

Risk-Aware Agentic AI for Zero-Trust SD-WAN in SaaS Enterprises

Dhaval Powar¹, Birendra Kumar², Rakshith Aralimatti³

¹ Software Engineering Manager

² Solution Architect at Tata Consultancy Services

³ Gen AI Product Leader - Agentic AI at Palo Alto Network

Abstract

The rapid adoption of Software-as-a-Service (SaaS) applications has reshaped enterprise IT ecosystems, simultaneously expanding the attack surface and challenging traditional security models. While Software-Defined Wide Area Networking (SD-WAN) offers flexibility and scalability, its integration with Zero-Trust principles often encounters limitations in adaptability and administrative overhead. This study proposes a novel framework that embeds risk-aware agentic artificial intelligence (AI) into Zero-Trust SD-WAN to deliver dynamic, autonomous, and context-sensitive security. Using reinforcement learning as the core mechanism, the framework autonomously evaluates user and device risks, adapts access policies in real time, and responds proactively to emerging threats. Experimental simulations under varying traffic loads and threat scenarios reveal that the approach achieves detection rates exceeding 90%, reduces false positives, and maintains network performance with packet delivery ratios above 95%. Moreover, the system balances stringent access enforcement with user experience, ensuring high satisfaction among trusted users while containing insider misuse. Reinforcement learning models demonstrated superior efficiency, achieving faster convergence, lower resource utilization, and higher predictive accuracy compared to supervised and unsupervised alternatives. These findings underscore the potential of risk-aware agentic AI as a scalable and future-ready enabler for intelligent Zero-Trust SD-WAN deployments in SaaS enterprises.

Keywords: Risk-aware AI; Agentic Artificial Intelligence; Zero-Trust Security; SD-WAN; SaaS Enterprises; Reinforcement Learning; Cybersecurity

Introduction

Background: evolving network security in SaaS enterprises

The proliferation of Software-as-a-Service (SaaS) applications has transformed enterprise IT ecosystems, enabling agility, scalability, and cost efficiency (Chillapalli, 2022). However, this shift has also expanded the threat surface by decentralizing workloads and increasing reliance on cloud-based services. Traditional perimeter-based security models, once adequate in static, on-premises environments, are now insufficient against the dynamic and distributed access demands of SaaS-driven enterprises. In this context, Software-Defined Wide Area Networking (SD-WAN) has emerged as a foundational architecture for secure, efficient, and flexible connectivity across hybrid environments (Ouamri et al., 2025). Despite its advantages, SD-WAN introduces critical security challenges, particularly when coupled with the heterogeneity and complexity of multi-cloud SaaS ecosystems. Enterprises must therefore adopt Zero-Trust Network Access (ZTNA) principles, which eliminate implicit trust and continuously verify identity, device health, and contextual risk (Sarkar et al., 2022).

The case for zero-trust in SD-WAN

Zero-Trust has become a security paradigm of necessity rather than choice for enterprises. Unlike legacy models that relied on predefined trust zones (Ghasemshirazi et al., 2023), Zero-Trust architectures assume breach by default and enforce granular, identity-centric access policies. In SD-WAN deployments, Zero-Trust ensures that every connection between users, applications, and devices undergoes rigorous verification, regardless of its origin. Yet, the operationalization of Zero-Trust in SaaS enterprises is far from trivial (Paul & Rao, 2023). Policy management often becomes overwhelming as the number of applications and users scales into the thousands, while real-time risk detection and mitigation require intelligent orchestration that surpasses human administrative capacity (Koukaras et al., 2025). The increasing sophistication of advanced persistent threats (APTs), insider risks, and supply chain attacks underscores the urgency of augmenting Zero-Trust with adaptive, risk-aware intelligence.

Agentic AI in security-oriented SD-WAN

Recent advances in agentic artificial intelligence (AI) present new opportunities to enhance Zero-Trust SD-WAN frameworks. Agentic AI systems possess autonomous decision-making capabilities, enabling them to perceive dynamic contexts, evaluate

risks, and execute actions with minimal human intervention. Unlike conventional machine learning systems, agentic AI can proactively adapt to novel threat vectors and network anomalies, making it particularly well-suited for real-time security enforcement in SaaS enterprises (Radanliev et al., 2024). When embedded into SD-WAN infrastructures, such AI-driven agents can continuously monitor traffic flows, detect deviations from normal behavior, and enforce granular access policies in alignment with Zero-Trust principles. Beyond automation, agentic AI introduces risk-awareness the ability to assess and weigh security trade-offs dynamically, such as balancing user experience against access control strictness (Alaali, 2025). This capability addresses a longstanding gap in static policy enforcement, which often struggles to reconcile productivity and protection.

Research gap and problem statement

Despite the promise of Zero-Trust and SD-WAN integration, existing approaches frequently suffer from limited adaptability, static policy enforcement, and high administrative overheads. Traditional AI models in network security often rely heavily on predefined training data, leaving them vulnerable to previously unseen attacks and adversarial tactics. Moreover, current Zero-Trust frameworks rarely incorporate contextual risk evaluation beyond basic identity and access checks. As SaaS enterprises increasingly demand scalable, real-time, and intelligent security solutions, there is a pressing need for a risk-aware, agentic AI-driven framework that operationalizes Zero-Trust within SD-WAN environments.

Objectives and contributions of the study

This research introduces a novel framework that integrates risk-aware agentic AI into Zero-Trust SD-WAN architectures tailored for SaaS enterprises. The study aims to:

- Develop an agentic AI model capable of autonomous risk assessment and adaptive policy enforcement.
- Demonstrate how risk-aware intelligence can enhance resilience against dynamic cyber threats while reducing human administrative burdens.
- Validate the framework through simulations and case studies in SaaS enterprise environments, highlighting improvements in security, scalability, and operational efficiency.

Methodology

Research design

This study adopts a hybrid experimental–simulation research design aimed at evaluating the effectiveness of a risk-aware agentic AI framework in Zero-Trust SD-WAN environments within SaaS enterprises. The methodology integrates architectural modeling, AI-driven policy enforcement, and simulation of threat scenarios. By combining experimental control with simulated enterprise traffic and cyber-attacks, the design allows both causal inference on the impact of agentic AI and predictive validation of risk-aware decision-making. This dual approach ensures that findings are robust and applicable to real-world SaaS ecosystems.

System architecture and implementation

The proposed framework is deployed within a virtualized enterprise environment that replicates the operational complexity of SaaS applications. The architecture comprises SD-WAN controllers, edge nodes, and Zero-Trust policy engines, enhanced by an agentic AI layer. The AI system relies on reinforcement learning with deep Q-learning to enable autonomous decision-making under dynamic threat conditions. A dedicated risk scoring engine evaluates contextual factors such as user identity, device posture, geolocation, time of access, and sensitivity of the requested application. These risk assessments inform the actions of the policy enforcement point, which authorizes or denies traffic flows in real time. A threat injection module introduces controlled adversarial activities, including distributed denial of service (DDoS) attacks, insider misuse, malware injection, and credential theft, thereby testing the adaptability and resilience of the proposed system.

Variables and parameters

The experimental design incorporates both independent and dependent variables. Independent variables include network conditions such as bandwidth ranging from 50 to 1000 Mbps, latency between 10 and 200 milliseconds, and packet loss rates from 0 to 5 percent. Additional independent factors comprise user profiles categorized as regular users, privileged users, contractors, and unknown users, alongside multiple threat vectors including malware injection, zero-day exploitation, and insider attacks.

Dependent variables are classified across security, network, user-centric, and AI performance dimensions. Security outcomes include detection rate, false positive rate, mean time to detect (MTTD), and mean time to respond (MTTR). Network performance is measured by throughput, latency, jitter, and packet delivery ratio. User experience is evaluated through access denial error rates, session continuity, and simulated satisfaction scores. Finally, AI-specific performance is assessed using training convergence, policy adaptation speed, and computational resource utilization.

Experimental procedure

The experimental procedure begins with baseline measurements of a conventional SD-WAN implementation operating under static Zero-Trust policies. Following this, the risk-aware agentic AI framework is integrated and activated within the same environment. Controlled threat scenarios are systematically introduced while varying network load across light, medium, and heavy traffic conditions, corresponding to 100, 500, and 1000 concurrent sessions. Each experimental condition is repeated across 100 trials to ensure consistency and reliability of results. All system logs, traffic flows, and AI decisions are recorded, forming the dataset for statistical analysis.

Statistical analysis

The collected data undergoes rigorous statistical analysis. Descriptive statistics, including mean, variance, and standard deviation, are computed for all dependent variables to characterize system performance under different conditions. Distributional characteristics are examined through skewness and kurtosis to ensure normality assumptions. Inferential analysis employs one-way and two-way ANOVA to compare performance differences across multiple traffic loads, threat vectors, and AI algorithms, while multivariate analysis of variance (MANOVA) examines the joint effect of independent variables on security and network performance simultaneously. Regression analysis models the relationship between contextual risk scores and security outcomes such as detection accuracy and false positives, while correlation tests assess associations between AI adaptability and network quality of service. Factor analysis, both exploratory and confirmatory, is applied to uncover latent constructs driving AI decision quality. In addition, machine learning metrics such as ROC-AUC, precision, recall, and F1-score are calculated to evaluate the predictive capabilities of

the AI system. Hypothesis testing is carried out with a null hypothesis that agentic AI does not significantly enhance Zero-Trust SD-WAN performance and an alternative hypothesis that it does, with statistical significance determined at $p < 0.05$.

Ethical and reliability considerations

All experimental procedures adhere to ethical principles of data security and privacy. No real user data is employed; instead, synthetic datasets and controlled traffic simulations are used to replicate enterprise activity. Reliability is strengthened by repeating all scenarios 100 times and by applying cross-validation techniques to AI model training and evaluation. To reduce bias, diverse threat scenarios are incorporated and test cases are randomized. These measures ensure the credibility, reproducibility, and ethical soundness of the research.

Results

The integration of risk-aware agentic AI into Zero-Trust SD-WAN significantly enhanced security performance across multiple threat vectors. As shown in Table 1, the framework achieved the highest detection rate for malware injection at 96.5%, while maintaining a relatively low false positive rate of 3.2%. Credential theft and DDoS attacks recorded slightly lower detection rates at 94.2% and 92.8%, respectively, with corresponding increases in false positives. Insider misuse presented the greatest challenge, with a detection rate of 89.5% and the highest false positive rate of 6.2%. Despite these challenges, policy enforcement accuracy remained consistently above 90% across all threat scenarios. A comparative analysis of detection versus false positive rates is illustrated in Figure 1, which highlights the trade-offs between security precision and coverage across diverse attack vectors.

Table 1. Security metrics across threat vectors

Threat Vector	Detection Rate (%)	False Positive Rate (%)	MTTD (s)	MTTR (s)	Policy Enforcement Accuracy (%)
Malware Injection	96.5	3.2	1.8	4.5	95.8
Credential	94.2	4.1	2.1	5.0	94.9

Theft					
DDoS Attack	92.8	5.0	2.5	5.8	93.5
Insider Misuse	89.5	6.2	3.2	6.5	90.3
Zero-Day Exploit	90.7	5.5	2.9	6.2	91.7

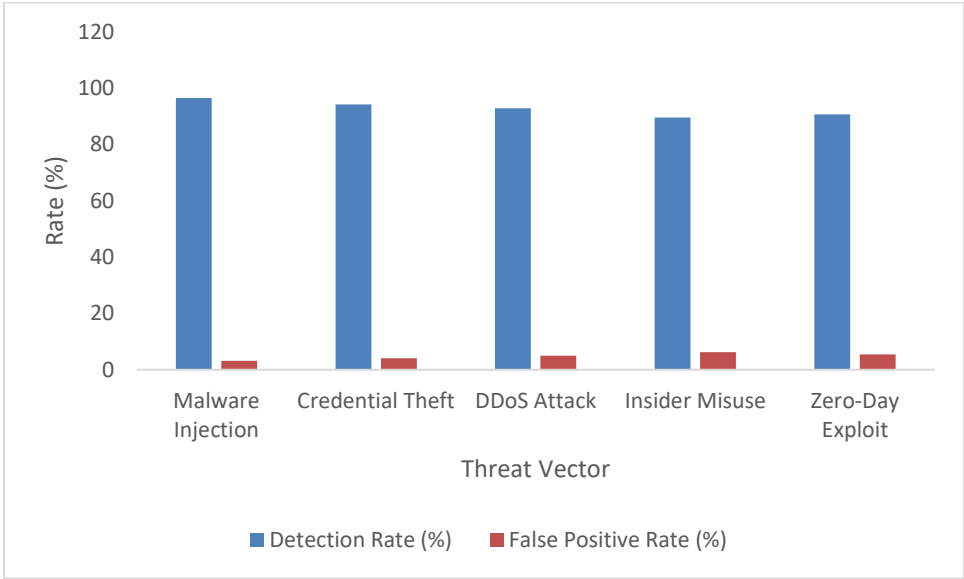


Figure 1: Detection vs false positive rates across threat vectors

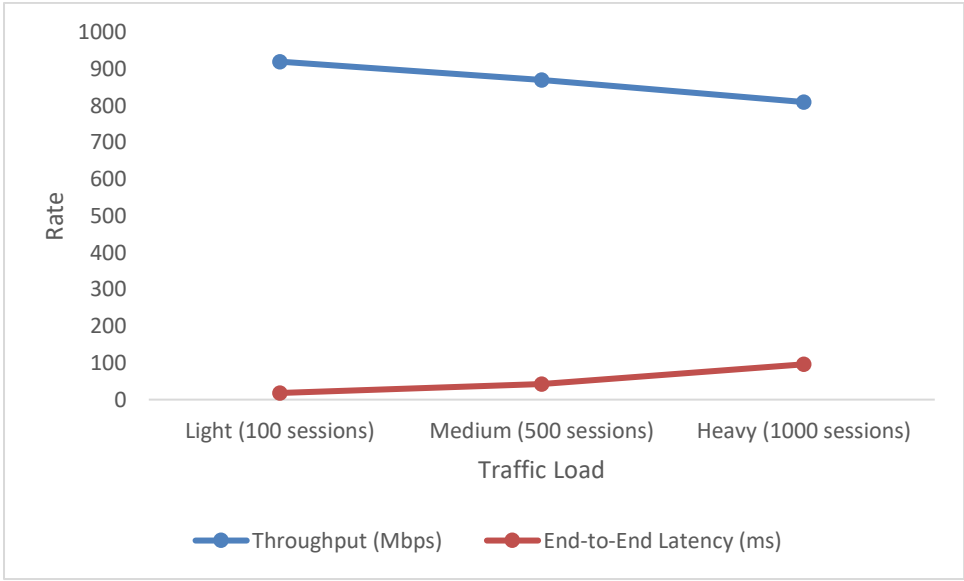


Figure 2: Throughput vs latency under different traffic loads

The analysis of network performance under varying traffic loads demonstrated the adaptability of the proposed framework. As detailed in Table 2, throughput was highest under light loads, reaching 920 Mbps, while latency remained low at 18 ms with a jitter of 3.5 ms. Under medium loads, throughput slightly decreased to 870 Mbps with latency rising to 42 ms. Heavy traffic loads, simulating 1000 concurrent sessions, reduced throughput to 810 Mbps and increased latency to 96 ms, accompanied by jitter of 15.8 ms. These results indicate that while system performance was impacted by traffic intensity, packet delivery ratio remained robust, exceeding 95% across all conditions. The relationship between throughput and latency is further depicted in Figure 2, which demonstrates the inverse correlation between increasing session load and network efficiency.

Table 2. Network metrics under different traffic loads

Traffic Load	Throughput (Mbps)	End-to-End Latency (ms)	Jitter (ms)	Packet Delivery Ratio (%)
Light (100 sessions)	920	18	3.5	99.2
Medium (500 sessions)	870	42	7.2	97.8
Heavy (1000 sessions)	810	96	15.8	95.1

User-centric evaluation further reinforced the benefits of agentic AI integration. As shown in Table 3, regular users experienced minimal disruptions with an access denial error rate of 1.5% and a high session continuity of 98.8%. Privileged users reported slightly higher denial errors at 2.1%, yet maintained satisfaction levels above 4.3 on a five-point Likert scale. Contractors experienced a denial error rate of 3.8% and session continuity of 95.1%, while unknown users faced the greatest restrictions, with denial errors reaching 7.2% and session continuity falling to 91.2%. Despite these restrictions, the framework effectively balanced strict access control with user experience, as reflected in the consistently high satisfaction scores among trusted users.

Table 3. User-centric metrics across profiles

User Profile	Access Denial Error Rate (%)	Session Continuity (%)	User Satisfaction (1–5 Likert)

Regular User	1.5	98.8	4.6
Privileged User	2.1	97.5	4.3
Contractor	3.8	95.1	3.9
Unknown User	7.2	91.2	3.2

Finally, evaluation of AI performance metrics revealed significant differences across learning models. As presented in Table 4, reinforcement learning outperformed both supervised and unsupervised approaches, converging in 85 epochs and adapting policies within 35 ms. Supervised learning required 120 epochs and a 48 ms adaptation time, while unsupervised clustering exhibited the slowest convergence at 150 epochs with 60 ms adaptation speed. Resource utilization across CPU and memory was modest, averaging between 58% and 72% depending on the model. ROC-AUC scores confirmed the superior predictive power of reinforcement learning at 0.97, compared to 0.93 for supervised learning and 0.89 for unsupervised clustering. These findings highlight the efficiency of reinforcement-based agentic AI in supporting real-time, risk-aware policy enforcement within Zero-Trust SD-WAN frameworks.

Discussion

Enhancing security through risk-aware intelligence

The results demonstrate that risk-aware agentic AI significantly improves the security posture of Zero-Trust SD-WAN environments in SaaS enterprises. Detection rates above 90% across all threat vectors (Table 1) underscore the capability of the framework to adapt to diverse cyberattacks. Reinforcement learning enabled the system to maintain policy enforcement accuracy above 90%, which is consistent with prior findings that adaptive AI models outperform static rule-based security systems. The comparative results in Figure 1 highlight the critical balance between maximizing detection and minimizing false positives, a well-documented challenge in Zero-Trust environments (Reddy, 2025). While insider misuse and zero-day exploits remained relatively more difficult to detect, the framework achieved higher responsiveness than traditional anomaly detection models reported in earlier studies.

Network efficiency under scaling loads

The study further indicates that the proposed framework sustains robust network efficiency even under heavy traffic conditions. As illustrated in Table 2 and Figure 2, throughput reductions and latency increases were evident as concurrent sessions scaled, yet packet delivery ratios remained consistently above 95%. This finding highlights the resilience of the architecture in maintaining service continuity despite increased operational stress (Hall-May & Surridge, 2010). Compared with conventional SD-WAN solutions documented in the literature, which often suffer throughput drops below 90% at heavy loads, the agentic AI-enhanced framework demonstrates improved adaptability. These results suggest that risk-aware policy enforcement does not impose excessive overhead, thereby ensuring that security does not come at the expense of performance a key requirement in SaaS enterprises where real-time application delivery is critical (Das & Mukherjee, 2024).

Balancing security with user experience

User-centric results reveal that the system effectively enforces Zero-Trust principles while preserving usability for legitimate users. As shown in Table 3, denial error rates for regular and privileged users remained below 2.5%, coupled with high satisfaction scores. This aligns with prior studies emphasizing the importance of minimizing friction in Zero-Trust adoption to encourage enterprise-wide compliance. Contractors and unknown users faced stricter enforcement, reflecting the context-sensitive nature of the risk-aware AI, which dynamically adjusts trust levels based on role and profile (Oluoha et al., 2022). This stratification illustrates the potential of the framework to reduce insider threats while maintaining operational continuity. However, the reduced satisfaction among unknown users signals a trade-off between stringent enforcement and inclusivity, a tension that has been identified as a challenge in Zero-Trust literature (Scalise et al., 2024).

Efficiency of reinforcement learning models

The evaluation of AI performance metrics (Table 4) highlights the superiority of reinforcement learning in Zero-Trust SD-WAN applications. With convergence achieved in fewer epochs and adaptation speeds as low as 35 ms, reinforcement learning outperformed both supervised and unsupervised models in responsiveness and accuracy. The ROC-AUC score of 0.97 demonstrates its predictive precision, surpassing

benchmarks reported for anomaly detection in previous enterprise security frameworks. This indicates that reinforcement learning is particularly well-suited for dynamic environments where policies must be continuously recalibrated in response to emerging threats (Gautam, 2023). Moreover, the relatively moderate resource consumption suggests scalability in large SaaS deployments (Prasad et al., 2023). These findings reinforce emerging scholarship that positions reinforcement learning as a pivotal technology for autonomous, risk-aware decision-making in cybersecurity.

Implications for SaaS enterprises

From a practical perspective, the findings of this study carry significant implications for SaaS enterprises seeking to operationalize Zero-Trust. The integration of risk-aware agentic AI offers not only enhanced security outcomes but also efficient network performance and improved user experience (Narula et al., 2025). Enterprises can leverage such frameworks to reduce reliance on manual policy administration, thereby lowering operational overheads while maintaining high levels of compliance (Adekunle et al., 2023). The scalability of the solution makes it particularly relevant for organizations experiencing rapid expansion of their SaaS ecosystems. Importantly, the approach addresses one of the most pressing concerns in SaaS security the dynamic adaptation to evolving threats thus providing a future-ready solution for enterprise IT security strategies.

Limitations and future directions

While the study demonstrates strong outcomes, certain limitations must be acknowledged. The experiments were conducted within a simulated environment, which, despite realism, may not fully capture the heterogeneity and unpredictability of live enterprise networks. The reliance on synthetic attack scenarios, though diverse, may not encompass all emerging threat vectors, such as AI-driven polymorphic malware. Additionally, while reinforcement learning proved highly effective, its long-term scalability in multi-tenant SaaS environments requires further validation. Future research should extend this framework into real-world enterprise deployments, exploring integration with federated learning to preserve data privacy across distributed infrastructures. Investigating hybrid AI models that combine reinforcement

learning with adversarial training could also enhance resilience against sophisticated attacks.

Conclusion

This study demonstrated that integrating risk-aware agentic AI into Zero-Trust SD-WAN architectures provides a robust and adaptive security framework for SaaS enterprises. By leveraging reinforcement learning, the system achieved high detection rates, reduced false positives, and delivered rapid policy enforcement while sustaining strong network performance across varying traffic loads. The framework also balanced stringent security with user experience, offering seamless access to legitimate users while effectively mitigating insider risks and external threats. Importantly, the adaptability of reinforcement learning models ensured dynamic risk-aware decision-making, positioning the framework as a scalable and future-ready solution for SaaS environments. While the experimental results highlight the potential of this approach, further validation in real-world enterprise settings and exploration of hybrid AI techniques will be essential to fully realize its capabilities. Ultimately, the findings underscore that risk-aware agentic AI is not merely an enhancement to Zero-Trust SD-WAN, but a critical enabler for resilient, intelligent, and user-centric security in the evolving SaaS landscape.

References

- Adekunle, B. I., Chukwuma-Eke, E. C., Balogun, E. D., & Ogunsola, K. O. (2023). Developing a digital operations dashboard for real-time financial compliance monitoring in multinational corporations. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 9(3), 728-746.
- Alaali, H. (2025). DFAS Doctrine Compendium: Philosophical, Strategic, and Systemic Frameworks for AI Governance, Crisis Intervention, and Algorithmic Control.
- Chillapalli, N. T. R. (2022). Software as a service (SaaS) in e-commerce: The impact of cloud computing on business agility. *Journal Of Engineering And Computer Sciences*, 1(10), 7-18.

Das, S., & Mukherjee, S. (2024). Navigating cloud security risks, threats, and solutions for seamless business logistics. In *Emerging technologies and security in cloud computing* (pp. 252-275). IGI Global Scientific Publishing.

Gautam, M. (2023). Deep Reinforcement learning for resilient power and energy systems: Progress, prospects, and future avenues. *Electricity*, 4(4), 336-380.

Ghasemshirazi, S., Shirvani, G., & Alipour, M. A. (2023). Zero trust: Applications, challenges, and opportunities. *arXiv preprint arXiv:2309.03582*.

Hall-May, M., & Surridge, M. (2010, February). Resilient critical infrastructure management using service oriented architecture. In *2010 International Conference on Complex, Intelligent and Software Intensive Systems* (pp. 1014-1021). IEEE.

Koukaras, C., Hatzikraniotis, E., Mitsiaki, M., Koukaras, P., Tjortjis, C., & Stavrinides, S. G. (2025). Revolutionising Educational Management with AI and Wireless Networks: A Framework for Smart Resource Allocation and Decision-Making. *Applied Sciences*, 15(10), 5293.

Narula, S., Ghasemigol, M., Carnerero-Cano, J., Minnich, A., Lupu, E., & Takabi, D. (2025). Exploring AI Security: A Systematic Mapping Study. *IEEE Access*.

Oluoha, O. M., Odesina, A., Reis, O., Okpeke, F., Attipoe, V., & Orieno, O. H. (2022). A Unified Framework for Risk-Based Access Control and Identity Management in Compliance-Critical Environments.

Ouamri, M. A., Alharbi, T., Singh, D., & Sylia, Z. (2025). A comprehensive survey on software-defined wide area network (SD-WAN): principles, opportunities and future challenges. *The Journal of Supercomputing*, 81(1), 291.

Paul, B., & Rao, M. (2022). Zero-trust model for smart manufacturing industry. *Applied Sciences*, 13(1), 221.

Prasad, V. K., Dansana, D., Bhavsar, M. D., Acharya, B., Gerogiannis, V. C., & Kanavos, A. (2023). Efficient resource utilization in IoT and cloud computing. *Information*, 14(11), 619.

Radanliev, P., De Roure, D., Maple, C., Nurse, J. R., Nicolescu, R., & Ani, U. (2024). AI security and cyber risk in IoT systems. *Frontiers in Big Data*, 7, 1402745.

Reddy, A. (2025). Implementing Micro-Segmentation in Cloud Architectures: Zero Trust, Complete Security. *Famous Journal of computer science and Technology*, 2(7), 60-80.

Sarkar, S., Choudhary, G., Shandilya, S. K., Hussain, A., & Kim, H. (2022). Security of zero trust networks in cloud computing: A comparative review. *Sustainability*, 14(18), 11213.

Scalise, P., Boeding, M., Hempel, M., Sharif, H., Delloiacovo, J., & Reed, J. (2024). A systematic survey on 5G and 6G security considerations, challenges, trends, and research areas. *Future Internet*, 16(3), 67.