



BioTrust-IoT: A Biologically Inspired Adaptive Trust System for Detecting Man-in-the-Middle Attacks

Banupriya. D

Department of Computer Applications, Nehru Arts and Science College (Autonomous), Coimbatore, Tamil Nadu, India.

banupriya170393@gmail.com

K. Prathapchandran

Department of Computer Applications, Nehru Arts and Science College (Autonomous), Coimbatore, Tamil Nadu, India.

✉ kprathapchandran@gmail.com

Received: 09 July 2025 / Revised: 23 September 2025 / Accepted: 23 October 2025 / Published: 30 December 2025

Abstract – The growing use of Internet of Things (IoT) devices has made them increasingly vulnerable to cyber threats, with Man-in-the-Middle (MitM) attacks standing out as one of the most dangerous. In these attacks, a malicious entity secretly intercepts and alters the communication between two trusted IoT devices. This can result in stolen data, unauthorized command execution, and disruption of normal device functions. Traditional security measures often fall short in detecting such threats, especially given the decentralized nature and limited processing power of many IoT systems. To address this challenge, we propose an innovative trust-based security framework inspired by biological defense systems. Drawing from the adaptive behavior of the human immune system, our model monitors device communications in real-time, identifies abnormal patterns, and isolates potentially compromised nodes similar to how immune cells identify and respond to harmful intruders. Trust scores are dynamically computed using key indicators such as behavioral consistency, message integrity, unusual latency, and data authenticity. Additionally, reinforcement learning is embedded into the system, enabling it to learn from past threats and continuously improve its detection strategies, much like how neural systems adapt over time. Both theoretical analysis and simulation results confirm the effectiveness of this framework in identifying and mitigating MitM attacks. By combining biological inspiration with intelligent learning, the proposed system offers a resilient and scalable solution to secure IoT networks against evolving cyber threats.

Index Terms – Internet of Things (IoT), Man-in-the-Middle (MitM) Attack, Biologically Inspired Security, Trust Evaluation, Anomaly Detection, Reinforcement Learning.

1. INTRODUCTION

The Internet of Things (IoT) has revolutionized the way devices communicate, connecting billions of smart devices such as wearable health trackers, home automation gadgets, industrial controllers, and intelligent transportation systems.

This extensive interconnection enables seamless data exchange and supports real-time decision-making across vital sectors including healthcare, agriculture, urban development, manufacturing, and logistics [1]. However, as IoT networks grow in scale and complexity, they face significant security challenges that can compromise their reliability and operational efficiency. Most IoT devices operate under strict resource limitations, with restricted processing capabilities, memory, and energy reserves. These constraints make it difficult to deploy traditional security measures like advanced encryption or sophisticated intrusion detection systems [2]. The problem is further compounded by the lack of uniform security standards across different devices and manufacturers, leaving networks vulnerable to advanced cyber-attacks. Among the most concerning threats are Man-in-the-Middle (MitM) attacks, where attackers secretly intercept communication between trusted devices. Such intrusions allow attackers to manipulate, eavesdrop, or impersonate devices, often without being detected [3].

The impact of MitM attacks can be devastating. In healthcare, altered sensor readings could lead to incorrect diagnoses or inappropriate treatments. In industrial or smart infrastructure systems, tampered commands may cause equipment failures, safety hazards, or production downtime. Conventional security mechanisms, such as static encryption keys, SSL/TLS protocols, or centralized authentication, often fall short in IoT contexts due to their high computational demands, increased latency, and inability to adapt to dynamic network conditions [4]. To tackle these challenges, this study introduces BioTrust-IoT, a biologically inspired adaptive trust system specifically designed to detect MitM attacks in IoT networks. Inspired by the human immune system, BioTrust-IoT continuously monitors device behavior to detect and neutralize malicious activities in real time. The system

RESEARCH ARTICLE

evaluates multiple behavioral indicators including communication consistency, message authenticity, abnormal packet delays, and repeated authentication failures to calculate dynamic trust scores for each device [5]. Devices exhibiting suspicious behavior can be flagged or temporarily isolated, preventing potential threats from propagating across the network.

BioTrust-IoT also integrates reinforcement learning, allowing it to learn from past attack patterns and adaptively improve its detection capabilities. This combination of biologically inspired trust mechanisms and machine learning results in a lightweight, decentralized, and adaptive security framework that enhances the resilience, scalability, and reliability of IoT networks against evolving MitM threats [6]. By focusing on trust-based, real-time detection, BioTrust-IoT fills a critical gap in IoT security and offers a practical solution for securing resource-constrained and heterogeneous IoT environments.

1.1. Problem Statement

The rapid growth of the Internet of Things (IoT) has resulted in the deployment of vast networks of interconnected devices across essential sectors, including healthcare, industrial automation, smart cities, and logistics. While this connectivity facilitates real time data exchange and intelligent decision making, it also introduces significant security challenges. The heterogeneous nature of IoT networks, coupled with limited device resources and the absence of standardized security protocols, makes them particularly susceptible to advanced cyber-attacks. Among these, Man-in-the-Middle (MitM) attacks are especially concerning, as they allow attackers to secretly intercept, alter, or impersonate communications between trusted devices.

Traditional security approaches, such as static encryption, SSL/TLS certificates, and centralized authentication, are often insufficient for IoT environments due to high computational demands, increased latency, and limited adaptability to dynamic network conditions. Likewise, conventional intrusion detection and trust-based methods either place excessive overhead on resource constrained devices or fail to deliver real time, adaptive protection against evolving threats. As a result, IoT networks remain vulnerable to MitM attacks, jeopardizing data integrity, privacy, and the safe operation of critical applications. To overcome these challenges, there is a need for a lightweight, adaptive, and decentralized security framework capable of continuously assessing device trustworthiness, detecting anomalies in real time, and responding effectively to emerging attack patterns, all while minimizing the burden on constrained IoT devices.

1.2. Objectives of the Proposed Work

- Develop a biologically inspired adaptive trust framework for real-time detection and mitigation of Man-in-the-Middle (MitM) attacks in IoT networks.

- Design a dynamic trust evaluation mechanism using behavioral indicators and reinforcement learning to improve detection accuracy and adaptability.
- Implement a lightweight, decentralized security solution suitable for resource constrained IoT devices.
- Evaluate the framework's effectiveness in enhancing detection performance, network resilience, and scalability while minimizing false positives.

1.3. Contribution of the Proposed Work

This research presents BioTrust-IoT, a biologically inspired adaptive trust system for detecting Man-in-the-Middle (MitM) attacks in IoT networks, with the following key contributions:

- Introduces a novel trust-based framework inspired by the human immune system for real-time MitM attack detection.
- Develops a dynamic trust evaluation mechanism using behavioral indicators to accurately assess device reliability.
- Integrates reinforcement learning to enhance detection accuracy and adapt to evolving attack patterns.
- Provides a lightweight, decentralized solution tailored for resource constrained IoT devices.
- Strengthens network resilience and scalability by isolating compromised devices while maintaining efficient operation.

The paper is organized as follows: Section 2 reviews the related work. Section 3 presents the proposed methodology. Section 4 provides the mathematical formulation and proof of the proposed model. Section 5 discusses the results and analysis. Finally, Section 6 concludes the study and highlights future research directions.

2. RELATED WORK

Securing Internet of Things (IoT) networks particularly against advanced threats such as Man-in-the-Middle (MitM) attacks has become a central research priority over the last two decades. IoT ecosystems are characterized by heterogeneous devices, dynamic topologies, and stringent resource limitations, all of which complicate the design of effective defense strategies. As a result, researchers have explored a wide range of techniques, including traditional cryptographic protocols, trust and reputation systems, bio-inspired models, and machine learning or reinforcement learning (RL) frameworks. Early studies concentrated on classical encryption techniques to safeguard IoT communications. The authors [7] examined the use of established cryptographic standards such as SSL/TLS and Public Key Infrastructure (PKI) within distributed IoT

RESEARCH ARTICLE

settings. These protocols deliver strong confidentiality, integrity, and authentication guarantees, effectively reducing the risk of eavesdropping or impersonation. Nevertheless, their reliance on computationally expensive encryption operations and complex key management makes them unsuitable for real-time MitM detection in resource constrained IoT nodes. To alleviate these overheads, the authors [8] introduced lightweight alternatives such as Elliptic Curve Cryptography (ECC) and symmetric key schemes. These methods improve energy efficiency compared to SSL/TLS but still struggle to adapt to rapidly evolving attack patterns and may introduce latency in large scale deployments.

To overcome the limitations of heavy encryption, several researchers turned to trust management frameworks. The authors [9] proposed a behavioral trust model that continuously evaluates the reliability of nodes based on interaction histories. The decentralized nature of this system lowers computational demands and suits heterogeneous IoT environments. However, the use of fixed trust thresholds limits its ability to counter stealthy or dynamic MitM attacks. The authors [10] extended this idea by incorporating a reputation-based mechanism, where nodes exchange recommendations to assess reliability. Although this method supports diverse IoT networks, its detection latency allows skilled adversaries to exploit vulnerabilities before defensive actions can be applied.

Nature-inspired defense strategies have also been investigated for IoT security. The authors [11] applied Artificial Immune Systems (AIS) with negative selection algorithms to identify anomalous behavior. Similar to a biological immune response, AIS can detect previously unknown attacks and adapt to changing environments. Despite these strengths, AIS requires significant computation and faces scalability challenges in low-power IoT networks. Building on this foundation, the authors [12] incorporated immune memory into AIS, enabling recognition of evolving threats over time. While this enhancement improves adaptability, it introduces additional memory and processing demands, making it unsuitable for lightweight, real time IoT operations.

With the rise of intelligent networking, machine learning (ML) and reinforcement learning (RL) have emerged as powerful tools for adaptive IoT security. The authors [13] developed an RL driven routing mechanism for mobile ad hoc networks that dynamically adjusts routing strategies based on observed threats.

The system exhibits excellent adaptability but incurs high energy costs and poses integration challenges for IoT platforms. The authors [14] advanced this direction with an RL-powered honeypot that employs Markov Decision Processes (MDPs) to model attacker behavior and anticipate potential intrusions. Although effective for tracking MitM

activities, the method remains resource-intensive and less suitable for real time, low power applications.

More recent efforts attempt to blend cryptographic, ML, and biologically inspired strategies to balance security strength and resource efficiency. The authors [15] introduced a polymorphic key mutation scheme, where cryptographic credentials are periodically changed to mimic immune system adaptability. This mechanism enhances credential security but is not explicitly targeted at MitM attacks and faces scalability issues in large deployments. The authors [16] combined RL, MTD, and Multipath TCP in a fog-computing architecture to detect MitM threats. The hybrid design improves detection accuracy and supports distributed IoT environments but introduces considerable processing overhead that can drain device energy. Block chain has also emerged as a promising decentralized security layer. The authors [17] pioneered a block chain-based trust framework for authentication, while the authors [18] refined the approach to achieve adaptive, tamper resistant IoT networks. Despite their robustness, block chain methods typically demand high computation and storage, which limits applicability in energy constrained devices. The authors [19] presented an in-depth study on multichannel Man-in-the-Middle (MitM) attacks, focusing on how these threats are becoming more sophisticated across various communication layers. Their research provides a clear understanding of the evolving attack patterns but stops short of offering a concrete strategy or framework for mitigation. On the other hand, the authors [20] explored an adaptive security model using Q-learning and Deep Q-Network (DQN) to detect and respond to threats in real time. Although their approach shows promise in improving threat adaptability, it suffers from high computational demands and practical integration difficulties, making it less suitable for deployment in resource limited IoT and edge computing environments.

Recent research trends emphasize the integration of machine bio-learning (ML), block chain, and inspired techniques to improve the resilience of Internet of Things (IoT) networks against advanced cyber threats such as Man-in-the-Middle (MitM) attacks. The authors [21] introduced a machine learning block chain architecture that unites the predictive capabilities of ML with the immutability of block chain ledgers. In this design, block chain is used to record network activities securely and verify trust relationships, while ML algorithms analyze traffic patterns to detect anomalies in real time. Experiments revealed notable improvements in attack detection and auditability compared to systems that rely solely on ML. Nevertheless, the combination of cryptographic block chain operations and real-time ML inference introduces significant computational and bandwidth demands, making deployment difficult on low power IoT devices or edge nodes with limited resources. The authors [22] proposed an adaptive Zero Trust Manager (ZTM) to enforce continuous authentication and strict access control across IoT networks.

RESEARCH ARTICLE

Following the “never trust, always verify” principle, the ZTM authenticates every device interaction even internal communication while dynamically adjusting verification frequency according to real-time risk levels. This strategy reduces insider threats and conserves energy through optimized authentication cycles. However, the framework struggles to scale when managing large, heterogeneous device populations and depends on block chain-based trust propagation, which introduces additional latency and storage overhead.

The authors [23] presented a distributed reinforcement learning (RL) framework aimed at decentralized, real-time intrusion detection. RL agents deployed across network nodes learn defense strategies by interacting with their environment, allowing the system to adapt autonomously to emerging attack patterns without a central controller. This approach enhances both network adaptability and resilience, but requires high processing capacity for training and frequent policy updates. Integration with diverse IoT protocols adds another layer of complexity, limiting large-scale deployment. The authors [24] combined supervised machine learning with bio-inspired mechanisms to design a hybrid intrusion detection system. Supervised models accurately classify known attack signatures, while bio-inspired components drawing on principles from the artificial immune system (AIS) provide the flexibility to detect previously unseen or zero-day attacks. Although this dual layer model delivers strong detection performance, its computational demands increase sharply in highly dynamic IoT settings, where device heterogeneity and fluctuating traffic patterns complicate training and scalability.

The authors [25] proposed a deep learning and bio-inspired hybrid model focused specifically on securing data transmission and countering MitM threats. Deep learning modules capture complex traffic features, while bio inspired

algorithms strengthen anomaly detection and improve resilience to unknown attacks. This architecture achieves higher detection reliability than traditional ML systems, but remains resource-intensive, requiring substantial memory and processing power, which limits its suitability for lightweight, real time IoT applications. The authors [26] explored the application of Deep Reinforcement Learning (DRL) to dynamically manage IoT security policies. DRL agents continuously refine defense strategies by learning from ongoing network feedback, enabling an effective trade-off between detection accuracy and energy consumption. Simulation results highlight the method’s adaptability and policy optimization capabilities. Despite these advantages, DRL entails complex training procedures, careful reward design, and the need for high-performance computing hardware, making it difficult to deploy on energy-constrained IoT devices or in latency-sensitive environments. The authors [27] proposed a hybrid trust evaluation model that integrates static metrics and behavioral analysis to assess IoT node reliability. The approach enhances accuracy and reduces false positives compared to single-factor models. However, it involves complex weight tuning, limited scalability, and added monitoring overhead. The authors [28] introduced a behavioral trust-based reintegration mechanism for isolated IoT nodes, enabling re-entry through updated trust scores and probation monitoring. The method improves network resilience and adaptability but remains vulnerable to on–off attacks and depends heavily on accurate behavioral data. The authors [29] developed a context-aware zero-trust hybrid framework for IoT-enabled self-driving vehicles, combining continuous authentication and contextual anomaly detection. The system strengthens security and adapts to dynamic vehicular conditions but faces challenges such as high complexity, latency issues, and intensive context data requirements. The table 1 depicts the summarization of the related work.

Table 1 Summarization Table of the Related Work

Reference	Methodology / Algorithm	Pros	Cons
[7]	SSL/TLS, PKI	Standard protocols; strong security	High computational cost; not real-time adaptive
[8]	ECC, Symmetric Key	Lightweight alternatives	Still resource intensive; limited MitM adaptability
[9]	Trust-based evaluation	Energy-efficient; decentralized	Static thresholds; slow adaptation
[10]	Reputation-based trust	Supports heterogeneous IoT	Limited real time response
[11]	AIS: Negative selection	Adaptive anomaly detection	High computational demand; limited scalability
[12]	AIS: Immune memory	Detects unknown attacks	Generic security; resource-heavy
[13]	RL-based adaptive routing	Learning-based adaptation	High energy cost; integration complexity

RESEARCH ARTICLE

[14]	RL honeypot, MDP	Models attacker behavior	Resource-intensive; real-time issues
[15]	Polymorphic key mutation	Adaptive credentialing	Not MitM focused; limited scalability
[16]	RL + MTD + Multipath TCP	Enhanced detection	High overhead; complex
[17]	Block chain-based trust	Tamper resistant, decentralized	High computation & storage cost
[18]	Block chain + IoT trust	Secure & adaptive	Not lightweight; resource-heavy
[19]	Multi-channel MitM analysis	Highlights evolving threats	Descriptive only; no mitigation framework
[20]	Q-learning, DQN	Adaptive threat learning	High computational cost; integration issues
[21]	Machine learning and block chain	Enhanced detection and prevention	Implementation complexity; resource requirements
[22]	Adaptive Zero Trust Manager	Real time authentication; energy-efficient	Limited scalability; dependency on block chain
[23]	Distributed RL	Enhanced adaptability; decentralized	High computational demand; integration challenges
[24]	Supervised ML and bio-inspired algorithms	Novel intrusion detection approach	Limited real-time response; scalability issues
[25]	Deep learning and bio-inspired techniques	Enhanced data transmission security	High computational cost; real-time adaptation challenges
[26]	Deep RL for security	Dynamic optimization of security policies	Resource-intensive; integration complexity
[27]	Hybrid trust model using static metrics + behavior analysis	Higher accuracy; reduces false positives; improves resilience	Complex weighting; limited scalability; added overhead
[28]	Behavioral trust-based reintegration for isolated nodes	Enables node recovery; improves network resilience	Vulnerable to on-off attacks; data accuracy dependent
[29]	Context-aware zero trust hybrid framework for vehicular IoT	Continuous verification; context adaptability; layered defense	High complexity; latency issues; heavy context management

2.1. Research Gap

Despite these encouraging developments, several key challenges remain. Many cryptographic and hybrid defense mechanisms demand significant computational resources, limiting their practicality on energy constrained IoT devices. Traditional trust-based models, which often rely on rigid thresholds and fixed rules, lack the flexibility needed to counter stealthy threats like MitM attacks that initially exhibit normal behavior before turning malicious. While RL and biologically inspired methods have independently demonstrated strong potential, their integrated application for MitM detection in IoT networks remains largely underexplored. Scalability also presents a persistent concern.

Most biologically inspired systems struggle to manage the complexity of large-scale, heterogeneous IoT deployments, particularly when trust decisions must be computed in a decentralized manner. Furthermore, existing trust evaluation models are often narrow in scope, focusing on a limited set of parameters. To enhance detection accuracy, there is a growing need for more comprehensive frameworks that incorporate multiple trust metrics. Finally, although RL offers the advantage of adaptive learning, integrating it with biologically inspired self-healing mechanisms [30] remains a promising yet underdeveloped direction. Such integration could enable continuous learning and proactive mitigation strategies, offering more robust protection against evolving MitM threats in IoT environments.

RESEARCH ARTICLE

3. PROPOSED WORK: BIOTRUST-IOT

3.1. Network Assumption and Initial Conditions

The proposed model is based on key assumptions and network conditions to develop a biologically inspired trust framework for detecting Man-in-the-Middle (MitM) attacks in IoT networks. The network is represented as a directed graph $G(V, E)$, where ' V ' is the set of autonomous nodes $N = \{n1, n2... nk\}$ and ' E ' represents wireless links. Each node maintains a dynamic trust score $T_i(t) \in [0, 1]$, calculated from weighted behavioral metrics protocol compliance, data integrity, latency, and message authenticity using adaptive coefficients $\alpha, \beta, \gamma, \delta, \epsilon$. Nodes with trust scores below a threshold ' θ ' are considered untrustworthy and isolated. Designed to detect stealthy MitM attacks where attackers behave normally before malicious activity, each node employs a Reinforcement Learning (RL) [31][32] agent to continuously adjust trust parameters and classification thresholds based on feedback, improving detection accuracy.

The model includes a recovery mechanism [33][34] that allows nodes to regain trust after sustained good behavior, reflecting immune system adaptability. Operating as a closed-loop, self-regulating system, the decentralized framework lets nodes or clusters perform local evaluations and share updated policies, enabling real-time adaptation and resilience against evolving threats, mirroring biological immune feedback mechanisms. The figure1 depicts the block diagram of the proposed model.

The proposed model consists of the following phases;

Phase 1: Initialization

Phase 2: Continuous Monitoring (Immune Surveillance)

Phase 3: Anomaly Detection (Antigen Detection)

Phase 4: Dynamic Trust Computation (T-cell Response)

Phase 5: Node Classification and Isolation (Immune Cytotoxic Action)

Phase 6: Reinforcement Learning-Based Trust Updating (Immune Memory & Neural

Plasticity)

Phase 7: Feedback Loop and Adaptation (Analogous to Adaptive Immune Regulation and

Systemic)

Phase 8: Continuous Operation

3.2. Phase 1: Initialization

At the outset of the proposed trust-based security framework for IoT networks, each node n_i is assigned an initial trust value $T_i(0)$, typically set to a neutral baseline $T_0=0.5$. This

neutral starting point avoids premature judgments about a node's trustworthiness before enough behavioral data has been observed. The trust evaluation mechanism then applies a multi-metric model that considers several behavioral indicators: behavioral stability ($S_{Behaviour}$), data integrity ($S_{Integrity}$), latency anomalies ($S_{Latency}$), and message authenticity ($S_{Authenticity}$). Each of these metrics is normalized to the range $[0, 1]$, where values closer to '1' represent more reliable behavior. To compute the overall trust score, a weighted aggregation is performed using coefficients $\alpha, \beta, \gamma, \delta, \epsilon$, which correspond to the relative importance of prior trust, behavioral stability, data integrity, latency anomalies, and message authenticity, respectively.

These weights are constrained such that $\alpha + \beta + \gamma + \delta + \epsilon = 1$, ensuring the final trust score remains within the normalized range. The choice of weight values is context-dependent. For example, in healthcare-related IoT systems, data integrity (γ) and message authenticity (ϵ) may be given higher priority due to their critical importance, whereas in latency-sensitive applications, greater emphasis might be placed on δ , which captures delay-related anomalies.

This initial trust assignment and adaptive evaluation process provide the foundation for dynamic monitoring and early anomaly detection, which are especially important in identifying stealthy MitM attacks that evolve over time.

3.3. Phase 2: Continuous Monitoring (Immune Surveillance)

The second phase of the proposed framework emphasizes real-time monitoring of IoT communications, inspired by immune surveillance in biological systems [35]. It ensures continuous observation of device interactions to detect anomalies, particularly those linked to Man-in-the-Middle (MitM) attacks. During each communication between source node ' i ' and destination node ' j ', the system passively captures metadata such as packet payloads (checked for unauthorized modifications), timestamps (used to detect replay attacks), transmission delays (to identify traffic manipulation), and authentication outcomes (to catch forgery or impersonation via cryptographic checks).

This metadata is continuously compared with historical baselines and learned behavioral profiles. Deviations like increased authentication failures, abnormal delays, or unexpected packet changes are flagged for further analysis. The low-latency surveillance mechanism enables early detection of MitM threats by recognizing subtle behavioral shifts. It also enhances network-wide situational awareness and dynamically updates trust scores.

Temporal analysis modules further strengthen the system by detecting slow-acting or stealthy attacks through pattern correlation over time.

RESEARCH ARTICLE

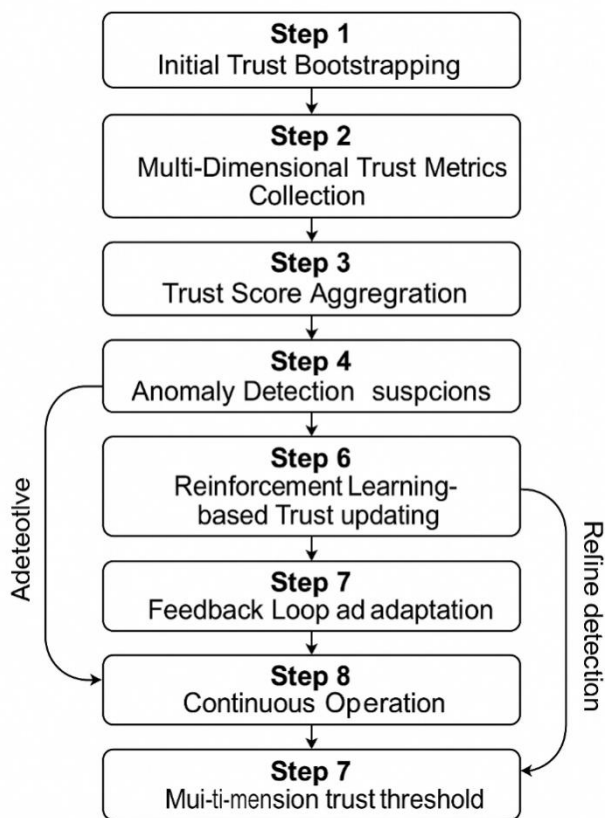


Figure 1 Block Diagram of Proposed Work

3.4. Phase 3: Anomaly Detection (Antigen Detection)

In this phase, the system emulates the antigen detection process of biological immune systems [35] by identifying irregularities in communication behavior that may signal the presence of a Man-in-the-Middle (MitM) attack. By continuously monitoring the interaction patterns between IoT nodes, the framework aims to detect subtle anomalies introduced by adversaries who intercept, alter, or inject packets into the communication stream.

These disruptions often manifest as deviations across key trust indicators. To identify such threats effectively, the system evaluates the following behavioral metrics:

3.4.1. Data Integrity

To ensure data integrity, the system compares the hash or checksum of the received data packet d_{recv} with that of the expected packet d_{exp} .

This comparison is quantified by calculating the normalized Hamming distance between the two hash values, providing a measure of similarity that helps detect any unauthorized modifications. The equation (1) is used to calculate the integrity.

$$S_{Integrity,i} = \frac{HammingDistance(Hash(d_{recv}), Hash(d_{exp}))}{HashLength} \quad (1)$$

A score near '1' signifies that the data is intact, while lower scores indicate potential tampering or modification, which could be a sign of a Man-in-the-Middle attack.

3.4.2. Latency Irregularities

Man-in-the-Middle attackers often cause extra delays in packet transmission because of interception and forwarding. To identify this, the system measures the latency deviation for node i as follows and it is calculated based on equation (2).

$$\Delta L_i = |L_{current,i} - \mu L_i| \quad (2)$$

In this context, $L_{current,i}$ represents the observed latency for the current packet transmission from node ' i ', while μL_i denotes the historical average latency for the same node. The latency anomaly score is then calculated using an exponential decay function which is mentioned in the equation (3) as follows:

$$S_{latency,i} = e^{-\kappa \Delta L_i} \quad (3)$$

Where $\kappa > 0$ is a sensitivity constant. Larger deviations reduce the score sharply, signaling possible interception delays.

3.4.3. Message Authenticity

The system evaluates the validity of each received message by checking digital signatures or authentication tokens. The result is binary and it is mentioned in the following equation (4).

$$S_{authentication,i} = \begin{cases} 1, & \text{if Authentication Succeed} \\ 0, & \text{if Authentication Fails} \end{cases} \quad (4)$$

A failed authentication, represented by a score of zero, is a strong sign that a message may have been spoofed or tampered with common characteristics of Man-in-the-Middle attacks.

3.4.4. Behavioral Stability

This metric measures the consistency of a node i 's behavior over time by analyzing observable features such as transmission frequency, packet size, encryption usage, and communication partners. Let $\vec{A}^i_{curr}$ and $\vec{A}^i_{prev}$ denote the current and historical behavioral feature vectors, respectively. The similarity between these vectors is quantified using cosine similarity, calculated as in equation (5):

$$S_{behaviour,i} = \frac{\vec{A}^i_{curr} \cdot \vec{A}^i_{prev}}{||\vec{A}^i_{curr}|| ||\vec{A}^i_{prev}||} \quad (5)$$

A cosine similarity value close to '1' reflects stable and consistent behavior, whereas lower values indicate behavioral changes that could be caused by adversarial influence. Taken together, these metrics provide a comprehensive means to detect anomalous behavior in nodes that may signal Man-in-

RESEARCH ARTICLE

the-Middle attacks. Nodes exhibiting poor scores in one or more of these areas are flagged for additional scrutiny or may be isolated to prevent further compromise within the network.

3.4.5. Dynamic Trust Computation (T-cell Response)

This phase simulates the adaptive T-cell response found in biological immune systems. The framework updates each node's trust score by combining past trust values with real-time behavioral data. This balanced approach helps smooth out temporary anomalies while emphasizing consistent long-term patterns, enhancing both robustness and adaptability. Specifically, for a node 'i' at time 't', the trust score $T_i(t)$ is updated using the following weighted aggregation equation (6):

$$T_i(t) = \alpha \cdot T_i(t-1) + \beta \cdot S_{behaviour,i} + \gamma \cdot S_{integrity,i} + \delta \cdot S_{latency,i} + \epsilon \cdot S_{authenticity,i} \quad (6)$$

where:

- $T_i(t-1)$ is the trust score at the previous time step.
- $S_{behaviour,i}, \gamma \cdot S_{integrity,i}, \delta \cdot S_{latency,i}, \epsilon \cdot S_{authenticity,i}$ are the real-time anomaly scores for behavior, data integrity, latency, and message authenticity, respectively.
- $\alpha, \beta, \gamma, \delta, \epsilon \in [0,1]$ are weight coefficients assigned to each component such that: $\alpha + \beta + \gamma + \delta + \epsilon = 1$
- ' α ' governs the influence of historical trust (memory retention), while the remaining weights emphasize current behavioral metrics.

The updated trust score for each node is derived from the combined evaluation of various behavioral metrics, embodying an adaptive immune-like response. When a node consistently behaves correctly showing authentic communication without data tampering or unusual delay its trust score is increased.

If anomalies such as failed authentications or data corruption are detected, the relevant metric lowers the trust score, resulting in an overall decrease. Over time, nodes that maintain trustworthy behavior sustain or improve their trust, while compromised nodes experience a gradual decline in trustworthiness.

3.5. Phase 5: Node Classification and Isolation (Immune Cytotoxic Action)

In biological immune systems, cytotoxic T-cells remove infected or malfunctioning cells to prevent the spread of infection. Drawing on this principle, the proposed biologically inspired trust framework for IoT isolates any node whose trust score drops below a specified threshold, labeling it as compromised. This isolation acts as a critical containment strategy to help mitigate threats like Man-in-the-Middle (MitM) attacks within the network.

3.5.1. Trust-Based Classification Rule

Each node 'i' is continuously observed and given a dynamic trust score $T_i(t)$. The node's status is then determined according to the following classification rule:

If $T_i(t) < \theta$,

then node 'i' is classified as suspicious or malicious.

Where:

- $T_i(t) \in [0, 1]$ is the trust score of node 'i' at time 't'.
- $\theta \in (0, 1)$ is the trust threshold that determines when a node should be flagged for isolation.

The threshold θ plays a key role in balancing detection sensitivity and false alarm rates. Setting a lower threshold reduces sensitivity to attacks, resulting in fewer false positives but a higher chance of missing actual threats (false negatives). Conversely, a higher threshold makes the system more sensitive, which helps catch more attacks but may also lead to legitimate nodes being wrongly flagged as suspicious (increased false positives).

3.5.2. Dynamic Threshold Tuning

Instead of relying on a fixed threshold θ , the system can adaptively adjust it based on the statistical characteristics of the trust score distribution which is in the equation (7).

$$\theta(t) = \mu T(t) - \lambda \cdot \sigma T(t) \quad (7)$$

Where:

- $\mu T(t)$ is the mean trust score across all nodes at time t,
- $\sigma T(t)$ is the standard deviation,
- λ is a confidence multiplier

Adaptive thresholding enables the system to accommodate contextual factors such as network congestion or firmware updates that may temporarily affect multiple nodes and cause anomalous behavior.

3.5.3. Isolation and Containment Strategy

When a node is identified as malicious, the system enforces a series of countermeasures: it is immediately quarantined by restricting or revoking its communication capabilities such as blocking it at the routing layer or isolating it in a monitored environment; neighboring nodes are alerted through broadcast messages or routing table updates to avoid interacting with the compromised node; the node's reduced trust level may be shared across the network via a controlled trust dissemination protocol to prevent the spread of misinformation; and finally, a temporary penalty is applied, after which the node may be reassessed manually or automatically and allowed to rejoin the network if its trust metrics improve.

RESEARCH ARTICLE

3.5.4. Mitigation of MitM Implications

In the case of Man-in-the-Middle (MitM) attacks:

- A compromised node that acts as a relay or spoofs other devices typically fails authentication checks, causes increased latency, or alters data, all of which contribute to a decreased trust score.
- Once such a node is isolated, it can no longer impersonate legitimate devices or relay malicious traffic, effectively containing and preventing the further spread of the attack.

3.5.5. Optional Reinforcement Feedback (Immune Memory)

To enhance robustness, the system may implement a memory mechanism where:

- The frequency of isolation incidents is tracked per node.
- Nodes repeatedly flagged below the threshold may be permanently blacklisted as per in the following equation (8).

$$B_i = \sum_{t=1}^T 1_{T_i(t) < \theta} \quad (8)$$

Where $1_{T_i(t) < \theta}$ is an indicator function equal to '1' when $T_i(t) < \theta$.

If ' B_i ' exceeds a maximum threshold B_{max} , the node is permanently blacklisted.

3.6. Phase 6: Reinforcement Learning-Based Trust Updating (Immune Memory & Neural Plasticity)

In biological systems, adaptive immunity and neural plasticity enable organisms to learn from previous encounters by storing immune memory and adjusting synaptic connections. This capacity allows for faster and more effective responses to familiar pathogens and environmental changes. Drawing on this principle, the proposed IoT trust framework incorporates Reinforcement Learning (RL) to continuously refine trust evaluation strategies based on the outcomes of past decisions. This adaptive learning empowers the system to respond effectively to evolving threats, such as stealthy Man-in-the-Middle (MitM) attacks, as well as to varying environmental conditions. To formalize this trust adaptation, the system models the update process as a Markov Decision Process (MDP), defined by the tuple (S, A, R, T, π) .

Here, ' S ' represents the set of possible states, ' A ' is the set of available actions, ' R ' is the reward function guiding learning, ' T ' is the state transition function (which may be omitted in model-free approaches), and ' π ' is the policy that determines the action selection based on the current state. The state at time t represents the current trust and security context in the IoT network and is defined as in equation (9).

$$S_x = [T_i(t), S_{behaviour,i}, S_{integrity,i}, S_{latency,i}, S_{authenticity,i}, \theta(t)] \quad (9)$$

Where, $T_i(t)$ is the current trust score of node ' i ', $S_{behaviour,i}$, $S_{integrity,i}$, $S_{latency,i}$, $S_{authenticity,i}$ represent the real-time trust metric scores, $\theta(t)$ is the current threshold for classifying a node as trustworthy or malicious. The action at time ' t ', denoted as ' a_t ', corresponds to adjustments in trust model parameters mentioned in the equation (10).

$$a_t = \{\Delta\alpha, \Delta\beta, \Delta\gamma, \Delta\delta, \Delta\epsilon, \Delta\theta\} \quad (10)$$

The possible actions taken by the system include:

- Increasing or decreasing the weight β assigned to behavioral metrics,
- Adjusting the classification threshold θ up or down,
- Rebalancing other weights such as α , γ , δ , ϵ to better detect emerging types of anomalies.

These adjustments may be implemented as discrete, stepwise changes or as continuous updates such as gradient-based tuning depending on the reinforcement learning method employed. Following each action, the system receives feedback reflecting the accuracy of its classification, which is quantified through a reward function. The following equation (11) is used to calculate the reward function.

$$r_t = \begin{cases} +1, & \text{if a malicious node is correctly isolated (true positive)} \\ 1, & \text{if a legitimate node is falsely isolated (false positive)} \\ 2, & \text{if a malicious node is missed (false negative)} \\ 0, & \text{otherwise (true negative)} \end{cases} \quad (11)$$

This reward system helps the agent learn strategies that improve security without excessively penalizing benign nodes. The framework uses Q-learning, a model-free, value-based reinforcement learning algorithm, to learn the optimal policy $\pi(s)$ that selects the best action in each state. The Q-value update rule is defined as equation (12):

$$Q(S_t, a_t) + \gamma[r_t + \gamma \cdot \max_a Q(S_{t+1}, a) - Q(S_t, a_t)] \quad (12)$$

Where, $\eta \in [0, 1]$ is the learning rate, which determines how quickly the model updates its estimates, $\gamma \in [0, 1]$ is the discount factor, balancing the importance of immediate vs. future rewards, $Q(s, a)$ represents the expected value (quality) of performing action ' a ' in state ' s '. The agent selects actions that maximize the expected future reward using equation (13):

$$\pi(s) = \arg \max_a Q(s, a) \quad (13)$$

3.7. Phase 7: Feedback Loop and Adaptation (Analogous to Adaptive Immune Regulation and Systemic Feedback in Biological Systems)

In the proposed biologically inspired framework, the reinforcement learning (RL) agent is designed to continuously evolve rather than remain static. This ongoing evolution is

RESEARCH ARTICLE

facilitated by a feedback loop that enables the system to dynamically adjust to new attack methods, environmental changes, and emerging anomalies in network behavior. This process mirrors how the adaptive immune system regulates its response based on systemic feedback, optimizing defense while minimizing unintended harm.

3.7.1. Real-Time Feedback Integration

Following each classification decision and trust score update, the RL agent integrates performance feedback from multiple sources. These include classification accuracy metrics (such as true positives and false positives), environmental factors like node density and channel noise, temporal variations such as time-of-day traffic shifts, and changes in network topology caused by node mobility or link disruptions. This rich feedback reinforces decisions that successfully detect attacks, penalizes incorrect classifications or unwarranted isolation of nodes, and supports the identification of behavioral drift or subtle, evolving attack patterns.

3.7.2. MitM Pattern Recognition

Over time, the RL agent learns to identify complex Man-in-the-Middle (MitM) attack behaviors by correlating subtle signals that individually may seem harmless but collectively indicate malicious activity. For example, an attacker might begin as a passive observer without altering data, then progressively introduce slight packet modifications. The system tracks such evolving behaviors by comparing long-term historical metrics with current anomalies, employing techniques like time-series feature extraction or sliding window analyses to detect these gradual changes.

3.7.3. Adaptive Sensitivity Adjustment

To maintain optimal detection performance, the system autonomously fine-tunes several key parameters, including anomaly score thresholds for flagging malicious behavior, the trust threshold ' θ ' for isolating nodes, and the weights α , β , γ , δ , ϵ used in the trust aggregation formula. These adaptive adjustments help minimize false positives by preventing overreactions to temporary benign anomalies, reduce false negatives by enhancing the detection of subtle, low-and-slow Man-in-the-Middle (MitM) tactics, and avoid unnecessary isolation of nodes impacted by transient instabilities such as mobility or lossy communication links.

Updates to these parameters occur in real time through reinforcement signals and are further refined by periodic batch training sessions triggered either at regular intervals or in response to significant shifts in classification performance.

3.7.4. Policy Refinement

The policy ' π ', which maps states to actions within the reinforcement learning framework, undergoes continuous refinement using advanced techniques such as reward shaping

and meta-reinforcement learning. This ongoing optimization improves both immediate classification accuracy and the model's ability to generalize across new or evolving attack patterns. Over time, the system becomes increasingly resilient to environmental noise and variability, more effective at detecting stealthy or polymorphic MitM attacks, and better aligned with overarching goals such as preserving network availability, accuracy, and trustworthiness.

3.8. Phase 8: Continuous Operation

(Analogous to Immune System Surveillance and Homeostasis)

To provide persistent protection, the trust framework operates as an autonomous, always-active process that continuously learns and adapts much like biological immune surveillance that maintains homeostasis.

3.8.1. Persistent Monitoring Cycle

The system perpetually executes the following operations:

- Observing real-time node behavior and gathering communication metadata.
- Calculating anomaly metrics related to data integrity, latency, authenticity, and behavioral consistency.
- Dynamically aggregating trust scores using adaptively evolving weights.
- Classifying nodes by comparing trust scores against the current threshold θ .
- Isolating or restricting nodes deemed malicious.
- Receiving reward-based feedback on classification outcomes.
- Updating the reinforcement learning model and associated parameters accordingly.

Each monitoring cycle is optimized for speed and efficiency, ensuring suitability for resource-constrained IoT environments with strict latency, computation, and bandwidth limits.

3.8.2. Auto-Healing & Recovery Support

If a previously isolated node resolves its issues such as through software updates or improved environmental conditions the system gradually restores its trust score $T_i(t)$ based on consistent benign behavior.

The reinforcement learning agent re-learns from these corrected classifications, allowing the node to be reintegrated into the network once it meets predefined recovery thresholds. This mechanism promotes fairness and prevents permanent penalties for transient faults.

RESEARCH ARTICLE

3.8.3. Scalability and Real-Time Operation

Scalability is achieved through decentralized trust computations, where individual nodes or cluster heads locally evaluate interactions, share learned policies or trust tables, and contribute to global attack detection frameworks such as federated learning. Real-time responsiveness is maintained by lightweight local computations, offloading resource-intensive tasks to edge or cloud resources as needed, and employing efficient data structures to manage reinforcement learning models and Q-tables.

3.8.4. Long-Term Learning

To enhance adaptability over extended periods, the system archives historical data and periodically retrains or fine-tunes its models. This process establishes a form of institutional "immune memory," capturing known MitM attack signatures, differentiating normal from abnormal behavior profiles, and identifying which trust aggregation strategies have proven effective or ineffective in past scenarios.

The following algorithm 1 demonstrates the proposed methodology.

Aim: Detect and isolate MitM attacks in an IoT network using trust metrics and reinforcement learning with feedback adaptation.

Inputs:

N IoT nodes ($i = 1$ to N)

Initial Trust Score: T_0

Default weights: $\alpha, \beta, \gamma, \delta, \epsilon$

RL Parameters: learning rate η , discount factor γ , exploration rate ϵ_{greedy}

Thresholds: trust threshold θ , blacklist limit B_{max}

Metrics per node: $S_{\text{behavior}}, S_{\text{integrity}}, S_{\text{latency}}, S_{\text{authenticity}}$

Q-Table: $Q(s, a)$

Output:

Classified node status: Trusted or Malicious

Updated Trust Scores: $T_i(t)$

Updated Q-values: $Q(s, a)$

MitM attack alerts & isolated node list

Step 1: Initialization

For each node i in the IoT network:

Initialize Trust Score $T_i \leftarrow T_0$ // e.g., $T_0 = 0.5$

Initialize Metrics $\alpha, \beta, \gamma, \delta, \epsilon \leftarrow$ default weights

Set Threshold $\theta \leftarrow 0.4$

Initialize Q-Table $Q(s, a) \leftarrow 0$

Set learning rate η , discount factor γ , exploration rate ϵ_{greedy}

Step 2: Behavior Monitoring

Loop at each time interval t :

For each node i :

Measure:

- Behavior score S_{behavior}

- Integrity score $S_{\text{integrity}}$

- Latency score S_{latency}

- Authenticity score $S_{\text{authenticity}}$

Record metrics for trust computation

Step 3: Trust Score Computation

For each node i :

Compute Trust Score:

$T_i(t) = \alpha * S_{\text{behavior}} + \beta * S_{\text{integrity}} + \gamma * S_{\text{latency}} + \delta * S_{\text{authenticity}} + \epsilon * \text{PreviousTrust}$

Step 4: Anomaly Detection and Scoring

For each node i :

If any trust component deviates from baseline significantly:

Label node i as suspicious

Store anomaly score $A_i(t)$ based on deviation magnitude

Step 5: Node Classification and Isolation

If $T_i(t) < \theta$:

Classify node i as malicious

Isolate node i :

- Block routing

- Inform neighboring nodes

- Broadcast reputation update

Update blacklist score $B_i \leftarrow B_i + 1$

If $B_i > B_{\text{max}}$:

Permanently ban node i

Step 6: Reinforcement Learning-Based Trust Updating

Observe current state $s_t \leftarrow [T_i(t), S_{\text{behavior}}, S_{\text{integrity}}, \dots, \theta(t)]$

Choose action a_t using ϵ -greedy:

RESEARCH ARTICLE

- Adjust α , β , γ , δ , ϵ or threshold θ

Apply a_t and re-evaluate trust

Classify nodes again $\rightarrow$ observe outcome

Assign reward r_t based on:

- True positive $\rightarrow +1$

- False positive $\rightarrow -1$

- False negative $\rightarrow -2$

- True negative $\rightarrow 0$

Update Q-Table:

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \eta * [r_t + \gamma * \max_a Q(s_{t+1}, a) - Q(s_t, a_t)]$$

Step 7: Feedback Loop and Adaptation

After each classification:

Analyze detection accuracy

Identify misclassifications and adjust reward signals

Adjust θ and metric weights based on feedback trends

Detect evolving MitM patterns using time-series correlation

Step 8: Continuous Operation and Recovery

Repeat all steps continuously:

If previously isolated node shows improved behavior:

Increment $T_i(t)$ gradually

Allow reintegration if $T_i(t) >$ recovery threshold

Periodically:

- Retrain RL model with historical data

- Share updated trust policies across nodes or clusters

Algorithm1 Biologically-Inspired Trust Framework for MitM Detection in IoT

4. MATHEMATICAL PROOF

The proposed Biologically-Inspired Trust-Based MitM Detection Framework has mathematically proven with following example.

Assume an IoT environment with 10 nodes (Node1 to Node10). The following assumption and initial condition have made.

4.1. Assumptions and Initial Setup

Before applying any trust evaluation, it's essential to define the foundational parameters such as the number of nodes, initial trust, trust threshold, and the weights assigned to behavioral features. This setup ensures a consistent basis for evaluating and comparing nodes over time. The following table 2 represents the assumptions and initial conditions.

Table 2 Assumptions and Initial Conditions

Parameter	Value
Nodes	10 (Node1 to Node10)
Initial Trust	$T_0=0.5, T_{-0} = 0.5, T_0=0.5$
Threshold	$\theta=0.4$
Weights	$\alpha=0.25, \beta=0.25, \gamma=0.2, \delta=0.2, \delta=0.2, \epsilon=0.1$

4.1.1. Phase 1: Trust Metric Input Values (Sample for 5 Nodes)

Among the 10 nodes, 5 nodes will take as sample input. The following table 3 demonstrates the calculated trust values from node 1 to node 5 which is discussed in phase 2 and these trust values are assumed.

Table 3 Calculated Trust Values

Node	$S_{Behaviour}$	$S_{Integrity}$	$S_{Latency}$	$S_{Authenticity}$	PrevTrust
Node1	0.90	0.95	0.85	0.90	0.50
Node2	0.60	0.55	0.40	0.50	0.50
Node3	0.20	0.30	0.80	0.20	0.50
Node4	0.85	0.88	0.90	0.90	0.50
Node5	0.40	0.45	0.70	0.40	0.50
Node6	0.30	0.35	0.60	0.30	0.50
Node7	0.75	0.80	0.85	0.78	0.50
Node8	0.65	0.70	0.80	0.70	0.50
Node9	0.80	0.85	0.90	0.85	0.50
Node10	0.70	0.75	0.80	0.75	0.50

RESEARCH ARTICLE

4.1.2. Phase 2: Trust Score Computation

Trust evaluation in IoT networks is based on multiple dynamic metrics, such as behavior, integrity, latency, and authenticity. The following table 4 denotes the trust score table. This step involves collecting the relevant observed values from each node for the current time window to compute its trustworthiness. Using the formula: $T(i) = \alpha \cdot S_b + \beta \cdot S_i + \gamma \cdot S_l + \delta \cdot S_a + \varepsilon \cdot \text{PrevTrust}$

Table 4 Trust Score Table

Node	Calculation	Trust Score T_i
Node1	$0.225 + 0.2375 + 0.17 + 0.18 + 0.05$	0.8625
Node2	$0.15 + 0.1375 + 0.08 + 0.10 + 0.05$	0.5175
Node3	$0.05 + 0.075 + 0.16 + 0.04 + 0.05$	0.375
Node4	$0.2125 + 0.22 + 0.18 + 0.18 + 0.05$	0.8425
Node5	$0.10 + 0.1125 + 0.14 + 0.08 + 0.05$	0.4825
Node6	$0.075 + 0.0875 + 0.12 + 0.06 + 0.05$	0.3925
Node7	$0.1875 + 0.2 + 0.17 + 0.156 + 0.05$	0.7635
Node8	$0.1625 + 0.175 + 0.16 + 0.14 + 0.05$	0.6875
Node9	$0.2 + 0.2125 + 0.18 + 0.17 + 0.05$	0.8125
Node10	$0.175 + 0.1875 + 0.16 + 0.15 + 0.05$	0.7225

4.1.3. Phase 3: Classification Decision (Threshold $\theta = 0.4$)

The core of the framework lies in aggregating multiple trust dimensions into a single score using a weighted linear combination. The following table 5 denotes the decision table. This score enables an objective comparison between nodes. Labels each node as Trusted or Malicious based on whether the trust score exceeds the defined threshold $\theta = 0.4$.

Table 5 Decision Table

Node	Trust Score	Verdict
Node1	0.8625	✔ Trusted
Node2	0.5175	✔ Trusted
Node3	0.375	✘ Malicious
Node4	0.8425	✔ Trusted
Node5	0.4825	✔ Trusted
Node6	0.3925	✘ Malicious
Node7	0.7635	✔ Trusted
Node8	0.6875	✔ Trusted
Node9	0.8125	✔ Trusted
Node10	0.7225	✔ Trusted

4.1.4. Phase 4: RL Reward Feedback Based on Ground Truth

Once trust scores are computed, this step uses a threshold mechanism to decide whether a node is Trusted or Malicious. Any node with a trust value below the threshold θ is flagged as suspicious. The table 6 denotes the ground truth table which represents malicious or not.

Table 6 Ground Truth Table

Node	Ground Truth	Verdict	Reward
Node3	Malicious	Isolated	+1
Node6	Malicious	Isolated	+1
Others	Legitimate	Trusted	0

4.1.5. Phase 5: Q-Learning Update (for Node3)

To simulate learning and adaptation, reinforcement signals are introduced based on the ground truth classification of each node. This helps the system learn from both correct and incorrect decisions.

Let: $Q=0.2, \eta=0.1, r=1, \gamma=0.9, \max(Q)=0.3$

Node3:

$$Q = 0.2 + 0.1(1 + 0.27 - 0.2)$$

$$= 0.307$$

Node6:

$$Q = 0.2 + 0.1(1 + 0.27 - 0.2)$$

$$= 0.307$$

Thus, after this learning phase, the updated Q-values for both Node3 and Node6 are 0.307, reflecting improved expected utility based on the recent experience.

4.1.6. Phase 6: Adaptive Threshold and Weight Adjustment (RL Agent Actions)

Q-learning allows the system to adjust its future decisions by updating Q-values based on the received reward and expected future values. This ensures continuous improvement in malicious detection accuracy.

Lower threshold to $\theta = 0.38$

Adjust weights:

$$\alpha = 0.2 \text{ (↓ behavior)}$$

$$\beta = 0.3 \text{ (↑ integrity)}$$

4.1.7. Phase 7 & 8: Continuous Operation

The system should be biologically inspired and self-healing, meaning it must adjust thresholds and weights based on historical detection success. If certain features (like integrity) are more useful, they get a higher weight; the threshold may

RESEARCH ARTICLE

also be fine-tuned. IoT networks are dynamic, and nodes may change behavior over time. Continuous monitoring enables re-evaluation of nodes based on updated behavior and reward feedback, allowing malicious nodes to be rehabilitated if

behavior improves. The following table 7 denotes the final table of Node 3 and Node 6. The table 8 denotes the final prediction.

Table 7 Final table of Node 3 and Node 6

Node	Timestep	S _b	S _i	S _l	S _a	PrevTrust	Trust Score	Verdict	Reward	Q-Value
Node3	T1	0.20	0.30	0.80	0.20	0.50	0.375	✗	+1	0.307
	T2	0.40	0.50	0.60	0.40	0.375	0.497	✓	+1	0.41
	T3	0.50	0.60	0.65	0.50	0.497	0.574	✓	+1	0.50
Node6	T1	0.30	0.35	0.60	0.30	0.50	0.3925	✗	+1	0.307
	T2	0.45	0.55	0.65	0.40	0.3925	0.545	✓	+1	0.41
	T3	0.55	0.60	0.70	0.45	0.545	0.605	-	-	-

Table 8 Final Prediction

Node	Final Trust (T3)	Verdict	Q-Value
Node1	~0.84	Trusted	
Node2	~0.44	Trusted	
Node3	0.574	Trusted	0.50
Node4	~0.88	Trusted	
Node5	~0.44	Trusted	
Node6	0.605	Trusted	0.50
Node7	~0.76	Trusted	
Node8	~0.68	Trusted	
Node9	~0.81	Trusted	
Node10	~0.72	Trusted	

5. RESULTS AND DISCUSSION

The developed trust-based security framework was implemented and tested within the Contiki Cooja simulation environment, a well-established tool for analyzing IoT protocol performance in constrained and realistic scenarios. The simulation environment which is mentioned in table 9 was carefully designed to emulate real-world IoT networks by integrating relevant elements such as randomized network layouts, standardized communication protocols, simulated attack behaviors, and dynamic trust evaluation mechanisms.

In the simulated network, 30 Z1 motes were randomly distributed within the virtual field and communicated using the IEEE 802.15.4 standard, with RPL employed as the routing protocol. To evaluate the system's resilience to security threats, approximately 16.7% of the nodes were intentionally configured as malicious. These nodes performed Man-in-the-Middle (MitM) attacks by intercepting and

tampering with packet transmissions, thereby degrading communication by increasing latency and compromising data integrity and authenticity. Each node in the network continuously observed the behavior of its neighboring nodes and computed trust scores at regular intervals. These scores were based on multiple criteria: behavior score (evaluated from packet forwarding success), integrity score (derived from hash comparison), latency score (based on measured communication delay), and authenticity score (verified through digital signatures). A weighted combination of these metrics produced a composite trust value, which informed secure routing decisions. To further enhance trust assessment, the framework incorporated reinforcement learning via Q-learning, utilizing a ϵ -greedy strategy. This learning model adaptively adjusted the weights of individual trust metrics and fine-tuned detection thresholds to improve the accuracy of identifying malicious activities while reducing the incidence of false positives. Key learning parameters such as the

RESEARCH ARTICLE

exploration rate, learning rate, and discount factor were optimized for stable and efficient convergence during simulation. The outcomes of the simulation highlighted the

viability and strength of the proposed trust-based approach in safeguarding IoT networks, even in the presence of adversarial threats.

Table 9 Simulation Parameters

Aspect	Description	Example Values
Network Topology and Devices	- 10 to 50 IoT nodes (Sky or Z1 motes), arranged randomly or in grid. - Communication via IEEE 802.15.4. - Routing protocol: RPL.	- Number of nodes: 30 - Device: Z1 motes - Topology: Random - Radio standard: IEEE 802.15.4 - Routing: RPL
Attack Model	- Some nodes act as Man-in-the-Middle (MitM) attackers. - Malicious nodes intercept and modify packets. - Affect latency, integrity, authenticity.	- Malicious nodes: 5 out of 30 ($\approx 16.7\%$) - Attack type: Packet interception and modification - Latency increase: +50 ms on attacked links
Trust Evaluation Metrics	- Behavior Score (S_{behavior}): message forwarding success. - Integrity Score ($S_{\text{integrity}}$): packet hash match. - Latency Score (S_{latency}): delay observed. - Authenticity Score ($S_{\text{authenticity}}$): digital signature verification. - Combined with weights to compute Trust Score $T_i(t)T_i(t)T_i(t)$.	- S_{behavior}: 0.0 to 1.0 (e.g., 0.85 average) - $S_{\text{integrity}}$: 0.0 to 1.0 (e.g., 0.9 average) - S_{latency}: 0 to 1 (normalized; e.g., 0.8 for low latency) - $S_{\text{authenticity}}$: 0.0 or 1.0 (valid signature = 1) - Weights example: $w_{\text{behavior}}=0.4$, $w_{\text{integrity}}=0.3$, $w_{\text{latency}}=0.2$, $w_{\text{authenticity}}=0.1$
Reinforcement Learning Integration	- Q-learning with ϵ-greedy policy. - Agent optimizes weights and detection thresholds. - States: trust scores and metrics. - Actions: adjust weights and thresholds. - Goal: maximize detection accuracy, minimize false positives.	- ϵ (exploration rate): starts 0.1, decays to 0.01 - Learning rate α: 0.05 - Discount factor γ: 0.9 - Initial thresholds: trust score threshold = 0.7

5.1. Packet Delivery Ratio

The proposed biologically inspired trust framework achieves a packet delivery ratio (PDR) of 95.3%, outperforming prior approaches such as [27] (92.8%), [28] (93.1%), and the recent context-aware zero trust-based hybrid model by Khan et al. [29] (94.2%).

As illustrated in Figure 2, which plots the number of MitM nodes against PDR, our approach demonstrates superior resilience in maintaining reliable data transmission. In [27], the trust evaluation process is predominantly static and reactive, leading to delayed identification of MitM nodes. This reliance on compromised paths results in higher packet drops and reduced delivery performance.

In [28], although biologically inspired, the approach lacks fine-grained reinforcement learning-based adaptation, slowing its responsiveness to dynamic malicious behavior and slightly degrading PDR. The recent work by Khan et al. [29] integrates a context-aware zero-trust model for IoT-based self-driving vehicle networks, which enhances security

through hybrid verification mechanisms. While this improves detection over traditional trust models, its context verification overhead and reliance on pre-defined trust rules slightly limit real-time adaptability under fast-evolving attack scenarios, thus marginally reducing its packet delivery performance compared to our proposed framework.

By contrast, the proposed model integrates continuous trust monitoring and reinforcement learning (RL) with real-time feedback, enabling rapid identification and isolation of malicious nodes. This proactive mechanism minimizes reliance on compromised paths, mitigates attack impact promptly, and significantly preserves route integrity. As a result, packet loss is reduced, ensuring more stable and reliable transmission across IoT networks.

5.2. Detection Accuracy

The proposed biologically inspired trust framework achieves a detection accuracy of 97.4%, which is higher than [27] (94.5%), [28] (95.2%), and the recent context-aware zero-trust hybrid model by Khan et al. [29] (96.1%). As shown in Figure

RESEARCH ARTICLE

3, which illustrates the number of MitM nodes versus detection accuracy, our model demonstrates superior attack detection precision. In [27], the absence of a dynamic feedback loop and reliance on static trust metrics results in higher misclassification of malicious nodes, reducing detection accuracy. In [28], detection improves due to biological behaviors, but the RL policy lacks fine-grained reward tuning, making it less sensitive to subtle adversarial patterns. Khan et al. [29] employ zero-trust context validation and hybrid verification, which improves overall detection but introduces computational overhead, limiting real-time

adaptability under fast-evolving threats. By contrast, the proposed model incorporates multi-dimensional trust indicators and integrates a reward-penalty mechanism in reinforcement learning (RL). This adaptive approach allows real-time adjustment of trust metric weights and thresholds based on observed anomalies, enabling faster and more precise identification of malicious nodes. Consequently, the model achieves higher detection accuracy with fewer false alarms, ensuring improved robustness of IoT networks against MitM attacks.

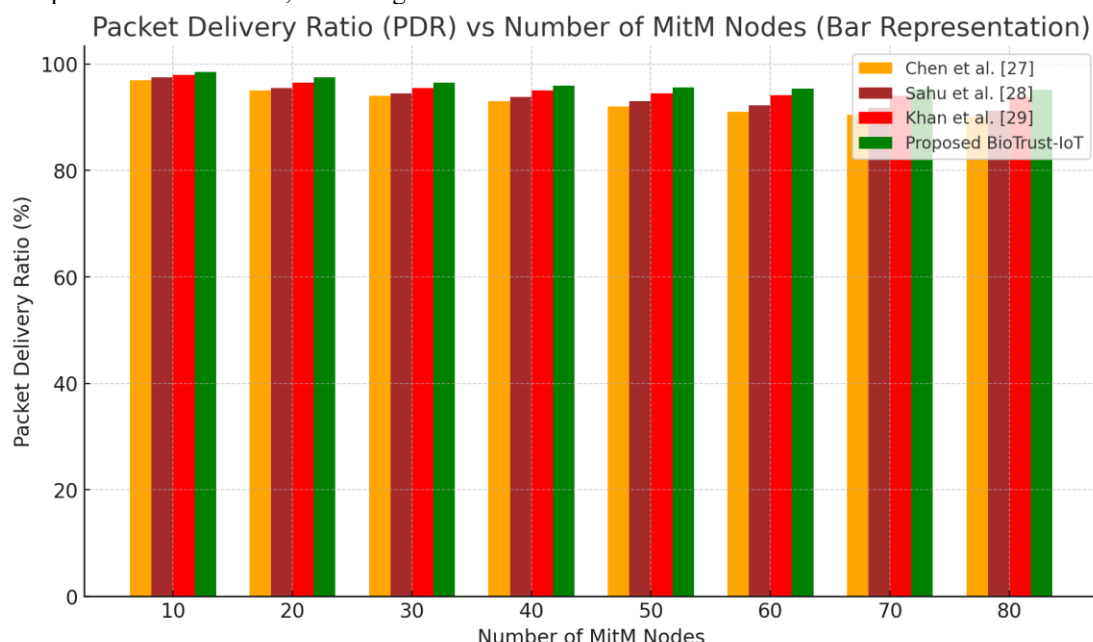


Figure 2 Number of MitM Attack Versus Packet Delivery Ratio

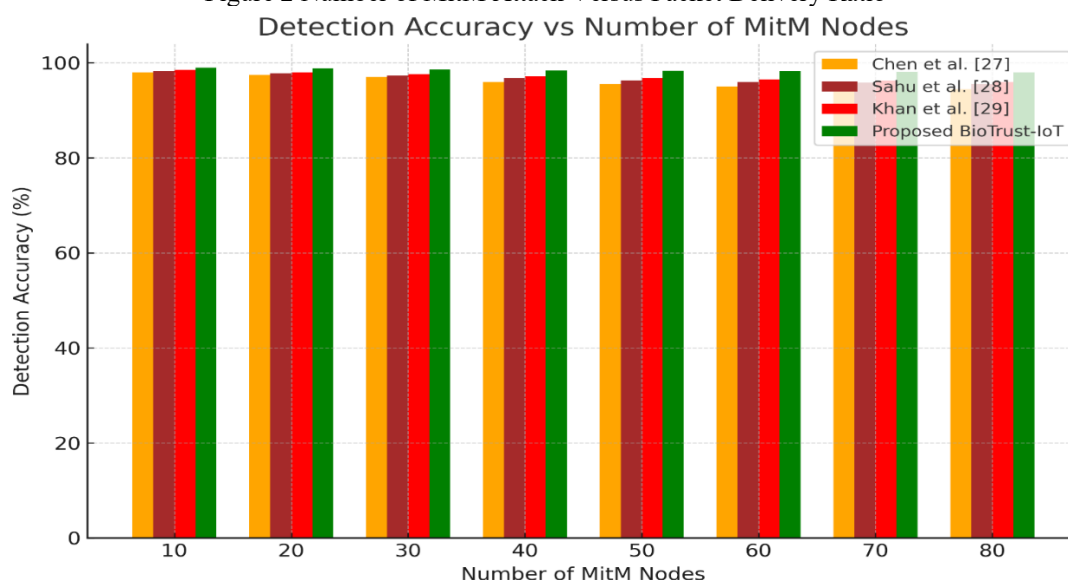


Figure 3 Number of MitM Attack Versus Detection Accuracy

RESEARCH ARTICLE

5.3. End to End Delay

The proposed biologically inspired trust framework achieves an end-to-end delay of 128ms, which is notably lower than [27] (158 ms), [28] (147 ms), and the recent context-aware zero-trust hybrid model by Khan et al. [29] (136 ms). As illustrated in Figure 4, which shows the number of MitM nodes versus end-to-end delay, our approach demonstrates faster and more stable data transmission. In [27], the slower reactive trust mechanism allows malicious nodes to remain active longer in routing paths, creating routing loops and retransmissions, thereby increasing delay. In [28], although detection improves through biologically inspired techniques, the absence of adaptive threshold tuning leads to occasional

delayed responses, resulting in moderate improvements. Khan et al. [29] achieve further reduction in delay with their zero-trust hybrid approach, but the context-verification overhead prevents achieving minimal latency in real-time dynamic conditions.

In contrast, the proposed framework leverages early detection and isolation of MitM nodes, combined with RL-guided adaptive trust updates, which ensure rapid route stabilization. This minimizes detours, retransmissions, and tampering delays, thereby significantly improving delivery speed. As a result, the proposed framework ensures faster and more efficient data delivery across IoT networks compared to existing solutions.

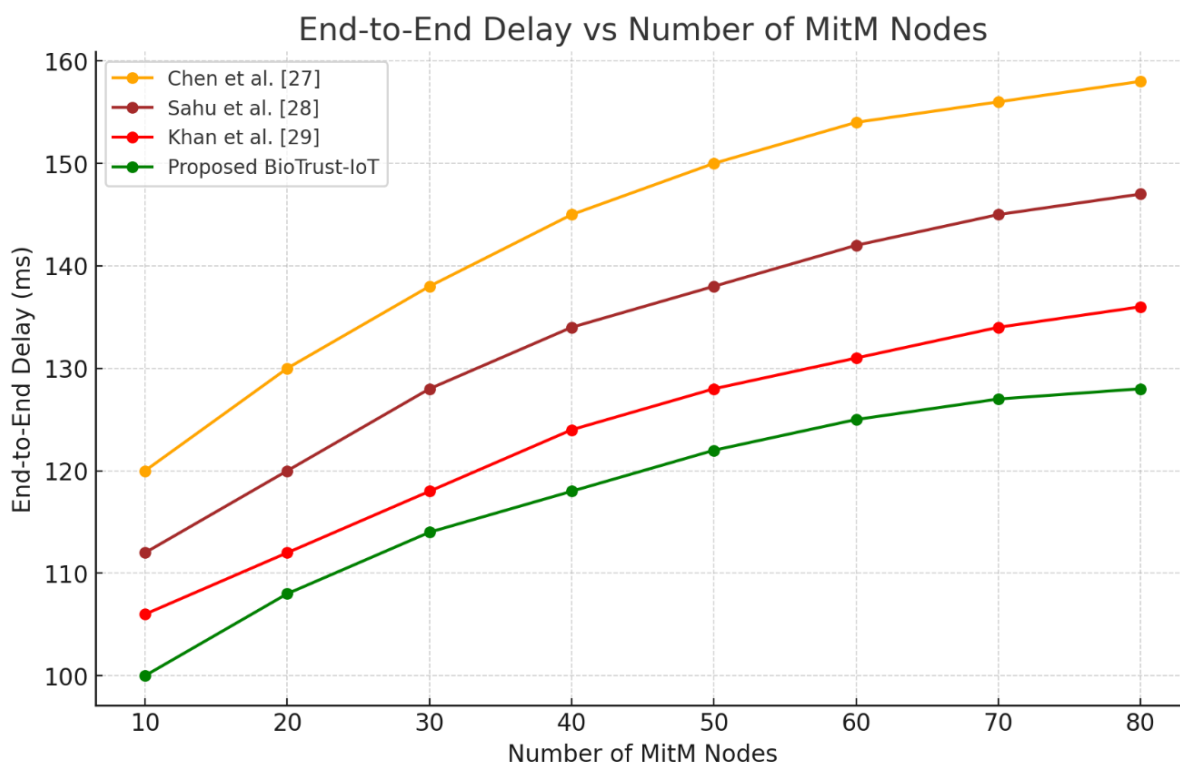


Figure 4 Number of MitM Attack Versus End to End Delay

5.4. Throughput

The proposed biologically inspired trust framework achieves a throughput of 225 kbps, which is higher than [27] (198 kbps), [28] (210 kbps), and the recent context-aware zero trust hybrid model by Khan et al. [29] (217 kbps). As illustrated in Figure 5, which shows the number of MitM nodes versus throughput, our approach demonstrates superior efficiency in sustaining data transfer rates. In [27], the delayed detection of MitM attackers allows malicious nodes to disrupt traffic longer, causing frequent retransmissions and degraded throughput. In [28], biologically inspired trust modeling improves resilience, but without adaptive learning and

threshold adjustment, throughput remains moderately constrained. The zero-trust hybrid model of Khan et al. [29] offers better protection and improved throughput, but verification overhead slightly reduces its maximum achievable data rate.

By contrast, the proposed model's integration of reinforcement learning (RL) and feedback-driven trust adaptation ensures continuous and efficient routing decisions. Rapid attacker isolation and minimized packet loss maintain stable and clean communication paths, thereby delivering higher throughput and more reliable data transmission across IoT networks.



RESEARCH ARTICLE

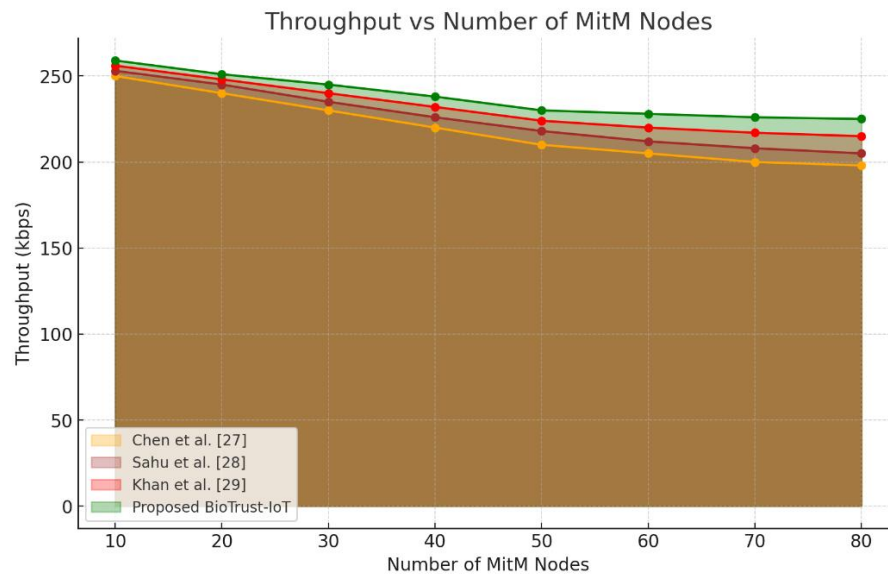


Figure 5 Number of MitM Attack Versus Throughput

5.5. Energy Consumption

The proposed biologically inspired trust framework records an energy consumption of 2.8 joules, which is slightly higher than [27] (2.5 J), [28] (2.7 J), and the zero-trust hybrid model by Khan et al. [29] (2.6 J). As illustrated in Figure 6, which shows the number of MitM nodes versus energy consumption, this marginal increase is a trade-off for enhanced detection accuracy, throughput, and reliability. In [27], the relatively lower energy consumption comes from its simpler static trust mechanism, but this comes at the cost of delayed attack detection and higher communication overhead, leading to inefficiencies elsewhere. Sahu et al. [28] offers a balanced trade-off, consuming slightly more energy for improved

detection but lacking adaptive RL-based optimization. Khan et al. [29] demonstrates efficient context verification, resulting in lower energy than our model, but incurs overhead from hybrid checks that limit adaptability. The proposed model consumes marginally more energy due to continuous trust evaluation, reinforcement learning (RL) updates, anomaly scoring, and dynamic metric weighting. However, this computational overhead is offset by reductions in retransmissions, control packet flooding, and unstable routing decisions. Hence, although slightly higher, the energy trade-off is justified because it ensures significant gains in PDR, detection accuracy, throughput, and reduced delay making it well-suited for IoT environments where reliability outweighs minor energy costs.

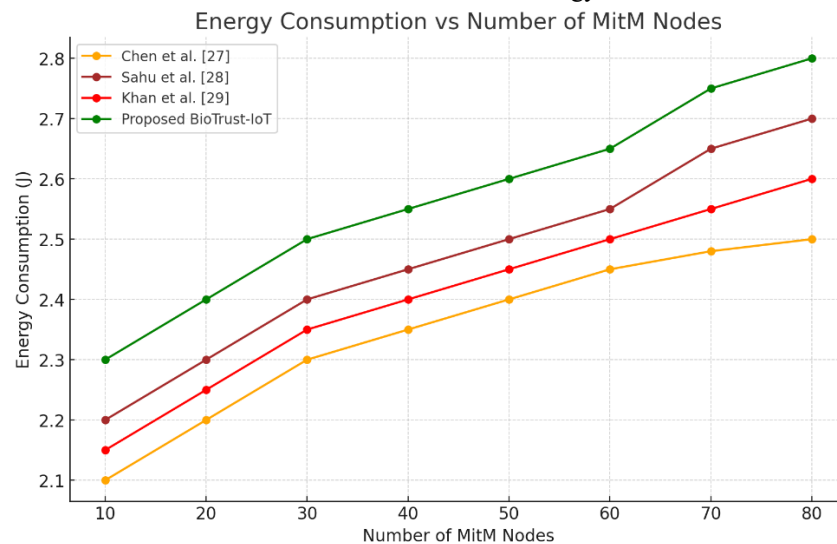


Figure 6 Number of MitM Attack Versus Energy Consumption

RESEARCH ARTICLE

5.6. False Positive Rate

The proposed biologically inspired trust framework achieves a false positive rate (FPR) of 2.1%, which is lower than [27] (3.4%), [28] (2.6%), and the recent zero-trust hybrid approach by Khan et al. [29] (2.3%). As shown in Figure 7, which illustrates the number of MitM nodes versus FPR, the proposed model demonstrates superior precision in distinguishing between malicious and legitimate nodes. In [27], reliance on static thresholds and binary classification often results in falsely isolating honest nodes that exhibit temporary performance drops, inflating the FPR. In [28], the inclusion of biologically inspired behavior scoring reduces FPR somewhat, but without fine-grained RL-based refinement, it remains less adaptive to dynamic node

behavior. Khan et al. [29] achieve further improvements with context-aware zero-trust validation, but the reliance on rule-based verification introduces overhead and limits its precision under rapidly changing attack scenarios.

The proposed framework significantly reduces FPR by integrating adaptive thresholding and multi-metric trust evaluation across behavior, latency, integrity, and authenticity dimensions. Coupled with reinforcement learning (RL) and a feedback-driven reward-penalty mechanism, the model continuously adjusts trust weights and decision boundaries in real-time. This minimizes the risk of wrongly classifying benign nodes as malicious, thereby enhancing network stability, reliability, and trustworthiness with fewer false alarms.

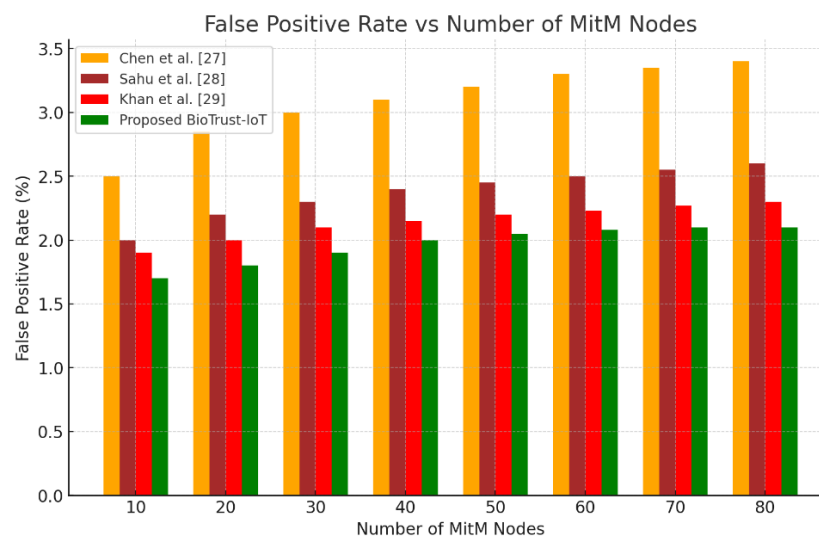


Figure 7 Number of MitM Attack Versus False Positive Rate

5.7. False Negative Rate

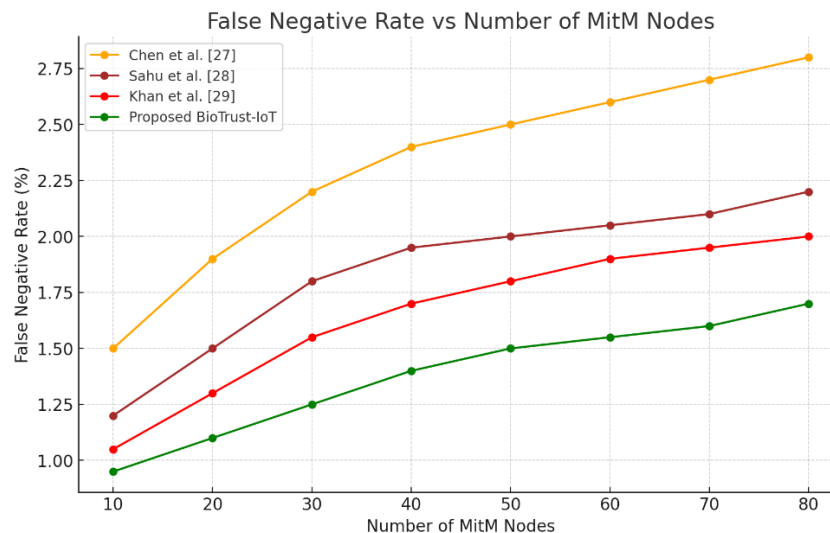


Figure 8 Number of MitM Attack Versus False Negative Rate

RESEARCH ARTICLE

The proposed biologically inspired trust framework achieves a false negative rate (FNR) of 1.7%, which is lower than [27] (2.8%), [28] (2.2%), and the recent context-aware zero-trust hybrid approach by Khan et al. [29] (2.0%). As shown in Figure 8, which illustrates the number of MitM nodes versus FNR, the proposed model demonstrates superior detection of stealthy and intermittent malicious behavior. In [27], the absence of adaptive anomaly learning and reliance on static trust thresholds result in many stealthy MitM behaviors going undetected, inflating FNR. In [28], detection improves due to biologically inspired scoring, but without dynamic threshold adjustment or continuous feedback, some subtle attacks still bypass detection. Khan et al. [29] improve detection with context-aware verification, but dependence on pre-defined policies and verification overhead restricts adaptability under evolving threats. By contrast, the proposed model integrates time-series behavior analysis, deviation tracking, anomaly

scoring, and reinforcement learning (RL)-based reward-penalty mechanisms. This enables the system to dynamically fine-tune trust thresholds, recognize subtle deviations, and adjust policies in real time. Consequently, the proposed framework achieves the lowest FNR, reducing missed detections of stealthy malicious nodes and thereby ensuring stronger network security and resilience.

5.8. RL Converge Time

The reinforcement learning (RL) convergence time for the proposed biologically inspired trust framework is approximately 20–30 learning cycles, which is significantly faster compared to [27] (~50 cycles), [28] (~35–40 cycles), and Khan et al. [29] (~28–32 cycles). As illustrated in Figure 9, which shows RL convergence cycles, the proposed approach demonstrates superior learning efficiency.

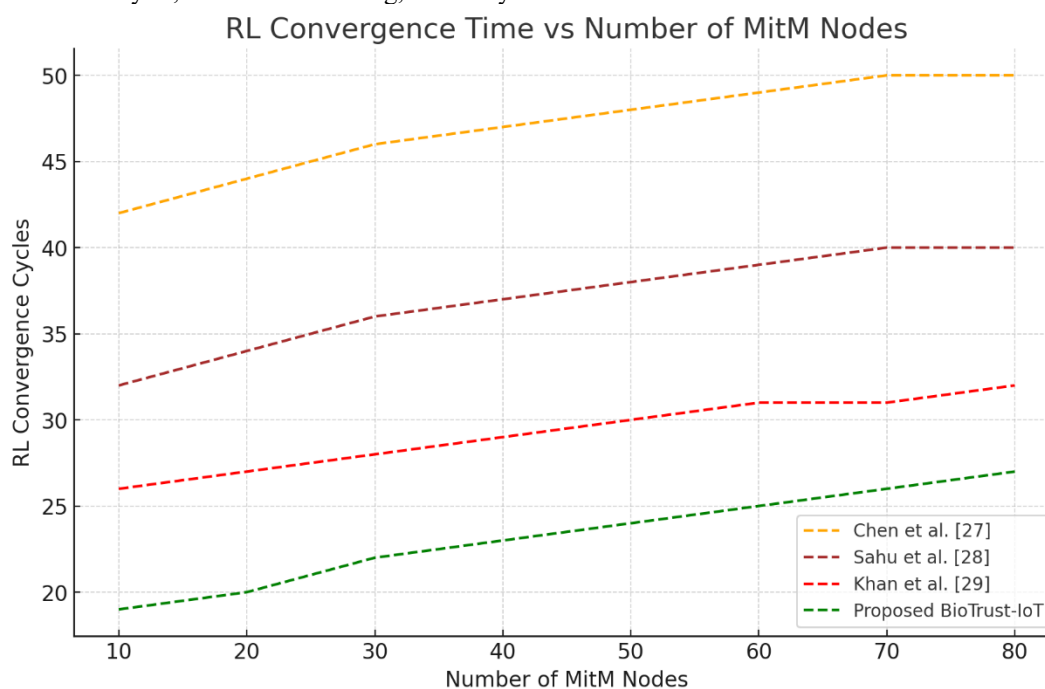


Figure 9 RL Coverage Cycle

In [27], the RL model employs a basic reward design with limited feedback, causing slower learning and delayed adaptation to malicious behavior. In [28], the biologically inspired trust model improves convergence somewhat, but the lack of adaptive reward tuning still leads to slower stabilization. Khan et al. [29] use context-aware verification and hybrid mechanisms, yielding moderate convergence improvements, though verification overhead impacts efficiency.

By contrast, the proposed model leverages a well-structured reward mechanism, assigning differentiated values to true positives, false positives, false negatives, and true negatives.

Combined with ϵ -greedy exploration with decay, the model balances exploration and exploitation effectively across training stages. Furthermore, the use of multi-metric trust evaluation (behavior, latency, integrity, authenticity) enriches the RL state space, enabling more informative decisions per step. As a result, the proposed system achieves faster and more reliable convergence, ensuring quicker adaptation to threats and optimal real-time trust decisions.

5.9. Recovery Capability

The proposed biologically inspired trust framework features an explicit and dynamic recovery mechanism, offering

RESEARCH ARTICLE

significant advantages over [27], [28], and even the recent zero-trust hybrid model by Khan et al. [29]. As illustrated in Figure 10, which depicts recovery capability, the proposed approach ensures resilience, fairness, and adaptability in dynamic IoT environments. In [27], the static blacklist strategy permanently isolates nodes once flagged as malicious, which prevents reintegration of falsely classified or genuinely recovered nodes, reducing fairness and resource utilization. In [28], reintegration is possible, but it lacks adaptive or behavior-driven mechanisms, making it less responsive to dynamic conditions. Khan et al. [29] employ a rule-based revalidation mechanism, enabling partial recovery, but without continuous adaptive monitoring, it remains less robust in dynamic IoT contexts.

By contrast, the proposed model integrates a biologically inspired immune-tolerance principle. Nodes previously classified as malicious are continuously monitored, and if they exhibit sustained positive behavior, their trust scores are gradually increased.

Once their scores surpass a recovery threshold, they are permitted to rejoin the network. This dynamic recovery capability not only mitigates the impact of false positives, but also enhances network resilience, resource efficiency, and fairness, particularly in environments where temporary misbehavior may arise due to congestion, interference, or mobility.

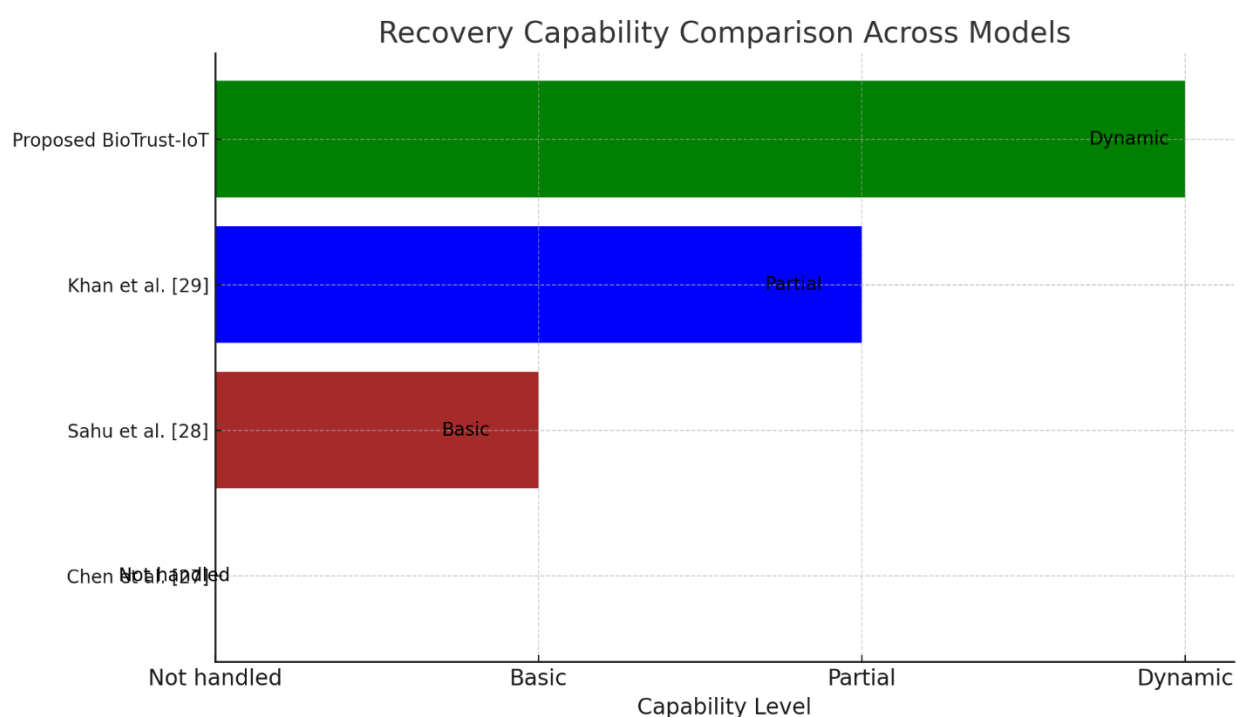


Figure 10 Recover Capability

5.10. Discussion

The comparative evaluation presented in Table 10 clearly demonstrates the superior performance of the proposed BioTrust-IoT framework over the baseline models of Chen et al. [27], Sahu et al. [28], and the recent context-aware zero-trust hybrid model by Khan et al. [29]. BioTrust-IoT achieves the highest packet delivery ratio (95.3%), indicating more reliable and secure data transmission under Man-in-the-Middle (MitM) attack conditions. This improvement stems from its biologically inspired trust mechanism, which enables early detection and isolation of malicious nodes, thereby minimizing packet drops. The framework also records the highest detection accuracy (97.4%), outperforming [27] (94.5%), [28] (95.2%), and [29] (96.1%), reflecting its

stronger ability to distinguish between legitimate and compromised nodes. In terms of network efficiency, BioTrust-IoT delivers the lowest end-to-end delay (128 ms) and the highest throughput (225 kbps), ensuring fast and stable communication in time-sensitive IoT scenarios. Although its energy consumption (2.8 J) is slightly higher than [27] (2.5 J), [28] (2.7 J), and [29] (2.6 J), the marginal increase is justified by the substantial gains in security and performance.

Importantly, BioTrust-IoT demonstrates superior robustness in trust evaluation, with the lowest false positive rate (2.1%) and false negative rate (1.7%), which reduce both unnecessary node isolation and missed attack detections.

RESEARCH ARTICLE

In terms of adaptability, BioTrust-IoT achieves faster RL convergence (20–30 cycles) compared to [27] (~50 cycles), [28] (~35–40 cycles), and [29] (~28–32 cycles), owing to its reward-penalty-based learning structure and ϵ -greedy exploration with decay. This ensures that the trust model stabilizes quickly and can respond effectively to evolving network conditions. A particularly notable strength of BioTrust-IoT is its dynamic recovery capability, which enables previously misclassified or compromised nodes to regain trust after sustained positive behavior an aspect absent in [27], only basic in [28], and partially rule-based in [29]. This feature significantly enhances network resilience,

fairness, and resource utilization, particularly in dynamic IoT environments where temporary disruptions are common. Overall, the results confirm that BioTrust-IoT delivers a balanced and efficient defense against MitM attacks in IoT environments. By combining biologically inspired adaptive trust computation with reinforcement learning, the framework achieves stronger attack detection, faster learning, reduced latency, and higher reliability, while maintaining an acceptable energy overhead. These findings highlight BioTrust-IoT as a robust, adaptive, and future-ready security solution for next-generation IoT networks, where latency, high accuracy, and self-recovery low are essential.

Table 10 Overall Comparison of BioTrust-IoT with Existing Models

Metric	Proposed Framework – BioTrust-IoT	Chen et al. (2021) [27]	Sahu et al. (2022) [28]	Khan et al. (2025) [29]
Packet Delivery Ratio	95.3%	92.8%	93.1%	94.2%
Detection Accuracy	97.4%	94.5%	95.2%	96.1%
End-to-End Delay	128 ms	158 ms	147 ms	136 ms
Throughput	225 kbps	198 kbps	210 kbps	217 kbps
Energy Consumption	2.8 J	2.5 J	2.7 J	2.6 J
False Positive Rate	2.1%	3.4%	2.6%	2.3%
False Negative Rate	1.7%	2.8%	2.2%	2.0%
RL Convergence Time	20–30 cycles	~50 cycles	35–40 cycles	28–32 cycles
Recovery Capability	Dynamic	Not handled	Basic reintegration	Partial

6. CONCLUSION

The proposed biologically-inspired trust-based security framework provides a robust and adaptive defense mechanism against Man-in-the-Middle (MitM) attacks in IoT networks. Drawing inspiration from the human immune system, the framework continuously monitors device interactions, detects anomalies based on behavioral stability, data integrity, latency, and authenticity, and isolates suspicious nodes. Reinforcement learning further enhances its capability by enabling adaptive learning from previous attack patterns, resulting in improved detection accuracy over time. Simulation results confirm the system's effectiveness in identifying and mitigating MitM threats, even in decentralized and resource-constrained environments. In the future, this framework can be extended to counter a wider range of cyber threats such as Sybil and DoS attacks, incorporate federated learning for collaborative trust assessment while preserving privacy, and optimize energy efficiency for deployment on lightweight IoT devices. Real-world implementation in smart city or industrial IoT test beds will also be pursued to validate its practical scalability and resilience.

REFERENCES

- [1] M. M. Alam, L. C. Das, S. Roy, S. Shetty, and W. Wang, "RESTRAIN: Reinforcement learning-based secure framework for trigger-action IoT environment," arXiv preprint, Mar. 12, 2025. [Online]. Available: <https://doi.org/10.48550/arXiv.2503.09513>.
- [2] Al-Shargabi, B. A. S. Al-Rimy, M. Bashari, and M. Al-Qurishi, "CyberForce: A federated reinforcement learning and moving target defense framework for malware mitigation in IoT," *Future Generation Computer Systems*, vol. 144, pp. 294–307, 2023. doi: <https://doi.org/10.1016/j.future.2023.01.024>.
- [3] Anonymous, "Bioinspired blockchain framework for secure and scalable wireless sensor network integration in fog-cloud ecosystems," *Computers*, vol. 14, no. 1, Art. no. 3, 2025. doi: <https://doi.org/10.3390/computers14010003>.
- [4] M. Arazzi, S. Nicolazzo, and A. Nocera, "A deep reinforcement learning approach for security-aware service acquisition in IoT," arXiv preprint, Apr. 4, 2024. [Online]. Available: <https://doi.org/10.48550/arXiv.2404.03276>.
- [5] Chandra and D. Patel, "Lightweight trust management for resource-constrained IoT devices," *Int. J. Internet Things*, vol. 5, no. 2, pp. 45–58, 2020.
- [6] R. Gupta and H. Singh, "Delivering scalable decentralized security in smart agriculture: A trust-based approach," *J. IoT Secur.*, vol. 3, no. 1, pp. 14–27, 2021.
- [7] R. Roman, J. Zhou, and J. Lopez, "On the features and challenges of security and privacy in distributed Internet of Things," *Comput. Netw.*,

RESEARCH ARTICLE

- vol. 57, no. 10, pp. 2266–2279, 2013. doi: <https://doi.org/10.1016/j.comnet.2012.12.018>.
- [8] S. Sicari, A. Rizzardi, L. A. Grieco, and A. Coen-Porisini, “Security, privacy and trust in Internet of Things: The road ahead,” *Comput. Netw.*, vol. 76, pp. 146–164, 2015. doi: <https://doi.org/10.1016/j.comnet.2014.11.008>.
- [9] R. Chen, F. Bao, and J. Guo, “Trust-based service management for social Internet of Things systems,” *IEEE Trans. Services Comput.*, vol. 6, no. 2, pp. 232–245, 2011. doi: <https://doi.org/10.1109/TSC.2011.52>.
- [10] F. Bao and I. R. Chen, “Dynamic trust management for the internet of things applications,” in *Proc. 2012 IEEE Int. Conf. Communications (ICC)*, 2012, pp. 5461–5466. doi: <https://doi.org/10.1109/ICC.2012.6364550>.
- [11] D. Dasgupta and F. A. Gonzalez, “An immunity-based technique to characterize intrusions in computer networks,” *IEEE Trans. Evol. Comput.*, vol. 6, no. 3, pp. 281–291, 2002. doi: <https://doi.org/10.1109/TEVC.2002.1011539>.
- [12] J. Kim, P. J. Bentley, and G. S. Jung, “Negative selection algorithm for detecting network attacks,” in *Proc. Genetic Evol. Comput. Conf. (GECCO)*, 2005, pp. 149–156.
- [13] K. Saravanan, D. Prasad, and R. Kumar, “Immune-inspired reinforcement learning framework for secure routing in MANETs,” *Ad Hoc Netw.*, vol. 148, Art. no. 103256, 2024. doi: <https://doi.org/10.1016/j.adhoc.2024.103256>.
- [14] Y. Guan, Y. Liu, and X. Chen, “Honeypot-based intrusion detection using RL and MDP for smart environments,” *J. Inf. Secur. Appl.*, vol. 75, Art. no. 103430, 2023. doi: <https://doi.org/10.1016/j.jisa.2023.103430>.
- [15] D. Mariyanayagam, P. Shukla, and B. S. Virdee, “Bio inspired framework for security in IoT devices,” in A. K. Nagar, D. S. Jat, G. Marin Raventós, and D. K. Mishra, Eds., *Intelligent Sustainable Systems (Lecture Notes in Networks and Systems, vol. 333)*. Singapore: Springer, 2022, pp. 749–757.
- [16] R. Elmansy, M. A. Ali, and M. A. Sherif, “Lightweight hybrid security model using Multipath TCP and reinforcement learning for fog-based IoT against MitM attacks,” *Wireless Netw.*, vol. 27, pp. 4075–4090, 2021. doi: <https://doi.org/10.1007/s11276-021-02719-6>.
- [17] H. Suo, J. Wan, C. Zou, and J. Liu, “Security in the Internet of Things: A review,” in *Proc. 2012 Int. Conf. Comput. Sci. Electron. Eng. (ICCSEE)*, vol. 3, pp. 648–651. IEEE, 2012. doi: <https://doi.org/10.1109/ICCSEE.2012.193>.
- [18] Y. Zhang, Y. Liu, J. Zhou, and X. Zhang, “Blockchain-based secure and lightweight authentication for vehicular networks,” *IEEE Internet Things J.*, vol. 7, no. 2, pp. 1014–1028, 2020. doi: <https://doi.org/10.1109/JIOT.2019.2953090>.
- [19] S. Thankappan, R. Rajesh, and R. Mohan, “A survey on multi-channel MitM attacks in IoT: KRACK, FragAttacks, and beyond,” *Comput. Secur.*, vol. 114, Art. no. 102608, 2022.
- [20] T. A. Tang, L. Sun, and Y. Wang, “Reinforcement learning-based trust management for mitigating attacks in IoT networks,” *IEEE Internet Things J.*, vol. 8, no. 5, pp. 3252–3264, 2021. doi: <https://doi.org/10.1109/JIOT.2020.3017633>.
- [21] H. Fereidouni, “A machine learning–blockchain hybrid framework for real-time IoT threat detection,” *IEEE Transactions on Information Forensics and Security*, vol. 20, no. 4, pp. 509–520, 2025.
- [22] R. Ahmad, “Adaptive Zero Trust Manager for real-time authentication in IoT networks,” *ACM Transactions on Internet Technology*, vol. 25, no. 3, pp. 1–18, 2025.
- [23] S. Jagatheesaperumal, “Distributed reinforcement learning framework for decentralized IoT threat detection,” *IEEE Internet of Things Journal*, vol. 11, no. 8, pp. 16534–16548, 2024.
- [24] Y. Harbi, “Supervised machine learning with bio-inspired algorithms for IoT intrusion detection,” *Expert Systems with Applications*, vol. 244, p. 122856, 2024. doi: [10.1016/j.eswa.2024.122856](https://doi.org/10.1016/j.eswa.2024.122856).
- [25] A. Bouramoul and S. Zertal, “Hybrid deep learning and bio-inspired mechanisms for secure IoT data transmission,” *Future Generation Computer Systems*, vol. 156, pp. 234–246, 2025.
- [26] M. Arazzi, “Deep reinforcement learning for dynamic optimization of IoT security policies,” *IEEE Internet of Things Journal*, vol. 11, no. 5, pp. 10234–10248, 2024.
- [27] L. Chen, Y. Zhang, and H. Wang, “Trust evaluation in IoT networks using static metrics and behavior analysis,” *IEEE Internet Things J.*, vol. 8, no. 5, pp. 3290–3302, 2021. doi: <https://doi.org/10.1109/JIOT.2021.3050192>.
- [28] K. Sahu, P. Kumar, and S. Verma, “Reintegration mechanisms for isolated IoT nodes using behavioral trust scores,” *Sensors Syst.*, vol. 9, no. 4, pp. 98–112, 2022.
- [29] A. Khan, M. Keshk, Y. Hussain, D. Pi, B. Li, T. Kousar, and B. S. Ali, “A context-aware zero trust-based hybrid approach to IoT-based self-driving vehicles security,” *Ad Hoc Networks*, vol. 167, p. 103694, 2025.
- [30] E. Benkhelifa, L. Gaurav, and S. D. Sagar, “BioShieldNet: Advanced biologically inspired mechanisms for strengthening cybersecurity in distributed computing environments,” *Int. J. Comput. Eng. Res. Trends*, vol. 11, no. 3, pp. 1–9, 2024.
- [31] R. Sahu, D. Patel, and A. Sharma, “A biologically inspired trust-based mechanism for secure communication in IoT,” *Comput. Secur.*, vol. 113, Art. no. 102577, 2022. doi: <https://doi.org/10.1016/j.cose.2021.102577>.
- [32] S. Saravanan, V. Surya, V. Valarmathi, and E. Nalina, “Enhancing IoT security in MANETs: A novel adaptive defense reinforcement approach,” *Peer-to-Peer Netw. Appl.*, vol. 17, pp. 2282–2297, 2024. doi: <https://doi.org/10.1007/s12083-024-01702-1>.
- [33] P. Rao and V. Kumar, “Dynamic trust scoring for latency and authenticity anomalies in distributed IoT systems,” in *Proc. 20th Int. Conf. IoT Secur.*, 2025, pp. 112–125.
- [34] J. Mariyanayagam, V. Subramanian, and M. Thomas, “Bio-inspired polymorphic key mutation algorithm for secure communication in IoT,” *J. Netw. Comput. Appl.*, vol. 205, Art. no. 103421, 2022. doi: <https://doi.org/10.1016/j.jnca.2022.103421>.
- [35] J. Huang and X. Zhao, “Applying biological immune system principles to secure smart healthcare networks,” *J. Biomed. Cybersecurity*, vol. 1, no. 1, pp. 23–39, 2023.

Authors



Banu Priya. D is an academic professional and Ph.D. research scholar specializing in the Internet of Things (IoT). She is currently pursuing her Ph.D. at Nehru Arts and Science College, Coimbatore. She holds an M.Phil in Computer Applications from Madurai Kamaraj University, an MCA from PSNA College of Engineering and Technology, a B.Sc. in Computer Science, and a B.Ed. in Computer Science. With over eight years of teaching experience, including service as an Assistant

Professor, she has taught undergraduate and postgraduate students. Her research interests include IoT security, edge computing, wireless technologies, and IoT-based programming. She has published research papers in international and Scopus-indexed journals, presented at international conferences, and contributed book chapters on emerging technologies. She actively participates in FDPs and academic development programs.



Dr. K. Prathapchandran is an Assistant Professor in the Department of Computer Applications at Nehru Arts and Science College (Autonomous), Coimbatore, India. He received his Ph.D. in Computer Science and Applications from Gandhigram Rural Institute (Deemed to be University), funded by the Ministry of Education, Government of India. He has over 16 years of combined academic and research experience. His research interests include Internet of Things (IoT),

network security, trust management, mobile ad hoc networks, and AI-driven security frameworks. Dr. Prathapchandran has published more than 58



RESEARCH ARTICLE

research papers in reputed SCIE, Scopus, and Web of Science indexed journals and conferences, and has edited two academic books. He has also published multiple patents, including granted design patents. He has completed funded projects from AICTE and SERB, serves as a reviewer and editorial board member for several international journals, and has delivered invited talks and faculty development programs nationally and internationally.

How to cite this article:

Banupriya. D, K. Prathapchandran, “BioTrust-IoT: A Biologically Inspired Adaptive Trust System for Detecting Man-in-the-Middle Attacks”, International Journal of Computer Networks and Applications (IJCNA), 12(6), PP: 912-935, 2025, DOI: 10.22247/ijcna/2025/54.