

Explainable Federated Learning for Privacy-Preserving Cybersecurity in Industrial IoT and Smart Factories

Bekzod Karimov

Department of Information Security and Artificial
Intelligence, Central Asian Institute of Technology Tashkent, Uzbekistan

ARTICLE INFO

Article history:

Submission: May 01, 2026

Accepted: July 31, 2026

Published: August 27, 2026

VOLUME: Vol.11 Issue 08 2026

Keywords:

Explainable Artificial Intelligence;
Federated Learning; Industrial
Internet of Things; Smart Factory;
Privacy-Preserving Cybersecurity;
Intrusion Detection; Feature
Selection; Dimensionality Reduction;
Anomaly Detection; Cyber-Physical
Systems

ABSTRACT

The rapid integration of Industrial Internet of Things (IIoT), cyber-physical systems, intelligent sensors, and connected manufacturing platforms has created highly distributed industrial environments in which cybersecurity must operate under stringent requirements for privacy, low latency, interpretability, and operational continuity. Conventional centralized machine-learning-based intrusion detection requires the aggregation of potentially sensitive industrial data, creating additional privacy and data-governance risks. Federated learning (FL) provides an alternative paradigm in which participating industrial nodes collaboratively train a shared model without directly exchanging their raw observations. However, privacy preservation alone is insufficient for industrial cybersecurity because security operators must understand why a model has classified an industrial event as malicious or anomalous. This paper develops a research-oriented framework for explainable federated learning (XFL) for privacy-preserving cybersecurity in IIoT and smart-factory environments. The proposed methodology integrates distributed model training, feature-selection and dimensionality-reduction principles, explainable decision analysis, and evaluation using classification-oriented performance measures. The framework is conceptually grounded in established studies of feature selection, dimensionality reduction, machine learning, and evaluation metrics. Mutual-information-based feature selection is positioned as a mechanism for reducing redundant industrial telemetry while preserving discriminative information, whereas dimensionality reduction is used to address high-dimensional and heterogeneous observations. Federated aggregation enables collaborative learning while limiting direct data exchange. An explanation layer subsequently identifies the principal features influencing individual security decisions, improving analyst trust and operational interpretability. The resulting framework is particularly relevant to heterogeneous smart factories where data distributions, device capabilities, and attack patterns may vary considerably across participating sites. The analysis demonstrates that explainability, privacy preservation, dimensionality management, and predictive performance should be treated as interconnected design objectives rather than isolated components. The study also identifies limitations associated with heterogeneous client distributions, computational overhead, explanation fidelity, and the absence of direct empirical validation in a deployed industrial environment.

INTRODUCTION

The Industrial Internet of Things has transformed manufacturing environments from relatively isolated production systems into interconnected computational ecosystems containing sensors, controllers, machines, edge devices, gateways, analytics platforms, and enterprise services. Such connectivity improves monitoring and automation but simultaneously increases the number of potential security observation points. Industrial cybersecurity systems must therefore identify anomalous traffic and potentially malicious behavior without imposing unacceptable latency or exposing sensitive operational information.

Machine learning offers a mechanism for learning complex relationships among network, device, and operational features. Traditional centralized learning, however, generally assumes that relevant data can be transferred to a common training location. In smart factories, this assumption is problematic because industrial telemetry may reveal production states, machine behavior, network configurations, operational patterns, or other commercially sensitive information. A distributed learning architecture is consequently attractive because it allows multiple industrial participants to contribute to a shared model without requiring routine transfer of raw training datasets.

Federated learning addresses this requirement by moving model computation toward participating clients. Each client trains or updates a local model using locally available observations, after which model parameters or updates can be combined to produce a global model. For industrial cybersecurity, this architecture can theoretically enable multiple factories, production cells, or edge domains to learn collectively from heterogeneous security observations while maintaining greater control over local data.

Nevertheless, federated learning does not automatically solve the interpretability problem. A cybersecurity model may accurately identify a suspicious industrial communication pattern while providing insufficient information about the reason for its decision. In operational environments, such opacity can delay incident investigation and make it difficult for security analysts to distinguish legitimate process changes from genuine attacks. An explainable layer is therefore necessary to translate model predictions into interpretable evidence.

Feature engineering is particularly important in this context. Feature-selection research emphasizes that reducing irrelevant or redundant variables can improve learning efficiency and potentially improve model generalization (Chandrashekar and Sahin, 2014). Mutual information provides a theoretical basis for evaluating the relationship between candidate features and target information, while more recent work has investigated theoretical aspects of mutual-information-based feature selection (Beraha et al., 2019). These principles are relevant to IIoT because industrial telemetry can contain numerous correlated and heterogeneous variables.

Dimensionality reduction provides a complementary mechanism for transforming high-dimensional observations into lower-dimensional representations. Research on anomaly detection in multivariate time series demonstrates the importance of understanding how dimensionality-reduction techniques influence detection performance (Altin and Cakir, 2024). Similarly, dimensionality-reduction evaluation in machine-learning applications illustrates the broader importance of selecting representations appropriate to the underlying prediction problem (Pham et al., 2022).

The theoretical foundation for machine-learning classification also derives from established statistical learning principles (Bishop, 2006). Practical machine-learning workflows emphasize systematic data preparation, model construction, validation, and evaluation (Geron, 2019). These principles are incorporated into the proposed XFL architecture rather than treating federated learning as an isolated algorithmic component.

The primary objective of this study is to formulate a unified framework for privacy-preserving and explainable cybersecurity learning in heterogeneous IIoT environments. The paper specifically examines how feature selection, dimensionality management, federated collaboration, and explanation mechanisms can be integrated into a coherent security architecture. A secondary objective is to establish a theoretically grounded evaluation strategy capable of examining both predictive effectiveness and interpretability.

The significance of this research is reinforced by the broader security-risk perspective represented by SecureRiskNet, which emphasizes intelligent security-risk detection across heterogeneous computing environments (Govindarajan et al., 2026). Although that work addresses cloud-fog computing rather than the exact IIoT setting considered here, its heterogeneous security perspective provides a useful conceptual motivation for treating distributed industrial environments as collaborative but operationally diverse security domains.

LITERATURE REVIEW

2.1 Feature Selection and Information Preservation

Feature selection is a fundamental preprocessing problem in machine learning because the presence of unnecessary variables can increase computational complexity and introduce irrelevant information into predictive models. Chandrashekar and Sahin (2014) provide a broad examination of feature-selection approaches and demonstrate that selecting informative variables is an important component of machine-learning system design. For IIoT cybersecurity, this is particularly relevant because industrial networks can generate large numbers of measurements from sensors, communication channels, device states, and traffic statistics.

Mutual information is especially relevant when the relationship between a feature and a target variable is not adequately represented by simple linear dependence. Vergara and Estevez (2014) examined feature-selection methods based on mutual information, establishing its relevance for identifying informative variables. Beraha et al. (2019) further investigated theoretical insights into mutual-information-based feature selection. Together, these works support the use of information-oriented feature selection as a principled stage before collaborative cybersecurity learning.

In a federated architecture, feature selection has an additional role. If each client processes a very large feature space, local training becomes more computationally demanding and differences between clients may become harder to manage. A carefully designed feature-selection stage can therefore reduce communication and computation requirements while preserving security-relevant information.

2.2 Dimensionality Reduction for Anomaly Detection

Dimensionality reduction differs from feature selection because it can transform the original feature space rather than simply retaining a subset of original variables. Altin and Cakir (2024) specifically investigated the influence of dimensionality reduction on anomaly detection performance in multivariate time series. This is directly relevant to IIoT because industrial observations frequently possess temporal structure and substantial variable interdependence.

Pham et al. (2022) similarly demonstrate the importance of evaluating dimensionality-reduction techniques according to their impact on machine-learning prediction performance. The central implication for cybersecurity is that dimensionality reduction should not be selected solely because it reduces the number of variables. A transformation is useful only if the resulting representation retains information relevant to distinguishing normal and abnormal behavior.

Within the proposed framework, dimensionality reduction is consequently positioned as an adaptive preprocessing layer rather than an unconditional optimization. Industrial clients may first perform feature analysis and subsequently evaluate whether a reduced representation maintains sufficient discrimination for security classification.

2.3 Machine Learning and Classification Foundations

Bishop (2006) provides the statistical and probabilistic foundations required for understanding machine-learning models, including classification and pattern recognition. Geron (2019) complements these theoretical principles with practical machine-learning workflows involving data preparation, model development, training, and evaluation. These references collectively establish the methodological basis for constructing the local models used by participating industrial clients.

Pal (2021) illustrates the practical use of logistic regression as a relatively interpretable classification technique. While logistic regression alone may not capture every complex industrial attack pattern, its interpretability makes it useful as a baseline for evaluating more sophisticated models. A cybersecurity framework should therefore compare not only predictive performance but also the transparency of decision-making.

2.4 Evaluation and Research Gap

Evaluation metrics are critical because cybersecurity datasets can contain class imbalance, meaning that high overall accuracy may conceal poor detection of minority attack classes. Hicks et al. (2022) emphasize the importance of carefully selecting evaluation metrics for artificial-intelligence applications. The scikit-learn metric documentation likewise provides operational definitions for precision, recall, and F-score-oriented evaluation (Metric and scoring: quantifying the quality of predictions).

The CIC-DDoS 2019 Dataset provides a concrete dataset context for examining network-level denial-of-service detection. However, a dataset representing a particular attack category cannot by itself establish generalization to every industrial cybersecurity scenario. Smart factories contain heterogeneous devices and operational conditions, so an effective research design must consider distributional variation between clients.

The principal research gap identified from the provided literature is therefore not simply the absence of another classifier. Instead, the gap lies in integrating four requirements: privacy-preserving collaborative learning, industrial heterogeneity, dimensionality-aware representation, and explainable security decisions. SecureRiskNet further motivates the need for intelligent security-risk detection in heterogeneous computing infrastructures (Govindarajan et al., 2026). The proposed framework extends this conceptual direction toward IIoT and smart-factory cybersecurity.

METHODOLOGY

3.1 Proposed Explainable Federated Architecture

The proposed architecture consists of five functional layers: industrial data acquisition, local preprocessing, federated local learning, global aggregation, and explainable security analysis. Each participating industrial client represents a logical security domain, such as a production line, factory, edge gateway, or manufacturing facility.

The first layer collects locally generated observations. These may include network-flow attributes, packet-level statistics, device behavior indicators, temporal measurements, and operational variables. Raw observations remain within the local domain. This design is central to the privacy objective because the architecture does not require centralized storage of the complete industrial dataset.

The second layer performs preprocessing. Missing-value handling, normalization, categorical encoding, and data-quality assessment are conducted locally. Feature-selection mechanisms can then identify variables carrying meaningful predictive information. Mutual-information-based selection is appropriate when nonlinear relationships may be present (Vergara and Estevez, 2014; Beraha et al., 2019).

A dimensionality-reduction stage may follow feature selection. The purpose is not simply to minimize dimensionality but to construct a representation that maintains information needed for anomaly or intrusion detection. This distinction is important because aggressive reduction can remove subtle signals associated with rare attacks. Findings from dimensionality-reduction research indicate that the impact of reduction should be empirically evaluated rather than assumed to be beneficial (Altin and Cakir, 2024; Pham et al., 2022).

3.2 Local Model Training

Each client trains a cybersecurity model using its own processed data. The model may perform binary classification, multiclass attack classification, or anomaly detection. Logistic regression can provide an interpretable baseline (Pal, 2021), while more complex machine-learning models can be investigated when nonlinear patterns require greater representational capacity.

The local training process follows standard statistical-learning principles in which model parameters are optimized against local observations (Bishop, 2006). Practical implementation can incorporate established machine-learning workflows involving training, validation, feature processing, and model assessment (Geron, 2019).

The central federated mechanism is that local clients communicate model updates rather than their raw observations. An aggregation process combines the participating updates into a shared model. The global model is then redistributed, allowing clients to improve their local detection capability using information learned collectively across participating domains.

3.3 Handling Heterogeneous Industrial Clients

A major methodological challenge is non-identically distributed data. One factory may contain robotic assembly systems, another may operate process-control equipment, and a third may have substantially different network traffic. Their normal operating distributions can therefore differ significantly.

The framework addresses this by retaining local preprocessing and validation while using collaborative model aggregation at the learning layer. Client-specific evaluation is maintained alongside global evaluation so that a strong aggregate score does not conceal poor performance for a particular industrial domain.

Heterogeneity also influences feature availability. A feature present at one site may not exist at another. Consequently, the feature-selection process should distinguish between globally compatible variables and client-specific variables. Mutual-information analysis can help identify locally informative variables, but a shared representation may require a common feature schema.

3.4 Explain ability Layer

Explain ability is integrated after prediction rather than treated as a purely visualization-oriented component. For each security decision, the system should identify which features contributed most strongly to the predicted class or anomaly score. This allows an analyst to inspect whether the decision was associated with meaningful security evidence.

For example, a model may classify a traffic instance as suspicious because of an unusual combination of connection frequency, packet characteristics, and temporal behavior. Instead of presenting only an attack label, the explanation layer can identify the dominant contributing variables. Such information is valuable for distinguishing model-driven alerts from operational anomalies.

Feature-selection literature provides a natural theoretical foundation for this mechanism because selected features represent variables considered informative for prediction (Chandrashekar and Sahin, 2014). However, selected features should not automatically be interpreted as causal explanations. A variable may correlate with an attack without being the underlying cause. Therefore, explanations should be treated as decision-support evidence rather than definitive causal statements.

The importance of interpretability is particularly strong in heterogeneous security environments. As emphasized by Govindarajan et al. (2026), intelligent risk detection across heterogeneous computing infrastructures requires mechanisms capable of supporting security-oriented decision-making. The proposed framework similarly treats explainability as an operational requirement rather than a cosmetic model feature.

3.5 Privacy and Security Considerations

Federated learning reduces the requirement for direct raw-data sharing, but privacy preservation should not be interpreted as absolute immunity from information leakage. Model updates may still contain information about local training patterns. Therefore, a production implementation should evaluate the privacy properties of the communication and aggregation mechanisms independently from the predictive model.

The framework consequently defines privacy as a system-level objective involving local data retention, controlled communication, access management, and secure aggregation. The present study does not claim that federated learning alone eliminates every privacy threat.

3.6 Evaluation Protocol

Evaluation should occur at three levels: predictive performance, operational efficiency, and explanation quality. Predictive evaluation should include precision, recall, F-score, and related metrics rather than relying exclusively on accuracy. This is consistent with the broader importance of carefully selecting AI evaluation metrics (Hicks et al., 2022).

The precision-recall relationship is particularly important in cybersecurity. Excessive false positives can overload analysts, while excessive false negatives can permit attacks to remain undetected. Consequently, the appropriate operating point depends on the security objective and the operational cost of errors.

The CIC-DDoS 2019 Dataset can be used to establish an experimental baseline for distributed attack-detection research. Nevertheless, dataset-based evaluation should be complemented by client-partition experiments that simulate heterogeneous industrial domains. Such partitioning allows the study to examine whether collaborative training remains effective when clients possess different attack distributions.

3.7 Experimental Comparison

A rigorous experiment should compare at least four configurations: centralized learning, independent local learning, federated learning without explainability, and explainable federated learning. This design separates the contribution of collaboration from the contribution of explanation.

A second experimental dimension should compare models before and after feature selection and dimensionality reduction. The objective is to determine whether computational simplification causes unacceptable degradation in detection performance. The literature indicates that dimensionality reduction can affect anomaly-detection outcomes, making this comparison methodologically necessary (Altin and Cakir, 2024).

Finally, the evaluation should compare global performance with client-level performance. A global improvement is meaningful only when it does not result in substantial degradation for specific industrial participants.

RESULTS

The synthesis of the provided literature supports several findings relevant to the proposed framework. First, privacy-preserving collaborative learning is most valuable when industrial data are distributed across organizational or operational boundaries. Keeping raw observations local can reduce the need for centralized data pooling, while federated model collaboration allows information learned by one client to influence the shared model. However, the privacy benefit should be interpreted as a reduction in raw-data exposure rather than a complete privacy guarantee.

Second, feature selection and dimensionality management emerge as essential supporting mechanisms rather than optional preprocessing steps. Feature-selection research indicates that irrelevant and redundant variables can complicate learning, while mutual-information approaches provide a principled basis for identifying informative relationships (Chandrashekar and Sahin, 2014; Vergara and Estevez, 2014). Theoretical work on mutual-information selection further strengthens the argument for information-oriented feature reduction (Beraha et al., 2019). In a federated environment, reducing unnecessary dimensions can also lower local computational requirements and potentially reduce the amount of model information that must be communicated.

Third, dimensionality reduction must be evaluated against predictive preservation. The evidence considered in this paper indicates that lower-dimensional representations can influence anomaly-detection performance, particularly in multivariate data (Altin and Cakir, 2024). Therefore, the proposed framework treats dimensionality reduction as a controlled experimental variable rather than assuming that fewer dimensions necessarily produce a better cybersecurity system.

Fourth, evaluation should prioritize multiple complementary measures. Precision, recall, and F-score provide more informative evidence than accuracy alone when attack and normal classes are unevenly

represented. This follows the broader evaluation principles emphasized by Hicks et al. (2022) and the operational metric definitions provided by the scikit-learn reference. A cybersecurity model that produces high accuracy but weak recall for attacks cannot be considered operationally reliable.

Fifth, explainability provides a distinct contribution beyond predictive accuracy. A model that identifies a malicious event but cannot communicate influential evidence may have limited usefulness for analysts investigating an industrial incident. The proposed XFL design therefore treats explanation as an additional evaluation dimension. The explanations should reveal influential features while avoiding the unsupported interpretation that statistical feature importance represents causality.

Sixth, heterogeneous clients represent a central challenge. A model trained collaboratively across diverse factories may learn broader attack patterns than isolated local models, but distribution differences can also cause uneven performance. The heterogeneous-security perspective identified in Govindarajan et al. (2026) reinforces the importance of evaluating security intelligence across different computational environments. Accordingly, the proposed framework recommends client-level evaluation rather than relying exclusively on one global metric.

Finally, the literature supports the feasibility of constructing a unified architecture but does not provide sufficient evidence to claim a specific numerical performance improvement for the proposed XFL system. The current findings are therefore methodological and analytical rather than experimentally validated. Actual improvements in detection, communication efficiency, privacy, or interpretability require implementation and controlled empirical testing.

DISCUSSION

The central theoretical implication of this study is that privacy, explainability, and predictive performance should not be optimized independently. Federated learning addresses the data-distribution problem, feature selection addresses representation complexity, dimensionality reduction addresses computational structure, and explainability addresses human interpretability. Their integration creates a more comprehensive cybersecurity architecture than any single component can provide.

The first major trade-off concerns privacy and model performance. Local training preserves data locality, but clients may possess only limited examples of particular attack types. Centralized learning can theoretically exploit a larger pooled dataset, whereas federated learning must obtain collaborative benefits indirectly through model updates. The effectiveness of collaboration therefore depends strongly on the diversity and quality of participating clients.

The second trade-off concerns dimensionality and information preservation. Feature reduction can make local learning faster and models easier to analyze, but removing variables may eliminate subtle attack indicators. The findings of Altin and Cakir (2024) are therefore particularly relevant: dimensionality reduction should be treated as an empirical design choice rather than a universally beneficial operation. Pham et al. (2022) similarly illustrate why reduced representations must be evaluated according to their effect on downstream prediction.

A third issue concerns explainability itself. Explanations may improve analyst confidence, but an explanation that is unstable, misleading, or disconnected from the underlying prediction process can create false confidence. The framework therefore recommends evaluating explanation consistency and relevance alongside predictive metrics. Feature-selection results can help organize explanations, but selected features should not automatically be interpreted as causal factors.

A fourth consideration is the tension between model complexity and interpretability. Logistic regression offers a relatively transparent baseline (Pal, 2021), whereas more sophisticated nonlinear models may capture complex attack behavior more effectively. Bishop (2006) provides the broader theoretical basis for understanding this relationship between model representation and predictive learning. The appropriate solution is not necessarily to select the simplest model, but to establish whether increased complexity produces sufficient security benefit to justify the additional interpretability burden.

The framework also has practical implications for smart-factory operators. A distributed industrial organization could maintain local telemetry while participating in collaborative cybersecurity training. A local security team could receive a global model and subsequently inspect alerts through an explanation layer. This configuration could support cross-site learning without requiring continuous centralization of raw operational data.

The research also aligns conceptually with heterogeneous security-risk detection frameworks such as SecureRiskNet, where distributed and diverse computing environments create challenges for intelligent security analysis (Govindarajan et al., 2026). The present framework differs in its emphasis on industrial IoT, federated collaboration, and explainable intrusion detection, but the underlying principle of combining intelligent analysis with heterogeneous-environment awareness is shared.

Several limitations remain. The provided literature does not directly establish the performance of the proposed XFL architecture in a real smart factory. The CIC-DDoS 2019 Dataset can support network-attack experimentation, but it cannot fully represent the diversity of industrial cyber-physical processes. Additionally, federated learning introduces communication and synchronization overhead. Client heterogeneity can complicate convergence, while privacy mechanisms beyond basic data locality may further influence performance.

Another limitation is that explainability is difficult to reduce to a single numerical measure. An explanation can be technically faithful to a model yet operationally difficult for an analyst to interpret. Future empirical studies should therefore evaluate explanation usefulness through both quantitative stability measures and structured analyst assessments.

Overall, the proposed architecture should be regarded as a research framework rather than a claim of experimentally proven superiority. Its principal contribution is the integration of complementary principles identified across the supplied literature into a coherent privacy-preserving and explainable cybersecurity methodology.

CONCLUSION

This study presented an explainable federated-learning framework for privacy-preserving cybersecurity in Industrial IoT and smart-factory environments. The analysis established that distributed industrial cybersecurity requires more than a high-performing classifier. Raw-data locality, feature selection, dimensionality management, collaborative learning, explainability, and rigorous evaluation must be designed as interconnected components.

The literature indicates that feature selection can reduce irrelevant information, with mutual information providing a theoretically grounded approach for identifying informative relationships (Vergara and Estevez, 2014; Beraha et al., 2019). Dimensionality reduction can further simplify complex industrial observations, but its influence on anomaly detection must be empirically assessed rather than presumed beneficial (Altin and Cakir, 2024). Established machine-learning principles provide the foundation for local model construction, while evaluation research highlights the importance of using metrics capable of revealing performance under imbalanced classification conditions (Bishop, 2006; Hicks et al., 2022).

The proposed XFL architecture contributes a unified perspective in which federated collaboration provides the privacy-oriented learning mechanism and explainability provides the analyst-oriented decision layer. The framework is particularly suited to heterogeneous environments because it allows local processing and evaluation while enabling collaborative model improvement. The broader heterogeneous-security perspective represented by Govindarajan et al. (2026) further supports the importance of intelligent security analysis across distributed computational environments.

Future work should implement the framework using real or realistically partitioned industrial datasets, evaluate multiple federated aggregation strategies, investigate client heterogeneity, quantify communication overhead, and compare alternative explainability techniques. Additional research should also examine stronger privacy guarantees and the robustness of explanations under adversarial conditions.

Ultimately, the success of explainable federated cybersecurity should be measured not only by classification performance but by its ability to provide reliable, privacy-conscious, and actionable security intelligence in operational smart factories.

REFERENCES

1. M. Altin and A. Cakir, "Exploring the influence of dimensionality reduction on anomaly detection performance in multivariate time series," Mar. 2024, arXiv:2403.04429.
2. M. Beraha, A. M. Metelli, M. Papini, A. Trinzoni, and M. Rostelli, "Feature selection via mutual information: New theoretical insight," in Int. Joint Conf. Neural Netw. (IJCNN), Budapest, Hungary, Jul.
3. C. M. Bishop, Pattern Recognition and Machine Learning. New York, NY, USA: Springer, 2006.
4. G. Chandrashekar and F. Sahin, "A survey on feature selection methods," Comput. Electr. Eng., vol. 40, no. 1, pp. 16–28, Jan. 2014. 10.1016/j.compeleceng.2013.11.024.
5. CIC-DDoS 2019 Dataset. Accessed: Feb. 5, 2024. [Online]. Available: <https://www.unb.ca/cic/datasets/ddos-2019.html>
6. A. Geron, Hands-on Machine Learning with Scikit-Learn. Oreilly, Boston: Keras & TensorFlow, 2019.
7. S. A. Hicks et al., "On evaluation metrics for medical applications of artificial intelligence," Nature Sci. Rep., vol. 12, no. 5979, pp. 1–9, Apr. 2022. 10.1038/s41598-022-09954-8.
8. Metric and scoring: quantifying the quality of predictions. Accessed: Mar. 10, 2024. [Online]. Available: https://scikitlearn.org/stable/modules/generated/sklearn.metrics.precision_recall_fscore_support.html
9. A. Pal, "Logistic Regression: A Simple Primer," Cancer Res. Stat. Treat. vol. 4, no. 3, pp. 551–554, Jul.
10. H. T. Pham, J. Awange, and M. Kuhn, "Evaluation of three feature dimension reduction techniques for machine-learning based crop yield prediction models," Sensors, vol. 22, no. 17, pp. 1–18, Sep. 2022. 10.3390/s22176609.
11. J. R. Vergara and P. Estevez, "A review of feature selection methods based on mutual information," Neural Comput. Appl., vol. 24, no. 1, pp. 10–23, Jan. 2014. 10.1007/s00521-013-1368-0.
12. K. Ramamurthy, R. K. Konduru and N. Amanmadov, "EvoGraphCoder: An Evolutionary Graph-Reasoning Framework for Self-Adaptive Software Engineering," in IEEE Access, vol. 14, pp. 63063–63076, 2026, doi: 10.1109/ACCESS.2026.3686019.
13. Philip, P. G. (2024). Artificial Intelligence-Driven Project Risk Prediction Models: Enhancing Decision-Making Accuracy in Large-Scale Infrastructure Projects . American Journal of Technology, 3(1), 52–69. <https://doi.org/10.58425/ajt.v3i1.571>
14. Govindarajan, V., Ahmed, F., Kamaluddin, K. et al. SecureRiskNet: An Advanced AI-Driven Framework for Intelligent Security Risk Detection in Heterogeneous Cloud-Fog Computing Networks. Int J Comput Intell Syst 19, 74 (2026).