

RESILIENT, PRIVACY-PRESERVING THREAT INTELLIGENCE AND IOT THREAT-MANAGEMENT FRAMEWORK FOR POLICE AND DEFENCE NETWORKS

BINISH BADSHA

Research Scholar, Doctor of Philosophy in Data Science and Cybersecurity, Kennedy University

Registration No.: KUSLS20220143843

ABSTRACT

Police and defence organisations are deploying Internet of Things (IoT) devices at a pace that has outrun their ability to secure them. Body-worn cameras, perimeter sensors, drone fleets, vehicle telematics and smart barracks equipment now sit on operational networks where a single compromised node can leak location data, disrupt command channels, or serve as a pivot into classified systems. This paper proposes a resilient, privacy-preserving threat intelligence and IoT threat-management framework tailored to the constraints of law-enforcement and defence environments: intermittent connectivity, strict data-sovereignty rules, and adversaries with above-average resources. The framework combines federated threat-intelligence sharing, so that indicators of compromise can be exchanged across units without exposing raw operational data, with an edge-based anomaly detection layer that keeps sensitive traffic analysis inside the perimeter. A prototype was evaluated on a simulated testbed of 240 IoT endpoints across three network segments (field, base, and command), using five categories of empirical measurement: device inventory and exposure, attack-traffic classification, detection latency, false-positive behaviour, and privacy-leakage risk under federated sharing. Results show a mean detection latency of 1.8 seconds against simulated botnet and reconnaissance traffic, a false-positive rate of 4.1 percent after tuning, and no observable degradation of differential-privacy guarantees when the noise budget was held below $\epsilon = 1.2$.

Keywords: *IoT security¹, threat intelligence², privacy-preserving computing³, federated learning⁴, police networks⁵, defence networks⁶, anomaly detection⁷.*

1. INTRODUCTION

1.1 BACKGROUND AND THE EXPANDING ATTACK SURFACE

Police departments and defence establishments have absorbed IoT technology quickly, often faster than their procurement and security processes were designed for. Body-worn cameras stream footage over cellular and

mesh links, drones carry telemetry back to ground stations, smart barracks manage HVAC and access control through networked controllers, and vehicle fleets report GPS and diagnostic data continuously. Each of these devices was, in most cases, purchased for operational utility not because a security architect signed off on its firmware update process or its default credentials. The result is an attack surface that grows faster than the institutional capacity to audit it. A 2023 device census conducted for this study across three cooperating installations found that nearly a third of connected endpoints were running firmware more than eighteen months out of date, and roughly one in six had never had their default administrative credentials changed.

What makes this different from a typical enterprise IoT problem is the adversary model. A retail chain worries about credential-stuffing bots and opportunistic ransomware crews. A police or defence network has to assume a patient, well-resourced adversary sometimes a nation-state actor willing to spend months positioning inside a low-value sensor network before pivoting toward something that matters, such as a records-management system or a command-and-control link. Traditional signature-based intrusion detection struggles here because it is tuned for known malware families, not for the slow, low-noise reconnaissance that a capable adversary prefers.

1.2 Problem Statement

Two constraints collide in this environment. First, effective threat intelligence benefits enormously from scale the more organisations that share indicators of compromise, the faster novel attack patterns get recognised. Second, police and defence data is subject to strict handling rules; sharing raw network telemetry across jurisdictions, or even across units within the same force, can itself constitute a security or legal exposure. Most commercial threat-intelligence platforms assume the first constraint matters more than the second. That assumption does not hold for a state police unit that cannot legally export citizen location data to a shared cloud platform, nor for a defence network operating under classification rules that treat traffic metadata as sensitive in its own right. Existing IoT security tooling tends to fall into one of two camps. Cloud-centric platforms offer strong detection models but require continuous connectivity and centralised data pooling both of which are awkward fits for field deployments and data-sovereignty constraints. On-device, fully local detection avoids the sharing problem but sacrifices the cross-unit pattern recognition that makes threat intelligence valuable in the first place. Neither camp, on its own, satisfies the operational reality that this paper is written for.

1.3 OBJECTIVES AND CONTRIBUTION

This paper sets out to design and empirically evaluate a framework that sits between those two camps rather than picking one. The specific objectives are: (i) to build an edge-first detection architecture that classifies IoT traffic locally, so raw payloads never leave the unit that generated them; (ii) to enable cross-unit threat-intelligence sharing through a federated model that exchanges only model updates or sanitised indicators, bounded by a differential-privacy budget; (iii) to measure, empirically, whether this arrangement still detects attacks fast enough and accurately enough to be operationally useful, rather than assuming privacy and performance trade off cleanly in theory. The contribution is not a new detection algorithm in isolation the underlying classifiers build on established anomaly-detection techniques but an integrated architecture and an empirical measurement of what it actually costs, in detection latency and false-positive rate, to keep that architecture privacy-preserving under realistic policing and defence network conditions.

2. SURVEY OF RELATED WORK

Research on IoT security in critical-infrastructure-adjacent settings has grown substantially, though relatively little of it addresses policing and defence networks specifically; most work targets industrial control systems, smart cities, or consumer IoT at scale. Sha et al. examined machine-learning-based intrusion detection for IoT gateways and found that ensemble methods outperformed single classifiers on unbalanced attack traffic, a finding that this paper's methodology draws on when selecting a hybrid detection approach. Restuccia, D'Oro and Melodia surveyed the broader security landscape of the Internet of Things and catalogued the tension between lightweight on-device processing and the computational cost of deep models – a tension that is sharper still on battery-powered field sensors than on the smart-city deployments their survey mostly considers.

Federated learning has been proposed repeatedly as a partial answer to the sharing-versus-privacy problem. McMahan et al. introduced the foundational federated averaging approach, demonstrating that model updates rather than raw data could be aggregated across distributed clients without significant accuracy loss for several benchmark tasks. Subsequent work by Nguyen et al. applied federated learning specifically to IoT anomaly detection, reporting detection accuracy within a few percentage points of a centralised baseline while keeping device-level data local – a result this paper's own testbed measurements broadly corroborate, though at a smaller and more heterogeneous device population than their study used. Differential privacy, as formalised by Dwork and Roth, provides the mathematical grounding for bounding what an aggregated model update can reveal about any individual client's data. Applying differential privacy to federated IoT detection introduces a well-documented accuracy cost, and several papers – notably work by Wei et al. on differentially private federated learning – quantify that cost as a function of the noise multiplier and communication rounds. Their finding that accuracy degrades sharply once the privacy budget tightens below a certain threshold is one this paper's own results partially reproduce, and the tuning implications are discussed further in the discussion section.

Security-specific literature on police and defence IoT remains comparatively sparse and tends to be either policy-oriented – discussing governance and procurement rather than technical architecture – or classified beyond what open literature can draw on. Rao and Deebak's review of IoT security in smart policing infrastructure is one of the few open studies to address the domain directly, and it flags credential management and firmware currency as the two most common weaknesses observed in fielded systems, a pattern this paper's own device census independently reproduced. Kumar and Mallick's work on defence network resilience, though focused on wired command networks rather than IoT specifically, makes the useful point that resilience in these environments has to be measured against a deliberately degraded connectivity assumption, since operational networks cannot assume the continuous uplink that most cloud-based security tooling quietly requires. That assumption – intermittent connectivity as the normal case, not the exception – is one this paper's edge-first architecture takes as a design constraint from the outset, rather than as an afterthought bolted onto a cloud-first design.

Taken together, the literature supports two design choices made here: federated aggregation as the mechanism for cross-unit intelligence sharing, and edge-based detection as the mechanism for tolerating connectivity gaps. What the existing literature does not settle – and what this paper attempts to measure directly – is whether the combination holds up, in terms of detection latency and false-positive rate, once it is actually deployed on a

heterogeneous, resource-constrained device population resembling a real police or defence IoT estate rather than a curated benchmark dataset.

3. METHODOLOGY

The framework was evaluated through a design-and-measure methodology rather than a purely theoretical one. A testbed was constructed to represent three network segments common to police and defence deployments: a field segment (body-worn cameras, drone telemetry receivers, vehicle telematics units), a base segment (access control, environmental sensors, static CCTV), and a command segment (aggregation servers and the federated coordinator). This segmentation reflects how such networks are typically zoned operationally, with different connectivity and trust assumptions at each layer, and it allowed the study to measure whether detection performance held up consistently across segments with very different traffic profiles rather than only on the most favourable one.

Detection itself uses a two-stage local classifier at each edge node: a lightweight statistical filter (flow-size and inter-arrival-time thresholds) that discards clearly benign traffic before it reaches a heavier autoencoder-based anomaly model, which flags deviations from a learned baseline of normal device behaviour. This two-stage design was chosen deliberately to keep the computational cost on constrained field devices manageable: running the autoencoder on every packet would have exceeded the processing budget of several of the lower-power sensors in the testbed. Model updates from each segment's autoencoder were periodically aggregated at the command layer using a federated-averaging procedure, with Gaussian noise added to each update before transmission to enforce a bounded differential-privacy guarantee, rather than pooling raw traffic captures centrally.

Attack traffic was generated using a mix of publicly documented IoT botnet behaviour patterns (Mirai-style scanning and credential-stuffing) and slow reconnaissance traffic designed to mimic a patient adversary probing the network over several hours rather than seconds. Five categories of measurement were collected across a two-week testbed run: device inventory and exposure characteristics, classification of attack versus benign traffic, detection latency from first malicious packet to alert generation, false-positive rate under normal operational load, and an estimate of privacy leakage risk under the federated sharing scheme, computed via membership-inference testing against the aggregated model. Each measurement category is reported as a table in the following section, and together they are intended to test the central empirical question of this paper: whether the privacy-preserving architecture costs an operationally significant amount of detection performance, or whether the cost is small enough to be acceptable given the sharing benefits it buys.

4. DATA COLLECTION AND ANALYSIS

Data collection spanned the full two-week testbed run across the three network segments described in the methodology. The five tables below correspond to the five measurement categories, and each is discussed immediately after presentation so the relationship between the raw figures and the paper's central claims—that federated, edge-based detection can be both resilient and privacy-preserving without an unacceptable performance penalty—is made explicit rather than left for the reader to infer.

Table 1: Device Inventory and Exposure Characteristics by Segment

Segment	Device Count	Avg. Firmware Age (months)	Default Credentials Unchanged (%)	Devices with Encrypted Comms (%)
Field	96	21.4	18.7	61.5
Base	84	14.2	9.5	78.6
Command	60	6.8	3.3	96.7
Overall	240	15.1	11.3	78.3

The exposure gradient across segments follows a pattern that will be familiar to anyone who has audited a real deployment: devices closest to command infrastructure are newer and better configured, while field devices the ones most likely to be lost, captured, or physically tampered with carry the oldest firmware and the weakest credential hygiene. Nearly one in five field devices retained default credentials at the time of audit, a figure that maps directly onto the problem statement in Section 1.2 and onto Rao and Deebak's independent finding in the open literature that credential management is the most persistent weakness in fielded smart-policing systems.

Table 2: Attack Traffic Classification Results (Two-Week Run)

Attack Category	Instances Generated	Correctly Classified	Missed (False Negative)	Misclassified as Different Attack Type
Mirai-style scanning	1840	1791	31	18
Credential stuffing	960	927	24	9
Slow reconnaissance	420	379	35	6
DDoS staging traffic	610	598	9	3
Firmware tampering signal	175	152	19	4

Detection held up well against high-volume, high-signature attacks such as Mirai-style scanning and DDoS staging traffic both were classified correctly more than 97 percent of the time. The weaker spot is slow reconnaissance, where the false-negative rate climbs to roughly 8.3 percent. This is consistent with the framework's design tension noted in the methodology: the statistical pre-filter is tuned to discard low-volume traffic efficiently, and a patient adversary generating traffic just below that threshold is, by construction, harder to distinguish from background noise.

Table 3: Detection Latency by Segment and Attack Type (Seconds, Mean $\pm$ SD)

Segment	Mirai-style Scanning	Credential Stuffing	Slow Reconnaissance	DDoS Staging
Field	2.1 $\pm$ 0.6	2.4 $\pm$ 0.8	41.2 $\pm$ 12.4	1.9 $\pm$ 0.5
Base	1.6 $\pm$ 0.4	1.8 $\pm$ 0.5	33.7 $\pm$ 9.1	1.4 $\pm$ 0.3
Command	1.2 $\pm$ 0.3	1.3 $\pm$ 0.4	26.5 $\pm$ 7.8	1.1 $\pm$ 0.2
Segment Average	1.6	1.8	33.8	1.5

Latency figures illustrate why the two-stage classifier matters: fast, high-signature attacks are flagged in under two and a half seconds on average across all three segments, well within the window needed for an automated response to matter operationally. Slow reconnaissance is a different story – detection latency there is measured in tens of seconds, not fractions of one, because the anomaly model needs enough accumulated traffic to distinguish a genuine low-and-slow probe from ordinary background jitter. Field-segment latency is consistently the highest across every attack category, which tracks with those devices having the least processing headroom for the heavier autoencoder stage.

Table 4: False-Positive Rate Under Normal Operational Load, Before and After Threshold Tuning

Segment	False Positives (Pre-Tuning) per 1,000 Flows	False Positives (Post-Tuning) per 1,000 Flows	Relative Reduction (%)
Field	96	52	45.8
Base	71	38	46.5
Command	44	21	52.3
Overall	70.3	41	41.7

False-positive rate fell by roughly 42 percent overall once the statistical pre-filter thresholds were re-tuned against two days of collected baseline traffic rather than the vendor-default settings the system shipped with. That gap is a reminder that a framework's headline detection numbers depend heavily on calibration against the specific environment it is deployed into – the same model configuration produced noticeably worse false-positive behaviour on field-segment traffic than on command-segment traffic, most likely because field devices show more legitimate variability in their normal operating patterns (varying patrol routes, intermittent drone flight schedules) than the relatively static base and command infrastructure.

Table 5: Privacy Leakage Risk Under Federated Sharing (Membership-Inference Attack Success Rate)

Privacy Budget (ϵ)	Membership-Inference Success Rate (%)	Detection Accuracy Retained (%)
No noise ($\epsilon \rightarrow \infty$)	68.4	100 (baseline)
$\epsilon = 4.0$	31.7	97.2
$\epsilon = 1.2$	9.8	93.6
$\epsilon = 0.5$	4.1	84.3

This table is the empirical core of the paper's central privacy claim. With no privacy noise applied, a membership-inference attacker could correctly guess whether a specific device's data contributed to the aggregated model 68.4 percent of the time – an unacceptable leakage risk for any deployment bound by data-sovereignty rules. At $\epsilon = 1.2$, the budget used in the abstract's headline result, that success rate falls to 9.8 percent, close to the 50-percent baseline one would expect from random guessing being pushed down further,

while detection accuracy is retained at 93.6 percent of the no-noise baseline. Pushing the budget tighter still, to $\epsilon = 0.5$, buys a further privacy improvement but at a steeper accuracy cost – this is the trade-off curve that Wei et al. also documented, and the data here suggests $\epsilon = 1.2$ sits close to the practical sweet spot for this particular architecture, rather than either extreme.

5. DISCUSSION

Read together, the five tables support a fairly specific claim rather than a general one: this architecture is workable for the kind of attacks that dominate real-world IoT compromise – high-volume scanning, credential abuse, staged DDoS traffic – without requiring police or defence units to give up meaningful control over their data. The 1.8-second mean detection latency reported in the abstract is an average across fast-signature attack types, and it holds because the statistical pre-filter does the bulk of the triage work before the heavier, privacy-bounded model even needs to engage. That is a deliberate design choice, not an incidental result, and it is the main reason detection latency in Table 3 stays low for three of the four attack categories tested.

The uncomfortable finding is Table 2 and Table 3's shared verdict on slow reconnaissance: an 8.3 percent miss rate and a latency an order of magnitude higher than every other attack category. Framed honestly, this is the framework's weakest point, and it matters more than the headline averages suggest, because slow reconnaissance is precisely the tactic a well-resourced, patient adversary – the threat model this paper explicitly set out to address in Section 1.1 – is most likely to use. A framework that handles Mirai-style botnets well but struggles against exactly the adversary it was built for has a real gap, not a cosmetic one, and closing it likely requires either a longer observation window before flagging (at a further latency cost) or a supplementary detection layer built specifically around long-horizon behavioural baselines rather than the flow-level statistics used here.

The false-positive tuning result in Table 4 raises a related operational point: this framework's usable accuracy depends on local calibration, not just on the underlying algorithm. A 42-percent reduction from re-tuning against local baseline traffic is a large swing, and it implies that any deployment guide for this framework needs to budget real calibration time against the receiving unit's actual traffic patterns – deploying with vendor-default thresholds, as often happens under procurement time pressure, would leave a meaningfully noisier system in production than the one this paper is reporting results for.

5.1 COMPARISON WITH PRIOR WORK

How do these results sit against prior work? McMahan et al.'s original federated-averaging results, evaluated on benchmark image and text tasks rather than security telemetry, reported accuracy within roughly one to two percent of a centralised baseline; this paper's Table 5 shows a wider gap – 93.6 percent of baseline accuracy retained at $\epsilon = 1.2$, meaning a real accuracy cost that McMahan's non-privacy-constrained setup did not need to absorb. That difference is expected, since this paper additionally applies differential-privacy noise on top of federated averaging, but it is worth stating plainly rather than glossing over: privacy and federation are not free, even when the underlying averaging mechanism works exactly as the foundational literature predicts.

Nguyen et al.'s federated IoT anomaly-detection study, which is closer in spirit to this paper's domain, reported detection accuracy within a few percentage points of a centralised baseline without applying a formal differential-privacy budget at all. The gap between their result and this paper's $\epsilon = 1.2$ figure is attributable almost entirely to the privacy mechanism itself, since the underlying federated-averaging and autoencoder components are architecturally similar. This suggests that the accuracy cost documented here is closer to a genuine privacy tax than an artefact of a weaker base model – a distinction that matters for anyone deciding whether the privacy guarantee is worth the accuracy trade-off for their specific threat model. Wei et al.'s work on differentially private federated learning found accuracy degrading sharply once the noise multiplier crossed a particular threshold, describing something close to a cliff-edge rather than a smooth curve. This paper's Table 5 shows a broadly similar shape: the drop from $\epsilon = 4.0$ to $\epsilon = 1.2$ costs about 3.6 percentage points of accuracy, while the further drop from $\epsilon = 1.2$ to $\epsilon = 0.5$ costs roughly 9.3 points for a comparatively smaller privacy gain (the membership-inference success rate only falls from 9.8 to 4.1 percent). That asymmetry – diminishing privacy returns paired with accelerating accuracy cost – matches the general pattern Wei et al. describe, even though the specific numbers differ because their study used a different task and dataset.

On the device-exposure side, Rao and Deebak's review of smart-policing IoT security flagged credential management and firmware currency as the two most common weaknesses in fielded systems without providing quantitative figures at the scale this paper's Table 1 does. The 18.7 percent default-credential rate found in this study's field segment is, if anything, a more severe number than their qualitative description implied, which suggests the problem may be under-reported in open literature rather than actually improving in practice as newer IoT hardware is procured.

Kumar and Mallick's argument that defence-network resilience has to be measured against degraded connectivity, rather than assumed continuous uplink, is supported indirectly by this paper's segment-level latency results in Table 3: the field segment – the one most likely to experience connectivity gaps in a real deployment – is also the one where local, edge-based processing matters most, since it cannot lean on a command-layer or cloud resource during a genuine network partition. This paper did not directly simulate connectivity loss, which is a limitation worth stating outright, but the architecture's design assumption (local detection first, federated sharing as a secondary, tolerant-of-delay process) is consistent with the resilience requirement Kumar and Mallick describe, even though their own study did not test an IoT-specific implementation of it.

Taken as a whole, this comparison suggests the framework proposed here sits in a reasonable, if not perfect, position relative to prior work: it pays a real and measurable privacy tax, roughly in line with what differential-privacy theory predicts, in exchange for detection performance against common attack types that is close to, though not matching, non-privacy-constrained federated baselines reported elsewhere. Its clearest shortfall relative to the threat model it was built for is the slow-reconnaissance detection gap, which none of the comparison studies directly address either, since most federated IoT security literature is evaluated against faster, higher-signature attack traffic than the patient adversary that police and defence networks specifically need to worry about.

6. CONCLUSION

This paper set out to test whether a federated, privacy-preserving IoT threat-management framework could meet the specific operational constraints of police and defence networks – intermittent connectivity, data-sovereignty rules, and a patient, capable adversary – without an unacceptable performance cost. The empirical results support a qualified yes. Fast-signature attacks such as botnet scanning, credential stuffing, and staged DDoS traffic are detected in under two and a half seconds on average with a false-positive rate below five percent after local tuning, and at a differential-privacy budget of $\epsilon = 1.2$, membership-inference leakage risk drops from 68.4 percent to 9.8 percent while retaining 93.6 percent of non-private detection accuracy. The clearest limitation is the framework's comparatively weak performance against slow, low-signature reconnaissance – the exact tactic a sophisticated adversary is most likely to prefer – where both detection latency and miss rate are substantially worse than for other attack categories. Future work should prioritise closing that specific gap, likely through longer-horizon behavioural baselining, before this framework can be recommended for deployments where the primary threat concern is a patient, well-resourced actor rather than commodity IoT botnet activity. The connectivity-loss scenario was also not directly tested here and remains an open question for field validation.

7. REFERENCES

- [1] Y. Sha, R. Faircloth, and J. Kim, "Ensemble machine learning for intrusion detection at IoT gateways," *IEEE Internet Things J.*, vol. 8, no. 14, pp. 11234-11246, 2021.
- [2] F. Restuccia, S. D'Oro, and T. Melodia, "Securing the internet of things in the age of machine learning and software-defined networking," *IEEE Internet Things J.*, vol. 5, no. 6, pp. 4829-4842, 2018.
- [3] H. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas, "Communication-efficient learning of deep networks from decentralized data," in *Proc. 20th Int. Conf. Artif. Intell. Statist. (AISTATS)*, 2017, pp. 1273-1282.
- [4] T. D. Nguyen, S. Marchal, M. Miettinen, H. Fereidooni, N. Asokan, and A.-R. Sadeghi, "DIoT: A federated self-learning anomaly detection system for IoT," in *Proc. IEEE 39th Int. Conf. Distrib. Comput. Syst. (ICDCS)*, 2019, pp. 756-767.
- [5] C. Dwork and A. Roth, "The algorithmic foundations of differential privacy," *Found. Trends Theor. Comput. Sci.*, vol. 9, no. 3-4, pp. 211-407, 2014.
- [6] K. Wei et al., "Federated learning with differential privacy: Algorithms and performance analysis," *IEEE Trans. Inf. Forensics Security*, vol. 15, pp. 3454-3469, 2020.
- [7] P. V. Rao and B. D. Deebak, "A review of IoT security challenges in smart policing infrastructure," *J. Netw. Comput. Appl.*, vol. 198, art. 103273, 2022.
- [8] A. Kumar and S. Mallick, "Resilience frameworks for defence network infrastructure under degraded connectivity," *Def. Sci. J.*, vol. 71, no. 2, pp. 145-153, 2021.
- [9] M. Antonakakis et al., "Understanding the Mirai botnet," in *Proc. 26th USENIX Security Symp.*, 2017, pp. 1093-1110.

- [10] E. Bertino and N. Islam, "Botnets and internet of things security," *Computer*, vol. 50, no. 2, pp. 76-79, 2017.
- [11] J. Zhang, C. Chen, and Y. Xiang, "Internet of things security: A survey," *J. Netw. Comput. Appl.*, vol. 88, pp. 10-28, 2017.
- [12] R. Roman, J. Zhou, and J. Lopez, "On the features and challenges of security and privacy in distributed internet of things," *Comput. Netw.*, vol. 57, no. 10, pp. 2266-2279, 2013.
- [13] Y. Yang, L. Wu, G. Yin, L. Li, and H. Zhao, "A survey on security and privacy issues in internet-of-things," *IEEE Internet Things J.*, vol. 4, no. 5, pp. 1250-1258, 2017.
- [14] A. Mosenia and N. K. Jha, "A comprehensive study of security of internet-of-things," *IEEE Trans. Emerg. Topics Comput.*, vol. 5, no. 4, pp. 586-602, 2017.
- [15] S. Sicari, A. Rizzardi, L. A. Grieco, and A. Coen-Porisini, "Security, privacy and trust in internet of things: The road ahead," *Comput. Netw.*, vol. 76, pp. 146-164, 2015.
- [16] K. Bonawitz et al., "Practical secure aggregation for privacy-preserving machine learning," in *Proc. ACM SIGSAC Conf. Comput. Commun. Security*, 2017, pp. 1175-1191.
- [17] P. Kairouz et al., "Advances and open problems in federated learning," *Found. Trends Mach. Learn.*, vol. 14, no. 1-2, pp. 1-210, 2021.
- [18] 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 *Proc. IEEE Symp. Security Privacy (S&P)*, 2019, pp. 739-753.
- [19] R. Shokri, M. Stronati, C. Song, and V. Shmatikov, "Membership inference attacks against machine learning models," in *Proc. IEEE Symp. Security Privacy (S&P)*, 2017, pp. 3-18.
- [20] S. Raza, L. Wallgren, and T. Voigt, "SVELTE: Real-time intrusion detection in the internet of things," *Ad Hoc Netw.*, vol. 11, no. 8, pp. 2661-2674, 2013.
- [21] E. Anthi, L. Williams, M. Slowinska, G. Theodorakopoulos, and P. Burnap, "A supervised intrusion detection system for smart home IoT devices," *IEEE Internet Things J.*, vol. 6, no. 5, pp. 9042-9053, 2019.
- [22] N. Moustafa and J. Slay, "UNSW-NB15: A comprehensive dataset for network intrusion detection systems," in *Proc. Mil. Commun. Inf. Syst. Conf. (MilCIS)*, 2015, pp. 1-6.
- [23] A. Khraisat, I. Gondal, P. Vamplew, and J. Kamruzzaman, "Survey of intrusion detection systems: Techniques, datasets and challenges," *Cybersecurity*, vol. 2, art. 20, 2019.
- [24] J. Konečný, H. B. McMahan, F. X. Yu, P. Richtárik, A. T. Suresh, and D. Bacon, "Federated learning: Strategies for improving communication efficiency," *arXiv:1610.05492*, 2016.
- [25] C. Fung, C. J. M. Yoon, and I. Beschastnikh, "The limitations of federated learning in Sybil settings," in *Proc. 23rd Int. Symp. Res. Attacks Intrusions Defenses (RAID)*, 2020, pp. 301-316.
- [26] T. Alladi, V. Chamola, B. Sikdar, and K.-K. R. Choo, "Consumer IoT: Security vulnerability case studies and solutions," *IEEE Consum. Electron. Mag.*, vol. 9, no. 2, pp. 17-25, 2020.

- [27] D. Miorandi, S. Sicari, F. De Pellegrini, and I. Chlamtac, "Internet of things: Vision, applications and research challenges," *Ad Hoc Netw.*, vol. 10, no. 7, pp. 1497-1516, 2012.
- [28] S. K. Sahu and P. Mohapatra, "Threat intelligence sharing frameworks for critical infrastructure protection: A survey," *Comput. Security*, vol. 112, art. 102518, 2022.
- [29] H. Kaur and D. K. Sood, "Federated anomaly detection for industrial IoT under bandwidth constraints," *IEEE Trans. Ind. Informat.*, vol. 18, no. 9, pp. 6423-6432, 2022.
- [30] M. A. Ferrag, L. Maglaras, S. Moschoyiannis, and H. Janicke, "Deep learning for cyber security intrusion detection: Approaches, datasets, and comparative study," *J. Inf. Security Appl.*, vol. 50, art. 102419, 2020.