

Supplementary Information for “Selective Knowledge Sharing for Privacy-Preserving Federated Distillation”

Jiawei Shao¹, Fangzhao Wu^{2,*}, and Jun Zhang^{1,*}

¹Hong Kong University of Science and Technology, Hong Kong

²Microsoft Research Asia, Beijing

*correspondce: Jun Zhang (eejzhang@ust.hk), Fangzhao Wu (wufangzhao@gmail.com)

7

8

9 Supplementary Material

10 Proof of Theorem 2

11 Prior to proving Theorem 2, we first present two Lemmas.

12 **Lemma 1.** Given hypothesis spaces $\mathcal{H} := \{\hat{\mathbf{h}}: \mathcal{X} \rightarrow V(\Delta^{C-1})\}$ and $\mathcal{G} := \{g: \mathcal{X} \rightarrow \{0, 1\}\}$ with $g(\mathbf{x}) = \frac{1}{2} \|\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}'(\mathbf{x})\|_1$
 13 for $\hat{\mathbf{h}}, \hat{\mathbf{h}}' \in \mathcal{H}$, we have $|\mathcal{L}_{\mathcal{D}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}'}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| \leq d_{\mathcal{G}}(\mathcal{D}, \mathcal{D}')$ for $\hat{\mathbf{h}}, \hat{\mathbf{h}}' \in \mathcal{H}$.

Proof.

$$\begin{aligned} d_{\mathcal{G}}(\mathcal{D}, \mathcal{D}') &= 2 \sup_{g \in \mathcal{G}} |\Pr_{\mathcal{D}}[g(\mathbf{x}) = 1] - \Pr_{\mathcal{D}'}[g(\mathbf{x}) = 1]|, \\ &= \sup_{\hat{\mathbf{h}}, \hat{\mathbf{h}}' \in \mathcal{H}} 2 \left| \frac{1}{2} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}} [\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}'(\mathbf{x})] - \frac{1}{2} \mathbb{E}_{\mathbf{x} \sim \mathcal{D}'} [\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}'(\mathbf{x})] \right| \geq |\mathcal{L}_{\mathcal{D}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}'}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')|. \end{aligned}$$

14

Lemma 2. For any $\delta \in (0, 1)$, with probability at least $1 - \delta$ over the choice of the samples, we have

$$\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) \leq \mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}) + \sqrt{\left(\frac{2\alpha^2}{m_k} + \frac{2(1-\alpha)^2}{m_{\text{proxy}}} \right) \log \frac{2}{\delta}}. \quad (1)$$

Proof. The loss $\mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}})$ can be written as

$$\begin{aligned} \mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}) &= \alpha \mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}) + (1 - \alpha) \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}), \\ &= \frac{1}{m_k + m_{\text{proxy}}} \left[\sum_{\mathbf{x} \in \hat{\mathcal{D}}_k} \frac{\alpha(m_k + m_{\text{proxy}})}{m_k} \|\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}^*(\mathbf{x})\|_1 + \sum_{\mathbf{x} \in \hat{\mathcal{D}}_{\text{proxy}}} \frac{(1 - \alpha)(m_k + m_{\text{proxy}})}{m_{\text{proxy}}} \|\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}^*(\mathbf{x})\|_1 \right]. \end{aligned} \quad (2)$$

Let $X_1^{(k)}, \dots, X_{m_k}^{(k)}$ and $X_1^{(\text{proxy})}, \dots, X_{m_{\text{proxy}}}^{(\text{proxy})}$ be independent random variables that take on the values of $\frac{\alpha(m_k + m_{\text{proxy}})}{m_k} \|\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}^*(\mathbf{x})\|_1$ for $\mathbf{x} \in \hat{\mathcal{D}}_k$ and $\frac{(1 - \alpha)(m_k + m_{\text{proxy}})}{m_{\text{proxy}}} \|\hat{\mathbf{h}}(\mathbf{x}) - \hat{\mathbf{h}}^*(\mathbf{x})\|_1$ for $\mathbf{x} \in \hat{\mathcal{D}}_{\text{proxy}}$, respectively. We define $\bar{X}$ as the mean value of these variables, which represents the empirical loss in (2). By linearity of expectations, $\mathbb{E}[\bar{X}]$ is equal to the loss $\mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}})$. According to Hoeffding's inequality, for any $\varepsilon > 0$, we have

$$\Pr \left[|\mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}) - \mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}})| \geq \varepsilon \right] \leq 2 \exp \left(\frac{-\varepsilon^2}{\frac{2\alpha^2}{m_k} + \frac{2(1-\alpha)^2}{m_{\text{proxy}}}} \right). \quad (3)$$

15 Let the right-hand side of (3) be δ , we can derive the inequality in Theorem 2. □

16 We are now ready to prove Theorem 2.

Proof. The following derives the upper bound of $|\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}})|$:

$$\begin{aligned} &|\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}})| = |\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \alpha \mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - (1 - \alpha) \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*)|, \\ &\leq \alpha |\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*)| + (1 - \alpha) |\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*)|, \\ &\leq \alpha \left[|\mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| + |\mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| + |\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*)| \right], \\ &\quad + (1 - \alpha) \left[|\mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| + |\mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| + |\mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*)| \right], \\ &\leq \alpha \left[\mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') + |\mathcal{L}_{\hat{\mathcal{D}}_k}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| + \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) \right], \end{aligned} \quad (4)$$

$$+ (1 - \alpha) \left[\mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) + |\mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}') - \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}')| + \mathcal{L}_{\mathcal{D}_{\text{test}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) \right], \quad (5)$$

$$\leq \alpha [\lambda_k + d_{\mathcal{G}_k}(\mathcal{D}_k, \mathcal{D}_{\text{test}})] + (1 - \alpha) \left[\lambda_{k, \text{proxy}} + d_{\mathcal{G}_k}(\mathcal{D}_{\text{proxy}}, \mathcal{D}_{\text{test}}) + \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) - \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) \right], \quad (6)$$

$$\leq \alpha [\lambda_k + d_{\mathcal{G}_k}(\mathcal{D}_k, \mathcal{D}_{\text{test}})] + (1 - \alpha) \left[\lambda_{k, \text{proxy}} + d_{\mathcal{G}_k}(\mathcal{D}_{\text{proxy}}, \mathcal{D}_{\text{test}}) + \mathcal{L}_{\hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}}, \hat{\mathbf{h}}^*) \right], \quad (7)$$

where

$$\mathcal{L}_{\mathcal{D}_{\text{proxy}}}(\mathbf{h}_{\text{proxy}}^*, \hat{\mathbf{h}}^*) = p_{\text{proxy}}^{(1)} \mathcal{L}_{\mathcal{D}_{\text{proxy}}^{(1)}}(\hat{\mathbf{h}}^*, \mathbf{h}_{\text{proxy}}^*) + p_{\text{proxy}}^{(2)} \mathcal{L}_{\mathcal{D}_{\text{proxy}}^{(2)}}(\hat{\mathbf{h}}_{\text{proxy}}^*, \mathbf{h}_{\text{proxy}}^*). \quad (8)$$

The triangle inequality gives rise to (4), (5), and (7), while Lemma 1 underlies the inequality (6). By combining (7), (8), and Lemma 2, we obtain the upper bound stated in Theorem 2. $\square$

More Details of Experiments

In the experiments, we use stochastic gradient descent (SGD) with a learning rate of 0.1 as the optimizer. During federated distillation, instead of directly minimizing the empirical loss $\mathcal{L}_{\hat{\mathcal{D}}_k \cup \hat{\mathcal{D}}_{\text{proxy}}}(\hat{\mathbf{h}})$ in Theorem 2, clients begin by training their local models independently based on their datasets for 200 SGD steps. In each communication round, clients then fine-tune the models on local samples and proxy samples for 1 and 10 SGD steps, respectively. This training scheme makes the knowledge distillation process more stable. Table 1 provides a summary of datasets, and the architectures of local models are shown in Table 2, Table 3, and Table 4. The convolutional layer with output channel o , kernel size k , and padding p is denoted as $\text{Conv}(o, k, p)$. The fully-connected layer with output dimension o is defined as $\text{linear}(o)$, and the max-pooling layer with kernel size k is denoted as $\text{MaxPool}(k)$. ReLU represents the rectified linear unit function.

In addition, our method has adopted a portion of the local data as a validation set to determine the threshold τ_{client} of the client-side selectors. The value of τ_{client} corresponds to the quantile of the estimated ratio over the aforementioned validation set. In the weak non-IID scenario, clients are given local training sets containing two classes. As modeling the joint distribution of these classes is challenging and results in high sampling complexity, we utilize an alternative approach where we construct two density-ratio estimators, denoted as $\mathbf{w}_k^{(1)}$ and $\mathbf{w}_k^{(2)}$, for each class. To determine whether a proxy sample should be treated as an in-distribution sample, we apply the same method for selecting the thresholds $\tau_k^{(1)}$ and $\tau_k^{(2)}$ for $\mathbf{w}_k^{(1)}$ and $\mathbf{w}_k^{(2)}$, respectively. If any of the estimators output a density ratio higher than its corresponding threshold, the local prediction of the proxy sample will not be filtered out.

Table 1. Summary of datasets.

	COVIDx	MNIST	Fashion MNIST	CIFAR-10
Number of local data	1000 ~ 2000	5,400	5,400	4,000
Local training batch size	64	64	64	64
Number of proxy data	1,500	6,000	6,000	10,000
Distillation batch size	128	512	512	32

Table 2. The architectures of local models on the pneumonia detection task.

Client 1	Client 2	Client 3	Client 4
Conv(10, 5, 0) + ReLU MaxPool(2)	Conv (10, 3, 1) + ReLU MaxPool(2)	Linear(1024) + ReLU Linear(512) + ReLU	Linear(1024) + ReLU Linear(1024) + ReLU
Conv(20, 5, 0) + ReLU MaxPool(2)	Conv (20, 3, 1) + ReLU MaxPool(2)	Linear(256) + ReLU Linear(3)	Linear(3)
Linear(50) + ReLU Linear(3)	Linear(128) + ReLU Linear(3)		

Table 3. The architectures of local models on the MNIST and Fashion MNIST datasets. Note that the MNIST and Fashion MNIST datasets have images of different sizes, which consequently leads to variations in the input dimensions for the fully-connected layers.

Client 1	Client 2	Client 3	Client 4	Client 5	Client 6	Client 7	Client 8	Client 9	Client 10
Conv(10, 5, 0)	Conv(10, 5, 0)	Conv(10, 3, 1)	Conv(10, 3, 1)	Conv(10, 5, 0)	Conv(10, 5, 0)	Linear(1024)	Linear(1024)	Linear(1024)	Linear(1024)
ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU
MaxPool(2)	MaxPool(2)	MaxPool(2)	MaxPool(2)	MaxPool(2)	MaxPool(2)	Linear(512)	Linear(512)	Linear(1024)	Linear(1024)
Conv(20, 5, 0)	Conv(20, 5, 0)	Conv(20, 3, 1)	Conv(20, 3, 1)	Conv(20, 3, 1)	Conv(20, 3, 1)	ReLU	ReLU	ReLU	ReLU
ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	Linear(256)	Linear(256)	Linear(10)	Linear(10)
MaxPool(2)	MaxPool(2)	MaxPool(2)	MaxPool(2)	MaxPool(2)	MaxPool(2)	ReLU	ReLU		
Linear(50)	Linear(50)	Linear(128)	Linear(128)	Linear(64)	Linear(64)	Linear(10)	Linear(10)		
ReLU	ReLU	ReLU	ReLU	ReLU	ReLU				
Linear(10)	Linear(10)	Linear(10)	Linear(10)	Linear(10)	Linear(10)				

Table 4. The architectures of local models on the CIFAR-10 dataset. ResNet* represents the layers of ResNet before the fully connected layer.

Client 1	Client 2	Client 3	Client 4	Client 5	Client 6	Client 7	Client 8	Client 9	Client 10
ResNet*-18	ResNet*-18	ResNet*-18	ResNet*-18	ResNet*-34	ResNet*-34	ResNet*-34	ResNet*-34	ResNet*-50	ResNet*-50
Linear(128)	Linear(128)	Linear(256)	Linear(256)	Linear(128)	Linear(128)	Linear(256)	Linear(256)	Linear(256)	Linear(256)
ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU	ReLU
Linear(64)	Linear(64)	Linear(10)	Linear(10)	Linear(64)	Linear(64)	Linear(10)	Linear(10)	Linear(10)	Linear(10)
ReLU	ReLU			ReLU	ReLU				
Linear(10)	Linear(10)			Linear(10)	Linear(10)				