{DP}}^2)$ and $\sigma_{\text{DP}}^2 \triangleq 8L^2 \ln(1.25/\delta)/\epsilon^2$ is (ϵ, δ) -differentially private for $0 < \epsilon \leq 1$ and $0 < \delta \leq 1$.

D. Mathematical Tools

Lemma 2: Assume independent RVs $X_1, X_2, \dots, X_n$ have $\text{Var}[X_i] = \sigma_i^2$. If $X = \alpha_1 X_1 + \alpha_2 X_2 + \dots + \alpha_n X_n$, where

$\sum_{i=1}^n \alpha_i = 1$, then X has the minimum variance when the weights α_i are selected as $\alpha_i^* = \frac{1}{\sigma_i^2 \sum_{j=1}^n \frac{1}{\sigma_j^2}}$, and

$$\min_{\alpha_1, \dots, \alpha_n} \text{Var}[X] = \sum_{i=1}^n (\alpha_i^*)^2 \sigma_i^2 = \frac{1}{\sum_{i=1}^n \frac{1}{\sigma_i^2}} \leq \min_{i \in [n]} \sigma_i^2.$$

Lemma 2 provides an intuition for the whole approach: having access to RVs with the same mean and even very high variances still allows to decrease the total variance of the estimator X , provided the weights α_i are properly chosen.

Assume we have two Gaussian RVs $X \sim \mathcal{N}(\mu, \sigma^2)$ and $Y \sim \mathcal{N}(\nu, \tau^2)$, where the variances σ^2 and τ^2 are known, but the means μ and ν are not. Here, we construct a simple hypothesis test for checking if $\mu = \nu$ as follows,

$$\mathcal{H}_0 : \mu = \nu \text{ and } \mathcal{H}_1 : \mu \neq \nu.$$

First, let $Z = (X - Y) / \sqrt{\sigma^2 + \tau^2} \sim \mathcal{N}((\mu - \nu) / \sqrt{\sigma^2 + \tau^2}, 1)$. Then, if we have a predefined confidence level $\theta \in [0, 1]$ and $z \triangleq \Phi^{-1}(1 - \theta/2)$, we

accept $\mathcal{H}_0$ if $|Z| < z$ and reject $\mathcal{H}_0$ otherwise.

We call the situation when $\mu = \nu$ but $\mathcal{H}_0$ is rejected, a *type-I error*. The probability of type-I error is θ . And if $\mu \neq \nu$ but $\mathcal{H}_0$ is accepted, we call this a *type-II error*.

III. OUR APPROACH

There are three ingredients of our approach that should be addressed:

- private release mechanism $Z_{b \rightarrow a}^{(t)}$ that adds the minimum amount of noise sufficient for (ϵ, δ) -DP of the sample by the agent b (see Section IV),
- identification by the agent a of the agents with the same distribution mean (*decision rule*, see Section III-B), and
- using the information obtained from these agents in order to improve the local mean estimate (statistic $T_{b \rightarrow a}$, see Section III-A).

Formally, we define the class of agents having the same distribution mean as a by $\mathcal{C}_a \triangleq \{b \in [M] : \mathcal{D}_b = \mathcal{D}_a\}$. The classes are not known, which makes the problem nontrivial. We denote also by $\mathcal{C}_a^{(t)} \triangleq \{b \in [M] : \chi_a^{(t)}(b; \theta_t) = 1\}$ the estimate of the class $\mathcal{C}_a$ by the agent a at time t , where $\chi_a^{(t)}(b; \theta_t)$ is some decision rule at time t , i.e., $\chi_a^{(t)}(b; \theta_t) = 1$ if at time t agent a believes that agent b is in $\mathcal{C}_a$. Here, θ_t is a prescribed confidence level that depends on t .

A. Linear Statistic $T_{b \rightarrow a}$

Let $\sum_{i=1}^{\kappa_{b \rightarrow a}} w_i = 1$, denote by $t_1, t_2, \dots$ the times at which agent b is queried by agent a , and let

$$T_{b \rightarrow a} = \sum_{i=1}^{\kappa_{b \rightarrow a}} w_i (\bar{X}_b^{(t_i)} + Z_{b \rightarrow a}^{(t_i)}) \quad (1)$$

be the current statistic of the received noise-corrupted sample means by agent a from agent b after the $\kappa_{b \rightarrow a}$ -th query, $\kappa_{b \rightarrow a} = 1, 2, \dots$. Then (with $t_0 \triangleq 0$),

$$\begin{aligned} \text{Var}[T_{b \rightarrow a}] &= \sigma_b^2 \sum_{i=1}^{\kappa_{b \rightarrow a}} (t_i - t_{i-1}) \left(\sum_{j=i}^{\kappa_{b \rightarrow a}} \frac{w_j}{t_j} \right)^2 \\ &\quad + \text{Var} \left[\sum_{i=1}^{\kappa_{b \rightarrow a}} w_i Z_{b \rightarrow a}^{(t_i)} \right]. \end{aligned} \quad (2)$$

The weights $\{w_i\}$ depend on $\kappa_{b \rightarrow a}$, but we omit this dependency for simpler notation. Picking $w_1 = \dots = w_{\kappa_{b \rightarrow a}-1} = 0$ and $w_{\kappa_{b \rightarrow a}} = 1$ corresponds to keeping the last update as in [13], while picking $w_1 = \dots = w_{\kappa_{b \rightarrow a}} = 1/\kappa_{b \rightarrow a}$ corresponds to what we refer to as the mean-of-mean (MoM) statistic. For simplicity, we refer to the former approach as non-MoM.

B. Decision Rule

The sum $\sum_{i=1}^{\kappa_{b \rightarrow a}} w_i \bar{X}_b^{(t_i)}$ is a weighted average of the independent and identically distributed (i.i.d.) RVs $X_b^{(1)}, X_b^{(2)}, \dots, X_b^{(\kappa_{b \rightarrow a})}$. If this sum is Gaussian distributed, it follows from (1) that $T_{b \rightarrow a}$ is also Gaussian. Hence, we pick a decision rule based on the hypothesis test outlined above in Section II-D, i.e., we let $\chi_a^{(t)}(b; \theta_t) = 1$, for $b \neq a$, if

$$\begin{aligned} |\bar{X}_a^{(t)} - T_{b \rightarrow a}| &< \Phi^{-1}\left(1 - \frac{\theta_t}{2}\right) \sqrt{\text{Var}[\bar{X}_a^{(t)}] + \text{Var}[T_{b \rightarrow a}]} \\ &= \Phi^{-1}\left(1 - \frac{\theta_t}{2}\right) \sqrt{\frac{\sigma_a^2}{t} + \text{Var}[T_{b \rightarrow a}]} \end{aligned}$$

and 0, otherwise, where $\theta_t \in [0, 1]$ and $t = t_{\kappa_{b \rightarrow a}}$. Additionally, $\chi_a^{(t)}(a; \theta_t) = 1$ always, and we set $\chi_a^{(t)}(b; \theta_t) = 1$ before agent a receives from agent b for the first time.

If we keep the last update, i.e., $w_1 = \dots = w_{\kappa_{b \rightarrow a}-1} = 0$ and $w_{\kappa_{b \rightarrow a}} = 1$, the sum above becomes $\bar{X}_b^{(t_{\kappa_{b \rightarrow a}})}$, which is asymptotically Gaussian by the central limit theorem. But in general, this decision rule leads to an asymptotically correct estimate of the class $\mathcal{C}_a$, even when asymptotic Gaussness cannot be proved, but more general conditions are satisfied.

C. Algorithm

We summarize the proposed algorithm in Algorithm 1, which is based on the linear statistic of Section III-A. The crucial step that differentiates Algorithm 1 from the COLME algorithm in [13, Alg. 1] is the design of a new decision rule $\chi_a^{(t)}(b; \delta)$ in Line 12, which is based on hypothesis testing as outlined in Section III-B. Moreover, we consider a more general linear statistic $T_{b \rightarrow a}$ in Line 10 (the COLME algorithm corresponds to fixing the last weight $w_{\kappa_{b \rightarrow a}}$ equal to one). The selection of an agent b to query is done according to some schedule, e.g., round-robin (RR), denoted by `choose_agent`. Finally, both $T_{b \rightarrow a}$ and $\mathcal{C}_a^{(t)}$ are updated and the statistics $T_{b \rightarrow a}$, for $b \in \mathcal{C}_a^{(t)}$, are linearly combined in Line 13 in order to obtain an improved estimate $\mu_a^{(t)}$ of the mean of agent a at time step t . The linear combination coefficients are optimized based on Lemma 2, i.e., selected according to

$$\alpha_{b \rightarrow a}^{(t)} = \begin{cases} \frac{t}{\sigma_a^2 \left(\sum_{b' \in \mathcal{C}_a^{(t)} \setminus \{a\}} \frac{1}{\text{Var}[T_{b' \rightarrow a}]} + \frac{t}{\sigma_a^2} \right)} & \text{if } b = a, \\ \frac{1}{\text{Var}[T_{b \rightarrow a}] \left(\sum_{b' \in \mathcal{C}_a^{(t)} \setminus \{a\}} \frac{1}{\text{Var}[T_{b' \rightarrow a}]} + \frac{t}{\sigma_a^2} \right)} & \text{otherwise,} \end{cases} \quad (3)$$

while in [13, Alg. 1] several different (heuristic) linear combination schemes are considered.

From Line 13 of Algorithm 1 and Lemma 2 it follows that

$$\text{Var}[\mu_a^{(t)}] = \frac{1}{\sum_{b' \in \mathcal{C}_a^{(t)} \setminus \{a\}} \frac{1}{\text{Var}[T_{b' \rightarrow a}]} + \frac{t}{\sigma_a^2}}.$$

Algorithm 1: Private-ColME

Input: agent a
Output: $\mu_a^{(t_{\max})}$

- 1 $\forall b \in [M] \setminus \{a\} : T_{b \rightarrow a} \leftarrow 0, \kappa_{b \rightarrow a} \leftarrow 0$
- 2 $\mathcal{C}_a^{(0)} \leftarrow [M]$
- 3 **for** $t = 1, 2, \dots, t_{\max}$ **do**
- 4 // Receive
- 5 Receive sample $X_a^{(t)} \sim \mathcal{D}_a$
- 6 $\bar{X}_a^{(t)} \leftarrow \bar{X}_a^{(t-1)} \times \frac{t-1}{t} + X_a^{(t)} \times \frac{1}{t}$
- 7 // Query
- 8 $b \leftarrow \text{choose_agent}(\mathcal{C}_a^{(t-1)}, [M])$
- 9 $\kappa_{b \rightarrow a} \leftarrow \kappa_{b \rightarrow a} + 1$
- 10 $T_{b \rightarrow a} \leftarrow \sum_{i=1}^{\kappa_{b \rightarrow a}} w_i (\bar{X}_b^{(t_i)} + Z_{b \rightarrow a}^{(t_i)})$
- 11 // Estimate
- 12 $\mathcal{C}_a^{(t)} \leftarrow \{b \in [M] : \chi_a^{(t)}(b; \delta) = 1\}$
- 13 $\mu_a^{(t)} \leftarrow \alpha_{a \rightarrow a}^{(t)} \bar{X}_a^{(t)} + \sum_{b \in \mathcal{C}_a^{(t)} \setminus \{a\}} \alpha_{b \rightarrow a}^{(t)} T_{b \rightarrow a}$
- 14 **return** $\mu_a^{(t_{\max})}$

If there has been no values received from an agent b , we assume as a convention that $T_{b \rightarrow a} = 0$ and $\text{Var}[T_{b \rightarrow a}] = +\infty$.

D. Schedules

As mentioned above in Section III-C, the agents are queried according to some schedule. In this work, we study a simple RR schedule in which agents are queried in order of their indices (but skipping the agent a itself). Additionally, we consider a *restricted* RR (rRR) schedule in which the agents are queried in the same order as in RR, but at any current time step t the agents not in $\mathcal{C}_a^{(t-1)}$ are skipped.

IV. PRIVACY

Below, we present two private release mechanisms giving different trade-offs between the variance of the linear statistic $T_{b \rightarrow a}$ (see Line 10 of Algorithm 1) and the overall privacy level for each individual sample $X_b^{(t)}$ of agent b when releasing noise-corrupted sample means to agent a .

Both mechanisms are based on the following idea, which can be seen as a generalization of Lemma 1. To construct a privatized version of $\bar{X}_b^{(t)}$ for release to agent a , agent b splits the corresponding sum of values with indices $[1 : t]$ into κ subsums (called *p-sums* in [16]) with indices $[1 : \tau_1], [\tau_1 + 1 : \tau_2], \dots, [\tau_{\kappa-1} + 1 : t]$:

$$\begin{aligned} \bar{X}_b^{(t)} &= \frac{X_b^{(1)} + \dots + X_b^{(t)}}{t} \\ &= \frac{\sum_{i=1}^{\tau_1} X_b^{(i)} + \sum_{i=\tau_1+1}^{\tau_2} X_b^{(i)} + \dots + \sum_{i=\tau_{\kappa-1}+1}^t X_b^{(i)}}{t} \end{aligned}$$

and adds independent noise with the same variance $\sigma_{\text{DP}}^2 = 8L^2 \ln(1.25/\delta)/\epsilon^2$ to each of the subsums:

$$\begin{aligned} &\frac{\sum_{i=1}^{\tau_1} X_b^{(i)} + Z_{b \rightarrow a}^{(1:\tau_1)}}{t} + \frac{\sum_{i=\tau_1+1}^{\tau_2} X_b^{(i)} + Z_{b \rightarrow a}^{(\tau_1+1:\tau_2)}}{t} \\ &+ \dots + \frac{\sum_{i=\tau_{\kappa-1}+1}^t X_b^{(i)} + Z_{b \rightarrow a}^{(\tau_{\kappa-1}+1:t)}}{t} = \bar{X}_b^{(t)} + Z_{b \rightarrow a}^{(t)}, \end{aligned}$$

where all $Z_{b \rightarrow a}^{(1:\tau_1)}, Z_{b \rightarrow a}^{(\tau_1+1:\tau_2)}, \dots, Z_{b \rightarrow a}^{(\tau_{\kappa-1}+1:t)}$ are i.i.d. according to $\mathcal{N}(0, \sigma_{\text{DP}}^2)$, and

$$Z_{b \rightarrow a}^{(t)} \triangleq \frac{1}{t} \sum_{i=1}^{\kappa} Z_{b \rightarrow a}^{(\tau_{i-1}+1:\tau_i)} \sim \mathcal{N}(0, \kappa \sigma_{\text{DP}}^2 / t^2),$$

where $\tau_0 = 0$ and $\tau_{\kappa} = t$. Note that the variance of the noise $Z_{b \rightarrow a}^{(t)}$ depends only on the time of release, t , the number of subsums in the split, κ , and the desired (ϵ, δ) .

Similar to Lemma 1, this release mechanism provides (ϵ, δ) -DP for each $X_b^{(1)}, X_b^{(2)}, \dots, X_b^{(t)}$. However, a subsum with the corresponding noise can be re-used by the agent b further for constructing privatized versions of the means $\bar{X}_b^{(t')}$ for $t' > t$, thus reducing the amount of “fresh” noise added. If for example $(X_b^{(1)} + \dots + X_b^{(\tau_1)} + Z_{b \rightarrow a}^{(1:\tau_1)})/t$ is released several times (with the same value of $Z_{b \rightarrow a}^{(1:\tau_1)}$) as a subsum of different sums, the privacy of each of the $X_b^{(1)}, \dots, X_b^{(\tau_1)}$ stays the same.

The only difference between the two mechanisms below is how we split into the subsums. In both mechanisms, if the same subsum (i.e., with the same interval of indices) needs to be used for the calculation of several privatized means by agent b for release to agent a , we actually *require* that the exact same value of noise is used for this subsum.

We stress here that in both mechanisms below, the RVs $Z_{b \rightarrow a}^{(t)}$ are dependent for $t = t_1, t_2, \dots$, and the variance of the linear statistic $T_{b \rightarrow a}$, needed for the implementation of the decision rule $\chi_a^{(t)}(b; \theta_t)$ and the computation of the coefficients $\alpha_{b \rightarrow a}^{(t)}$ in (3), depends on it and the weights $\{w_i\}$ used. An analytical expression for the variance (specific to the weights and the privacy mechanism) can be derived from (2), but is omitted due to lack of space.

A. Privacy Mechanism I (PM-I)

This mechanism is inspired by the Simple Counting Mechanism II in [16]. The split of sums into subsums as above now exactly corresponds to the times $t_1, t_2, \dots$ when agent b is queried by agent a , i.e., $[1 : t_{\kappa}]$ is split into $[1 : t_1], [t_1 + 1 : t_2], \dots, [t_{\kappa-1} + 1 : t_{\kappa}]$. Hence, $\bar{X}_b^{(t_{\kappa})} + Z_{b \rightarrow a}^{(t_{\kappa})} = \bar{X}_b^{(t_{\kappa})} + 1/t_{\kappa} \sum_{i=1}^{\kappa} Z_{b \rightarrow a}^{(t_{i-1}+1:t_i)}$, and $Z_{b \rightarrow a}^{(t_{\kappa})} \sim \mathcal{N}(0, \kappa \sigma_{\text{DP}}^2 / t_{\kappa}^2)$.

Note that PM-I allows for efficient implementation by the agent b . Indeed, it can keep only the current value of $\sum_{i=1}^{\kappa} Z_{b \rightarrow a}^{(t_{i-1}+1:t_i)}$. At the next release time $t = t_{\kappa+1}$, it updates this cumulative noise by adding “fresh” noise $Z_{b \rightarrow a}^{(t_{\kappa}+1:t_{\kappa+1})}$ and releases $\bar{X}_b^{(t_{\kappa+1})}$ privatized with this updated noise (divided by $t_{\kappa+1}$). In particular, agent b does not need to keep $t_1, \dots, t_{\kappa}$. Hence, agent b needs $O(1)$ memory to implement PM-I.

Since every $X_b^{(t)}$, $1 \leq t \leq t_{\kappa}$, participates in an exactly one subsum (defined by the aforementioned split), PM-I allows for a constant privacy level as we query agent b from agent a , rather than getting weaker and weaker over time, while at the same time having a decreasing DP noise variance due to the factor κ/t_{κ}^2 .

Lemma 3: Consider RR or rRR scheduling and an oracle class estimator, i.e., $C_a^{(t)} = C_a$ for all t . Next, let agent b be the last agent queried by agent a in a single round. Then, selecting $w_1 = \dots = w_{\kappa b \rightarrow a - 1} = 0$ and $w_{\kappa b \rightarrow a} = 1$ minimizes the variance of $T_{a \rightarrow b}$.

Hence, based on Lemma 3, keeping the last update (as proposed in [13]) is a good strategy for PM-I.

B. Privacy Mechanism II (PM-II)

This mechanism is inspired by the Binary Counting Mechanism in [16]. When an agent b releases data to agent a for the κ -th time (at time instant t_{κ}), we construct the split into subsums in two steps as follows. First, consider the same split of $[1 : t_{\kappa}]$ into $[1 : t_1], [t_1 + 1 : t_2], \dots, [t_{\kappa-1} + 1 : t_{\kappa}]$ as for PM-I. Second, we now join the corresponding subsums into larger subsums according to the binary representation of κ . More precisely, let $\kappa = 2^{s_1} + 2^{s_2} + \dots + 2^{s_{w_H(\kappa)}}$, $s_i > s_{i+1}$, be the unique representation of κ based on positions of ones in the binary representation of κ .¹ Then, we join the first 2^{s_1} aforementioned subsums into the first larger subsum, the next 2^{s_2} subsums into the second larger subsum, etc. In total, these two steps result in splitting $[1 : t_{\kappa}]$ as follows:

$$\begin{aligned} [1 : t_{2^{s_1}}] & \quad \text{(first subsum),} \\ [t_{2^{s_1}} + 1 : t_{2^{s_1}+2^{s_2}}] & \quad \text{(second subsum),} \\ [t_{2^{s_1}+2^{s_2}} + 1 : t_{2^{s_1}+2^{s_2}+2^{s_3}}] & \quad \text{(third subsum),} \\ & \vdots \\ [t_{2^{s_1}+2^{s_2}+\dots+2^{s_{w_H(\kappa)-1}}} + 1 : t_{\kappa}] & \quad (w_H(\kappa)\text{-th subsum).} \end{aligned}$$

Finally, we add independent noise with variance σ_{DP}^2 to each of the corresponding subsums and construct the noisy mean $\bar{X}_b^{(t_{\kappa})} + Z_{b \rightarrow a}^{(t_{\kappa})}$ as previously. This mechanism gives $Z_{b \rightarrow a}^{(t_{\kappa})} \sim \mathcal{N}(0, w_H(\kappa) \sigma_{\text{DP}}^2 / t_{\kappa}^2)$.

The noise term $Z_{b \rightarrow a}^{(t_{\kappa})}$ has variance at most $(\lfloor \log_2 \kappa \rfloor + 1) / t_{\kappa}^2$ times σ_{DP}^2 as $w_H(\kappa) \leq \lfloor \log_2 \kappa \rfloor + 1$. The variance is a factor of at most $(\lfloor \log_2 \kappa \rfloor + 1) / \kappa$ of the variance of PM-I and this factor approaches zero in κ . On the other hand, in contrast to PM-I, a value $X_b^{(t)}$ for $1 \leq t \leq t_{\kappa}$ can be used in up to $\lfloor \log_2 \kappa \rfloor + 1$ different subsums. Therefore, from the composition theorem of DP (see, e.g., [18, Thm. 3.1]), PM-II will give $((\lfloor \log_2 \kappa \rfloor + 1)\epsilon, (\lfloor \log_2 \kappa \rfloor + 1)\delta)$ -DP with respect to each individual sample $X_b^{(1)}, \dots, X_b^{(t_{\kappa})}$.

Compared to PM-I, the privacy parameters of this mechanism grow with κ , and hence the privacy level gets weaker over time. On the other hand, the DP noise variance decreases faster with κ . As shown below in Section VI, this may result in reaching a given target average mean squared error for a given *fixed* overall privacy level faster in some scenarios. On the other hand, in contrast to PM-I, agent b needs $O(\lfloor \log_2 \kappa \rfloor)$ memory to implement PM-II.

Interesting, keeping only the last update does not minimize $\text{Var}[T_{b \rightarrow a}]$ as for PM-I (cf. Lemma 3). For this particular mechanism we will illustrate in Section VI below that a *windowed* MoM (wMoM) approach where $w_1 = \dots = w_{2^{\lfloor \log_2 \kappa_{b \rightarrow a} \rfloor - 1}} = 0$ and $w_{2^{\lfloor \log_2 \kappa_{b \rightarrow a} \rfloor}} = \dots = w_{\kappa_{b \rightarrow a}} = 1 / (\kappa_{b \rightarrow a} - 2^{\lfloor \log_2 \kappa_{b \rightarrow a} \rfloor + 1})$ performs better.

V. PERFORMANCE ANALYSIS

We first present the performance of a pure local approach in which each agent $a \in [M]$ estimates its mean solely based

¹E.g., if $\kappa = 13 = 1101$, we represent it as $\kappa = 2^3 + 2^2 + 2^0$.

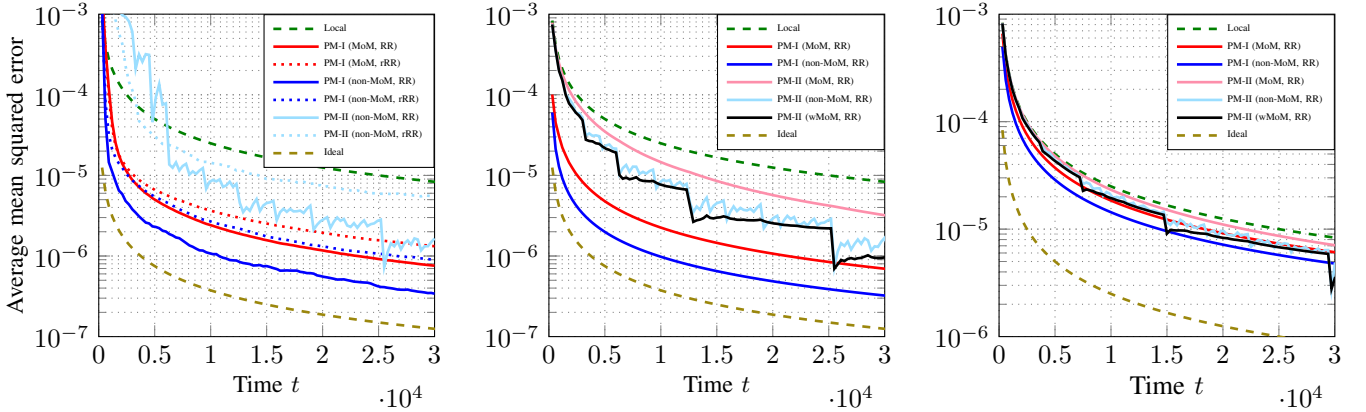


Fig. 1. Average mean squared error of Algorithm 1 with $\sigma = 1/2$. The left plot shows simulation results with RR and rRR scheduling for the case of $M = 200$ agents forming three classes, with an overall privacy level of $\epsilon = 1$ with $\delta = 10^{-6}$. The middle and right plots show the corresponding (analytical) performance with an oracle class estimator and with RR scheduling for $M = 200$ and 30 agents, respectively. The curves are for uniform data and $L = \sigma\sqrt{3}$.

on its own data. For the local approach, the privacy is perfect as there is no sharing of data.

Proposition 1 (Local): The average mean squared error of a pure local approach is

$$\frac{1}{M} \sum_{a \in [M]} \mathbb{E} \left[\left(\mu_a^{(t)} - \mu_a \right)^2 \right] = \frac{1}{Mt} \sum_{a \in [M]} \sigma_a^2.$$

It is good to understand the limits of what can be achieved. If privacy is ignored, and agent a knows $\mathcal{C}_a$ and has access to all the data of all the agents in $\mathcal{C}_a$ at time t , the agent a virtually has one large sample of size $|\mathcal{C}_a|t$ (as opposed to the sample of size t for the pure local approach). With this, the average mean squared error is $1/Mt \sum_{a \in [M]} \frac{\sigma_a^2}{|\mathcal{C}_a|}$ and no approach can perform better than this *ideal* performance.

The following theorem shows that the average mean squared error from Algorithm 1 converges to zero as $t \rightarrow \infty$.

Theorem 1: Let $1/\theta_t = o(e^t)$ and $1/\theta_t \rightarrow \infty$ as $t \rightarrow \infty$. For any distributions $\mathcal{D}_a$, $a \in [M]$, for both PM-I and PM-II and with both non-MoM and MoM weights, Algorithm 1 with RR scheduling produces $\mu_a^{(t)}$ that is asymptotically unbiased and

$$\frac{1}{M} \sum_{a \in [M]} \mathbb{E} \left[\left(\mu_a^{(t)} - \mu_a \right)^2 \right] \rightarrow 0 \text{ as } t \rightarrow \infty.$$

VI. NUMERICAL RESULTS

We consider the case of $M = 200$ agents forming three classes. The agents are placed uniformly at random within the classes, giving roughly balanced class sizes. The agents' data distributions are uniform (to model tabular data) on a range of size $2L = 2\sigma\sqrt{3}$ (giving a standard deviation of σ), with $\sigma = 1/2$, but with different class-dependent means; $1/5$, $2/5$, and $4/5$, respectively, for the three classes. In order to have a fair comparison between PM-I and PM-II, (ϵ, δ) of PM-II is scaled by $\lceil \log_2(t_{\max}) \rceil + 1$ so that both mechanisms provide the same privacy level. For the decision rule, we use $\theta_t = 0.05/\ln(t+1)$.

In Fig. 1 (left plot), we show simulation results for RR and rRR scheduling in Algorithm 1 (with $t_{\max} = 3 \cdot 10^4$) for both PM-I and PM-II with the MoM and non-MoM approach. The overall privacy level is $\epsilon = 1$ with $\delta = 10^{-6}$. As expected, for PM-I the non-MoM approach outperforms the MoM approach (cf. Lemma 3), while there is still some gap

to ideal performance. Moreover, PM-I with non-MoM weights outperforms PM-II for the range of squared errors shown in the plot. Note that RR performs better than rRR, which was not the case when privacy is ignored (see [13, Fig. 1]). This can be explained by the fact that sample means are released less often with RR than with rRR (the gaps $t_i - t_{i-1}$ are larger for RR). The collaborative schemes (ultimately) perform significantly better than a pure local approach, while PM-II with non-MoM weights and rRR scheduling being a notable exception; the reason being that the decision rule type-I error increases over time rather than converging to a close-to-zero value. In the middle plot, the corresponding (analytical) performance with an oracle class estimator and with RR scheduling is provided which shows qualitatively the same behavior as in the left plot. While PM-I performs better than PM-II for the range of squared errors shown, asymptotically (for very low squared errors) the curves will cross (not shown). Note that PM-II with wMoM weights smoothens the corresponding curve with non-MoM weights (the oscillations are due to the factor $w_H(\kappa)$ in the variance of $Z_{b \rightarrow a}^{(t, \kappa)}$) and shows that having non-MoM weights for PM-II is not optimal. Moreover, PM-II with wMoM weights performs significantly better than a pure MoM approach. In the right plot, we show the corresponding oracle performance for $M = 30$. In contrast to the middle plot, PM-II with wMoM weights shows the best performance for a low squared error. On the other hand, compared to $M = 200$, the performance gap to ideal performance is far larger.

VII. CONCLUSION

We presented a private collaborative algorithm for personalized mean estimation in an online setting where agents continuously receive samples according to arbitrary unknown distributions. An approach based on DP was proposed and significant gains compared to a pure local approach where the agents do not share their data were demonstrated. For a considered scenario, the best scheme performs comparably to an ideal scheme where all data is public. We also provided a convergence analysis of the proposed algorithm. The communication cost was not considered and its implication on the accuracy/privacy will be studied next as future work.

REFERENCES

- [1] H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. Agüera y Arcas, "Communication-efficient learning of deep networks from decentralized data," in *Proc. 20th Int. Conf. Artif. Intell. Statist. (AISTATS)*, Ft. Lauderdale, FL, USA, Apr. 20–22, 2017, pp. 1273–1282.
- [2] J. Konečný, H. B. McMahan, F. X. Yu, P. Richtárik, A. T. Suresh, and D. Bacon, "Federated learning: Strategies for improving communication efficiency," in *NeurIPS Workshop Private Multi-Party Mach. Learn. (PMPML)*, Barcelona, Spain, Dec. 9, 2016.
- [3] T. Li, A. K. Sahu, A. Talwalkar, and V. Smith, "Federated learning: Challenges, methods, and future directions," *IEEE Signal Process. Mag.*, vol. 37, no. 3, pp. 50–60, May 2020.
- [4] S. Warnat-Herresthal *et al.*, "Swarm learning for decentralized and confidential clinical machine learning," *Nature*, vol. 594, no. 7862, pp. 265–270, Jun. 2021.
- [5] V. Smith, C.-K. Chiang, M. Sanjabi, and A. Talwalkar, "Federated multi-task learning," in *Proc. Neural Inf. Process. Syst. (NeurIPS)*, Long Beach, CA, USA, Dec. 4–9, 2017, pp. 4424–4434.
- [6] P. Vanhaesebrouck, A. Bellet, and M. Tommasi, "Decentralized collaborative learning of personalized models over networks," in *Proc. 20th Int. Conf. Artif. Intell. Statist. (AISTATS)*, Ft. Lauderdale, FL, USA, Apr. 20–22, 2017, pp. 509–517.
- [7] F. Hanzely and P. Richtárik, "Federated learning of a mixture of global and local models," Feb. 2020, arXiv:2002.05516v3 [cs.LG]. [Online]. Available: <https://arxiv.org/abs/2002.05516>
- [8] U. Madhushani, A. Dubey, N. E. Leonard, and A. Pentland, "One more step towards reality: Cooperative bandits with imperfect communication," in *Proc. Neural Inf. Process. Syst. (NeurIPS)*, Online, Dec. 6–14, 2021, pp. 1–12.
- [9] C. Tao, Q. Zhang, and Y. Zhou, "Collaborative learning with limited interaction: Tight bounds for distributed exploration in multi-armed bandits," in *Proc. 60th Annu. IEEE Symp. Found. Comp. Sci. (FOCS)*, Baltimore, MD, USA, Nov. 9–12, 2019, pp. 126–146.
- [10] E. Hillel, Z. Karnin, T. Koren, R. Lempel, and O. Somekh, "Distributed exploration in multi-armed bandits," in *Proc. Neural Inf. Process. Syst. (NeurIPS)*, Lake Tahoe, NV, USA, Dec. 5–10, 2013, pp. 854–862.
- [11] N. Karpov and Q. Zhang, "Collaborative best arm identification with limited communication on non-IID data," Jul. 2022, arXiv:2207.08015v2 [cs.LG]. [Online]. Available: <https://arxiv.org/abs/2207.08015>
- [12] C. Shi, C. Shen, and J. Yang, "Federated multi-armed bandits with personalization," in *Proc. 24th Int. Conf. Artif. Intell. Statist. (AISTATS)*, Online, Apr. 13–15, 2021, pp. 2917–2925.
- [13] M. Asadi, A. Bellet, O.-A. Maillard, and M. Tommasi, "Collaborative algorithms for online personalized mean estimation," *Trans. Mach. Learn. Res.*, pp. 1–23, Dec. 2022.
- [14] C. Dwork, "Differential privacy," in *Proc. 33rd Int. Coll. Automata, Lang. Program. (ICALP), part II*, Venice, Italy, Jul. 10–14, 2006, pp. 1–12.
- [15] C. Dwork, F. McSherry, K. Nissim, and A. Smith, "Calibrating noise to sensitivity in private data analysis," in *Proc. Theory Cryptography Conf. (TCC)*, New York, NY, USA, Mar. 4–7, 2006, pp. 265–284.
- [16] T.-H. H. Chan, E. Shi, and D. Song, "Private and continual release of statistics," *ACM Trans. Inf. Syst. Secur.*, vol. 14, no. 3, pp. 26:1–26:24, Nov. 2011.
- [17] C. Dwork, M. Naor, T. Pitassi, and G. N. Rothblum, "Differential privacy under continual observation," in *Proc. 42nd ACM Symp. Theory Comput. (STOC)*, Cambridge, MA, USA, Jun. 6–8, 2010, pp. 715–724.
- [18] P. Kairouz, S. Oh, and P. Viswanath, "The composition theorem for differential privacy," *IEEE Trans. Inf. Theory*, vol. 63, no. 6, pp. 4037–4049, Jun. 2017.