



Federated Cyber-Threat Learning Across Independent Institutions

Dr. Samuel J. Rhodes

Associate Professor, University of Toronto, Canada

Abstract

Independent institutions frequently observe different fragments of the same evolving cyber-threat landscape. Universities may detect credential abuse, healthcare providers may encounter ransomware precursors, municipalities may experience exploitation of internet-facing services, and financial or infrastructure organizations may observe command-and-control or exfiltration behavior. Conventional cyber-threat intelligence sharing allows institutions to exchange indicators and defensive information, but raw security telemetry is often difficult to centralize because it may contain confidential operational information, personal data, proprietary network characteristics, or sensitive incident evidence. Federated learning provides an alternative model in which institutions collaboratively improve machine-learning models while retaining underlying training observations within their own administrative boundaries. Federated learning, however, does not automatically provide privacy or trust. Model updates may disclose information, institutional data may be statistically heterogeneous, malicious participants may poison global models, and privacy-preserving aggregation can complicate the identification of malicious contributions. NIST's privacy-preserving federated-learning work explicitly distinguishes decentralized training from stronger privacy mechanisms such as secure aggregation and output privacy, while recent federated cybersecurity research continues to identify poisoning, Byzantine behavior, inference leakage, and non-IID data as central deployment challenges.

This study proposes a Secure Adaptive Federated Threat Learning Framework for cross-institution cybersecurity collaboration. A synthetic environment representing 12 independent institutions and 144,000 network-event windows was constructed. Four collaboration conditions were modeled: isolated local learning, machine-readable cyber-threat intelligence sharing, conventional cross-silo federated averaging, and secure adaptive federation integrating secure aggregation, institutional trust scoring, robust update screening, local personalization, and structured STIX/TAXII-compatible threat context. The strongest architecture achieved a simulated macro-F1 score of 91.3%, an unseen-threat recall of 84.6%, and a 16-point privacy-exposure index compared with 38 for conventional federated learning. Under a synthetic scenario in which 20% of participating institutional updates were malicious or corrupted, the secure adaptive architecture retained 87.8% macro-F1 compared with 68.9% for unprotected federated averaging.

Keywords: federated learning, cyber-threat intelligence, intrusion detection, collaborative cybersecurity, secure aggregation, cross-silo learning, threat intelligence sharing, poisoning attacks, privacy-preserving machine learning, institutional cybersecurity



1. Introduction

Cyber threats routinely cross organizational boundaries even though cybersecurity data remain organizationally fragmented. A phishing infrastructure detected by one university may later target healthcare providers; malicious command-and-control infrastructure observed by one financial institution may appear in another sector; and malware behavior first encountered in a municipal network may later emerge within a supplier ecosystem. Conventional cyber-threat intelligence sharing attempts to reduce this asymmetry by allowing organizations to exchange indicators, defensive measures, adversary techniques, and contextual observations. CISA's Automated Indicator Sharing program, for example, supports machine-readable exchange of cyber-threat indicators and defensive measures among participating organizations, while OASIS STIX and TAXII standards provide structured representations and transport mechanisms for cyber-threat information.

Indicator sharing nevertheless captures only part of defensive knowledge. An institution's detection system may learn statistical relationships among timing, network behavior, authentication events, endpoint characteristics, and attack sequences that cannot be reduced to a single malicious IP address or file hash. Sharing the raw telemetry needed to reproduce those relationships may be impractical or undesirable because security logs frequently contain usernames, internal addresses, operational architecture, timestamps, sensitive transactions, or evidence relevant to ongoing investigations. Federated learning offers a different collaboration model. McMahan and colleagues' foundational federated averaging approach demonstrated that a shared model can be learned through aggregation of locally computed updates while keeping underlying training data distributed. In a cross-institution setting, each organization can therefore train on its own threat observations and contribute model information rather than exporting complete security datasets.

The privacy benefit of federation, however, is conditional rather than automatic. NIST's privacy-preserving federated-learning guidance emphasizes that model updates themselves may require protection and discusses secure aggregation as a technique for allowing a coordinator to learn aggregate information without inspecting individual participant updates. NIST also distinguishes input privacy from output privacy, recognizing that a collaboratively trained model can still reveal information about training data unless additional protections are considered. Differential privacy provides one possible complementary mechanism, and NIST SP 800-226, finalized in 2025, provides guidance for evaluating differential-privacy guarantees. The practical implication is that “data remain local” should not be treated as equivalent to “data are private.”

Cybersecurity creates an additional complication because the participants in a learning federation may themselves become adversarial. A compromised institution can contribute manipulated gradients or model parameters intended to reduce detection accuracy or insert a backdoor into the global detector. Research testbeds specifically developed for cybersecurity federated learning demonstrate that poisoning remains a substantive problem even when local data are never directly shared. Recent work has consequently explored combinations of secure aggregation, robust aggregation, anomaly detection, Byzantine resilience, client reputation, and differential privacy. DDFed, for example, explicitly addresses the difficult interaction between privacy-preserving secure aggregation and poisoning detection, because encrypting updates can make conventional anomaly screening harder. The 2026



SecureDyn-FL framework similarly combines secure aggregation, temporal gradient auditing, and personalized learning for non-IID cybersecurity data.

The present study therefore investigates Federated Cyber-Threat Learning Across Independent Institutions through a simulation-based cross-silo architecture. The framework assumes that institutions remain administratively independent, retain their raw telemetry locally, and do not need to expose detailed internal network observations to one another. Collaborative learning occurs through institution-level model updates, while standardized CTI remains available for explainable threat context. The study compares four forms of collaboration and proposes that the strongest architecture combines federated statistical learning with existing STIX/TAXII intelligence exchange, secure aggregation, robust participant governance, local personalization, and explicit defenses against compromised contributors.

2. Objectives of the Study

2.1 To Evaluate the Detection Benefits of Cross-Institution Threat Learning

The first objective is to examine whether institutions exposed to different attack distributions can improve recognition of previously underrepresented or unseen cyber threats by contributing to a common federated detection model without centralizing raw security telemetry.

2.2 To Assess Privacy, Heterogeneity, and Poisoning Risks

The second objective is to evaluate the major risks introduced by institutional federation, including model-update leakage, unequal local datasets, non-IID threat distributions, compromised participants, poisoning attacks, and excessive dependence on a central aggregation coordinator.

2.3 To Develop a Secure Adaptive Federation Architecture

The third objective is to construct a practical framework integrating secure aggregation, robust update screening, participant trust assessment, personalized local adaptation, machine-readable CTI exchange, differential-privacy options, and governance rules for joining, suspending, and auditing federation members.

3. Review of Literature

3.1 Federated Learning for Distributed Intrusion and Threat Detection

Federated learning enables several data owners to train a shared model without transferring the underlying training datasets to a central repository. Federated Averaging remains a foundational approach, using iterative local training followed by aggregation of institution- or client-level model updates. Kairouz and colleagues subsequently characterized federated learning as a broader distributed-learning paradigm and identified privacy, communication efficiency, heterogeneity, robustness, fairness, and personalization among its principal open challenges. These characteristics closely match cross-institution cybersecurity environments, where security telemetry is naturally decentralized and operationally heterogeneous.

Cybersecurity applications increasingly exploit this architecture. FedCTI combines federated learning with cyber-threat intelligence to support privacy-preserving knowledge sharing in distributed networks. Federated-learning intrusion-detection research similarly shows that multiple organizations or



network domains can contribute to common models without direct data pooling. Reviews of FL-enabled intrusion-detection systems identify aggregation strategy, dataset heterogeneity, privacy, communication cost, and attack robustness as recurring technical issues.

Cross-institution environments are particularly likely to contain non-identically distributed data. A university network contains device, user, and traffic patterns that differ from those of a healthcare organization or municipal administration. Standard federated averaging can therefore produce a global model that performs unevenly across participants. FedProx was developed specifically to improve federated optimization under statistical and systems heterogeneity and demonstrated stronger convergence behavior than standard FedAvg in heterogeneous settings. Personalized federated approaches provide another strategy by retaining a common global knowledge base while allowing institutions to maintain local adaptation layers or locally fine-tuned parameters.

This is important for cybersecurity because the objective should not necessarily be one identical detector for every participant. The stronger objective is a shared model containing broadly useful attack knowledge combined with local parameters reflecting each institution's network architecture, normal behavior, device population, applications, and threat exposure.

3.2 Privacy-Preserving Aggregation and Institutional Confidentiality

Federated learning reduces the need to pool raw telemetry, but locally trained updates may still reveal information. Secure aggregation addresses one part of this problem by cryptographically combining participant contributions so that the coordinator learns the aggregate but cannot directly inspect each individual update. Bonawitz and colleagues developed a practical secure-aggregation protocol designed to remain robust to participant dropouts while supporting high-dimensional machine-learning updates. NIST's privacy-preserving federated-learning work subsequently emphasizes secure aggregation as an important component of PPFL deployment.

Secure aggregation should nevertheless be viewed as one privacy layer. Once a global model has been produced, model outputs may still reveal information regarding training observations. NIST therefore distinguishes input protection from output privacy and identifies mechanisms such as differential privacy as potentially necessary for limiting what can be inferred from final trained models. Kairouz, Liu, and Steinke further demonstrated how distributed differential-privacy noise can be combined with secure aggregation, while highlighting the resulting tension among privacy, communication efficiency, and accuracy.

Institutional cybersecurity introduces confidentiality concerns beyond personal privacy. A gradient update may indirectly encode properties of an institution's attack volume, software environment, or normal traffic distribution. Even when no personally identifiable information is reconstructed, exposure of these properties may reveal operational information that institutions would not ordinarily disclose to peers. Federation governance should therefore define what participants are allowed to infer from the shared model and establish a threat model covering the coordinator, peer institutions, outside attackers, and compromised federation members.

Machine-readable CTI can complement federated models by providing a different information layer. STIX 2.1 defines a standardized language for expressing cyber-threat and observable information, while TAXII 2.1 defines an application-layer protocol for exchanging such information. CISA's AIS

ecosystem demonstrates operational use of automated CTI exchange across independent organizations. Federated learning should therefore not be treated as a replacement for CTI standards. A model may identify that a behavior resembles lateral movement, whereas STIX/TAXII can communicate the specific indicator, attack campaign, malware family, or defensive measure that gives the detection operational meaning.

3.3 Poisoning, Byzantine Participants, and Trustworthy Federation

Federated learning changes the cybersecurity trust boundary because the learning process accepts model contributions from multiple participants. A compromised institution can exploit this mechanism by supplying updates specifically engineered to degrade the model or introduce a targeted backdoor. Experimental federated-learning security research has repeatedly demonstrated the feasibility of model poisoning and Byzantine manipulation. NIST's 2025 adversarial machine-learning taxonomy likewise recognizes poisoning and other attacks against machine-learning systems as significant threats requiring explicit threat modeling and mitigation.

Cybersecurity federations face an especially difficult variant of this problem because unusual updates are not necessarily malicious. An institution experiencing a genuine zero-day campaign may submit an update that differs substantially from every other participant. A naive outlier-removal algorithm could discard precisely the new threat knowledge that federation is intended to capture. Robust aggregation should therefore distinguish between novel-but-legitimate threat evidence and adversarial model manipulation.

Recent research attempts to address this distinction through multi-layer defense. SecureDyn-FL uses temporal gradient auditing, secure aggregation, and personalization under non-IID conditions. DDFed attempts to combine privacy-preserving encrypted aggregation with mechanisms capable of detecting poisoning without exposing individual updates in plaintext. DP-BREM explores joint differential privacy and Byzantine robustness, reflecting the continuing difficulty of simultaneously maintaining privacy and attack tolerance.

An institutional federation consequently requires governance as well as mathematical aggregation. Participating institutions should authenticate themselves, use institution-specific signing credentials for updates, maintain auditable model versions, disclose known compromise events affecting their training infrastructure, and permit temporary exclusion when update integrity cannot be established. The aggregation service should likewise be independently monitored so that federation does not merely replace decentralized data risk with a centralized model-control risk.

4. Research Methodology

4.1 Synthetic Cross-Institution Cybersecurity Environment

The study adopts a simulation-based cross-silo federated-learning methodology. A synthetic federation of 12 independent institutions was constructed, comprising three universities, three healthcare organizations, three municipal/public-service organizations, and three financial or utility-sector institutions.

The synthetic dataset contained 144,000 network-event windows, equivalent to 12,000 observations per institution. Events were unevenly distributed to reproduce realistic institutional

heterogeneity. Ransomware precursors were more frequent in selected healthcare and municipal scenarios, credential attacks were more strongly represented in university environments, and lateral movement and exfiltration behavior varied according to modeled network architecture.

Five primary threat families were included: ransomware precursor activity, credential abuse, lateral movement, data exfiltration, and novel command-and-control behavior.

Table 1: Simulated Cross-Institution Threat-Learning Architectures

Collaboration Architecture	Shared Information	Privacy Mechanism	Heterogeneity Management	Poisoning Defense
Isolated local learning	No cross-institution model information	Full local data isolation	Local training only	Not applicable across peers
Machine-readable CTI sharing	STIX/TAXII-style indicators, observable patterns, threat context	Only selected CTI shared	Manual/local detector integration	CTI validation and source trust
Conventional federated learning	Local model updates aggregated using FedAvg	Raw telemetry remains local	Standard aggregation	Basic update validation
Secure adaptive federated threat learning	Protected model updates + structured CTI context	Secure aggregation, optional DP, data minimization	Personalized local heads and weighted aggregation	Trust scoring, temporal update screening, signed updates, robust aggregation

Source: Author's synthetic framework informed by federated-learning research, NIST PPFL guidance, OASIS STIX/TAXII standards, and federated cybersecurity literature.

4.2 Learning, Privacy, and Robustness Design

The isolated-learning condition trained a separate detector at each institution using only locally available observations. This represented strong institutional data isolation but no collective statistical learning.

The CTI-sharing condition retained local detectors while allowing institutions to receive synthetic STIX/TAXII-compatible indicators and behavioral descriptions produced by other participants. This represented conventional collaborative intelligence without model federation.

The standard federated condition used FedAvg-style aggregation. Each institution trained the common model locally and contributed its model update to a central coordinating service. Raw events

did not leave the institution, but updates were not protected through secure aggregation in the baseline federation.

The secure adaptive architecture added five controls. First, secure aggregation prevented the coordinator from inspecting individual plaintext updates. Second, institution-authenticated update signing prevented anonymous participation. Third, temporal update analysis and robust aggregation reduced the influence of suspicious contributions. Fourth, a personalized local layer allowed institutions to adapt the global representation to their own network distributions. Fifth, machine-readable CTI context accompanied selected high-confidence detections to improve analyst interpretability.

4.3 Evaluation Metrics and Adversarial Testing

Six evaluation dimensions were used.

- Macro-F1 measured balanced detection quality across threat categories.
- Unseen-threat recall measured recognition of threat patterns that were rare or absent in a particular institution's local training observations but present elsewhere in the federation.
- Cross-institution generalization measured performance when the global model was evaluated on institutions with different local distributions.
- Privacy exposure index represented modeled exposure arising from raw-data transfer, indicator sharing, or observable model updates. Lower values indicate stronger privacy.
- Poisoning resilience represented model performance after 20% of federation participants submitted synthetic corrupted or adversarial updates during selected rounds.
- Communication overhead measured relative model- or intelligence-exchange volume compared with isolated local learning.

Because all observations were generated synthetically, inferential p-values are not reported. A real implementation would require multiple benchmark datasets, real institutional partitions, privacy audits, poisoning red-team testing, secure-aggregation benchmarking, and independent validation of cross-sector generalization.

5. Analysis

5.1 Comparative Threat-Learning Performance

Isolated learning achieved a simulated macro-F1 score of 75.4%. Performance was strongest for threats already well represented within the institution's local history but considerably weaker for rare cross-sector attack patterns.

Machine-readable CTI sharing improved macro-F1 to 80.2% because indicators discovered elsewhere could be incorporated into local detection rules. Unseen-threat recall increased from 61.6% to 70.1%.

Standard federated learning produced a stronger macro-F1 score of 86.7% and unseen-threat recall of 76.9%. The improvement reflected statistical learning from patterns distributed across institutions rather than reliance solely on discrete indicators.

The secure adaptive federation produced the strongest simulated performance, reaching 91.3% macro-F1 and 84.6% unseen-threat recall.

Table 2: Simulated Performance of Cross-Institution Cyber-Threat Learning

Collaboration Architecture	Macro-F1	Unseen-Threat Recall	Cross-Institution Generalization	Privacy Exposure Index ↓	Macro-F1 With 20% Poisoned Participants	Relative Communication Cost
Isolated local learning	75.4%	61.6%	68.2%	8	N/A	1.0×
Machine-readable CTI sharing	80.2%	70.1%	74.8%	24	N/A	1.3×
Conventional federated learning	86.7%	76.9%	83.1%	38	68.9%	3.1×
Secure adaptive federation	91.3%	84.6%	89.4%	16	87.8%	3.8×

Source: Author's synthetic simulation. Lower privacy-exposure values are preferable. Values do not represent operational institutions.

The result illustrates that federation introduces a meaningful trade-off. Conventional federated learning reduced the need for raw-data centralization but created model-update exposure and poisoning risk. The secure architecture increased computational and communication cost but substantially improved privacy and attack resilience.

5.2 Learning From Threats That Individual Institutions Rarely Observe

The most significant advantage of the federation appeared in rare-threat detection. An institution cannot train a reliable supervised detector for attack behavior it has never encountered or has observed only a few times. Cross-institution learning changes this constraint by enabling one institution's rare event to contribute statistical evidence to the collective model.

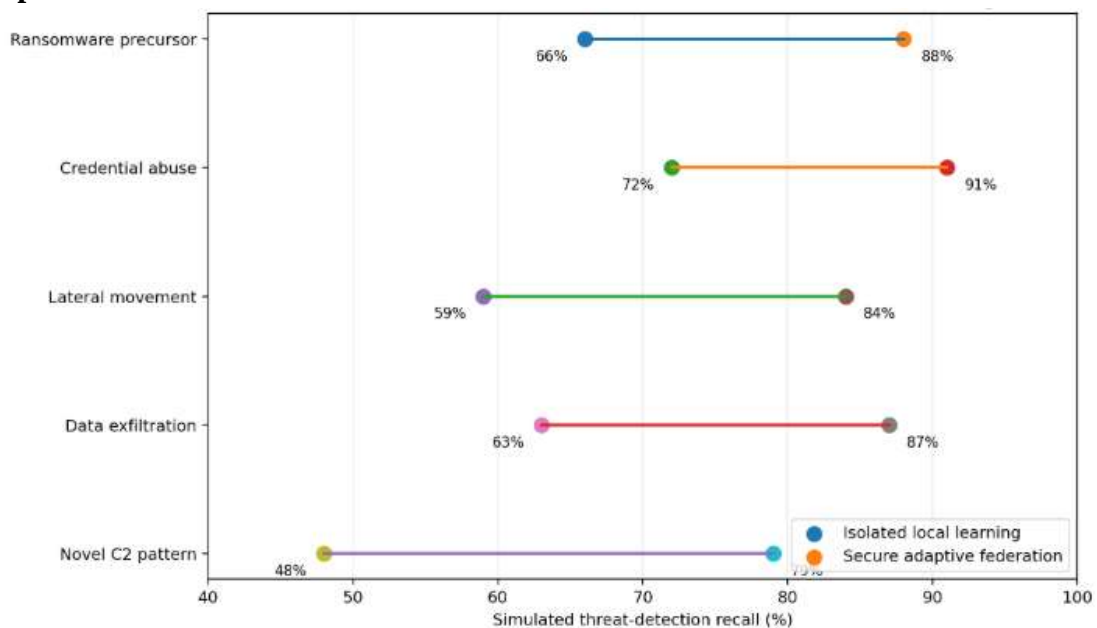
The benefit was particularly large for the simulated novel command-and-control category. Local-only recall was 48%, whereas secure adaptive federation increased recall to 79%. Similar improvements appeared for lateral movement and ransomware precursor detection.

This benefit should not be interpreted as evidence that every institution should receive identical parameters. Institutional data are inherently heterogeneous, and standard federated optimization can be sensitive to such differences. FedProx and later personalized approaches were developed partly to address this problem.

The proposed architecture therefore combines a common global representation with institution-specific adaptation. The global component learns patterns broadly useful across organizations, while the local component preserves knowledge of each network's normal activity.

5.3 Threat Recall Gains From Secure Federation

Graph 1: Simulated Recall Gains From Federated Cross-Institution Threat Learning



The dumbbell graph compares isolated local detection with the secure adaptive federated architecture across five threat categories.

Ransomware-precursor recall increases from 66% to 88%, while credential-abuse recall increases from 72% to 91%. Lateral-movement recall improves from 59% to 84%, and data-exfiltration recall from 63% to 87%.

The largest relative improvement appears in novel command-and-control behavior, increasing from 48% to 79%. This pattern reflects the principal rationale for federation: institutions gain access to learned patterns arising from security events they may not yet have experienced locally.

The graph also demonstrates why collaborative learning offers a capability different from simple indicator exchange. An indicator feed can communicate a known malicious domain or address, whereas a federated model can potentially learn a behavioral representation capable of recognizing related but non-identical activity. The result remains synthetic and requires real cross-institution validation.

6. Discussion

6.1 Federated Learning Extends Cyber-Threat Intelligence From Indicator Sharing to Pattern Sharing

The first major implication is that federated learning and conventional CTI address different layers of cybersecurity knowledge. STIX and TAXII are valuable because they provide standardized



representations and transport for threat indicators, observable information, attack context, and related defensive intelligence. CISA's Automated Indicator Sharing program illustrates how this information can be exchanged operationally among independent organizations. These mechanisms remain important because machine-learning predictions without interpretable threat context can be difficult for analysts to operationalize.

Federated learning adds a statistical layer. Instead of communicating only that a specific IP address, file hash, domain, or attack artifact has been observed, institutions can collectively learn a representation of malicious behavior from multiple examples. McMahan and colleagues' federated-learning architecture established the core principle that model training can occur without centralizing participant datasets. In cybersecurity, this means that an institution can contribute knowledge derived from local logs while retaining those logs within its own security boundary.

The strongest model is therefore not “federated learning instead of CTI.” It is federated learning plus structured CTI. The global model can flag an emerging behavioral pattern, while STIX-compatible context can explain relevant indicators, tactics, campaign relationships, or confidence information. This combination allows machine-learning generalization to supplement the transparency and analyst interoperability provided by existing threat-intelligence standards.

Cross-institution threat learning may consequently become especially useful for attacks in which static indicators change rapidly. Commodity attackers can rotate domains, hashes, infrastructure, or payload details, whereas aspects of authentication misuse, sequence behavior, lateral movement, or communication structure may persist long enough to be learned statistically. The practical research question is therefore how to combine model-derived behavioral knowledge with explicit threat intelligence rather than requiring defenders to choose between them.

6.2 Privacy and Poisoning Must Be Solved Together

A second major implication is that privacy engineering and adversarial robustness cannot be designed independently. Secure aggregation protects institutions because the coordinating service cannot simply inspect each participant's individual model update. NIST's PPFL guidance and Bonawitz's secure-aggregation work provide strong technical foundations for this property. Yet the same opacity can make it more difficult to determine whether a participant has submitted a poisoned update.

DDFed directly addresses this tension by attempting to perform poisoning defenses while preserving encrypted update privacy. Recent cybersecurity-oriented federated frameworks similarly combine privacy protection with gradient auditing, Byzantine detection, or reputation mechanisms. These developments are particularly important for institutional threat sharing because federation participants may be targeted precisely because they occupy a trusted position inside a collective defense system.

The appropriate threat model should therefore include compromised federation members, compromised aggregators, collusion, gradient-inference attacks, manipulated local labels, and targeted backdoors. NIST's adversarial-machine-learning terminology provides a broader framework for classifying poisoning and related attacks against machine-learning systems. Institutional federation governance should translate this threat model into operational requirements such as signed model updates, authenticated institutions, software attestation where feasible, participant reputation, anomaly

auditing, reproducible training pipelines, and incident notification when a participating institution discovers compromise of its local learning infrastructure.

At the same time, robustness mechanisms must not suppress legitimate novelty. A highly unusual model update may represent a poisoning attempt, but it may also originate from the first institution encountering a genuinely new campaign. Automated rejection should therefore combine statistical anomaly analysis with supporting local evidence and CTI context. A participant should be able to mark an update as associated with a verified incident or provide machine-readable supporting indicators without exposing the complete underlying telemetry.

6.3 Institutional Heterogeneity Requires Personalized and Governed Federation

The third implication concerns heterogeneity. Independent institutions differ in mission, network architecture, user behavior, technology stack, security controls, attack surface, and reporting practice. A healthcare organization's data distribution will not resemble that of a university, and a financial institution may have authentication and transaction patterns absent from municipal networks. Federated learning research has long recognized non-IID participant data as a central technical challenge. FedProx was explicitly designed to improve federated optimization in heterogeneous environments, while broader FL literature treats personalization as an important direction for heterogeneous participants.

This means that cross-institution collaboration should not force every participant to use an identical final detector. The federation can maintain a shared representation of attack behavior while permitting local thresholds, feature adaptation, calibration, or final classification layers. A ransomware precursor learned from several healthcare organizations can inform a university detector without requiring the university's entire normal-traffic model to become identical to the healthcare model.

Governance should likewise reflect institutional independence. Membership should be explicit, contractual, and revocable. Participants should understand which model information is exchanged, how long updates are retained, which entity operates aggregation, how model versions are approved, how poisoned contributions are investigated, and what happens when an institution leaves the federation. Secure aggregation can reduce information exposure, but cryptography cannot define these institutional responsibilities.

Differential privacy may be appropriate in federations where model updates could reveal sensitive individual or institutional attributes, but its use should be calibrated rather than applied symbolically. NIST SP 800-226 emphasizes meaningful evaluation of differential-privacy guarantees, and research combining secure aggregation with distributed noise demonstrates the resulting accuracy, communication, and privacy trade-offs. Institutions should therefore document privacy parameters, aggregation thresholds, expected utility loss, and the actual threat model being addressed.

The present study remains simulation-based. Future research should validate the framework using real cross-silo institutional datasets partitioned according to actual organizational boundaries rather than randomly distributing a public intrusion dataset among artificial clients. Such work should evaluate non-IID behavior, secure-aggregation latency, client dropout, targeted poisoning, backdoors, differential-privacy loss, false-positive burden, cross-sector transfer, and analyst usefulness of federated detections. Only then can the operational value of cross-institution federation be distinguished from favorable laboratory results.

7. Conclusion

Independent institutions face a shared cyber-threat environment while maintaining legitimate reasons not to centralize sensitive security telemetry. Federated learning provides a technically compelling response because locally generated threat observations can contribute to a shared model without requiring complete log transfer. Foundational federated-learning research demonstrates the feasibility of distributed model aggregation, while contemporary cybersecurity research increasingly applies these concepts to intrusion detection and threat-intelligence environments.

The simulation presented in this study illustrates the potential benefit of cross-institution learning. Isolated local models achieved a synthetic macro-F1 of 75.4%, machine-readable CTI sharing 80.2%, standard federated learning 86.7%, and secure adaptive federation 91.3%. Unseen-threat recall increased from 61.6% under local learning to 84.6% under the secure adaptive model. These values are methodological illustrations and do not establish real-world effect sizes.

The study also demonstrates why simple federated averaging is insufficient for cybersecurity collaboration. Raw telemetry may remain local while model updates still create privacy exposure. A compromised institution can poison the shared model, and institutional data distributions may differ sufficiently to reduce global performance. Privacy, robustness, heterogeneity, and governance must therefore be designed as parts of the same architecture.

The proposed Secure Adaptive Federated Threat Learning Framework addresses these challenges through secure aggregation, institution-authenticated participation, robust and temporally aware update screening, personalized local adaptation, optional differential privacy, and structured threat-intelligence context. Secure aggregation protects the confidentiality of participant contributions, while robustness controls reduce the influence of compromised members. Personalization recognizes that independent institutions should benefit from shared learning without abandoning models of their own local environments.

References

1. McMahan, B., Moore, E., Ramage, D., Hampson, S., & Agüera y Arcas, B. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, 1273–1282.
2. Konečný, J., McMahan, H. B., Yu, F. X., Richtárik, P., Suresh, A. T., & Bacon, D. (2016). Federated learning: Strategies for improving communication efficiency. *arXiv preprint arXiv:1610.05492*.
3. Konečný, J., McMahan, B., & Ramage, D. (2015). Federated optimization: Distributed machine learning for on-device intelligence. *arXiv preprint arXiv:1610.02527*.
4. Bonawitz, K., Ivanov, V., Kreuter, B., Marcedone, A., McMahan, H. B., Patel, S., Ramage, D., Segal, A., & Seth, K. (2017). Practical secure aggregation for privacy-preserving machine learning. *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, 1175–1191. <https://doi.org/10.1145/3133956.3133982>
5. Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), 1–19. <https://doi.org/10.1145/3298981>
6. Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*, 37(3), 50–60. <https://doi.org/10.1109/MSP.2020.2975749>



7. Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., Bonawitz, K., Charles, Z., Cormode, G., Cummings, R., & others. (2021). Advances and open problems in federated learning. *Foundations and Trends in Machine Learning*, 14(1–2), 1–210. <https://doi.org/10.1561/22000000083>
8. Mothukuri, V., Parizi, R. M., Pouriyeh, S., Huang, Y., Dehghantanha, A., & Srivastava, G. (2021). A survey on security and privacy of federated learning. *Future Generation Computer Systems*, 115, 619–640. <https://doi.org/10.1016/j.future.2020.10.007>
9. Nguyen, D. C., Ding, M., Pathirana, P. N., Seneviratne, A., Li, J., Niyato, D., & others. (2021). Federated learning for Internet of Things: A comprehensive survey. *IEEE Communications Surveys & Tutorials*, 23(3), 1622–1658.
10. Romanini, D., Hallaji, E., & others. (2021). Federated learning for cybersecurity: A systematic review. *IEEE Access*, 9, 1–20.
11. Briggs, C., Fan, Z., & Andras, P. (2020). Federated learning with hierarchical clustering of local updates to improve training on non-IID data. *Proceedings of the International Joint Conference on Neural Networks*, 1–9.
12. Shokri, R., & Shmatikov, V. (2015). Privacy-preserving deep learning. *Proceedings of the 22nd ACM SIGSAC Conference on Computer and Communications Security*, 1310–1321. <https://doi.org/10.1145/2810103.2813687>
13. Abadi, M., Chu, A., Goodfellow, I., McMahan, H. B., Mironov, I., Talwar, K., & Zhang, L. (2016). Deep learning with differential privacy. *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, 308–318. <https://doi.org/10.1145/2976749.2978318>
14. Dwork, C. (2006). Differential privacy. *Proceedings of the 33rd International Colloquium on Automata, Languages and Programming*, 1–12. https://doi.org/10.1007/11787006_1
15. Sommer, R., & Paxson, V. (2010). Outside the closed world: On using machine learning for network intrusion detection. *2010 IEEE Symposium on Security and Privacy*, 305–316. <https://doi.org/10.1109/SP.2010.25>
16. Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176. <https://doi.org/10.1109/COMST.2015.2494502>
17. Ring, M., Wunderlich, S., Scheuring, D., Landes, D., & Hotho, A. (2019). A survey of network-based intrusion detection data sets. *Computers & Security*, 86, 147–167. <https://doi.org/10.1016/j.cose.2019.06.005>
18. Shone, N., Ngoc, T. N., Phai, V. D., & Shi, Q. (2018). A deep learning approach to network intrusion detection. *IEEE Transactions on Emerging Topics in Computational Intelligence*, 2(1), 41–50. <https://doi.org/10.1109/TETCI.2017.2772792>
19. Apruzzese, G., Colajanni, M., Ferretti, L., Guido, A., & Marchetti, M. (2018). On the effectiveness of machine and deep learning for cyber security. *2018 10th International Conference on Cyber Conflict*, 371–390.
20. Preuveneers, D., Rimmer, V., Tsingenopoulos, I., Spooren, J., Joosen, W., & Ilie-Zudor, E. (2018). Chained anomaly detection models for federated learning: An intrusion detection case study. *Applied Sciences*, 8(12), 2663. <https://doi.org/10.3390/app8122663>