

Federated Learning in IoT Environments: Examining the Three-way See-saw for Privacy, Model-Performance, and Network-Efficiency

Roufaida Laidi , Nassima Merabtine, , Djamel Djenouri , Shahid Latif , *Member, IEEE*,
Hemin Ali Qadir , Youcef Djenouri , *Senior Member, IEEE*,
and Ilanko Balasingham , *Senior Member, IEEE*.

Abstract—This survey paper provides an in-depth exploration of Federated Learning (FL) in Internet of Things (IoT) environments, focusing on privacy-preserving techniques and their influence on model performance and network efficiency. It highlights key challenges and opportunities at the intersection of these technologies by offering a comprehensive review of FL applications in IoT. First, a customized taxonomy is introduced to evaluate the privacy levels, quality of service (QoS) and network efficiency of various Privacy-Preserving FL (PPFL) solutions in IoT configurations. Furthermore, the survey investigates strategies to improve FL accuracy while addressing resource and network constraints, both independently and together with privacy preservation techniques. Our findings underscore the complexity of optimizing resource utilization, learning performance, and privacy resilience, revealing that no single PPFL solution universally applies. The paper further identifies future research directions, including the integration of advanced technologies beyond 5G networks, and discusses standards, protocols, real-world PPFL projects from world-renowned industries for potential IoT applications.

Index Terms—Federated Learning, Internet of Things, Privacy-Preserving Federated Learning, Network Efficiency, Data Utility, Cybersecurity, Network security.

I. INTRODUCTION

THE IoT systems and platforms allow devices equipped with sensors to collect and transfer data with little or no human intervention [1]. They are deployed for a broad field of applications such as health care, smart buildings and cities,

manufacturing, and transportation. These applications continuously generate a large amount of data that must be processed, often by ML algorithms, to extract valuable information and gain more insight and intelligence [2], [3]. Traditionally, ML algorithms are trained on a cloud, where all training data is recorded. However, with the tremendous growth of IoT data and the limitations of IoT devices, it becomes infeasible to transfer all the data to remote servers [4]. Moreover, many IoT applications require data to be analyzed in near real-time for rapid decision-making, which is unattainable due to the round-trip delay from IoT devices to the cloud and back. In addition to time and bandwidth constraints, the transmission of sensitive data collected by IoT devices, such as patient electronic health records [5] and other personal information, raises serious privacy concerns.

The concept of FL [6]–[8] has been proposed to provide intelligence-enhanced learning systems while improving their data privacy [9]. FL is an ML approach that enables data owners (IoT devices in our case) to collaboratively train ML models without communicating the data to third parties. Instead, data remains stored on the devices that recorded or generated them. FL can be implemented in a centralized (client-server) or decentralized (peer-to-peer) way. Centralized FL is the architecture most commonly used by state-of-the-art technologies. A typical centralized FL system is triggered by the server that initializes the ML model parameters and hyperparameters and sends them to k selected clients. The latter train their local models and return the calculated parameters to the server. Upon receiving all local parameters from the clients, the server aggregates them and sends the updated parameters again to the clients. This process continues until the termination condition is met. It has already been applied in various IoT and distributed systems applications and services, such as IIoT [10], medical domains and wearable IoT [11], sentiment analysis [12], etc.

Although FL is designed as a privacy-first framework that avoids exposing users' raw data, the regular exchange of model parameters between IoT devices and the server makes it still vulnerable to privacy attacks. Recent works have confirmed that exchanging models' parameters and the final model can leak sensitive information about the user's private data [13]–[17]. Therefore, data privacy in FL systems must be strengthened by using privacy-preserving mechanisms

R. Laidi is with the Norwegian University of Science and Technology (NTNU) N-7491, Trondheim, Norway. E-mail: laidi.roufaida@ntnu.no

N. Merabtine is with Ecole nationale Supérieure d'Informatique, Algiers, Algeria. E-mail: nassimane@gmail.com

D. Djenouri is with the University of the West of England, Bristol, UK. E-mail: Djamel.Djenouri@uwe.ac.uk

S. Latif is with the University of the West of England, Bristol, UK. E-mail: Shahid.Latif@uwe.ac.uk

Hemin Ali Qadir is with Oslo University Hospital, N-0372, Oslo, Norway. E-mail: hemqad@ous-hf.no

Y. Djenouri is with the University of South-Eastern Norway, and NORCE Research Centre, Norway. E-mail: youcef.djenouri@usn.no, yodj@norceresearch.no

I. Balasingham is with Oslo University Hospital and NTNU, Norway. E-mail: ilanko.balasingham@ntnu.no

This work is supported by the EU CHISTER-ERA project (Grant EP/Y036301/1 from EPSRC, UK), and in part by the AGYA Accademy (grant 01DL20003 from BMBF, Germany). This work was carried out during the tenure of an ERCIM 'Alain Bensoussan' Fellowship Programme and is also a contribution to the International Alliance for Strengthening Cybersecurity and Privacy in Healthcare (CybAlliance, Project no. 337316).

such as differential privacy and encryption. However, these privacy-preserving techniques tend to degrade data utility and network efficiency, leading to complex challenges and raising the following questions:

- Do the PPFL solutions proposed for IoT environments respond to all the privacy attacks that FL systems suffer from?
- Are these techniques applicable in real IoT deployments? If so, then what is the cost on the IoT network?
- Could privacy-preserving techniques affect the data utility and the performances of the FL system?

This paper aims to address these open questions by (i) *comprehensively reviewing* PPFL solutions designed for IoT systems, (ii) *proposing* a novel multidimensional evaluation framework to assess these solutions in terms of privacy, quality of service, and network efficiency, and (iii) *identifying* open challenges and future research directions to improve the effectiveness and applicability of PPFL in resource-constrained IoT environments. Table I is the nomenclature of the abbreviations used throughout the paper.

A. Related Work and Research Gap

To highlight the contribution of this work, recent surveys and studies in IoT, FL, and data privacy are reviewed and classified into four categories, as shown in Table II. About 87% of the reviewed papers are published between 2020 and 2024, reflecting the rapid growth and interest in this domain. Below, we discuss each category in more detail:

Cat. A : FL in IoT. Numerous surveys have examined the benefits, potential applications, and challenges of applying FL in resource-limited IoT environments [10], [18]–[23]. These works primarily focus on the technical aspects of FL implementation, such as model design, communication overhead, and computational constraints. However, they often overlook the privacy threats and attacks specific to FL systems.

Cat. B : Privacy in IoT. Data privacy is a critical requirement in IoT applications that handle sensitive information. Surveys in this category focus on privacy issues and enhancement techniques within IoT ecosystems and introduce enhancement techniques to protect user data [24]–[28]. However, most of these works do not address FL-specific privacy concerns or how federated models can be compromised or protected in IoT scenarios.

Cat. C : Privacy in FL. With the advent of FL, new privacy challenges have emerged, leading to a vast body of literature exploring PPFL. These surveys discuss various defense mechanisms as well as future directions [17], [29]–[38]. While they provide a solid foundation for understanding PPFL, they often do not consider the unique constraints and requirements of IoT environments.

Cat. D : Privacy, FL, and IoT. A relatively smaller set of works addresses the intersection of all three domains—privacy, FL, and IoT [39]. These studies acknowledge the need to safeguard sensitive IoT data through FL while maintaining performance and network

TABLE I: Nomenclature of the used terms.

Term	Description
ADASYN	Adaptive Synthetic Sampling
BGV	Brakerski-Gentry-Vaikuntanathan
CNN	Convolutional Neural Networks
DFL	Distributed Federated Learning
DL	Deep Learning
DNN	Deep Neural Network
DP	Differential Privacy
DT	Decision Trees
FedAvg	Federated Averaging
FHE	Fully Homomorphic Encryption
FL	Federated Learning
FNN	Feed-forward Neural Networks
FTL	Federated Transfer Learning
GAN	Generative Adversarial Network
GDPR	General Data Protection Regulation
HBC	Honest but Curious
HE	Homomorphic Encryption
ICT	Information and Communication Technology
IDS	Intrusion Detection Systems
IID	Independent and Identically Distributed
IIoT	Industrial Internet of Things
IoT	Internet of Things
LightGBM	Light Gradient Boosting Machine
MC	Meta Classifier
MEA	Model Extraction Attacks
MIA	Membership Inference Attacks
MKT	Model Knowledge Transfer
ML	Machine Learning
NOMA	Non-Orthogonal Multiple Access
P2P	Peer-to-Peer
PAC	Probably Approximately Correct
PFL	Personalized Federated Learning
PHE	Partially Homomorphic Encryption
PII	Personally Identifiable Information
PPFL	Privacy-Preserving Federated Learning
QoS	Quality of Service
RA	Reconstruction Attacks
RCI	Randomized Client Identifiers
RFID	Radio Frequency Identification
SMOTE	Synthetic Minority Oversampling Technique
SMPC	Secure Multi-Party Computation
SVM	Support Vector Machines
TDMA	Time Division Multiple Access
TTP	Trusted Third Party
VAE	Variational Autoencoder
WSAN	Wireless Sensor and Actuator Networks

efficiency. However, they typically do not provide an in-depth analysis of how privacy-preserving methods influence federated learning outcomes (e.g., accuracy, communication overhead) in resource-constrained IoT settings.

The PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) flow diagram in Fig. 1 illustrates the data collection and screening process for this survey. The data collection process identified 250 records from five scientific databases (IEEE Xplore, ScienceDirect, ACM, Springer, and ArXiv) published between 2017 and 2024 in English. Eligible records included journal articles, conference papers, and book chapters focused on FL, privacy, performance, and network efficiency, with an emphasis on IoT applications.

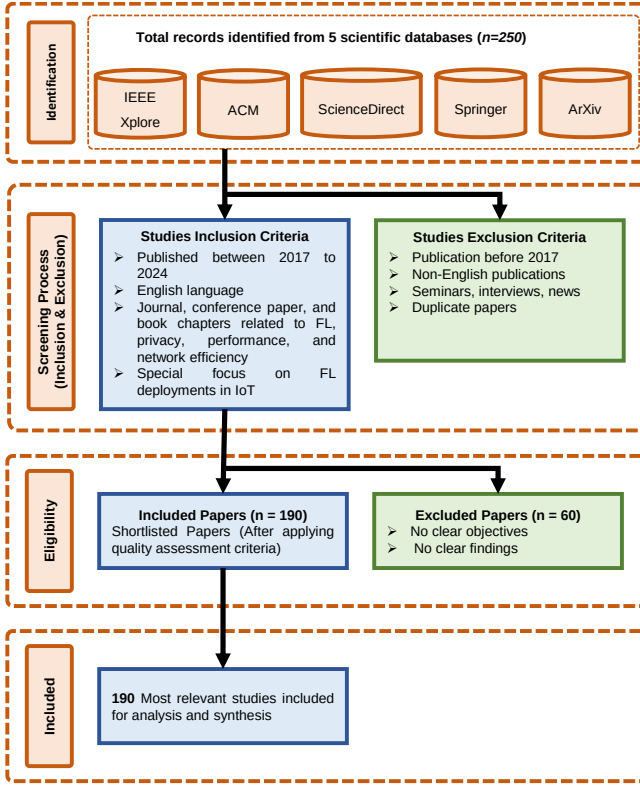


Fig. 1: PRISMA flow diagram for data collection and screening process.

After that, inclusion criteria have been applied, while and excluding duplicates, non-English works, and irrelevant formats such as news or interviews. Further quality assessments excluded 60 studies due to the lack of clear objectives or findings, leaving 190 high-quality studies for analysis and synthesis.

B. Organization

The organizational structure of this survey article is demonstrated in Fig. 2. Section II introduces the FL principles and their implications in IoT, discussing the FL fundamentals, challenges, and trade-offs, as well as a multidimensional evaluation system as described in the paper. Section III explores strategies to improve privacy in IoT FL, detailing privacy vulnerabilities and approaches to preserving privacy. Section IV examines accuracy and learning performance in IoT FL, focusing on techniques to enhance accuracy and balance trade-offs. Section V delves into network efficiency in IoT FL, addressing resource management, scalability, and techniques to enhance network efficiency. Section VI evaluates the proposed PPFL techniques and associated trade-offs, and Section VII discusses technological innovations and future directions. Section VIII covers standards, protocols, and real-world PPFL projects. Finally, Section IX concludes the survey.

II. FL PRINCIPLES AND IOT IMPLICATIONS

A. FL Fundamentals

FL is a distributed ML paradigm, allowing multiple clients to train a global model collaboratively without sharing their

local data, thus preserving data privacy and security. FL addresses data governance and privacy concerns associated with centralized data collection by keeping data on-device and only exchanging model parameters or updates (e.g., gradients). Studies have shown that FL-trained models can achieve performance comparable to centrally trained models and outperform models trained on isolated datasets [40], [41].

In IoT environments, FL architectures can be broadly classified into centralized (client-server) and decentralized (peer-to-peer and blockchain). The centralized and decentralized FL architectures are illustrated in Fig. 3.

1) *Centralized FL Architecture*: In the centralized FL architecture, a central server coordinates the learning process with the IoT devices (clients). The FL process typically involves multiple rounds, each consisting of the following steps [18], [19]:

- **Initialization**: The server selects a subset of k clients based on criteria such as data distribution, device capabilities, and network conditions. The server then distributes the initial global model parameters to these clients.
- **Local Training**: Each selected client trains the received model on its local dataset and computes updated model parameters or gradients.
- **Aggregation**: Clients return their local updates to the server. The server aggregates these updates using algorithms like *FedAvg* [42], *FedProx* [43], or *FedNova* [44] to update the global model.
- **Iteration**: The updated global model is redistributed to clients for further training, and the process repeats until convergence.

2) *Decentralized FL Architecture*: Decentralized FL eliminates the need for a central server, allowing clients to collaboratively train a model in a distributed manner. This architecture can be categorized into two main approaches: P2P decentralized FL and blockchain-based decentralized FL. Both approaches eliminate the need for a central server in the FL process, but differ in how clients communicate, coordinate, and ensure security and trust within the network.

a) *P2P Decentralized FL*: In P2P decentralized FL, clients directly exchange model updates directly with each other without relying on a central server. Communication protocols such as cyclic and random transfer are utilized to disseminate model information throughout the network [45]. In cyclic transfer, clients are organized in a circular chain, where each client sends its model updates to the next client, facilitating a structured flow of information. In random transfer, clients select others to share their model updates, promoting a more dynamic and potentially robust aggregation process. Clients aggregate the received models with their local updates to collaboratively refine the global model.

However, the P2P decentralized FL faces several challenges that need to be addressed. Synchronization among clients is a significant concern due to the absence of a central coordinator, which can lead to inconsistencies and hinder the convergence of the global model [46]. Establishing secure communication channels is crucial to prevent data leakage and ensure the integrity of exchanged model updates, as direct peer-to-peer interactions may expose vulnerabilities. In addition, scalability

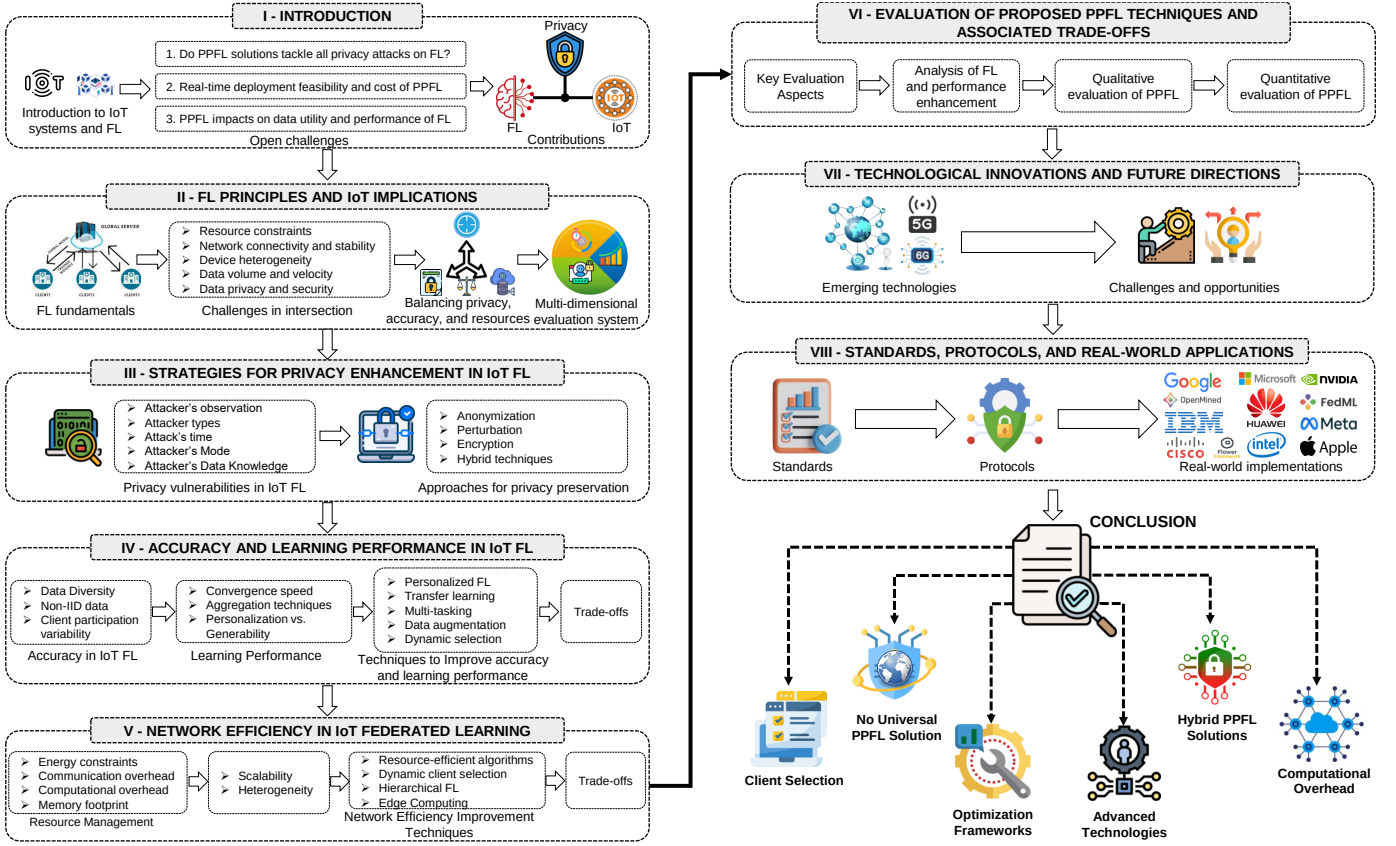


Fig. 2: Organizational layout of the survey.

issues arise as the number of clients increases, resulting in increased communication overhead and potential network congestion, which can degrade the overall efficiency of the learning process [47].

b) Blockchain-Based Decentralized FL: Blockchain technology provides a secure and transparent platform for decentralized FL. In this approach, clients submit their encrypted model updates as transactions to the blockchain network. Consensus mechanisms validate these transactions, ensuring that only legitimate updates are incorporated into the global model. Smart contracts can automate the aggregation process, streamlining collaboration without needing a central coordinator [48]–[50]. Using blockchain, decentralized FL benefits from enhanced security through immutability and consensus features inherent in blockchain systems. This setup allows for transparent tracking of contributions from each client, fostering trust among participants. Furthermore, it opens up possibilities for implementing incentive mechanisms, encouraging clients to participate actively and honestly.

However, blockchain-assisted decentralized FL also introduces new challenges that must be addressed. One significant concern is plagiarism attacks, where malicious clients, also called “lazy clients,” replicate others’ updates and submit them as their own to gain undeserved rewards or reputation. This issue undermines the integrity of the FL process and can degrade model performance. Weng et al. [51] discussed how such attacks can harm blockchain-based FL systems and propose mechanisms to detect and prevent them. Build-

ing on these insights, BLADE-FL [52] examined plagiarism attacks from a resource allocation perspective. Because the presence of lazy clients exacerbates training inefficiencies by plagiarizing others’ models and introducing artificial noise to obscure their behavior, BLADE-FL provides a theoretical convergence analysis that yields an optimal strategy for balancing the frequency of model updates (blocks) with the computation time spent on training. To further mitigate plagiarism, T-BFL (Trustworthy Blockchain-Assisted Federated Learning) [53] introduced a decentralized reputation management (DRM) mechanism, where nodes’ contributions are independently tested and mapped to dynamic reputation scores. These scores directly influence aggregation weights and consensus difficulty—malicious clients with low reputation face higher mining barriers, disincentivizing plagiarism. Similarly, CPoC (Context-aware Proof-of-Contribution) [54] countered lazy clients by computing a “global contribution” score that holistically evaluates clients’ efforts across training, verification, and consensus steps. By enforcing encryption and public-key-based model sharing for verification, CPoC limits plagiarized submissions while aligning rewards with verifiable contributions, which improves fairness and efficiency over traditional schemes.

Furthermore, the resource constraints of IoT devices pose a hurdle, as the handling of blockchain operations can be computationally intensive. The latency introduced by consensus protocols can affect the efficiency of the learning process, particularly in large-scale deployments [55], [56]. Deng et al.

TABLE II: Existing studies in the IoT, FL, and privacy research fields.

Cat.	Covered Topics	Related Work	Key Contributions	Limitations
A	FL, IoT	[10]	Study of challenges of applying FL in resource-constrained IoT devices.	Don't consider privacy threats and attacks of FL systems.
		[18]	Utilization of FL in IoT and the opportunities FL creates for IoT services and applications.	
		[19]	A brief overview of the applications and challenges of FL in IoT.	
		[20]	Study the use and the challenges of FL for IDSs.	
		[21]	Challenges and future directions of applying FL in vehicular IoT.	
		[22]	The use of FL in key IIoT applications and services	
		[23]	Recent advances and open research challenges of FL applied to IoT + proposition of a set of evaluation metrics for FL systems.	
B	Privacy, IoT	[24]	A study on privacy concerns in IoT environments.	These papers discuss privacy concerns in IoT environments, not the privacy attacks raised by FL systems nor the corresponding privacy-preserving techniques.
		[25]	Privacy issues and privacy-enhancing techniques on smart metering applications.	
		[26]	Privacy threat issues, privacy legislation, privacy enhancing technologies in the IoT.	
		[27]	A study and comparison of privacy-preserving techniques proposed for the IoT ecosystem + analysis vs. the EU's GDPR.	
		[28]	A study on security and privacy problems in the cyber-physical world.	
C	Privacy, FL	[29], [30], [31]	Security and privacy of FL: challenges, defense mechanisms and future directions.	The challenges of PPFL solutions are not discussed when applied to the particular context of IoT applications.
		[32]	Privacy-preserving FL + systematic taxonomy of privacy leakage risks in the FL systems.	
		[33]	A brief overview of the challenges and solutions of data privacy in FL.	
		[34]	Privacy-preserving techniques in FL regarding the EU/UK GDPR.	
		[35], [17], [36], [37], [38]	Threats and attacks of FL, discussion on the privacy-preserving techniques.	
D	Privacy, FL, IoT	[39]	Privacy preservation in FL and the challenges of applying PPFL to the IoT.	Does not discuss the potential effect of privacy-preserving techniques on the effectiveness of FL systems and the privacy attacks they deal with. It does not study PPFL solutions proposed for IoT environment.

[57] addressed this by integrating training and block mining at the client side, optimizing resource allocation via a Lyapunov-based framework (DRACS) to maximize training efficiency under energy constraints. By addressing these challenges, ongoing research continues to explore efficient consensus mechanisms [9], robust incentive models [53], [54], [58], and scalable architectures to enhance security and trust in decentralized FL while maintaining the efficiency necessary for IoT environments.

FL systems can also be categorized on the basis of data partitioning into:

- Horizontal FL: Clients possess datasets with the same feature space but different samples. FL focuses on learning among different users with similar data types.
- Vertical FL: Clients have datasets with the same sample IDs but different feature spaces. Useful when institutions have different attributes of the same individuals.
- Federated Transfer Learning: Combines horizontal and vertical FL to handle datasets with little overlap in both samples and features.

A comprehensive exploration of horizontal, vertical, and federated transfer learning, along with their respective challenges and potential solutions, can be found in [59].

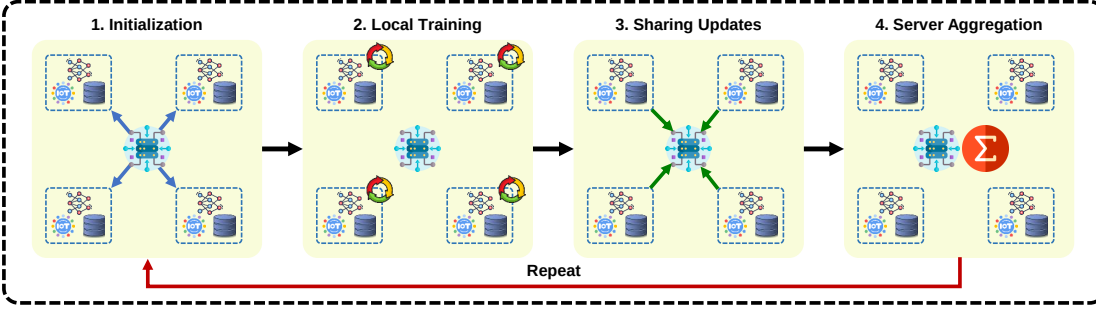
B. FL in IoT

FL addresses several critical challenges that IoT deployments face, including privacy, network performance, and learning efficiency, making it an essential technology to take full advantage of the full potential of distributed IoT ecosystems. FL fundamentally transforms data privacy and security protocols within IoT networks [60]. IoT devices often collect and process sensitive data from users or environments, such as personal health metrics from wearable devices or operational data from industrial machinery. FL allows these devices to learn from shared models without exposing or centralizing the data, thus greatly enhancing privacy and reducing the risk of data breaches [61].

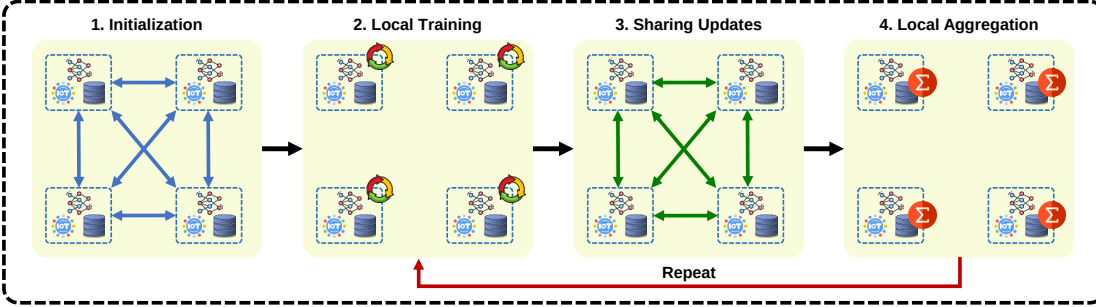
FL also addresses the technical constraints of the IoT infrastructure, notably in terms of network bandwidth and latency. Unlike traditional cloud-based models, FL processes data locally on IoT devices and only exchanges small model updates [62], typically much smaller than the raw data itself. This significantly reduces bandwidth required and decreases latency, which is crucial for IoT applications that often rely on real-time or near-real-time data processing.

Furthermore, scalability becomes a formidable challenge with the exponential growth of IoT devices. FL offers a scalable ML solution that does not require proportional increases

1. Centralized Federated Learning



2. Peer to Peer Federated Learning



3. Blockchain Federated Learning

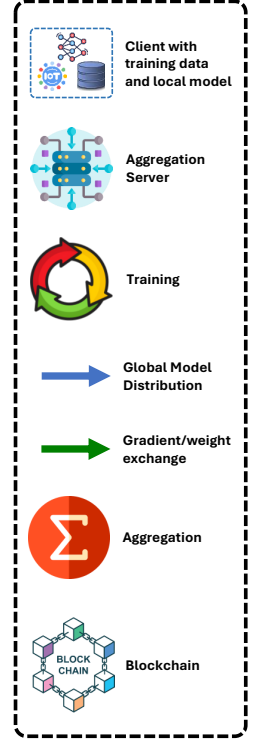
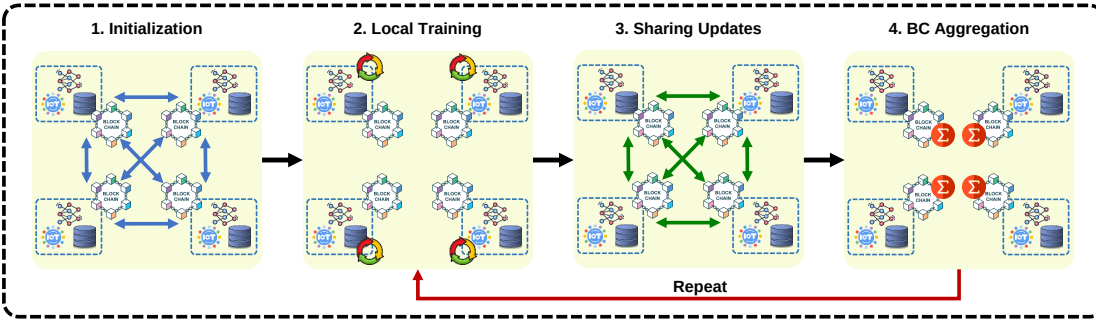


Fig. 3: FL architectures.

in central infrastructure [63]. By distributing computations across numerous devices, FL uses their collective power efficiently. This distributed approach supports scalability and conserves energy [64], which is particularly beneficial for battery-dependent devices in remote or inaccessible areas.

IoT FL also increases the diversity of data and the robustness of the model. The diversity of IoT deployment contexts gives FL a unique advantage; models trained through FL benefit from varied data reflective of real-world environments and scenarios. This diversity helps to develop robust generalizable models that are more adept at handling different operational conditions [65]. Continuous learning through regular updates further enables these models to adapt over time, enhancing their accuracy and relevance as environmental conditions evolve [66].

FL represents a promising paradigm shift for deploying ML in IoT environments. It aligns well with the distributed nature of IoT and addresses many of its fundamental challenges, enabling smarter, more private, and efficient use of IoT data. However, the intersection of FL and IoT results in unique

challenges detailed in the following section.

C. Challenges in the Intersection

Despite its benefits, the integration of FL into the IoT introduces specific challenges. Issues such as device heterogeneity, security vulnerabilities specific to distributed networks, and managing dynamic device populations must be addressed to fully realize FL's potential [23]. Fig. 4 illustrates the array of challenges that FL encounters in the context of IoT environments.

The following section describes IoT-specific challenges and IoT data characteristics and handling challenges.

1) IoT-Specific Challenges for FL:

a) *Resource Constraints*: IoT devices typically suffer significant limitations in computational power, storage capacity, and battery life. These resource constraints critically impact the design and training of FL models [67]. FL requires devices to perform computations locally, which can be challenging for devices with limited processing capabilities. Furthermore, storing intermediate model updates during the

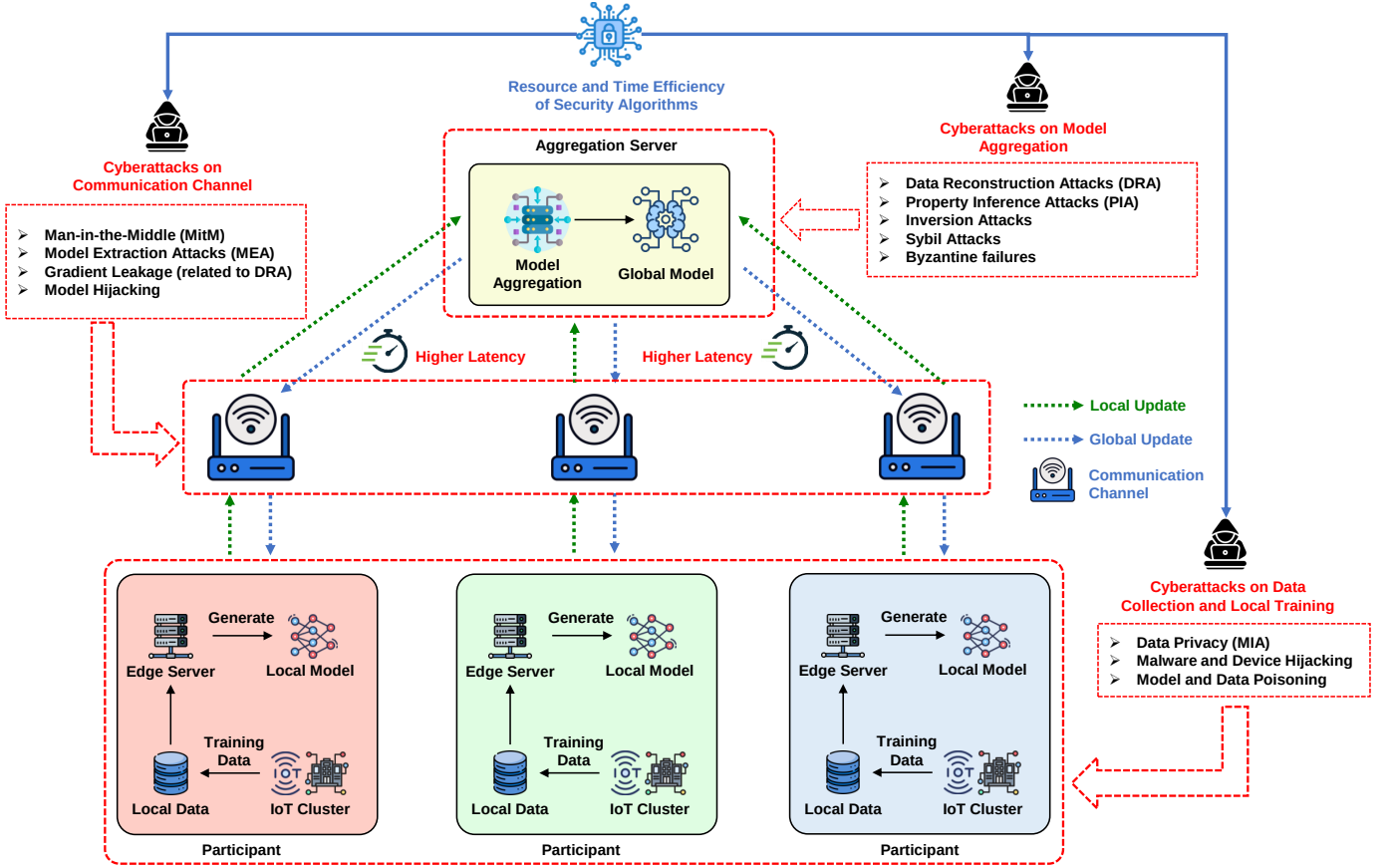


Fig. 4: FL security and privacy challenges in IoT.

learning process demands substantial memory, which may only be feasible for some IoT devices. Lastly, continuous communication required to send model updates to a central aggregator or other devices can rapidly drain battery life, affecting the device's primary functionalities [67].

b) Network Connectivity and Stability: The effectiveness of FL is highly dependent on network conditions, which can be extremely variable in IoT settings [68], [69]. High latency and intermittent connectivity can severely disrupt the synchronization for successful model training across devices. All devices must contribute simultaneously in synchronous FL, which becomes problematic with unstable network connections [70]. Asynchronous FL methods can mitigate some connectivity issues but introduce challenges in maintaining model consistency and convergence, as outdated or delayed updates from some devices can skew the aggregate model learning process [71].

c) Device Heterogeneity: The wide variety of IoT devices, each with different hardware capabilities, operating systems, and data handling protocols, poses another significant challenge for FL. This heterogeneity affects the uniformity of the learning process and the performance of the model. For example, differences in data collection methods and formats can result in non-IID (independent and identically distributed) data, complicating the aggregation of model updates and leading to biased or suboptimal learning outcomes [72]. Moreover, the variance in computational capabilities and storage options

requires the development of FL models that can adapt to the least capable device, which can compromise the efficiency and effectiveness of the learning process [72].

2) IoT Data Characteristics and Handling Challenges:

a) Data Volume and Velocity: IoT environments are characterized by their ability to generate large volumes of data at high velocities. Managing this data flood in an FL context is challenging due to the need for rapid aggregation and model updating to keep pace with incoming data [73]. The sheer volume complicates the handling and processing of data on local devices, while the speed demands quick turnaround times for the learning to be effective and relevant. This can strain the network and device resources, potentially leading to bottlenecks or delays in model convergence [74].

b) Data Privacy and Security: Data often includes sensitive personal or operational information in IoT configurations, making privacy and security paramount. The distributed nature of FL, where data is supposed to remain on local devices, inherently supports privacy. However, transmitting model updates across networks introduces vulnerabilities in which data can be intercepted or inferred. Furthermore, the diverse and widespread deployment of IoT devices increases the risk of security breaches, where a compromised device could affect the integrity of the entire federated model.

c) Levels of Privacy: The privacy measures can be categorized into four distinct levels of protection, offering a structured approach to assess the robustness of various

techniques.

- 1) Minimal Protection: FL solutions ensure that raw data remains localized on devices at this foundational level. Although this inherently provides privacy by not centralizing data, the exchanged model updates or gradients might still be relatively transparent, making the system vulnerable to basic attacks.
- 2) Moderate Protection: Solutions in this category focus on protecting against inference attacks. They may add noise to model updates or use simple anonymization methods. These steps block basic inference attempts, but may not stop more advanced attacks.
- 3) High Protection: At this level, advanced mechanisms tailored to counteract inference attacks are implemented. Techniques such as differential privacy introduce calibrated noise to model updates. Secure aggregation ensures that individual updates are combined in an encrypted space, significantly reducing the chances of successful attacks.
- 4) Ultimate Protection: Solutions at this pinnacle offer the most robust defense against privacy attacks. They integrate state-of-the-art privacy-preserving techniques, regularized training, and adaptive mechanisms that respond quickly to detected inference attempts. These solutions are designed to make it computationally infeasible, if not impossible, for adversaries to succeed.

D. Navigating Trade-offs in FL for IoT: Balancing Privacy, Accuracy and Device Limitations

This section explores the trade-offs inherent in FL deployment in IoT environments, emphasizing the complex interaction between privacy, learning accuracy, and the operational limitations of IoT devices. Each subsection is structured to highlight how different elements affect each other.

1) *Privacy vs. FL Learning Accuracy*: FL inherently promotes privacy by training on local data and only sharing updates (often gradients) rather than raw data. However, gradients can still leak information about the data. Thus, to increase privacy, gradients can be aggregated, masked, encrypted, or quantized (e.g., FL based on signs [75]). However, these processes diminished the fidelity of the gradient and may affect the accuracy of the learning. For example, differential privacy adds noise to the data or gradients, which can impact the accuracy of the model [76]. Finally, data heterogeneity is prevalent in IoT settings, and device data are often non-IID. Although privacy-enhancing techniques can help protect individual data contributions, the global model synthesized from such data may not perform optimally on local datasets. This introduces a trade-off between the global utility of the model and its local efficacy on individual devices.

2) *IoT Limitations vs. FL Learning Accuracy*: IoT devices typically possess limited computational power and storage capacity, which may prevent direct training of complex models on devices. Simpler models, while less computationally demanding, often fail to capture the complexities of the underlying data as effectively as more sophisticated models. Furthermore, computational limitations prevent all devices from

computing and synchronizing updates. This asynchronicity can lead to issues with stale gradients, negatively impacting the model's convergence and overall accuracy. In addition, model compression or pruning is employed to accommodate device limitations. Although these approaches help fit the model within the device's capabilities, excessive compression can strip away critical information, diminishing the model's accuracy.

Many IoT environments suffer from limited connectivity, which complicates the transmission of large gradients. Techniques that reduce the size of these updates, such as gradient quantization or implementing sparse updates, help conserve bandwidth, but at the potential cost of reduced model accuracy. In addition, connectivity issues can cause devices to temporarily disconnect or introduce significant delays. Such disruptions can affect the timeliness and quality of model updates, affecting both the convergence and accuracy of the FL process.

3) *IoT Limitations vs. Privacy*: Although encryption techniques significantly improve data privacy, they also introduce substantial computational overhead. The limited processing power of IoT devices often makes the application of intensive cryptographic techniques impractical. Furthermore, encrypted data typically resist efficient compression, so transmitting encrypted gradient updates requires more bandwidth. This can be particularly challenging for IoT devices operating under bandwidth restrictions.

Fig. 5 visualizes the trade-offs between key dimensions such as Privacy, Learning Accuracy, Resource Consumption, Bandwidth Consumption, and Energy Consumption in various FL scenarios tailored for IoT applications. Each scenario is plotted to demonstrate how it manages these critical dimensions for optimizing FL deployments in IoT environments. The chart highlights the balance between maintaining high data privacy and achieving efficient and accurate model performance under varying resource constraints.

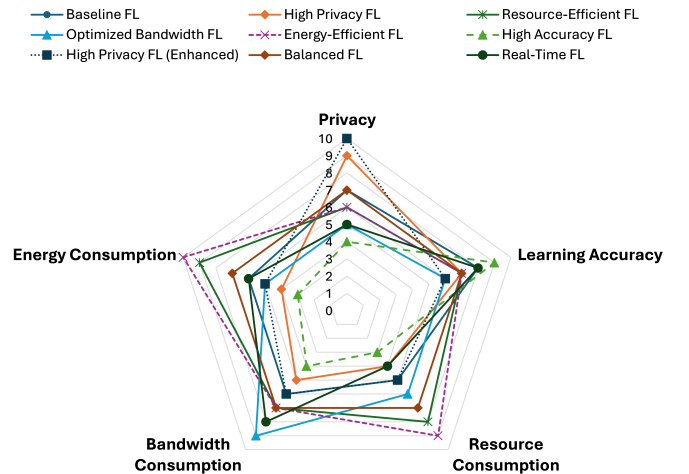


Fig. 5: Comparative analysis of FL scenarios in IoT environments.

As the adoption of FL in IoT settings increases, an intricate balancing act must be maintained. Enhancing privacy comes

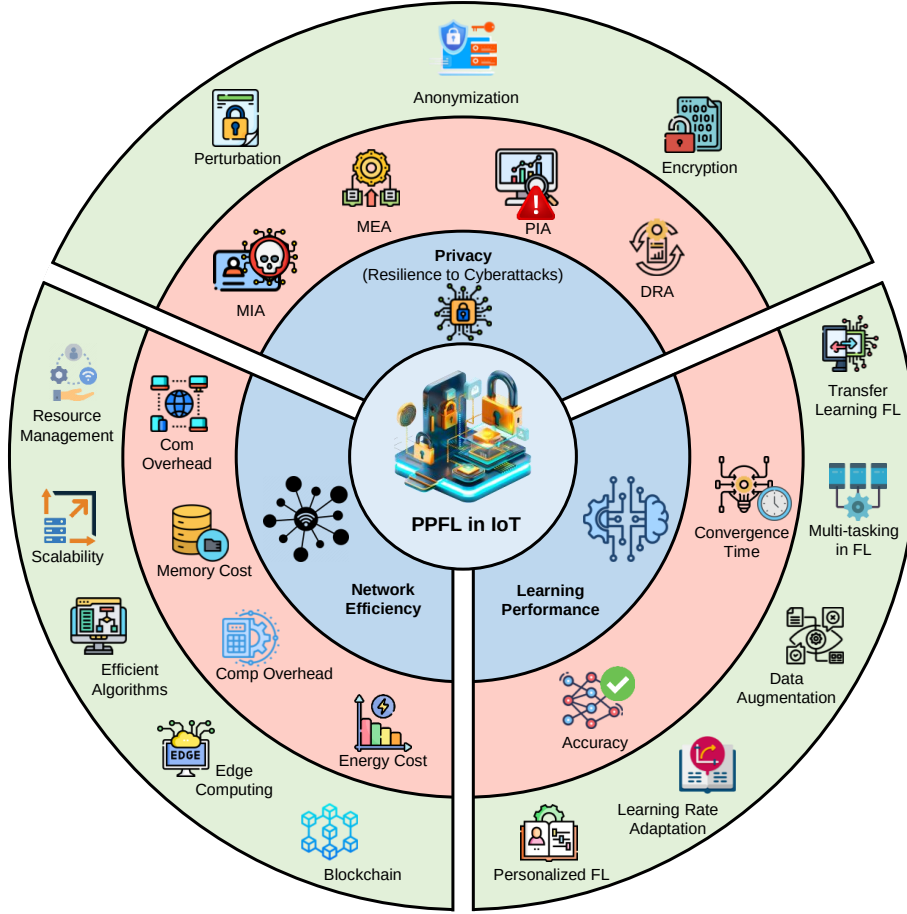


Fig. 6: Proposed multidimensional evaluation system of PPFL in IoT.

at the cost of learning accuracy, and catering to IoT limitations could compromise both. The ideal scenario would ensure robust privacy without significant degradation in model accuracy, while respecting the resource constraints of IoT devices.

E. Multidimensional Evaluation System to FL in IoT: Privacy, Performance, and Efficiency

As previously discussed, exchanging model parameters can inadvertently leak sensitive data. Therefore, several privacy-preserving techniques have been developed using different and sometimes combined defense mechanisms (as detailed in Section III). However, applying these privacy techniques can deteriorate FL performance and IoT network efficiency [32]. The literature discussed in this paper is analyzed from three perspectives: 1) privacy, 2) FL performance, and 3) network efficiency. This framework is a tool for theoretically analyzing existing PPFL in IoT and a guide for designing efficient FL systems with a balanced privacy, service/quality, and efficiency trade-off. Privacy refers to the resilience against privacy attacks the FL system might face. FL performance indicates the convergence error and time, and network efficiency refers to the communication and computation overhead generated by the PPFL and the required memory footprint. Fig. 6 illustrates the paper's proposed evaluation system.

As demonstrated in [77] through the application of the No-Free-Lunch (NFL) theorem to FL, it is unrealistic to expect

an FL algorithm to simultaneously excel in privacy, utility, and efficiency under certain conditions. Consequently, most solutions that incorporate a privacy preservation mechanism for IoT FL focus on optimizing learning utility or network efficiency, but not both. Building on this foundation, the remainder of this paper will first explore privacy threats and their countermeasures in FL for IoT. Subsequent sections will dive into the effects of these privacy preservation techniques on 1) accuracy and learning performance and 2) network efficiency in IoT FL environments.

III. STRATEGIES FOR PRIVACY ENHANCEMENT IN IoT FL

Personal data collected by IoT devices, including location, physical activities, and medical information such as heart rate or blood pressure, are sensitive and can cause privacy concerns if accessed without authorization. In the context of IoT FL, ensuring privacy is both a priority and a challenge. This section dives into the intricacies of privacy concerns in FL, presenting the spectrum of potential attacks and their corresponding mitigation techniques. Tailored defenses are detailed for each attack, providing a robust understanding of available countermeasures.

A. Privacy Vulnerabilities in IoT FL

Numerous recent studies [13], [30], [32], [78]–[84] addressed privacy concerns in FL systems, showing that sensitive

data can be acquired through intermediate gradients, model parameters, or the final model. This section provides an overview of privacy attacks in FL and their main categories. To better understand these concepts, consider an ML model f_θ trained in an FL environment with N participants, where each participant i has a unique training dataset D_i , P_i data features or properties, and local parameter θ_i .

1) *Attacker's Observation: Black-box vs. White-box Inference* : The adversary can only see the model's output in the black-box setting on different inputs. It cannot access the model's parameters or intermediate steps of computation. Therefore, for any data point d_i^j , the attacker can only obtain $f(d_i^j; \theta)$. An example of this type of setting can be found in ML-as-a-service platforms. In white-box attacks, the attacker can access the model parameters θ , input, output, and architecture. The attacker can compute all intermediate states given a data point d_i^j . Finally, gray-box access describes situations in between.

2) *Attacker Types*: The server or a client can carry out attacks. A curious server can collect updates from each participant represented by W_i^t to gain insight into the training set of each participant. In addition, a malicious server can manipulate the view of each participant to extract more information about their training data. Similarly, a participant can act adversarially by observing global parameters W_i^t and uses its updates W_i^t to gain information about the union of training data from other participants. In both cases, the adversary can analyze multiple versions of the target model over time, revealing sensitive information about the training data. Finally, attackers can also be outsiders, that is, final users and eavesdroppers on the communication channel.

3) *Attack's Time: Training vs. Inference*: During the training process, an attacker can access the model's updates, which can result in privacy breaches through the embedding and fully connected layers, as well as the gradients. In addition, the attacker can manipulate intermediate gradient updates to uncover the training data of the participants. However, during the inference or prediction/deployment phase, the attacker can only access the final model.

4) *Attacker's Mode: Active vs. Passive*: During active attacks, the adversary involved in the training process can manipulate the target model to extract more information about its training data. A server or curious participant can create malicious parameter updates to prepare for a future attack. However, passive attacks generally occur during the inference phase, as no updates can be made to the model.

5) *Attacker's Data Knowledge: Supervised vs. Unsupervised* : During the training process, if an adversary has access to a dataset D_i that overlaps with the global dataset D , they can utilize this information to train their supervised attack models, targeting the remaining training data. However, when the adversary lacks samples from the target training set, they have two primary avenues for training their attack models: supervised training using shadow models and unsupervised training.

- **Supervised Training with Shadow Models**: A shadow model is a privacy attack technique in which the adversary trains a new model that mimics the behavior of

the target model. To create a shadow model, the attacker generates a training dataset D_s that is sampled from the same distribution as the target training data D . This can be accomplished by using publicly available data or synthesizing data using data augmentation or generative modeling techniques. The shadow model is then trained on D_s , ensuring that it does not have any known overlap with the target training set.

Once the shadow model is trained, the attacker analyzes its behavior, particularly how it generalizes to unseen data. Since the shadow model's architecture mirrors the target model's, the attacker can infer valuable information about its functionality by observing how the model performs on its own training data. Empirical studies indicate that attack models developed based on shadow model behavior tend to be effective against the target model [85].

- **Unsupervised Attacks**: In cases where the adversary has access to a dataset D' that partially overlaps with the target training set D , unsupervised attacks may be feasible. In this context, the attacker may not know which specific data points belong to the intersection $D' \cap D$. However, the attacker can obtain information on the distribution of data using clustering techniques or anomaly detection methods on D' . This information can be used to craft attacks that exploit potential vulnerabilities in the target model, even without explicit labels or direct examples from the target training data.
- **Privacy Implications of the Attacks**: Through supervised training with shadow models and unsupervised attacks, the adversary can derive varying levels of privacy information. In supervised scenarios, the adversary can extract the target model's specific attributes and decision boundaries, potentially revealing sensitive data characteristics. Although the information gleaned may be less precise in unsupervised attacks, it can still provide insight into the data distribution and model behavior, leading to privacy breaches through targeted inference attacks.

6) *Privacy Attack Types*: FL systems can face different types of attack that fall into the following categories: Membership Inference Attacks [86], Model Extraction Attacks [87], Property Inference Attacks [88], and Data Reconstruction Attacks [89]. MIA and DRA aim to uncover individual data samples d_i^j of the training dataset D_i , while PIA targets certain properties p_i^j . MEA involve stealing or approximating the parameters θ of the trained model f_θ . A comparison of privacy vulnerabilities in IoT FL is represented in Table III.

a) *Membership Inference Attacks*: MIA attacks are primarily used to determine whether a particular data record was part of the training dataset. It can also determine whether a sample belongs to a specific class in the model. For instance, an attacker may utilize an MIA attack to determine if a patient's clinical record was used to train a disease detection model, thereby exposing whether the patient is affected by the disease.

To successfully attack, MIA assumes that the attacker has access to two elements. Firstly, the trained target model is f_θ . The more information the attacker has about the model, the easier it is to attack. The attacker must have query access

at a minimum. Secondly, the attacker needs a query dataset D' , which should ideally include some training data samples d_i^j that were potentially used to train f_θ . This means that d_i^j must be present in both D_i and D' . The attacker must have a dataset that contains samples d' in a distribution similar to that of the original D_i . The ultimate goal of MIA is to identify which of the samples, d' , were used to train the target model, f_θ . Overfitting and poor generalization can greatly affect a model's vulnerability, making it easier for MIAs to succeed, even with just black-box access.

Large corporations utilize user data and implement ML models on a wide scale. This poses a danger as users may be wrongly identified or re-identified if the model is accessed. The threat is further amplified if the trained models are available in open or semi-open formats. Although GDPR safeguards user privacy, it does not extend to ML models [90]. Nevertheless, MIAs may be used to locate individual records employed to train open-access ML models, thereby risking user data privacy. Finally, from the data owner's perspective, these attacks may be able to audit black-box models to check if their data were used without permission.

b) Model Extraction Attacks: Model extraction is a black-box attack in which an attacker aims to fully reconstruct a targeted model f_θ by creating a substitute model f'_ω that exhibits similar behavior. The process of creating this substitute model typically involves the attacker querying the target model multiple times to gather input-output pairs. By systematically selecting a diverse set of input samples, the attacker can collect responses from the target model, which serve as training data for the substitute model [91].

There are two primary goals for developing substitute models:

- **Task-Specific Accuracy:** In this scenario, the attacker endeavors to construct a substitute model that matches or exceeds the accuracy of the target model on a specific test set. This involves using the input-output pairs collected from the target model to train the substitute model on the same learning task, ensuring that it learns to produce similar outputs for the same inputs.
- **Decision Boundary Approximation:** Alternatively, the attacker may aim to replicate the decision boundary of f_θ as closely as possible. This can involve sampling inputs unrelated to the original learning task, allowing the attacker to focus on capturing the overall behavior of the target model rather than achieving high accuracy on any specific dataset.

In both cases, the effectiveness of the substitute model can depend on the complexity of the chosen architecture. While knowledge of the target model's architecture can enhance the attack's effectiveness, it is only sometimes necessary. An adversary can select a substitute model that is equally complex or even more complex than the target model to increase the chances of successful replication. In addition to creating substitute models, attackers may utilize other techniques to extract information from the target model, including obtaining hyper-parameters of the objective function and details about the neural network architecture, such as activation functions, optimization algorithms, and the number of layers.

c) Property Inference Attacks: PIA are a class of white-box attacks that aim to extract sensitive information that a model has learned, such as latent characteristics of the training dataset that were not explicitly included as features, but may correlate with the learning task [92]. This leakage poses significant privacy risks, as it allows adversaries to gain insight into the training data, potentially enabling them to replicate similar models or exploit system vulnerabilities.

In executing a PIA, the attackers aim to create an MC to discern whether a target model f_θ includes a specific attribute p_i^* . To construct the MC, the attackers first generate a set of shadow classifiers trained on a dataset similar to the original dataset but containing only a subset of instances with the attribute p_i^* . These shadow classifiers do not directly learn the attribute p_i^* ; instead, they implicitly capture it due to inherent biases in the training dataset.

The process begins by selecting a diverse training set that mimics the characteristics of the original dataset D . The shadow classifiers are trained on this set, where the attribute p_i^* is present in some instances and absent in others. By training the shadow models in this way, the attackers can observe how variations in the presence of p_i^* affect the output of these classifiers. Once the shadow classifiers have been trained, the attackers collect the output predictions and the corresponding model parameters (weights and biases) from f_θ when evaluated on a set of inputs. These outputs serve as features for training the meta-classifier. The MC learns to classify the target model f_θ as possessing or lacking the attribute p_i^* , based on the relationships observed during the training phase. This enables attackers to effectively infer sensitive properties within the target model, highlighting the vulnerabilities associated with PIA.

d) Data Reconstruction Attacks: DRA seek to rebuild training samples accurately and related labels used during training. One of the well-known attacks of this category is the deep leakage of the gradient [93], which aims to reveal the private training data from the gradients, which can obtain the training inputs and the labels in only a few iterations. The core idea of this kind of attack is to synthesize pairs of "dummy" inputs and labels by matching their "dummy" gradients close to the real ones.

B. Approaches for Privacy Preservation in IoT FL

Various privacy-preserving solutions have been proposed and their effectiveness has been evaluated in several studies [27], [94]–[97]. These approaches are summarized and classified into four main types: anonymization-based, perturbation-based, cryptography-based, and hybrid. In the following these approaches are presented, and we also describe how they counteract the attacks introduced in the previous section.

1) Anonymization-based: Anonymization is usually the first line of defense, and it includes two main techniques used on sensitive data: i) pseudonymization and, ii) anonymization. Pseudonymization replaces sensitive information with new data that can be used to re-identify it later. On the other hand, anonymization involves removing PII, such as names and identification numbers, while still maintaining the usefulness of

TABLE III: Comparison of privacy vulnerabilities in IoT FL.

Attack Type	Attack Visibility	Attacker Type	Prior Knowledge Needed	Objective
MIA	Black-box	Malicious clients or external observers	Access to model output predictions	Determine if a specific data point was part of the training dataset
MEA	Black-box	Malicious clients or external users	Knowledge of the API or direct access to model output	Reconstruct an approximate or exact copy of the model
PIA	White-box	Malicious clients or curious participants	Knowledge of the target model's training setup and specific property characteristics	Infer properties of the training dataset, even if not directly linked to the prediction task
DRA	White-box	Malicious clients or adversarial collaborators	Full or partial access to gradients or model updates	Reconstruct original training data based on gradients or model updates

the data. Anonymization is developed for structured data with three attributes: i) unique identifiers, ii) sensitive attributes, and iii) nonsensitive attributes. Popular anonymization techniques include k-anonymity, l-diversity, and t-closeness [98], [99]. K-anonymity protects the privacy of unique identifier-based records but may be vulnerable to inference attacks on sensitive attributes. L-diversity adds diversity within a group to sensitive attributes but may be vulnerable to attribute linkage attacks. T-closeness maintains the distribution of sensitive attributes, making it more effective for numeric attributes. However, enforcing t-closeness can degrade the usefulness of the data.

Resilience Against Attacks:

- **MIA:** Anonymization techniques such as k-anonymity ensure that each record is indistinguishable from at least $k - 1$ other records in the dataset. Blurring the lines between individual data points makes it harder for attackers to confidently assert whether a specific data point was used in model training.
- **PIA:** Techniques such as l-diversity and t-closeness enhance privacy by ensuring that the sensitive attributes of any group of records are well represented or closely follow the overall distribution of the dataset. This distributional cloaking prevents attackers from making accurate inferences about subgroup properties based on the output model.

2) *Perturbation-based:* These techniques involve adding noise to the original data to make the statistical information calculated from the perturbed data statistically similar to the original data. Common types of perturbation techniques include DP, additive perturbation, and multiplicative perturbation. DP modifies a dataset or algorithm to protect individual identity while maintaining the overall statistical distribution of the data. This means that an outside observer should not be able to determine whether a specific individual's data were used to obtain a result from the dataset. Essentially, DP allows for statistical reasoning without compromising individual privacy. To establish this, a function is used to add selected random noise to the original response of a ML model. This creates a consistent uncertainty across all records (privacy budget) and reduces the likelihood of exposing any particular one. DP uses various methods such as randomization, Laplace method, and exponential mechanism [100], [101]. DP protects

against reidentification attacks, like linkage or set differencing. It can be applied to input data (local DP), computation results (global DP), or the algorithm during training or inference. However, applying DP to image data might be challenging, as manipulating the data may degrade its quality, which could negatively affect the algorithm's performance. There is still uncertainty about implementing DP in imaging data, unlike tabular data, which could be easily shuffled. Therefore, the effects of perturbing images are unpredictable.

Additive perturbation involves adding random noise from a distribution (such as uniform or Gaussian). Although simple and able to maintain statistical properties, it may reduce the usefulness of the data and be vulnerable to noise reduction. However, multiplicative perturbation involves multiplying data with noise from a distribution, which results in the transformation of data points to a specific space. This technique is more effective than additive perturbation, as it is difficult to reconstruct the original data from the perturbed data.

Resilience Against Attacks:

- **MIA:** By adding controlled noise to data or model parameters, DP ensures that the system's output is virtually indistinguishable, whether a specific record is part of the training data. This obfuscation hampers the attacker's ability to ascertain data membership. Furthermore, introducing noise to the model's output can hinder an attacker's efforts to deduce whether a specific input was part of the training set.
- **MEA:** DP ensures that the specifics of any local dataset remain confidential by adding noise to model updates, which hampers the attacker's ability to recreate an accurate copy of the central model.
- **PIA:** Adding noise makes it considerably difficult for attackers to identify unique properties or characteristics inherent in training data.
- **DRA:** It introduces uncertainty in the data, making precise reconstruction of the original data challenging for attackers.

3) *Encryption-based:* The process of defending the appropriate encryption algorithms and parameters is complex, as well as their implementation. Brute force attacks cannot crack current cryptographic algorithms, which can be used to secure both the algorithm and data for secure joint computation. Stan-

standard cryptographic techniques for privacy-preserving machine learning include homomorphic encryption, secret sharing, and secure multiparty computation.

Homomorphic encryption (HE) is the most widespread cryptographic solution, allowing calculations to be carried out directly on encrypted data without the need to decrypt it. HE schemes can be classified as partially or fully homomorphic. Partially HE supports limited operations, such as addition and multiplication, while fully HE schemes support additional computations, such as quadratic functions. There is a trade-off between efficiency and security when using HE, with computational performance being the main concern [102], [103]. Despite this, HE has shown success in CNN and has benefits in “ML-as-a-service” scenarios. It can also securely aggregate encrypted algorithm updates in FL scenarios, with or without additional DP.

Secret sharing is a scheme in which a secret key comprising several shares can only be reconstructed if enough shares are combined. SMPC is a method of secure computation in which data is split among multiple parties to prevent any single party from accessing the entire dataset. The result of the computation can be announced without exposing the whole data, which can only be accessed through consensus. This method is useful for “secret sharing” in semi-trusted or low-trust environments. By performing analysis on encrypted datasets, SMPC can increase the amount of available data without revealing individual identities or risking information leakage [104]. However, SMPC requires continuous online availability and data transfer between parties. Moreover, as the complexity of algorithms increases, concerns regarding scalability and computational efficiency arise, especially for algorithms with a larger number of parameters or layers. However, cutting-edge neural network algorithms are currently being developed to address these concerns, particularly for implementation within SMPC frameworks.

Finally, participant authentication may be included in this category, although it is not strictly a privacy preservation method, as it often involves cryptographic processes to verify the identity of devices or users in a network [105]. Ensuring that every participant in the FL system is verified reduces the risk of intrusion by potential adversaries.

Resilience Against Attacks:

- MEA: Encrypting local updates through secure aggregation ensures that these updates remain incomprehensible even if intercepted. Only the aggregated update can be decipherable by the central server. In addition, HE facilitates computation on encrypted data, and update aggregations can occur without exposing data samples.
- DRA: Similarly to MEA, HE allows computations to be performed without decrypting the data, thereby safeguarding the reconstruction of the raw data. When using SMPC, the data is segmented into numerous parts, and computations occur on these fragments. The original data remain unexposed, as the computations do not require a full reconstruction.

4) *Hybrid Privacy Preservation Techniques:* In response to the growing need for privacy, researchers combine various preservation techniques, such as merging encryption with

DP, to develop hybrid methods that effectively address this concern.

In conclusion, advancing these privacy-preserving techniques becomes more paramount as IoT and FL technologies evolve and integrate more deeply into everyday life. This ensures that the technology serves its users without compromising their privacy and security. This holistic approach to privacy preservation is essential to maintain trust in IoT systems and the sustainable growth of FL applications across diverse sectors. The following section delves into how these privacy-preserving techniques are judiciously balanced with other strategies to achieve an optimal privacy/utility trade-off.

IV. ACCURACY AND LEARNING PERFORMANCE IN IOT FL

Achieving high model accuracy and efficiency in the context of FL-IoT presents significant challenges due to the diverse capabilities of devices and the variability in network quality. This section explores the complexities of assessing accuracy and performance within the FL-IoT ecosystem. We then introduce several techniques to improve accuracy and learning performance in FL-IoT environments. In addition, we examine solutions that aim to balance accuracy and privacy in FL-IoT, addressing the delicate interplay between securing data and optimizing model functionality. The section discusses these solutions, evaluating their effectiveness and implications in real-world scenarios.

A. Accuracy in FL on IoT

The accuracy of a federated model is a critical measure of its effectiveness, but several IoT-specific factors can impact this metric.

1) *Data Diversity:* The variety of data collected by different devices can enrich the learning process but also pose challenges in maintaining a consistent model accuracy across all nodes. Ensuring that the global model accurately represents this diversity is essential for its applicability across all devices and contexts [106].

2) *Non-IID Data:* IoT devices often generate data that needs to be identically and independently distributed. This non-IID nature of the data can lead to skewed model learning, where the model might perform well for some data distributions but poorly for others [107]. Effective strategies are needed to mitigate these effects and improve the robustness of the model.

3) *Client Participation Variability:* The intermittent connectivity and varying degrees of availability of IoT devices mean that not all nodes participate uniformly in the learning process. This variability can lead to inconsistencies in the model, affecting overall accuracy.

Enhancing model accuracy can also contribute to increased privacy protection. Employing L1 and L2 regularization helps mitigate overfitting, preventing the model from capturing and revealing specific details about the training data. By reducing overfitting, these fine-grained details are obscured, effectively thwarting MIA. Moreover, a model not overly tailored to the unique characteristics of the training data is less vulnerable to PIA, further enhancing privacy security.

B. Learning Performance

In the decentralized setting of FL, particularly in IoT, the efficiency of data processing and learning speed are paramount. These aspects are influenced by the methods for aggregating updates from devices and the strategies to accommodate the wide variability in device capabilities and data characteristics.

1) *Convergence Speed*: The speed with which a federated model reaches a stable solution is very important. Faster convergence reduces the need for communication between devices, reducing bandwidth and energy, critical factors for the IoT.

2) *Model Aggregation Techniques*: Aggregating model updates effectively in an environment with heterogeneous devices is a challenge. Techniques like FedAvg are popular, but often need to be adapted or improved to better handle the specific demands of IoT environments [108]. Efficient data aggregation techniques are crucial to handle the scale and speed of data. Employing strategies like edge processing to filter and preprocess data before it is used in federated models can reduce the load on the network and speed up the learning process.

3) *Personalization vs. Generability*: Tailoring models to specific groups of devices or even individual devices can significantly improve performance in particular applications. However, this approach must be balanced against the benefits of a generalized model that performs well across the entire network. Furthermore, models that generalize well across various datasets are inherently resistant to MIA, since they do not rely heavily on the specifics of the training data.

Finally, in IoT, devices vary widely in terms of computational abilities and types of collected data, contributing to diverse and sometimes sparse datasets across the network. This diversity necessitates sophisticated approaches to model training and aggregation to ensure that the learned models truly reflect the collective data and are not biased toward the characteristics of any subset of devices. Some of these techniques are described in the next section.

C. Techniques to Enhance Accuracy and Learning Performance

This section outlines techniques employed to improve FL's accuracy and learning performance within IoT environments.

1) *Personalized FL*: PFL focuses on customizing the global model to individual devices or users, enhancing relevance and efficiency. Various personalized FL strategies have been explored in recent research [109]. FedAMP [110] uses an attentive message passing mechanism, particularly effective for clients with similar data types, especially when dealing with non-IID data. FedRep [111] blends global representation learning with personalized local model heads. pFedMe [112] utilizes the Moreau envelope for loss function regularization. FedCAC [113] takes an aggressive approach by focusing on critical parameters to leverage data distribution similarities. PerFED-GAN [114] is a personalized FL method leveraging GANs. PerFED-GAN enables each client to independently design and train its model, without revealing its architecture or parameters, by sharing generated samples instead of model information.

2) *Transfer Learning FL*: Transfer learning can significantly accelerate the performance of FL models in IoT by applying knowledge from similar tasks to new contexts. This is particularly useful when new devices join the network or data distribution shifts, helping maintain high model accuracy and faster convergence.

Def-KT [115] integrates MKT into DFL to improve model performance by leveraging the distinct expertise of local models trained on diverse datasets. This approach helps each model gain insights from others' specialized knowledge, improving generalization capabilities and preventing catastrophic forgetting. [116] allows for sharing knowledge between labeled and unlabeled networks, enhancing their learning capabilities. The process involves using labeled data to train a model, which helps predict the outcomes for an unlabeled network. The solution aims to leverage the overlapped data samples and adjust the network parameters accordingly, improving the generalizability of the model in different datasets and enhancing its predictive accuracy. [117] revolves around the use of the TrAdaBoost algorithm [118], which leverages public data as a source domain to enhance training in target domains represented by client devices. This method optimizes model performance by dynamically adjusting the weights of instances based on their classification accuracy, which helps filter reusable instances and boost the model's overall performance.

3) *Federated Multi-Task Learning*: Expanding on personalized learning, federated multitask learning allows the model to simultaneously learn multiple tasks that may not be identical but share common features. MOCHA [119] introduces a way to extend collaborative learning of a shared prediction model to multi-task scenarios where different nodes may have different learning tasks. This technique can be very effective in IoT scenarios where devices perform different tasks but can benefit from the characteristics of the shared model. [120] incorporates non-federated Batch-Normalization layers into federated DNNs. This approach allows personalized model training on users' devices, enhancing user model accuracy and convergence speed. Furthermore, it explores the use of FedAvg-Adam (FedAvg combined with a distributed form of Adam optimization), showing that this combination can further improve the convergence speed of FL algorithms. AFL [121] used an iterative pruning network to create a sparse structure that allows efficient sharing of parameters between different tasks and devices. In addition, customized task mask layers facilitate training specialized subnets optimized for specific tasks, enhancing the model's performance on diverse tasks. Finally, an adaptive loss function dynamically adjusts the priority between tasks during training, helping to balance the training focus according to the needs of each task.

4) *Data Augmentation*: In FL, especially with limited data on each device, data augmentation can effectively increase the diversity and volume of training data. Techniques such as synthetic data generation or transformations that alter data while preserving its core characteristics can help enhance the robustness and performance of the model. Moreover, enhancing the dataset with additional synthetic or real instances can dilute unique properties, making it difficult for attackers to distinguish genuine dataset properties. Fed-ZDA [122] used

zero-shot data augmentation techniques to generate synthetic data that the model had never seen during training, which addresses data imbalances. FedM-UNE [123] incorporates a classic data augmentation technique, MixUp, adapted for federated settings without transferring raw data. Furthermore, to make this approach suitable for regression tasks, MixUp is modified to MixUp-BNE (bilateral neighborhood expansion), leading to another variant called FedM-BNE. In [124], each client independently optimizes its local model for its specific subset of data, and these models are then aggregated to enhance the performance of the global model. The augmentation strategy involves the clients enhancing their local datasets, particularly focusing on the minority class (anomalies) to balance the dataset. Techniques like random oversampling, SMOTE, ADASYN, and GAN generate synthetic data instances. This helps mitigate class imbalance and enriches the dataset without significantly changing the distribution of data features.

5) *Dynamic Client and Data Selection*: Optimizing which devices or “clients” are active in each round of model training can significantly enhance the learning process. FL can be more efficient and effective by dynamically selecting clients based on their data quality, availability, and relevance to the current model training phase. However, dynamic data sampling refers to selectively choosing which data samples to use for training during each iteration or phase of model training. This approach can optimize learning efficiency and model performance by focusing on the most informative or relevant data samples at different stages of the training process. It can be based on various criteria, such as sample difficulty, novelty, or representativeness of the overall data distribution.

AUCTION [125] incorporates a neural network that uses reinforcement learning to adapt its client selection policy based on real-time feedback and performance metrics. This network leverages an encoder-decoder architecture with an attention mechanism, allowing it to handle dynamic changes in the number of clients and to make informed selection decisions. FED GS [126] uses a gradient-based binary permutation algorithm to select and cluster factory devices into super nodes, which are more homogeneous and thus better suited for efficient FL training. A novel synchronization protocol helps coordinate training within and across these super nodes, with the aim of improving data security and robustness against data heterogeneity. The framework is specifically optimized for dynamic environments, such as those enabled by 5G technology. DSS-Edge-FL [127] incorporates the PAC learning theory to optimize the trade-offs between data size, model complexity, and accuracy. This theory helps to determine the necessary number of training samples to achieve the desired accuracy and reliability of the model under diverse data distributions.

6) *Learning Rate Adaptation*: Adjusting the learning rate dynamically during model training can help address the issue of non-IID data in federated settings. Adaptive learning rates can ensure faster convergence and better model performance by fine-tuning how quickly a model learns from diverse data sources across the network. FedLALR [128] addressed the inefficiency of standard FL techniques such as FedAvg and FedAdam, particularly in scenarios with heterogeneous data. Unlike these methods that use a uniform learning rate for all

clients, FedLALR allows each client to adjust its learning rate dynamically based on the squares of its local historical gradients and synchronized learning rates. This approach aims to enhance the optimization performance without compromising the convergence speed. The paper presents a theoretical analysis that supports the idea that FedLALR can achieve linear speedup with the number of clients, thus promising scalability. [129] tackled the challenges posed by device heterogeneity in FL systems, particularly the synchronization issues that arise when devices with different capabilities attempt to train a shared model simultaneously. The paper proposes an adaptive approach to batch size adjustment during model training to minimize idle waiting times on IoT devices. They establish a theoretical relationship between batch size and learning rate, deriving a scaling rule that helps to set the learning rate based on the batch size to maintain stability and convergence of the global model. Theoretical analysis determines the convergence rate and establishes an upper bound for convergence. Based on these insights, they developed an algorithm that dynamically adjusts batch sizes and learning rates between heterogeneous devices. The effectiveness of this approach is validated through extensive simulations and practical experiments on a testbed, showing promising results in reducing synchronization delays and enhancing energy efficiency.

However, despite these advances, these methods often depend on coarse-grained aggregation and direct gradient uploads for model updates. Although the raw data remain localized, the shared model updates might inadvertently contain sensitive information derived from the data. Sophisticated attacks could potentially reverse engineer these updates to extract or infer private data characteristics – consequently, the solutions discussed in the subsequent section aim to balance the trade-off between privacy and accuracy.

D. Trade-offs and Balancing Act

This section explores the critical balance between enhancing privacy protections and maintaining learning performance in FL environments. Specifically, when implementing privacy-preserving mechanisms such as DP, the information loss following data perturbation must be minimized to prevent degradation in data utility and, consequently, learning accuracy. Likewise, the accuracy levels achieved when using encrypted gradients or model parameters should be comparable to those obtained with original data to ensure practical applicability. It is also essential for many IoT applications that require rapid decision making to support near-real-time data processing. Therefore, the chosen privacy-preserving solutions must operate efficiently to prevent extended convergence times. This section presents various studies that have attempted to balance accuracy and privacy.

1) *Privacy Enhancements with Minimal Impact on Accuracy*: Zhang et al. [130] introduced a PPFL mechanism geared to DL within IoT healthcare systems. They adjusted the ElGamal encryption algorithm to achieve additive homomorphism, facilitating the aggregation of local models by encrypting a single variable for each client in each training round. Although homomorphic encryption traditionally introduces significant

computational overhead, adaptations such as ElGamal encryption allow efficient aggregation of encrypted data, striving to reduce the performance impact while maintaining strong privacy guarantees. A dropout-tolerant scheme was also proposed, which allowed the FL process to continue as long as the number of online clients met a preset threshold. This scheme utilizes the Diffie-Hellman key exchange and Shamir's secret sharing algorithm, allowing the server to reconstruct keys for dropped clients and continue training without compromising privacy. Their security evaluation affirmed the system's compliance with data privacy standards, and extensive experiments on real skin cancer datasets confirmed the scheme's efficacy in preserving privacy while maintaining communication and time-cost efficiency.

Wang et al. [131] proposed a privacy-enhancing method for disease diagnosis using FL. They used a VAE to reconstruct patient data, protecting against reconstruction attacks, and introduced differential privacy by adding Laplace noise to the reconstructed data. An incentive mechanism encouraged active participation in local model training, with a trusted third-party evaluation server assessing training quality and calculating rewards. Experiments on the MIT and BIH datasets demonstrated the method's ability to balance accuracy and privacy, albeit at the cost of increased computational overhead.

Zhou et al. [132] developed a privacy-preserving FL approach for fog computing environments. Each fog node collects data from IoT devices and performs learning tasks, significantly improving training efficiency. Their approach combines homomorphic DP, blinding, and Paillier encryption to protect data and model privacy against honest but serious servers and collusion attacks among untrustworthy nodes. The challenge with DP is to determine the amount of noise to add. Too much noise can degrade the model's accuracy, while too little noise may not offer enough privacy. The researchers experimented with different noise levels to find a balance, ensuring that the model's utility was not significantly compromised.

2) *Handling Non-IID Data and Model Generalization Across Heterogeneous Data*: Zhao et al. [133] explored the use of blockchain and DP in smart home applications. They replaced the central server with blockchain to generate the global FL model, improving security against malicious model updates from participants. DP was incorporated by adding Laplacian noise to normalized features, and a batch-normalization-based technique was proposed to improve accuracy. Simulations on the MNIST dataset and a Raspberry Pi proof-of-concept implementation highlighted the solution's effectiveness.

Stephanie et al. [134] combined blockchain technology with SMPC to ensure encrypted inference and model verification, preserving high accuracy in the FL models in the set for IoT health. Hospitals could develop unique model architectures, employing ensemble weights to enhance model generalization. The accuracy of this approach was validated against traditional FL methods using medical image datasets, with additional measurements for the time taken for ensemble weight tuning, encrypted inference, and blockchain contract execution.

3) *Optimizing Computational Efficiency and Real-time Processing*: Li et al. [135] addressed privacy and trust issues in cross-silo FL for IoT. They proposed a blockchain-assisted

privacy-enhanced FL protocol that uses homomorphic encryption to protect quantized gradients based on local sign and a smart contract for secure self-aggregation. This protocol offers privacy and public verifiability, effectively resisting gradient inversion attacks without compromising learning accuracy. Performance evaluation involved time consumption and communication overhead assessments for each participating client.

4) *Discussion*: The reviewed studies highlight the significant challenges and innovative solutions in balancing privacy and performance in FL. Although techniques such as homomorphic encryption, DP, and blockchain offer robust privacy guarantees, they also introduce computational overhead and potential accuracy loss [136]. Achieving an optimal balance between privacy protection and learning performance requires careful consideration and adaptation of these techniques.

Privacy-preserving techniques, such as homomorphic encryption, provide robust security by allowing encrypted model aggregation, though they introduce computational overhead that can limit their applicability in real-time scenarios. Alternatively, DP methods that add noise to the data protect against data inference and reconstruction attacks but may reduce model accuracy, particularly in applications where high data utility is critical, such as healthcare diagnostics. Hybrid approaches, which combine DP with encryption and local aggregation at edge or fog nodes, address the challenges of non-IID data and enhance scalability by reducing dependency on a central server. However, these approaches often require careful tuning of privacy parameters, as excessive noise can degrade model performance. In decentralized architectures such as blockchain combined with SMPC, privacy and security are enhanced as data are processed and validated across distributed nodes, preventing single points of failure. However, these solutions can also add computational demands and latency, particularly in complex ensemble models or large-scale networks. Blockchain-based protocols with homomorphic encryption address privacy and trust by securely aggregating model updates without a central server. However, validation steps can slow response times in time-sensitive applications.

In summary, while HE maximizes privacy at computational cost, DP offers flexible privacy with some accuracy trade-offs, and decentralized methods improve trust and scalability but may impact responsiveness. For optimal adaptability, a dynamic approach that adjusts privacy parameters based on network conditions and application requirements may provide the most effective balance between privacy, accuracy, and efficiency in IoT-based FL [137]. Table IV categorizes the discussed PPFL solutions to balance accuracy and privacy, highlighting their goals, privacy-preserving techniques, FL optimizations, and target applications.

V. NETWORK EFFICIENCY IN IoT FL

Implementing FL in IoT environments presents unique challenges, mainly due to the limited resources of IoT devices. This section explores key issues such as energy constraints, computational and communication overhead, memory limitations, device heterogeneity, and scalability challenges. Addressing these issues is crucial for optimizing FL deployment and maintaining network efficiency.

TABLE IV: Summary of PPFL solutions proposed for accuracy/privacy balance.

Ref.	Goal	Privacy-Preserving Technique	FL Optimization	Target Application
[134]	Model generalization & efficiency	SMPC, Blockchain	Model heterogeneity, weighted ensemble DL	Healthcare
[135]	Efficiency and real-time performance	ElGamal encryption, Blockchain (smart contracts)	Decentralized architecture	Generic IoT System
[130]	Privacy enhancements with minimal impact on accuracy	Masks, homomorphic encryption	Quality-driven contribution, dropout-tolerable scheme	Healthcare
[133]	Handling data heterogeneity	DP	New normalization technique	Smart home system
[131]	Privacy enhancements with minimal impact on accuracy	Data reconstruction via VAE, DP via Laplace Noise	Incentive mechanism for FL participation	Disease diagnosis
[132]	Privacy enhancements with minimal impact on accuracy	DP, blinding, and Paillier homomorphic encryption	Fog nodes for data aggregation	Smart homes and Smart healthcare

A. Resource Management

1) *Energy Constraints*: IoT devices, particularly those that are battery-powered, face significant energy constraints. The iterative nature of FL, compounded by the energy demands of privacy-preserving methods, can rapidly deplete the batteries of devices. Strategies to improve energy efficiency in FL processes are critical and include optimizing model communication and computation techniques to prolong the life of the device [10], [39].

2) *Communication Overhead*: Limited bandwidth and variable network conditions pose significant challenges for FL in IoT. Strategies to mitigate these issues include optimizing data transmission protocols, reducing the size of model updates, and utilizing techniques to minimize the number of communication rounds without compromising the model's accuracy.

3) *Computational Overhead*: The disparity in computational power among IoT devices requires adaptive approaches in FL. Employing model quantization and simpler yet effective learning models can help reduce computational strain on less capable devices, facilitating broader participation in the FL process.

4) *Memory Footprint*: IoT devices often have limited memory, making it impossible to store large models or datasets. Techniques such as model pruning and federated transfer learning are essential for managing memory resources effectively while maintaining satisfactory model performance.

B. Scalability and Device Heterogeneity

1) *Device Heterogeneity*: The diverse range of device capabilities within IoT networks introduces complexity in maintaining a uniform FL process. Effectively addressing this heterogeneity is crucial for the scalability of FL systems, ensuring that they can handle large numbers of devices without performance degradation.

2) *Scalability*: As IoT networks can scale up to millions of devices, the FL framework must accommodate such large-scale operations. This includes efficiently aggregating updates from numerous devices and ensuring that the system remains responsive even with the continuous addition of new nodes.

Although FL offers a promising framework for learning across distributed IoT devices, it introduces several efficiency challenges. Addressing these effectively requires advanced technological solutions and strategic planning, ensuring that

network efficiency is maintained without compromising learning outcomes. The next section delves into some proposed techniques to overcome these challenges.

C. Techniques to Enhance Network Efficiency in IoT FL

Several techniques have been developed to enhance network efficiency, including model compression, dynamic client selection, and edge computing integration, to address resource limitations and heterogeneous attributes of IoT devices.

1) *Resource-Efficient and Lightweight Learning Algorithms*: Lightweight learning algorithms are essential for minimizing the resource consumption of IoT devices during FL. These algorithms require minimal computational resources, making them suitable for deployment on IoT devices. Model compression techniques such as pruning and quantization help manage the resource constraints of IoT devices by reducing the computational demand. Pruning eliminates minimal weights from the model, reducing its size, and improving operational efficiency. This technique helps to make the model lightweight and easier to deploy on resource-constrained devices. Quantization reduces the precision of the numerical values in the model, decreasing the size of the model and the computational requirements. This method allows the model to run efficiently on devices with limited computational power. GWEP [138] is a model compression-based FL method that combines quantization and model pruning to reduce FL's computational, memory, and network demands of FL, making it feasible for low-end IoT devices. Theoretical guarantees of FL convergence are provided and empirical evaluations show that GWEP outperforms baseline algorithms, achieving up to 10.23 times faster performance with 11 times fewer communication rounds while enhancing model compression and energy efficiency. Zhang et al. [139] employed local adaptive optimizers such as Adam and a cross-round learning rate scheduler. This approach improves FL performance by identifying diverse attack types without significant overheads. This method improves network efficiency by optimizing learning rates and reducing unnecessary updates.

2) *Dynamic Client Selection*: Dynamic client selection involves selecting clients based on their current status and capabilities to optimize the trade-off between learning performance and resource consumption. FedMCCS [140] dynamically selects clients for training to balance resource usage and model performance. Considering factors such as data quality,

computational power, and connectivity, this method ensures that the most suitable clients participate in each training round.

3) *Hierarchical FL*: Hierarchical FL structures can reduce communication overhead by aggregating local models at intermediate nodes before sending them to the central server. This hierarchical approach reduces the volume of data transmitted over long distances and can help to manage network bandwidth more effectively. In [141], the authors propose a hierarchical FL architecture optimized for user assignment and resource allocation in heterogeneous IoT systems. This architecture aims to enhance model training efficiency and reduce communication overhead while dealing with non-uniformly distributed data across users. Using two real-world datasets, the proposed system demonstrates a 4 – 6% increase in classification accuracy compared to existing hierarchical FL methods. It achieves a 75 – 85% reduction in communication rounds between the edge nodes and the central server.

4) *Edge Computing Integration*: Integrating edge computing with FL enhances data processing capabilities and reduces latency by leveraging the proximity of edge nodes to IoT devices [142]. Edge nodes preprocess data locally, minimizing the amount of data that must be transmitted to FL and improving overall efficiency. Furthermore, edge computing allows for local decision making, which is crucial for applications requiring immediate responses, such as autonomous driving and real-time monitoring systems.

5) *Energy-efficient Communication Protocols*: Energy-efficient communication protocols reduce energy consumption during data transmission by minimizing the frequency of communication rounds and compressing transmitted data. The study in [143] focuses on optimizing system energy efficiency through joint communication and optimization of computation resources. Specifically, it considers two transmission protocols for uploading model parameters: NOMA and TDMA. The objective is to minimize the total energy consumption of the edge devices during a finite training period while maintaining a specified training accuracy. Optimization involves adjusting transmission power, uploading model parameters, and CPU frequencies for local updates. The authors propose algorithms that significantly enhance energy efficiency by balancing the energy trade-offs between communication and computation using convex optimization techniques. The numerical results indicate that this approach outperforms the benchmark schemes, substantially improving the energy efficiency of the federated edge learning systems.

Enhancing network efficiency in FL IoT involves addressing various challenges, such as resource constraints, network connectivity, device heterogeneity, and data processing. Using techniques such as model compression and adaptive algorithms, dynamic client selection, and edge computing integration, FL can be effectively deployed in IoT environments, ensuring robust and efficient learning while preserving device resources and maintaining high performance. These challenges became more difficult when privacy preservation techniques were included. The next section discusses solutions that try to balance these two aspects.

D. Trade-offs and Balancing Act

Various methods have been proposed to improve privacy while mitigating the challenges imposed by the limitations of IoT resources and the heterogeneity of devices without significantly compromising learning performance. This section categorizes and summarizes notable approaches based on the techniques used to manage this trade-off and their impact on increasing network efficiency.

1) *Model Compression and Lightweight Algorithms*: Model compression techniques such as pruning and quantization help manage resource constraints of IoT devices. These techniques reduce the size of the model and the computational demand, enabling efficient deployment on resource-constrained devices while maintaining privacy and accuracy. Fu et al. [144] proposed verifiable FL using Lagrange interpolation to verify the correctness of aggregated gradients and blinding technology to protect data privacy. This method achieves smaller attack probabilities and constant verification overhead regardless of the number of participants. Blinding technology reduces the amount of data that need to be encrypted, thereby minimizing communication overhead, which is beneficial for network efficiency. Privacy is preserved by ensuring that individual gradient updates remain confidential through blinding, making it difficult for adversaries to infer private data.

Kong et al. [145] introduced FedLoc, which utilizes a homomorphic threshold cryptosystem for key configurations and updates, and the bounded Laplace mechanism for securing model hyperparameters. This method allows for encrypted computations, reducing the need to transmit raw data. This technique reduces the communication load by only requiring encrypted updates, thus enhancing network efficiency. Privacy is maintained through homomorphic encryption, which allows computations on encrypted data without exposing the raw data, and DP, which adds noise to the data.

2) *Dynamic Client Selection and Two-Phase Protocols*: Dynamic client selection and two-phase protocols optimize resource usage by involving the most suitable clients in each training round. These techniques reduce communication costs and improve scalability, making the FL process more efficient. Kanagavelu et al. [146] developed a two-phase MPC protocol to reduce communication overhead in large-scale IoT networks. This method enhances scalability and reduces execution time without sacrificing model accuracy. By electing a subset of FL participants as the model aggregation committee, communication costs are significantly reduced, thus improving network efficiency. Privacy is preserved by secret sharing, where data are split into shares and distributed among multiple participants, ensuring that no single participant can access complete data.

Chen et al. [147] proposed DWFL, an algorithm that uses the superposition property of wireless channels to parallelize communication while incorporating differential privacy using Gaussian noise. This method ensures privacy protection and effective communication in a decentralized topology, which reduces the communication rounds needed for aggregation, enhancing network efficiency. Privacy is protected by adding Gaussian noise to the data, ensuring that individual data points cannot be easily re-identified.

3) *Edge Computing and Hierarchical Structures*: Integrating edge computing with FL leverages the proximity of edge nodes to IoT devices, enhancing data processing capabilities and reducing latency. This integration allows for local decision making and preprocessing of data, minimizing the amount of data transmitted, and improving network efficiency. Yin et al. [148] addressed potential breaches at data and content levels during the FL training phase. They introduced a privacy-preserving method that combines multi-input function encryption (MIFE) and Bayesian differential privacy, allowing users to adjust privacy budgets. Sparse Differential Gradient Enhancement minimizes communication encryption overhead by only transmitting substantial gradient changes, thus improving network efficiency. Privacy is maintained by using MIFE to hide data functions and Bayesian differential privacy to add noise dynamically based on data sensitivity.

Zhang et al. [139] focused on detecting data anomalies in IoT environments using the FedIoT platform and the FedDetect algorithm, which employs local adaptive optimizers like Adam and a multi-round learning rate scheduler. This approach improves FL performance by identifying diverse attack types without significant overheads. This method improves network efficiency by optimizing learning rates and reducing unnecessary updates. Privacy is preserved by optimizing local models without sharing raw data, ensuring that data stays within the local environment.

4) *Blockchain and Cryptographic Techniques*: Blockchain and cryptographic techniques improve security and privacy in FL by providing robust verification and encryption mechanisms. These methods ensure data integrity and reduce the frequency of communication, improving overall network efficiency. Qu et al. [149] proposed FL-Block, which integrates blockchain to verify updates to the local model through the Proof-of-Work consensus process, enhancing privacy and resilience to poisoning attacks. This method avoids single-point failures and ensures robust model updates in a distributed manner, reducing communication frequency, and improving network efficiency.

5) *Discussion*: The solutions presented in this section introduce various innovative techniques to address the critical trade-offs between privacy and network performance in FL for IoT environments. Each method employs a unique strategy to address computational and communication challenges, but their effectiveness varies depending on the specific context and constraints of IoT networks.

Model compression and lightweight algorithms emphasize the reduction of computational burden by minimizing model size and complexity through pruning, quantization, and homomorphic encryption techniques. These methods are advantageous in highly resource-constrained environments, as they reduce computational costs and network load. However, while these techniques are generally effective in improving efficiency, there is often a trade-off in model accuracy or privacy, as aggressive compression may reduce the robustness of the model, especially in highly dynamic IoT environments. Homomorphic encryption, for example, preserves data privacy but introduces latency due to the complexity of encrypted computations, which may not be ideal for latency-sensitive

applications.

Dynamic client selection and two-phase protocols optimize resource usage by selecting clients that are better suited for each training round. These methods significantly reduce communication costs by focusing only on the most relevant clients, which improves scalability in large networks. However, while this approach enhances scalability and efficiency, it may compromise the diversity of data sources, potentially affecting the model's generalizability. Additionally, while secret sharing and differential privacy ensure data protection, the noise added in differential privacy can diminish the accuracy of the aggregated model, especially in smaller IoT networks where data is sparse.

Edge computing and hierarchical structures leverage edge nodes to process data closer to the source, reducing latency and network load. This hierarchical approach is particularly effective in environments with varied device capabilities, enabling local data processing and decision-making. However, edge computing requires robust synchronization and resource management, as data must be securely transmitted between layers. Reliance on edge nodes can also create points of vulnerability, where failures or attacks at the edge level could disrupt the function of the broader network.

Blockchain and cryptographic techniques add an additional layer of privacy and verification. Blockchain provides decentralized verification, improving resilience against attacks. However, the Proof-of-Work consensus mechanism may introduce latency and require substantial computational resources, making it less suitable for low-power IoT devices. Although blockchain benefits data integrity, its high computational cost makes it more practical for applications prioritizing data security over real-time responsiveness.

In summary, integrating these techniques into FL frameworks addresses the inherent challenges of IoT environments. It ensures a balanced trade-off between privacy and network performance, paving the way for more robust and efficient FL deployments. Table V provides a comprehensive overview of the PPFL solutions discussed in this section. For every solution, the table highlights the specific techniques used to ensure privacy, the techniques implemented to optimize FL performance and efficient utilization of IoT resources, and the primary application for which the solution was designed. This structured summary compares how different approaches manage the trade-off between privacy and learning performance while improving network efficiency in various IoT scenarios.

VI. EVALUATION OF EXISTING PPFL TECHNIQUES AND ASSOCIATED TRADE-OFFS

Combining FL with emerging technologies enabled a move from centralized data processing to more distributed and privacy-focused systems. However, this shift addresses growing concerns about data privacy, security, and efficiency, which will be discussed below.

A. Key Evaluation Aspects

In evaluating the effectiveness of PPFL solutions for IoT applications, three primary aspects are considered, as highlighted in the taxonomy presented in Figure 6:

TABLE V: Summary of PPFL solutions proposed for IoT applications.

Ref.	Performance Optimization Technique	Privacy Technique	Network Efficiency Improvement	Target Application
[144]	Model Compression (Pruning, Quantization)	Blinding Technology	Reduces encryption communication overhead	Industrial
[146]	Dynamic Client Selection, Two-Phase MPC	Secret Sharing	Reduces communication costs, enhances scalability	Smart manufacturing
[147]	Communication Parallelization	Differential Privacy (Gaussian Noise)	Reduces communication rounds needed	Generic IoT System
[148]	Sparse Differential Gradient Enhancement	Function Hiding MIFE, Bayesian Differential Privacy	Minimizes encryption overhead, reduces communication	Generic IoT System
[145]	Homomorphic Encryption	Homomorphic Threshold Cryptosystem, Bounded Laplace DP	Encrypted computations reduce data transmission	Vehicular navigation
[149]	Blockchain and Cryptographic Techniques	Blockchain with Proof-of-Work Consensus	Reduces communication frequency, enhances robustness	Generic IoT System
[139]	Edge Computing, Local Adaptive Optimizers	Local Adaptive Optimizers, Differential Privacy	Optimizes learning rates, reduces unnecessary updates	Anomaly detection in IoT

a) *Privacy*: It is a core concern in PPFL solutions, particularly in regard to resilience to various privacy attacks in FL. Solutions are evaluated based on their ability to defend against potential vulnerabilities, such as MIA, MEA, PIA, and DRA. High resilience in this category ensures that data privacy is maintained without compromising model performance, even when malicious entities attempt to exploit data during or after training.

b) *Network Efficiency*: It encompasses optimizing communication, computation, memory, and energy costs associated with IoT devices that participate in FL. Effective PPFL solutions must minimize these costs to ensure the feasibility of implementation in resource-constrained IoT systems:

- **Communication Overhead**: Evaluation considers how solutions manage the data transmitted between devices and the central server. Techniques that reduce communication frequency and volume help maintain bandwidth and improve speed.
- **Computation Overhead**: Solutions are assessed for their computational efficiency to ensure that they can operate on devices with limited processing capabilities without significantly affecting response times or energy consumption.
- **Memory Cost**: Evaluating memory usage is crucial, as many IoT devices have limited storage. Solutions that optimize memory usage allow for smoother, more efficient FL processes.
- **Energy Cost**: This sub-metric considers the energy required for devices to participate in FL. Solutions with lower energy consumption are preferable for prolonged deployment in battery-operated IoT devices.

c) *Learning Performance*: To assess the learning performance of PPFL, the model accuracy and convergence time are important metrics:

- **Accuracy**: Solutions are evaluated based on the accuracy achieved by the federated model after training. A high degree of accuracy indicates that the model learns effectively from distributed data sources.

- **Convergence Time**: This aspect examines how quickly the model reaches its optimal accuracy. Faster convergence time is desirable as it implies reduced computational and communication costs and improved efficiency.

B. Analysis of Privacy-Preserving Techniques

a) *Homomorphic Encryption*: It is highly effective against privacy attacks, enabling computations on encrypted data without exposure. However, it introduces significant computational and communication overhead as a result of the complexity of the encryption and decryption processes, potentially leading to longer model convergence times. Examples of its application include [130], [132], and [145].

b) *Differential Privacy (DP)*: It provides robust protection by adding noise to the data, ensuring that individual data points remain anonymous. Although it effectively protects privacy, added noise can degrade model accuracy and increase computational overhead. The iterative addition of noise to the FL process can slow the convergence of the model and increase communication costs. This technique has been explored by Zhao et al. [133], Wang et al. [131], and Chen et al. [147].

c) *Blockchain and Cryptographic Techniques*: They ensure data integrity, transparency, and robust security through consensus mechanisms. However, blockchain transactions are time-consuming and the advanced cryptographic operations required add to the computational overhead, potentially increasing communication costs. These methods are highlighted in the studies by Qu et al. [149] and Stephanie et al. [134].

d) *Blinding and Lagrange Interpolation*: They offer a combination of lightweight and more complex techniques to protect data. Blinding provides efficient protection for encrypted gradients, while Lagrange interpolation verifies the integrity of aggregated gradients. Although blinding has minimal computational impact, Lagrange interpolation can introduce significant overhead, particularly in higher-dimensional data. The effectiveness of these combined techniques can vary based on the specific dataset. Fu et al. [144] provides an example of this approach.

e) Specialized Techniques: Function Hiding MIFE, Bayesian DP, and BGV encryption offer tailored privacy protections and robustness, often with high computational complexities affecting model convergence time and increasing communication overhead. These specialized techniques are discussed in the work of Yin et al. [148] and Meng et al. [150].

C. Analysis of Techniques for Enhancing FL Performance and Efficiency

a) Model Compression (Pruning and Quantization):

Techniques such as pruning and quantization reduce model size and computational demand, making it feasible for resources-constrained devices to participate in the FL process. This optimization improves network efficiency and maintains privacy and accuracy. Fu et al. [144] explored these techniques, demonstrating their effectiveness in industrial applications.

b) Dynamic Client Selection and Two-Phase Protocols:

These techniques optimize resource usage by involving the most suitable clients in each training round, reducing communication costs, and improving scalability. Kanagavelu et al. [146] and Chen et al. [147] implemented such methods to improve network efficiency while preserving privacy through secret sharing and differential privacy mechanisms.

c) Edge Computing and Local Adaptive Optimizers:

Integrating edge computing with FL leverages the proximity of edge nodes to IoT devices, enhancing data processing capabilities and reducing latency. This integration allows for local decision making and preprocessing of data, minimizing the amount of data transmitted, and improving network efficiency. Zhang et al. [139] and Yin et al. [148] used these techniques to optimize learning rates and reduce unnecessary updates, ensuring that the data remain within the local environment.

d) Crowdsourcing Aggregation and Quality-Driven Contribution: These methods reduce bandwidth usage and improve accuracy by focusing on data quality. Zhang et al. [151] introduced a secure crowdsourcing aggregation algorithm, which prioritizes quality over quantity, although it may face challenges in scalability.

D. Qualitative Evaluation of PPFL Solutions

Table VI summarizes the strengths and weaknesses of existing privacy preservation techniques, FL performance metrics, and IoT network optimization techniques used in PPFL solutions, highlighting the specific advantages and limitations of each technique in these critical areas. Table VII provides a structured evaluation of PPFL solutions applied to IoT applications, assessing them in terms of resilience to privacy attacks, FL performance metrics such as accuracy and convergence time, and network efficiency metrics that include communication, computation, and memory costs. In this table, a check mark (✓) in cell (i, j) indicates that the solution i addresses the evaluation aspect j as categorized by the taxonomy in Fig. 6, which organizes the aspects into three key areas: i) privacy (resilience to attacks in FL), ii) network efficiency (communication overhead, computation overhead, memory cost and energy cost), and iii) learning performance

(accuracy and convergence time). In contrast, a cross mark (✗) signifies that the aspect j is not addressed in solution i , while a dashed (—) suggests that the efficacy of this aspect may vary depending on the specific data or the application context.

Achieving the right balance between robust privacy protection and model performance is crucial for effective FL deployments in IoT environments. Enhanced privacy often introduces complexities that affect model learning accuracy and extend the convergence time. Techniques that improve data privacy sometimes compromise the speed of communication and computation in the FL network. This trade-off between security and efficiency is a persistent challenge, especially in IoT scenarios where device resources and network bandwidth are limited. Future research should optimize these privacy-preserving methods to ensure computational efficiency and high learning accuracy, making FL a viable option for real-time IoT applications. Continuous evaluation and refinement of these methods will be essential to address the evolving challenges of privacy and performance in FL.

E. Quantitative Evaluation of PPFL Techniques

To enhance the utility of the paper as a practical guide, we incorporated an experimental evaluation of the discussed PPFL approaches.

The existing PPFL were evaluated separately in platforms, making it difficult to numerically compare them. To tackle this, we built a small-scale FL architecture and used a real-time dataset, “IDSIoT2024”. This is an open source that can be accessed from IEEE DataPort [152]. The detailed class distribution of the dataset is presented in Table. VIII, while the details of the experimental setup are presented in Table IX.

1) Experimental Setup: We implemented a lightweight FL setup in the Google Colab Pro Platform, a cloud-based platform that provides access to powerful computational resources. This platform offers a NVIDIA T4 GPU, a high-performance computing unit renowned for its Tensor Cores, which significantly accelerates ML/DL operations. We incorporated the LightGBM technique [153] for local training in FL. LightGBM is a highly efficient and scalable gradient boosting framework for fast and accurate training on tabular datasets. It is particularly well-suited for handling large-scale datasets and high-dimensional features. LightGBM optimizes training speed and model performance, reducing memory usage and speeding up split finding. Additionally, it supports leaf-wise tree growth, which improves accuracy for complex datasets and includes built-in handling for categorical features, further simplifying pre-processing. In an FL framework, LightGBM is an excellent choice for local training due to its computational efficiency and resource-efficient design. This makes it ideal for edge devices or distributed environments with limited processing power or memory. Its ability to work well with tabular data aligns perfectly with many real-world applications of FL, such as healthcare records analysis and IoT data processing. In experiments, we kept the basic architecture of LightGBM and FL constant for a better comparison. The results of experiments are presented in Table X.

TABLE VI: Strengths and weaknesses of the proposed privacy-preserving, FL performance, and IoT network optimization techniques used in the PPFL solutions.

	Technique	Paper	Advantages	Weaknesses
Privacy Preservation	BGV HE	[150]	Simplifies computations and enlarges the plaintext space	Requires careful parameter selection
	A-LWE	[150]	Enables PP during aggregation without additional interaction	Requires significantly larger keys and ciphertexts
	SMPC	[134]	Supports computation on encrypted data	Computationally intensive
	ElGamal encryption	[135]	Ensures the privacy of quantized gradients	Not ideal for long messages, requires modular exponentiation
	Masks	[130]	Simple, versatile, low computational requirements	Not always robust and end-to-end secure
	Homomorphic Encryption	[130], [132], [145]	Enables computations on encrypted data with strong security	Significant computational overhead, increases storage and bandwidth use
	Blockchain	[134], [135]	Ensures data integrity, transparency, traceability	Scalability and storage issues
	Lagrange interpolation	[144]	Verifies integrity of aggregated gradients	Computational complexity
	Blinding technology	[132], [144]	Protects encrypted gradients from inversion	Risk of collusion attacks
	Function Hiding MIFE	[148]	Enables decryption of only aggregated results	High computational overhead, dependence on a TTP, scalability issues
	Bayesian DP	[148]	Incorporates prior knowledge for flexible privacy guarantees	Dependence on prior knowledge, potential over-reliance
	Data reconstruction via VAE	[131]	Provides a level of abstraction for data	High complexity, reconstruction accuracy challenges
	DP (Laplace/Gaussian noise)	[131]–[133], [145], [147], [150]	Ensures individual data points cannot be singled out	Trade-off between privacy and utility, choosing the right amount of noise
FL Performance	Crowdsourcing Aggregation	[151]	Reduces bandwidth usage, maintains accuracy	Scalability concerns
	Model heterogeneity	[134]	Allows different model structures for participants	Coordination challenges
	Weighted ensemble DL	[134]	High-quality models have more influence	Risk of oversimplification
	Quality-driven contribution	[130]	Prioritizes data quality over quantity	Ambiguity in quality assessment
	Dropout-tolerable scheme	[130]	Robust to client dropouts	Assumes sufficient online clients
	Normalization technique	[133]	Improves model generalization across heterogeneous data	-
	Local adaptive optimizer	[139]	Enhances model updates from each device	Risk of overfitting, increased memory overhead
	Cross-round learning rate scheduler	[139]	Facilitates quick convergence	Additional complexity
	Incentive mechanism for FL participation	[131]	Encourages data sharing and participation	Fairness in reward distribution, scalability concerns
	Decentralized architecture	[134], [135], [147]	Eliminates single-point failures, speeds up training	Risks from malicious or faulty nodes
IoT Resources	Fog nodes for data aggregation	[132]	Enhances data distribution efficiency	Risks of fog node failure or compromise
	Two-phase MPC	[146]	Minimizes direct communications	Committee bias and trust issues, election overhead
	Sparse Differential Gradient	[148]	Enhances communication/storage efficiency, focuses on important features	Risk of information loss, potential over-sparsity
	Lightweight communication	[139]	Efficient for IoT messaging	Limited security
	Modular design	[139]	Minimizes unnecessary processes	Integration challenges
	Direct implementation on edge devices	[139]	Reduces data travel	Limited resources
	Communication parallelization	[147]	Reduces communication rounds, uses over-the-air computation	Vulnerable to noise and interference

2) *Discussion on Quantitative Results:* The experimental evaluation of various PPFL techniques provides valuable insight into their trade-offs across multiple critical parameters, including accuracy, computational efficiency, convergence time, and memory footprint. DP techniques such as Gaussian, Laplace, and Poisson noise methods demonstrated a strong balance between accuracy and resource demands. Laplace noise achieved the highest accuracy at 99.37%, closely followed by the Gaussian and Poisson noise distributions, which scored 94.18% and 91.41%, respectively. These methods introduced moderate computational overhead, with average client

training times around 0.6 seconds and convergence times near 68 seconds. Their memory footprints were consistent at approximately 2MB, reflecting their efficiency and suitability for scenarios where computational and memory resources are restricted, but privacy remains essential.

SMPC methods prioritized robust security through techniques such as Shamir’s secret sharing, additive secret sharing, and clumping circuits. However, this came at the expense of lower accuracy, ranging from 63.58% to 69.36%, and significantly higher convergence times, exceeding 240 seconds in some cases. The memory requirements for the SMPC tech-

TABLE VII: Evaluation of PPFL solutions proposed for IoT applications.

PPFL	Privacy Attacks				FL Performances		Efficiency Metrics		
	MIA	MEA	PIA	DRA	Accuracy	Convergence Time	Communication	Computation	Memory
[134]	High	Mod. to High	Moderate	High	✓	✓	✗	✗	✗
[135]	High	Moderate	Moderate	High	-	-	✗	✗	✗
[130]	High	High	Mod. to High	High	✓	✓	✗	✗	✗
[144]	High	Moderate	Moderate	High	✓	✗	✓	✓	✗
[146]	High	Moderate	Mod. to High	High	✓	✓	✓	✗	✗
[133]	High	Moderate	Mod. to High	High	✗	✓	✗	✗	✗
[131]	High	Moderate	Mod. to High	High	✓	✓	✗	✗	✗
[147]	High	Moderate	Mod. to High	High	-	-	✓	✓	✓
[148]	High	High	High	High	✓	✗	✓	✓	✗
[150]	High	High	Mod. to High	High	✓	✗	✓	✓	✗
[132]	High	High	Mod. to High	High	✓	✓	✗	✓	✗
[145]	High	High	M. to High	High	✗	✗	✗	✓	✗
[151]	High	High	Moderate	High	✓	✗	✓	✓	✗
[139]	Moderate	Low	Low	Moderate	✓	✓	✗	✗	✓
[149]	High	Moderate	Moderate	High	✓	✓	✗	✗	✗

TABLE VIII: Class distribution of IDSIoT2024 dataset.

Class Name	Count
DoS	40000
Injection	1280
Malware	10000
MITM	5000
Normal	10000
Routing	10000
Vulnerability_Analysis	20000

TABLE IX: Specifications of experimentation platform.

Parameter/Specification	Value/Description
Platform	Google Colab Pro
Hardware Accelerator	NVIDIA T4
Runtime Environment	Python (3.10.12)
Framework	TensorFlow (2.17.0)
Number of Clients	10
Number of Rounds	10
Data Distribution	IID
Local Model	LightGBM
loss	multi_logloss
Learning Rate	0.10
num_leaves	31
max_depth	-1
verbosity	-1
num_boost_round	50

niques were slightly higher than those for differential privacy, with footprints ranging from 2.11 MB to 2.36 MB. Despite their reduced predictive performance, SMPC methods excel in applications demanding the highest levels of data security, where ensuring that the data of no individual participant are exposed outweighs the computational costs.

HE has emerged as a promising approach to balance accuracy, efficiency, and resource demands. PHE delivered high accuracy at 99.15% with the smallest memory footprint of all evaluated methods at only 0.93 MB. FHE slightly improved accuracy to 99.20% but required more memory at 1.18 MB and exhibited marginally slower convergence times. Secure aggregation techniques, including random masking and partial sum masking, also achieved high accuracy (99.10%–99.13%),

with faster convergence times (55–56 seconds) and lightweight memory usage of around 1.15 MB, making them highly practical for scenarios requiring both security and efficiency.

Anonymization techniques further demonstrated various strengths, with RCI achieving the highest overall accuracy at 99.41%, although with a higher memory footprint of 2.77 MB and slower convergence times. Rotating IDs offered a more resource efficient alternative, achieving a memory footprint of just 0.85 MB while maintaining reasonable accuracy (95.32%) and the fastest convergence time among all methods (49 seconds). These findings highlight the intricate balance between privacy, performance, and resource efficiency in PPFL methodologies, highlighting the need to tailor the choice of approach to the specific requirements of the application. Whether the priority lies in achieving top-tier accuracy, minimizing memory usage, or ensuring the strongest privacy guarantees, these experimental outcomes provide a comprehensive guide for informed decision making in real-time FL deployments.

VII. TECHNOLOGICAL INNOVATIONS AND FUTURE DIRECTIONS

A. Emerging Technologies

Rapid advancements in emerging technologies are expected to revolutionize the use of FL in IoT systems, improving both their efficiency and privacy. The following discusses how these technologies, particularly the evolving 6G networks, can support and optimize privacy-preserved FL in IoT:

1) *Enhanced Bandwidth and Lower Latency*: 5G and 6G networks promise significantly higher bandwidth and lower latency than previous generations. This improvement is crucial for FL, as it involves frequent exchanges of model updates across a network of distributed IoT devices. With 6G, ultra-low latency and near instantaneous data transmission will reduce the risk of delays that could expose model updates to adversarial attacks during transmission. Moreover, a higher bandwidth allows for the secure and faster transmission of encrypted model parameters, supporting privacy-preserving FL techniques such as secure aggregation [154].

2) *Increased Connection Density*: 6G networks are designed to support an unprecedented density of interconnected devices, far exceeding the capabilities of 5G. This increased

TABLE X: Quantitative comparison of PPFL approaches through experimental evaluation.

PPFL Technique	Method	Accuracy (%)	Average Client Training Time (sec)	Average Aggregation Time (ms)	Total Convergence Time (sec)	Memory Footprint (MBs)
Differential Privacy	Gaussian Noise	94.18	0.5974	3.8	67.1942	2.1170
	Laplace Noise	99.37	0.6130	3.6	68.7546	2.1156
	Poisson Noise	91.41	0.5957	3.5	67.0437	2.2086
SMPC	Shamir's Secret Sharing	69.36	0.6052	0.0038	238.7399	2.3606
	Additive Secret Sharing	68.24	0.6078	0.1	240.0294	2.1170
	Garbled Circuits	63.58	0.6095	0.1	248.8431	2.1170
Homomorphic Encryption	PHE	99.15	0.6532	0.004	84.6845	0.9352
	FHE	99.20	0.6364	0.003	87.1462	1.1870
Secure Aggregation	Random Masking	99.10	0.5421	0.243	55.3322	1.1482
	Partial Sum Masking	99.13	0.5570	0.225	56.8228	1.1659
Anonymization	RCI	99.41	0.6821	3.6	75.8817	2.7704
	Rotating ID	95.32	0.4339	3.7	48.9923	0.8542

capacity allows greater participation in FL by a wide array of IoT devices, ensuring more diverse and representative model training. From a privacy perspective, the inclusion of more devices enables stronger differential privacy mechanisms, as noise can be distributed across a larger number of participants, improving both privacy and utility [155].

3) *Reliable Communication*: Communication reliability is a cornerstone of FL implementations in critical IoT applications such as healthcare and autonomous systems. Ultra-reliable low-latency communication (URLLC) of 6G minimizes the likelihood of packet loss or data corruption during model updates, reducing vulnerabilities to man-in-the-middle attacks and ensuring the integrity of privacy preservation techniques such as HE [156].

4) *Edge Computing and Distributed AI Integration*: 6G networks will natively integrate edge computing capabilities with distributed AI, enabling IoT devices to process data and train models closer to the data source. By reducing the need to send raw data over the network, this approach inherently enhances privacy. For FL, this evolution will support advanced privacy-preserving frameworks such as SMPC and federated unlearning, which benefit from decentralized data processing and localized model updates [157].

5) *Intelligent Privacy Management*: 6G introduces the concept of "native AI" in its infrastructure, embedding AI in the core of network operations. This allows dynamic optimization of privacy-preserving FL protocols, such as adaptive noise injection for differential privacy based on real-time traffic conditions. Furthermore, AI-driven anomaly detection mechanisms in 6G can proactively identify and mitigate potential privacy threats during model aggregation [158].

6) *Quantum-Safe Security Enhancements*: 6G networks are likely to incorporate quantum-resistant cryptographic algorithms to counteract the threat of quantum computing, which could potentially break traditional encryption methods. This advancement is critical for FL, where encrypted model updates and aggregated results need to remain secure over extended periods. The integration of such quantum-safe protocols ensures long-term privacy preservation in FL [159].

7) *Privacy-Preserving Localization and Context Awareness*: 6G's capability for precise localization and context awareness offers new opportunities for FL in IoT. By securely and accurately determining device locations and contexts, FL systems can optimize model updates without revealing sensi-

tive location information. This feature is particularly beneficial for location-sensitive IoT applications, such as smart cities and personalized healthcare [160].

The evolution from 5G to 6G represents a significant boon for FL in the IoT by addressing critical aspects such as speed, connectivity, reliability, and security. These technological advances are poised to overcome many of the current limitations of FL deployments, which will enable more dynamic, responsive, and efficient IoT systems.

B. Future Challenges and Opportunities

Although advances in technologies such as 5G and 6G offer substantial benefits, several challenges remain for future research and development in FL for IoT:

1) *FL in Heterogeneous IoT Networks*: IoT networks often consist of devices with diverse computational capabilities, communication protocols, and energy constraints [161]. Developing adaptive FL frameworks that account for this heterogeneity is critical. Techniques such as device clustering [162], personalized FL [163], and resource-aware adaptive FL [164], [165] are promising directions to address these challenges. Heterogeneous FL models may also represent an option in which a different model may be adapted according to the devices capabilities. However, the accuracy of such models represents a big challenge.

2) *Multimodal FL and Privacy in IoT*: With the proliferation of multimodal IoT data (e.g., audio, video, and sensor data), FL must effectively integrate diverse data types while ensuring privacy preservation. Research into multimodal FL techniques, such as cross-modal learning and shared representation models, can enhance the robustness of FL systems [166]. However, multimodal data also introduce unique privacy challenges, which require advanced cryptographic and anonymization techniques [167].

3) *Federated Unlearning for Privacy Mitigation*: Federated unlearning, a process to selectively remove specific data contributions from a trained model, has gained attention as a method to address privacy concerns [168]. This technique is particularly relevant in IoT scenarios where users can withdraw their consent to use data [169]. The design of efficient and scalable unlearning algorithms compatible with FL frameworks will be an important research area [170].

4) *Collusion Attacks*: Despite robust privacy measures, the threat of collusion attacks, in which multiple malicious entities cooperate to compromise the system, remains unresolved. The design of systems that counteract these advanced threats is an open research question [171].

5) *Computational Overhead*: While bolstering security, integrating cryptographic techniques might introduce significant computational and communication overhead. Balancing these with the need for real-time responses will be crucial in applications such as real-time navigation and healthcare [172].

6) *Data Imbalance*: Uneven data distribution among participants is a recurrent issue in FL. Methods for ensuring efficient learning, even with data disparity, will remain a focal point of future research. Techniques such as federated transfer learning and adaptive sampling are promising areas to explore [173].

7) *Transparency vs. Privacy*: While blockchain provides immutable record keeping and ensures the integrity of data and the model, it introduces another tension between transparency and privacy. Striking a balance between these aspects in various application scenarios will be an important research trend [174], [175].

8) *Interoperability*: Ensuring interoperability between various IoT devices and systems is critical to the widespread adoption of FL. Developing standardized protocols and interfaces seamlessly integrating devices from different manufacturers will be essential [176].

9) *Scalability*: As IoT networks grow, scaling FL to accommodate millions of devices without affecting performance is a significant challenge. Research in hierarchical FL and efficient model aggregation techniques will be necessary to address this [177].

10) *Policy and Regulation*: With the growing focus on data privacy and security, compliance with regulations such as GDPR and CCPA will be essential. Developing FL frameworks that comply with these regulations while maintaining high performance will be a key challenge [178].

11) *Explainability and Interpretability*: As FL becomes more prevalent, the need for models that can explain their predictions and decisions becomes critical, especially in sectors such as healthcare and finance. Developing techniques to ensure that FL models are interpretable will be an important area of research [179].

12) *Integrating AI with FL*: Combining artificial intelligence (AI) techniques with FL can improve learning. For example, reinforcement learning can optimize client selection and resource allocation in FL systems, improving efficiency and performance [180].

13) *Generative AI*: Generative AI models, such as GANs, Diffusion models, and Large language models, can be leveraged to create synthetic data, which can help address data imbalance and enhance the robustness of FL models. These models can generate realistic and diverse data samples that can be used to augment training datasets and improve the generalizability of FL systems. Furthermore, Generative AI can help anonymize data, preserving privacy while providing useful synthetic data for training purposes [178].

In conclusion, while emerging technologies such as 5G, 6G, and edge computing significantly enhance FL in IoT, address-

ing the associated challenges will require continued innovation and research. Future work will need to focus on optimizing the balance between privacy, efficiency, and scalability to fully realize FL's potential in diverse IoT applications.

VIII. STANDARDS, PROTOCOLS, AND REAL-WORLD APPLICATIONS

This section explores the foundational technologies that enable effective FL implementations within IoT environments. These technologies address core challenges such as security, interoperability, and scalability, which are critical to ensure efficient operation of FL systems on diverse and distributed IoT platforms. Key standards that provide the guidelines and specifications necessary for network efficiency are discussed. In addition, we explore various protocols that facilitate reliable and secure data exchanges between IoT devices and FL systems. Finally, some real-world applications of PPFL from world-renowned companies are highlighted to assess the full potential of FL in IoT. Each component, standards, protocols, and frameworks plays a vital role in shaping the landscape of FL, making it a viable and robust solution to the complex demands of IoT ecosystems.

A. Standards and Protocols

1) *Network Efficiency Standards*: This section explores the critical standards and protocols that improve network efficiency, essential for managing the substantial data flows inherent in FL deployments across numerous IoT devices. It discusses various strategies and technologies to optimize bandwidth usage, reduce latency, and improve communication efficiency between disparate devices. These measures support the robust performance of FL applications and ensure that network resources are utilized judiciously to prevent bottlenecks and maximize the throughput of IoT networks.

a) *IEEE P1932.1*: IEEE P1932.1¹ is an emerging standard that establishes interoperability within wireless mobile networks. Although it acknowledges the growing relevance of FL, its primary objective is not to directly facilitate federated machine learning systems. Instead, the standard defines a robust framework of interfaces and interactions that promote seamless collaboration among diverse devices and platforms. By focusing on interoperability, IEEE P1932.1 ensures that devices from various manufacturers can communicate and function together effectively across different applications, including but not limited to FL. This interoperability framework enables diverse technologies to participate in complex environments without compatibility challenges, supporting innovative use cases across multiple domains.

b) *3GPP and 5G Standards*: 3GPP² has been instrumental in defining standards that facilitate enhanced mobile broadband, ultra-reliable low-latency communications, and massive machine-type communications – all beneficial for IoT applications, including FL. The 5G network standard is crucial for IoT and FL due to its ability to support high data

¹<https://standards.ieee.org/ieee/1932.1/7042/>

²<https://www.3gpp.org/>

rate transmissions, reduced latency, and increased connectivity density, which are essential for effective FL operation in IoT networks.

2) *Privacy Standards*: This section delves into the standards and protocols to protect data privacy and secure communication in FL deployments.

a) *IEEE P7000 Series*: This series of standards focuses on ethical concerns and privacy issues related to developing and implementing autonomous systems and artificial intelligence. For example, IEEE P7002 on Data Privacy Process addresses how to manage privacy issues within systems and software engineering processes. This standard is relevant because it helps developers of IoT devices and FL systems ensure privacy is built into the design and deployment phases.

b) *IETF Protocols (DTLS, TLS)*: The Internet Engineering Task Force (IETF) has developed several security protocols, such as Datagram Transport Layer Security (DTLS) and Transport Layer Security (TLS), which are crucial for securing communications in IoT networks. These protocols help protect the privacy of data as it moves between devices and servers, which is essential for FL applications where data can be particularly sensitive.

c) *ISO/IEC 27001*: This international standard describes best practices for an information security management system (ISMS) and is vital for organizations that manage information related to FL and IoT. Implementing ISO/IEC 27001 helps ensure user data's confidentiality, integrity, and availability and can be particularly useful in managing the privacy of data distributed across FL IoT systems.

d) *GDPR and Related Frameworks*: While not a technical standard, the General Data Protection Regulation (GDPR) profoundly impacts how data privacy must be handled by any FL and IoT system operating within or targeting users in the European Union. It emphasizes users' privacy rights and the secure processing of personal data, influencing how standards and protocols are implemented in IoT and FL environments.

e) *FIDO Alliance (Fast Identity Online)*: The FIDO Alliance has developed standards for stronger authentication practices that reduce password reliance, enhancing privacy and security for online services, including IoT devices. These standards ensure secure and private user authentication in FL systems.

B. Real-world Implementations of PPFL

This section provides an overview of real-world implementations of PPFL in IoT environments, addressing associated privacy issues and showcasing examples of effective use. An overview of some practical PPFL projects from world-renowned companies is presented in Fig. 7.

1) *Google (TensorFlow Federated (TFF))*: Google's TFF³ is a widely adopted FL framework, designed to preserve privacy while allowing the training of decentralized models. TFF leverages differential privacy to mask individual contributions and secure aggregation to ensure data is encrypted during computations [181]. Use cases include predictive text input in Gboard and training healthcare AI models without sensitive

medical data leaving user devices, addressing concerns about privacy leakage. Applying FL to Gboard allows the model to learn from user interactions directly on their devices, ensuring that personal data remain local.

In [182], researchers trained more than 20 Gboard language models with (ρ -zCDP) privacy guarantees where $\rho \in (0.2, 2)$. Two of the models were trained with secure aggregation. All models achieved stronger privacy guarantees (smaller zCDP) than the zCDP > 2.6 standard used by the US Census Bureau. The authors found that system configuration parameters like report goal and minimum client separation have a significant impact on the privacy-utility trade-off. However, they noted that training was "notably slower" with secure aggregation, indicating potential areas for future optimization. The results of this research demonstrate that it is possible to train production-ready language models with strong privacy guarantees using FL and DP. This research was impactful enough that Gboard has implemented a new policy requiring differential privacy to be used for all future training of Gboard language models.

2) *Microsoft (Azure Confidential Computing and Project Florida)*: Microsoft leverages PPFL within Azure Confidential Computing by combining secure enclaves, HE, and SMPC for multiparty collaborations [183]. The latest enhancements include scalable privacy-preserving workflows for financial fraud detection and clinical AI trials. Innovations focus on enabling organizations to train collaborative models while adhering to strict data protection standards such as GDPR and HIPAA, particularly in healthcare diagnostics and sensitive legal analytics.

Project Florida [184] aims to simplify FL deployment, enabling FL as a Service (FLaaS) and facilitating the broader adoption of this privacy-preserving machine learning approach. Project Florida has been tested on production workloads with up to 70000 connected devices, demonstrating its capability for real-world applications. In [184], an experiment using Project Florida demonstrated spam classification with FL. Using BERT Tiny and the Enron Spam dataset, the study simulated 32 clients on 8 AzureML nodes, with training conducted over 10 iterations using FedAvg. Variations included adding DP with $\epsilon = 2$, asynchronous learning to reduce the duration of the iteration, and overparticipation by increasing client nodes to 16, further optimizing training time. The results showed high test accuracy, though DP caused slight accuracy drops and convergence challenges. Scaling tests validated Florida's ability to handle more than 1000 concurrent clients with potential scalability to hundreds of thousands per iteration.

3) *NVIDIA (Clara FL)*: NVIDIA's Clara platform uses FL to enable collaborative development of AI models in the healthcare sector. By allowing institutions to train models on their local data and share only model updates, Clara facilitates the creation of robust medical AI applications without exposing sensitive patient information. PPFL techniques such as HE and secure aggregation to protect patient data during federated medical AI model training [185]. Recent advances focus on large-scale collaborative efforts in genomics, early cancer detection, and real-time analysis of medical imaging. These updates improve model accuracy while meeting

³<https://github.com/tensorflow/federated>

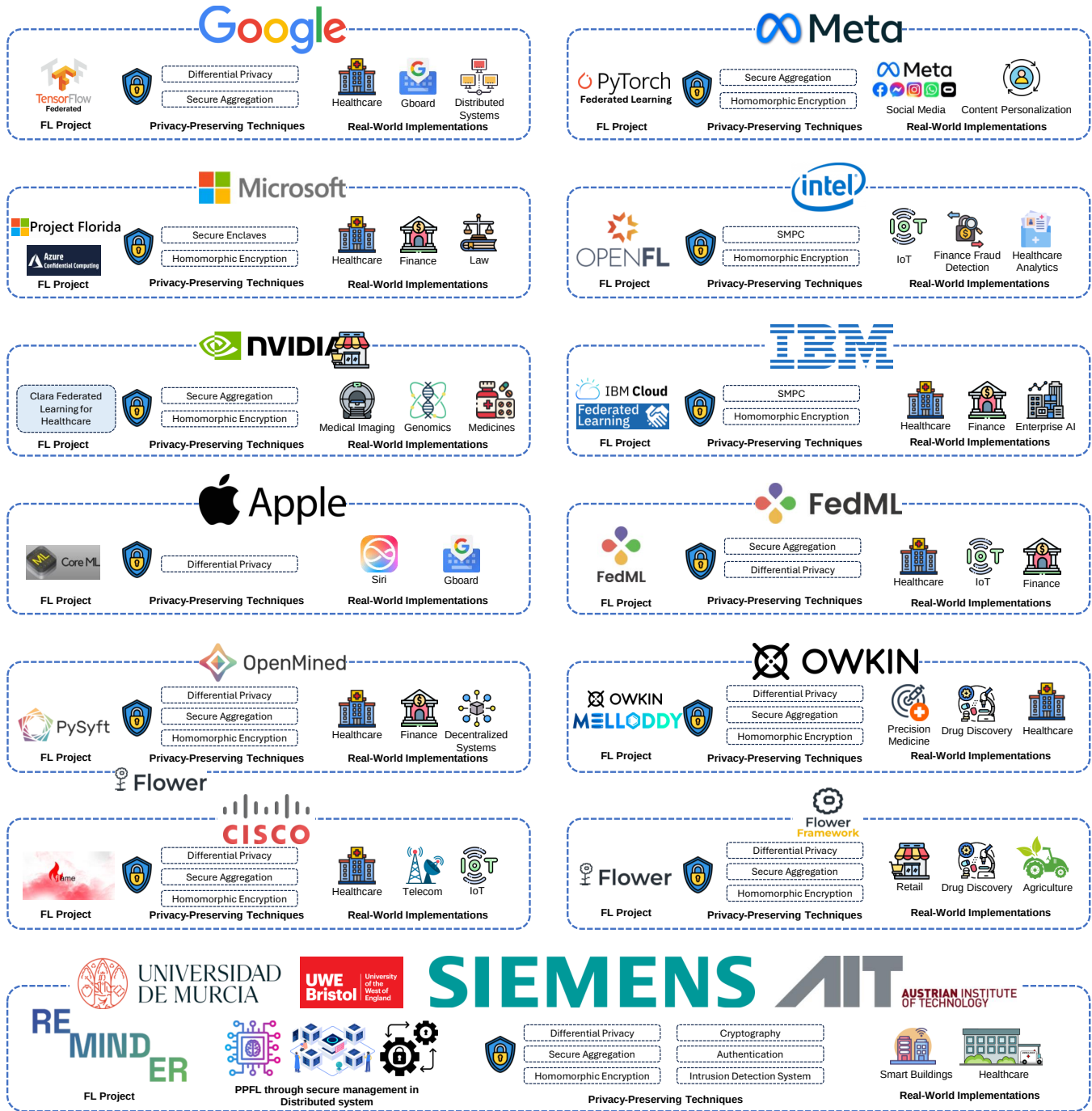


Fig. 7: Examples of some PPFL implementations from world-renowned enterprises and research institutions.

stringent data privacy requirements in global decentralized healthcare systems.

4) *Apple (Core ML)*: Apple's Core ML implements PPFL by training models locally on user devices, preventing raw data from being transmitted to central servers [186]. Enhanced privacy techniques include DP and local encryption, enabling secure training for Siri, keyboard predictions, and health monitoring. Recent advances focus on FL in Apple Health and biometric data, ensuring end-to-end privacy on personal and wearable devices.

5) *OpenMined (PySyft)*: PySyft⁴ enhances PPFL implementation with open source tools for secure computation, including SMPC, DP, and HE. It allows organizations to collaborate securely on sensitive data in multiple domains [187]⁵. PySyft has been applied in various domains, including healthcare and finance, to develop models that respect data privacy regulations. Recent updates focus on scaling privacy-

⁴<https://github.com/OpenMined/PySyft>

⁵<https://blog.openmined.org/federated-privacy-preserving-analytics-for-secure-collaboration-among-telco-and-partners-to-improve-customer-engagement/>

preserving collaborations for AI in fraud detection, predictive patient diagnostics, and cross-institutional research on large datasets.

6) *Meta (PyTorch FL)*: Meta applies PPFL in PyTorch FL by using secure aggregation for privacy-preserving updates in large-scale models. It enables decentralized learning across billions of users while protecting individual data [188]. New developments include enhanced federated techniques for personalized recommendations, privacy-compliant ad targeting, and secure AI-driven analytics for social media platforms such as Facebook, Instagram, and WhatsApp.

7) *Intel (Open Federated Learning (OpenFL))*: Intel's OpenFL⁶ implements PPFL through SMPC and DP for decentralized training in industries such as healthcare and IoT. Recent updates include enhanced multi-tier workflows that allow secure AI training for predictive maintenance in smart devices and healthcare analytics. The platform ensures privacy-compliant collaboration between organizations on sensitive data while maintaining robust encryption standards [189].

8) *IBM (Federated Learning on AI Cloud)*: IBM integrates PPFL into its AI Cloud platform, enabling enterprises to collaboratively build AI models without exposing sensitive data [190]. Privacy measures include homomorphic encryption for secure computations and SMPC for data sharing between different organizations. Key applications are healthcare diagnostics and fraud detection in financial systems.

9) *Flower*: Flower⁷ is a flexible and friendly FL framework to facilitate the development and experimentation of FL algorithms and applications. Unlike many FL frameworks that focus on specific technologies or platforms, Flower is designed to be framework-agnostic, which means that it can be used with a variety of ML libraries and environments, including popular ones like TensorFlow, PyTorch, and scikit-learn.

The Flower framework demonstrates exceptional scalability and adaptability in FL [191], handling up to 15 million clients, with efficient training even when 1000 clients participate per round. It outperformed other FL frameworks like *FedJax* and *TFF* in computationally intensive settings, completing training significantly faster in larger local epochs. Tests on heterogeneous devices, including Nvidia Jetson, Raspberry Pi, and Android smartphones, highlighted Flower's low framework overhead (under 100 ms per round) and energy efficiency, with Jetson Nano-CPU reducing energy usage by 60% compared to Raspberry Pi despite longer convergence times. Experiments with varying client network speeds revealed significant increases in training time with slower clients, although customized client selection strategies reduced convergence time by up to 30%. Furthermore, Flower's implementation of secure aggregation protocols proved robust, with linear computational overhead and resilience to client dropouts. These results show Flower as a highly scalable, efficient and adaptable framework for real-world FL deployments.

10) *Cisco (Flame)*: Flame, developed by Cisco, is an open source platform designed to streamline the FL life cycle, including tasks such as managing compute resources and

datasets, scheduling jobs, preserving privacy and monitoring the system. In addition, it emphasizes edge computing applications. Flame experiments validate its flexibility, scalability, and performance in various FL scenarios [192]. Hybrid FL, a Flame-enabled topology, achieved the target accuracy 1.77 faster than Classical FL (C-FL) and Hierarchical FL (H-FL) under straggler node conditions, while H-FL outperformed other topologies during client node failure scenarios. Flame's TAG framework (Topology Abstraction Graph) demonstrated scalability, expanding to support 100000 trainers in less than 32 seconds. Further tests on MNIST and CIFAR-10 datasets showed that Hybrid FL achieved up to 2.21 faster training with significantly lower bandwidth usage compared to C-FL. Coordinated FL (CO-FL), another topology supported by Flame, efficiently handled straggler aggregators and reduced iteration times compared to static hierarchical setups. Comparisons with existing frameworks such as FedML and Flower highlighted Flame's superiority in supporting various topologies while achieving comparable accuracy with lower execution times, underscoring its utility for complex and large-scale FL deployments, such as IoT scenarios.

11) *FedML (Open-Source Federated Learning Platform)*: FedML enhances PPFL implementation with scalable privacy-preserving techniques such as DP and secure aggregation [193]. It supports various applications, including federated IoT, predictive maintenance, and disease diagnosis. Recent updates focus on lightweight PPFL solutions for edge devices and hierarchical FL to facilitate collaboration across multiple organizational tiers.

12) *REMINDER*: The REMINDER Project [194] is dedicated to develop a platform enhancing privacy-preserving ML, by ensuring secure data management throughout its life cycle within decentralized and distributed systems. By integrating FL with cryptographic techniques, REMINDER enables collaborative AI model training without sharing raw data, addressing critical privacy and security concerns in dynamic, heterogeneous environments such as IoT networks and 5G/6G infrastructures. Its objectives include the development of advanced authentication protocols, edge-based FL architectures, and adaptive systems that mitigate new threat vectors while complying with regulations such as GDPR. The practical applications of REMINDER are demonstrated through use cases in the eHealth and smart building ecosystems. These examples underscore the transformative potential of FL in enabling secure, privacy-centric AI solutions. In addition, the project aligns with the Sustainable Development Goals (SDGs) of the United Nations, promoting sustainable and responsible technological innovation.

13) *OWKIN (MELLODDY project)*: OWKIN is a leading AI platform that uses FL to enable secure and decentralized model training in sensitive industries such as healthcare. Its framework integrates differential privacy to anonymize individual data contributions and homomorphic encryption to ensure secure computations. OWKIN's use cases include the MELLODDY project, where pharmaceutical companies collaboratively train drug discovery models without exposing proprietary data, and hospital collaborations that train healthcare AI models while patient data remain on local servers, ad-

⁶<https://github.com/securefederatedai/openfl>

⁷<https://github.com/adap/flower>

addressing stringent privacy and regulatory requirements [195].

IX. CONCLUSION

Equipping IoT applications with advanced AI marks a pivotal step in the evolution of ICT. This enables advanced applications and services that will significantly influence modern economies and improve citizen's life. However, the acceptance and integration of AI models in personal and industrial IoT devices faces many challenges, including severe resource constraints, strict requirements for convergence speed and accuracy, and the pressing need to ensure user data privacy. FL has emerged as a promising paradigm to address these challenges in IoT environments, but it still faces challenges that need to be addressed. This paper has provided a comprehensive overview and taxonomy of PPFL techniques tailored to IoT systems. The taxonomy enables a structured evaluation of PPFL methods, examining their privacy guarantees, quality of service, and efficiency in IoT networks.

Our study reveals several key insights:

- 1) No Universal Solution: PPFL solutions must navigate complex trade-offs between resource utilization, learning performance, and privacy resilience. Effective solutions require tailored approaches to balance these conflicting demands.
- 2) Client Selection: Optimal client selection based on energy reserves, memory, bandwidth, and computational capacity is critical for efficient FL rounds, enhancing learning performance while minimizing resource waste.
- 3) Hybrid Privacy-Preserving Techniques: Combining methods such as encryption, anonymization, and DP can harness their strengths. For example, DP can preserve data utility while ensuring robust privacy protection.
- 4) Reducing Computational Overhead: Lowering the resource demands of techniques like HE and secure multi-party computation is essential for their practical adoption in FL-IoT.
- 5) Optimization Frameworks: Mathematical tools such as game theory can help balance data quality and privacy, offering strategies to manage competing demands for resource efficiency and security.
- 6) Advanced Network Technologies: 5G, 6G, and edge computing provide critical improvements in bandwidth, latency, and scalability, supporting real-time FL in IoT systems.

In conclusion, while PPFL has made remarkable strides for IoT, there are still persistent challenges that current research attempts to overcome. Future efforts should focus on optimizing the balance between privacy, efficiency, and scalability. Advances in network technologies and sophisticated optimization techniques will be instrumental in unlocking FL's full potential across diverse IoT applications.

REFERENCES

- [1] M. Centenaro, C. E. Costa, F. Granelli, C. Sacchi, and L. Vangelista, "A survey on technologies, standards and open challenges in satellite iot," *IEEE Communications Surveys Tutorials*, vol. 23, no. 3, pp. 1693–1720, 2021.
- [2] Z.-H. Zhan, J.-Y. Li, and J. Zhang, "Evolutionary deep learning: A survey," *Neurocomputing*, vol. 483, pp. 42–58, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0925231222001345>
- [3] D. Djenouri, R. Laidi, Y. Djenouri, and I. Balasingham, "Machine learning for smart building applications: Review and taxonomy," *ACM Comput. Surv.*, vol. 52, no. 2, pp. 24:1–24:36, 2019.
- [4] A. K. Alnaim and A. M. Alwakeel, "Machine-learning-based iot-edge computing healthcare solutions," *Electronics*, vol. 12, no. 4, p. 1027, 2023.
- [5] A. Lakhan, Q.-u.-a. Mastoi, M. A. Dootio, F. Alqahtani, I. R. Alzahrani, F. Baothman, S. Y. Shah, S. A. Shah, N. Anjum, Q. H. Abbasi *et al.*, "Hybrid workload enabled and secure healthcare monitoring sensing framework in distributed fog-cloud network," *Electronics*, vol. 10, no. 16, p. 1974, 2021.
- [6] K. M. J. Rahman, F. Ahmed, N. Akhter, M. Hasan, R. Amin, K. E. Aziz, A. K. M. M. Islam, M. S. H. Mukta, and A. K. M. N. Islam, "Challenges, applications and design aspects of federated learning: A survey," *IEEE Access*, vol. 9, pp. 124 682–124 700, 2021.
- [7] S. Latif, D. Djenouri, J. L. Hernandez-Ramos, A. Skarmeta, and J. Ahmad, "A lightweight integrity-driven federated learning approach to mitigate poisoning attacks in iot," in *2024 IEEE Future Networks World Forum (FNWF)*. IEEE, 2024, pp. 1–6, <https://uwe-repository.worktribe.com/output/12895988>.
- [8] M. Aledhari, R. Razzak, R. M. Parizi, and F. Saeed, "Federated learning: A survey on enabling technologies, protocols, and applications," *IEEE Access*, vol. 8, pp. 140 699–140 725, 2020.
- [9] Y. Li, C. Chen, N. Liu, H. Huang, Z. Zheng, and Q. Yan, "A blockchain-based decentralized federated learning framework with committee consensus," *IEEE Network*, vol. 35, no. 1, pp. 234–241, 2021.
- [10] A. Imteaj, U. Thakker, S. Wang, J. Li, and M. H. Amini, "A survey on federated learning for resource-constrained iot devices," *IEEE Internet of Things Journal*, vol. 9, no. 1, pp. 1–24, 2022.
- [11] W. Ni, H. Ao, H. Tian, Y. C. Eldar, and D. Niyato, "Fedsl: Federated split learning for collaborative healthcare analytics on resource-constrained wearable iomt devices," *IEEE Internet Things J.*, vol. 11, no. 10, pp. 18 934–18 935, 2024.
- [12] S. I. Ahsan, D. Djenouri, and R. Haider, "Privacy-enhanced sentiment analysis in mental health: Federated learning with data obfuscation and bidirectional encoder representations from transformers," *Electronics*, vol. 13, no. 23, 2024.
- [13] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *2017 IEEE Symposium on Security and Privacy (SP)*, 2017, pp. 3–18.
- [14] L. T. Phong, Y. Aono, T. Hayashi, L. Wang, and S. Moriai, "Privacy-preserving deep learning via additively homomorphic encryption," *IEEE Transactions on Information Forensics and Security*, vol. 13, no. 5, pp. 1333–1345, 2018.
- [15] L. Su and J. Xu, "Securing distributed gradient descent in high dimensional statistical learning," *Proceedings of the ACM on Measurement and Analysis of Computing Systems*, vol. 3, no. 1, pp. 1–41, 2019.
- [16] L. Zhu, Z. Liu, and S. Han, "Deep leakage from gradients," *Advances in neural information processing systems*, vol. 32, 2019.
- [17] L. Lyu, H. Yu, X. Ma, L. Sun, J. Zhao, Q. Yang, and P. S. Yu, "Privacy and robustness in federated learning: Attacks and defenses," *CoRR*, vol. abs/2012.06337, 2020. [Online]. Available: <https://arxiv.org/abs/2012.06337>
- [18] D. C. Nguyen, M. Ding, P. N. Pathirana, A. Seneviratne, J. Li, and H. Vincent Poor, "Federated learning for internet of things: A comprehensive survey," *IEEE Communications Surveys Tutorials*, vol. 23, no. 3, pp. 1622–1658, 2021.
- [19] T. Zhang, L. Gao, C. He, M. Zhang, B. Krishnamachari, and A. S. Avestimehr, "Federated learning for the internet of things: applications, challenges, and opportunities," *IEEE Internet of Things Magazine*, vol. 5, no. 1, pp. 24–29, 2022.
- [20] E. M. Campos, P. F. Saura, A. González-Vidal, J. L. Hernández-Ramos, J. B. Bernabé, G. Baldini, and A. Skarmeta, "Evaluating federated learning for intrusion detection in internet of things: Review and challenges," *Computer Networks*, vol. 203, p. 108661, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1389128621005405>
- [21] Z. Du, C. Wu, T. Yoshinaga, K.-L. A. Yau, Y. Ji, and J. Li, "Federated learning for vehicular internet of things: Recent advances and open issues," *IEEE Open Journal of the Computer Society*, vol. 1, pp. 45–61, 2020.
- [22] D. C. Nguyen, M. Ding, P. N. Pathirana, A. Seneviratne, J. Li, D. Niyato, and H. V. Poor, "Federated learning for industrial internet

- of things in future industries,” *IEEE Wireless Communications*, vol. 28, no. 6, pp. 192–199, 2021.
- [23] L. U. Khan, W. Saad, Z. Han, E. Hossain, and C. S. Hong, “Federated learning for internet of things: Recent advances, taxonomy, and open challenges,” *IEEE Communications Surveys Tutorials*, vol. 23, no. 3, pp. 1759–1799, 2021.
 - [24] S. I. Al-Sharekh and K. H. A. Al-Shqeerat, “An overview of privacy issues in iot environments,” in *2019 International Conference on Advances in the Emerging Computing Technologies (AECT)*, 2020, pp. 1–6.
 - [25] M. R. Asghar, G. Dán, D. Miorandi, and I. Chlamtac, “Smart meter data privacy: A survey,” *IEEE Communications Surveys Tutorials*, vol. 19, no. 4, pp. 2820–2835, 2017.
 - [26] S.-C. Cha, T.-Y. Hsu, Y. Xiang, and K.-H. Yeh, “Privacy enhancing technologies in the internet of things: Perspectives and challenges,” *IEEE Internet of Things Journal*, vol. 6, no. 2, pp. 2159–2187, 2019.
 - [27] S. Imtiaz, R. Sadre, and V. Vlassov, “On the case of privacy in the iot ecosystem: A survey,” in *2019 International Conference on Internet of Things (iThings) and IEEE Green Computing and Communications (GreenCom) and IEEE Cyber, Physical and Social Computing (CPSCom) and IEEE Smart Data (SmartData)*, 2019, pp. 1015–1024.
 - [28] Z. Yu, Z. Kaplan, Q. Yan, and N. Zhang, “Security and privacy in the emerging cyber-physical world: A survey,” *IEEE Communications Surveys Tutorials*, vol. 23, no. 3, pp. 1879–1919, 2021.
 - [29] M. Asad, A. Moustafa, and C. Yu, “A critical evaluation of privacy and security threats in federated learning,” *Sensors*, vol. 20, no. 24, 2020. [Online]. Available: <https://www.mdpi.com/1424-8220/20/24/7182>
 - [30] A. Blanco-Justicia, J. Domingo-Ferrer, S. Martínez, D. Sánchez, A. Flanagan, and K. E. Tan, “Achieving security and privacy in federated learning systems: Survey, research challenges and future directions,” 2020.
 - [31] V. Mothukuri, R. M. Parizi, S. Pouriyeh, Y. Huang, A. Dehghantanha, and G. Srivastava, “A survey on security and privacy of federated learning,” *Future Generation Computer Systems*, vol. 115, pp. 619–640, 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0167739X20329848>
 - [32] X. Yin, Y. Zhu, and J. Hu, “A comprehensive survey of privacy-preserving federated learning: A taxonomy, review, and future directions,” *ACM Comput. Surv.*, vol. 54, no. 6, jul 2021. [Online]. Available: <https://doi.org/10.1145/3460427>
 - [33] Z. Li, V. Sharma, and S. P. Mohanty, “Preserving data privacy via federated learning: Challenges and solutions,” *IEEE Consumer Electronics Magazine*, vol. 9, no. 3, pp. 8–16, 2020.
 - [34] N. Truong, K. Sun, S. Wang, F. Guitton, and Y. Guo, “Privacy preservation in federated learning: An insightful survey from the gdpr perspective,” *Computers and Security*, vol. 110, p. 102402, 2021. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S01674048211002261>
 - [35] N. Rodríguez-Barroso, D. J. López, M. V. Luzón, F. Herrera, and E. Martínez-Cámara, “Survey on federated learning threats: concepts, taxonomy on attacks and defences, experimental study and challenges,” 2022.
 - [36] A. Wainakh, E. Zimmer, S. Subedi, J. Keim, T. Grube, S. Karuppayah, A. S. Guinea, and M. Mühlhäuser, “Federated learning attacks revisited: A critical discussion of gaps, assumptions, and evaluation setups,” 2022.
 - [37] L. Lyu, H. Yu, and Q. Yang, “Threats to federated learning: A survey,” *arXiv preprint arXiv:2003.02133*, 2020.
 - [38] X. X. . W. W. Liu, P., “Threats, attacks and defenses to federated learning: issues, taxonomy and perspectives,” *Cybersecurity*, vol. 5, 4, 2022.
 - [39] C. Briggs, Z. Fan, and P. Andras, *A Review of Privacy-Preserving Federated Learning for the Internet-of-Things*. Cham: Springer International Publishing, 2021, pp. 21–50.
 - [40] W. Li, F. Milletari, D. Xu, N. Rieke, J. Hancox, W. Zhu, M. Baust, Y. Cheng, S. Ourselin, M. J. Cardoso, and A. Feng, “Privacy-preserving federated brain tumour segmentation,” in *Machine Learning in Medical Imaging: 10th International Workshop, MLMI 2019, Held in Conjunction with MICCAI 2019, Shenzhen, China, October 13, 2019, Proceedings*. Berlin, Heidelberg: Springer-Verlag, 2019, p. 133–141. [Online]. Available: https://doi.org/10.1007/978-3-030-32692-0_16
 - [41] M. J. Sheller, G. A. Reina, B. Edwards, J. Martin, and S. Bakas, “Multi-institutional deep learning modeling without sharing patient data: A feasibility study on brain tumor segmentation,” in *Brainlesion: Glioma, Multiple Sclerosis, Stroke and Traumatic Brain Injuries: 4th International Workshop, BrainLes 2018, Held in Conjunction with MICCAI 2018, Granada, Spain, September 16, 2018, Revised Selected Papers, Part I 4*. Springer, 2019, pp. 92–104.
 - [42] B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y. Arcas, “Communication-Efficient Learning of Deep Networks from Decentralized Data,” in *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, ser. Proceedings of Machine Learning Research, A. Singh and J. Zhu, Eds., vol. 54. PMLR, 20–22 Apr 2017, pp. 1273–1282.
 - [43] T. Li, A. K. Sahu, M. Zaheer, M. Sanjabi, A. Talwalkar, and V. Smith, “Federated optimization in heterogeneous networks,” 2018. [Online]. Available: <https://arxiv.org/abs/1812.06127>
 - [44] J. Wang, Q. Liu, H. Liang, G. Joshi, and H. V. Poor, “Tackling the objective inconsistency problem in heterogeneous federated optimization,” in *Proceedings of the 34th International Conference on Neural Information Processing Systems*, ser. NIPS ’20. Red Hook, NY, USA: Curran Associates Inc., 2020.
 - [45] A. Lalitha, O. C. Kilinc, T. Javidi, and F. Koushanfar, “Peer-to-peer federated learning on graphs,” 2019. [Online]. Available: <https://arxiv.org/abs/1901.11173>
 - [46] L. Yuan, Z. Wang, L. Sun, S. Y. Philip, and C. G. Brinton, “Decentralized federated learning: A survey and perspective,” *IEEE Internet of Things Journal*, 2024.
 - [47] X. Zhou, W. Liang, I. Kevin, K. Wang, Z. Yan, L. T. Yang, W. Wei, J. Ma, and Q. Jin, “Decentralized p2p federated learning for privacy-preserving and resilient mobile robotic systems,” *IEEE Wireless Communications*, vol. 30, no. 2, pp. 82–89, 2023.
 - [48] H. Kim, J. Park, M. Bennis, and S.-L. Kim, “Blockchain on-device federated learning,” *IEEE Communications Letters*, vol. 24, no. 6, pp. 1279–1283, 2020.
 - [49] H. Zhang, S. Jiang, and S. Xuan, “Decentralized federated learning based on blockchain: concepts, framework, and challenges,” *Computer Communications*, vol. 216, pp. 140–150, 2024.
 - [50] E. T. M. Beltrán, M. Q. Pérez, P. M. S. Sánchez, S. L. Bernal, G. Bovet, M. G. Pérez, G. M. Pérez, and A. H. Celdrán, “Decentralized federated learning: Fundamentals, state of the art, frameworks, trends, and challenges,” *IEEE Communications Surveys & Tutorials*, 2023.
 - [51] J. Weng, J. Weng, J. Zhang, M. Li, Y. Zhang, and W. Luo, “Deepchain: Auditable and privacy-preserving deep learning with blockchain-based incentive,” *IEEE Transactions on Dependable and Secure Computing*, vol. 18, no. 5, pp. 2438–2455, 2021.
 - [52] J. Li, Y. Shao, K. Wei, M. Ding, C. Ma, L. Shi, Z. Han, and H. V. Poor, “Blockchain assisted decentralized federated learning (blade-fl): Performance analysis and resource allocation,” *IEEE Transactions on Parallel and Distributed Systems*, vol. 33, no. 10, pp. 2401–2415, 2021.
 - [53] W. Zhu, L. Shi, J. Li, B. Cao, K. Wei, Z. Wang, and T. Huang, “Trust-worthy blockchain-assisted federated learning: Decentralized reputation management and performance optimization,” *IEEE Internet of Things Journal*, vol. 12, no. 3, pp. 2890–2905, 2025.
 - [54] Y. Zhao, Y. Qu, Y. Xiang, F. Chen, and L. Gao, “Context-aware consensus algorithm for blockchain-empowered federated learning,” *IEEE Transactions on Cloud Computing*, vol. 12, no. 2, pp. 491–503, 2024.
 - [55] J. Kang, Z. Xiong, D. Niyato, H. Yu, Y.-C. Liang, and D. I. Kim, “Incentive design for efficient federated learning in mobile networks: A contract theory approach,” in *2019 IEEE VTS Asia Pacific Wireless Communications Symposium (APWCS)*, 2019, pp. 1–5.
 - [56] J. Kang, Z. Xiong, D. Niyato, Y. Zou, Y. Zhang, and M. Guizani, “Reliable federated learning for mobile networks,” *IEEE Wireless Communications*, vol. 27, no. 2, pp. 72–80, 2020.
 - [57] X. Deng, J. Li, C. Ma, K. Wei, L. Shi, M. Ding, W. Chen, and H. V. Poor, “Blockchain assisted federated learning over wireless channels: Dynamic resource allocation and client scheduling,” *IEEE Transactions on Wireless Communications*, vol. 22, no. 5, pp. 3537–3553, 2023.
 - [58] L. Witt, U. Zafar, K. Shen, F. Sattler, D. Li, S. Wang, and W. Samek, “Decentralized and incentivized federated learning: A blockchain-enabled framework utilising compressed soft-labels and peer consistency,” *IEEE Transactions on Services Computing*, 2023.
 - [59] Q. Yang, Y. Liu, T. Chen, and Y. Tong, “Federated machine learning: Concept and applications,” *ACM Trans. Intell. Syst. Technol.*, vol. 10, no. 2, Jan. 2019. [Online]. Available: <https://doi.org/10.1145/3298981>
 - [60] T. Alam and R. Gupta, “Federated learning and its role in the privacy preservation of iot devices,” *Future Internet*, vol. 14, no. 9, 2022. [Online]. Available: <https://www.mdpi.com/1999-5903/14/9/246>
 - [61] X. Xie, C. Hu, H. Ren, and J. Deng, “A survey on vulnerability of federated learning: A learning algorithm perspective,” *Neurocomputing*, p. 127225, 2024.

- [62] J. Konečný, H. B. McMahan, F. X. Yu, P. Richtárik, A. T. Suresh, and D. Bacon, "Federated learning: Strategies for improving communication efficiency," *arXiv preprint arXiv:1610.05492*, 2016.
- [63] S. Savazzi, M. Nicoli, and V. Rampa, "Federated learning with cooperating devices: A consensus approach for massive iot networks," *IEEE Internet of Things Journal*, vol. 7, no. 5, pp. 4641–4654, 2020.
- [64] Z. Yang, M. Chen, W. Saad, C. S. Hong, and M. Shikh-Bahaei, "Energy efficient federated learning over wireless communication networks," *IEEE Transactions on Wireless Communications*, vol. 20, no. 3, pp. 1935–1949, 2021.
- [65] M. J. Sheller, B. Edwards, G. A. Reina, J. Martin, S. Pati, A. Kotrotsou, M. Milchenko, W. Xu, D. Marcus, R. R. Colen *et al.*, "Federated learning in medicine: facilitating multi-institutional collaborations without sharing patient data," *Scientific reports*, vol. 10, no. 1, p. 12598, 2020.
- [66] J. Le, X. Lei, N. Mu, H. Zhang, K. Zeng, and X. Liao, "Federated continuous learning with broad network architecture," *IEEE Transactions on Cybernetics*, vol. 51, no. 8, pp. 3874–3888, 2021.
- [67] L. G. F. da Silva, D. F. Sadok, and P. T. Endo, "Resource optimizing federated learning for use with iot: A systematic review," *Journal of Parallel and Distributed Computing*, vol. 175, pp. 92–108, 2023.
- [68] M. Doudou, D. Djenouri, J. M. Barcelo, and N. Badache, "Delay-efficient mac protocol with traffic differentiation and run-time parameter adaptation for energy-constrained wireless sensor networks," *Wireless networks*, vol. 22, no. 2, pp. 467–490, 02 2016. [Online]. Available: <http://link.springer.com/article/10.1007/s11276-015-0965-5>
- [69] T. M. Mengistu, T. Kim, and J.-W. Lin, "A survey on heterogeneity taxonomy, security and privacy preservation in the integration of iot, wireless sensor networks and federated learning," *Sensors*, vol. 24, no. 3, p. 968, 2024.
- [70] M. G. Herabad, "Communication-efficient semi-synchronous hierarchical federated learning with balanced training in heterogeneous iot edge environments," *Internet of Things*, vol. 21, p. 100642, 2023.
- [71] C. Xu, Y. Qu, Y. Xiang, and L. Gao, "Asynchronous federated learning on heterogeneous devices: A survey," *Computer Science Review*, vol. 50, p. 100595, 2023.
- [72] K. Pfeiffer, M. Rapp, R. Khalili, and J. Henkel, "Federated learning for computationally constrained heterogeneous devices: A survey," *ACM Computing Surveys*, vol. 55, no. 14s, pp. 1–27, 2023.
- [73] J. Pei, S. Li, Z. Yu, L. Ho, W. Liu, and L. Wang, "Federated learning encounters 6g wireless communication in the scenario of internet of things," *IEEE Communications Standards Magazine*, vol. 7, no. 1, pp. 94–100, 2023.
- [74] B. Farahani and A. K. Monsefi, "Smart and collaborative industrial iot: A federated learning and data space approach," *Digital Communications and Networks*, vol. 9, no. 2, pp. 436–447, 2023.
- [75] Z. Tang, Y. Wang, and T.-H. Chang, "z-signfedavg: A unified stochastic sign-based compression for federated learning," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 38, no. 14, 2024, pp. 15 301–15 309.
- [76] A. El Ouadrhiri and A. Abdelhadi, "Differential privacy for deep and federated learning: A survey," *IEEE access*, vol. 10, pp. 22 359–22 380, 2022.
- [77] X. Zhang, Y. Kang, K. Chen, L. Fan, and Q. Yang, "Trading off privacy, utility, and efficiency in federated learning," vol. 14, no. 6, nov 2023. [Online]. Available: <https://doi.org/10.1145/3595185>
- [78] J. Geiping, H. Bauermeister, H. Dröge, and M. Moeller, "Inverting gradients-how easy is it to break privacy in federated learning?" *Advances in neural information processing systems*, vol. 33, pp. 16 937–16 947, 2020.
- [79] H. Hu, Z. Salic, L. Sun, G. Dobbie, and X. Zhang, "Source inference attacks in federated learning," in *2021 IEEE International Conference on Data Mining (ICDM)*. IEEE, 2021, pp. 1102–1107.
- [80] M. S. Jere, T. Farnan, and F. Koushanfar, "A taxonomy of attacks on federated learning," *IEEE Security Privacy*, vol. 19, no. 2, pp. 20–28, 2021.
- [81] M. Naseri, J. Hayes, and E. De Cristofaro, "Local and central differential privacy for robustness and privacy in federated learning," *arXiv preprint arXiv:2009.03561*, 2020.
- [82] M. Song, Z. Wang, Z. Zhang, Y. Song, Q. Wang, J. Ren, and H. Qi, "Analyzing user-level privacy attack against federated learning," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 10, pp. 2430–2444, 2020.
- [83] O. Thakkar, S. Ramaswamy, R. Mathews, and F. Beaufays, "Understanding unintended memorization in federated learning," *arXiv preprint arXiv:2006.07490*, 2020.
- [84] W. Wei, L. Liu, M. Loper, K.-H. Chow, M. E. Gursoy, S. Truex, and Y. Wu, "A framework for evaluating gradient leakage attacks in federated learning," 2020.
- [85] M. Nasr, R. Shokri, and A. Houmansadr, "Comprehensive privacy analysis of deep learning: Passive and active white-box inference attacks against centralized and federated learning," in *2019 IEEE Symposium on security and privacy (SP)*. IEEE, 2019, pp. 739–753.
- [86] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *2017 IEEE Symposium on Security and Privacy (SP)*, 2017, pp. 3–18.
- [87] H. Masuda, K. Kita, Y. Koizumi, J. Takemasa, and T. Hasegawa, "Model fragmentation, shuffle and aggregation to mitigate model inversion in federated learning," in *2021 IEEE International Symposium on Local and Metropolitan Area Networks (LANMAN)*, 2021, pp. 1–6.
- [88] Z. Wang, Y. Huang, M. Song, L. Wu, F. Xue, and K. Ren, "Poisoning-assisted property inference attack against federated learning," *IEEE Transactions on Dependable and Secure Computing*, 2022.
- [89] C. Chen, L. Lyu, H. Yu, and G. Chen, "Practical attribute reconstruction attack against federated learning," *IEEE Transactions on Big Data*, 2022.
- [90] P. Voigt and A. Von dem Bussche, "The eu general data protection regulation (gdpr)," *A Practical Guide, 1st Ed., Cham: Springer International Publishing*, vol. 10, no. 3152676, pp. 10–5555, 2017.
- [91] V. Mothukuri, P. Khare, R. M. Parizi, S. Pouriyeh, A. Dehghantanha, and G. Srivastava, "Federated-learning-based anomaly detection for iot security attacks," *IEEE Internet of Things Journal*, vol. 9, no. 4, pp. 2545–2554, 2021.
- [92] J. Zhang, H. Zhu, F. Wang, J. Zhao, Q. Xu, and H. Li, "Security and privacy threats to federated learning: Issues, methods, and challenges," *Security and Communication Networks*, vol. 2022, no. 1, p. 2886795, 2022.
- [93] H. Ren, J. Deng, and X. Xie, "Grnn: generative regression neural network—a data leakage attack for federated learning," *ACM Transactions on Intelligent Systems and Technology (TIST)*, vol. 13, no. 4, pp. 1–24, 2022.
- [94] L. Zhang, Y. Zhang, Q. Wu, Y. Mu, and F. Rezaeiabgha, "A secure and efficient decentralized access control scheme based on blockchain for vehicular social networks," *IEEE Internet of Things Journal*, pp. 1–1, 2022.
- [95] X. Cheng, W. Tian, F. Shi, M. Zhao, S. Chen, and H. Wang, "A blockchain-empowered cluster-based federated learning model for blade icing estimation on iot-enabled wind turbine," *IEEE Transactions on Industrial Informatics*, pp. 1–1, 2022.
- [96] A. Riahi Sfar, E. Natalizio, S. Mazlout, Y. Challal, and Z. Chtourou, "Privacy preservation using game theory in e-health application," *Journal of Information Security and Applications*, vol. 66, p. 103158, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2214212622000461>
- [97] B. D. Deebak, F. H. Memon, K. Dev, S. A. Khowaja, W. Wang, and N. M. F. Qureshi, "Tab-sapp: A trust-aware blockchain-based seamless authentication for massive iot-enabled industrial applications," *IEEE Transactions on Industrial Informatics*, pp. 1–1, 2022.
- [98] A. Fathalizadeh, V. Moghtadaiee, and M. Alishahi, "On the privacy protection of indoor location dataset using anonymization," *Computers & Security*, p. 102665, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0167404822000645>
- [99] K. Kita, Y. Koizumi, and T. Hasegawa, "Private retrieval of location-related content using k-anonymity and application to icn," *Computer Networks*, p. 108908, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1389128622001013>
- [100] W. Xue, Y. Shen, C. Luo, W. Xu, W. Hu, and A. Seneviratne, "A differential privacy-based classification system for edge computing in iot," *Computer Communications*, vol. 182, pp. 117–128, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0140366421004187>
- [101] Y. Sei and A. Ohsuga, "Private true data mining: Differential privacy featuring errors to manage internet-of-things data," *IEEE Access*, vol. 10, pp. 8738–8757, 2022.
- [102] W. Liu, J. Cheng, X. Wang, X. Lu, and J. Yin, "Hybrid differential privacy based federated learning for internet of things," *Journal of Systems Architecture*, vol. 124, p. 102418, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1383762122000200>
- [103] Z. Shao, L. Ma, Q. Lin, J. Li, M. Gong, and A. K. Nandi, "Pmcdm: Privacy-preserving multiresolution community detection in multiplex networks," *Knowledge-Based Systems*, vol. 244, p. 108542, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0950705122002386>

- [104] Y. Xu, Y. Mao, S. Li, J. Li, and X. Chen, "Privacy-preserving federal learning chain for internet of things," *IEEE Internet of Things Journal*, vol. 10, no. 20, pp. 18 364–18 374, 2023.
- [105] A. Vyas, P.-C. Lin, R.-H. Hwang, and M. Tripathi, "Privacy-preserving federated learning for intrusion detection in iot environments: A survey," *IEEE Access*, 2024.
- [106] Y. Zhao, M. Li, L. Lai, N. Suda, D. Civin, and V. Chandra, "Federated learning with non-iid data," 2018. [Online]. Available: <https://arxiv.org/abs/1806.00582>
- [107] T. Li, A. K. Sahu, A. Talwalkar, and V. Smith, "Federated learning: Challenges, methods, and future directions," *IEEE Signal Processing Magazine*, vol. 37, no. 3, pp. 50–60, may 2020. [Online]. Available: <https://doi.org/10.1109%2Fmmsp.2020.2975749>
- [108] B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y. Arcas, "Communication-Efficient Learning of Deep Networks from Decentralized Data," in *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, ser. Proceedings of Machine Learning Research, A. Singh and J. Zhu, Eds., vol. 54. PMLR, 20–22 Apr 2017, pp. 1273–1282. [Online]. Available: <https://proceedings.mlr.press/v54/mcmahan17a.html>
- [109] A. Z. Tan, H. Yu, L. Cui, and Q. Yang, "Towards personalized federated learning," *IEEE Transactions on Neural Networks and Learning Systems*, 2022.
- [110] Y. Huang, L. Chu, Z. Zhou, L. Wang, J. Liu, J. Pei, and Y. Zhang, "Personalized cross-silo federated learning on non-iid data," in *Proceedings of the AAAI conference on artificial intelligence*, vol. 35, no. 9, 2021, pp. 7865–7873.
- [111] L. Collins, H. Hassani, A. Mokhtari, and S. Shakkottai, "Exploiting shared representations for personalized federated learning," in *International conference on machine learning*. PMLR, 2021, pp. 2089–2099.
- [112] C. T. Dinh, N. Tran, and J. Nguyen, "Personalized federated learning with moreau envelopes," *Advances in Neural Information Processing Systems*, vol. 33, pp. 21 394–21 405, 2020.
- [113] X. Wu, X. Liu, J. Niu, G. Zhu, and S. Tang, "Bold but cautious: Unlocking the potential of personalized federated learning through cautiously aggressive collaboration," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2023, pp. 19 375–19 384.
- [114] X. Cao, G. Sun, H. Yu, and M. Guizani, "Perfed-gan: Personalized federated learning via generative adversarial networks," *IEEE Internet of Things Journal*, vol. 10, no. 5, pp. 3749–3762, 2023.
- [115] C. Li, G. Li, and P. K. Varshney, "Decentralized federated learning via mutual knowledge transfer," *IEEE Internet of Things Journal*, vol. 9, no. 2, pp. 1136–1147, 2022.
- [116] T. V. Khoa, D. T. Hoang, N. L. Trung, C. T. Nguyen, T. T. T. Quynh, D. N. Nguyen, N. V. Ha, and E. Dutkiewicz, "Deep transfer learning: A novel collaborative learning model for cyberattack detection systems in iot networks," *IEEE Internet of Things Journal*, vol. 10, no. 10, pp. 8578–8589, 2023.
- [117] J. Zhang, C. Luo, M. Carpenter, and G. Min, "Federated learning for distributed iiot intrusion detection using transfer approaches," *IEEE Transactions on Industrial Informatics*, vol. 19, no. 7, pp. 8159–8169, 2023.
- [118] W. Dai, Q. Yang, G.-R. Xue, and Y. Yu, "Boosting for transfer learning," in *Proceedings of the 24th International Conference on Machine Learning*, ser. ICML '07. New York, NY, USA: Association for Computing Machinery, 2007, p. 193–200. [Online]. Available: <https://doi.org/10.1145/1273496.1273521>
- [119] V. Smith, C.-K. Chiang, M. Sanjabi, and A. S. Talwalkar, "Federated multi-task learning," in *Advances in Neural Information Processing Systems*, vol. 30. Curran Associates, Inc., 2017. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2017/file/6211080fa89981f66b1a0c9d55c61d0f-Paper.pdf
- [120] J. Mills, J. Hu, and G. Min, "Multi-task federated learning for personalised deep neural networks in edge computing," *IEEE Transactions on Parallel and Distributed Systems*, vol. 33, no. 3, pp. 630–641, 2022.
- [121] C. Zhao, Z. Gao, Q. Wang, K. Xiao, and Z. Mo, "Afi: An adaptively federated multitask learning for model sharing in industrial iot," *IEEE Internet of Things Journal*, vol. 9, no. 18, pp. 17 080–17 088, 2022.
- [122] W. Hao, M. El-Khamy, J. Lee, J. Zhang, K. J. Liang, C. Chen, and L. C. Duke, "Towards fair federated learning with zero-shot data augmentation," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) Workshops*, June 2021, pp. 3310–3319.
- [123] H. Zhang, Q. Hou, T. Wu, S. Cheng, and J. Liu, "Data-augmentation-based federated learning," *IEEE Internet of Things Journal*, vol. 10, no. 24, pp. 22 530–22 541, 2023.
- [124] B. Weinger, J. Kim, A. Sim, M. Nakashima, N. Moustafa, and K. J. Wu, "Enhancing iot anomaly detection performance for federated learning," *Digital Communications and Networks*, vol. 8, no. 3, pp. 314–323, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2352864822000190>
- [125] Y. Deng, F. Lyu, J. Ren, H. Wu, Y. Zhou, Y. Zhang, and X. Shen, "Auction: Automated and quality-aware client selection framework for efficient federated learning," *IEEE Transactions on Parallel and Distributed Systems*, vol. 33, no. 8, pp. 1996–2009, 2022.
- [126] Z. Li, Y. He, H. Yu, J. Kang, X. Li, Z. Xu, and D. Niyato, "Data heterogeneity-robust federated learning via group client selection in industrial iot," *IEEE Internet of Things Journal*, vol. 9, no. 18, pp. 17 844–17 857, 2022.
- [127] M. A. Serhani, H. G. Abreha, A. Tariq, M. Hayajneh, Y. Xu, and K. Hayawi, "Dynamic data sample selection and scheduling in edge federated learning," *IEEE Open Journal of the Communications Society*, vol. 4, pp. 2133–2149, 2023.
- [128] H. Sun, L. Shen, S. Chen, J. Sun, J. Li, G. Sun, and D. Tao, "Fedlral: Client-specific adaptive learning rates achieve linear speedup for non-iid data," *arXiv preprint arXiv:2309.09719*, 2023.
- [129] Z. Ma, Y. Xu, H. Xu, Z. Meng, L. Huang, and Y. Xue, "Adaptive batch size for federated learning in resource-constrained edge computing," *IEEE Transactions on Mobile Computing*, vol. 22, no. 1, pp. 37–53, 2023.
- [130] L. Zhang, J. Xu, P. Vijayakumar, P. K. Sharma, and U. Ghosh, "Homomorphic encryption-based privacy-preserving federated learning in iot-enabled healthcare system," *IEEE Transactions on Network Science and Engineering*, 2022.
- [131] X. Wang, J. Hu, H. Lin, W. Liu, H. Moon, and M. J. Piran, "Federated learning-empowered disease diagnosis mechanism in the internet of medical things: From the privacy-preservation perspective," *IEEE Transactions on Industrial Informatics*, vol. 19, no. 7, pp. 7905–7913, 2023.
- [132] C. Zhou, F. Liu, and J. Xiong, "Privacy-preserving federated learning in fog computing," *IEEE Transactions on Industrial Informatics*, vol. 16, no. 7, pp. 5622–5633, 2020.
- [133] W. Zhao, K. Zhang, and J. Zheng, "Enhancing smart home security with blockchain and differential privacy in federated learning," *IEEE Internet of Things Journal*, vol. 8, no. 6, pp. 4180–4190, 2021.
- [134] V. Stephanie, I. Khalil, M. Atiquzzaman, and X. Yi, "Trustworthy privacy-preserving hierarchical ensemble and federated learning in healthcare 4.0 with blockchain," *IEEE Transactions on Industrial Informatics*, vol. 19, no. 7, pp. 7936–7945, 2023.
- [135] H. Li, Y. Sun, Y. Yu, D. Li, Z. Guan, and J. Liu, "Privacy-preserving cross-silo federated learning atop blockchain for iot," *IEEE Internet of Things Journal*, pp. 1–1, 2023.
- [136] E. Moore, A. Imteaj, S. Rezapour, and M. H. Amini, "A survey on secure and private federated learning using blockchain: Theory and application in resource-constrained computing," *IEEE Internet of Things Journal*, 2023.
- [137] Q. Xie, S. Jiang, L. Jiang, Y. Huang, Z. Zhao, S. Khan, W. Dai, Z. Liu, and K. Wu, "Efficiency optimization techniques in privacy-preserving federated learning with homomorphic encryption: A brief survey," *IEEE Internet of Things Journal*, vol. 11, no. 14, pp. 24 569–24 580, 2024.
- [138] P. Prakash, J. Ding, R. Chen, X. Qin, M. Shu, Q. Cui, Y. Guo, and M. Pan, "Iot device friendly and communication-efficient federated learning via joint model pruning and quantization," *IEEE Internet of Things Journal*, vol. 9, no. 15, pp. 13 638–13 650, 2022.
- [139] T. Zhang, C. He, T. Ma, L. Gao, M. Ma, and S. Avestimehr, "Federated learning for internet of things," in *Proceedings of the 19th ACM Conference on Embedded Networked Sensor Systems*, ser. SenSys '21. New York, NY, USA: Association for Computing Machinery, 2021, p. 413–419. [Online]. Available: <https://doi.org/10.1145/3485730.3493444>
- [140] S. Abdulrahman, H. Tout, A. Mourad, and C. Talhi, "Fedmccs: Multi-criteria client selection model for optimal iot federated learning," *IEEE Internet of Things Journal*, vol. 8, no. 6, pp. 4723–4735, 2021.
- [141] A. A. Abdellatif, N. Mhaisen, A. Mohamed, A. Erbad, M. Guizani, Z. Dawy, and W. Nasreddine, "Communication-efficient hierarchical federated learning for iot heterogeneous systems with imbalanced data," *Future Generation Computer Systems*, vol. 128, pp. 406–419, 2022. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0167739X2100412X>
- [142] L. U. Khan, S. R. Pandey, N. H. Tran, W. Saad, Z. Han, M. N. H. Nguyen, and C. S. Hong, "Federated learning for edge networks: Re-

- source optimization and incentive mechanism,” *IEEE Communications Magazine*, vol. 58, no. 10, pp. 88–93, 2020.
- [143] X. Mo and J. Xu, “Energy-efficient federated edge learning with joint communication and computation design,” *Journal of Communications and Information Networks*, vol. 6, no. 2, pp. 110–124, 2021.
- [144] A. Fu, X. Zhang, N. Xiong, Y. Gao, H. Wang, and J. Zhang, “Vfl: A verifiable federated learning with privacy-preserving for big data in industrial iot,” *IEEE Transactions on Industrial Informatics*, vol. 18, no. 5, pp. 3316–3326, 2022.
- [145] Q. Kong, F. Yin, R. Lu, B. Li, X. Wang, S. Cui, and P. Zhang, “Privacy-preserving aggregation for federated learning-based navigation in vehicular fog,” *IEEE Transactions on Industrial Informatics*, vol. 17, no. 12, pp. 8453–8463, 2021.
- [146] R. Kanagavelu, Z. Li, J. Samsudin, Y. Yang, F. Yang, R. S. Mong Goh, M. Cheah, P. Wiwatphonthana, K. Akkarajitsakul, and S. Wang, “Two-phase multi-party computation enabled privacy-preserving federated learning,” in *2020 20th IEEE/ACM International Symposium on Cluster, Cloud and Internet Computing (CCGRID)*, 2020, pp. 410–419.
- [147] S. Chen, D. Yu, Y. Zou, J. Yu, and X. Cheng, “Decentralized wireless federated learning with differential privacy,” *IEEE Transactions on Industrial Informatics*, vol. 18, no. 9, pp. 6273–6282, 2022.
- [148] L. Yin, J. Feng, H. Xun, Z. Sun, and X. Cheng, “A privacy-preserving federated learning for multiparty data sharing in social iots,” *IEEE Transactions on Network Science and Engineering*, vol. 8, no. 3, pp. 2706–2718, 2021.
- [149] Y. Qu, L. Gao, T. H. Luan, Y. Xiang, S. Yu, B. Li, and G. Zheng, “Decentralized privacy using blockchain-enabled federated learning in fog computing,” *IEEE Internet of Things Journal*, vol. 7, no. 6, pp. 5171–5183, 2020.
- [150] M. Hao, H. Li, X. Luo, G. Xu, H. Yang, and S. Liu, “Efficient and privacy-enhanced federated learning for industrial artificial intelligence,” *IEEE Transactions on Industrial Informatics*, vol. 16, no. 10, pp. 6532–6542, 2020.
- [151] K. Zhang, K. Chen, and J. Liu, “Secure crowdsourcing aggregation for federated learning,” *IEEE Internet of Things Journal*, vol. 7, no. 10, pp. 9550–9561, 2020.
- [152] M. Koppula and L. J. L.M.I., “A real time dataset ‘idsiot2024,’” 2024. [Online]. Available: <https://dx.doi.org/10.21227/gfaz-t124>
- [153] G. Ke, Q. Meng, T. Finley, T. Wang, W. Chen, W. Ma, Q. Ye, and T.-Y. Liu, “Lightgbm: A highly efficient gradient boosting decision tree,” in *Advances in Neural Information Processing Systems*, I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, Eds., vol. 30. Curran Associates, Inc., 2017. [Online]. Available: https://proceedings.neurips.cc/paper_files/paper/2017/file/6449f44a102fde848669bdd9eb6b76fa-Paper.pdf
- [154] D. Niyato, P. Wang, Y. Liang, D. I. Kim, and Z. Han, “Federated learning in mobile edge computing: A comprehensive survey,” *IEEE Communications Surveys Tutorials*, vol. 23, no. 3, pp. 1622–1658, 2021.
- [155] W. Saad, M. Bennis, and M. Chen, “A vision of 6g wireless systems: Applications, trends, technologies, and open research problems,” *IEEE Network*, vol. 34, no. 3, pp. 134–142, 2020.
- [156] H. Tataria, M. Shafi, A. F. Molisch, M. Dohler, H. Sjöland, and F. Tufvesson, “6g wireless systems: Vision, requirements, challenges, insights, and opportunities,” *Proceedings of the IEEE*, vol. 109, no. 7, pp. 1166–1199, 2021.
- [157] S. Abdelwahab, B. Hamdaoui, M. Guizani, and A. Rayes, “Edge computing: A primer,” *IEEE Internet of Things Journal*, vol. 9, no. 1, pp. 573–590, 2022.
- [158] R. Shafin, L. Liu, M. R. Khandaker, K.-K. Wong, Z. Lin, G. Zhang, and Z. Han, “6g vision: An ai-driven decentralized and software-defined paradigm,” *IEEE Communications Magazine*, vol. 58, no. 8, pp. 28–34, 2020.
- [159] D. Javeed, M. S. Saeed, I. Ahmad, M. Adil, P. Kumar, and A. N. Islam, “Quantum-empowered federated learning and 6g wireless networks for iot security: Concept, challenges and future directions,” *Future Generation Computer Systems*, vol. 160, pp. 577–597, 2024. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0167739X24003236>
- [160] T. Zhu, Q. Zhang, H. Song, B. Fang, and T. Wang, “Secure federated learning for iot networks,” *IEEE Wireless Communications*, vol. 28, no. 5, pp. 50–56, 2021.
- [161] M. Ye, X. Fang, B. Du, P. C. Yuen, and D. Tao, “Heterogeneous federated learning: State-of-the-art and research challenges,” *ACM Comput. Surv.*, vol. 56, no. 3, Oct. 2023. [Online]. Available: <https://doi.org/10.1145/3625558>
- [162] Z. Lin, Z. Chen, Z. Fang, X. Chen, X. Wang, and Y. Gao, “Fedsn: A federated learning framework over heterogeneous leo satellite networks,” *IEEE Transactions on Mobile Computing*, pp. 1–15, 2024.
- [163] X. Liu, T. Ratnarajah, M. Sellathurai, and Y. C. Eldar, “Adaptive model pruning and personalization for federated learning over wireless networks,” *IEEE Transactions on Signal Processing*, vol. 72, pp. 4395–4411, 2024.
- [164] S. M. S. Mohammadabadi, S. Zawad, F. Yan, and L. Yang, “Speed up federated learning in heterogeneous environments: A dynamic tiering approach,” *IEEE Internet of Things Journal*, pp. 1–1, 2024.
- [165] Y. Zhang, H. Xia, S. Xu, X. Wang, and L. Xu, “Adaptfl: Adaptive federated learning framework for heterogeneous devices,” *Future Generation Computer Systems*, p. 107610, 2024. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S0167739X24005740>
- [166] D. Gao, H. Wang, X. Guo, L. Wang, G. Gui, W. Wang, Z. Yin, S. Wang, Y. Liu, and T. He, “Federated learning based on ctc for heterogeneous internet of things,” *IEEE Internet of Things Journal*, vol. 10, no. 24, pp. 22 673–22 685, 2023.
- [167] C. Feng, D. Feng, G. Huang, Z. Liu, Z. Wang, and X.-G. Xia, “Robust privacy-preserving recommendation systems driven by multimodal federated learning,” *IEEE Transactions on Neural Networks and Learning Systems*, 2024.
- [168] Z. Liu, Y. Jiang, J. Shen, M. Peng, K.-Y. Lam, X. Yuan, and X. Liu, “A survey on federated unlearning: Challenges, methods, and future directions,” *ACM Comput. Surv.*, vol. 57, no. 1, Oct. 2024. [Online]. Available: <https://doi.org/10.1145/3679014>
- [169] Z. Liu, Y. Jiang, W. Jiang, J. Guo, J. Zhao, and K.-Y. Lam, “Guaranteeing data privacy in federated unlearning with dynamic user participation,” *arXiv preprint arXiv:2406.00966*, 2024.
- [170] X. Zuo, M. Wang, T. Zhu, L. Zhang, S. Yu, and W. Zhou, “Federated learning with blockchain-enhanced machine unlearning: A trustworthy approach,” *arXiv preprint arXiv:2405.20776*, 2024.
- [171] R. Neisse, G. Steri, and I. Nai-Fovino, “Towards a framework for addressing collusion attacks in federated learning,” *IEEE Transactions on Information Forensics and Security*, vol. 15, pp. 3213–3222, 2020.
- [172] Y. Lu, T. Huang, B. Dai, and K. Liang, “Federated learning: Challenges, methods, and future directions,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 31, no. 8, pp. 2642–2652, 2020.
- [173] T. Li, A. K. Sahu, A. Talwalkar, and V. Smith, “A comprehensive survey on federated learning: Opportunities and challenges,” *Journal of Machine Learning Research*, vol. 22, no. 1, pp. 1–42, 2021.
- [174] S. Wang, Y. Li, W. Cao, J. Zhang, and X. Lin, “Blockchain for federated learning: A data privacy perspective,” *IEEE Network*, vol. 34, no. 4, pp. 70–76, 2020.
- [175] M. Ali, H. Karimipour, and M. Tariq, “Integration of blockchain and federated learning for internet of things: Recent advances and future challenges,” *Computers & Security*, vol. 108, p. 102355, 2021.
- [176] G. T. Nguyen, T.-D. Hoang, P. N. Pathirana, and A. Seneviratne, “Interoperability in iot: A review,” *IEEE Internet of Things Journal*, vol. 8, no. 1, pp. 948–963, 2021.
- [177] A. Yousefpour, R. Ishii, D. Choffnes, P. Shenoy, A. Bourgeois, and A. Mislove, “A survey on federated learning systems: Vision, enabling technologies, and future directions,” *IEEE Communications Surveys Tutorials*, vol. 23, no. 2, pp. 1466–1497, 2021.
- [178] P. Vepakomma, O. Gupta, T. Swedish, and R. Raskar, “Federated learning for healthcare: Privacy-aware data integration for medical applications,” *arXiv preprint arXiv:1812.02903*, 2018.
- [179] C. Rudin, “Stop explaining black box machine learning models for high-stakes decisions and use interpretable models instead,” *Nature Machine Intelligence*, vol. 1, pp. 206–215, 2019.
- [180] H. Zhu, H. Jin, and J. Luo, “Reinforcement learning for federated learning: A framework for client selection and resource allocation,” *IEEE Internet of Things Journal*, vol. 7, no. 10, pp. 9275–9284, 2020.
- [181] B. McMahan and D. Ramage, “Federated learning: Collaborative machine learning without centralized training data,” *Google Research Blog*, vol. 3, 2017.
- [182] Z. Xu, Y. Zhang, G. Andrew, C. A. Choquette-Choo, P. Kairouz, H. B. McMahan, J. Rosenstock, and Y. Zhang, “Federated learning of gboard language models with differential privacy,” *arXiv preprint arXiv:2305.18465*, 2023.
- [183] J. Omhover, “Federated learning with azure machine learning: Powering privacy-preserving innovation in ai,” 2023, <https://techcommunity.microsoft.com/blog/machinelearningblog/federated-learning-with-azure-machine-learning-powering-privacy-preserving-innov/3824720>.

- [184] D. M. Diaz, A. Manoel, J. Chen, N. Singal, and R. Sim, "Project florida: Federated learning made easy," *arXiv preprint arXiv:2307.11899*, 2023.
- [185] NVIDIA, "Federated learning for healthcare using nvidia clara, white paper," 2021, <https://developer.download.nvidia.com/CLARA/Federated-Learning-Training-for-Healthcare-Using-NVIDIA-Clara.pdf>.
- [186] M. Pelikan, S. S. Azam, V. Feldman, J. Silovsky, K. Talwar, T. Likhomanenko *et al.*, "Federated learning with differential privacy for end-to-end speech recognition," *arXiv preprint arXiv:2310.00098*, 2024.
- [187] V. B. Bivin, "Case study - federated privacy preserving analytics for secure collaboration among telco and partners to improve customer engagement," 2022.
- [188] B. Stojkovic, J. Woodbridge, Z. Fang, J. Cai, A. Petrov, S. Iyer, D. Huang, P. Yau, A. S. Kumar, H. Jawa *et al.*, "Applied federated learning: Architectural design for robust and efficient learning in privacy aware settings," *arXiv preprint arXiv:2206.00807*, 2022.
- [189] Intel, "Case study: Secure federated learning for a better world," 2022, https://download.intel.com/newsroom/2022/security/secure-federated-learning_tech-brief.pdf.
- [190] H. Ludwig, N. Baracaldo, G. Thomas, Y. Zhou, A. Anwar, S. Rajamoni, Y. Ong, J. Radhakrishnan, A. Verma, M. Sinn *et al.*, "Ibm federated learning: an enterprise framework white paper v0. 1," *arXiv preprint arXiv:2007.10987*, 2020.
- [191] D. J. Beutel, T. Topal, A. Mathur, X. Qiu, J. Fernandez-Marques, Y. Gao, L. Sani, K. H. Li, T. Parcollet, P. P. B. de Gusmão *et al.*, "Flower: A friendly federated learning research framework," *arXiv preprint arXiv:2007.14390*, 2020.
- [192] H. Daga, J. Shin, D. Garg, A. Gavrilovska, M. Lee, and R. R. Kompella, "Flame: Simplifying topology extension in federated learning," in *Proceedings of the 2023 ACM Symposium on Cloud Computing*, 2023, pp. 341–357.
- [193] C. He, S. Li, J. So, X. Zeng, M. Zhang, H. Wang, X. Wang, P. Vepakomma, A. Singh, H. Qiu *et al.*, "Fedml: A research library and benchmark for federated machine learning," *arXiv preprint arXiv:2007.13518*, 2020.
- [194] A. R. Group, "Reminder: Privacy-preserving machine learning through secure management of data's lifecycle in distributed systems," 2024, <https://ants.inf.um.es/en/reminder>.
- [195] M. Oldenhof, G. Ács, B. Pejó, A. Schuffenhauer, N. Holway, N. Sturm, A. Dieckmann, O. Fortmeier, E. Boniface, C. Mayer *et al.*, "Industry-scale orchestrated federated learning for drug discovery," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, no. 13, 2023, pp. 15 576–15 584.