

GPT-FL: Generative Pre-trained Model-Assisted Federated Learning

Tuo Zhang^{1,*}, Tiantian Feng^{1,*}, Samiul Alam²,
 Dimitrios Dimitriadis³, Sunwoo Lee⁴,

Mi Zhang², Shrikanth S. Narayanan¹, Salman Avestimehr¹

¹University of Southern California, ²The Ohio State University, ³Amazon, ⁴Inha University

tuozhang@usc.edu *

Abstract

In this work, we propose GPT-FL, a generative pre-trained model-assisted federated learning (FL) framework. At its core, GPT-FL leverages generative pre-trained models to generate diversified synthetic data. These generated data are used to train a downstream model on the server, which is then fine-tuned with private client data under the standard FL framework. We show that GPT-FL consistently outperforms state-of-the-art FL methods in terms of model test accuracy, communication efficiency, and client sampling efficiency. Through comprehensive ablation analysis, we discover that the downstream model generated by synthetic data plays a crucial role in controlling the direction of gradient diversity during FL training, which enhances convergence speed and contributes to the notable accuracy boost observed with GPT-FL. Also, regardless of whether the target data falls within or outside the domain of the pre-trained generative model, GPT-FL consistently achieves significant performance gains, surpassing the results obtained by models trained solely with FL or synthetic data.

1. Introduction

Federated learning (FL) is a privacy-preserving machine learning paradigm that allows clients to collaboratively train a machine learning model without sharing their private data [44]. Most existing FL studies such as [4, 25] follow the standard FL architecture, where each participating client trains a local model using its own private data and a central server aggregates these locally trained models to update a global model and send it back to the clients for the next round of training. However, although many efforts have been made [16, 30, 32], the performance of standard FL is still constrained by client drift caused by the heterogeneity in private data distribution across the clients.

To enhance the performance of FL, recent studies propose to incorporate data collected from public spaces such as the internet into the FL process [7, 15, 21, 22]. However, the performance of such public data-based approaches is heavily dependent on the quality of the collected public data. Unfortunately, obtaining the desired public data can be extremely challenging in practice and there is a lack of principled guidance on how to obtain them. To address the issues of public data-based approaches, FL methods based on synthetic data emerge [28, 41, 43, 46]. In [43, 46], a generative model is trained through knowledge distillation (KD) and the synthetic data are generated from the generative model in an *interleaved* manner *throughout* the federated training iterations. Unfortunately, these approaches are confronted with two limitations: (1) since the training of the generative model and the federated training process interleave, the quality of the synthetic data generated by the generative model before it converges can be extremely unstable. Such low-quality synthetic data would in turn jeopardize the federated training process; (2) given that KD requires clients to report model weights as teachers to transfer knowledge, they are incompatible with secure aggregation protocols [3, 36], which limits their privacy guarantee compared to standard FL.

In this work, we propose GPT-FL, a generative pre-trained model-assisted FL framework that effectively addresses the issues of existing methods. The key idea behind GPT-FL is to leverage the knowledge from the generative pre-trained models and to *decouple* synthetic data generation from the federated training process. Specifically, GPT-FL prompts the generative pre-trained models to generate diversified synthetic data. These generated data are used to train a downstream model on the server in the centralized manner, which is then fine-tuned with the private client data under the standard FL framework. By doing this, the proposed GPT-FL is able to combine the advantages of previous methods while addressing their limitations.

The proposed GPT-FL exhibits multifold merits compared to prior arts (Table 1): (1) In contrast to public data-

*Equal contribution

Table 1. Comparison of GPT-FL with existing FL methods.

	External Data	Limited to Smaller Client Model	Generate Data during FL	Data Generator Location	Client Access to Public/Generated Data	Support Data Modality	Compatibility with Secure Aggregation
FedAvg [25] FedOpt [30] FedProx [32] SCAFFOLD [16]	No	No	N/A	N/A	N/A	Any	Yes
FedDF [22] DS-FL [15] Fed-ET [7]	Public Data	No	N/A	N/A	Not Required Required Not Required	Any Any Any	No
FedGen [46] FedFTG [43] DynaFed [28] GPT-FL (Ours)	Generated Data	Yes Yes No	Yes Yes No	Client Server Server Server	Required Not Required Not Required Not Required	Only Image Only Image Only Image Any	No No Yes Yes

based FL methods, GPT-FL gets rid of the dependency on the availability of the desired public data, offering much more flexibility in its applications. (2) Compared to other generative data-based approaches, the leverage of generative pre-trained models and the decoupling between synthetic data generation from the federated training process make the generated synthetic data in GPT-FL not impacted by private data distribution on the clients and the structure of the model to be trained. (3) By leveraging the computational resources on the server, GPT-FL provides a much more efficient way to utilize external data by incorporating them into the pre-training of the downstream model, which significantly reduces the communication and computation costs of FL. (4) The generation of downstream models using synthetic data takes place on the server. As such, it thereby eliminates the need for clients to bear any additional computational burden. (5) Lastly, as GPT-FL does not alter the standard FL framework, it is fully compatible with secure aggregation protocols as in standard FL methods. More importantly, GPT-FL does not introduce any additional hyper-parameters beyond the standard FL framework. This significantly simplifies the hyper-parameter optimization process, making GPT-FL much more practically useful.

We evaluate the performance of GPT-FL by comparing it against state-of-the-art FL methods under three categories: standard FL methods, public data-based methods, and generated data-based methods on five datasets that cover both image and audio data modalities. We highlight five of our findings: (1) GPT-FL consistently outperforms state-of-the-art FL methods under both low and high data heterogeneity scenarios with significant advantages in communication and client sampling efficiency. (2) Under a zero-shot setting, *i.e.* no real-world data is available, the downstream model after centralized training with synthetic images as part of GPT-FL achieves higher performance compared to the global model based on standard FL training

with private data. On the contrary, the centralized training with synthetic audio performs worse than FL setups due to the impact of data modality and the quality of the generative pre-trained models. (3) GPT-FL does not fully rely on generated data. Regardless of whether the target data falls within or outside the domain of the pre-trained generative model, GPT-FL can largely improve model performance beyond relying solely on private data in a standard FL framework. (4) The downstream model generated by synthetic data controls gradient diversity during FL training, improving convergence speed and leading to significant accuracy gains with GPT-FL. (5) GPT-FL effectively leverages existing pre-trained downstream models to improve performance in the FL setting, similar to methods under the standard FL framework.

2. Related Work

FL with Public Data. To further mitigate client drift, recent studies propose to utilize public data (e.g., collected from the internet) in the process of federated training. For example, FedDF [22] leverages public data at the server to aggregate client models through knowledge distillation (KD). DS-FL [15] proposes a similar approach based on semi-supervised FL. Fed-ET [7] introduces a weight consensus distillation scheme using public data to train a large server model with smaller client models. However, utilizing public data for FL has several limitations: the performance of FL heavily relies on the selected public data. However, it is unclear to which extent should the public data be related to the training data to guarantee effective knowledge distillation, making it challenging to find appropriate public data for every use case [1, 37, 43]. Moreover, the involvement of KD requires clients to send model weights to the server. This requirement makes it incompatible with secure aggregation protocols, making them vulnerable to backdoor attacks [39]. Furthermore, some proposed methods [21, 22]

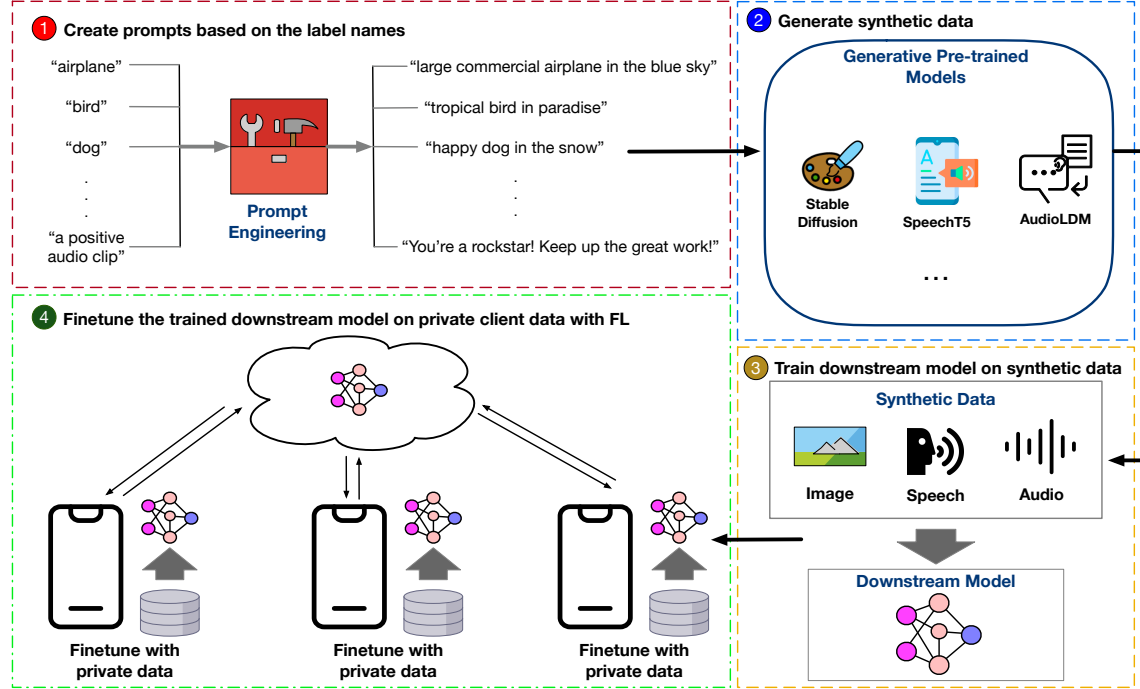


Figure 1. System overview of GPT-FL.

require clients to process the public data. Such requirement adds an extra computational burden to clients.

FL with Synthetic Data. To address the issues of public data-based approaches, FL methods based on synthetic data have been proposed [28, 41, 43, 46]. In particular, FedGen [46] proposes to train a lightweight generator on the server using an ensemble of local models in a data-free manner. The generator is then sent to the clients to regularize local training. FedFTG [43] trains a GAN-based generator where the global model acts as the discriminator. The generated data are then used to fine-tune the global model on the server after model aggregation. However, training of the generator relies heavily on the global model, which can lead to poor performance under high data heterogeneity. Additionally, the quality of training the generator is impacted by the structure of the global model [17], making the quality of the synthetic data unstable during training. Furthermore, these approaches are limited to image-related tasks, restricting their applicability to other data modalities. Specifically, both FedGen and FedFTG rely on training MLP-based or GAN-based lightweight generator networks to ensemble user information in a data-free manner, where the lightweight generator may have limitations in generating high-fidelity data. In addition, the MLP-based model is impractical to model temporal structures to signals such as audio and speech. Finally, some approaches [43, 46] could not support secure aggregation protocols due to the KD-based training, which could compromise the privacy of client data. As an alternative, DynaFed [28] proposes to

generate synthetic data via gradient inversion by applying multi-step parameter matching on global model trajectories and using the synthesized data to help aggregate the deflected clients into the global model. However, using gradient inversion for generating synthetic data could encounter limitations when dealing with high-resolution images [14]. In addition, this approach could not be directly used for other data modalities such as audio [9]. In this work, we propose GPT-FL as a solution to address these limitations.

3. GPT-FL: Generative Pre-trained Model-Assisted Federated Learning

The overall architecture of GPT-FL is illustrated in Figure 1. As shown, GPT-FL consists of four steps. First, prompts are created based on the label names at the server. These prompts are then utilized to guide the generative pre-trained models to generate synthetic data. The server uses these generated synthetic data to train a downstream model and distributes the trained model to the clients. Lastly, the clients use the trained model as the starting point, and fine-tune the model with their private data under the standard FL framework until it converges. In the following, we describe the details in each step.

3.1. Create Prompts based on Label Names

As the first step of GPT-FL, a prompt that describes the desired content of the data is required to guide the synthetic data generation process. To do so, GPT-FL requires

the clients to provide the set of label names¹ of their private local data to generate prompts. However, prior research [13, 34] shows that using only label names to generate prompts could restrict the quality and diversity of the generated synthetic data. Moreover, in FL, the server does not have access to detailed descriptions of the private data. To address these issues, GPT-FL incorporates large language models (LLMs) such as GPT-3 to expand each input class’s details and use them as prompts for synthetic data generation. As an example, for the label name “airplane”, GPT-FL uses the following query for the LLM to generate the prompt as follows:

Q: " _ _ _ _ airplane _ _ _ _ "
 Please fill in the blank and
 make it as a prompt to
 generate the image
 A: Large commercial airplane in
 the blue sky.

Moreover, inspired by [34], we randomly set the unconditional guidance scale of the Stable Diffusion model between 1 and 5 to further enrich the data diversity. In addition to the aforementioned techniques, it is worth noting that GPT-FL is flexible and compatible with other prompt engineering techniques that can be used to generate diversified synthetic data.

It should be noted that GPT-FL can employ Invertible Bloom Lookup Tables (IBLT) to encode label names before sending them to the server so that the label information of each client is not leaked to the server [10]. Specifically, each client locally encodes its unique label names into IBLT, a probabilistic data structure that can encode items in an open domain efficiently. The server linearly aggregates these IBLTs via the secure aggregation [5] and decodes the aggregated table for the union of unique label names without revealing individual label information. More details about IBLT in GPT-FL are provided in Appendix 6 via an illustrative experiment.

3.2. Generate Synthetic Data from Generative Pre-trained Model

Next, the generated prompts are used as the inputs to the generative pre-trained models to generate synthetic data. In this work, we utilize the state-of-the-art Latent Diffusion Model [31] loaded with Stable Diffusion V2.1 weights to generate synthetic images for image-based FL applications; and we utilize the state-of-the-art SpeechT5 model [2] for text-to-speech and AudioLDM model [23] for text-to-audio to generate synthetic speech and audio data, respectively. It should be noted that the proposed GPT-FL is a general

¹To protect user data privacy in FL setting, GPT-FL only requests the set of distinct label names instead of detailed label name distributions, and generates a uniform number of prompts for each label name.

framework that supports other generative pre-trained models and data modalities beyond images and audio.

3.3. Train Downstream Model on Generated Synthetic Data

With the generated synthetic data, GPT-FL trains a downstream model on the server in a centralized manner, and distributes the trained model to the clients participated in FL. This trained model acts as the initialized model for the following federated training process. One note should be emphasized from our empirical experiences is that training with synthetic data is prone to overfitting, as synthetic data tend to be highly patternized compared to real data. To mitigate the effects of overfitting, we adopt relatively large weight decay hyperparameters and small learning rates compared to training with real data. The detailed hyper-parameter selections are listed in Appendix 6.

3.4. Finetune Trained Downstream Model on Private Client Data with FL

Lastly, the clients use the trained model distributed from the server as the starting point, and finetune the model with their private data under the standard FL framework until the finetuning converges. As such, GPT-FL does not alter the standard FL framework, making it fully compatible with secure aggregation protocols as in standard FL methods. More importantly, unlike existing generated data-based approaches [28, 43, 46], GPT-FL does not introduce any additional hyper-parameters beyond the standard FL framework. This significantly simplifies the hyper-parameter optimization process, making GPT-FL much more practically useful.

4. Experiments

Datasets, Models, and Tasks. We evaluate the performance of GPT-FL on five datasets from three FL applications: image classification, speech keyword spotting, and environmental sound classification. For image classification, we conduct experiments on CIFAR-10, CIFAR-100 [18], and Oxford 102 Flower [27] using ConvNet [28], ResNet18, ResNet50 [12], and VGG19 [35]. Among them, CIFAR-10 and CIFAR-100 contain images from diverse objects whereas Oxford 102 Flower only contains images of flowers but with higher resolutions for fine-grained classification. For audio-related tasks, we choose the Google Command speech dataset [40] for keyword spotting and ESC-50 dataset [29] for environmental sound classification. We followed the previous work [45] to use the same model for these two datasets.

Data Heterogeneity. For CIFAR-10 and CIFAR-100, the training dataset is partitioned heterogeneously amongst 100 clients using the Dirichlet distribution $Dir_K(\alpha)$ with α equal to 0.1 and 0.5 following the previous work [7]. With the same method, we partition Flowers102 into 50 subsets

Table 2. Model accuracy comparison between GPT-FL and existing FL methods. For public data-based methods MOON, FedDF, DS-FL and Fed-ET, the results on CIFAR-10 and CIFAR-100 are obtained from [7], and the results on Flowers102 are marked as N/A given the practical challenge on finding a set of suitable public data that can boost its performance.

Method	Training Model	High Data Heterogeneity ($\alpha = 0.1$)			Low Data Heterogeneity ($\alpha = 0.5$)		
		CIFAR-10	CIFAR-100	Flowers102	CIFAR-10	CIFAR-100	Flowers102
FedAvg	VGG19	71.19 (± 0.27)	30.21 (± 0.32)	30.30 (± 0.16)	74.82 (± 0.23)	33.12 (± 0.13)	34.75 (± 0.90)
FedProx		72.45 (± 0.13)	31.51 (± 0.11)	33.23 (± 0.24)	75.24 (± 0.19)	33.64 (± 0.08)	40.56 (± 0.19)
SCAFFOLD		75.12 (± 0.20)	30.61 (± 0.57)	26.75 (± 0.50)	78.69 (± 0.15)	34.91 (± 0.61)	33.21 (± 0.41)
MOON	VGG19	75.68 (± 0.51)	33.72 (± 0.89)	N/A	81.17 (± 0.41)	42.15 (± 0.72)	N/A
FedDF		73.81 (± 0.42)	31.87 (± 0.46)	N/A	76.55 (± 0.32)	37.87 (± 0.31)	N/A
DS-FL		65.27 (± 0.53)	29.12 (± 0.51)	N/A	68.44 (± 0.47)	33.56 (± 0.55)	N/A
Fed-ET		78.66 (± 0.31)	35.78 (± 0.45)	N/A	81.13 (± 0.28)	41.58 (± 0.36)	N/A
FedGen	ConvNet ²	42.05 (± 0.93)	26.64 (± 0.66)	Not Converged	54.86 (± 0.13)	34.03 (± 0.42)	Not Converged
DynaFed		71.59 (± 0.10)	36.08 (± 0.15)	Not Converged	75.66 (± 0.21)	43.82 (± 0.30)	Not Converged
GPT-FL	VGG19	82.16 (± 0.13)	47.80 (± 0.32)	70.56 (± 0.34)	82.17 (± 0.20)	48.39 (± 0.17)	74.84 (± 0.43)
	ConvNet	72.62 (± 0.24)	42.66 (± 0.19)	37.91 (± 0.43)	77.18 (± 0.21)	47.89 (± 0.28)	48.61 (± 0.51)

due to its relatively small size. For audio datasets, Google Speech Command is partitioned over speaker IDs, making the dataset naturally non-IID distributed. It contains a total of 105,829 audio recordings collected from 2,618 speakers. The training set includes the recordings from 2,112 speakers and the test set includes the rest. To create non-IID data distributions on ESC-50, we followed the previous work [45] to partition ESC-50 into 100 subsets using $Dir_K(\alpha)$ with α equal to 0.1.

Baselines and Evaluation Metrics. We compare GPT-FL against three categories of baselines: 1) standard FL methods without the use of public or generated synthetic data – FedAvg, FedProx, and Scaffold; 2) FL methods that involve the use of public data – MOON, FedDF, DS-FL, and Fed-ET; and 3) FL methods that utilize generated synthetic data – FedGen and DynaFed³. We use the test accuracy of the trained model as our evaluation metric. We run experiments with three different random seeds and report the average and standard deviation. The details of the hyper-parameter selection of each dataset and experiment are described in Appendix.

4.1. Performance Comparison with State-of-the-Art FL Methods

First, we compare the performance of GPT-FL with state-of-the-art FL methods. To enforce fair comparisons, in this experiment, we choose to evaluate on the three image datasets (CIFAR-10, CIFAR-100 and Flowers102) since baseline methods MOON, FedGen and DynaFed only support image data. Moreover, we used the same models (VGG19 and ConvNet) and experiment settings as previous work [7, 28]. In each communication round, we randomly sample 10 clients from 100 clients for CIFAR and use all 50 clients for Flowers102. We choose FedAvg as the FL optimizer. All the training starts from random initialization and

total number of communication rounds is set to 500.

Overall Performance. Table 2 summarizes our results. We make three key observations: (1) GPT-FL consistently outperforms all the baselines we selected in Table 2 under both low and high data heterogeneity scenarios across all three datasets. (2) In direct comparison with state-of-the-art generated data-based FL methods, although FedGen and DynaFed perform reasonably well on CIFAR-10 and CIFAR-100, they do not converge on Flowers102 whose images have higher resolutions than CIFAR. Moreover, both FedGen and DynaFed fail to converge when training a larger VGG19 model on Flowers102 and even lower-resolution CIFAR-10/100. In contrast, GPT-FL not only converges but also achieves state-of-the-art accuracy on Flowers102. More importantly, GPT-FL is able to support larger model, and its accuracy is significantly higher than the smaller ConvNet. (3) For Flowers102, as both public data-based and generated data-based FL methods are confronted with challenges, the only viable options are standard FL methods and GPT-FL. As shown, with the same model, GPT-FL outperforms standard FL methods by a significant margin.

Communication Efficiency. Besides model accuracy, we also compare the communication costs of GPT-FL with existing FL methods on CIFAR-10/100 under high data heterogeneity, where communication cost is measured as the total number of model parameters communicated between the server and clients during federated training until reaching a target model test accuracy. Specifically, Figure 2 shows the communication cost comparison between standard FL methods, public data-based methods⁴, and GPT-FL under VGG19; and Figure 3 shows the communication cost comparison between generative data-based methods and GPT-FL under ConvNet. The target test ac-

³We did not compare with FedFTG because its code is not open-source, and we could not reproduce their results following the paper.

⁴We do not compare with FedDF and DS-FL as they do not achieve competitive model accuracy.

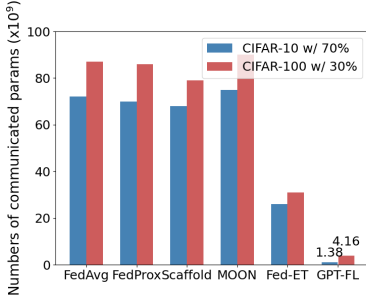


Figure 2. Communication costs of standard FL methods, public data-based methods and GPT-FL to achieve the target accuracy.

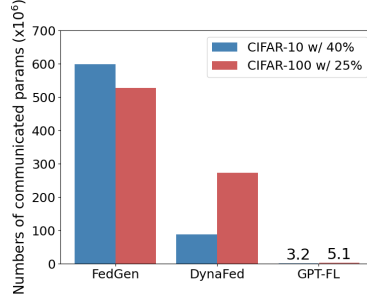


Figure 3. Communication costs of generated data-based methods and GPT-FL to achieve the target test accuracy.

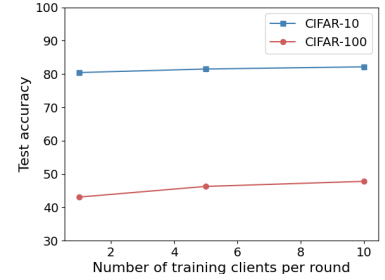


Figure 4. Test accuracy of GPT-FL for CIFAR-10/100 under different client sampling rates.

curacies in Figure 3 are set to be lower given the low accuracies achieved by FedGen. As shown, GPT-FL has the least communication cost among all the methods, achieving up to 94% communication reduction compared to the best-performed public data-based baseline Fed-ET and 98% communication reduction compared to the best-performed generated data-based baseline DynaFed. These results highlight the significant advantage of GPT-FL in communication reduction over state-of-the-art FL methods.

Client Sampling Efficiency. One critical hyperparameter of FL is the client sampling rate in each communication round during the federated training process. In Figure 4, we plot the test model accuracies obtained by GPT-FL under low, medium, and high client sampling rates on CIFAR-10/100 with VGG19 under high data heterogeneity. As shown, even with a single participating client per round, GPT-FL is able to achieve 80.44% and 43.07% test accuracy on CIFAR-10 and CIFAR-100 respectively. This performance already surpasses all the other FL methods listed in Table 2, which employs 9 times more clients for training per round. These results highlight the significant advantage of GPT-FL in client sampling efficiency over state-of-the-art FL methods, making GPT-FL a very attractive solution in challenging scenarios where not many clients can participate at the same time.

4.2. Understanding GPT-FL

(1) Can we only rely on centralized training with synthetic data to achieve competitive results compared to Federated Learning with private data?

To answer this question, we compare the model performance between generated downstream model by centralized training with synthetic data and the global model by standard FL training with private data on both image and audio benchmark datasets. Different from the previous section, we select ResNet18 and ResNet50 models for CIFAR-10 and CIFAR-100 dataset, respectively. We choose the models proposed in the FedAudio Benchmark [45] for audio tasks. We report the best F1 score for the audio datasets.

The results are summarized in Table 3.

Impact of Out-of-Domain Data Generation. We choose the ESC-50 and Google Speech Commands datasets to examine the impact of out-of-domain data generation for the generative pre-trained model. We did not conduct a similar analysis for the image datasets as the LAION-5B [33] open-source dataset for training the Stable Diffusion model we used is a vast collection of publicly available datasets, including nearly all relevant ones for our experiments.

Our experiments show that synthetic image outperforms synthetic audio regarding model performance when using centralized training. We observed that centralized training with synthetic images achieves higher accuracy than FL setups for all three image benchmark datasets. In contrast, centralized training with synthetic audio performs worse than FL setups for ESC-50 and Google Speech Command datasets. The finding from the Google Speech Command experiments aligns with the previous study [20] that utilizes pure synthetic speech data to train the automatic speech recognition system leading to substantial performance degradation. One plausible explanation is related to the relatively small training data sizes (approximately 400M sentences) and constrained domain knowledge (book corpus) compared to training other generative pre-trained models like Stable Diffusion. For example, using human inspection, we discovered that the TTS model fails to synthesize simple spoken words like "house". This deficiency may originate from the lack of short-spoken utterance samples in training data. In addition, the synthesized speech often lacks diversity due to the limited range of speakers represented in the training dataset. On the other hand, there is insufficient knowledge of the audio generation models, making the model performance of using the synthesized audio data as training data remains unknown. However, our manual inspection revealed that the model frequently encounters difficulties in generating audio samples, such as generated audio related to water sounds often sounds like music. This issue could be largely associated with the relatively small data size in pre-training compared to other models.

Table 3. Accuracy performance of the generated downstream model and standard FL on benchmark datasets. "1x Synthetic" represents the size of synthetic data is one time as the real data.

	Dataset	1x Synthetic	2x Synthetic	3x Synthetic	FedAvg	FedOpt
Image Data	CIFAR-10	61.48 (± 0.08)	67.41 (± 0.40)	75.65 (± 0.09)	64.48 (± 0.13)	72.68 (± 0.22)
	CIFAR-100	24.70 (± 0.00)	33.41 (± 0.01)	41.76 (± 0.03)	25.89 (± 0.67)	20.85 (± 0.14)
	Flowers102	24.94 (± 0.57)	28.26 (± 0.14)	31.29 (± 0.18)	30.30 (± 0.16)	26.43 (± 0.09)
Audio Data	Google Command	24.78 (± 0.04)	25.65 (± 0.07)	26.24 (± 0.01)	73.68 (± 0.49)	83.01 (± 0.23)
	ESC-50	6.89 (± 0.29)	8.68 (± 0.35)	12.72 (± 0.31)	22.76 (± 1.01)	32.49 (± 0.57)

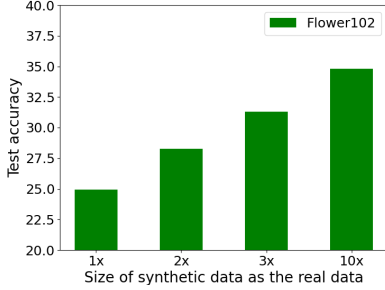


Figure 5. Impact of synthetic data sample number to the generated downstream model.

Impact of Numbers of Synthetic Data. With both image and audio data, one commonality is centralized training with synthetic data can benefit from increasing the number of synthetic data. To validate this finding, we test the impact of numbers of synthetic data on the performance of the generated model on the Flowers102 dataset, where we increase the size of the synthetic data up to ten times that of the real data. As shown in Figure 5, our experimental results demonstrate that as we enlarge the amount of synthetic data, the performance of the model improves. One justification for this finding is that enlarging the number of synthetic data enriches the diversity and increases overlap between the synthetic and real data, allowing the model to learn more robust and generalizable features. Even if the data is generated randomly by the label name without any other diversity-enriching guidance from the real data, with more synthetic data, there is an increasing chance that some of these additional synthetic data overlap with the real data, allowing the model to perform better on the real test data.

(2) What benefits does GPT-FL bring?

We explore the benefits that GPT-FL provides for custom models that are built on top of downstream models generated from synthetic data. Specifically, we want to examine how fine-tuning these downstream models with private data under the FL framework can lead to performance improvements. To demonstrate how GPT-FL can be integrated with existing FL server optimizers, we evaluate its performance with both FedAvg and FedOpt as the server aggregator. Our

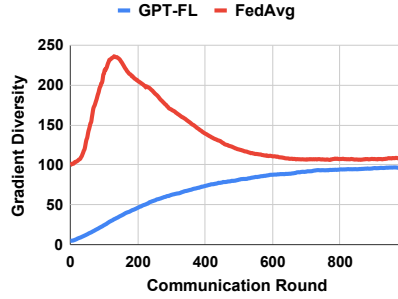


Figure 6. Smoothed Gradient diversity of client updates during training on Google speech commands dataset.

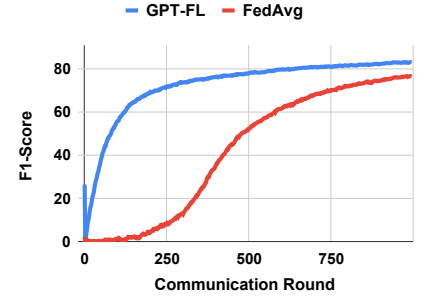


Figure 7. Learning curve of the global model during training on Google speech commands dataset.

experimental results are presented in Table 4.

Effectiveness of Private Data. Our experiments demonstrate the effectiveness of incorporating private data with FL into the finetuning process of the downstream model generated from synthetic data. As shown in Table 4, regardless of the modality and quality of the synthetic data used to generate the downstream model, FL fine-tuning leads to significant performance gains, outperforming the ones trained solely with FL or CL combined with synthetic training by a large margin. Furthermore, we observe that fine-tuning with private data can especially benefit the cases for out-of-domain synthetic data, such as in the audio data. For example, GPT-FL with FedOpt could achieve 43.46 test accuracy in the ECS-50 dataset, which nearly provides two times increment than standard FL and three times increment than centralized training by synthetic data. These results suggest that leveraging private data with FL in the fine-tuning process can greatly enhance the performance of synthetic data-generated models, making them more suitable for real-world applications.

Generated Downstream Model Helps FL Optimization. To gain a comprehensive understanding of why the custom models built using GPT-FL provide benefits to performance improvements, we decided to compare the gradient diversity between model weights initialized by GPT-FL and random initialization. Specifically, we apply the definition of the gradient diversity introduced from [42] by adapting the gradients g_i to client update Δ_i :

Table 4. Accuracy comparison between generated downstream model, standard FL and GPT-FL. Differ from the experiments shown in Table 1, the CIFAR-10 and Flowers102 datasets are trained with ResNet18 model and the CIFAR-100 dataset is trained with ResNet50 model. "ΔMetric" represents the accuracy increment by GPT-FL on top of the generated downstream model.

Dataset	3x Synthetic	FedAvg	FedOpt	GPT-FL w/ FedAvg	GPT-FL w/ FedOpt	ΔMetric
CIFAR-10	75.65 (± 0.09)	64.48 (± 0.13)	72.68 (± 0.22)	81.38 (± 0.05)	79.08 (± 0.17)	↑ 5.73
CIFAR-100	41.76 (± 0.03)	25.89 (± 0.67)	20.85 (± 0.14)	62.83 (± 0.31)	48.80 (± 0.12)	↑ 21.07
Flowers102	31.29 (± 0.18)	30.30 (± 0.16)	26.43 (± 0.09)	70.56 (± 0.34)	77.57 (± 0.03)	↑ 46.28
Google Command	26.24 (± 0.01)	73.68 (± 0.49)	83.01 (± 0.23)	81.90 (± 0.20)	83.46 (± 0.11)	↑ 57.22
ESC-50	12.72 (± 0.31)	22.76 (± 1.01)	32.49 (± 0.57)	41.80 (± 0.32)	43.46 (± 0.30)	↑ 30.74

Table 5. Accuracy performance comparison between generated downstream model, standard federated learning and GPT-FL. All the training is initialized by ImageNet-based pre-train model.

Dataset	3x Synthetic	FedAvg	FedOpt	GPT-FL w/ FedAvg	GPT-FL w/ FedOpt
CIFAR-10	72.65 (± 0.05)	66.10 (± 0.03)	79.08 (± 0.39)	75.87 (± 0.73)	82.20 (± 0.61)
CIFAR-100	42.30 (± 0.01)	62.83 (± 0.03)	45.27 (± 0.10)	66.84 (± 0.05)	66.03 (± 0.02)
Flowers102	41.05 (± 0.26)	80.73 (± 0.01)	87.33 (± 0.29)	86.18 (± 0.04)	88.66 (± 0.40)

$$\Delta_S = \frac{\sum_{i \in S} \|\Delta_i\|^2}{\|\sum_{i \in S} \Delta_i\|^2} \quad (1)$$

where S is the set of sampled clients in each communication round and i represents the client index. As shown in Figure 6, the gradient diversity plot for FedAvg reveals that GPT-FL displays lower initial gradient diversity compared to random initialization. Over training time, both GPT-FL and random initialization converge to similar gradient diversity levels, consistent with the performance curve in Figure 7, where a larger Δ_S corresponds to slower convergence rate. This aligns with prior findings [26], indicating that starting from a pre-trained model leads to less variation in local client updates, potentially addressing the client drift issue.

Harmonization With Existing Pre-train Model. As the standard FL framework, GPT-FL could also benefit from other existing pre-train models. Specifically, besides training from scratch, GPT-FL could utilize the existing pre-train model to start training the synthetic data to generate downstream model and finetune it again with private data in FL. Table 5 presents the performance evaluation of GPT-FL on top of the pre-trained models for image datasets. We follow the approach from prior work [26] and use the ImageNet pre-trained model available in the PyTorch Torchvision library. Our experiments show that GPT-FL achieves better results compared to training solely with FL or synthetic data, as reported in Table 4. Notably, the improvement in performance is consistent across

three image benchmark datasets, with a gain ranging from 1% to 11% compared to the results in Table 4. These results demonstrate that GPT-FL can effectively leverage pre-trained models to improve performance in the FL setting.

5. Conclusion

We present GPT-FL, a generative pre-trained model-assisted federated learning framework. GPT-FL leverages the generative pre-trained model to generate diversified synthetic data for a wide range of data modalities before FL training. This synthetic data is then utilized to construct a downstream model, which undergoes fine-tuning with private data within a standard FL framework. Our experimental results showcase the remarkable performance of GPT-FL when compared to state-of-the-art FL methods. Through detailed ablation studies, we demonstrate that GPT-FL is a flexible and applicable framework solution to the challenges associated with cross-device FL scenarios.

Limitations and Future Works. Due to the limitations in our computational resources, we cannot further scale up the synthetic data volume in our study, as it may take several weeks for the generation. We do not investigate the larger model sizes in our current study, which we will pursue it as our future work. We also want to explore the expansions of the GPT-FL framework. GPT-FL seamlessly integrates with the vanilla FL framework, allowing for harmonization with most of the existing FL methods. We are interested in exploring the combination of the public-data-based FL aggregation scheme and the GPT-FL framework by replacing the public data.

References

- [1] Samiul Alam, Luyang Liu, Ming Yan, and Mi Zhang. Fedrolex: Model-heterogeneous federated learning with rolling sub-model extraction. *ArXiv*, abs/2212.01548, 2022. 2
- [2] Junyi Ao, Rui Wang, Long Zhou, Shujie Liu, Shuo Ren, Yu Wu, Tom Ko, Qing Li, Yu Zhang, Zhihua Wei, Yao Qian, Jinyu Li, and Furu Wei. Speecht5: Unified-modal encoder-decoder pre-training for spoken language processing. In *Annual Meeting of the Association for Computational Linguistics*, 2021. 4
- [3] Keith Bonawitz, Vladimir Ivanov, Ben Kreuter, Antonio Marcedone, H. B. McMahan, Sarvar Patel, Daniel Ramage, Aaron Segal, and Karn Seth. Practical secure aggregation for privacy-preserving machine learning. *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, 2017. 1
- [4] Keith Bonawitz, Hubert Eichner, Wolfgang Grieskamp, Dzmitry Huba, Alex Ingerman, Vladimir Ivanov, Chloé Kidon, Jakub Konečný, Stefano Mazzocchi, H. B. McMahan, Timon Van Overveldt, David Petrou, Daniel Ramage, and Jason Roselander. Towards federated learning at scale: System design. *ArXiv*, abs/1902.01046, 2019. 1
- [5] K. A. Bonawitz, Vladimir Ivanov, Ben Kreuter, Antonio Marcedone, H. Brendan McMahan, Sarvar Patel, Daniel Ramage, Aaron Segal, and Karn Seth. Practical secure aggregation for federated learning on user-held data. In *NIPS Workshop on Private Multi-Party Machine Learning*, 2016. 4
- [6] Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101—mining discriminative components with random forests. In *Computer Vision—ECCV 2014: 13th European Conference, Zurich, Switzerland, September 6–12, 2014, Proceedings, Part VI 13*, pages 446–461. Springer, 2014. 2
- [7] Yae Jee Cho, Andre Manoel, Gauri Joshi, Robert Sim, and Dimitrios Dimitriadis. Heterogeneous ensemble knowledge transfer for training large models in federated learning. In *International Joint Conference on Artificial Intelligence*, 2022. 1, 2, 4, 5, 3
- [8] Muhammad EH Chowdhury, Tawsifur Rahman, Amith Khandakar, Rashid Mazhar, Muhammad Abdul Kadir, Zaid Bin Mahbub, Khandakar Reajul Islam, Muhammad Salman Khan, Atif Iqbal, Nasser Al Emadi, et al. Can ai help in screening viral and covid-19 pneumonia? *Ieee Access*, 8:132665–132676, 2020. 2
- [9] Trung Dang, Om Thakkar, Swaroop Indra Ramaswamy, Rajiv Mathews, Peter Chin, and Francoise Beaufays. A method to reveal speaker identity in distributed asr training, and how to counter it. *ICASSP 2022 - 2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 4338–4342, 2021. 3
- [10] Adrià Gascón, Peter Kairouz, Ziteng Sun, and Ananda Theertha Suresh. Federated heavy hitter recovery under linear sketching. *ArXiv*, abs/2307.13347, 2023. 4, 1
- [11] Michael T. Goodrich and Michael Mitzenmacher. Invertible bloom lookup tables. *2011 49th Annual Allerton Conference on Communication, Control, and Computing (Allerton)*, pages 792–799, 2011. 1
- [12] Kaiming He, X. Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 770–778, 2015. 4, 2, 3
- [13] Ruifei He, Shuyang Sun, Xin Yu, Chuhui Xue, Wenqing Zhang, Philip H. S. Torr, Song Bai, and Xiaojuan Qi. Is synthetic data from generative models ready for image recognition? *ArXiv*, abs/2210.07574, 2022. 4
- [14] Yangsibo Huang, Samyak Gupta, Zhao Song, Kai Li, and Sanjeev Arora. Evaluating gradient inversion attacks and defenses in federated learning. In *Neural Information Processing Systems*, 2021. 3
- [15] Sohei Itahara, Takayuki Nishio, Yusuke Koda, Masahiro Morikura, and Koji Yamamoto. Distillation-based semi-supervised federated learning for communication-efficient collaborative training with non-iid private data. *IEEE Transactions on Mobile Computing*, 22:191–205, 2020. 1, 2
- [16] Sai Praneeth Karimireddy, Satyen Kale, Mehryar Mohri, Sashank J. Reddi, Sebastian U. Stich, and Ananda Theertha Suresh. Scaffold: Stochastic controlled averaging for federated learning. In *International Conference on Machine Learning*, 2019. 1, 2
- [17] Jang-Hyun Kim, Jinuk Kim, Seong Joon Oh, Sangdoo Yun, Hwanjun Song, Joonhyun Jeong, Jung-Woo Ha, and Hyun Oh Song. Dataset condensation via efficient synthetic-data parameterization. In *International Conference on Machine Learning*, 2022. 3
- [18] Alex Krizhevsky. Learning multiple layers of features from tiny images. 2009. 4, 2
- [19] Abhinav Kumar, Vishal Purohit, Vandana Bharti, Rishav Singh, and Sanjay Kumar Singh. Medisecfed: Private and secure medical image classification in the presence of malicious clients. *IEEE Transactions on Industrial Informatics*, 18:5648–5657, 2022. 2
- [20] Jason Li, Ravi Gadde, Boris Ginsburg, and Vitaly Lavrukhin. Training neural speech recognition systems with synthetic speech augmentation. *arXiv preprint arXiv:1811.00707*, 2018. 6
- [21] Qinbin Li, Bingsheng He, and Dawn Xiaodong Song. Model-contrastive federated learning. *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10708–10717, 2021. 1, 2
- [22] Tao Lin, Lingjing Kong, Sebastian U. Stich, and Martin Jaggi. Ensemble distillation for robust model fusion in federated learning. *ArXiv*, abs/2006.07242, 2020. 1, 2
- [23] Haohe Liu, Zehua Chen, Yiitan Yuan, Xinhao Mei, Xubo Liu, Danilo P. Mandic, Wenwu Wang, and Mark D. Plumbley. Audioldm: Text-to-audio generation with latent diffusion models. *ArXiv*, abs/2301.12503, 2023. 4
- [24] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *International Conference on Learning Representations*, 2017. 4
- [25] H. B. McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Agüera y Arcas. Communication-efficient

- learning of deep networks from decentralized data. In *International Conference on Artificial Intelligence and Statistics*, 2016. 1, 2
- [26] John Nguyen, Jianyu Wang, Kshitiz Malik, Maziar Sanjabi, and Michael G. Rabbat. Where to begin? on the impact of pre-training and initialization in federated learning. *ArXiv*, abs/2210.08090, 2022. 8
- [27] Maria-Elena Nilsback and Andrew Zisserman. Automated flower classification over a large number of classes. *2008 Sixth Indian Conference on Computer Vision, Graphics & Image Processing*, pages 722–729, 2008. 4, 2
- [28] Renjie Pi, Weizhong Zhang, Yueqi Xie, Jiahui Gao, Xiaoyu Wang, Sunghun Kim, and Qifeng Chen. Dynafed: Tackling client data heterogeneity with global dynamics. *arXiv preprint arXiv:2211.10878*, 2022. 1, 2, 3, 4, 5
- [29] Karol J. Piczak. ESC: Dataset for Environmental Sound Classification. In *Proceedings of the 23rd Annual ACM Conference on Multimedia*, pages 1015–1018. ACM Press, 2015. 4, 3
- [30] Sashank J. Reddi, Zachary B. Charles, Manzil Zaheer, Zachary Garrett, Keith Rush, Jakub Konečný, Sanjiv Kumar, and H. B. McMahan. Adaptive federated optimization. *ArXiv*, abs/2003.00295, 2020. 1, 2
- [31] Robin Rombach, A. Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10674–10685, 2021. 4
- [32] Anit Kumar Sahu, Tian Li, Maziar Sanjabi, Manzil Zaheer, Ameet Talwalkar, and Virginia Smith. Federated optimization in heterogeneous networks. *arXiv: Learning*, 2018. 1, 2
- [33] Christoph Schuhmann, Romain Beaumont, Richard Vencu, Cade Gordon, Ross Wightman, Mehdi Cherti, Theo Coombes, Aarush Katta, Clayton Mullis, Mitchell Wortsman, Patrick Schramowski, Srivatsa Kundurthy, Katherine Crowson, Ludwig Schmidt, Robert Kaczmarczyk, and Jenia Jitsev. Laion-5b: An open large-scale dataset for training next generation image-text models. *ArXiv*, abs/2210.08402, 2022. 6
- [34] Jordan Shipard, Arnold Wiliem, Kien Nguyen Thanh, Wei Xiang, and Clinton Fookes. Diversity is definitely needed: Improving model-agnostic zero-shot classification via stable diffusion. 2023. 4
- [35] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. *CoRR*, abs/1409.1556, 2014. 4, 2, 3
- [36] Jinhyun So, Chaoyang He, Chien-Sheng Yang, Songze Li, Qian Yu, Ramy E. Ali, Basak Guler, and Salman Avestimehr. Lightsecagg: a lightweight and versatile design for secure aggregation in federated learning. In *Conference on Machine Learning and Systems*, 2021. 1
- [37] Samuel Stanton, Pavel Izmailov, P. Kirichenko, Alexander A. Alemi, and Andrew Gordon Wilson. Does knowledge distillation really work? *ArXiv*, abs/2106.05945, 2021. 2
- [38] TensorFlow. Private heavy hitters, 2023. 1
- [39] Hongyi Wang, Kartik K. Sreenivasan, Shashank Rajput, Harit Vishwakarma, Saurabh Agarwal, Jy yong Sohn, Kangwook Lee, and Dimitris Papailopoulos. Attack of the tails: Yes, you really can backdoor federated learning. *ArXiv*, abs/2007.05084, 2020. 2
- [40] Pete Warden. Speech commands: A dataset for limited-vocabulary speech recognition. *ArXiv*, abs/1804.03209, 2018. 4, 3
- [41] Achintha Wijesinghe, Songyang Zhang, and Zhi Ding. Ps-fedgan: An efficient federated learning framework based on partially shared generative adversarial networks for data privacy. *ArXiv*, abs/2305.11437, 2023. 1, 3
- [42] Dong Yin, Ashwin Pananjady, Max Lam, Dimitris Papailopoulos, Kannan Ramchandran, and Peter Bartlett. Gradient diversity: a key ingredient for scalable distributed learning. In *International Conference on Artificial Intelligence and Statistics*, pages 1998–2007. PMLR, 2018. 7
- [43] Lin Zhang, Li Shen, Liang Ding, Dacheng Tao, and Ling-Yu Duan. Fine-tuning global model via data-free knowledge distillation for non-iid federated learning. *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10164–10173, 2022. 1, 2, 3, 4
- [44] Tuo Zhang, Lei Gao, Chaoyang He, Mi Zhang, Bhaskar Krishnamachari, and Salman Avestimehr. Federated learning for the internet of things: Applications, challenges, and opportunities. *IEEE Internet of Things Magazine*, 5:24–29, 2021. 1
- [45] Tuo Zhang, Tiantian Feng, Samiul Alam, Sunwoo Lee, Mi Zhang, Shrikanth S. Narayanan, and Salman Avestimehr. Fedaudio: A federated learning benchmark for audio tasks. In *ICASSP 2023 - 2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 1–5, 2023. 4, 5, 6, 3
- [46] Zhuangdi Zhu, Junyuan Hong, and Jiayu Zhou. Data-free knowledge distillation for heterogeneous federated learning. *Proceedings of machine learning research*, 139:12878–12889, 2021. 1, 2, 3, 4