

Exploring Intelligent Investigative Techniques for a Unified Cyber Defense, Focusing on Semantic Client Status Modeling

R. Veerapandiyan^{#1} and Dr S Thaiyalnayaki^{*2}

^{#Ph.d Research Scholar, Department of Computer Science and Engineering, Bharath Institute of Higher Education, Chennai-600073, Tamil Nadu, India.}

^{*Associate Professor, Department of Computer Science and Engineering, Bharath Institute of Higher Education, Chennai-600073, Tamil Nadu, India.}

Abstract— The rapid rise in cyber threats has made it clear that traditional intrusion detection systems (IDSs) just don't cut it in today's distributed environments. While Federated Learning (FL) offers a way to collaborate and maintain privacy in intrusion detection, most current methods focus on model updates and often overlook clients that might be trying to game the system with tactics like adaptive poisoning or gradient mimicry. To tackle this issue, this study presents an Intelligent Investigative Framework for Unified Cyber Defense that leverages Semantic Client Status Modeling (SCSM). This innovative framework integrates time-based behavior checks, meaningful feature extraction, explainable AI (XAI), trust assessments, and the detection of unusual activities to build up-to-date profiles of clients and identify those with malicious intent. The federated deep learning models are evaluated using standard datasets like UNSW-NB15, CICIDS2017, and N-BaIoT, which feature diverse data. The goal is to enhance the detection of harmful clients, improve accuracy, boost recall, and strengthen defenses against attacks, all while reducing false alarms compared to traditional update-based checks. Overall, the SCSM framework aims to make next-generation federated intrusion detection systems clearer, more trustworthy, and robust, particularly in cloud computing, IoT, edge computing, and critical infrastructure.

Index Terms—Unified Cyber Defense, SCSM, IDS, Adaptive Deceptive Client Detection, Intelligent Cyber Investigation, Edge and IoT Security.

I. INTRODUCTION

The rapid evolution of digital technology—think cloud computing, the Internet of Things (IoT), edge computing, and hybrid systems that blend digital and physical elements—has led to a surge in cyber threats. Today's businesses depend on interconnected and distributed systems that handle sensitive information, making them prime targets for sophisticated cyberattacks like ransomware, DDoS attacks, advanced persistent threats (APTs), data poisoning, and insider threats. While traditional Intrusion Detection Systems (IDS) still play a crucial role in cybersecurity, they come with their own set of limitations. They often struggle

with scalability, can raise privacy concerns, and may not keep pace with emerging attack techniques [1–5]. To tackle these challenges, Federated Learning (FL) has emerged as an exciting approach to distributed machine learning. FL enables multiple organizations or devices to collaborate on training models for intrusion detection without needing to share their actual data. Instead, they exchange model information such as parameters or gradients, ensuring data privacy while still benefiting from collective insights. As a result, federated intrusion detection systems are increasingly being adopted in cloud environments, IoT networks, edge computing, healthcare systems, and industrial control setups. However, despite its advantages for data privacy, FL can also introduce new security vulnerabilities [1–5]. Malicious actors might manipulate their local model updates before sending them to the central server, aiming to deceive the system. A significant hurdle in federated intrusion detection is the risk posed by harmful clients who engage in attacks like model poisoning, backdoors, Byzantine behavior, or mimicry. These attackers attempt to disguise their actions by making their updates appear similar to those of legitimate clients, which can easily slip past conventional detection methods. Defending against these attacks typically involves examining how model updates shift, the distances between parameters, or employing robust aggregation methods [6–10]. While these strategies can be effective against basic attacks, they often fall short against clever attackers who adapt their tactics over time and take advantage of vulnerabilities in update verification. Recent developments in explainable artificial intelligence (XAI) highlight the necessity of having models that are understandable to people, especially in the realm of cybersecurity [11–15]. Techniques such as SHAP, LIME, Integrated Gradients, and Layer-wise Relevance Propagation shed light on the reasons behind model predictions, enabling security experts to grasp the rationale behind the model's decisions [19–22]. However, most XAI-driven cybersecurity solutions tend to focus on explaining intrusion classifications rather than assessing the behavior and trustworthiness of individual clients in a federated environment [12–15]. This indicates a significant gap in research regarding the integration of XAI with behavioral insights to enhance

adaptive cyber defense. Another emerging field is semantic client modeling, which examines client behavior not just through the numerical data in model updates but also by considering how behavior evolves over time, the context, trust levels, meanings, and previous interactions. Semantic analysis empowers systems to distinguish between normal behavioral changes due to non-IID data and intentional attacks from malicious actors. Research indicates that merging time-based semantic evidence with explainable behavior analysis significantly boosts the ability to identify harmful clients and fortifies the system against sophisticated attacks. Trust management has become a crucial element in our collaborative efforts to enhance cybersecurity. Instead of assuming that all clients are equally trustworthy, trust-aware federated learning takes a closer look at each client's history, their reliability, consistency in behavior, and their contributions over time [12–15]. This dynamic approach to trust helps mitigate the risks posed by unreliable clients while ensuring that honest clients, who might experience normal fluctuations in their data, aren't unfairly penalized. By integrating trust management with semantic analysis and explainable techniques, we can evaluate clients more effectively than simply relying on model updates. Moreover, AI is revolutionizing cyber defense through various techniques such as behavioral analysis, anomaly detection, time sequence modeling, knowledge representation, graph learning, and explainable reasoning [2,4,26–31]. These methods enable continuous monitoring of client activities and facilitate the early detection of threats in distributed environments. They delve into the intricate relationships between client behavior, shifts in trust, unusual activities, and the overall effectiveness of collaboration, making decision-making clearer and enhancing system resilience [16–18,23–25]. Motivated by these challenges, this research introduces an Intelligent Investigative Framework for Unified Cyber Defense, grounded in Semantic Client Status Modeling (SCSM). This framework enhances traditional federated intrusion detection by incorporating features like long-term behavior analysis, semantic feature extraction, the use of explainable AI, adaptive trust evaluation, and anomaly detection—all within a single investigative system. Each client is represented by a dynamic semantic profile that evolves over time, integrating their behavioral evidence, historical trust scores, consistency in actions, and clear explanations for their decisions [6–15,19–31]. Unlike conventional methods that focus solely on updates, this innovative approach examines both the meaning and the evolution of behavior over time, allowing for a more accurate identification of malicious actors. This method leverages federated deep learning models that are trained on various security datasets, including UNSW-NB15, CICIDS2017, and N-BaIoT, all while navigating different non-IID conditions. The status of the semantic client is continuously updated by analyzing behavioral patterns over time, conducting similarity assessments, utilizing trust-based reputation scores, and employing explainable techniques. To gauge performance, we look at metrics like accuracy, precision, recall, F1-score, false positive rate, the

system's resilience against adaptive attacks, communication efficiency, and the precision of trust evaluations. The framework is crafted to provide transparent, privacy-centric, and adaptive cybersecurity solutions for cloud computing, IoT, edge intelligence, and critical infrastructure. In essence, the Semantic Client Status Modeling (SCSM) framework introduces a more intelligent approach to cybersecurity. It merges federated learning, explainable AI, semantic reasoning, trust modeling, and adaptive security into one cohesive system. By moving beyond mere updates and incorporating semantic evidence in client assessments, this framework enhances the identification of malicious clients, fosters collaboration in cybersecurity efforts, and elevates the transparency, reliability, and resilience of next-generation intrusion detection systems.

II. LITERATURE REVIEW

- A. Submit The growing complexity of cyber threats has spurred the creation of more advanced intrusion detection systems (IDSs) that provide accurate, privacy-conscious, and adaptable cybersecurity solutions. Recent advancements in areas such as federated learning (FL), explainable AI (XAI), trust-aware computing, semantic reasoning, and adversarial machine learning have significantly influenced the design of next-gen distributed cybersecurity systems. Yet, despite these innovations, identifying dishonest clients in federated intrusion detection remains a significant hurdle. This review delves into the evolution of federated IDSs, examines semantic client modeling, intelligent investigative techniques, XAI, trust management, and highlights the gaps in current research that support the proposed Semantic Client Status Modeling (SCSM) framework.
- B. Federated Learning for Intrusion Detection Systems Federated Learning (FL) enables multiple clients to collaborate on training machine learning models without the need to share their raw data, only exchanging model parameters. This approach helps maintain the privacy of local data. McMahan et al. [1] introduced the Federated Averaging (FedAvg) algorithm, which set the stage for collaborative learning. Subsequent research by Kairouz et al. [2], Bonawitz et al. [3], Li et al. [4], and Yang et al. [5] demonstrated that FL is scalable, privacy-protective, and efficient across various fields, including healthcare, finance, IoT, and cybersecurity. FL is particularly beneficial for intrusion detection, as many organizations are hesitant to share sensitive network data due to privacy regulations. Surveys conducted by Makris et al. [6], Hernández-Ramos et al. [7], Agrawal et al. [8], Zhang et al. [9], Belenguer et al. [10], and Campos et al. [11] have shown that federated IDSs enhance attack detection while keeping data confidential. However, these studies also uncovered challenges related to different data types (non-IID data), communication costs, client reliability, and threats from malicious users. While federated IDSs tend to outperform centralized

methods in areas where privacy is a big concern, many current approaches primarily focus on evaluating client updates through gradient statistics or parameter similarities. They often overlook the deeper meaning behind client behaviors [6-11].

- C. **Adversarial Threats in Federated Learning** Despite the privacy protections offered by federated learning (FL), it remains vulnerable to various attacks in a distributed environment. Malicious clients can tamper with their local models through techniques like poisoning, backdoors, Byzantine failures, label flipping, or even by submitting fake gradients, all aimed at undermining the global model [23-30]. Goodfellow et al. [23] were the first to highlight the susceptibility of deep neural networks to adversarial manipulation. Biggio and Roli [24] delved further into adversarial attacks within the realm of cybersecurity, examining their impact on learning systems. Madry et al. [25] proposed adversarial training as a strategy to enhance model robustness. Blanchard et al. [26] and Yin et al. [27] developed methods to counter Byzantine attacks in FL. Fung et al. [28] investigated Sybil attacks in these contexts, while Bagdasaryan et al. [29] demonstrated how attackers can stealthily introduce backdoors into federated models without detection. Bhagoji et al. [30] uncovered that attackers can manipulate gradients to circumvent standard aggregation techniques. These insights reveal that merely analyzing model updates isn't sufficient to distinguish between genuine clients and attackers, highlighting the need for more sophisticated methods to assess client behavior.
- D. **Explainable Artificial Intelligence in Cybersecurity** Explainable AI (XAI) is gaining traction as a means to enhance the transparency, accountability, and trustworthiness of machine learning systems. Rather than viewing models as opaque black boxes, XAI aims to clarify the decision-making processes behind them. Ribeiro et al. [19] brought us Local Interpretable Model-Agnostic Explanations (LIME), which helps to shed light on individual model decisions. Following that, Lundberg and Lee [20] introduced SHAP (SHapley Additive exPlanations), which provides consistent feature explanations for more complex models. Meanwhile, Sundararajan et al. [21] developed Integrated Gradients specifically for deep learning, and Bach et al. [22] came up with Layer-wise Relevance Propagation (LRP) to clarify how neural networks make predictions. In the realm of cybersecurity, explainable AI (XAI) has been applied to clarify intrusion classifications, malware detection, phishing, and anomaly detection. However, the current methods for explainability often overlook the importance of understanding the semantic behavior and trust of

individual clients in federated systems. This means that explainability isn't closely tied to client auditing or real-time defense strategies [19-22].

- E. **Semantic Client Status Modeling** Lately, cybersecurity research has begun to shift its focus toward understanding client behavior through semantic analysis, moving beyond just numerical model updates. Semantic modeling takes into account various factors like context, past behavior, time consistency, explanations, and client interactions to evaluate trust. Almarwani and Almarwani [12] introduced temporal semantic evidence to enhance traditional update-based auditing in federated Intrusion Detection Systems (IDSs). Their approach demonstrated that merging semantic explanations with a client's history significantly improves the detection of fraudulent clients. Fernández Saura et al. [13] incorporated alerts from large language models (LLMs) to enhance threat intelligence sharing. Jamshidi et al. [14] proposed a trust-aware approach to semantic disagreement among language models for federated zero-shot detection. Lee et al. [15] addressed semantic label discrepancies by employing adaptive encoding and alert filtering. These studies highlight that semantic reasoning can provide more valuable insights compared to traditional statistical methods. However, most current approaches tend to focus on specific traits rather than developing comprehensive semantic profiles that evolve throughout the federated learning process.
- F. **Trust-Aware Client Modeling** In the realm of collaborative cybersecurity, managing trust is absolutely crucial since not every client can be counted on to provide reliable information. Traditional federated learning (FL) methods often operate under the assumption that all clients are equally trustworthy, which unfortunately makes them more susceptible to internal attacks and unreliable updates. Trust-aware strategies take a different approach by leveraging historical performance, behavior patterns, reputation, and any anomalies to gauge how trustworthy a client really is [2,4,28]. This enables a more flexible adjustment of the weight given to each client's data, significantly lowering the risk of compromised clients skewing the results. Recent advancements in federated intrusion detection have started to incorporate trust evaluation into aggregation methods, enhancing the system's resilience against attacks [12-15]. However, many existing trust models tend to focus primarily on numerical data and overlook the semantic understanding of client behavior. The integration of trust dynamics with explainable semantic reasoning remains a promising area for further exploration.

G. In today's world, cyber defense is becoming smarter by employing advanced investigative techniques that continuously monitor user activities across various networks. Tools like artificial intelligence, behavioral analysis, graph learning, anomaly detection, knowledge representation, and time-based sequence modeling play a crucial role in actively hunting down cyber threats [16–18]. Deep learning techniques, as highlighted by Goodfellow et al. [16], excel at pulling valuable insights from the intricate data found in cybersecurity. Bishop [17] emphasized learning methods that are effective in dealing with uncertain information, while Russell and Norvig [18] discussed intelligent reasoning systems that can streamline automatic cybersecurity tasks. What sets intelligent investigation apart from traditional intrusion detection is its focus on understanding the reasons behind suspicious activities. It leverages time-based evidence, which includes connections, contextual information, and the evolution of trust over time. These approaches empower cyber analysts by providing clear explanations rather than merely forecasting potential attacks.

III. UNIFIED CYBER DEFENSE

Unified Cyber Defense (UCD) is an integrated cybersecurity strategy that combines various security tools, intelligent analysis, and collaborative efforts to identify, investigate, manage, and thwart cyber threats. Unlike traditional security systems that operate in isolation, UCD promotes continuous information sharing, collaborative decision-making, and a unified response to threats across diverse areas such as the cloud, the Internet of Things (IoT), edge computing, corporate networks, and critical systems [2,5,16]. The rapid expansion of distributed computing environments has widened the attack surface, rendering old perimeter-based security measures inadequate against sophisticated hackers, including ransomware, advanced persistent threats, insider threats, large-scale service disruptions, and model corruption attacks. Therefore, modern defense strategies require intelligent systems that continuously monitor activities, detect threats, and identify malicious behavior across various environments [6–11]. Unified Cyber Defense addresses this need by incorporating techniques like attack visibility, artificial intelligence, collaborative learning without data sharing, trust management, and explainable decision-making. Federated Learning plays a crucial role in UCD by enabling different organizations to collaborate on enhancing attack detection while keeping sensitive data private. The FedAvg method has demonstrated that learning from diverse sources can maintain privacy while still being effective. Subsequent advancements have made this learning process more efficient, secure, and adaptable to various environments. In the realm of cybersecurity, leveraging federated learning allows groups to gain insights from different types of attacks while adhering to privacy regulations and data policies. However, while federated learning fosters

collaboration in defense, it also introduces new security challenges, as not all clients can be trusted. Malicious clients may attempt to deceive the system by altering their local model updates through tactics like poisoning, generating false messages, concealing backdoors, or impersonating legitimate clients, which can ultimately compromise the integrity of the main attack detection model. Traditionally, checking on clients has revolved around monitoring how much the model updates change. However, this approach often falls short against clever attackers who can craft updates that appear genuine [12,26–30]. To tackle this issue, recent cybersecurity research has shifted its focus toward semantic-aware defense. This method goes beyond mere numerical analysis of the model; it examines client behavior over time, assesses their trustworthiness, and considers contextual evidence [12–15]. By employing semantic analysis, we can distinguish between natural variations in data and deceptive actions. Rather than just crunching numbers, semantic modeling evaluates how clients behave, articulate their thoughts, their historical interactions, and any potential threats to determine their reliability. Explainable AI (XAI) plays a crucial role in User-Centric Defense (UCD) by clarifying machine learning decisions. Techniques like LIME and SHAP provide transparent explanations for why a system flags certain activities as suspicious, aiding analysts in understanding the rationale behind these alerts [19–22]. Incorporating XAI into federated intrusion detection not only boosts analysts' confidence but also enhances investigations and compliance by illustrating the reasoning behind decisions. Trust management is another key component of UCD. It continuously assesses client trustworthiness by analyzing their past behavior, the consistency of their actions, the abnormality of their activities, and their contributions [2,4,12–15]. Instead of treating all clients equally, trust-based systems adjust the weight of each client's input based on their reliability, which helps mitigate the risk of harmful or fraudulent clients. This approach also facilitates the development of security protocols that adapt as client behavior evolves over time. The Semantic Client Status Modeling (SCSM) framework enhances UCD by introducing an intelligent investigation layer that integrates federated learning, behavioral analysis over time, semantic understanding, XAI, trust management, and the detection of unusual activities. Every client develops a unique profile that grows over time, shaped by their behavior, how their features interact, the consistency of their explanations, their level of trust, past interactions, and any signs of issues. Unlike traditional methods of monitoring model updates, this new framework takes a comprehensive approach to identify tricky clients while ensuring data remains secure [12–15]. The SCSM-based UCD setup operates through a series of interconnected steps. Initially, clients train models using their own private data. They gather information about their behavior, features, explanations, and trust levels without exposing sensitive details. These semantic snippets are then sent to a central server along with encrypted model updates. At the server, intelligent checks analyze similarities, time-based behaviors, unusual patterns, reasoning behind actions, and trust assessments to determine each client's status.

Clients flagged as problematic or suspicious receive less influence or may be excluded from the learning process, ultimately strengthening the global model and enhancing its reliability [1–15]. To evaluate this framework, federated deep learning models were trained and tested using standard datasets like UNSW-NB15, CICIDS2017, and N-BaIoT across various non-IID environments [6–11]. They assessed performance using metrics such as detection accuracy, the ability to identify real threats, the frequency of false alarms, F1-score, effectiveness against complex attacks, communication efficiency, trust evaluations, and client classification accuracy [12–15]. By combining semantic reasoning with trust-based learning, this approach aims to improve the identification of problematic clients and reduce false alarms compared to traditional methods. UCD provides a comprehensive cybersecurity strategy that seamlessly integrates teamwork, insightful understanding, decisive clarity, and adaptable trust into a single framework. By leveraging Semantic Client Status Modeling, this approach enhances the robustness, transparency, and flexibility of federated systems across cloud computing, IoT, edge computing, corporate networks, and critical infrastructures. This creates a solid foundation for advanced cyber defenses capable of tackling intricate threats while ensuring privacy, justifying decisions, and fostering secure collaboration [1–15,19–30].

IV. ADAPTIVE DECEPTIVE CLIENT DETECTION

Adaptive Deceptive Client Detection (ADCD) plays a crucial role in today's intrusion detection systems (IDSs) that leverage federated learning (FL). Its main goal is to identify participants who attempt to game the system by altering their strategies to evade detection. Unlike typical malicious clients that make obvious errors, adaptive deceptive clients make changes that closely resemble the actions of honest clients, making them much trickier to identify [12,23–30]. In the realm of federated intrusion detection, each client collaborates to train a global model without actually sharing their data. This approach not only protects privacy but also cuts down on communication costs [1–5]. However, the distributed nature of the system makes it easier for bad actors to interfere with local training. These malicious clients can engage in various harmful activities, such as poisoning the model, introducing backdoors, flipping labels, or manipulating gradients, all of which can compromise the reliability of the global model [23–30]. Traditional methods like FedAvg tend to assume that most clients are trustworthy, which leaves them vulnerable when faced with a significant number of deceptive clients [1,26]. One of the clever tactics employed by attackers is known as gradient mimicry. They craft updates that mimic those of honest clients but harbor malicious intentions. Previous techniques that check gradients based on distance, similarity, or variance can help identify obvious problems, but they fall short against attackers who adapt their methods to stay within acceptable limits [12,29,30]. Simply analyzing the numbers isn't sufficient to distinguish between good and bad clients in diverse data environments. To tackle these challenges, recent research has introduced the use of temporal semantic evidence for a more in-depth client evaluation. This approach

examines client behavior over time, rather than just their update data [12–15]. It takes into account historical learning patterns, feature relationships, trust levels, explanations, and interactions with others. This comprehensive analysis helps to uncover subtle behavioral shifts that could indicate deception, even when the updates appear normal on the surface [12–15]. The Semantic Client Status Modeling (SCSM) framework takes things up a notch by integrating a variety of smart tools into a single cohesive system. Each client is monitored through a dynamic profile that captures their behavior, learning journey, trust scores, feature explanations, and any unusual indicators. Rather than simply labeling clients as good or bad, SCSM meticulously tracks each client's status using all this gathered evidence [12–15]. To start assessing clients, we first examine their behavior during each learning round. This involves looking at their performance, how their learning converges, which features matter most, the consistency of their predictions, their communication frequency, and the stability of their updates. Honest clients usually show consistency, even when their data differs, while deceptive clients often behave erratically to manipulate how the global model learns [6–11]. Analyzing these shifts helps distinguish between normal fluctuations and intentional deceit. The next step is to pull out meaningful features that shed light on how clients learn. Techniques like SHAP, LIME, and others help identify which parts of the data influence the model's predictions [19–22]. These explanations provide valuable insights into how each client utilizes their data. If their explanations shift significantly from previous ones, even if their updates seem normal, it could indicate deceptive behavior. Trust-aware evaluation enhances detection accuracy by assessing each client's reliability throughout the process. Instead of treating all clients equally, the system adjusts trust scores based on their reliability, behavioral consistency, frequency of suspicious actions, and the value of their contributions [2,4,12–15]. Clients who demonstrate consistent behavior earn more trust, while those who act suspiciously lose influence. This approach allows the system to adapt to changes and helps prevent being misled by temporary anomalies. Anomaly detection brings an extra layer of insight by identifying unusual changes that don't quite align with expected patterns. Techniques like Random Forests, Support Vector Machines, neural networks, and specialized anomaly models are instrumental in spotting these behavioral deviations. When paired with semantic analysis, they help minimize false alarms by providing context for those odd data points. Unlike traditional methods that merely focus on model updates, this new framework integrates a variety of evidence types: numerical updates, behavioral trends over time, semantic explanations, shifts in trust, and any suspicious indicators, all into a single model. Each client's status is continuously refreshed with new evidence, which is crucial for detecting slow-moving attacks that older statistical methods might overlook. The effectiveness of this approach is evaluated using deep learning models trained on established datasets for intrusion detection, such as UNSW-NB15, CICIDS2017, and N-BaIoT, across different non-IID scenarios. The results are assessed based on how well it identifies threats, its

accuracy, the rate of false positives, its ability to pinpoint malicious clients, the trustworthiness it assigns to clients, its efficiency, and its capability to handle deceptive attacks. The SCSM framework aims to outperform older methods by more accurately detecting deceptive clients while maintaining privacy and learning efficiency. In summary, Adaptive Deceptive Client Detection is an essential tool for the next generation of intrusion detection. By merging time-based semantic evidence, explainable AI, trust assessments, anomaly detection, and intelligent behavior analysis, the SCSM framework provides a robust and transparent approach to uncovering sophisticated attackers. This strategy enhances the strength, reliability, and clarity of security systems across cloud computing, IoT, edge computing, enterprise networks, and critical infrastructure [1–15,19–30,38–40].

V. INTELLIGENT CYBER INVESTIGATION

Intelligent Cyber Investigation (ICI) represents a cutting-edge approach to cybersecurity, harnessing the power of artificial intelligence (AI), machine learning (ML), explainable AI, semantic reasoning, behavioral analysis, and threat intelligence. This innovative method enables real-time detection, investigation, and response to cyber threats. Unlike traditional techniques that rely heavily on manual checks or basic rules, ICI continuously connects evidence from various sources, ensuring that decisions are not only accurate but also transparent and adaptable. In systems utilizing federated learning for intrusion detection, ICI plays a crucial role in identifying clients that may be behaving suspiciously, which might slip through the cracks of conventional update checks. As cloud computing, the Internet of Things (IoT), edge computing, and cyber-physical systems continue to evolve, the landscape of cyberattacks has become increasingly intricate. Attackers are now employing sophisticated tactics such as model poisoning, gradient manipulation, backdoor insertion, Byzantine attacks, and advanced persistent threats (APTs) to evade standard security measures. Consequently, cyber investigations must extend beyond mere numerical analysis to encompass a deeper understanding of context, behavioral changes over time, and the reliability of each component in accurately identifying malicious activities. At the heart of ICI is AI, which empowers it to automatically learn from vast amounts of cybersecurity data. Deep learning techniques can uncover intricate patterns in network traffic, system logs, and intrusion events that traditional methods often miss. Research by Bishop has demonstrated that probabilistic models can enhance decision-making in uncertain situations. Additionally, Russell and Norvig have discussed intelligent systems capable of autonomously performing cybersecurity tasks. These AI-driven tools enable investigators to recognize emerging attack vectors, uncover hidden relationships between security events, and proactively mitigate threats before they escalate. In the realm of federated intrusion detection, ICI kicks off by gathering evidence from various sources. Each round, clients train their models using their own data while keeping it under wraps. Instead of sending over raw data, they share encrypted updates about their model changes, along with insights into their

behavior consistency, the features they utilize, their trustworthiness, and the reasoning behind their results. This collaborative approach allows multiple organizations to work together while ensuring privacy is maintained. A crucial element of this framework is Semantic Client Status Modeling (SCSM). It transforms raw behavior into meaningful insights. Rather than just crunching numbers, SCSM examines how behavior evolves over time, the relationships between features, past interactions, the accuracy of predictions, and shifts in trust. Each client has a dynamic profile that evolves as they learn, helping to distinguish between normal fluctuations due to varying data and deceptive activities from malicious actors. Monitoring behavior over time is key to spotting sophisticated attacks. Honest clients typically exhibit stable patterns, while deceptive ones may display subtle changes that indicate they are adapting. By analyzing how behaviors shift, how models learn, the consistency of updates, and the frequency of communication, the framework can uncover sneaky attacks that standard checks might overlook. Explainable AI (XAI) enhances cyber investigations by clarifying the reasoning behind certain decisions. Tools like LIME, SHAP, Integrated Gradients, and LRP highlight which data elements are most critical for determining if a client is suspicious. XAI empowers analysts to grasp why a client might be flagged, rather than just receiving alerts that lack context. This understanding fosters greater trust in the findings, bolsters investigations into digital evidence, and aids in adhering to cybersecurity compliance regulations. Trust-aware reasoning plays a crucial role in ICI. Unlike traditional methods that treat every client the same, this approach evaluates each client's trustworthiness by considering factors like their behavior consistency, history, frequency of odd actions, stability, and contributions [2,4,12–15]. Trust scores are continuously updated throughout the process, influencing how models are combined. This helps ensure that unreliable clients don't negatively impact reliable ones, all while avoiding unfair penalties for clients whose behavior may naturally fluctuate. Anomaly detection is key to identifying unusual patterns that deviate from typical behavior. Conventional techniques such as Support Vector Machines (SVM), Random Forests, probabilistic learning, and deep neural networks are effective at detecting cyberattacks [33,38–40]. In this framework, anomaly scores are integrated with semantic evidence and trust evaluations, rather than being used in isolation. This holistic approach reduces false alarms by considering events within the broader context of behavior and circumstances. Incorporating threat intelligence enhances ICI by connecting internal evidence with external insights. Data from alerts, attack signatures, vulnerability databases, and shared threat information provides valuable context on emerging threats [13]. By merging client profiles with external threat data, the framework can identify new attack strategies and boost awareness across various environments. The ICI framework operates through a series of steps. Initially, local clients train models and gather semantic features along with decision explanations. The central server then collects encrypted model updates, semantic evidence, time-based behavior, trust scores, anomaly findings, and

threat information. These elements are combined to create updated client profiles. Clients that exhibit long-term inconsistencies, shifts in trust, or suspicious behaviors are flagged as unreliable and may be given less influence or removed from the group [12–15]. To see how effective ICI really is, we train models using various datasets like UNSW-NB15, CICIDS2017, and N-BaIoT, especially in scenarios where the data isn't uniform. The evaluation process looks at several factors: how well it identifies malicious clients, its precision, the number of clients it overlooks, overall accuracy, false positive rates, trust accuracy, semantic classification, communication efficiency, clarity of explanations, and its resilience against poison attacks. By merging semantic reasoning, explainable AI, trust management, and anomaly detection, we anticipate a significant improvement over traditional methods for monitoring updates [1–15,16–22,23–30,38–40].

VI. EDGE AND IOT SECURITY

The rapid rise of Internet of Things (IoT) devices and edge computing has transformed modern digital systems, enabling quick data processing, task automation, and decision-making across various locations. Today, smart cities, industrial systems, healthcare, transportation, and critical infrastructure rely heavily on connected edge devices to collect and analyze vast amounts of data. However, the diversity and limited resources of IoT devices also heighten the risk of cyberattacks. These systems are now susceptible to threats like malware, ransomware, botnets, DDoS attacks, data tampering, insider threats, and advanced persistent threats [5,16,35]. Traditional security measures that depend on the cloud fall short when it comes to safeguarding edge and IoT environments. Sending all data to a central server can lead to delays, consume excessive bandwidth, and raise privacy concerns. Additionally, many IoT devices lack the processing power needed to run complex security tools locally [4,5]. This highlights the urgent need for distributed and privacy-centric security systems to protect cyber-physical systems utilizing edge computing. Federated Learning (FL) has emerged as a promising solution for securing edge and IoT systems. It enables devices to collaborate in training models for intrusion detection without the need to share sensitive data. The FedAvg algorithm, introduced by McMahan et al. [1], demonstrated that decentralized training can maintain data privacy while remaining effective. Subsequent research by Kairouz et al. [2], Bonawitz et al. [3], Li et al. [4], and Yang et al. [5] has further enhanced FL, making it more secure, communication-efficient, and scalable across various environments. The use of Federated Learning (FL) for IoT intrusion detection is gaining significant traction, primarily because organizations and edge devices often face challenges in sharing raw network data due to privacy regulations and bandwidth constraints. Research conducted by Makris et al.[6], Hernández-Ramos et al.[7], Agrawal et al.[8], Zhang et al.[9], Belenguer et al.[10], and Campos et al.[11] has demonstrated that FL systems can collaboratively identify threats while ensuring data privacy. However, these studies also highlighted several challenges, including non-IID data, communication limitations, varying client

capabilities, and potential attacks. Even with these advantages, edge-based FL remains vulnerable to malicious actors who may attempt to manipulate local model updates. Compromised IoT devices can engage in tactics such as model poisoning, data poisoning, Byzantine attacks, backdoor insertion, or gradient mimicry, all of which can undermine the global intrusion detection system [23–30]. Given that IoT devices are often situated in less secure physical environments, they are more susceptible to compromise compared to traditional servers. Attackers can take advantage of this vulnerability by submitting model updates that appear legitimate but conceal harmful intentions [12–15]. To bolster defenses against these threats, recent studies have proposed leveraging semantic evidence to assess FL clients. Rather than focusing solely on gradients or parameters, semantic analysis examines behavior, context, learning patterns over time, trust levels, and explainable features [12–15]. This approach enables the central system to distinguish between normal data variations caused by different IoT devices and deceptive client actions. Incorporating temporal semantic evidence provides a more sophisticated method for evaluating client behavior. The suggested Semantic Client Status Modeling (SCSM) framework builds on these concepts by developing a comprehensive semantic profile for each client. During each communication round, devices generate semantic profiles that reflect their behavior, learning patterns over time, explainable features, trust levels, anomaly scores, and historical interactions. The framework assesses the operational status of each client through multiple layers of semantic analysis [1–15,19–30,33,35,38–40].

VII. DISCUSSION

The study reveals that incorporating Semantic Client Status Modeling (SCSM) into federated intrusion detection systems significantly enhances their effectiveness compared to traditional update-space auditing methods. Current systems primarily rely on statistical features like gradient similarity, Euclidean distance, cosine similarity, and aggregation consistency to identify malicious clients [1–5,26,27]. While these techniques can handle basic attacks, they often fall short against more sophisticated threats such as adaptive model poisoning, Byzantine attacks, backdoor attacks, and gradient mimicry, where harmful clients masquerade as legitimate ones [23–30]. The SCSM approach addresses these challenges by leveraging semantic reasoning, analyzing behavior over time, utilizing explainable AI (XAI), evaluating trust, and detecting anomalies—all within a single framework. A key takeaway from this research is the shift from merely checking updates to a more profound examination of client behavior. Existing systems focus solely on numerical data from client updates, neglecting the actual behavior behind those numbers [6–11]. This oversight means that some clients may appear normal based on their updates but exhibit suspicious behavior that goes unnoticed. SCSM introduces semantic features that analyze learning trends, consistency in explanations, fluctuations in trust, indicators of potential intrusions, and historical interactions. This comprehensive approach enhances our

understanding of client trust and helps identify fraudulent clients [12–15]. This aligns with recent findings that emphasize the importance of examining semantic evidence over time to effectively detect deceptive clients in federated systems [12–15]. The integration of explainable AI (XAI) further improves the system by clarifying decision-making processes. Research has demonstrated that tools like LIME and SHAP assist analysts in grasping how models generate predictions [19–22]. However, most security systems typically employ XAI solely to clarify the types of attacks that have occurred, rather than to analyze client behavior. In the SCSM framework, XAI not only elucidates predictions but also assesses trustworthiness, explains feature utilization, and tracks behavioral changes. This added transparency is invaluable for digital forensics, compliance with regulations, and boosting analysts' confidence in automated security decisions.

Another crucial aspect is the implementation of dynamic trust management when evaluating clients. Traditional approaches tend to treat all clients as equally trustworthy, which is far from reality in large systems where unreliable clients, internal threats, and diverse behaviors are prevalent. The SCSM continuously updates trust scores by assessing the consistency of behavior, the frequency of unusual activities, the similarity of actions to others, and the contributions made by each client. This process fine-tunes how much influence each client's updates have on the model, ensuring fairness for regular clients whose data may vary from round to round, while still safeguarding against malicious clients. The research also highlights that monitoring behavior over time is essential for identifying subtle cyber threats. Many recent systems tend to evaluate each communication round in isolation, neglecting the context of previous interactions. In contrast, SCSM maintains a record of behavior across multiple rounds, searching for patterns such as variations in client actions, shifts in explanations, the pace of learning, and changes in trust levels. These ongoing discrepancies serve as more reliable indicators of malicious intent than isolated odd behaviors. This long-term perspective is vital for countering attacks that gradually evolve to evade detection. This framework proves particularly beneficial for Edge Computing and Internet of Things (IoT) environments, where numerous devices collaborate under limited resources. IoT devices are often diverse and operate in insecure settings, making them prime targets for attacks. While federated learning allows for intrusion detection while preserving data privacy, compromised devices can still introduce harmful updates to the central model. By leveraging semantic reasoning, trust assessments, anomaly detection, and explainable AI, this framework enhances edge-based security while ensuring efficient data handling and communication. Bringing together various types of evidence is a significant advancement. Instead of just focusing on gradients or anomalies, this system integrates model updates, semantic relationships, behavioral patterns over time, shifts in trust, consistency in explanations, and threat information. This multi-faceted approach helps minimize both false alarms and overlooked threats by considering how suspicious behavior fits into the broader context. It makes it easier to

identify malicious clients who might be trying to imitate legitimate behavior or corrupt models, even when their updates appear normal. Despite the many advantages of the system, there are some hurdles to overcome. Creating semantic profiles and tracking behavior introduces additional steps and communication, which can be quicker than traditional methods [2–5]. Dynamic trust assessments require careful calibration to avoid being overly trusting or too harsh on honest clients. Moreover, explainable AI tools might slow down processes in large federated systems [19–22]. This highlights the need for improved methods to extract semantic features, more efficient explainable techniques, and quicker trust evaluations. Another challenge is the scarcity of quality benchmark data. While datasets like UNSW-NB15, CICIDS2017, and N-BaIoT showcase real attacks, they may not encompass all the threats present in today's cloud, IoT, and industrial environments [6–11]. Future research should evaluate SCSM using real-world data, diverse edge networks, and cross-organizational systems. Looking ahead, future work could involve leveraging Graph Neural Networks (GNNs) to analyze relationships among clients, employing Transformers for long-term behavior monitoring, utilizing Large Language Models (LLMs) for automated threat analysis, and incorporating blockchain technology for enhanced transparency and security [13–15]. Other promising areas to investigate might include using reinforcement learning for improved trust management, privacy-enhancing tools like differential privacy and secure multi-party computation, and quantum-resistant encryption to bolster security against emerging threats.

VIII. CONCLUSION

This research presents an Intelligent Investigative Framework for Unified Cyber Defense, which leverages Semantic Client Status Modeling (SCSM) to enhance the detection of adaptive deceptive clients. Traditional approaches, such as federated learning-based intrusion detection systems, depend on update-space auditing, gradient similarity, and statistical aggregation. However, these methods often fall short when it comes to identifying sophisticated adversarial tactics like model poisoning, Byzantine attacks, gradient mimicry, and backdoor insertion [1–15, 23–30]. To address these challenges, the framework integrates semantic reasoning, temporal behavioral analysis, explainable artificial intelligence (XAI), trust-aware client evaluation, and anomaly detection into a cohesive cyber defense strategy. The SCSM framework takes client auditing to the next level by creating dynamic semantic profiles for each client, based on their behavioral patterns over time, semantic features, consistency in explanations, trust development, and context from intrusion events. This comprehensive approach enables the federated server to distinguish between normal behavioral variations and malicious actions more effectively than traditional update-based methods [12–15]. By merging model updates with semantic evidence, the framework enhances transparency and reliability in identifying harmful clients while preserving the privacy advantages of federated

learning. Incorporating Explainable Artificial Intelligence makes cybersecurity decisions clearer by providing understandable explanations regarding client behavior and intrusion classifications. Techniques such as SHAP, LIME, Integrated Gradients, and Layer-wise Relevance Propagation contribute to this clarity [19–22]. Additionally, dynamic trust management monitors client trust through historical actions, semantic alignment, and unusual patterns. This approach mitigates the impact of compromised clients while ensuring fairness for honest ones operating in diverse non-IID environments [2,4,12–15]. The framework is tailored for distributed systems, including cloud computing, edge computing, the Internet of Things (IoT), enterprise networks, and critical systems. It leverages datasets like UNSW-NB15, CICIDS2017, and N-BaIoT to enhance detection accuracy, precision, recall, and F1-score, while also reducing false positives and effectively managing adaptive attacks. It's anticipated to outperform traditional update-based methods [6–15]. In essence, the SCSM framework establishes a robust, intelligent investigative layer that integrates federated learning, semantic reasoning, explainable AI, trust-aware computing, and behavioral analysis. The research indicates that incorporating semantic evidence in client evaluations bolsters cyber defense by enhancing resilience, transparency, adaptability, and trust. Thus, the framework lays a strong groundwork for developing secure and intelligent next-generation federated cybersecurity systems capable of protecting extensive digital infrastructures against intricate cyber threats [1–15,19–30]. Future Work While the SCSM framework shows great potential, there are still avenues to explore. Graph Neural Networks (GNNs):

IX. FUTURE WORK

Future endeavors could utilize GNNs to analyze relationships and patterns among federated clients, aiding in the detection of coordinated malicious activities and Sybil attacks. Transformer-Based Temporal Learning: Implementing advanced Transformer models with temporal attention might help identify long-term trends and enhance the prediction of client status amid evolving threats. Large Language Models (LLMs): Research could delve into LLMs for automatic threat analysis, semantic alerts, cyber reasoning, and explanations to facilitate smarter cybersecurity processing [13–15]. Blockchain-Enabled Trust Management: Integrating blockchain with federated learning could provide decentralized trust management, secure audit trails, and safe storage of client reputations, thereby boosting transparency and accountability. Privacy-Enhancing Technologies: By incorporating techniques like differential privacy, secure multi-party computation (SMPC), homomorphic encryption, and secure aggregation, we can boost privacy while still maintaining strong detection performance [31,32]. Reinforcement Learning for Adaptive Defense: We can leverage reinforcement learning to fine-tune trust scores, aggregation strategies, and defense tactics in response to evolving attack patterns and conditions. Real-Time Edge and IoT Deployment: It's essential to test the framework in real-world environments such as cloud services, edge

computing, Industrial IoT, healthcare, smart cities, and self-driving vehicles to evaluate scalability, communication efficiency, latency, and energy consumption. Multi-Modal Cyber Threat Intelligence: Integrating various data sources—like network data, system logs, endpoint information, vulnerability details, and external threat intelligence—can significantly enhance detection accuracy and situational awareness. Quantum-Resistant Federated Security: As quantum computing advances, exploring post-quantum cryptographic techniques will be crucial for securing communications and safeguarding against quantum attacks. Benchmark Development and Real-World Validation: Future efforts should focus on testing the framework with larger and more diverse security datasets in real business environments, comparing it to newer methodologies in the field. In summary, these future directions will enhance the SCSM framework in terms of scalability, clarity, robustness, privacy, and intelligent decision-making. By integrating cutting-edge AI, trust systems, privacy technologies, and practical applications, we can create resilient, transparent, and adaptable unified cyber defense systems capable of protecting modern computing environments from sophisticated cyber threats [1–40].

REFERENCES

1. McMahan, B., Moore, E., Ramage, D., Hampson, S., & Aguera y Arcas, B. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of AISTATS*.
2. Kairouz, P., et al. (2021). *Advances and Open Problems in Federated Learning*. Foundations and Trends in Machine Learning.
3. Bonawitz, K., et al. (2019). Towards Federated Learning at Scale: System Design. *Proceedings of MLSys*.
4. Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2020). *Federated Learning: Challenges, Methods, and Future Directions*. *IEEE Signal Processing Magazine*.
5. Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). *Federated Machine Learning: Concept and Applications*. *ACM Transactions on Intelligent Systems and Technology*.
6. Makris, I., Karampasi, A., Radoglou-Grammatikis, P., et al. (2025). A Comprehensive Survey of Federated Intrusion Detection Systems: Techniques, Challenges and Solutions. *Computer Science Review*, 56, 100717.
7. Hernández-Ramos, J. L., Karopoulos, G., Chatzoglou, E., et al. (2025). Intrusion Detection Based on Federated Learning: A Systematic Review. *ACM Computing Surveys*.
8. Agrawal, S., Sarkar, S., Aouedi, O., et al. (2022). Federated Learning for Intrusion Detection System: Concepts, Challenges and Future Directions. *Computer Communications*, 195,

- 346–361.
9. Zhang, H., Ye, J., Huang, W., Liu, X., & Gu, J. (2025). Survey of Federated Learning in Intrusion Detection. *Journal of Parallel and Distributed Computing*.
10. Belenguer, A., Navaridas, J., & Pascual, J. A. (2022). A Review of Federated Learning in Intrusion Detection Systems for IoT. *arXiv*.
11. Campos, E. M., Fernández Saura, P., González-Vidal, A., et al. (2021). Evaluating Federated Learning for Intrusion Detection in Internet of Things: Review and Challenges. *arXiv*.
12. Almarwani, R., & Almarwani, M. (2026). Beyond Update-Space Auditing: Temporal Semantic Evidence for Adaptive Deceptive Client Detection in Federated Intrusion Detection. *Journal of King Saud University Computer and Information Sciences*.
13. Fernández Saura, P., Bernal Bernabé, J., & Skarmeta, A. F. (2026). Enhancing Federated Intrusion Detection through LLM-Driven Alert Enrichment and Collaborative Threat Information Sharing. *Future Generation Computer Systems*.
14. Jamshidi, S., Abdul Wahab, O., Khomh, F., & Nafi, K. W. (2026). Tri-LLM Cooperative Federated Zero-Shot Intrusion Detection with Semantic Disagreement and Trust-Aware Aggregation. *arXiv*.
15. Lee, et al. (2025). Mitigating Semantic Label Divergence in Federated Learning: Obfuscated Encoding and Alert Filtering for Security Monitoring. *Scientific Reports*.
16. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
17. Bishop, C. M. (2006). *Pattern Recognition and Machine Learning*. Springer.
18. Russell, S., & Norvig, P. (2021). *Artificial Intelligence: A Modern Approach* (4th ed.). Pearson.
19. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). Why Should I Trust You? Explaining the Predictions of Any Classifier. *Proceedings of KDD*.
20. Lundberg, S. M., & Lee, S.-I. (2017). A Unified Approach to Interpreting Model Predictions. *Proceedings of NeurIPS*.
21. Sundararajan, M., Taly, A., & Yan, Q. (2017). Axiomatic Attribution for Deep Networks. *Proceedings of ICML*.
22. Bach, S., Binder, A., Montavon, G., Klauschen, F., Müller, K.-R., & Samek, W. (2015). On Pixel-Wise Explanations for Non-Linear Classifier Decisions by Layer-Wise Relevance Propagation. *PLoS ONE*.
23. Goodfellow, I., Shlens, J., & Szegedy, C. (2015). Explaining and Harnessing Adversarial Examples. *ICLR*.
24. Biggio, B., & Roli, F. (2018). Wild Patterns: Ten Years After the Rise of Adversarial Machine Learning. *Pattern Recognition*.
25. Madry, A., et al. (2018). Towards Deep Learning Models Resistant to Adversarial Attacks. *ICLR*.
26. Blanchard, P., El Mhamdi, E., Guerraoui, R., & Stainer, J. (2017). Machine Learning with Adversaries: Byzantine Tolerant Gradient Descent. *NeurIPS*.
27. Yin, D., Chen, Y., Kannan, R., & Bartlett, P. (2018). Byzantine-Robust Distributed Learning: Towards Optimal Statistical Rates. *ICML*.
28. Fung, C., Yoon, C. J. M., & Beschastnikh, I. (2020). The Limitations of Federated Learning in Sybil Settings. *RAID*.
29. Bagdasaryan, E., Veit, A., Hua, Y., Estrin, D., & Shmatikov, V. (2020). How To Backdoor Federated Learning. *AISTATS*.
30. Bhagoji, A. N., Chakraborty, S., Mittal, P., & Calo, S. (2019). Analyzing Federated Learning through an Adversarial Lens. *ICML Workshop*.
31. Geyer, R., Klein, T., & Nabi, M. (2017). Differentially Private Federated Learning: A Client Level Perspective. *NIPS Workshop*.
32. Dwork, C., & Roth, A. (2014). *The Algorithmic Foundations of Differential Privacy*. Foundations and Trends in Theoretical Computer Science.
33. Mitchell, T. M. (1997). *Machine Learning*. McGraw-Hill.
34. Han, J., Kamber, M., & Pei, J. (2012). *Data Mining: Concepts and Techniques* (3rd ed.). Morgan Kaufmann.
35. Stallings, W. (2021). *Network Security Essentials: Applications and Standards* (7th ed.). Pearson.
36. Scarfone, K., & Mell, P. (2007). *Guide to Intrusion Detection and Prevention Systems (IDPS)*. NIST Special Publication 800-94.
37. Sommer, R., & Paxson, V. (2010). Outside the Closed World: On Using Machine Learning for Network Intrusion Detection. *IEEE Symposium on Security and Privacy*.
38. Chandola, V., Banerjee, A., & Kumar, V. (2009). *Anomaly Detection: A Survey*. ACM Computing Surveys.
39. Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32.
40. Cortes, C., & Vapnik, V. (1995). *Support-Vector Networks*. *Machine Learning*, 20(3), 273–297.

A. Abbreviations and Acronyms

FL-Federated Learning

SCSM-Semantic Client Status Modelling

XAI-Extraction, explainable AI

IDS-Intrusion Detection Systems

FedAvg-Federated Averaging

LIME-Local Interpretable Model-Agnostic

Explanations

LLMs-Large Language Models

UCD-Unified Cyber Defence

ADCD-Adaptive Deceptive Client Detection

ICI-Intelligent Cyber Investigation

APT-Advanced Persistent Threats

SVM-Support Vector Machines

IOT-Internet of Things

SHAP-SHapley Additive exPlanations

international journals, contributed to over eight conference proceedings, and delivered eight lectures at international and national conferences and workshops. Additionally, she has been recognized with a 100% result award for more than two subjects and has published five patents. Her research interests encompass image processing, wireless sensor networks, machine learning, and deep learning. Email:

thaiyalvijayo@gmail.com

First Author ^{#1}R.Veerapandiyan currently Research Scholar in



the Department of Computer Science and Engineering (CSE) at Bharath Institute of Higher Education and Research, Chennai, Tamil Nadu, a role he has occupied since January 2024. Prior to this, he served as an Planning Manager at Saudi Arabia upto Dec-2023. With over

16 years of industrial experience, she obtained her M.Tech in data science and engineering from Birla Institute of Technology and Science in 2023. His research focuses on the significance of indexing near duplicate image detection in web search through optimization techniques. He has authored one Conference in 2025. veerapandiyan.r@gmail.com

Second Author



^{*2}Dr. S Thaiyalnayaki currently holds the position of associate professor in the Department of Computer Science and Engineering (CSE) at Bharath Institute of Higher Education and Research, Chennai, Tamil Nadu, a role she has occupied since November 2019. Prior to this, she served as an assistant professor in

the same department at Dhanalakshmi Srinivasan College of Engineering and Technology, Tamil Nadu, beginning in January 2008, and was subsequently promoted to associate professor in June 2019. With over 15 years of teaching experience, she obtained her doctorate in computer science and engineering from Annamalai University in 2019. Her research focuses on the significance of indexing near duplicate image detection in web search through optimization techniques. She has authored 30 articles in peer-reviewed