

A Comprehensive Survey of AI Security Threats in IoT and Edge Cloud Systems

Venkatesan Cherappa¹, Hsin-Hung Cho^{2*}, Yasir Abdullah Rabi³, Ananth John Patrick⁴

¹ Department of Electronics and Communication Engineering, HKBK College of Engineering, India

² Department of Computer Science and Information Engineering, National Ilan University, Taiwan

³ Department of Artificial Intelligence and Data Science, Dr. Mahalingam College of Engineering and Technology, India

⁴ Department of Computer Science and Engineering, Dayananda Sagar University, India
 venkatesanc.ec@hkbk.edu.in, hhcho@niu.edu.tw, ry.aids@drmcet.ac.in, ananth-cse@dsu.edu.in

Abstract

AI security is a key design criterion for IoT and edge cloud architectures where learning models are embedded near physical processes, receive streaming data from sensors, and are subject to continuous updates. This survey research work proposes a lifecycle-based understanding of AI security threats and proposed a unified taxonomy for the four major categories of threats observed in real-world settings, namely, data poisoning and backdoor attacks on learning model updates, adversarial attacks on model outputs through input manipulation, privacy leakage of confidential information through model-based queries, and model extraction for intellectual property theft and creation of rogue replicas of learning models. The research in surveying these defenses takes account of the limitations posed by the use of edge computing platforms from a computational standpoint, latency tolerance, network connections, and variety of hardware. These defense mechanisms include model provenance and data screening, robust training and backdoor attacks, monitoring and calibration, privacy-preserving learning and access control, and model protection through throttling, fingerprinting, watermarking, and attestation. Unlike prior surveys that treat these threats separately, this survey unifies them under a lifecycle-based taxonomy tailored to IoT and edge cloud deployments and emphasizes deployable defenses under latency, compute, and hardware constraints.

Keywords: AI security, IoT, Edge cloud, Data poisoning, Adversarial attacks

1 Introduction

The motivation for AI at the IoT edge and its relationship to the expansion of the security boundary can be explained by the constraints of day-to-day operation. Edge computing distributes AI models, feature pipelines, and decision logic across devices, gateways, and nearby edge servers [1].

The security boundary expands because there are new assets that need to be protected in addition to the

firmware and network keys. In fact, security guidance for AI systems has already shown that model integrity and decision reliability are security goals rather than quality goals [2]. Operational security for AI systems includes model rollback and auditable retraining in addition to patching and key rotation. Why conventional IoT security mechanisms are not fully applicable to AI-specific failure modes is evident when channel integrity is distinguished from decision integrity [3].

The NIST adversarial ML taxonomy provides a framework for understanding attacks and mitigations throughout the lifecycle and attacker objectives and demonstrates that integrity compromises are possible without compromising cryptography and identification mechanisms [4-5]. The ENISA taxonomy also identifies poisoning and evasion as attack classes that leverage statistical learning characteristics and are mitigated via threat model-aware testing [6]. These are important because they are systematic and replicable and therefore closer to exploitation than random noise. Poisoning is possible via compromised sensing, poor supervision, and incorrect labeling of training data and results in targeted failure that is obscured by average performance metrics [7]. In the inference phase, an adversary can generate inputs close to the decision boundary such that there are certain failures with confidence that lead to unsafe decisions within edge-controlled models [8]. In this work, we consider model poisoning attacks and backdoors in the context of data streaming pipelines, collaborative learning scenarios, and the reuse of pretrained modules in supply chains [9]. Adversarial attacks are taken into consideration both in the digital and the physical domain [10]. Privacy breaches are considered since black-box access to a model can provide an adversary with membership details of training data; such details are important for location data, biomedical data collected via Internet-of-things sensors, and industrial monitoring data [11]. Model extraction is considered as access to prediction interfaces and devices can enable the theft of deployed decision functions, leading to stronger attacker replicas and evasion techniques [12-13].

The federated learning approach is also covered if used to tackle data siloing and connectivity issues by leveraging existing frameworks to make assumptions

*Corresponding Author: Hsin-Hung Cho; Email: hhcho@niu.edu.tw
 DOI: <https://doi.org/10.70003/160792642026072704008>

about heterogeneity and unreliability of clients [14]. In terms of privacy leakage, the survey establishes the connection between membership inference risk and mitigation measures like differential privacy and interface control while acknowledging trade-offs between them, utility, and latency [15]. Additionally, the survey highlights the importance of MLOps' security in terms of signing and versioning, which is crucial for rollback and accountability [16]. The structure of survey research adopts a foundation-synthesis approach to ensure that the discussion remains pertinent. It begins with a description of architectures for IoT and edge cloud, AI workload models, and security characteristics that affect safety-critical decision making. Before analyzing AI-specific threats, one needs to first comprehend the architecture tiers, trust boundaries, and limitations associated with IoT and edge cloud technologies. These form the basis for the location of the learning capabilities, attack surfaces, and explain why some traditional security mechanisms might be inadequate.

The present study uses a systematic narrative review to analyze the current literature about AI-based threats in IoT and edge cloud technologies, where sources were chosen from reputable academic databases and standardization organizations. In particular, only articles that were relevant to the study of edge-AI threats and limitations as well as their lifecycle-based discussion were considered.

In contrast to other surveys on attacks against different attack classes or studies on security in artificial intelligence in the context of machine learning, this survey presents a unified approach through the lens of the lifecycle model which is specific for IoT and edge cloud deployments. The novelty consists of the integration of threat taxonomy, trust boundaries, deployment stages, and edge-specific limitations like latency, computing power, hardware

diversity, and intermittent network connectivity into one comprehensive framework.

There are four key contributions of this survey. First, we present a unified taxonomy of AI security threats within the IoT and edge cloud framework. Second, we tie these threats to assets, trust boundaries, and deployment stages. Third, we synthesize defense approaches based on the edge-centric view of deployment stage limitations in terms of latency and computing constraints. Fourth, we discuss useful evaluation criteria such as robustness, privacy risk, extraction resistance, and overheads.

2 AI Enabled IoT and Edge Cloud Security Foundations

AI-enabled IoT and edge cloud computing deployments have been found to be typically based on a three-tiered reference architecture for the separation of sensing, near device processing, and cloud scale management [17]. However, this also adds to the security scope because the trust scope now includes not just the devices and the network, but also the models, the preprocessing logic, the calibration parameters, and the updates, as highlighted by AI security analyses that have emphasized the fact that these assets can be compromised without the traditional need to break into the system, as the nature of the attack surfaces for machine learning and decision-making makes it easy to compromise these systems in a cyber-physical setting [18-19]. The foundations of IoT security have also highlighted the fact that trust is typically fragmented among different stakeholders and that end-to-end trust is based on the coordination of controls among different components, and not just the device and the cloud, as typically considered [20].

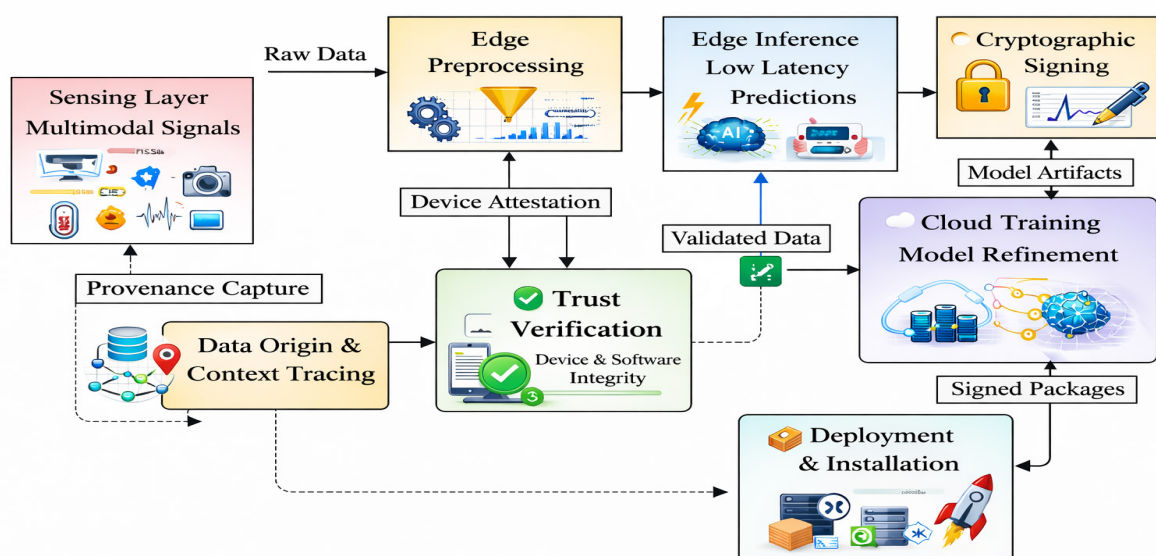


Figure 1. Secure IoT edge cloud AI lifecycle

Adversarial machine learning surveys have shown that, in terms of inference time attacks, attackers can create small but structured perturbations in the data, and the edge model may be more vulnerable due to noisy input channels and tight compute constraints [21]. This shifts the emphasis to security mechanisms that are light, measurable, and integrated into the model, as opposed to being added at the network layer [22].

In this context, federated learning is often used to minimize data movement and centralization, which in turn creates a new problem in terms of model updates, as clients can send malicious updates or poisoned data that can affect the global model performance in a negative way [23]. In terms of poisoning, there are different types of poisonings, as shown in the poisoning surveys, which include targeted poisoning and availability poisoning, in which performance is degraded to create outages or loss of trust in the model or its performance [24].

The research in membership inference has shown that even black box access can disclose information about the presence or absence of specific data in the training data, which is a direct privacy risk for IoT data such as health data, location data, and industrial process data [25-26]. The research in model extraction has shown that even a determined adversary can create a surrogate model mimicking the original model, leading to intellectual property theft and more potent evasion by tuning the attacks offline [27].

Figure 1 illustrates an end-to-end IoT edge cloud AI pipeline with in-built trust controls at all stages of its life cycle. Multimodal sensing streams are collected at the sensing layer and simultaneously connected to provenance capture to track the origin and context of the data for auditing purposes. Data cleansing, normalization, and preparation occur at the edge preprocessing layer before the data is sent to the inference layer for decision-making or prediction purposes. Conceptually, Figure 1 presents the end-to-end secure AI lifecycle and trust placement across system stages, while Figure 2 and Figure 3 focus more specifically on attack entry points, threat propagation, and corresponding control placement across the MLOps pipeline.

The basic security requirements of AI-enabled services should be specified at the system level and should be based on integrity, availability, confidentiality, safety, and accountability as operational objectives. These requirements will include integrity of data streams, labels, feature pipelines, model weights, and outputs and will be directly impacted by poisoning, backdoor, and adversarial attacks. The availability of AI-enabled services should be continuous and should be impacted by denial of service, resource starvation, and data availability attacks on edge bottlenecks.

3 Threat Model and Unified Taxonomy

Threat modeling for AI security of IoT and edge cloud systems starts with defining what is protected, what is under the control of the system, and what are the goals of

an attacker. A lifecycle approach is useful because it links each identified threat to a step of the ML pipeline and an attacker's goal and information, as is done in popular ML security taxonomies [28]. In edge cloud security, there is a need to track trust boundaries and ownership over time, as is done in risk management practices [29]. The assets of concern are not limited to model weights alone, but also include data assets, model artifacts, outputs, and operational components. These assets are important because integrity violations may be introduced upstream and remain latent until deployment, while confidentiality violations may occur even without direct exposure of raw data [30].

At the cloud MLOps layer, issues are centered on dataset curation, training configuration, dependency supply chain, and release channels, as poisoning and backdoors are introduced and then propagate throughout the entire fleet [31]. There are issues of poisoning and backdoor introduction during collection, curation, training, and federated aggregation of ML models, resulting in targeted misbehavior while still achieving expected accuracy on the entire dataset [32]. Adversarial issues occur during inference as crafted inputs are used to cause misclassifications during inference under realistic constraints [33]. There are also issues of privacy leakage as ML models leak information about training data through outputs, such as membership inference attacks that demonstrate that black-box access to a model leaks information about whether a specific record was used during training [34].

The threat of model extraction comes about as a result of repeated inquiries or limited access, which can be used to develop a substitute model that can reproduce the victim model's capabilities, leading to IP theft and more precise tuning of offline attacks [35]. In the context of black box attackers who can only see the top 1 labels or probabilities, there are still viable attacks and extraction strategies that can be implemented using systematic querying and transfer learning [36]. In the context of white box attackers who have access to the device or pipeline, there can be corruption, backdoors, and deactivation of runtime validation, leading to persistent effects due to weak provenance and validation [37]. For instance, poisoning and adversarial examples could lead to the blocking or false triggering of fault alarms that would consequently compromise the availability and safety of such systems [38]. A brief comment on the availability aspect being not just a purely classical notion in that edge AI services could potentially be compromised through attack vectors intimately associated with model extraction, poisoning, and overload.

From Table 1, it becomes evident that AI security threats within the IoT/edge cloud ecosystem vary according to the nature of the threat itself, but additionally, depending on the lifecycle phase, the capabilities of the adversary, and other factors. What can be emphasized is that a taxonomy is necessary since an attack could potentially compromise multiple aspects of availability at once.

Table 1. Threat taxonomy summary with assumptions capabilities and impact

Threat family	Primary lifecycle stage	Typical attacker knowledge	Typical attacker capability	Typical impact in IoT and edge cloud
Data poisoning [31]	Data collection, curation, training	Partial knowledge of features or labeling process	Inject or bias samples, corrupt labels, skew streams	Targeted errors & silent drift
Backdoor insertion [32]	Training, model supply chain, updates	Knowledge of training pipeline or model artifact path	Implant trigger patterns in data or weights	Stealthy misbehavior on trigger
Adversarial attacks [33-34]	Inference	Black box	Craft inputs under constraints to cross decision boundary	Evasion & false negatives
Privacy leakage [35]	Inference and model release	Black box	Membership or attribute inference via queries	Exposure of sensitive participation
Model extraction [36]	Deployment	Black box	Query at scale to train surrogate knockoff	IP theft, stronger offline attack tuning, replication
Availability disruption [28-29]	Any stage, often runtime and updates	System level knowledge of bottlenecks	Resource exhaustion	Service outages & delayed mitigation

4 Data Poisoning Threats in IoT and Edge Cloud

It is recommended that poisoning is considered an intentional manipulation of the training distribution and its impact is evaluated against attacker objectives and persistence, not average accuracy [39]. Another area of pressure is in the generation of labels, and this can be manipulated with targeted confusion, and once the labels are compromised, retraining can amplify the attacker's desired behavior across versions [40]. Using a window can make this worse because a small poisoning period can be reused to calibrate or tune thresholds that were set before the attack was detected [41]. Hence, the importance of supply chain controls, signing, and provenance checks becomes clear, along with the need to have disciplined dependency management for training and packaging pipelines [42-43].

Federated learning introduces another poisoning vector, with the attacker possibly using the model updates to their advantage. Although federated learning minimizes the need to transfer large amounts of data and keeps the raw data local, the possibility exists that malicious clients or Sybil participants will be able to submit gradients that can affect the overall model [44].

Training time defenses include influence limiting, filtering suspicious data, and robust aggregation in federated learning, followed by periodic server-side tests on trusted slices of data retained on the server [45]. Staged rollouts, canary releases, rollback preparation, and monitoring of model calibration and error rates can detect poisoning-induced drift, as in a threat taxonomy that is life cycle-oriented and applicable to AI security governance in general [46]. Figure 2 illustrates the major security risks of AI systems and their entry points in the IoT edge cloud

pipeline, as well as the entry points of the secure MLOps controls that can mitigate them. In this figure, sensor spoofing and replay affect the data sent to the pipeline, while poisoning and weight trojaning affect the learning artifacts on the cloud side.

The literature converges on the view that poisoning and backdoor attacks are difficult to detect using clean accuracy alone and therefore require provenance-aware training and validation practices. However, a clear open gap remains in lightweight poisoning defenses that can operate reliably under streaming, federated, and resource-constrained edge settings.

5 Adversarial Attacks Against Edge AI Inference

An attacker in IoT and edge cloud architectures does not necessarily need to take control of the edge node since modifying the input data at the sensor interface, over the wireless channel, or client application programming interfaces can work just fine when traditional security metrics on the device report everything is well [47]. Adversarial attack modifications are limited by sampling rates, quantization, and other physical constraints. It is possible to achieve significant reductions in anomaly detection scores or even assign specific classes by making small modifications. Due to the fact that edge models are usually trained with quantization and pruning techniques for achieving latency, memory, and energy requirements, the model robustness must be tested using the exact precision and compression settings of the model in deployment rather than the original settings on the cloud side.

In radio sensing and authentication, however, the attack can be transmitted as low-power interference and, as

such, is defined as a waveform rather than a pixel pattern [49]. Practical black-box attacks have demonstrated the possibility of training a substitute model with limited feedback and using it to generate transferable adversarial attacks on the target [53].

Quantization affects the precision of the model, and experimental results have demonstrated that robustness does not correlate with clean accuracy and should be evaluated at the precision used in deployment, rather than at the floating-point checkpoint [54]. Pruning can also affect robustness, improving or decreasing robustness,

especially when pruning is primarily used as a function of latency targets [55]. It can be seen from the literature that robustness to attacks should be considered within realistic deployment settings and not in ideal cloud environments. An important open problem is how robustness can be preserved when quantization, pruning, and compression are used without having to pay an unacceptable amount in terms of latency and power consumption. It can be seen from Table 2 that the effectiveness of the adversarial defense mechanism for the edge setting highly depends on the deployment requirements and the point of attack.

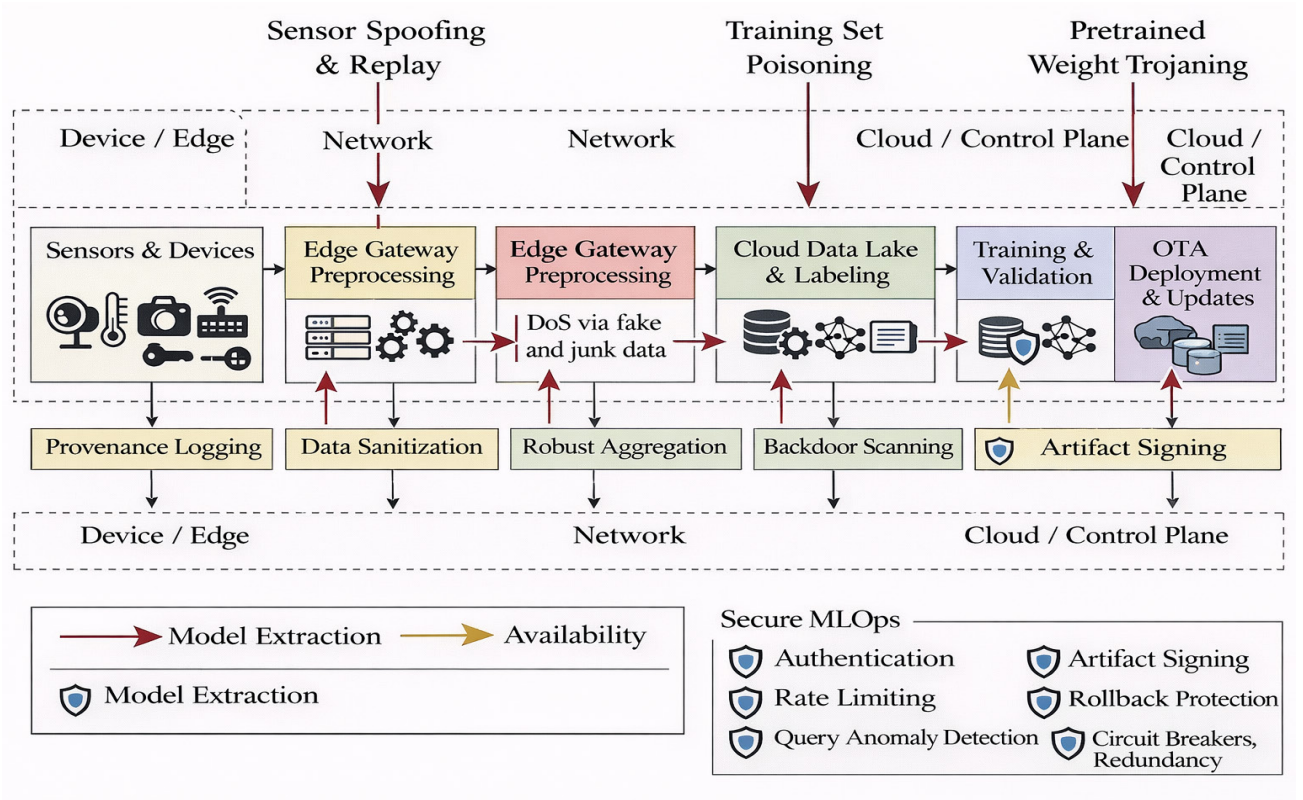


Figure 2. Integrated IoT edge cloud threat and control map

Table 2. Adversarial attack types versus defenses with edge feasibility notes

Attack and setting	Typical entry point	Representative defenses	Edge feasibility
Digital evasion on time-series and features [45, 50]	Sensor Stream, Feature Pipeline	Robust Training & Input Constraints	Medium
Black-box transfer and query-limited evasion [47, 46]	Prediction API, Confidence Outputs	Output Restriction & Randomized Smoothing	Medium
Physical perturbations on vision and wearables [48-49]	Camera and Wearable Sensing	Sensor Fusion & Environment-Aware Training	Medium
Channel-aware wireless perturbations [51-52]	Over-The-Air Interference	Channel-Aware Defenses, Noise Augmentation, Anomaly Detection	Low-Medium
Compression-induced shifts [54]	Pruned Runtime	Robustness Re-Testing Per Build	High

6 Privacy Leakage in IoT and Edge Cloud AI

Privacy leakage in IoT systems and edge cloud-based AI systems is not restricted to the usual breaches in stored data sets. In fact, sensitive information can leak via the model itself as well as its associated interfaces [56-57]. Membership inference attacks have already demonstrated that an attacker can make a determination about whether a certain instance was used for training based on model responses. This can be particularly dangerous because a particular instance can represent an individual, a household routine, a rare medical occurrence, or a unique industrial occurrence [58-59]. Model inversion attacks and attribute inferences can be considered as extensions of model leakage [60-61].

Classifier access can provide access to meaningful characteristics of the training sets and trade secrets that go beyond the intended model owner [62]. Edge devices can be deployed with multiple services that co-execute. They may not have the isolation and monitoring that cloud environments provide. This can be used to launch stronger attacks with faster model extraction. Even with basic latency measurement, attackers can gain insight into the model's depth, branching characteristics, or batching [63]. LIME is a popular model-agnostic explanation technique that makes the local decision behavior of models interpretable, and in untrusted environments, the outputs of the explanation should be treated as sensitive outputs with their own access control [64].

There are techniques that can be used to perform privacy-preserving learning and inference, and the tradeoffs matter more at the edge. One technique that can be employed to provide bounds on information leakage from a learning process is differential privacy. DP-SGD was an interesting technique for training deep neural networks with provable privacy guarantees, drawing on advancements in differential privacy. But improved privacy levels come at the expense of accuracy and efficiency [65]. Secure aggregation is another approach for reducing information disclosure for federated learning participants, where the server only sees aggregated updates from clients rather than client-by-client updates, mitigating potential client-level information disclosure for participants who are only partially trusted [66-67]. Existing studies consistently show that model outputs and interfaces can leak sensitive information even when raw data are never directly exposed. However, an important open challenge is how to balance privacy-preserving mechanisms with utility, interpretability, and low-overhead deployment in heterogeneous edge environments.

7 Model Extraction, Availability Disruption, and Defense Landscape

Model extraction is considered one of the major economic and security risks to AI-based IoT and edge cloud computing, as models are exposed to prediction

interfaces and duplicated in the fleet. In the case of a distributed environment, the attacker can steal the functionality by repeating the queries to the prediction interface and creating a surrogate model that is equivalent to the original model, as demonstrated in the black-box stealing studies and consolidations [68]. The strategies that the attacker may follow are based on the feedback that the attacker can get, such as probabilities or hard labels, and the creation of a transfer set using public or synthetic data to reduce the cost with high fidelity [69]. PRADA demonstrates that model extraction has behavioral patterns in the queries, which can be used to detect the attacks instead of relying on prevention [70].

Physical theft is a complementing method to the above-described API-based extraction. Physical access, firmware dumps, debug interfaces, and memory scraping can potentially reveal the deployed artifact along with its quantized weights and preprocessing functions. Compromise of the supply chain could also serve as a more extensive method of attacking, whereby a poisoned package or a malicious update routine can leak the weights post-deployment. This shows that the confidentiality of the model is something that should be guaranteed at run time with the help of protecting the model's artifact [71-72].

IP theft is not the only negative impact of model stealing. A surrogate model that is stolen will be a valuable asset in exploring the boundaries and fine-tuning evasions, and, in addition, cloning the service creates a problem with accountability in the context of safety-critical IoT environments [73-74]. The interface control methods may comprise user authentication, rate limiting, output minimization, and detecting anomalies in the distribution of requests. An advanced technique of detecting theft attacks is PRADA [75].

Techniques like DeepMarks are used for model attribution if there are more than one licensed copy of the models available. Enclaves and attested execution can be used in case of model weights in order to minimize direct weight exposure, if supported by hardware. Otherwise, encryption of storage, least privilege execution environment, and secure debug interfaces may help avoid opportunistic weight leakage [76].

Denial-of-service attacks may impact the inference services, telemetry, and updates. Moreover, the edge taxonomy defines other vulnerabilities, which include flooding of edge API, saturation of backhaul connectivity, and exploitation of orchestration services for generating restart storms. There is another category of attacks, which are called resource exhaustion attacks, and they are feasible since the adversaries may enforce the expensive model paths and generate retries causing CPU, memory, and battery exhaustion [77].

A realistic security model assumes that extraction resistance and availability resilience are secure MLOps requirements. Release packages should be signed, versioned, and verified during installation. Rollback attacks and key compromise should be addressed with update system research like TUF [76]. Supply chain-aware update systems can be used to limit the blast radius by improving delegations and trust relationships in multiple repositories.

This can be particularly relevant for container-based edge computing scenarios [78].

In reality, more secure AI controls may come at a cost to latency, computational burden, memory consumption, and operational difficulties. Techniques such as secure training, differential privacy, attestation, and query auditing can increase security, but their use must be weighed against the limitations of the device and the need for continued service availability.

The data-centric security principles outlined in AI security guidelines can be employed to minimize the chance of silent degradation of decision-making due to disruptions [79]. This Survey is able to bring together our insights on the interplay between these four types of attacks in real-world IoT edge cloud deployments. In the case of poisoning and backdoors, future decisions are impacted since they involve changes to streaming data, labels, and updates. On the other hand, adversarial inputs exploit the decision boundaries under latency constraints in their immediate environment. Likewise, privacy attacks and model theft attacks usually arise due to the same interface choices. The survey makes clear that individual solutions will not suffice, and layered approaches must

be taken in defense against origin, robust learning, and runtime/MLOps security issues. From the analysis of the paper, we can see that there is a close relationship between model extraction, disruption of availability, and security in edge-based AI systems. However, currently, there is no single framework that takes into account all three aspects mentioned when measuring the performance of an edge-based AI application. When measuring the security of an AI-enabled IoT device or edge-based computing, it is crucial to look at metrics other than only accuracy and consider aspects such as attack success probability under realistic threat scenarios, degradation of robustness under realistic deployment, risk of leaking private information through the output or update of the model, and model extraction or cloning.

Table 3 indicates that model extraction and availability disruption must be understood through the attacker's access mode, since remote, local, and supply-chain threats require different protective controls. The main insight is that secure edge AI deployment cannot rely on a single safeguard, but instead requires coordinated interface, device, and update-pipeline protections.

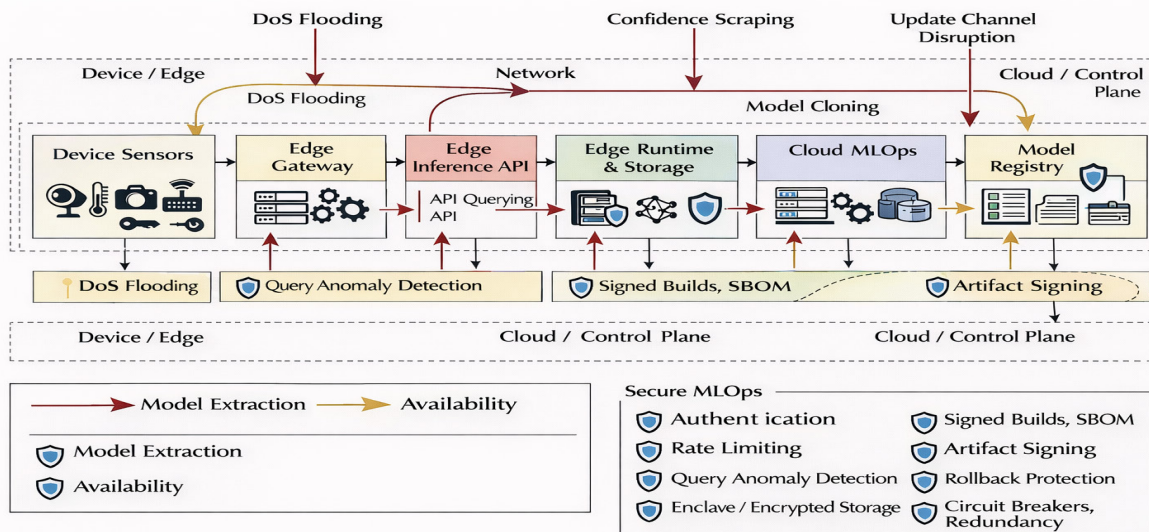


Figure 3. Integrated Threat-to-Control map across the IoT Edge-Cloud MLOps pipeline

Table 3. Unified summary of model extraction vectors

Threat vector or mode	Attacker access	Likely impact	Edge-feasible controls
API query extraction	remote, black-box	model cloning	throttling, output restriction, query anomaly alarms
On-device weight theft	local or insider	full artifact loss	debug hardening, encryption at rest, isolation and attestation
Supply-chain exfiltration	pipeline or update access	fleet-wide compromise	signed builds, delegated trust, rollback protection
DoS and resource exhaustion	remote or botnet	service outage	rate limits, circuit breakers, redundancy, staged rollout

8 Conclusion

This survey has placed the security of AI in IoT systems and edge cloud systems as a lifecycle process that includes data acquisition, training, deployment, inference, and updates. In this regard, the synthesis has unified poisoning and backdoors, adversarial attacks, privacy leakage, and model extraction as a taxonomy that makes it clear where attacks come in, what they touch, and what can be done about it at the edge. The synthesis has highlighted deployable defenses that include provenance and screening, robust learning and backdoors, runtime monitoring with calibration and out-of-distribution detection, privacy-preserving learning, model protection with throttling, fingerprinting, watermarking, and attestation. Application-wise, these results will be particularly important for applications like IoT in healthcare, where leakage and manipulation of privacy may impact diagnostics and monitoring capabilities. They are just as crucial for the industrial sector, where attacks such as poisoning or extraction might cause faults in predictive maintenance and process control. Smart city edge computing systems require reliable and secure AI deployment since faulty AI inference in nodes like traffic, surveillance, or public safety may cause cascading risks within the interconnected city services. Future research efforts on this topic would need to concentrate on creating benchmarks where there are tests for physical attacks, digital attacks, drifts in streaming, and realistic knowledge that attacker claims to prove their robustness and privacy are aligned. Better alignment of secure MLOps practices with governance, logging hygiene, and incident response could turn metrics into service level objectives. Another direction is to optimize robustness, privacy, and efficiency together, as edge security compromises on accuracy to prioritize safety and energy efficiency. Hardware-assisted isolation, including trusted execution and measured boot, should be compared to extraction and side channels. Red teaming can limit regression in production. Several focused research challenges remain open for secure AI deployment at the edge. These include secure federated update validation against poisoned or malicious client contributions, benchmark standardization for realistic edge-AI attack evaluation, hardware-assisted protection against extraction and side-channel risks, and attack-aware edge orchestration that can adapt inference, updates, and recovery policies under active adversarial conditions.

Acknowledgements

This research was partly funded by the National Science and Technology Council of the R.O.C. under Grants NSTC 114-2221-E-197-021, NSTC 112-2221-E-197-016-MY3, and NSTC 114-2634-F-004-002-MBK.

References

- [1] J.-Y. Zeng, M.-Y. Tsai, H.-H. Cho, Security of artificial intelligence, in: J. G. Webster (Ed.), *Wiley Encyclopedia of Electrical and Electronics Engineering*, John Wiley and Sons, 2026.
<https://doi.org/10.1002/047134608X.W8449>
- [2] W. Shi, J. Cao, Q. Zhang, Y. Li, L. Xu, Edge computing: vision and challenges, *IEEE Internet of Things Journal*, Vol. 3, No. 5, pp. 637-646, October, 2016.
<https://doi.org/10.1109/JIOT.2016.2579198>
- [3] S. Sicari, A. Rizzardi, L. A. Grieco, A. Coen-Porisini, Security, privacy and trust in internet of things: The road ahead, *Computer Networks*, Vol. 76, pp. 146-164, January, 2015.
<https://doi.org/10.1016/j.comnet.2014.11.008>
- [4] E. Tabassi, *Artificial intelligence risk management framework (AI RMF 1.0)*, National Institute of Standards and Technology, Report No. NIST AI 100-1, January, 2023.
<https://doi.org/10.6028/NIST.AI.100-1>
- [5] A. Vassilev, A. Oprea, A. Fordyce, H. Anderson, X. Davies, M. Hamin, *Adversarial machine learning: A taxonomy and terminology of attacks and mitigations*, National Institute of Standards and Technology, Report No. NIST AI 100-2e2025, March, 2025.
<https://doi.org/10.6028/NIST.AI.100-2e2025>
- [6] European Union Agency for Cybersecurity, *Securing machine learning algorithms*, ENISA Reports, December, 2021.
<https://www.enisa.europa.eu/sites/default/files/publications/ENISA%20Report%20-%20Securing%20Machine%20Learning%20Algorithms.pdf>
- [7] Z. Wang, J. Ma, X. Wang, J. Hu, Z. Qin, K. Ren, Threats to training: A survey of poisoning attacks and defenses on machine learning systems, *ACM Computing Surveys*, Vol. 55, No. 7, Article No. 134, July, 2023.
<https://doi.org/10.1145/3538707>
- [8] Y. Li, Y. Jiang, Z. Li, S.-T. Xia, Backdoor learning: A survey, *IEEE Transactions on Neural Networks and Learning Systems*, Vol. 35, No. 1, pp. 5-22, January, 2024.
<https://doi.org/10.1109/TNNLS.2022.3182979>
- [9] P. Kairouz, H. B. McMahan, B. Avent, A. Bellet, M. Bennis, A. N. Bhagoji, K. Bonawitz, Z. Charles, G. Cormode, R. Cummings, R. G. L. D'Oliveira, H. Eichner, S. El Rouayheb, D. Evans, J. Gardner, Z. Garrett, A. Gascón, B. Ghazi, P. B. Gibbons, M. Gruteser, Z. Harchaoui, C. He, L. He, Z. Huo, B. Hutchinson, J. Hsu, M. Jaggi, T. Javidi, G. Joshi, M. Khodak, J. Konečný, A. Korolova, F. Koushanfar, S. Koyejo, T. Lepoint, Y. Liu, P. Mittal, M. Mohri, R. Nock, A. Özgür, R. Pagh, H. Qi, D. Ramage, R. Raskar, M. Raykova, D. Song, W. Song, S. U. Stich, Z. Sun, A. T. Suresh, F. Tramèr, P. Vepakomma, J. Wang, L. Xiong, Z. Xu, Q. Yang, F. X. Yu, H. Yu, S. Zhao, Advances and open problems in federated learning, *Foundations and Trends in Machine Learning*, Vol. 14, No. 1-2, pp. 1-210, June, 2021.
<https://doi.org/10.1561/22000000083>
- [10] K. Ren, T. Zheng, Z. Qin, X. Liu, Adversarial attacks and defenses in deep learning, *Engineering*, Vol. 6, No. 3, pp. 346-360, March, 2020.
<https://doi.org/10.1016/j.eng.2019.12.012>
- [11] N. Carlini, D. Wagner, Towards evaluating the robustness of neural networks, *IEEE Symposium on Security and Privacy*, San Jose, CA, USA, 2017, pp. 39-57.
<https://doi.org/10.1109/SP.2017.49>
- [12] R. Shokri, M. Stronati, C. Song, V. Shmatikov, Membership inference attacks against machine learning models, *IEEE Symposium on Security and Privacy*, San Jose, CA, USA, 2017, pp. 3-18.

- <https://doi.org/10.1109/SP.2017.41>
- [13] D. Oliynyk, R. Mayer, A. Rauber, I know what you trained last summer: A survey on stealing machine learning models and defences, *ACM Computing Surveys*, Vol. 55, No. 14s, Article No. 324, December, 2023.
<https://doi.org/10.1145/3595292>
- [14] F. Tramèr, F. Zhang, A. Juels, M. K. Reiter, T. Ristenpart, Stealing machine learning models via prediction APIs, *USENIX Security Symposium*, Austin, TX, USA, 2016, pp. 601-618.
- [15] C. Dwork, Differential privacy, in: M. Bugliesi, B. Preneel, V. Sassone, I. Wegener (eds), *Automata, Languages and Programming, Lecture Notes in Computer Science, Vol. 4052*, Springer, Berlin, Heidelberg, 2006, pp. 1-12.
https://doi.org/10.1007/11787006_1
- [16] U. B. Nisha, A. Subair, R. Yasir Abdullah, An efficient algorithm for anomaly detection in wireless sensor networks, *International Conference on Smart Electronics and Communication*, Trichy, India, 2020, pp. 925-932.
<https://doi.org/10.1109/ICSECC49089.2020.9215258>
- [17] M. Chiang, T. Zhang, Fog and IoT: An overview of research opportunities, *IEEE Internet of Things Journal*, Vol. 3, No. 6, pp. 854-864, December, 2016.
<https://doi.org/10.1109/JIOT.2016.2584538>
- [18] L. Atzori, A. Iera, G. Morabito, The internet of things: A survey, *Computer Networks*, Vol. 54, No. 15, pp. 2787-2805, October, 2010.
<https://doi.org/10.1016/j.comnet.2010.05.010>
- [19] T. Li, A. K. Sahu, A. Talwalkar, V. Smith, Federated learning: Challenges, methods, and future directions, *IEEE Signal Processing Magazine*, Vol. 37, No. 3, pp. 50-60, May, 2020.
<https://doi.org/10.1109/MSP.2020.2975749>
- [20] K. T., B. N. U., Y. A. R., A unified threat intelligence framework using FTLMS for proactive IoT security in cloud environments, *International Conference on Computing Methodologies and Communication*, Erode, India, 2025, pp. 1105-1112.
<https://doi.org/10.1109/ICCMC65190.2025.11140656>
- [21] B. Wang, Y. Yao, S. Shan, H. Li, B. Viswanath, H. Zheng, B. Y. Zhao, Neural cleanse: Identifying and mitigating backdoor attacks in neural networks, *IEEE Symposium on Security and Privacy*, San Francisco, CA, USA, 2019, pp. 707-723.
<https://doi.org/10.1109/SP.2019.00031>
- [22] S. Rose, O. Borchert, S. Mitchell, S. Connelly, zero trust architecture, National Institute of Standards and Technology, Report No. NIST SP 800-207, August, 2020.
<https://doi.org/10.6028/NIST.SP.800-207>
- [23] X. Yuan, P. He, Q. Zhu, X. Li, Adversarial examples: Attacks and defenses for deep learning, *IEEE Transactions on Neural Networks and Learning Systems*, Vol. 30, No. 9, pp. 2805-2824, September, 2019.
<https://doi.org/10.1109/TNNLS.2018.2886017>
- [24] P. J. Sun, Privacy protection and data security in cloud computing: A survey, challenges, and solutions, *IEEE Access*, Vol. 7, pp. 147420-147452, October, 2019.
<https://doi.org/10.1109/ACCESS.2019.2946185>
- [25] R. Y. Abdullah, A. M. Posonia, U. B. Nisha, An enhanced anomaly forecasting in distributed wireless sensor network using fuzzy model, *International Journal of Fuzzy Systems*, Vol. 24, No. 7, pp. 3327-3347, October, 2022.
<https://doi.org/10.1007/s40815-022-01349-1>
- [26] M. Abadi, A. Chu, I. Goodfellow, H. B. McMahan, I. Mironov, K. Talwar, L. Zhang, Deep learning with differential privacy, *ACM SIGSAC Conference on Computer and Communications Security*, Vienna, Austria, 2016, pp. 308-318.
<https://doi.org/10.1145/2976749.2978318>
- [27] M. Bartock, M. Souppaya, R. Savino, T. Knoll, U. Shetty, M. Cherfaoui, R. Yeluri, A. Malhotra, D. Banks, M. Jordan, D. Pendarakis, J. R. Rao, P. Romness, K. Scarfone, *Hardware-enabled security: Enabling a layered approach to platform security for cloud and edge computing use cases*, National Institute of Standards and Technology, Report No. NISTIR 8320, May, 2022.
<https://doi.org/10.6028/NIST.IR.8320>
- [28] A. Vassilev, A. Oprea, A. Fordyce, H. Andersen, *Adversarial machine learning: A taxonomy and terminology of attacks and mitigations*, Report No. NIST AI 100-2e2023, January, 2024.
<https://doi.org/10.6028/NIST.AI.100-2e2023>
- [29] G. S. R. Emil Selvan, R. Ganeshan, I. D. J. Jingle, J. P. Ananth, FACVO-DNFN: Deep learning-based feature fusion and distributed denial of service attack detection in cloud computing, *Knowledge-Based Systems*, Vol. 261, Article No. 110132, February, 2023.
<https://doi.org/10.1016/j.knosys.2022.110132>
- [30] S. Batewela, P. Ranaweera, M. Liyanage, E. Zeydan, M. Ylianttila, Addressing security orchestration challenges in next-generation networks: A comprehensive overview, *IEEE Open Journal of the Computer Society*, Vol. 6, pp. 669-687, April, 2025.
<https://doi.org/10.1109/OJCS.2025.3564788>
- [31] M. Souppaya, K. Scarfone, D. Dodson, *Secure software development framework (SSDF) version 1.1*, National Institute of Standards and Technology, Report No. NIST SP 800-218, February, 2022.
<https://doi.org/10.6028/NIST.SP.800-218>
- [32] M. Souppaya, J. Morello, K. Scarfone, *Application container security guide*, National Institute of Standards and Technology, Report No. NIST SP 800-190, September, 2017.
<https://doi.org/10.6028/NIST.SP.800-190>
- [33] J. Boyens, A. Smith, N. Bartol, K. Winkler, A. Holbrook, M. Fallon, *Cybersecurity supply chain risk management practices for systems and organizations*, National Institute of Standards and Technology, Report No. NIST SP 800-161r1, May, 2022.
<https://doi.org/10.6028/NIST.SP.800-161r1>
- [34] National Institute of Standards and Technology, *Security and privacy controls for information systems and organizations*, National Institute of Standards and Technology, Report No. NIST SP 800-53 Rev. 5, September, 2020.
<https://doi.org/10.6028/NIST.SP.800-53r5>
- [35] P. Cichonski, T. Millar, T. Grance, K. Scarfone, *Computer security incident handling guide*, National Institute of Standards and Technology, Report No. NIST SP 800-61 Rev. 2, August, 2012.
<http://dx.doi.org/10.6028/NIST.SP.800-61r2>
- [36] T. Orekondy, B. Schiele, M. Fritz, Knockoff nets: Stealing functionality of black-box models, *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, Long Beach, CA, USA, 2019, pp. 4949-4958.
<https://doi.org/10.1109/CVPR.2019.00509>
- [37] K. Dempsey, N. S. Chawla, A. Johnson, R. Johnston, A. C. Jones, A. Orebaugh, M. Scholl, K. Stine, *Information security continuous monitoring (ISCM) for federal*

- information systems and organizations, National Institute of Standards and Technology, Report No. NIST SP 800-137, September, 2011.
- [38] B. Biggio, G. Fumera, F. Roli, Security evaluation of pattern classifiers under attack, *IEEE Transactions on Knowledge and Data Engineering*, Vol. 26, No. 4, pp. 984-996, April, 2014.
<https://doi.org/10.1109/TKDE.2013.57>
- [39] National Institute of Standards and Technology, *Guide for conducting risk assessments*, National Institute of Standards and Technology, Report No. NIST SP 800-30 Rev. 1, September, 2012.
- [40] P. A. Grassi, M. E. Garcia, J. L. Fenton, *Digital identity guidelines*, National Institute of Standards and Technology, Report No. NIST SP 800-63-3, June, 2017.
<https://doi.org/10.6028/NIST.SP.800-63-3>
- [41] N. Rieke, J. Hancox, W. Li, F. Milletari, H. R. Roth, S. Albarqouni, S. Bakas, M. N. Galtier, B. A. Landman, K. Maier-Hein, S. Ourselin, M. Sheller, R. M. Summers, A. Trask, D. Xu, M. Baust, M. J. Cardoso, The future of digital health with federated learning, *npj Digital Medicine*, Vol. 3, Article No. 119, September, 2020.
<https://doi.org/10.1038/s41746-020-00323-1>
- [42] I. Mironov, Rényi differential privacy, *IEEE Computer Security Foundations Symposium*, Santa Barbara, CA, USA, 2017, pp. 263-275.
<https://doi.ieeecomputersociety.org/10.1109/CSF.2017.11>
- [43] K. Kent, M. Souppaya, *Guide to computer security log management*, National Institute of Standards and Technology, Report No. NIST SP 800-92, September, 2006.
<https://doi.org/10.6028/NIST.SP.800-92>
- [44] D. Shivaramakrishna, M. Nagaratna, A novel hybrid cryptographic framework for secure data storage in cloud computing: integrating AES-OTP and RSA with adaptive key management and time-limited access control, *Alexandria Engineering Journal*, Vol. 84, pp. 275-284, December, 2023.
<https://doi.org/10.1016/j.aej.2023.10.054>
- [45] S. M. Moosavi-Dezfooli, A. Fawzi, O. Fawzi, P. Frossard, Universal adversarial perturbations, *IEEE Conference on Computer Vision and Pattern Recognition*, Honolulu, HI, USA, 2017, pp. 86-94.
<https://doi.ieeecomputersociety.org/10.1109/CVPR.2017.17>
- [46] M. Pawlicki, R. S. Choraś, Preprocessing pipelines including block-matching convolutional neural network for image denoising to robustify deep reidentification against evasion attacks, *Entropy*, Vol. 23, No. 10, Article No. 1304, October, 2021.
<https://doi.org/10.3390/e23101304>
- [47] N. Papernot, P. McDaniel, I. Goodfellow, S. Jha, Z. B. Celik, A. Swami, Practical black-box attacks against machine learning, *ACM Asia Conference on Computer and Communications Security*, Abu Dhabi, UAE, 2017, pp. 506-519.
<https://doi.org/10.1145/3052973.3053009>
- [48] K. Eykholt, I. Evtimov, E. Fernandes, B. Li, A. Rahmati, C. Xiao, A. Prakash, T. Kohno, D. Song, Robust physical-world attacks on deep learning visual classification, *IEEE/CVF Conference on Computer Vision and Pattern Recognition*, Salt Lake City, UT, USA, 2018, pp. 1625-1634.
<https://doi.org/10.1109/CVPR.2018.00175>
- [49] M. Sharif, S. Bhagavatula, L. Bauer, M. K. Reiter, Accessorize to a crime: Real and stealthy attacks on state-of-the-art face recognition, *ACM SIGSAC Conference on Computer and Communications Security*, Vienna, Austria, 2016, pp. 1528-1540.
<https://doi.org/10.1145/2976749.2978392>
- [50] H. Azath, M. D. Mani, G. K. D. P. Venkatesan, D. Sivakumar, J. P. Ananth, S. Kamalraj, Identification of IoT device from network traffic using artificial intelligence based capsule networks, *Wireless Personal Communications*, Vol. 123, No. 3, pp. 2227-2243, April, 2022.
<https://doi.org/10.1007/s11277-021-09236-y>
- [51] B. Kim, Y. E. Sagduyu, K. Davaslioglu, T. Erpek, S. Ulukus, Channel-aware adversarial attacks against deep learning-based wireless signal classifiers, *IEEE Transactions on Wireless Communications*, Vol. 21, No. 6, pp. 3868-3880, June, 2022.
<https://doi.org/10.1109/TWC.2021.3124855>
- [52] S. Zhang, Y. Yang, Z. Zhou, Z. Sun, Y. Lin, DIBAD: A disentangled information bottleneck adversarial defense method using Hilbert-Schmidt independence criterion for spectrum security, *IEEE Transactions on Information Forensics and Security*, Vol. 19, pp. 3879-3891, March, 2024.
<https://doi.org/10.1109/TIFS.2024.3372798>
- [53] J. Ma, J. Zhang, G. Shen, A. Marshall, C.-H. Chang, Adversarial attacks against deep learning-based radio frequency fingerprint identification, *IEEE Transactions on Mobile Computing*, Vol. 25, No. 6, pp. 7831-7844, June, 2026.
<https://doi.org/10.1109/TMC.2025.3646257>
- [54] M. Gorsline, J. Smith, C. Merkel, On the adversarial robustness of quantized neural networks, *Great Lakes Symposium on VLSI*, Virtual Event, 2021, pp. 189-194.
<https://doi.org/10.1145/3453688.3461755>
- [55] A. Jordao, H. Pedrini, On the effect of pruning on adversarial robustness, *IEEE/CVF International Conference on Computer Vision Workshops*, Montreal, Canada, 2021, pp. 1-11.
<https://doi.org/10.1109/ICCVW54120.2021.00007>
- [56] R. Chandramouli, *Security strategies for microservices-based application systems*, National Institute of Standards and Technology, Report No. NIST SP 800-204, August, 2019.
<https://doi.org/10.6028/NIST.SP.800-204>
- [57] S. Brooks, M. Garcia, N. Lefkovitz, S. Lightman, E. Nadeau, *An introduction to privacy engineering and risk management in federal systems*, National Institute of Standards and Technology, Report No. NISTIR 8062, January, 2017.
<https://doi.org/10.6028/NIST.IR.8062>
- [58] M. Nasr, R. Shokri, A. Houmansadr, Comprehensive privacy analysis of deep learning: Passive and active white-box inference attacks against centralized and federated learning, *IEEE Symposium on Security and Privacy*, San Francisco, CA, USA, 2019, pp. 739-753.
<https://doi.ieeecomputersociety.org/10.1109/SP.2019.00065>
- [59] M. Fredrikson, S. Jha, T. Ristenpart, Model inversion attacks that exploit confidence information and basic countermeasures, *ACM SIGSAC Conference on Computer and Communications Security*, Denver, CO, USA, 2015, pp. 1322-1333.
<https://doi.org/10.1145/2810103.2813677>
- [60] K. Ganju, Q. Wang, W. Yang, C. A. Gunter, N. Borisov, Property inference attacks on fully connected neural networks using permutation invariant representations, *ACM SIGSAC Conference on Computer and Communications Security*, Vienna, Austria, 2016, pp. 1528-1540.
<https://doi.org/10.1145/2976749.2978392>

- Security*, Toronto, Canada, 2018, pp. 619-633.
<https://doi.org/10.1145/3243734.3243834>
- [61] G. Ateniese, L. V. Mancini, A. Spognardi, A. Villani, D. Vitali, G. Felici, Hacking smart machines with smarter ones: How to extract meaningful data from machine learning classifiers, *International Journal of Security and Networks*, Vol. 10, No. 3, pp. 137-150, September, 2015.
<https://doi.org/10.1504/IJSN.2015.071829>
- [62] C. Song, T. Ristenpart, V. Shmatikov, Machine learning models that remember too much, *ACM SIGSAC Conference on Computer and Communications Security*, Dallas, TX, USA, 2017, pp. 587-601.
<https://doi.org/10.1145/3133956.3134077>
- [63] H. Chabanne, J.-L. Danger, L. Guiga, U. Kühne, Side channel attacks for architecture extraction of neural networks, *CAAI Transactions on Intelligence Technology*, Vol. 6, No. 1, pp. 3-16, March, 2021.
<https://doi.org/10.1049/cit2.12026>
- [64] M. T. Ribeiro, S. Singh, C. Guestrin, "Why should I trust you": Explaining the predictions of any classifier, *ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco, CA, USA, 2016, pp. 1135-1144.
<https://doi.org/10.1145/2939672.2939778>
- [65] Ú. Erlingsson, V. Pihur, A. Korolova, RAPPOR: Randomized aggregatable privacy-preserving ordinal response, *ACM SIGSAC Conference on Computer and Communications Security*, Scottsdale, AZ, USA, 2014, pp. 1054-1067.
<https://doi.org/10.1145/2660267.2660348>
- [66] F. McSherry, K. Talwar, Mechanism design via differential privacy, *IEEE Symposium on Foundations of Computer Science*, Providence, RI, USA, 2007, pp. 94-103.
<https://doi.org/10.1109/FOCS.2007.66>
- [67] K. Bonawitz, V. Ivanov, B. Kreuter, A. Marcedone, H. B. McMahan, S. Patel, D. Ramage, A. Segal, K. Seth, Practical secure aggregation for privacy-preserving machine learning, *ACM SIGSAC Conference on Computer and Communications Security*, Dallas, TX, USA, 2017, pp. 1175-1191.
<https://doi.org/10.1145/3133956.3133982>
- [68] M. Barreno, B. Nelson, A. D. Joseph, J. D. Tygar, The security of machine learning, *Machine Learning*, Vol. 81, No. 2, pp. 121-148, November, 2010.
<https://doi.org/10.1007/s10994-010-5188-5>
- [69] M. K. Ahmed, M. P. Kealoha, J. M. Mbongue, S. K. Saha, E. N. Tchinda, P. E. Mbua, C. Bobda, Multi-tenant cloud FPGA: A survey on security, trust, and privacy, *ACM Transactions on Reconfigurable Technology and Systems*, Vol. 18, No. 2, Article No. 23, June, 2025.
<https://doi.org/10.1145/3713078>
- [70] J. R. Correia-Silva, R. F. Berriel, C. Badue, A. F. De Souza, T. Oliveira-Santos, Copycat CNN: Are random non-labeled data enough to steal knowledge from black-box models, *Pattern Recognition*, Vol. 113, Article No. 107830, May, 2021.
<https://doi.org/10.1016/j.patcog.2021.107830>
- [71] M. Juuti, S. Szyller, S. Marchal, N. Asokan, PRADA: Protecting against DNN model stealing attacks, *IEEE European Symposium on Security and Privacy*, Stockholm, Sweden, 2019, pp. 512-527.
<https://doi.org/10.1109/EuroSP.2019.00044>
- [72] Y. Uchida, Y. Nagai, S. Sakazawa, S. Satoh, Embedding watermarks into deep neural networks, *ACM International Conference on Multimedia Retrieval*, Bucharest, Romania, 2017, pp. 269-277.
<https://doi.org/10.1145/3078971.3078974>
- [73] H. Chen, B. D. Rouhani, C. Fu, J. Zhao, F. Koushanfar, DeepMarks: A secure fingerprinting framework for digital rights management of deep learning models, *ACM International Conference on Multimedia Retrieval*, Ottawa, Canada, 2019, pp. 105-113.
<https://doi.org/10.1145/3323873.3325042>
- [74] R. Uddin, S. A. P. Kumar, V. Chamola, Denial of service attacks in edge computing layers: Taxonomy, vulnerabilities, threats and solutions, *Ad Hoc Networks*, Vol. 152, Article No. 103322, January, 2024.
<https://doi.org/10.1016/j.adhoc.2023.103322>
- [75] M. M. Salim, S. Rathore, J. H. Park, Distributed denial of service attacks and its defenses in IoT: A survey, *Journal of Supercomputing*, Vol. 76, No. 7, pp. 5320-5363, July, 2020.
<https://doi.org/10.1007/s11227-019-02945-z>
- [76] J. Samuel, N. Mathewson, J. Cappos, R. Dingleline, Survivable key compromise in software update systems, *ACM SIGSAC Conference on Computer and Communications Security*, Chicago, IL, USA, 2010, pp. 61-72.
<https://doi.org/10.1145/1866307.1866315>
- [77] R. Nagaraju, C. Venkatesan, J. Kalaivani, G. Manju, S. B. Goyal, C. Verma, C. O. Safirescu, T. C. Mihaltan, Secure routing-based energy optimization for IoT application with heterogeneous wireless sensor networks, *Energies*, Vol. 15, No. 13, Article No. 4777, July, 2022.
<https://doi.org/10.3390/en15134777>
- [78] C. Venkatesan, T. Thamaraimanalan, D. Balamurugan, J. Gowrishankar, R. Manjunath, A. Sivaramakrishnan, Hybrid machine learning technique to detect active botnet attacks for network security and privacy, *Journal of Machine and Computing*, Vol. 3, No. 4, pp. 523-533, October, 2023.
<https://doi.org/10.53759/7669/jmc202303044>
- [79] M. Moore, T. K. Kuppusamy, J. Cappos, Artemis: Defanging software supply chain attacks in multi-repository update systems, *Annual Computer Security Applications Conference*, Austin, TX, USA, 2023, pp. 83-97.
<https://doi.org/10.1145/3627106.3627129>

Biographies



and wireless sensor networks.

Venkatesan Cherappa is a Professor in Electronics and Communication Engineering at HKBK College of Engineering, Bangalore, India. He holds Ph.D. from Anna University, Chennai, and has published over 90 papers. His research focuses on AI, machine learning, future Internet technologies,



Hsin-Hung Cho is a Professor at National I-Lan University, Taiwan, with a Ph.D. in Computer Science. His research includes mobile networks, AI, and information security. He has published about 80 papers, received the IET Premium Award, and serves in editorial roles for several international journals and IEEE publications.



Yasir Abdullah Rabi is an Assistant Professor (SG) at Dr. Mahalingam College of Engineering and Technology, Pollachi. His work focuses on artificial intelligence, machine learning, and data analytics contributing through publications, student mentoring, and collaborative projects.



Ananth John Patrick is a Professor of Computer Science with over 24 years of experience. He has published 70+ papers and authored two books. A Senior IEEE, he serves as Director of IQAC at Dayananda Sagar University. His research focuses on computer vision, machine learning, and artificial

intelligence.