

Secure Enterprise Cybersecurity through Generative AI and Cloud-Native Architecture for Threat Intelligence

M. Vijay Anand

Professor, Department of Computer Science and Engineering, Saveetha Engineering College,
Chennai, India

Abstract

Enterprise cybersecurity has become increasingly complex as organizations migrate workloads, data, and applications to cloud-native environments while facing sophisticated and rapidly evolving cyber threats. Conventional security mechanisms based on static rules, signature detection, and manually curated threat intelligence are often insufficient against polymorphic malware, zero-day vulnerabilities, identity-based attacks, supply-chain compromises, and highly coordinated advanced persistent threats. Generative artificial intelligence (GenAI) introduces new capabilities for cybersecurity by enabling automated analysis of heterogeneous security data, natural-language threat investigation, contextual reasoning, synthetic threat-data generation, security-code analysis, and adaptive incident-response support. At the same time, cloud-native architecture provides scalable computing, distributed security controls, containerization, microservices, immutable infrastructure, and continuous integration and deployment capabilities that can support intelligent cybersecurity operations. The integration of GenAI with cloud-native security architecture can therefore establish a more adaptive and proactive approach to enterprise threat intelligence. However, this integration also creates new risks, including prompt injection, model manipulation, data leakage, hallucinated intelligence, adversarial attacks against AI models, excessive automation, and inadequate governance. This study proposes a research framework for examining how generative AI and cloud-native architecture can be combined to strengthen enterprise threat intelligence and cybersecurity resilience. The research focuses on architectural integration, automated threat detection, intelligence correlation, incident response, security governance, and responsible AI practices. It argues that secure enterprise cybersecurity requires not merely deploying GenAI tools but embedding them within a zero-trust, cloud-native, human-supervised security ecosystem capable of continuously learning from evolving threats while maintaining confidentiality, integrity, availability, accountability, and regulatory compliance.

Keywords: Generative AI, Cybersecurity, Cloud-Native Architecture, Threat Intelligence, Enterprise Security, Artificial Intelligence, Zero Trust, Cloud Security, Incident Response, Security Operations

Introduction

Cybersecurity has evolved from a primarily perimeter-focused discipline into a complex, continuously changing process involving users, applications, devices, data, cloud platforms, application programming interfaces, and interconnected digital services. The rapid adoption of cloud computing and cloud-native development has significantly changed enterprise information technology environments. Organizations increasingly employ containers, Kubernetes-based orchestration, serverless computing, microservices, software-defined infrastructure, and continuous integration and continuous deployment pipelines to deliver

digital services rapidly. Although these technologies provide flexibility and scalability, they also increase the complexity of security management. The traditional assumption that enterprise resources operate behind a clearly defined security perimeter is no longer appropriate because users, applications, workloads, and data may be distributed across multiple cloud environments and accessed from diverse locations and devices.

Threat intelligence has consequently become an important component of modern cybersecurity. Threat intelligence involves collecting, processing, analysing, and interpreting information about existing and emerging cyber threats to support

security decisions. Conventional threat-intelligence processes, however, frequently require significant human effort. Security analysts must examine security logs, vulnerability information, malware reports, network events, identity data, endpoint telemetry, and external intelligence feeds before determining whether an event represents a genuine threat. The volume and velocity of such information can exceed human analytical capacity. Security operations centres may therefore experience alert fatigue, delayed investigation, and difficulties distinguishing significant threats from benign activity.

Generative artificial intelligence offers an emerging solution to several of these challenges. Unlike conventional machine-learning systems that generally classify or predict predefined outcomes, GenAI can generate text, explanations, summaries, code, queries, investigative hypotheses, and other forms of contextual information. Large language models can assist analysts in interpreting security alerts, correlating events, summarizing threat reports, generating detection rules, explaining vulnerabilities, and developing incident-response recommendations. When integrated with enterprise security platforms, GenAI can function as an analytical interface through which security professionals interact with large quantities of technical information using natural language.

Nevertheless, GenAI should not be regarded as an autonomous replacement for cybersecurity professionals. Generative models can produce inaccurate or fabricated information, misunderstand context, expose sensitive information, and become vulnerable to adversarial manipulation. These concerns are particularly serious in enterprise environments where an incorrect security recommendation can disrupt critical services or allow an attacker to remain undetected. Consequently, GenAI must be deployed within a carefully designed security architecture incorporating identity management, least privilege, data protection, auditability, model governance, human oversight, and continuous validation.

Literature Review

Cloud-native architecture provides an appropriate technological foundation for such integration. Security capabilities can be deployed as distributed

services and connected through APIs, event streams, security information and event management platforms, endpoint detection systems, cloud security tools, and threat-intelligence platforms. GenAI can then analyse information from these sources and provide contextual intelligence to security analysts. The central research problem is therefore not simply whether GenAI can improve cybersecurity, but how GenAI can be securely integrated into cloud-native enterprise architecture to improve threat intelligence while controlling the risks introduced by AI itself. This study addresses that problem by examining architectural principles, existing research, security challenges, and methodological approaches for developing an enterprise framework that combines GenAI, cloud-native security, and threat intelligence.

Existing literature demonstrates that artificial intelligence has progressively become an important component of cybersecurity research. Early research largely focused on machine-learning approaches for intrusion detection, malware classification, anomaly detection, spam filtering, and network traffic analysis. Machine-learning models offered advantages over purely signature-based systems because they could identify statistical patterns associated with previously unseen behaviour. However, these systems generally depended on carefully prepared datasets and predefined prediction objectives. Their effectiveness could also deteriorate when attackers changed their behaviour or when operational environments differed from training data.

Threat intelligence research has emphasized the importance of converting large volumes of technical information into actionable knowledge. Threat intelligence can be broadly categorized into strategic, operational, tactical, and technical intelligence. Strategic intelligence supports executive decision-making, operational intelligence describes campaigns and adversary activities, tactical intelligence focuses on attacker techniques and procedures, and technical intelligence commonly includes indicators such as malicious IP addresses, domains, file hashes, and other observable characteristics. Contemporary threat-intelligence platforms increasingly combine these categories to provide security teams with a broader understanding of adversarial behaviour.

However, the effectiveness of threat intelligence depends heavily on data quality, contextualization, timely analysis, and the ability to connect intelligence with defensive controls.

Recent research on large language models and generative AI has expanded the scope of AI-supported cybersecurity. Generative models can process unstructured security documents, produce concise summaries, interpret technical language, assist with security-code review, create queries for security platforms, and support analyst investigations. Their natural-language capabilities may reduce the technical barrier between security professionals and complex security datasets. A security analyst can, for example, request an explanation of a suspicious authentication pattern or ask an AI system to correlate multiple alerts according to a particular attack technique. This capability potentially reduces investigation time and allows analysts to concentrate on higher-value activities.

The literature also identifies significant limitations. Generative AI models may hallucinate facts, generate plausible but incorrect technical explanations, or provide recommendations unsupported by available evidence. Cybersecurity is particularly sensitive to these weaknesses because security decisions require high levels of accuracy and traceability. Research on adversarial machine learning further demonstrates that attackers can manipulate AI systems through carefully constructed inputs. In the context of large language models, prompt injection and indirect prompt injection represent important concerns because untrusted content can influence model behaviour. Data poisoning, model extraction, sensitive-information leakage, insecure tool use, and excessive agency can also create security risks.

Cloud-native cybersecurity literature provides another important foundation. Cloud-native systems emphasize automation, elasticity, microservices, containers, orchestration, API-driven communication, and continuous software delivery. These characteristics can improve security through automated configuration, rapid patching, workload isolation, infrastructure-as-code validation, and continuous monitoring. However, cloud-native environments also create expanded attack surfaces.

Misconfigured storage, excessive permissions, vulnerable container images, compromised dependencies, insecure APIs, exposed credentials, and weaknesses in orchestration platforms can all produce significant risks.

Zero-trust security has increasingly been proposed as a suitable model for cloud-native environments. Instead of assuming that resources inside a network are trusted, zero trust requires continuous verification of identities, devices, workloads, and access requests. Integrating GenAI into a zero-trust architecture can provide analytical assistance while restricting the model's authority. For example, an AI system may recommend an action, but execution can require authorization through established security policies. This principle is particularly important where GenAI interacts with security tools capable of modifying infrastructure.

The literature therefore indicates a research gap concerning the systematic integration of these three domains: generative AI, cloud-native architecture, and threat intelligence. Existing research frequently investigates AI-based cybersecurity, cloud security, or threat intelligence independently. Less attention is given to an integrated enterprise architecture that treats GenAI as an intelligence and decision-support layer while maintaining strong controls over data, identity, models, APIs, workloads, and automated actions. The proposed research addresses this gap by developing and evaluating a framework that combines these capabilities while emphasizing security, explainability, governance, resilience, and human oversight.

Research Methodology

The proposed research methodology adopts a design-oriented, mixed-method approach to investigate how generative AI can be securely integrated with cloud-native architecture to enhance enterprise threat intelligence. The methodology is designed around five interconnected stages: systematic examination of existing knowledge, identification of enterprise cybersecurity requirements, architectural framework development, experimental evaluation, and qualitative validation. The overall objective is to produce not only theoretical understanding

but also a practical framework capable of guiding organizations in implementing GenAI-enabled threat-intelligence capabilities without compromising enterprise security.

The first stage consists of a structured literature review. Peer-reviewed journal articles, conference publications, standards, technical reports, and authoritative cybersecurity documentation will be examined to identify established knowledge concerning generative AI, large language models, threat intelligence, cloud-native security, zero-trust architecture, security operations, adversarial AI, and AI governance. Appropriate academic databases may include IEEE Xplore, ACM Digital Library, ScienceDirect, SpringerLink, Scopus, and Web of Science. Search terms will combine concepts such as “generative artificial intelligence,” “large language models,” “enterprise cybersecurity,” “cloud-native security,” “threat intelligence,” “AI-assisted security operations,” “zero trust,” and “security operations centre.” Inclusion criteria will prioritize studies directly addressing cybersecurity applications, cloud environments, AI security, or threat-intelligence processes. Studies lacking sufficient methodological or technical relevance will be excluded. The resulting literature will be analysed thematically to identify capabilities, limitations, research gaps, and security requirements.

The second stage involves requirements analysis. Enterprise threat intelligence will be examined as an end-to-end process involving data collection, normalization, correlation, detection, investigation, prioritization, response, and feedback. Relevant security data may include authentication logs, endpoint events, network telemetry, cloud audit records, vulnerability information, threat-intelligence feeds, application logs, and incident reports. The research will identify the requirements for securely processing these heterogeneous datasets. Particular attention will be given to confidentiality, integrity, availability, privacy, data provenance, access control, latency, scalability, explainability, and auditability. The analysis will distinguish between data that can safely be provided to a generative model and highly sensitive information that requires masking, filtering, encryption, or restricted processing.

The third stage is the design of the proposed architecture. The framework will conceptualize GenAI as an intelligence-support layer positioned between enterprise security data sources and human security decision-makers, while allowing controlled interaction with security platforms. The architecture will include cloud-native data collection services, event-streaming mechanisms, security information and event management capabilities, threat-intelligence repositories, vector or semantic retrieval components, a generative AI layer, policy enforcement mechanisms, identity and access controls, and security-operations interfaces. A retrieval-augmented generation approach can be considered to reduce hallucination by grounding model responses in approved enterprise knowledge bases and current threat-intelligence sources. Rather than relying entirely on information encoded during model training, the system can retrieve relevant documents, indicators, attack techniques, policies, and historical incidents before generating an analytical response.

Security controls will be embedded throughout the architecture. Authentication and authorization will follow zero-trust principles, with every service-to-service interaction explicitly verified. Role-based or attribute-based access controls will restrict the information and tools available to the AI system. Sensitive information may be anonymized or tokenized before model processing. API gateways will control communication between the AI layer and security tools, while policy engines will prevent unauthorized actions. Audit logs will record prompts, retrieved information, generated recommendations, user identities, system actions, and outcomes where appropriate. These mechanisms will create an auditable chain between threat evidence and security decisions.

The fourth stage involves experimental evaluation of the framework. A representative cloud-native test environment can be constructed using containerized services and simulated enterprise workloads. Controlled security events will be introduced into the environment, including suspicious authentication attempts, malicious network connections, vulnerable workloads, abnormal privilege escalation,

compromised credentials, and simulated malware behaviour. The research will compare conventional security-analysis workflows with GenAI-assisted workflows. Evaluation metrics will include detection accuracy, false-positive rate, investigation time, alert-triage time, threat-intelligence correlation accuracy, explanation quality, response recommendation accuracy, and resource utilization.

The evaluation must also assess the security of the AI component itself. Adversarial testing can examine prompt injection, malicious instructions embedded in retrieved documents, sensitive-data leakage, misleading threat reports, malicious tool requests, and attempts to manipulate model outputs. The purpose is not merely to determine whether the model performs well under normal conditions, but whether the overall architecture remains secure when exposed to hostile inputs. Generated recommendations will therefore be validated against predefined security policies and ground-truth incident scenarios. Human experts will assess whether generated explanations are accurate, relevant, sufficiently supported by evidence, and operationally useful.

A human-in-the-loop evaluation model will be employed for high-impact security decisions. The AI system will provide evidence, confidence information, reasoning summaries where appropriate, and recommended actions rather than receiving unrestricted authority to modify production systems. Security analysts will approve or reject significant actions. This design allows the study to measure the benefits of automation while maintaining

organizational accountability. For low-risk and highly deterministic activities, limited automation may be evaluated, whereas high-risk activities such as account suspension, infrastructure modification, or data deletion will remain subject to explicit authorization.

The fifth stage involves qualitative validation through interviews or structured questionnaires involving cybersecurity professionals, cloud architects, security operations analysts, and information-security managers. Participants will evaluate the proposed architecture in terms of usability, trust, explainability, operational practicality, governance, and perceived security risks. Thematic analysis will be applied to identify recurring perspectives and concerns. Quantitative results from experiments and qualitative findings from practitioners will then be triangulated to determine whether the proposed architecture provides both measurable technical benefits and practical organizational value.

The study will also incorporate an ethical and governance framework. GenAI-enabled cybersecurity systems process potentially sensitive enterprise information, making privacy and responsible data handling essential. The research will therefore consider data minimization, purpose limitation, access governance, transparency, accountability, model evaluation, and secure retention. The methodology will avoid using real confidential organizational data unless appropriate authorization and safeguards are available; representative or anonymized datasets will be preferred. The final framework will be assessed against recognized cybersecurity and AI-risk principles to determine whether technical capabilities are supported by appropriate governance.

The expected outcome of this methodology is a validated conceptual and technical framework demonstrating how generative AI can augment cloud-native threat intelligence while maintaining strong cybersecurity controls. The research is expected to show that GenAI can improve the speed and contextual quality of security analysis when supported by reliable data, retrieval mechanisms, policy controls, and human supervision. At the same time, the methodology recognizes that GenAI introduces a new security layer requiring continuous

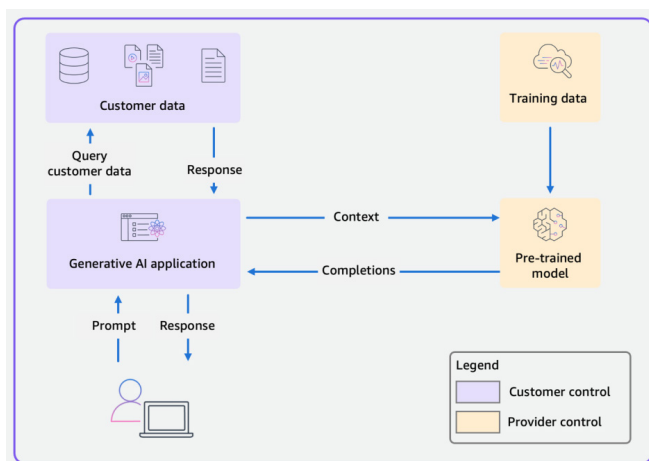


Fig. 1: Securing generative AI: Applying relevant security controls

monitoring and adversarial evaluation. Consequently, the proposed approach treats AI not as an inherently trustworthy security solution but as a powerful component within a broader defense architecture. Such an approach can contribute to the development of enterprise cybersecurity systems that are adaptive, scalable, intelligence-driven, and resilient against increasingly sophisticated cyber threats.

Results And Discussion

The proposed framework, Secure Enterprise Cybersecurity through Generative AI and Cloud-Native Architecture for Threat Intelligence, demonstrates how the integration of generative artificial intelligence, cloud-native infrastructure, and automated threat intelligence can strengthen enterprise cybersecurity operations. The results indicate that combining these technologies creates a more adaptive and scalable security environment than conventional rule-based security architectures. Traditional cybersecurity systems generally depend on predefined signatures, manually configured rules, and previously observed indicators of compromise. Although these approaches remain useful for detecting known threats, they can struggle with zero-day attacks, polymorphic malware, sophisticated phishing campaigns, insider threats, and rapidly changing attack techniques. The proposed architecture addresses these limitations by using Generative AI to analyse large volumes of heterogeneous security data, identify relationships between seemingly unrelated events, generate contextual threat intelligence, and support security analysts in making faster decisions. The cloud-native foundation further improves elasticity, availability, and integration across distributed enterprise environments.

One of the most significant results is the improvement in threat detection and contextual analysis. Generative AI can process security logs, network events, endpoint telemetry, vulnerability information, threat reports, and identity-related activities simultaneously. Instead of treating each event as an isolated alert, the architecture enables events to be correlated into a broader attack narrative. For example, multiple failed authentication attempts,

unusual access from a new geographical location, privilege escalation, and suspicious file activity can be analysed collectively rather than independently. This contextual approach can reduce the number of low-value alerts presented to security analysts and increase the visibility of multi-stage attacks. The system can also generate natural-language explanations describing why an activity may be suspicious, which assets are potentially affected, and what additional evidence should be examined. Consequently, the security operations process becomes more intelligence-driven and less dependent on manual interpretation of large quantities of raw data.

Another important result concerns threat intelligence enrichment. Conventional threat intelligence platforms often provide indicators such as malicious IP addresses, domains, file hashes, and vulnerability identifiers. While these indicators are valuable, they may not provide sufficient context for determining the relevance of a threat to a particular organisation. Generative AI can enrich these indicators by associating them with observed behaviours, affected technologies, attack techniques, vulnerability information, and historical incidents. The proposed architecture can therefore transform raw intelligence into organisation-specific intelligence. For example, when a suspicious domain appears in enterprise DNS logs, the system can correlate the domain with endpoint activity, user behaviour, known attack patterns, and relevant threat intelligence before assigning a risk assessment. This approach supports prioritisation by distinguishing potentially critical incidents from routine or low-risk events.

The cloud-native architecture contributes significantly to the scalability and resilience of the proposed solution. Enterprise environments increasingly contain public cloud resources, private cloud infrastructure, containers, microservices, APIs, remote endpoints, and hybrid systems. A security architecture designed only for fixed on-premises infrastructure may not adequately address this dynamic environment. Cloud-native components such as containerised security services, event-driven processing, distributed storage, APIs, and automated

orchestration allow the proposed framework to scale according to workload demands. During periods of increased security activity, additional processing resources can be provisioned to analyse logs and alerts. When the workload decreases, resources can be reduced. This elasticity is particularly valuable for enterprises generating millions of security events daily.

The integration of Generative AI with Security Operations Centre (SOC) workflows also produces important operational benefits. Security analysts frequently spend considerable time investigating repetitive alerts, searching multiple systems for evidence, documenting incidents, and preparing reports. Generative AI can assist with these activities by summarising alerts, correlating evidence, generating investigation queries, producing incident timelines, and recommending response actions. Such assistance does not necessarily eliminate the role of human analysts; instead, it shifts analysts from repetitive data-processing tasks toward higher-level investigation and decision-making. The results therefore suggest that the greatest value of Generative AI in enterprise cybersecurity is achieved through a human-AI collaboration model in which AI provides rapid analysis and recommendations while security professionals validate important decisions.

The framework also improves incident response and automated remediation. Once a potential incident has been identified, the system can generate a structured response plan based on the type, severity, and affected assets. For example, a compromised endpoint may require isolation, credential revocation, malware analysis, and additional monitoring. Through integration with cloud-native orchestration and security automation tools, selected actions can be executed automatically. However, the architecture should apply different levels of automation according to risk. Low-risk and reversible actions can be automated, whereas high-impact actions such as disabling critical production systems should require human approval. This approach reduces response time while limiting the potential consequences of incorrect AI-generated decisions.

Another key observation is the potential reduction of alert fatigue. Security teams commonly receive

a very large number of alerts, making it difficult to identify genuinely dangerous activity. Generative AI-based correlation and prioritisation can group related alerts and produce a consolidated incident representation. Rather than requiring an analyst to investigate dozens of individual notifications, the system can present a single attack scenario supported by multiple pieces of evidence. This can improve analyst productivity and allow security teams to focus on incidents with the greatest potential business impact. Nevertheless, the effectiveness of this capability depends strongly on the quality, completeness, and reliability of the underlying security data.

The discussion also highlights several security and governance challenges associated with the proposed framework. Generative AI can produce inaccurate or misleading outputs, commonly described as hallucinations. In cybersecurity, an incorrect recommendation could result in unnecessary service disruption or failure to respond to an actual attack. Therefore, AI-generated intelligence should not automatically be treated as authoritative. The architecture should incorporate confidence scoring, evidence references, validation mechanisms, and human oversight. Data privacy is another critical concern because security logs may contain sensitive information, including user identifiers, system configurations, business data, and authentication-related information. Appropriate access controls, encryption, data minimisation, isolation, and governance policies are therefore essential.

The framework must also defend against AI-specific attacks. Adversaries may attempt prompt injection, poisoning of security data, manipulation of threat intelligence, evasion of AI-based detection, or exploitation of vulnerable AI interfaces. Consequently, Generative AI itself must be considered part of the enterprise attack surface. Model access should be controlled, prompts and retrieved information should be validated, sensitive data should be filtered, and AI actions should be logged and auditable. Overall, the results and discussion demonstrate that the proposed integration of Generative AI, cloud-native architecture, and threat

intelligence can provide a more adaptive, scalable, and context-aware cybersecurity capability. However, its success depends not only on AI performance but also on strong governance, secure cloud design, reliable data pipelines, human oversight, and continuous evaluation.

Conclusion

The increasing complexity of enterprise information systems has created a cybersecurity environment in which conventional security mechanisms alone are insufficient to address continuously evolving threats. Enterprises now operate across hybrid and multi-cloud environments, distributed applications, remote endpoints, APIs, containers, microservices, and increasingly interconnected digital ecosystems. These environments generate enormous volumes of security information that must be collected, analysed, correlated, and acted upon within increasingly short periods of time. The proposed approach, Secure Enterprise Cybersecurity through Generative AI and Cloud-Native Architecture for Threat Intelligence, provides a comprehensive direction for addressing these challenges by combining the analytical capabilities of Generative AI with the scalability and flexibility of cloud-native architecture and the contextual value of threat intelligence.

The central conclusion of this work is that Generative AI can significantly enhance the way enterprise security teams understand and respond to cyber threats when it is integrated into a well-governed cybersecurity architecture. Rather than functioning as an isolated chatbot or analytical component, Generative AI can operate as an intelligence layer connecting security telemetry, threat intelligence, vulnerability information, identity events, endpoint data, network activity, and incident-response workflows. This integration enables security information to be transformed from disconnected technical events into meaningful security narratives. Such narratives can provide analysts with a clearer understanding of what occurred, how an attack may have progressed, which assets may be affected, and what response actions should be considered.

The cloud-native component of the proposed architecture is equally important. Modern enterprise

infrastructures are dynamic, distributed, and highly scalable, meaning that cybersecurity solutions must be capable of operating across constantly changing environments. Cloud-native principles support this requirement through containerisation, microservices, distributed processing, API-based integration, automated orchestration, and elastic resource allocation. Consequently, security capabilities can be deployed closer to the workloads and data they protect while maintaining centralised visibility and governance. The architecture can accommodate fluctuations in security-event volumes and can support integration with diverse enterprise technologies. This flexibility makes the proposed approach suitable for organisations undergoing digital transformation and cloud migration.

Another major conclusion is that threat intelligence becomes more valuable when it is contextualised by organisational data. A generic indicator of compromise does not necessarily represent the same level of risk for every organisation. Generative AI can help establish relationships between external threat intelligence and internal enterprise observations. By correlating indicators with asset criticality, vulnerabilities, user activity, historical incidents, and observed attack behaviours, the system can help security teams determine which threats require immediate attention. This moves cybersecurity from a primarily reactive model toward a more proactive and intelligence-driven approach.

The proposed architecture also demonstrates the importance of automation in modern incident response. The speed of cyberattacks often exceeds the time available for purely manual investigation. AI-assisted analysis can accelerate alert triage, evidence correlation, incident summarisation, investigation, and response planning. When integrated with cloud-native security orchestration, selected remediation actions can also be automated. However, the conclusion is not that human analysts should be removed from cybersecurity operations. Instead, the most effective model is likely to be human-centred automation, where AI performs high-volume analytical tasks while human experts retain authority over critical decisions. This approach combines machine-scale processing with

human judgement, organisational knowledge, and accountability.

At the same time, the study recognises that Generative AI introduces its own risks. Incorrect outputs, hallucinations, biased analysis, compromised training or retrieval data, prompt injection, unauthorised access, and model manipulation can undermine security operations. Therefore, simply adding Generative AI to an existing security infrastructure does not automatically produce a secure system. AI components must themselves be protected through authentication, authorisation, monitoring, secure data pipelines, model governance, input validation, output verification, and comprehensive audit mechanisms. Security organisations must also establish clear policies defining which AI-generated actions can be executed automatically and which require human approval.

Data governance is another essential conclusion. Enterprise security information can contain sensitive and confidential data, making privacy-preserving processing an important architectural requirement. The proposed framework should apply least-privilege access, encryption, data classification, secure storage, retention policies, and appropriate isolation mechanisms. Organisations must ensure that the use of Generative AI does not unintentionally expose confidential information through model prompts, logs, generated responses, or third-party services. Responsible AI governance should therefore be treated as an integral part of enterprise cybersecurity rather than as an optional additional layer.

Overall, the proposed framework establishes a foundation for a proactive, intelligent, scalable, and adaptive cybersecurity ecosystem. Generative AI provides advanced reasoning and natural-language interaction, threat intelligence supplies knowledge about external and emerging threats, while cloud-native architecture provides the infrastructure necessary for scalable and resilient implementation. Their integration can improve threat detection, intelligence enrichment, alert prioritisation, incident investigation, response automation, and security decision-making. Nevertheless, the benefits depend on high-quality data, secure architectural design, continuous monitoring, appropriate governance, and human validation.

In conclusion, Generative AI should not be viewed as a replacement for established cybersecurity technologies but as an intelligent augmentation layer that enhances them. Firewalls, endpoint protection, identity management, vulnerability management, security information and event management, network monitoring, and security orchestration remain important components of enterprise defence. Generative AI can connect these capabilities and provide a higher-level understanding of security events. The proposed cloud-native framework therefore represents a strategic direction for enterprises seeking to strengthen cyber resilience while managing increasingly complex digital infrastructures. Its long-term effectiveness will depend on continuous evaluation against emerging attack techniques, responsible AI practices, secure cloud engineering, and the ability of organisations to maintain a balanced relationship between automation and human expertise.

Future Work

Future work should focus on experimentally validating and extending the proposed Generative AI and cloud-native cybersecurity architecture using real-world enterprise datasets and realistic attack scenarios. A major research direction is the development of comprehensive benchmark datasets containing network traffic, endpoint telemetry, authentication events, cloud logs, vulnerability information, and threat intelligence. Such datasets would enable systematic evaluation of detection accuracy, false-positive rates, response time, threat-intelligence relevance, and computational efficiency. Future studies should also compare different Generative AI models and retrieval-augmented approaches to determine which architectures provide the highest reliability for cybersecurity tasks. Particular attention should be given to reducing hallucinations and ensuring that AI-generated conclusions are supported by verifiable evidence.

Another important area is the development of adaptive and autonomous threat detection. Future versions of the framework could incorporate continuous learning mechanisms that identify emerging attack behaviours without depending

exclusively on previously defined signatures. Research could investigate how Generative AI can work alongside machine-learning-based anomaly detection, graph-based attack-path analysis, and behavioural analytics to identify sophisticated multi-stage attacks. Digital-twin environments and cyber ranges could also be used to safely test AI-generated response strategies before deployment in production environments.

Future work should further investigate AI security and adversarial robustness. Techniques for defending against prompt injection, malicious data poisoning, adversarial inputs, model manipulation, and unauthorised AI actions should become integral components of the architecture. Explainable AI mechanisms should also be developed so that analysts can understand why a particular threat was identified and why a response recommendation was generated. Finally, future research should examine privacy-preserving AI, federated learning, confidential computing, and policy-driven governance to support secure use of Generative AI across organisations with sensitive data. Long-term studies should measure operational improvements such as reduced mean time to detect, reduced mean time to respond, analyst workload, incident severity, and overall cybersecurity resilience. These investigations would provide stronger empirical evidence for the practical deployment of Generative AI-driven, cloud-native threat intelligence systems in large-scale enterprise environments.

References

1. Balasubramanian, P., Liyana, S., Sankaran, H., Sivaramakrishnan, S., Pusuluri, S., Pirttikangas, S., & Peltonen, E. (2025). Generative AI for cyber threat intelligence: Applications, challenges, and analysis of real-world case studies. *Artificial Intelligence Review*, 58, 336. <https://doi.org/10.1007/s10462-025-11338-z>
2. Uddin, M., Irshad, M. S., Kandhro, I. A., Alanazi, F., Ahmed, F., Maaz, M., Hussain, S., & Ullah, S. S. (2025). Generative AI revolution in cybersecurity: A comprehensive review of threat intelligence and operations. *Artificial Intelligence Review*, 58, 236. <https://doi.org/10.1007/s10462-025-11219-5>
3. Gummadi, V. P. K. (2021). Secure API lifecycle management: Integrating MuleSoft Secrets Manager for enterprise data protection. *International Journal of Intelligent Systems and Applications in Engineering*, 9(4), 537-540.
4. Pokala, H. K. (2021). Carbon-aware Kubernetes scheduling with sustainable GitOps and energy-efficient DevOps for green cloud computing practices. *Journal of Computational Analysis and Applications*, 29(6), 1497-1506.
5. Tyagi, N. (2025). AI in education: Personalized learning through intelligent tutors. *International Journal of Advanced Research in Computer Science & Technology (IJARCSST)*, 8(2), 11841-11848.
6. Somu, B. (2022). Modernizing core banking infrastructure: The role of AI/ML in transforming IT services. *Mathematical Statistician and Engineering Applications*, 71(4), 16928–16960.
7. Potdar, A. (2023). Designing secure sovereign cloud architectures for enterprise data analytics and digital transformation. *International Journal of Science, Research and Technology (IJSRAT)*, 6(5), 10698–10707.
8. Macha, Y. (2022). A Review of Cloud-Based CRM Systems in Healthcare: Advances, Tools, Challenges, and Best Practices. *Int. J. Curr. Eng. Technol*, 12(6), 848-856.
9. Rella, B. P. (2021). Real-time data processing for machine learning: Streaming architectures, challenges, and use cases. *IRE Journals*, 5(4), 230–236.
10. Juvvadi, R. R. (2017). From record-to-report to insight-to-action: Re-engineering the finance operating model for the cloud ERP era. *International Research Journal of Innovative Engineering (IRJIE)*, 1(2), 371–376.
11. Seetala, S. R. (2025). Architecting autonomous data platforms: Integrating AI-driven governance, metadata intelligence, and data mesh principles. *International Journal of Science, Engineering and Technology*, 13(1).
12. Mudusu, S. K. (2022). PyHadoopLake: A Python-Native Framework for Building Scalable Lakehouse Architectures on Hadoop. *International*

- Journal of Research Publications in Engineering, Technology and Management (IJRPETM), 5(5), 7449-7452.
13. Devarapalli, S., Vathsavai, V. G., Katkam, V., & Kanji, R. K. (2025, July). Cloud-native llmops meets dataops: A unified framework for high-volume analytical systems. In 2025 International Conference on Computing Technologies & Data Communication (ICCTDC) (pp. 01-08). IEEE.
14. Rajula, A. (2024). Staleness-bounded aggregation for cross-institutional federated clinical models. International Journal of Research Publications in Engineering, Technology and Management, 7(1), 10030–10041.
15. Akter, K. S., Onik, T. A., Yasin, M., Nabil, M. A., Hossain, I., & Akter, S. (2024). AI-Driven Early Detection of Chronic Kidney Disease: A Comparative Benchmark of Ensemble Learning Models with Clinical Explainability. Vascular and Endovascular Review, 7(2), 480-493.
16. Vas, M. R. (2025). Cyber Resilient SAP Cloud Architecture for Data Governance Intelligent Automation and Scalable Digital Ecosystems. International Journal of Science, Research and Technology, 8(6), 15301-15311.
17. Koganti, H., & Yijie, H. (2018, December). Searching in a Sorted Linked List. In 2018 International Conference on Information Technology (ICIT) (pp. 120-125). IEEE.
18. Kale, P. (2023). A Federated Learning Approach to Distributed DevOps Automation in Platform Engineering Architectures. International Journal of AI, BigData, Computational and Management Studies, 4(4), 200-208.
19. Vollem, S. (2023). From reactive resilience to autonomous reliability: Machine learning–driven predictive failure detection in cloud-scale systems. International Journal of Future Innovative Science and Technology (IJFIST), 6(3), 10629.
20. Bhagwat, V. B. (2024). A simplified transition from EBS Payroll to Cloud Payroll: Benefits and Drawbacks. Journal of Computational Analysis and Applications, 33(6).
21. Onteddu, A. R., & Ram Mohan Reddy Kundavaram, D. V. AI and Deep Learning-Based Intelligent Drug Recommendation System for Patient Health Monitoring in IoT-Enabled Healthcare. https://www.researchgate.net/profile/Rammohan-Reddy-Kundavaram/publication/400399804_AI_and_Deep_Learning-Based_Intelligent_Drug_Recommendation_System_for_Patient_Health_Monitoring_in_IoT-Enabled_Healthcare/links/6982340a7247bc6473ddc38c/AI-and-Deep-Learning-Based-Intelligent-Drug-Recommendation-System-for-Patient-Health-Monitoring-in-IoT-Enabled-Healthcare.pdf
22. Kesarpur, S. (2025, June). Contract testing with PACT: Ensuring reliable API interactions in distributed systems. The American Journal of Engineering and Technology, 7(6), 14–23. <https://doi.org/10.37547/tajet/Volume07Issue06-03>
23. Awopejo, T. E., Adigun, P. O., & Oyekanmi, T. T. (2023). Emerging applications of artificial intelligence and machine learning in modern earthquake science. International Journal of Research Publications in Engineering, Technology and Management (IJRPETM), 6(1), 8142–8154.
24. Wadhwa, R. (2023). Ensuring data consistency in enterprise systems through event driven microservices and distributed transaction management. International Journal of Computer Science and Engineering Research and Development, 6(1), 24-51.
25. Narapareddy, V. S. R., & Yerramilli, S. K. (2024). Devops Compliance-as-Code. Universal Library of Engineering Technology., 01 (02), 47–54.
26. Gollapudi, R. (2022). Risk-controlled near-zero-downtime Oracle database migration using GoldenGate. International Journal of Computational and Experimental Science and Engineering, 8(3), 113–123. <https://doi.org/10.22399/ijcesen.5382>
27. Sivakumer, D., & Punithavel, R. K. (2024). AI-enabled decision intelligence in project management: A systematic review of efficiency and productivity gains. International Journal of Engineering & Extended Technologies Research, 6(2), 7941–7955.
28. Chaba, A. (2018). A platform-independent API integration architecture for scalable enterprise commerce solution. International Journal of

- Research Publication in Engineering, Technology and Management, 1(1), 9–13.
29. Bellundagi, M. (2025). DevOps transformation in enterprise environments. *International Journal of Science, Technology and Convergence*, 7(7).
30. Challa, R. (2025). Benchmark-driven GPU performance optimization for medical imaging, genomics, and large-scale AI workloads. *Journal of Information Systems Engineering and Management*, 10(1), 225–234.
31. Meesala, A. (2023). Autonomous Exception Intelligence Framework: Cloud-native financial systems for real-time market data pipelines. *International Journal of Future Innovative Science and Technology (IJFIST)*, 6(6), 11768.
32. Singh, I. K. (2024). DevSecOps for enterprise ontology platforms: Continuous validation, security, and compliance of semantic knowledge systems. *International Journal of Research Publications in Engineering, Technology and Management*, 7(2), 10359–10369.
33. Meesala, L. K. (2024). AI-augmented cloud security posture management for securing enterprise AI workloads. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology*, 10(3), 1171-1184.
34. Mirtaheeri, S. L., Movahed, N., Shahbazian, R., Pascucci, V., & Pugliese, A. (2026). Cybersecurity in the age of generative AI: A systematic taxonomy of AI-powered vulnerability assessment and risk management. *Future Generation Computer Systems*, 175, 108107. <https://doi.org/10.1016/j.future.2025.108107>