

The Internet of Things in the Era of Generative AI: Vision and Challenges

Xin Wang  and Zhongwei Wan , *The Ohio State University, USA, OH 43210*

Arvin Hekmati  and Mingyu Zong , *University of Southern California, USA, CA 90089*

Samiul Alam  and Mi Zhang , *The Ohio State University, USA, OH 43210*

Bhaskar Krishnamachari , *University of Southern California, USA, CA 90089*

Abstract—Advancements in generative AI hold immense promise to push the Internet of Things (IoT) to the next level. In this article, we share our vision on the IoT in the era of generative AI. We discuss some of the most important applications of generative AI in IoT-related domains. We also identify some of the most critical challenges and discuss current gaps as well as promising opportunities on enabling generative AI for the IoT. We hope this article can inspire new research on the IoT in the era of generative AI.

Keywords: IoT, Generative AI, Edge Computing, AI Agents.

INTRODUCTION

Today, Internet of Things (IoT) such as smartphones, wearables, smart speakers, and household robots have already become an integrated part of our daily lives. These devices can sense, communicate, and are empowered by modern artificial intelligence (AI) techniques.

Advancements in Generative AI¹ have enabled a new wave of AI revolution. Generative AI refers to AI models that can generate new content in the form of text, images, videos, codes, and many more. While Generative AI is not new, it is only until recently that large-scale generative models exemplified by Large Language Models (LLMs) (e.g., GPT, LLaMA, and Gemini)² and Multimodal Generative Models (e.g., GPT-4V, DALL-E, and Stable Diffusion)³ have made the breakthrough. Such breakthrough comes from their significantly large model sizes while being pre-trained

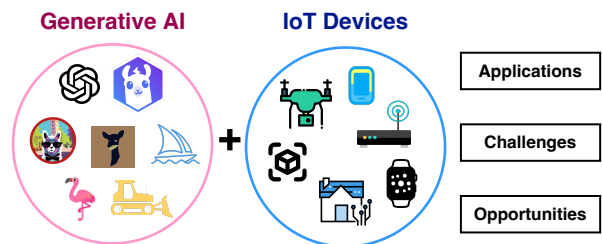


FIGURE 1. IoT in the era of Generative AI.

on massive amounts of data. These characteristics enable Generative AI to generate high-quality data, tackle complex tasks with human-level performance, and exhibit superior generalization ability on new tasks, all of which were not attainable before.

The implications of the advancements of Generative AI for IoT are profound. The distinctive capabilities of Generative AI bring pivotal benefits across the entire IoT pipeline, encompassing data generation, data processing, interfacing with IoT devices, and IoT system development and evaluation. These benefits position Generative AI as having substantial potential to revolutionize a wide range of IoT applications such as mobile networks, autonomous driving, metaverse, robotics, healthcare, and cybersecurity. At the same time, turning these applications into reality is, however, not trivial. Innovative techniques are needed to address

some of the most formidable challenges so as to realize the full potential of Generative AI for IoT.

In this article, we provide our vision and insights on the applications, challenges, and opportunities of what Generative AI brings to IoT (Figure 1). We start by explaining how Generative AI could benefit some of the most important IoT applications. Next, we discuss some of the most critical challenges that serve as impediments to enabling Generative AI for IoT, and share our views on the gaps as well as promising opportunities to address those challenges. We hope this article can act as a catalyst to inspire new research on IoT in the era of Generative AI.

APPLICATIONS OF GENERATIVE AI IN IOT-RELATED DOMAINS

Leveraging its distinctive capabilities, Generative AI holds the potential to revolutionize numerous critical IoT applications. In this section, we delve into a number of application domains (Figure 2) where Generative AI has already left its mark and others where its potential is just beginning to be recognized.

Mobile Networks

Generative AI has the considerable potential to revolutionize the design and operation of mobile networks⁴. For instance, Generative AI can generate simulations based on historical mobile network data to help network operators foresee potential bottlenecks and adjust resources dynamically. Moreover, Generative AI can create highly efficient codecs, resulting in smaller sizes of data to be exchanged between nodes to increase communication efficiency in mobile networks.

Autonomous Vehicles

The transformational journey of the automobile industry towards autonomous vehicles has been deeply influenced by Generative AI as well⁵. For example, AI agents empowered by LLMs like Grok can be integrated into vehicle systems such as Tesla Model-3 to enhance user interfaces and improve communication between the vehicle and its occupants, making interactions more responsive and user-friendly. Meanwhile, multimodal generative models are crucial for the development of the advanced driver-assistance systems (ADAS) in autonomous vehicles. ADAS will provide highly accurate prediction of traffic conditions, pedestrian behavior, and potential hazards, allowing for safer and more reliable vehicle automation. Furthermore, by generating realistic driving scenarios during

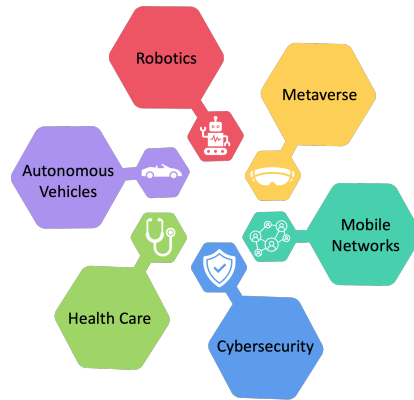


FIGURE 2. Representative application domains of Generative AI for IoT.

the testing phase, multimodal generative models will enable manufacturers to refine vehicle responses under various conditions, which significantly accelerates the safe deployment of autonomous vehicles.

Metaverse

Generative AI can significantly enhance the Metaverse by leveraging its ability to visualize and simulate based on multimodal sensor data, generating a vivid and immersive virtual realm⁶. Moreover, utilizing generative models to adeptly handle data from various sensors and human inputs can construct dynamic, responsive environments that closely mimic the real world or create entirely new, fantastical settings. For instance, multimodal generative models such as GPT-4V are capable of generating simulations based on user interactions and environmental changes captured through IoT devices such as smart glasses and head-mounted devices, enabling real-time adjustments and enhancements to the virtual landscape. This capability allows for a seamless and adaptive user experience to make the Metaverse not just a static backdrop but a living, evolving entity.

Robotics

In the rapidly evolving field of robotics, the integration of Generative AI has emerged as a cornerstone, significantly enhancing the capabilities and intelligence of robotic systems⁷. For example, LLMs can be utilized to enhance the interaction capabilities of robots, enabling them to understand dialogues from users and generate human-friendly responses in real-time. Additionally, multimodal generative models which can effectively process and integrate visual and auditory information,

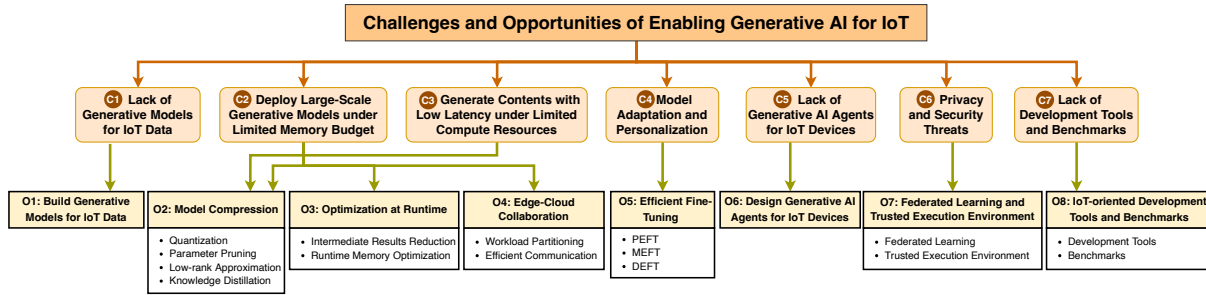


FIGURE 3. Challenges and opportunities of enabling Generative AI for IoT.

are crucial for robots to better understand their physical environment and context, allowing robots to adapt to new situations with greater autonomy. In doing so, these powerful generative models help robots learn from their surroundings and make informed decisions, paving the way for more adaptive and intelligent robotic systems in various settings.

Health Care

The healthcare sector, empowered by IoT devices, is experiencing a transformative paradigm shift with the integration of Generative AI to revolutionize patient care⁸. For example, LLMs can be used to automatically process and understand medical literature, and draft patient reports and personalized treatment plans based on patient history. Meanwhile, multimodal generative models are instrumental in analyzing diverse data modalities collected from IoT devices such as electrocardiogram (ECG), blood pressure monitor, and pulse oximeter; they are also capable of translating those raw sensor data into human-understandable reports⁹. These models altogether can integrate text data from electronic health records, numerical data from monitors, and visual data from scans to provide a comprehensive understanding of a patient's health. This integration allows for delivering more accurate diagnoses, preventive care, and tailored treatment regimens, thereby significantly enhancing patient outcomes and operational efficiencies in medical facilities.

Cybersecurity

IoT devices are particularly vulnerable to cyber-attacks due to their widespread deployment and the often minimal built-in security features. Generative AI can be employed to enhance security measures and address their security challenge¹⁰. For instance, LLMs can analyze and generate security protocols by learning from vast databases of cybersecurity incidents

and responses. Similarly, multimodal generative models can integrate and analyze data from multiple sources including network traffic, user behavior, and anomaly detection systems, can predict potential security breaches before they occur. These generative models are trained to identify patterns indicative of cyber threats, enabling proactive security measures and automatic updates to defense strategies, thus protecting the IoT devices against a wide range of cyber threats.

CHALLENGES AND OPPORTUNITIES OF ENABLING GENERATIVE AI FOR IOT

Turning these applications into reality is, however, not trivial. We have identified five challenges that act as key barriers to realizing Generative AI for IoT. In this section, we describe these challenges, and share our perspectives on promising opportunities to address these challenges (Figure 3).

Challenges

C1: Lack of Generative Models for IoT Data. Generative models today are predominantly developed using a few data modalities, including text, images, and videos. However, as summarized in Table 1, IoT-related applications encompass a much wider range of data modalities, including network traffic data, home-deployed sensor data such as temperature and humidity, as well as healthcare data collected from mobile and wearable devices such as heart rate, steps taken, and calories burned. Therefore, building generative models based on IoT-related data modalities is much needed, yet it remains a very challenging task.

C2: Deploy Large-Scale Generative Models under Limited Memory Budget. Generative models in general contain billions of parameters. Such large model

TABLE 1. Representative use cases and data modalities in different IoT application domains.

IoT Application Domain	Representative Use Cases	Data Modalities
Mobile Networks	Network Optimization, Fault Detection	Network Traffic, Network Logs, Multimedia Data
Autonomous Vehicles	Navigation, Adversarial Testing	Lidar, Radar, Point Cloud, GPS Coordinates, Control Signal
Metaverse	Virtual Meetings, Virtual Tourism	Gestures, Voices, Facial Expressions, Eye-Tracking Data
Robotics	Path Planning, Human-Robot Interaction	Images, Audios, Paths, Control Signal, Error Logs
Health Care	Disease Diagnosis, Report Generation	ECG, Blood Pressure, Steps Taken, Calories Burned, Food Intake
Cybersecurity	Incident Response, Malware Detection	Authentication Data, Attack Vectors, Breach Reports

sizes directly translate to their significant demands for memory resources¹¹. For example, LLaMA-7B requires about 14GB memory for storing its 7 billion parameters in half-precision floating-point format (FP16). However, IoT devices are known to be memory-constrained. This discrepancy between intensive memory requirements of generative models and limited memory budgets of IoT devices poses a considerable challenge on the deployment of large-scale generative models on IoT devices.

C3: Generate Contents with Low Latency under Limited Compute Resources. Many applications of Generative AI in IoT-related domains such as autonomous driving, Metaverse, and robotics are latency-sensitive. Generative models often involve computation-intensive operations during content generation. However, IoT devices are limited in their on-board compute resources, making it challenging to generate contents while meeting application-specific latency requirements. For example, Mixture of Experts (MoE) based LLM ST-128 requires six seconds to generate only one token on Raspberry Pi 4B¹², which is too slow for latency-sensitive applications.

C4: Model Adaptation and Personalization. As the environments in which IoT devices are deployed evolve, the newly collected data may deviate from prior distributions. Consequently, the post-deployment pre-trained generative models often necessitate fine-tuning on the devices to effectively adapt to this new data. Moreover, pre-trained generative models also need to adapt to the local data on the device for personalization to enhance user experience.

C5: Lack of Generative AI Agents for IoT Devices. One of the killer applications of Generative AI is AI agents, which are autonomous programs capable of generating new content, making decisions, and performing tasks based on their learned knowledge and user requests. Currently, most AI agents are designed for cloud platforms and operate as cloud services. However, due to the privacy concerns associated with personal data collected by IoT devices, along with the

high latency of delivering cloud-based services¹³, there is a strong need to develop IoT-based AI agents that can process personal data locally on the devices and deliver a wide range of IoT-oriented services promptly, which is not a trivial task due to the diverse and dynamic nature of IoT ecosystems.

C6: Privacy and Security Threats. Data collected by IoT devices such as personal instructions, dialogues, photos and videos often contain privacy-sensitive information that need to be securely stored and processed¹⁰. Dealing with privacy and security threats is challenging in the context of generative models. For example, during retrieval-augmented generation (RAG), the query sentences are frequently sent over the network to a remote vector database to find similar samples for augmentation. While being transmitted through the network, the query is vulnerable to leakages or attacks, presenting challenges for protecting the privacy and securing the generative models.

C7: Lack of Development Tools and Benchmarks. The implementation of Generative AI for IoT applications presents a wide range of challenges due to the unique characteristics of IoT devices and their deployment environments. To facilitate implementation and widespread adoption of these applications, development tools are essential. However, existing generic development tools such as PyTorch and TensorFlow are not designed for IoT-oriented scenarios; development tools designed for mobile and edge devices such as PyTorch Mobile and TFLite do not provide dedicated supports for large generative models. Moreover, the advancement of a field cannot be realized without established benchmarks. Unfortunately, there are still no benchmarks that are specifically designed for Generative AI for IoT-oriented applications.

Opportunities

O1: Build Generative Models for IoT Data (Address C1). The lack of generative models for IoT data hinders the development of generative AI-assisted IoT applications. Therefore, there is a significant opportunity for researchers and practitioners to develop new generative

models for analyzing or generating various IoT data. One promising approach to building generative models for IoT data is through multimodal LLMs. For instance, MEIT⁹ is a multimodal LLM designed to understand and analyze raw ECG sensor data, and generate the analysis reports using human-understandable language. Specifically, MEIT transforms raw ECG sensor data into tokens and stores them alongside text tokens for report generation. By fine-tuning the MEIT model with a small set of ECG sensor data, it has demonstrated superior performance in generating expert-level ECG reports for heart disease diagnosis.

O2: Model Compression (Address C2, C3). To address the challenges of model deployment and content generation under limited memory and compute resources on IoT devices, one effective approach is to compress the generative models¹⁴ before deployment. Conventional compression techniques designed for models that are much smaller than LLMs require retraining after compression. In contrast, many compression techniques for generative models are post-training-based, which avoid retraining after compression given that retraining generative models is expensive. In general, compression techniques for generative models fall into four categories: quantization, parameter pruning, low-rank approximation, and knowledge distillation. Quantization compresses the model by reducing the precision of the weights and/or activations. Nevertheless, even with the most advanced quantization technique, the highest model compression ratio is limited by the smallest bit width. Parameter pruning compresses the model by eliminating redundant model parameters. Pruning methods can be classified into structured pruning and unstructured pruning. Unstructured pruning has much more pruning flexibility and thus enjoys a lower accuracy drop compared to structured pruning. However, unstructured pruning incurs irregular sparsification, which makes the resulting pruned models difficult to be deployed on IoT devices due to lack of hardware support. Low-rank approximation compresses the model by approximating the weight matrix using the product of two or more smaller low-ranking ones with lower dimensions. Knowledge distillation (KD) transfers knowledge from a larger model (the teacher) to a smaller one (the student). Unlike the other three compression strategies, it requires training or fine-tuning for knowledge transfer, making it more expensive to apply. Lastly, it is worthwhile to note that these four categories of model compression methods are orthogonal to each other, allowing for their combinations to further compress the models. Existing methods are able to compress LLMs by up

to 80% with minimum accuracy degradation, enabling the deployment of generative models such as LLaMA-7B on IoT devices. As an example, the quantized 4-bit LLaMA-7B can already be deployed in smartphones while only consuming 4GB memory.

O3: Optimization at Runtime (Address C2). Model compression reduces the memory and compute resource demands of generative models before they are deployed on IoT devices. When performing inference at runtime, intermediate results such as activation outputs and Key-Value (KV) cache have to be computed and stored for further processing. These intermediate results consume a significant amount of computation and memory resources at runtime. Optimizing the runtime memory of KV cache is a unique challenge that conventional ML models do not have. Thus, techniques that reduce the intermediate results such as evicting unnecessary items from the cache or compressing the cache contents are fruitful optimization opportunities.

Another strategy for runtime optimization is allocating certain model parameters, intermediate results, and computational tasks at different memory hierarchy. Given the limited onboard GPU memory inside IoT devices, leveraging the larger capacity of CPU RAM and disk storage is a promising opportunity. This approach allows for the storage of larger intermediate results and part of model parameters outside GPU memory to enable the execution of LLMs on resource-constrained IoT devices. At the same time, a primary challenge of this strategy is the slow data transfer rate between different memory units. Innovative techniques that can reduce the frequency of data transfers or pipeline data transfer with other operations will have great promise to address this challenge.

O4: Edge-Cloud Collaboration (Address C2). Given the limited memory and compute resources, some IoT devices may not be able to even run the models that are already compressed. In such cases, it is necessary to offload the execution of part or even the whole model to nearby resourceful edge servers or the cloud¹⁵.

The key to such edge-cloud collaboration is the design of effective workload partitioning and efficient communication techniques. Workload partitioning is not trivial since IoT devices, edge servers, and cloud have very different compute, memory, and energy resources, and the available network bandwidths can be dynamic. Due to the NP-Hard nature of the workload partitioning problem, manually identifying the best-performing partition can be practically infeasible, especially for billion-parameter generative models where the number of potential choices of partition points can be extremely large. One promising opportunity lies at

designing a highly-efficient search-based strategy to automatically search for and identify the partition points that optimize the overall performance.

Communication between IoT devices and edge servers or the cloud is often conducted through wireless channel in which bandwidth can become the bottleneck. To ensure a timely exchange of migrated workloads while minimizing bandwidth usage and power consumption caused by wireless transmission, efficient communication is essential. At the same time, due to their large model sizes and the potential large amount of data they need to generate, generative models incur a significant burden on communication. In such cases, we envision that efficient communication techniques are highly demanded. For example, instead of directly transmitting the model or data, techniques that compress the model, embedding vectors, as well as raw data will become extremely useful.

O5: Efficient Fine-Tuning (Address C4). The needs for model adaptation and personalization underscore the importance of developing highly efficient on-device fine-tuning techniques for resource-constrained IoT devices. At a high level, efficient fine-tuning techniques fall into three categories: parameter-efficient fine-tuning (PEFT), memory-efficient fine-tuning (MEFT), and data-efficient fine-tuning (DEFT). PEFT reduces the computational cost of fine-tuning by selecting only a subset of model parameters for tuning. Among PEFT, low-rank adaptation (LoRA) is one of the most widely used methods. Instead of fine-tuning the full parameters of generative models, LoRA freezes the weights and injects trainable rank decomposition matrices into the model. Through fine-tuning these small matrices, the model is able to achieve comparable performance as full-parameter fine-tuning. Although PEFT is able to reduce the computational cost of the fine-tuning process, it can still incur large memory usage. Motivated by this limitation, MEFT focuses on reducing fine-tuning memory footprints by conducting model quantization before fine-tuning, utilizing optimizers that require less memory, or combining the gradient calculation and model parameter updating together. Different from PEFT and MEFT, DEFT achieves efficient fine-tuning from a data-centric perspective. By using a small fraction of the data, DEFT can achieve comparable performance to that obtained from fine-tuning with the entire dataset. Another benefit of DEFT is that it can be combined with PEFT or MEFT to further enhance the fine-tuning efficiency. At the same time, most of the existing DEFT methods heavily rely on manual selection of the small set of data for fine-tuning, which can be difficult to accomplish without domain

knowledge. Therefore, an automated data selection scheme would benefit more on applying DEFT for specific IoT application domains. Existing efforts on efficient fine-tuning for IoT devices are able to fine-tune OPT-1.3B using approximately 4GB and 6.5GB of memory on the OPPO Reno6 smartphone¹⁶. However, most generative models, especially LLMs, contain much more parameters than 1.3 billion. Fine-tuning large-scale generative models for IoT devices is still a challenging task but full of opportunities.

O6: Design Generative AI Agents for IoT Devices (Address C5). The privacy and latency issues of cloud-based AI agents motivate the design of IoT-based AI agents. One fundamental capability of AI agents is task planning, which involves breaking down complex tasks into a sequence of simpler steps that could be automatically accomplished. However, designing AI agents that can perform effective task planning can be challenging given the resource constraints of IoT devices as well as diversified IoT-related tasks including sensing, user interactions, data management and processing, information retrieval, control and activation. In such cases, we envision that opportunities lie at developing techniques that allow agents to make highly effective plans for diversified IoT-related tasks and transform the agent's plans into low-level instructions compatible with various IoT devices, and scheduling onboard resources to ensure the execution of plans in an efficient manner.

O7: Federated Learning and Trusted Execution Environment (Address C6). As a privacy-preserving machine learning paradigm, federated learning (FL) emerges as a solution that can improve the quality of the generative models through the personal data while mitigating privacy risks by keeping the data inside the IoT devices¹⁷. While FL has been intensively studied in recent years, most of the proposed techniques have been developed for models with much smaller scales. The emergence of billion-parameter generative models precludes their complete storage within IoT devices due to resource limitations, presenting new challenges in designing FL frameworks that were not previously encountered. To address this challenge, we envision that the opportunities lie in exploring partial training (PT)-based approaches where each IoT device trains a smaller sub-model extracted from the large generative model hosted on the cloud server, and this server model is updated by aggregating those trained sub-models. For example, Alam et al. introduces FedRolex¹⁸, which extracts sub-models from the large server model via a rolling window. Such rolling mechanism results in more stable convergence and

ensures that the global model is updated uniformly with superior model quality.

Trusted Execution Environment (TEE) has become a standard technology to enhance the security of IoT devices. TEE provides a secure area within a processor, ensuring that the data and operations executed within its confines are protected from external threats. In the context of generative AI, TEE can be leveraged in many innovative ways. For example, to secure the user input to the generative models, one can perform tokenization or embedding extraction of the input inside TEE. This ensures the user input is safeguarded from attacks on local devices before being sent to the generative models for inference.

O8: IoT-oriented Development Tools and Benchmarks (Address C7). New development tools that support generative models on mobile and edge devices such as llama.cpp and MLC LLM have recently been developed. However, these newly developed tools are still in their infancy. We envision that further refinement on better supporting IoT-related tasks such as runtime resource management, workload offloading, privacy and security enhancement, as well as efficient fine-tuning is much needed and hence holds great promise.

Lastly, although benchmarks such as MMLU, GSM8K, and MMMU are becoming standards to evaluate the performance of generative models for diverse tasks, there is still no benchmark that is specifically designed for generative AI for IoT-related applications. As the development of generative AI for IoT-related applications is advancing rapidly, we envision that a comprehensive and dedicated benchmark that covers a wide range of IoT-oriented data modalities, tasks, platforms, and evaluation metrics such as latency, memory footprint, and energy consumption will become critical and beneficial to the IoT community.

CONCLUDING REMARKS

Generative AI has shown immense promise in advancing the capabilities of IoT. In this article, we highlighted the key benefits and elaborated on important IoT applications empowered by Generative AI. We also presented the key challenges and opportunities to enable Generative AI for IoT. We hope this article can spark further research in this exciting field.

Acknowledgement

This material is based in part upon work supported by National Science Foundation under award NeTS-2312675 and Defense Advanced Research Projects

Agency under contract number HR001120C0160. Any views, opinions, and/or findings expressed are those of the authors and should not be interpreted as representing the official views or policies of the funding agency or the US government.

REFERENCES

1. Y. Cao et al., "A comprehensive survey of ai-generated content (aigc): A history of generative ai from gan to chatgpt," 2023, arXiv:2303.04226.
2. W. X. Zhao et al., "A survey of large language models," 2023, arXiv:2303.18223.
3. S. Yin et al., "A survey on multimodal large language models," 2023, arXiv:2306.13549.
4. Karapantelakis et al., "Generative AI in mobile networks: a survey," *Ann. Telecommun.* 79, 15–33, 2024.
5. C. Cui et al., "A Survey on Multimodal Large Language Models for Autonomous Driving," 2024, *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV) Workshops*.
6. V. Chamola et al., "Beyond Reality: The Pivotal Role of Generative AI in the Metaverse," 2023, arXiv:2308.06272.
7. F. Zeng et al., "Large language models for robotics: A survey," 2023, arXiv:2311.07226.
8. Peng, C. et al, "A study of generative large language model for medical research and healthcare," *npj Digit. Med.* 6, 210, 2023.
9. Wan, Zhongwei et al., "MEIT: Multi-Modal Electrocardiogram Instruction Tuning on Large Language Models for Report Generation," 2024, arXiv:2403.04945.
10. M.A. Ferrag et al., "Revolutionizing Cyber Threat Detection with Large Language Models: A privacy-preserving BERT-based Lightweight Model for IoT/IIoT Devices," 2024, arXiv:2306.14263.
11. Wan, Z. et al., "Efficient Large Language Models: A Survey," *Transactions on Machine Learning Research (TMLR)*, 2024.
12. Yi, R. et al., "EdgeMoE: Fast On-Device Inference of MoE-based Large Language Models," 2023, arXiv:2308.14352.
13. Li, Y. et al., "Personal LLM Agents: Insights and Survey about the Capability, Efficiency and Security," 2024, arXiv:2401.05459.
14. X. Zhu et al., "A survey on model compression for large language models," 2023, arXiv:2308.07633.
15. Z. Zhou et al., "Edge Intelligence: Paving the Last Mile of Artificial Intelligence With Edge Computing," *Proceedings of the IEEE*, vol. 107, no. 8, pp. 1738–1762, 2019.
16. Peng, D. et al., "PocketLLM: Enabling On-

Device Fine-Tuning for Personalized LLMs," 2024, arXiv:2407.01031.

17. Alam, S. et al., "FedAloT: A Federated Learning Benchmark for Artificial Intelligence of Things," Journal of Data-centric Machine Learning Research, 2024.
18. Alam, S. et al., "FedRolex: Model-Heterogeneous Federated Learning with Rolling Sub-Model Extraction," Conference on Neural Information Processing Systems (NeurIPS), 2022.

Xin Wang is a Ph.D. student in computer science and engineering at The Ohio State University, Columbus, OH 43210, USA, advised by Prof. Mi Zhang. His research interests include efficient large language models, machine learning systems, and edge AI. Wang received his master's degree in computer science and engineering from The Ohio State University. Contact him at wang.15980@osu.edu.

Zhongwei Wan is a Ph.D. student in computer science and engineering at The Ohio State University, Columbus, OH 43210, USA, advised by Prof. Mi Zhang. His research interests include efficient large language models, large multimodal models, and their applications. Wan received his master's degree in control science and engineering from University of Chinese Academy of Sciences, Beijing 100049, China. Contact him at wan.512@osu.edu.

Arvin Hekmati is a Ph.D. student in computer science at the University of Southern California, Los Angeles, CA 90089, USA, advised by Prof. Bhaskar Krishnamachari. His research interests include machine learning, data-driven algorithms, anomaly detection, and edge computing. Hekmati received his master degree in electrical and computer engineering from McMaster University, ON L8S 4L8, Canada. Contact him at hekmati@usc.edu.

Mingyu Zong is a Ph.D. student in computer science at the University of Southern California, Los Angeles, CA 90089, USA, advised by Prof. Bhaskar Krishnamachari. Her research interests include large language models for the Internet of Things (IoT) and federated learning. Zong received her master degree in spatial data science from University of Southern California. Contact her at mzung@usc.edu.

Samiul Alam is a Ph.D. student in computer science at The Ohio State University, Columbus, OH 43210, USA, advised by Prof. Mi Zhang. His research interests include federated learning and efficient large language models on IoT devices. Alam received

his master degree in computer science from Michigan State University, MI 48824, USA. Contact him at alam.140@osu.edu.

Mi Zhang is an associate professor in the Department of Computer Science and Engineering at The Ohio State University, Columbus, OH 43210, USA. He is a senior member of IEEE and ACM. His research interests include Artificial Intelligence of Things (AloT), machine learning systems, generative AI, mobile computing, and their domain-specific applications. Zhang received his Ph.D. in computer engineering from the University of Southern California, Los Angeles, CA 90089, USA. Contact him at mizhang.1@osu.edu.

Bhaskar Krishnamachari is a professor and Ming Hsieh Faculty Fellow in electrical engineering at the University of Southern California, Los Angeles, CA 90089, USA. He is a Fellow of IEEE. His research interests include the design and analysis of algorithms, protocols, and applications for next-generation wireless networks, the IoT, distributed systems, blockchain technologies, AI and machine learning, and network economics. Krishnamachari received his Ph.D. in electrical engineering from Cornell University, NY 14853, USA. Contact him at bkrishna@usc.edu.