

The Marriage of Foundation Models and Federated Learning: A Survey

Anonymous ACL submission

Abstract

The recent development of Foundation Models (FMs), represented by large language models, vision transformers, and multimodal models, has been making a significant impact on both academia and industry. Compared with small-scale models, FMs have a much stronger demand for high-volume data during the pre-training phase. Although general FMs can be pre-trained on data collected from open sources such as the Internet, domain-specific FMs need proprietary data, posing a practical challenge regarding the amount of data available due to privacy concerns. Federated Learning (FL) is a collaborative learning paradigm that breaks the barrier of data availability from different participants. Therefore, it provides a promising solution to customize and adapt FMs to a wide range of domain-specific tasks using distributed datasets whilst preserving privacy. This survey paper discusses the potentials and challenges of synergizing FL and FMs and summarizes core techniques, future directions, and applications.

1 Introduction

The landscape of Artificial Intelligence (AI) has been revolutionized by the emergence of *Foundation Models (FMs)* (Bommasani et al., 2021), such as BERT (Devlin et al., 2019), GPT series (Brown et al., 2020; OpenAI, 2022, 2023), and LLaMA series (Touvron et al., 2023a,b) in Natural Language Processing (NLP); ViTs (Dosovitskiy et al., 2021) and SAM (Kirillov et al., 2023) in Computer Vision (CV); CLIP (Radford et al., 2021), DALL-E (Ramesh et al., 2021), and Gemini (Google, 2023) in multimodal applications. These FMs have become pivotal in a myriad of AI applications across diverse domains. Their superb capability to generalize across tasks and domains stems from their pre-training on extensive datasets (Gunasekar et al., 2023), which imbues them with a profound understanding of language, vision, and multimodal data.

While general-purpose FMs can leverage openly accessible data from the Internet, domain-specific FMs require proprietary data. However, it is challenging to collect vast amounts of proprietary data and perform centralized pre-training or fine-tuning for domain-specific FMs, due to privacy restrictions (Jo and Gebru, 2020; GDPR, 2016; CCPA, 2023). Particularly in domains such as law, healthcare, and finance, where data is inherently privacy-sensitive, there is a pressing need for stringent privacy safeguards. Furthermore, given that data often constitutes a pivotal asset for enterprises, its widespread distribution is prohibitive. Consequently, there is an urgent need for novel strategies to handle data availability and facilitate model training, thereby unlocking the potential of domain-specific FMs whilst respecting data privacy.

To address the challenges associated with data privacy in model training, Federated Learning (FL) (McMahan et al., 2017) has emerged as a promising paradigm. FL facilitates collaborative model training across decentralized clients without the need for sharing raw data, thus ensuring privacy preservation. Concretely, FL encompasses periodic interactions between the server and decentralized clients for the exchange of trainable model parameters, without the requirement for private client data. Recognizing such a benefit, *integrating FMs with FL presents a compelling solution for domain-specific FMs* (Zhuang et al., 2023; Yu et al., 2023d).

Despite the potential synergies between FL and FMs, the field is still nascent, lacking a comprehensive understanding of challenges, methodologies, and directions. This survey aims to bridge this gap by providing a thorough exploration of the integration of FMs and FL. We delve into the motivations and challenges of combining these two paradigms, highlight representative techniques, and discuss future directions and applications. By elucidating the intersection of FL and FMs, we aim to catalyze further research and innovation in this burgeoning

ing area, ultimately advancing the development of privacy-aware, domain-specific AI models.

The paper continues as follows: The next section introduces background on FMs and FL. Section 3 presents the motivation and challenges for marrying FMs and FL. Section 4 highlights representative techniques. Before concluding, we discuss future directions and applications in Section 5.

2 Background

2.1 Foundation Models

An FM is a model that can be adapted to a wide array of tasks through fine-tuning after initial pre-training (Bommasani et al., 2021). The lifecycle of FMs typically involves pre-training on extensive generic data to establish the basis of their abilities (Bubeck et al., 2023), followed by adaptation to downstream tasks such as domain-specific question answering (Zhang et al., 2023b), and ultimately application in various domains.

FMs have sparked a significant paradigm shift in various fields of AI such as NLP, CV, speech and acoustics, and beyond. In the realm of NLP, the most prominent example is Large Language Models (LLMs) with substantial parameter sizes (Zhao et al., 2023). These models, such as ChatGPT and GPT-4 (OpenAI, 2022, 2023), demonstrate exceptional prowess in natural language understanding and generation, enabling them to comprehend and respond to user inputs with remarkable contextual relevance. This capability proves invaluable in applications like customer service, virtual assistants, and chatbots, where effective communication is paramount. Moreover, LLMs eliminate the need for training models from scratch for specific tasks, be it machine translation, document summarization, text generation, or other language-related tasks.

In the realm of CV and other modalities, FMs have also made remarkable progress. Vision Transformers (ViTs) (Dosovitskiy et al., 2021) segment images into distinct patches, which serve as inputs for transformer architectures. SAM (Kirillov et al., 2023) can segment anything in images according to the input prompts. CLIP (Radford et al., 2021) bridges the gap between text and images through contrastive learning. DALL-E, proposed by Ramesh et al. (2021), generates images from textual descriptions, expanding the possibilities of creative image generation. Additionally, models like Gato, introduced by Reed et al. (2022), exhibit versatility by being applicable across various tasks

such as conversational agents, robotic control, and gaming.

2.2 Federated Learning

FL (McMahan et al., 2017) is a distributed learning paradigm that involves periodic interactions between the server and decentralized clients for the exchange of trainable model parameters, without the need for private client data. FL offers a privacy-preserving and efficient way to train models on a large scale and diverse data (Kairouz et al., 2021), leading to its application across various domains such as healthcare (Lincy and Kowshalya, 2020; Rieke et al., 2020; Joshi et al., 2022), finance (Chatterjee et al., 2023; Liu et al., 2023b), and smart cities (Ramu et al., 2022; Pandya et al., 2023).

3 FM-FL Marriage: Motivation and Challenges

In this section, we first motivate the marriage of FMs and FL (§3.1), then summarize the key challenges stemming from the FM-FL marriage (§3.2).

3.1 Motivation

The AI sector widely concurs that the capabilities of FMs are fundamentally driven by large-scale and high-quality datasets (Bommasani et al., 2021; Kaplan et al., 2020; Gunasekar et al., 2023), which encompass both public and private sources. Leveraging private data for training FMs presents considerable challenges owing to privacy concerns. Privacy regulations and data protection laws often prohibit sharing sensitive information (GDPR, 2016; CCPA, 2023), limiting the feasibility of traditional data-centralized training processes.

The FM-FL marriage represents a compelling collaboration that utilizes the strengths of each to address their respective limitations, embodying a complementary relationship (Zhuang et al., 2023; Li and Wang, 2024).

FL expands data availability for FMs. By leveraging data from a wide range of sources in a privacy-preserving manner, FL makes it possible to build models on sensitive data in specific domains, such as healthcare (Lincy and Kowshalya, 2020; Joshi et al., 2022; Rieke et al., 2020) and finance (Chatterjee et al., 2023; Liu et al., 2023b). This enhances the diversity and volume of training data, improving model robustness and adaptability. Moreover, FL enables the integration of personal and task-specific data, allowing FMs to

be customized for personal applications. For instance, Google has trained next-word-prediction language models on mobile keyboard input data with FL to improve user experience (Xu et al., 2023c; Bonawitz et al., 2021).

FMs boost FL with knowledge and understanding capabilities. By pre-training on large-scale generic data, FMs acquire essential knowledge and understanding capabilities (Brown et al., 2020). Firstly, they benefit FL systems by offering advanced feature representations and learning capabilities from the outset. Secondly, leveraging the pre-learned knowledge of FMs accelerates the FL process, enabling efficient and effective adaptation to specific tasks with minimal additional training. Thirdly, FMs’ powerful generative capabilities could help FL overcome the data heterogeneity challenge by synthesizing extra data, thus enhancing model convergence (Huang et al., 2024).

3.2 Challenges

In this part, we discuss challenges emerging from the FM-FL marriage in three levels: *Task-Level*, *System-Level*, as well as *Governance-Level*. Due to the space limitation, we only list the most representative challenges for each level.

3.2.1 Task-Level Challenges

Task-level challenges stem from the adaptation of an FM to a specific downstream task (e.g., by fine-tuning) in a federated setting. Challenges include:

Data Heterogeneity. Performance degradation in FL, attributed to heterogeneous data distributions among clients, is a well-recognized issue (Kairouz et al., 2021; Li et al., 2022). A recent study (Babakniya et al., 2023a) has shown that such performance penalty is even more substantial when fine-tuning FMs.

Federated Alignment Tuning. Model alignment is the process of ensuring that FMs behave in line with human intentions and values (Ji et al., 2024). The distributed nature of FL, where data remains on local devices, and the diversity of data exacerbate the difficulty of ensuring fairness, transparency, and accountability in models (Ezzeldin et al., 2023).

3.2.2 System-Level Challenges

System-level challenges stem from the mismatch between the significant resource demands of FM

training and the limited, heterogeneous system resources (e.g., for mobile devices) within FL systems, such as communication bandwidth, computational power, and memory (Su et al., 2023a). This line of challenges include:

Communication Efficiency. In FL, the communication bottleneck is induced by frequently exchanging training information between the server and clients over limited bandwidth channels (Kairouz et al., 2021). The substantial number of parameters in FMs further exacerbates this burden, thus hindering the training process.

System Heterogeneity. The memory and computational resources of the devices for different participants may be diverse (Diao et al., 2021), which could cause delays for model synchronization and inactivation of some participants, i.e., stragglers, making it challenging to leverage the full potential of FMs in FL setting.

3.2.3 Trustworthiness Challenges

This level emphasizes the overarching concerns such as ethical considerations, privacy, security, and fairness in the entire lifecycle of FM-FL, from the pre-training and model adaptation to the application stages. We present two representative types of challenges from this perspective:

Intellectual Property. Intellectual property (IP) protection in FM-FL primarily involves attributing ownership rights for both models and data. FL complicates the identification of contributions from multiple participants, raising questions about who holds the IP for developed models and the FM-generated contents (Li et al., 2023a). From the server’s perspective, broadcasting a pre-trained model to multiple nodes for fine-tuning poses IP protection and security risks (e.g., model theft), necessitating measures to safeguard IP rights and ensure model integrity (Kang et al., 2024).

Privacy Leakage. Although FL does not immediately share data, studies have shown that it may not always guarantee sufficient privacy preservation (Geiping et al., 2020), as model parameters (e.g., weights or gradients) may leak sensitive information to malicious adversaries (Zhu et al., 2019). In terms of FMs, recent studies (Huang et al., 2022; Li et al., 2023c) have shown that generative models still suffer from privacy leakage threats through well-crafted prompts (e.g., jailbreak prompts), even though safety mechanisms are implemented.

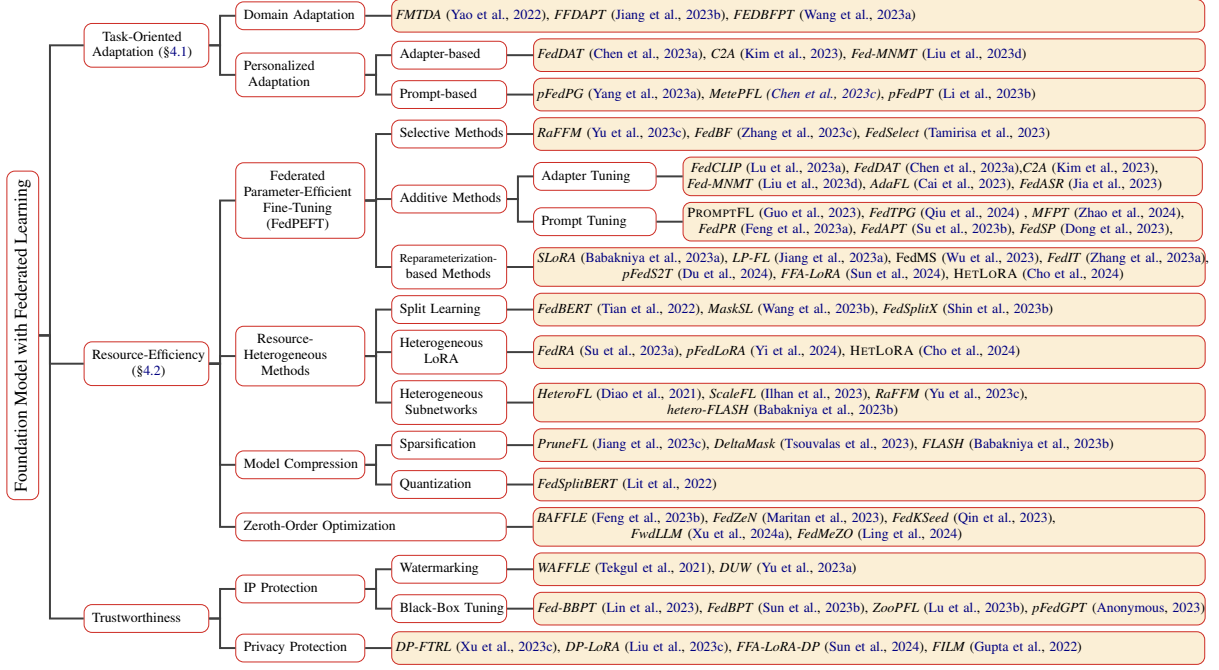


Figure 1: Taxonomy of research in foundation models with federated learning

4 Techniques

In this section, we survey FM-FL techniques, categorizing them as *Task-Oriented Adaptation* (§4.1), *Resource-Efficiency* (§4.2), and *Trustworthiness* (§4.3). As shown in Figure 1, we further refine them according to the key features of different methods.

4.1 Task-Oriented Adaptation

This part introduces major approaches that adapt FMs to handle specific tasks with FL, aiming to tackle task-level challenges.

4.1.1 Domain Adaptation

Despite being heavily reliant on large-scale, public datasets for their initial training, FMs often require further Domain-Adaptive Pre-Training (DAPT) with domain-specific data for tasks that necessitate specialized knowledge (Gururangan et al., 2020; Guo and Yu, 2022). In domains like healthcare, FL allows for the continued pre-training of these models using sensitive, domain-specific data without compromising privacy. Based on this idea, Jiang et al. (2023b) proposed *FFDAPT*, a computational-efficient further pre-training algorithm that freezes a portion of consecutive layers while optimizing the rest of the layers. Similarly, Wang et al. (2023a) proposed *FEDBFPT* that builds a local model for each client, progressively training the shallower layers of local models while sampling deeper lay-

ers, and aggregating trained parameters on a server to create the final global model.

4.1.2 Personalized Adaptation

Personalized adaptation refers to the process of tailoring a pre-trained FM to meet the specific needs or preferences of individual clients while leveraging the decentralized and privacy-preserving nature of FL. Particularly, we discuss two types of popular personalized methods as follows¹:

Adapter-based. Adapter-based methods introduce small, trainable modules (adapters) into the frozen pre-trained FMs, allowing for client-specific model adaptation without altering the original FL. *FedDAT* (Chen et al., 2023a) leverages a dual-adapter structure, with personalized adapters focusing on client-specific knowledge, and a global adapter maintaining client-agnostic knowledge. *FedDAT* executes bi-directional knowledge distillation between personalized adapters and the global adapter to regularize the client’s updates and prevent overfitting.

Prompt-based. Prompt-based methods involve using client-specific soft prompts to guide the model’s response. *pFedPG* (Yang et al., 2023a) trains a prompt generator to exploit underlying client-specific characteristics and produce personal-

¹While this classification intersects with FedPEFT (§4.2.1), which is detailed later, the focus here is on personalization aspects.

ized prompts for each client enabling efficient and personalized adaptation.

4.2 Resource-Efficiency

In response to the system-level challenges, there has been a considerable focus on developing resource-efficient approaches. This part describes techniques that improve resource efficiency.

4.2.1 Federated Parameter-Efficient Fine-Tuning

Federated Parameter-Efficient Fine-Tuning (FedPEFT), originating from the fine-tuning practices of FMs (Lester et al., 2021; Hu et al., 2022; Li and Liang, 2021), is a suite of techniques designed to reduce both the computational load and the associated communication overheads (Malaviya et al., 2023; Woisetschlager et al., 2024). In alignment with existing FM fine-tuning taxonomies (Lialin et al., 2023; Ding et al., 2023), we present FedPEFT methods in three categories: *Selective Methods*, *Additive Methods*, and *Reparameterization-Based Methods*. We depict this taxonomy and representative methods in Figure 2.

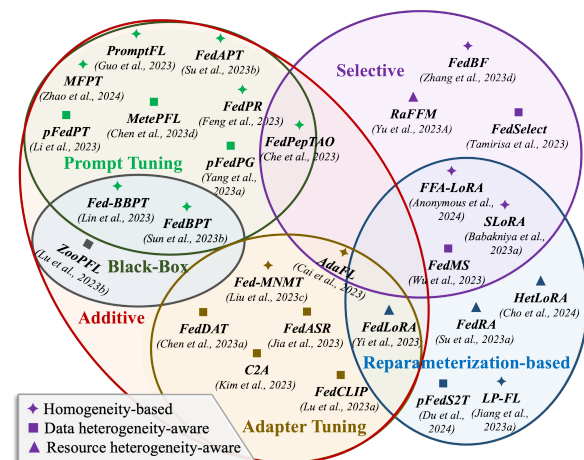


Figure 2: Taxonomy of Federated Parameter-Efficient Fine-Tuning (FedPEFT).

Selective Methods. Selective methods fine-tune a small subset of the parameters, leaving the majority unchanged during fine-tuning. In the field of LLMs, a prominent example of such methods is *BitFit* (Ben Zaken et al., 2022), which only fine-tunes the bias terms. *BitFit* has inspired a series of studies in FedPEFT (Bu et al., 2022; Sun et al., 2022a; Zhang et al., 2023c), demonstrating the superior communication efficiency of only updating the bias terms while still achieving competitive performance. More sophisticated methods strive

to find sparse subnetworks for partial fine-tuning. Among them, various methods (Seo et al., 2021; Li et al., 2021; Tamirisa et al., 2023) advocate for the Lottery Ticket Hypothesis (LTH) (Frankle and Carbin, 2019), positing that a dense network contains many subnetworks whose inference capabilities are as accurate as that of the original network. *FedSelect* (Tamirisa et al., 2023) is a representative method that encourages clients to find optimal subnetworks based on LTH and continually fine-tunes these derived subnetworks to encapsulate local knowledge. As another important aspect, *RaFFM* (Yu et al., 2023c) proposes to prioritize specialized salient parameters by ranking them using salience evaluation metrics such as ℓ_1 and ℓ_2 norms.

Additive Methods. Instead of fine-tuning a subset of model parameters, additive methods incorporate lightweight trainable blocks into frozen FMs and tune the additional parameters for model adaptation. These methods not only enhance computational and communicational efficiency but also introduce an extra benefit: personalization (Lu et al., 2023a), *i.e.*, the integration of these supplementary parameters allows for the customization of heterogeneous models tailored to specific local data characteristics or user preferences. Additive methods include the following representative branches:

Adapter Tuning. Adapter tuning integrates small-scale neural networks (known as “adapters”) into the pre-trained models (Houlsby et al., 2019; Hu et al., 2022). A straightforward implementation of adapter tuning is to collaboratively train a shared adapter among all clients in the *FedAvg* manner, as highlighted by Sun et al. (2022a). Based on *FedAvg*, *FedCLIP* (Lu et al., 2023a) incorporates an attention-based adapter for the image encoder in CLIP models (Radford et al., 2021). In the domain of multilingual machine translation, where different language pairs exhibit substantial discrepancies in data distributions, *Fed-MNMT* (Liu et al., 2023d) explores clustering strategies that group adapter parameters and makes inner-cluster parameters aggregation for alleviating the undesirable effect of data discrepancy. Another representative approach named C2A (Kim et al., 2023) employs hypernetworks (Ha et al., 2017) to generate client-specific adapters by conditioning on the client’s information, maximizing the utility of shared model parameters while minimizing the divergence caused by data heterogeneity.

Prompt Tuning. Prompt tuning incorporates trainable task-specific continuous prompt vectors at the input layer (Liu et al., 2023a; Dong et al., 2023). Compared to full fine-tuning, it achieves comparable performance but with $1000\times$ less parameter storage (Jia et al., 2022). A variation of prompt tuning, *FedPrefix* (Sun et al., 2023a) uses a local adapter to generate the prefixes and aggregate the original self-attention layers.

Reparameterization-based Methods. The hypothesis behind reparameterization-based methods is that fine-tuning adaptations can be reparameterized into optimization within low-rank subspaces (Aghajanyan et al., 2021). Low-Rank Adaptation (LoRA) (Hu et al., 2022), as a popular PEFT method from the area of LLMs, reduces the number of trainable parameters for downstream tasks by representing the weight updates with two smaller matrices (called update matrices) through low-rank decomposition (Ding et al., 2023). For instance, *FedIT* (Zhang et al., 2023a) leverages LoRA to improve the response quality of LLMs by utilizing diverse instructions from different clients. Noticeably, LoRA and its variants have also exhibited considerable potential in addressing the challenges inherent in data heterogeneity among clients in FL. *FedLoRA* (Yi et al., 2024) assigns a homogeneous small low-rank linear adapter for each clients local personalized heterogeneous local model.

4.2.2 Resource-Heterogeneous Methods

FL systems may consist of devices with varying levels of resources, leading to disparities where certain devices exhibit more efficient model training compared to others (Chen et al., 2023a). To address this, several methods have been developed to customize model architectures for heterogeneous clients.

Heterogeneous LoRA. LoRA-based FedPEFT exhibits unique flexibility for resource-limited mobile devices with natural system heterogeneity. Cho et al. (2023) applied heterogeneous LoRA ranks across clients via utilizing zero-padding and truncation for the aggregation and distribution of the LoRA modules. *FedRA* (Su et al., 2023a) integrates LoRA with randomly-allocated subnetworks for local fine-tuning with heterogeneous clients.

Heterogeneous Subnetworks. Some works train heterogeneous subnetworks selected from global models, tailored to the varying capabilities of in-

dividual clients. *HeteroFL* (Diao et al., 2021) appeared as the first method that adaptively allocates subsets of global model parameters for local training. *ScaleFL* (Ilhan et al., 2023) integrates a resource-adaptive 2-D model downscaling mechanism along the width and depth dimensions by leveraging early exits to find the best-fit models for resource-aware local training.

Split Learning. Split learning addresses the resource heterogeneity between servers and clients by splitting a large model into client-side and server-side components (Thapa et al., 2022). For the first time, *FedBERT* (Tian et al., 2022) leverages split training for training the BERT model, showing the feasibility of pre-training large FMs in FL settings. *FedSplitX* (Shin et al., 2023b) is a more fine-grained method that allows multiple partition points for model splitting, accommodating diverse client capabilities.

4.2.3 Model Compression

Model compression refers to the techniques used to reduce the size of models, thereby improving communication and computational efficiency (Shah and Lau, 2023).

Sparsification. Model sparsification methods reduce communication burden by only transmitting a subset of FM parameters across the network (Jiang et al., 2023c). Typical methods focus on identifying and cultivating high-potential subnetworks (Frankle and Carbin, 2019; Tsouvalas et al., 2023).

Quantization. Quantization is well-established in both the FM and FL domains (Xu et al., 2024b; Reisizadeh et al., 2020), which involves decreasing the precision of floating-point parameters for mitigating the storage, computational, and communication demands. Quantization is both effective and easy to implement, making it ideal for use with other resource-efficient methods (Lit et al., 2022).

4.2.4 Zeroth-Order Optimization

Distinct from the ubiquitous reliance on gradient descent in most FL optimization algorithms, a specific line of research advocates for the removal of backpropagation (BP) (Malladi et al., 2023) in favor of Zeroth-Order Optimization (ZOO) (Fang et al., 2022; Li and Chen, 2021). BP-free approaches conserve memory needed for computing gradients and minimize communication overhead for model aggregation (Qin et al., 2023), making FMs more accessible for lower-end devices,

thereby enhancing their applicability in diverse hardware environments of FL. Recent work based on perturbed inferences, such as that by [Xu et al. \(2023a\)](#); [Qin et al. \(2023\)](#), has initiated preliminary explorations into the deployment of both FedPEFT and full-model fine-tuning of billion-sized FMs, like LLaMA, on mobile devices.

4.3 Trustworthiness

This line of work aims to enhance trustworthiness throughout the FM-FL lifecycle, covering a variety of key areas including, but not limited to, *IP Protection*, *Attack Robustness*, and *Privacy Protection*.

IP Protection. Existing IP protection involves safeguarding ownership of FL models from unauthorized use (e.g., model theft) ([Tekgul et al., 2021](#)). Two common kinds of IP protection strategies are watermarking and black-box tuning.

Watermarking is a well-known deterrence method for model IP protection by providing the identities for model owners to demonstrate ownership of their models ([Adi et al., 2018](#)). [Tekgul et al. \(2021\)](#) proposed *WAFFLE*, the first solution that addresses the ownership problem by injecting a watermark into the global model in FL environments. Recently, [Yu et al. \(2023b\)](#) proposed *DUW* that embeds a client-unique key into each clients local model, aiming to identify the infringer of a leaked model while verifying the FL models ownership.

Black-Box Tuning is a set of gradient-free methods to drive large language models. ZOO allows for black-box fine-tuning in scenarios where direct access to model parameters is restricted, e.g., due to privacy concerns or proprietary limitations ([Sun et al., 2022b](#)). *Fed-BBPT* ([Lin et al., 2023](#)) is a general prompt tuning framework that facilitates the joint training of a global lightweight prompt generator across multiple clients. *FedBPT* ([Sun et al., 2023b](#)) adopts a classic evolutionary-based ZOO method, CMA-ES ([Hansen and Ostermeier, 2001](#)), for training an optimal prompt that improves the performance of the frozen FMs. *ZooPFL* ([Lu et al., 2023b](#)), on the other hand, applies coordinate-wise gradient estimate to learn input surgery that incorporates client-specific embeddings. Nevertheless, the pronounced slower convergence rates of ZOO compared to gradient-based approaches in high-dimensional settings ([Golovin et al., 2020](#)), underscore a significant research gap. The implications of these slower rates on convergence efficiency and computational burden in FL, especially

for large-scale FMs, remain insufficiently explored.

Differential Privacy. Differential Privacy (DP) is a theoretical framework that governs privacy boundaries and manages the tradeoff between privacy and model convergence ([Wei et al., 2020](#)). DP-based FL approaches often add artificial noise (e.g., Gaussian noise) to parameters at the clients side before aggregating to prevent information leakage ([Xu et al., 2023c](#)). Besides, DP is compatible with most FedPEFT methods. For instance, [Sun et al. \(2024\)](#) showed that DP noise can even be amplified by the locally "semi-quadratic" nature of LoRA-based methods, motivating the integration of LoRA with DP to improve resource efficiency while maintaining data privacy ([Liu et al., 2023c](#)). In terms of attack, [Gupta et al. \(2022\)](#) presented an attack that recovers private text data by extracting information from gradients transmitted during training, despite the employment of a naive DP mechanism.

5 Future Directions & Applications

We highlight potential research directions and future FL-FM applications in this section.

5.1 Future Directions

Personalization. FL on FMs can improve user profiling by capturing more granular and diverse data from individuals while preserving privacy. This can lead to more accurate and comprehensive user profiles, enabling personalized recommendations and services tailored to specific preferences, needs, and contexts ([Chen et al., 2023b](#)). Users can contribute their preferences, feedback, and insights, allowing the models to learn directly from their interactions and refine personalization algorithms accordingly. Future research directions may incorporate multi-modal data, including text, images, audio, and sensor data.

Model Compression. Future directions may involve designing more efficient and lightweight model compression techniques ([Deng et al., 2020](#)) specifically tailored for FL systems to reduce the computational and memory requirements of FMs while maintaining their performance. They may leverage multi-task learning approaches for model sharing and parameter reuse across different tasks or domains. Adaptive model compression techniques could dynamically adjust the compression level based on the available computing resources or application requirements ([Xu et al., 2023b](#)).

Split Learning. Split learning (Thapa et al., 2020) partitions the model, placing one part on the client device and the other on the server. Future developments may explore more sophisticated and adaptive methods for partitioning the FMs, such as adaptive model partitioning based on the computational capabilities and resources of the client devices. Dynamic model partitioning techniques may adjust the partitioning scheme at different stages of the FL process.

Mixture of Experts (MoE). MoE allows FL to incorporate multiple expert FMs, each specializing in different aspects or domains of the data. By combining the expertise of these models, FL can achieve higher model performance and accuracy. MoE also allows FL and FM models to adapt to local data characteristics present on individual client devices. Bridging MoE with FMs could strengthen generalization ability by balancing between larger overall model capacity and flexible per-instance specialization (Cong et al., 2023).

Privacy Preservation. Advanced model aggregation methods can be designed to incorporate privacy awareness when performing FL on FMs. This includes techniques to control the amount of information leaked during the aggregation process and mechanisms to enforce privacy guarantees while maintaining the accuracy and utility of the aggregated model (Nagy et al., 2023). As privacy concerns continue to grow, future developments may involve the establishment of privacy regulations and standards specifically tailored for FL and FMs.

Continual Learning. Continual learning enables models to adapt to new data over time, improving their performance and accuracy. By incorporating new data into the model training process, FL and FMs can continuously improve and adapt to changing environments and user needs (Yang et al., 2023b). Future directions may involve leveraging transfer learning techniques in continual learning for FL and FMs. Models can transfer knowledge from previous tasks or domains to new ones, enabling more efficient learning and adaptation (Good et al., 2023).

Resource-Efficiency. FL-FM enables collaborative training, model adaptation, and utilization of FMs for a wide range of novel and powerful applications on heterogeneous edge devices (Shen et al., 2024; Xu et al., 2024b). The design and adaptation

of the FM models, optimization of computation and communication, and coordination among heterogeneous edge devices and the cloud remain to be further explored in this new era.

5.2 Domain-Specific Applications

In this section, we discuss how FM-FL can be utilized in several representative domains.

Healthcare. FL and FMs enable the development of personalized medical applications. By training with massive individual patient data, such as medical history, genetics, and lifestyle factors, the models can learn global patterns and provide tailored recommendations for treatment, medication, and prevention strategies.

Law, Finance and Banking. FL can support the training of FMs on massive legal documents, cases, and statutes (Zhang et al., 2023b). The models can assist in identifying key legal arguments, summarizing case details, providing insights, and making predictive analytics into potential case outcomes (Yue et al., 2023). FL can build FMs that support risk management applications in banking and finance. By analyzing aggregated data from multiple sources, such as credit scores, market data, and economic trends, the models can support risk assessment, credit scoring, and investment management (Shin et al., 2023a). FM-FL models trained on historical investment data and market trends can support investment opportunities, analyze investment risks, and assist in portfolio optimization.

Education and Personal Agents. FL can be used to train FMs for intelligent tutoring systems to support individual student learning. Personalized foundation models can provide customized and personalized responses to users based on their individual interests, preferences, behavior, and history.

6 Conclusions

In this survey, we have meticulously surveyed the intersection of FM and FL. We identified three levels of challenges: task-level, system-level as well as trustworthiness challenges, and proposed a comprehensive taxonomy of techniques in response to these challenges. In addition, we discussed future directions and applications in this research field, hoping to attract more breakthroughs in future research.

Limitations

FM and FL are very fast-moving fields. We have put a lot of effort to include the latest research efforts in the community in this survey. The majority of the papers referenced in our taxonomy are indeed from 2023 which also demonstrates the importance of the integration of FM and FL. Therefore, we believe that our survey will help to inspire and push further research and innovation in this important areas. Our survey does not include any benchmarking of the available ideas and systems. We believe that would be an important next step that we are leaving to future work. It would, however, require some tools to support such an evaluation campaign and such tools are, to the best of our knowledge, not available yet.

References

Yossi Adi, Carsten Baum, Moustapha Cisse, Benny Pinkas, and Joseph Keshet. 2018. [Turning your weakness into a strength: Watermarking deep neural networks by backdooring](#). In *27th USENIX Security Symposium (USENIX Security 18)*, pages 1615–1631, Baltimore, MD. USENIX Association.

Armen Aghajanyan, Sonal Gupta, and Luke Zettlemoyer. 2021. [Intrinsic dimensionality explains the effectiveness of language model fine-tuning](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 7319–7328, Online. Association for Computational Linguistics.

Anonymous. 2023. [Personalized federated learning for text classification with gradient-free prompt tuning](#).

Sara Babakniya, Ahmed Elkordy, Yahya Ezzeldin, Qingfeng Liu, Kee-Bong Song, MOSTAFA EL-Khamy, and Salman Avestimehr. 2023a. [SLoRA: Federated parameter efficient fine-tuning of language models](#). In *International Workshop on Federated Learning in the Age of Foundation Models in Conjunction with NeurIPS 2023*.

Sara Babakniya, Souvik Kundu, Saurav Prakash, Yue Niu, and Salman Avestimehr. 2023b. [Revisiting sparsity hunting in federated learning: Why does sparsity consensus matter?](#) *Transactions on Machine Learning Research*.

Elad Ben Zaken, Yoav Goldberg, and Shauli Ravfogel. 2022. [BitFit: Simple parameter-efficient fine-tuning for transformer-based masked language-models](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 1–9, Dublin, Ireland. Association for Computational Linguistics.

Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. 2021. [On the opportunities and risks of foundation models](#). *arXiv preprint arXiv:2108.07258*.

Kallista Bonawitz, Peter Kairouz, Brendan McMahan, and Daniel Ramage. 2021. [Federated learning and privacy: Building privacy-preserving systems for machine learning and data science on decentralized data](#). *Queue*, 19(5):87114.

Tom Brown et al. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901. Curran Associates, Inc.

Zhiqi Bu, Yu-Xiang Wang, Sheng Zha, and George Karypis. 2022. [Differentially private bias-term only fine-tuning of foundation models](#). In *Workshop on Trustworthy and Socially Responsible Machine Learning, NeurIPS 2022*.

Sébastien Bubeck, Varun Chandrasekaran, Ronen Eldan, Johannes Gehrke, Eric Horvitz, Ece Kamar, Peter Lee, Yin Tat Lee, Yuanzhi Li, Scott Lundberg, Harsha Nori, Hamid Palangi, Marco Tulio Ribeiro, and Yi Zhang. 2023. [Sparks of artificial general intelligence: Early experiments with gpt-4](#).

Dongqi Cai, Yaozong Wu, Shangguang Wang, Felix Xiaozhu Lin, and Mengwei Xu. 2023. [Efficient federated learning for modern nlp](#). In *Proceedings of the 29th Annual International Conference on Mobile Computing and Networking, ACM MobiCom '23*, New York, NY, USA. Association for Computing Machinery.

CCPA. 2023. [California consumer privacy act \(ccpa\)](#).

Pushpita Chatterjee, Debashis Das, and Danda B Rawat. 2023. Use of federated learning and blockchain towards securing financial services. *arXiv preprint arXiv:2303.12944*.

Haokun Chen, Yao Zhang, Denis Krompass, Jindong Gu, and Volker Tresp. 2023a. [Feddat: An approach for foundation model finetuning in multi-modal heterogeneous federated learning](#). *arXiv preprint arXiv:2308.12305*.

Jin Chen, Zheng Liu, Xu Huang, Chenwang Wu, Qi Liu, Gangwei Jiang, Yuanhao Pu, Yuxuan Lei, Xiaolong Chen, Xingmei Wang, et al. 2023b. When large language models meet personalization: Perspectives of challenges and opportunities. *arXiv preprint arXiv:2307.16376*.

Shengchao Chen, Guodong Long, Tao Shen, and Jing Jiang. 2023c. [Prompt federated learning for weather forecasting: Toward foundation models on meteorological data](#). In *Proceedings of the Thirty-Second International Joint Conference on Artificial Intelligence, IJCAI-23*, pages 3532–3540. International Joint Conferences on Artificial Intelligence Organization. Main Track.

813	Yae Jee Cho, Luyang Liu, Zheng Xu, Aldi Fahrezi,	2024. Communication-efficient personalized federated learning for speech-to-text tasks.	869
814	Matt Barnes, and Gauri Joshi. 2023. Heterogeneous		870
815	loRA for federated fine-tuning of on-device founda-		
816	tion models. In <i>International Workshop on Federated</i>		
817	<i>Learning in the Age of Foundation Models in Con-</i>		
818	<i>junction with NeurIPS 2023.</i>		
819	Yae Jee Cho, Luyang Liu, Zheng Xu, Aldi Fahrezi, and	Yahya H. Ezzeldin, Shen Yan, Chaoyang He, Emilio	871
820	Gauri Joshi. 2024. Heterogeneous low-rank approxi-	Ferrara, and A. Salman Avestimehr. 2023. Fairfed:	872
821	mation for federated fine-tuning of on-device founda-	Enabling group fairness in federated learning. <i>Pro-</i>	873
822	tion models. <i>arXiv preprint arXiv:2401.06432.</i>	<i>ceedings of the AAAI Conference on Artificial Intelli-</i>	874
		<i>gence</i> , 37(6):7494–7502.	875
823	Wenyan Cong, Hanxue Liang, Peihao Wang, Zhiwen	Wenzhi Fang, Ziyi Yu, Yuning Jiang, Yuanming	876
824	Fan, Tianlong Chen, Mukund Varma, Yi Wang, and	Shi, Colin N. Jones, and Yong Zhou. 2022.	877
825	Zhangyang Wang. 2023. Enhancing nerf akin to en-	Communication-efficient stochastic zeroth-order op-	878
826	hancing llms: Generalizable nerf transformer with	timization for federated learning. <i>IEEE Transactions</i>	879
827	mixture-of-view-experts. In <i>2023 IEEE/CVF Interna-</i>	<i>on Signal Processing</i> , 70:5058–5073.	880
828	<i>tional Conference on Computer Vision (ICCV)</i> , pages		
829	3170–3181.	Chun-Mei Feng, Bangjun Li, Xinxing Xu, Yong Liu,	881
		Huazhu Fu, and Wangmeng Zuo. 2023a. Learning	882
830	Lei Deng, Guoqi Li, Song Han, Luping Shi, and Yuan	federated visual prompt in null space for mri recon-	883
831	Xie. 2020. Model compression and hardware accel-	struction. In <i>2023 IEEE/CVF Conference on Com-</i>	884
832	eration for neural networks: A comprehensive survey.	<i>puter Vision and Pattern Recognition (CVPR)</i> , pages	885
833	<i>Proceedings of the IEEE</i> , 108(4):485–532.	8064–8073.	886
		Haozhe Feng, Tianyu Pang, Chao Du, Wei Chen,	887
834	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	Shuicheng YAN, and Min Lin. 2023b. Does fed-	888
835	Kristina Toutanova. 2019. BERT: Pre-training of	erated learning really need backpropagation?	889
836	deep bidirectional transformers for language under-		
837	standing. In <i>Proceedings of the 2019 Conference of</i>	Jonathan Frankle and Michael Carbin. 2019. The lottery	890
838	<i>the North American Chapter of the Association for</i>	ticket hypothesis: Finding sparse, trainable neural	891
839	<i>Computational Linguistics: Human Language Tech-</i>	networks. In <i>International Conference on Learning</i>	892
840	<i>nologies, Volume 1 (Long and Short Papers)</i> , pages	<i>Representations.</i>	893
841	4171–4186, Minneapolis, Minnesota. Association for		
842	Computational Linguistics.	GDPR. 2016. Regulation (eu) 2016/679 of the european	894
		parliament and of the council.	895
843	Enmao Diao, Jie Ding, and Vahid Tarokh. 2021. Het-	Jonas Geiping, Hartmut Bauermeister, Hannah Dröge,	896
844	ero{fl}: Computation and communication efficient	and Michael Moeller. 2020. Inverting gradients - how	897
845	federated learning for heterogeneous clients. In <i>Inter-</i>	easy is it to break privacy in federated learning? In	898
846	<i>national Conference on Learning Representations.</i>	<i>Advances in Neural Information Processing Systems</i> ,	899
		volume 33, pages 16937–16947. Curran Associates,	900
847	Ning Ding, Yujia Qin, Guang Yang, Fuchao Wei,	Inc.	901
848	Zonghan Yang, Yusheng Su, Shengding Hu, Yulin	Daniel Golovin, John Karro, Greg Kochanski, Chansoo	902
849	Chen, Chi-Min Chan, Weize Chen, et al. 2023.	Lee, Xingyou Song, and Qiuyi Zhang. 2020. Gradi-	903
850	Parameter-efficient fine-tuning of large-scale pre-	entless descent: High-dimensional zeroth-order opti-	904
851	trained language models. <i>Nature Machine Intelli-</i>	mization. In <i>International Conference on Learning</i>	905
852	<i>gence</i> , 5(3):220–235.	<i>Representations.</i>	906
853	Chenhe Dong, Yuexiang Xie, Bolin Ding, Ying Shen,	Jack Good, Jimit Majmudar, Christophe Dupuy, Jix-	907
854	and Yaliang Li. 2023. Tunable soft prompts are mes-	uan Wang, Charith Peris, Clement Chung, Richard	908
855	sengers in federated learning. In <i>Findings of the</i>	Zemel, and Rahul Gupta. 2023. Coordinated replay	909
856	<i>Association for Computational Linguistics: EMNLP</i>	sample selection for continual federated learning. In	910
857	2023, pages 14665–14675, Singapore. Association	<i>Proceedings of the 2023 Conference on Empirical</i>	911
858	for Computational Linguistics.	<i>Methods in Natural Language Processing: Industry</i>	912
		<i>Track</i> , pages 331–342, Singapore. Association for	913
859	Alexey Dosovitskiy, Lucas Beyer, Alexander	Computational Linguistics.	914
860	Kolesnikov, Dirk Weissenborn, Xiaohua Zhai,	Gemini Team Google. 2023. Gemini: a family of	915
861	Thomas Unterthiner, Mostafa Dehghani, Matthias	highly capable multimodal models. <i>arXiv preprint</i>	916
862	Minderer, Georg Heigold, Sylvain Gelly, Jakob	<i>arXiv:2312.11805.</i>	917
863	Uszkoreit, and Neil Houlsby. 2021. An image		
864	is worth 16x16 words: Transformers for image	Suriya Gunasekar, Yi Zhang, Jyoti Aneja, Caio	918
865	recognition at scale. In <i>International Conference on</i>	César Teodoro Mendes, Allie Del Giorno, Sivakanth	919
866	<i>Learning Representations.</i>	Gopi, Mojan Javaheripi, Piero Kauffmann, Gustavo	920
867	Yichao Du, Zhirui Zhang, Linan Yue, Xu Huang,	de Rosa, Olli Saarikivi, et al. 2023. Textbooks are all	921
868	Yuqing Zhang, Tong Xu, Linli Xu, and Enhong Chen.	you need. <i>arXiv preprint arXiv:2306.11644.</i>	922

923	Tao Guo, Song Guo, Junxiao Wang, Xueyang Tang, and Wenchao Xu. 2023. Promptfl: Let federated participants cooperatively learn prompts instead of models - federated learning in age of foundation model . <i>IEEE Transactions on Mobile Computing</i> , pages 1–15.	979
924		980
925		981
926		982
927		983
928	Xu Guo and Han Yu. 2022. On the domain adaptation and generalization of pretrained language models: A survey . <i>arXiv preprint arXiv:2211.03154</i> .	984
929		985
930		986
931	Samyak Gupta, Yangsibo Huang, Zexuan Zhong, Tianyu Gao, Kai Li, and Danqi Chen. 2022. Recovering private text in federated learning of language models . In <i>Advances in Neural Information Processing Systems</i> , volume 35, pages 8130–8143. Curran Associates, Inc.	987
932		988
933		989
934		990
935		991
936		992
937	Suchin Gururangan, Ana Marasović, Swabha Swayamdipta, Kyle Lo, Iz Beltagy, Doug Downey, and Noah A. Smith. 2020. Don’t stop pretraining: Adapt language models to domains and tasks . In <i>Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics</i> , pages 8342–8360, Online. Association for Computational Linguistics.	993
938		994
939		995
940		996
941		997
942		
943		
944		
945	David Ha, Andrew M. Dai, and Quoc V. Le. 2017. Hypernetworks . In <i>International Conference on Learning Representations</i> .	998
946		999
947		1000
948	Nikolaus Hansen and Andreas Ostermeier. 2001. Completely derandomized self-adaptation in evolution strategies . <i>Evolutionary Computation</i> , 9(2):159–195.	1001
949		1002
950		1003
951	Neil Houlsby, Andrei Giurgiu, Stanislaw Jastrzebski, Bruna Morrone, Quentin De Laroussilhe, Andrea Gesmundo, Mona Attariyan, and Sylvain Gelly. 2019. Parameter-efficient transfer learning for NLP . In <i>Proceedings of the 36th International Conference on Machine Learning</i> , volume 97 of <i>Proceedings of Machine Learning Research</i> , pages 2790–2799. PMLR.	1004
952		1005
953		
954		
955		
956		
957		
958		
959	Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen. 2022. LoRA: Low-rank adaptation of large language models . In <i>International Conference on Learning Representations</i> .	1006
960		1007
961		1008
962		1009
963		1010
964		1011
965	Jie Huang, Hanyin Shao, and Kevin Chen-Chuan Chang. 2022. Are large pre-trained language models leaking your personal information? In <i>Findings of the Association for Computational Linguistics: EMNLP 2022</i> , pages 2038–2047, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.	1012
966		1013
967		1014
968		1015
969		1016
970	Xumin Huang, Peichun Li, Hongyang Du, Jiawen Kang, Dusit Niyato, Dong In Kim, and Yuan Wu. 2024. Federated learning-empowered ai-generated content in wireless networks . <i>IEEE Network</i> , pages 1–1.	1017
971		
972		
973		
974	Fatih Ilhan, Gong Su, and Ling Liu. 2023. Scalefl: Resource-adaptive federated learning with heterogeneous clients . In <i>2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)</i> , pages 24532–24541.	1018
975		1019
976		1020
977		1021
978		1022
		1023
		1024
		1025
		1026
		1027
		1028
		1029
		1030
		1031
		1032
		1033
		1034

1035	Yeachen Kim, Junho Kim, Wing-Lam Mok, Jun-Hyung Park, and SangKeun Lee. 2023. Client-customized adaptation for parameter-efficient federated learning . In <i>Findings of the Association for Computational Linguistics: ACL 2023</i> , pages 1159–1172, Toronto, Canada. Association for Computational Linguistics.	1090
1036		1091
1037		1092
1038		1093
1039		1094
1040		
1041	Alexander Kirillov, Eric Mintun, Nikhila Ravi, Hanzi Mao, Chloe Rolland, Laura Gustafson, Tete Xiao, Spencer Whitehead, Alexander C. Berg, Wan-Yen Lo, Piotr Dollár, and Ross Girshick. 2023. Segment anything .	1095
1042		1096
1043		1097
1044		1098
1045		
1046	Brian Lester, Rami Al-Rfou, and Noah Constant. 2021. The power of scale for parameter-efficient prompt tuning . In <i>Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing</i> , pages 3045–3059, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.	1109
1047		1100
1048		1101
1049		1102
1050		
1051		1103
1052		1104
1053	Ang Li, Jingwei Sun, Binghui Wang, Lin Duan, Sicheng Li, Yiran Chen, and Hai Li. 2021. Lotteryfl: Empower edge intelligence with personalized and communication-efficient federated learning . In <i>2021 IEEE/ACM Symposium on Edge Computing (SEC)</i> , pages 68–79.	1105
1054		1106
1055		
1056		1107
1057		1108
1058		1109
1059	Bowen Li, Lixin Fan, Hanlin Gu, Jie Li, and Qiang Yang. 2023a. Fedipr: Ownership verification for federated deep neural network models . <i>IEEE Transactions on Pattern Analysis and Machine Intelligence</i> , 45(4):4521–4536.	1110
1060		1111
1061		1112
1062		1113
1063		
1064	Guanghao Li, Wansen Wu, Yan Sun, Li Shen, Baoyuan Wu, and Dacheng Tao. 2023b. Visual prompt based personalized federated learning . <i>Transactions on Machine Learning Research</i> .	1114
1065		1115
1066		1116
1067		1117
1068		1118
1069	Haoran Li, Dadi Guo, Wei Fan, Mingshi Xu, Jie Huang, Fanpu Meng, and Yangqiu Song. 2023c. Multi-step jailbreaking privacy attacks on ChatGPT . In <i>Findings of the Association for Computational Linguistics: EMNLP 2023</i> , pages 4138–4153, Singapore. Association for Computational Linguistics.	1119
1070		1120
1071		1121
1072		1122
1073		
1074	Qinbin Li, Yiqun Diao, Quan Chen, and Bingsheng He. 2022. Federated learning on non-iid data silos: An experimental study . In <i>2022 IEEE 38th International Conference on Data Engineering (ICDE)</i> , pages 965–978.	1123
1075		1124
1076		1125
1077		1126
1078		
1079	Xi Li and Jiaqi Wang. 2024. Position paper: Assessing robustness, privacy, and fairness in federated learning integrated with foundation models .	1127
1080		1128
1081		1129
1082	Xiang Lisa Li and Percy Liang. 2021. Prefix-tuning: Optimizing continuous prompts for generation . In <i>Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)</i> , pages 4582–4597. Association for Computational Linguistics.	1130
1083		1131
1084		1132
1085		1133
1086		
1087		1134
1088		1135
1089		1136
		1137
	Zan Li and Li Chen. 2021. Communication-efficient decentralized zeroth-order method on heterogeneous data . In <i>2021 13th International Conference on Wireless Communications and Signal Processing (WCSP)</i> , pages 1–6.	1138
		1139
		1140
		1141
		1142
	Vladislav Lialin, Vijeta Deshpande, and Anna Rumshisky. 2023. Scaling down to scale up: A guide to parameter-efficient fine-tuning . <i>arXiv preprint arXiv:2303.15647</i> .	1143
		1144
		1145
		1146
		1147
		1148
		1149
		1150
		1151
		1152
		1153
		1154
		1155
		1156
		1157
		1158
		1159
		1160
		1161
		1162
		1163
		1164
		1165
		1166
		1167
		1168
		1169
		1170
		1171
		1172
		1173
		1174
		1175
		1176
		1177
		1178
		1179
		1180
		1181
		1182
		1183
		1184
		1185
		1186
		1187
		1188
		1189
		1190
		1191
		1192
		1193
		1194
		1195
		1196
		1197
		1198
		1199
		1200
		1201
		1202
		1203
		1204
		1205
		1206
		1207
		1208
		1209
		1210
		1211
		1212
		1213
		1214
		1215
		1216
		1217
		1218
		1219
		1220
		1221
		1222
		1223
		1224
		1225
		1226
		1227
		1228
		1229
		1230
		1231
		1232
		1233
		1234
		1235
		1236
		1237
		1238
		1239
		1240
		1241
		1242
		1243
		1244
		1245
		1246
		1247
		1248
		1249
		1250
		1251
		1252
		1253
		1254
		1255
		1256
		1257
		1258
		1259
		1260
		1261
		1262
		1263
		1264
		1265
		1266
		1267
		1268
		1269
		1270
		1271
		1272
		1273
		1274
		1275
		1276
		1277
		1278
		1279
		1280
		1281
		1282
		1283
		1284
		1285
		1286
		1287
		1288
		1289
		1290
		1291
		1292
		1293
		1294
		1295
		1296
		1297
		1298
		1299
		1300

1143	Shubham Malaviya, Manish Shukla, and Sachin Lodha. 2023. Reducing communication overhead in federated learning for pre-trained language models using parameter-efficient finetuning . In <i>Proceedings of The 2nd Conference on Lifelong Learning Agents</i> , volume 232 of <i>Proceedings of Machine Learning Research</i> , pages 456–469. PMLR.	Ilya Sutskever. 2021. Zero-shot text-to-image generation . In <i>International Conference on Machine Learning</i> , pages 8821–8831. PMLR.	1197
1144			1198
1145			1199
1146			
1147		Swarna Priya Ramu, Parimala Boopalan, Quoc-Viet Pham, Praveen Kumar Reddy Maddikunta, Thien Huynh-The, Mamoun Alazab, Thanh Thi Nguyen, and Thippa Reddy Gadekallu. 2022. Federated learning enabled digital twins for smart cities: Concepts, recent advances, and future directions. <i>Sustainable Cities and Society</i> , 79:103663.	1200
1148			1201
1149			1202
			1203
1150	Sadhika Malladi, Tianyu Gao, Eshaan Nichani, Alex Damian, Jason D. Lee, Danqi Chen, and Sanjeev Arora. 2023. Fine-tuning language models with just forward passes . In <i>Workshop on Efficient Systems for Foundation Models @ ICML2023</i> .	Scott Reed et al. 2022. A generalist agent . <i>Transactions on Machine Learning Research</i> . Featured Certification, Outstanding Certification.	1204
1151			1205
1152			1206
1153			
1154			1207
1155	Alessio Maritan, Subhrakanti Dey, and Luca Schenato. 2023. Fedzen: Towards superlinear zeroth-order federated learning via incremental hessian estimation .	Amirhossein Reisizadeh, Aryan Mokhtari, Hamed Hassani, Ali Jadbabaie, and Ramtin Pedarsani. 2020. Fedpaq: A communication-efficient federated learning method with periodic averaging and quantization . In <i>Proceedings of the Twenty Third International Conference on Artificial Intelligence and Statistics</i> , volume 108 of <i>Proceedings of Machine Learning Research</i> , pages 2021–2031. PMLR.	1208
1156			1209
1157			
			1210
1158	Brendan McMahan, Eider Moore, Daniel Ramage, Seth Hampson, and Blaise Aguera y Arcas. 2017. Communication-efficient learning of deep networks from decentralized data . In <i>Artificial intelligence and statistics</i> , pages 1273–1282. PMLR.		1211
1159			1212
1160			1213
1161			1214
1162			1215
			1216
1163	Balázs Nagy, István Hegedűs, Noémi Sándor, Balázs Egedi, Haaris Mehmood, Karthikeyan Saravanan, Gábor Lóki, and Ákos Kiss. 2023. Privacy-preserving federated learning and its application to natural language processing. <i>Knowledge-Based Systems</i> , 268:110475.	Nicola Rieke, Jonny Hancox, Wenqi Li, Fausto Milletari, Holger R Roth, Shadi Albarqouni, Spyridon Bakas, Mathieu N Galtier, Bennett A Landman, Klaus Maier-Hein, et al. 2020. The future of digital health with federated learning. <i>NPJ digital medicine</i> , 3(1):119.	1217
1164			1218
1165			1219
1166			1220
1167			1221
1168			1222
			1223
1169	OpenAI. 2022. Chatgpt .	Sejin Seo, Seung-Woo Ko, Jihong Park, Seong-Lyun Kim, and Mehdi Bennis. 2021. Communication-efficient and personalized federated lottery ticket learning . In <i>2021 IEEE 22nd International Workshop on Signal Processing Advances in Wireless Communications (SPAWC)</i> , pages 581–585.	1224
			1225
1170	OpenAI. 2023. Gpt-4 technical report . <i>arXiv</i> .		1226
			1227
1171	Sharnil Pandya, Gautam Srivastava, Rutvij Jhaveri, M. Rajasekhara Babu, Sweta Bhattacharya, Praveen Kumar Reddy Maddikunta, Spyridon Mastorakis, Md. Jalil Piran, and Thippa Reddy Gadekallu. 2023. Federated learning for smart cities: A comprehensive survey . <i>Sustainable Energy Technologies and Assessments</i> , 55:102987.	Suhail Mohmad Shah and Vincent K. N. Lau. 2023. Model compression for communication efficient federated learning . <i>IEEE Transactions on Neural Networks and Learning Systems</i> , 34(9):5937–5951.	1228
1172			1229
1173			
1174			1230
1175			1231
1176			1232
1177			1233
1178	Zhen Qin, Daoyuan Chen, Bingchen Qian, Bolin Ding, Yaliang Li, and Shuiguang Deng. 2023. Federated full-parameter tuning of billion-sized language models with communication cost under 18 kilobytes . <i>arXiv preprint arXiv:2312.06353</i> .	Yifei Shen, Jiawei Shao, Xinjie Zhang, Zehong Lin, Hao Pan, Dongsheng Li, Jun Zhang, and Khaled B Letaief. 2024. Large language models empowered autonomous edge AI for connected intelligence. <i>IEEE Communications Magazine</i> .	1234
1179			1235
1180			1236
1181			1237
1182			1238
1183	Chen Qiu, Xingyu Li, Chaithanya Kumar Mummadi, Madan Ravi Ganesh, Zhenzhen Li, Lu Peng, and Wan-Yi Lin. 2024. Text-driven prompt generation for vision-language models in federated learning . In <i>The Twelfth International Conference on Learning Representations</i> .	Jaemin Shin, Hyungjun Yoon, Seungjoo Lee, Sungjoon Park, Yunxin Liu, Jinho Choi, and Sung-Ju Lee. 2023a. FedTherapist: Mental health monitoring with user-generated linguistic expressions on smartphones via federated learning . In <i>Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing</i> , pages 11971–11988, Singapore. Association for Computational Linguistics.	1239
1184			1240
1185			1241
1186			1242
1187			1243
1188			1244
			1245
1189	Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. 2021. Learning transferable visual models from natural language supervision . In <i>International conference on machine learning</i> , pages 8748–8763. PMLR.	Jiyeun Shin, Jinhyun Ahn, Honggu Kang, and Joonhyuk Kang. 2023b. Fedsplitx: Federated split learning for computationally-constrained heterogeneous clients .	1246
1190			1247
1191			1248
1192			1249
1193			
1194			1250
			1251
1195	Aditya Ramesh, Mikhail Pavlov, Gabriel Goh, Scott Gray, Chelsea Voss, Alec Radford, Mark Chen, and	Shangchao Su, Bin Li, and Xiangyang Xue. 2023a. Fedra: A random allocation strategy for federated tuning to unleash the power of heterogeneous clients . <i>arXiv preprint arXiv:2311.11227</i> .	1252
1196			1253

1254	Shangchao Su, Mingzhao Yang, Bin Li, and Xiangyang	Azhar, et al. 2023a. Llama: Open and effi-	1309
1255	Xue. 2023b. Federated adaptive prompt tuning for	cient foundation language models . <i>arXiv preprint</i>	1310
1256	multi-domain collaborative learning .	<i>arXiv:2302.13971</i> .	1311
1257	Guangyu Sun, Matias Mendieta, Jun Luo, Shandong	Hugo Touvron, Louis Martin, Kevin Stone, Peter Al-	1312
1258	Wu, and Chen Chen. 2023a. Fedperfix: Towards par-	bert, Amjad Almahairi, Yasmine Babaei, Nikolay	1313
1259	tial model personalization of vision transformers in	Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti	1314
1260	federated learning . In <i>Proceedings of the IEEE/CVF</i>	Bhosale, et al. 2023b. Llama 2: Open founda-	1315
1261	<i>International Conference on Computer Vision</i> , pages	tion and fine-tuned chat models . <i>arXiv preprint</i>	1316
1262	4988–4998.	<i>arXiv:2307.09288</i> .	1317
1263	Guangyu Sun, Matias Mendieta, Taojiannan Yang, and	Vasileios Tsouvalas, Yuki Asano, and Aaqib Saeed.	1318
1264	Chen Chen. 2022a. Conquering the communication	2023. Federated fine-tuning of foundation models	1319
1265	constraints to enable large pre-trained models in fed-	via probabilistic masking .	1320
1266	erated learning . <i>arXiv preprint arXiv:2210.01708</i> .		
1267	Jingwei Sun, Ziyue Xu, Hongxu Yin, Dong Yang,	Xin’ao Wang, Huan Li, Ke Chen, and Lidan Shou.	1321
1268	Daguang Xu, Yiran Chen, and Holger R Roth. 2023b.	2023a. Fedbft: An efficient federated learning	1322
1269	Fedbft: Efficient federated black-box prompt tun-	framework for bert further pre-training . In <i>Proceed-</i>	1323
1270	ing for large language models . <i>arXiv preprint</i>	<i>ings of the Thirty-Second International Joint Con-</i>	1324
1271	<i>arXiv:2310.01467</i> .	<i>ference on Artificial Intelligence, IJCAI-23</i> , pages	1325
1272	Tianxiang Sun, Yunfan Shao, Hong Qian, Xuanjing	4344–4352. International Joint Conferences on Arti-	1326
1273	Huang, and Xipeng Qiu. 2022b. Black-box tuning for	ficial Intelligence Organization. Main Track.	1327
1274	language-model-as-a-service . In <i>Proceedings of the</i>	Zhousheng Wang, Geng Yang, Hua Dai, and Chun-	1328
1275	<i>39th International Conference on Machine Learning</i> ,	ming Rong. 2023b. Privacy-preserving split learn-	1329
1276	volume 162 of <i>Proceedings of Machine Learning</i>	ing for large-scaled vision pre-training . <i>IEEE</i>	1330
1277	<i>Research</i> , pages 20841–20855. PMLR.	<i>Transactions on Information Forensics and Security</i> ,	1331
1278	Youbang Sun, Zitao Li, Yaliang Li, and Bolin Ding.	18:1539–1553.	1332
1279	2024. Improving loRA in privacy-preserving feder-	Kang Wei, Jun Li, Ming Ding, Chuan Ma, Howard H.	1333
1280	ated learning . In <i>The Twelfth International Confer-</i>	Yang, Farhad Farokhi, Shi Jin, Tony Q. S. Quek, and	1334
1281	<i>ence on Learning Representations</i> .	H. Vincent Poor. 2020. Federated learning with dif-	1335
1282	Rishub Tamirisa, John Won, Chengjun Lu, Ron Arel,	ferential privacy: Algorithms and performance anal-	1336
1283	and Andy Zhou. 2023. Fedselect: Customized selec-	ysis . <i>IEEE Transactions on Information Forensics</i>	1337
1284	tion of parameters for fine-tuning during personalized	<i>and Security</i> , 15:3454–3469.	1338
1285	federated learning . In <i>Federated Learning and Ana-</i>	Herbert Woisetschlager, Alexander Isenko, Shiqiang	1339
1286	<i>lytics in Practice: Algorithms, Systems, Applications,</i>	Wang, Ruben Mayer, and Hans-Arno Jacobsen.	1340
1287	<i>and Opportunities</i> .	2024. A survey on efficient federated learning meth-	1341
1288	Buse G. A. Tekgul, Yuxi Xia, Samuel Marchal, and	ods for foundation model training . <i>arXiv preprint</i>	1342
1289	N. Asokan. 2021. Waffle: Watermarking in feder-	<i>arXiv:2401.04472</i> .	1343
1290	ated learning . In <i>2021 40th International Sympo-</i>	Panlong Wu, Kangshuo Li, Ting Wang, and Fangxin	1344
1291	<i>sium on Reliable Distributed Systems (SRDS)</i> , pages	Wang. 2023. Fedms: Federated learning with mix-	1345
1292	310–320.	ture of sparsely activated foundations models .	1346
1293	Chandra Thapa, Mahawaga Arachchige Pathum	Mengwei Xu, Dongqi Cai, Yaozong Wu, Xiang Li, and	1347
1294	Chamikara, and Seyit Camtepe. 2020. Splitfed:	Shangguang Wang. 2024a. Fwdllm: Efficient fedllm	1348
1295	When federated learning meets split learning . <i>CoRR</i> ,	using forward gradient .	1349
1296	abs/2004.12088.	Mengwei Xu, Yaozong Wu, Dongqi Cai, Xiang Li, and	1350
1297	Chandra Thapa, Pathum Chamikara Ma-	Shangguang Wang. 2023a. Federated fine-tuning of	1351
1298	hawaga Arachchige, Seyit Camtepe, and Lichao Sun.	billion-sized language models across mobile devices .	1352
1299	2022. Splitfed: When federated learning meets split	<i>arXiv preprint arXiv:2308.13894</i> .	1353
1300	learning . <i>Proceedings of the AAAI Conference on</i>	Mengwei Xu, Wangsong Yin, Dongqi Cai, Rongjie	1354
1301	<i>Artificial Intelligence</i> , 36(8):8485–8493.	Yi, Daliang Xu, Qipeng Wang, Bingyang Wu,	1355
1302	Yuanyishu Tian, Yao Wan, Lingjuan Lyu, Dezhong Yao,	Yihao Zhao, Chen Yang, Shihe Wang, Qiyang	1356
1303	Hai Jin, and Lichao Sun. 2022. Fedbert: When feder-	Zhang, Zhenyan Lu, Li Zhang, Shangguang Wang,	1357
1304	ated learning meets pre-training . <i>ACM Trans. Intell.</i>	Yuanchun Li, Yunxin Liu, Xin Jin, and Xuanzhe Liu.	1358
1305	<i>Syst. Technol.</i> , 13(4).	2024b. A survey of resource-efficient llm and multi-	1359
1306	Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier	modal foundation models .	1360
1307	Martinet, Marie-Anne Lachaux, Timothee Lacroix,	Yang Xu, Yunming Liao, Hongli Xu, Zhenguo Ma, Lun	1361
1308	Baptiste Roziere, Naman Goyal, Eric Hambro, Faisal	Wang, and Jianchun Liu. 2023b. Adaptive control of	1362
		local updating and model compression for efficient	1363

1364	federated learning . <i>IEEE Transactions on Mobile Computing</i> , 22(10):5675–5689.	1417
1365		1418
1366	Zheng Xu, Yanxiang Zhang, Galen Andrew, Christopher Choquette, Peter Kairouz, Brendan McMahan, Jesse Rosenstock, and Yuanbo Zhang. 2023c. Federated learning of gboard language models with differential privacy . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 5: Industry Track)</i> , pages 629–639, Toronto, Canada. Association for Computational Linguistics.	1419
1367		1420
1368		1421
1369		1422
1370		1423
1371		1424
1372		
1373		
1374	Fu-En Yang, Chien-Yi Wang, and Yu-Chiang Frank Wang. 2023a. Efficient model personalization in federated learning via client-specific prompt generation . In <i>2023 IEEE/CVF International Conference on Computer Vision (ICCV)</i> , pages 19102–19111.	1425
1375		1426
1376		1427
1377		1428
1378		1429
1379	Xin Yang, Hao Yu, Xin Gao, Hao Wang, Junbo Zhang, and Tianrui Li. 2023b. Federated continual learning via knowledge fusion: A survey. <i>arXiv preprint arXiv:2312.16475</i> .	1430
1380		1431
1381		1432
1382		
1383	Chun-Han Yao, Boqing Gong, Hang Qi, Yin Cui, Yukun Zhu, and Ming-Hsuan Yang. 2022. Federated multi-target domain adaptation . In <i>2022 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)</i> , pages 1081–1090.	1433
1384		1434
1385		1435
1386		1436
1387		1437
1388		1438
1389	Liping Yi, Han Yu, Gang Wang, Xiaoguang Liu, and Xiaoxiao Li. 2024. pfedlora: Model-heterogeneous personalized federated learning with lora tuning .	1439
1390		1440
1391	Shuyang Yu, Junyuan Hong, Yi Zeng, Fei Wang, Ruoxi Jia, and Jiayu Zhou. 2023a. Who leaked the model? tracking ip infringers in accountable federated learning .	1441
1392		1442
1393		1443
1394		1444
1395	Shuyang Yu, Junyuan Hong, Yi Zeng, Fei Wang, Ruoxi Jia, and Jiayu Zhou. 2023b. Who leaked the model? tracking IP infringers in accountable federated learning . In <i>NeurIPS 2023 Workshop on Regulatable ML</i> .	1445
1396		1446
1397		1447
1398		
1399	Sixing Yu, J Pablo Muñoz, and Ali Jannesari. 2023c. Bridging the gap between foundation models and heterogeneous federated learning . <i>arXiv preprint arXiv:2310.00247</i> .	1448
1400		1449
1401		1450
1402		1451
1403	Sixing Yu, J Pablo Muñoz, and Ali Jannesari. 2023d. Federated foundation models: Privacy-preserving and collaborative learning for large models . <i>arXiv preprint arXiv:2305.11414</i> .	
1404		
1405		
1406		
1407	Linan Yue, Qi Liu, Yichao Du, Weibo Gao, Ye Liu, and Fangzhou Yao. 2023. Fedjudge: Federated legal large language model . <i>arXiv preprint arXiv:2309.08173</i> .	
1408		
1409		
1410		
1411	Jianyi Zhang, Saeed Vahidian, Martin Kuo, Chunyuan Li, Ruiyi Zhang, Tong Yu, Guoyin Wang, and Yiran Chen. 2023a. Towards building the federatedGPT: Federated instruction tuning . In <i>International Workshop on Federated Learning in the Age of Foundation Models in Conjunction with NeurIPS 2023</i> .	
1412		
1413		
1414		
1415		
1416		
	Zhuo Zhang, Xiangjing Hu, Jingyuan Zhang, Yating Zhang, Hui Wang, Lizhen Qu, and Zenglin Xu. 2023b. FEDLEGAL: The first real-world federated learning benchmark for legal NLP . In <i>Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 3492–3507, Toronto, Canada. Association for Computational Linguistics.	
	Zhuo Zhang, Yuanhang Yang, Yong Dai, Qifan Wang, Yue Yu, Lizhen Qu, and Zenglin Xu. 2023c. FedPETuning: When federated learning meets the parameter-efficient tuning methods of pre-trained language models . In <i>Findings of the Association for Computational Linguistics: ACL 2023</i> , pages 9963–9977, Toronto, Canada. Association for Computational Linguistics.	
	Wanru Zhao, Yihong Chen, Royson Lee, Xinchu Qiu, Yan Gao, Hongxiang Fan, and Nicholas Donald Lane. 2024. Breaking physical and linguistic borders: Multilingual federated prompt tuning for low-resource languages . In <i>The Twelfth International Conference on Learning Representations</i> .	
	Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, et al. 2023. A survey of large language models . <i>arXiv preprint arXiv:2303.18223</i> .	
	Ligeng Zhu, Zhijian Liu, and Song Han. 2019. Deep leakage from gradients . In <i>Advances in Neural Information Processing Systems</i> , volume 32. Curran Associates, Inc.	
	Weiming Zhuang, Chen Chen, and Lingjuan Lyu. 2023. When foundation model meets federated learning: Motivations, challenges, and future directions . <i>arXiv preprint arXiv:2306.15546</i> .	