

Evolving Cyber Threat Intelligence: A Systematic Review and Comparative Analysis

Hilalah Alturkistani¹, Abdul Ghafar Jaafar¹, Suriayati Chuprat¹, Nur Hanis Sabrina Suhaimi²

¹Faculty of Artificial Intelligence, Universiti Teknologi Malaysia, Kuala Lumpur, Malaysia

²Center of Cybersecurity, Faculty of science information and Technology, National University of Malaysia

DOI: <https://doi.org/10.51244/IJRSI.2026.1307000326>

Received: 05 August 2026; Accepted: 10 August 2026; Published: 18 August 2026

ABSTRACT

To improve cybersecurity across industries, Cyber Threat Intelligence (CTI) is becoming increasingly crucial. This systematic review explores how CTI practices are evolving in response to advancements in Artificial Intelligence (AI), particularly in the context of Large Language Models (LLMs). We examined 61 peer-reviewed studies using the PRISMA methodology, which demonstrates a strict selection procedure founded on specified inclusion, exclusion, and quality standards. This approach aligns with the scope of similar systematic reviews in the field of cyber threat intelligence. The review provides a comparative synthesis of CTI research capabilities across threat detection and prediction, attribution, forecasting, and automated reporting. We classify these approaches into three categories: conventional methods, those enhanced by AI and Machine Learning, and those based on LLMs. Our findings indicate that LLMs offer significant advantages in contextual reasoning, processing unstructured threat intelligence, and generating actionable mitigation plans. However, challenges such as model explainability, data privacy, system interoperability, and standardization impede their integration into operational environments. In addition to highlighting the potential and practical limitations of LLMs in CTI, this study identifies research gaps and proposes methods to create scalable, secure, and flexible CTI systems that support real-time cyber defense.

Keywords: Cyber Threat Intelligence; Artificial Intelligence; Cybersecurity; Threat Detection; Threat Prediction; Large Language Models

INTRODUCTION

Everyday life has been profoundly altered by the development of information and communication technologies, exemplified by the Internet. The widespread use of Internet technology has increased cybersecurity vulnerabilities while also making things more convenient. System security has declined due to the increased complexity introduced by advances in IT infrastructure and interconnected systems [1]. From casual hacking to well-planned campaigns by actors with financial or political motives, the nature of cyberattacks has changed drastically. Sophisticated malware, such as Advanced Persistent Threats (APTs), poses significant challenges because they often evade firewalls and antivirus programs, which still primarily rely on static signatures [2]. According to a study [3] that highlights the growing involvement of organized crime and nation-state actors in strategic cyber operations, organizations still face increasing risks despite significant cybersecurity investments. Cybersecurity has become a strategic issue rather than merely a technical priority. According to a PwC survey, 77% of CEOs consider cybersecurity a major threat to growth, which suggests that cybersecurity is becoming increasingly important at the executive level [5]. Despite the existence of organized cybersecurity frameworks, many remain reactive and struggle with the accuracy of proactive detection [4]. Given this, CTI has become a vital tool for enabling more proactive and informed cybersecurity solutions.

In 2011, Hutchins et al. [74] defined CTI as the “systematic gathering, analysis, and sharing of threat data in order to generate useful intelligence that focuses on understanding adversary behavior to support proactive cybersecurity decision-making.” Therefore, CTI refers to any information that supports cybersecurity decision-

making, ranging from structured indicators of compromise to unstructured reports of phishing activity or vulnerability exploits [3]. Sources of CTI span both internal systems (e.g., firewall logs, honeynets) and external channels (e.g., OSINT, commercial feeds, government advisories). Structured data is relatively easy to parse using conventional tools, whereas unstructured sources, such as social media and forums, require advanced natural language processing and machine learning techniques to extract actionable insights promptly [51]. By leveraging threat-informed information such as Tactics, Techniques, and Procedures (TTPs) and malicious IP addresses, CTI facilitates early detection, attacker profiling, and defensive system configuration [6]. The recent integration of Artificial Intelligence (AI), particularly Machine Learning (ML) and Deep Learning (DL), has improved CTI capabilities by allowing systems to automatically learn threat patterns and increase detection accuracy [9,10]. These innovations set the stage for future breakthroughs, particularly through the use of LLMs, which are examined in detail in this study. CTI is emerging as a key component of contemporary cybersecurity defense strategies, shifting from reactive analysis toward predictive, data-driven intelligence.

Study Objectives and Contribution

To successfully address the inherent difficulties of combining CTI data from heterogeneous sources, this SLR examines recent advancements in CTI methods, techniques, objectives, data handling, limitations, CTI tasks, and evaluation. The ultimate goal is to determine how well LLMs that can process both structured and unstructured CTI data and reports can advance threat detection, recognition, and mitigation while also improving CTI operations in areas such as data collection from heterogeneous sources, filtering, analysis, visualization, and the production of actionable steps. This research's main contribution is the presentation of cutting-edge ML and AI methods, with a focus on LLMs. This research proposes a unified framework that integrates various CTI tasks, traditionally managed by separate systems, into a single, adaptable pipeline. By utilizing a fine-tuned large language model, this framework enables end-to-end automation across core CTI functions, including data processing, threat analysis, and the generation of mitigation strategies. This integration aims to address long-standing challenges in scalability, semantic understanding, and real-time responsiveness that are often encountered in conventional and machine-learning-based CTI approaches.

The scope of this work primarily focuses on progress over the past 6 years, particularly in ML-, AI-, and LLM-based CTI methods. The main goal of this study is to provide a reference for scholars and industry professionals interested in developing data-driven automated cybersecurity systems to enhance CTIs. Key contributions of this research include:

1. Categorize CTI research into three groups: Conventional CTI methods, AI in CTI, and LLMs in CTI processing, allowing for effective comparison and analysis of current methodologies.
2. Provide a comprehensive list of state-of-the-art CTI methods, detailing their contributions, study objectives, and limitations.
3. Analyze the adaptability of various ML, AI, and LLM paradigms in real-world cybersecurity contexts, highlighting their potential to drive a shift towards proactive cyber defense.
4. Discuss the challenges of existing CTI techniques, the pros and cons of various AI methods, and the transformative impact of LLMs on analyzing unstructured CTI, paving the way for future research.

The structure of this article is as follows: Section 1 describes how primary research studies in the SLR are chosen. Section 2 presents a comparative analysis of the relevant work. To identify research gaps, Section 3 compares studies, discusses trends, and summarizes traditional, AI/ML, and LLM-based CTI methods. An improved LLM-based CTI is developed in Section 4. To recommend future research directions, Section 5 examines obstacles and unresolved problems.

MATERIALS AND METHODS

This review follows the systematic review guidelines established by Kitchenham et al. [7] and the PRISMA (Preferred Reporting Items for Systematic Reviews and Meta-Analyses) framework [8] to ensure transparency, replicability, and methodological rigor. A structured and iterative process was applied across four main stages: planning, identification, screening, and analysis. To construct a robust search strategy, we selected seven major academic databases and repositories known for cybersecurity and AI research: IEEE Xplore, ACM Digital

Library, ScienceDirect, Elsevier, Web of Science (WoS), MDPI, and leading university-affiliated preprint repositories such as arXiv. These databases were chosen for their breadth, peer-reviewed content, and relevance to CTI, AI, and cybersecurity research. While MDPI was included to capture emerging contributions, it was not over-relied upon; papers from other high-impact venues were prioritized. The search was conducted using combinations of keywords and Boolean operators:

- (“Cyber Threat Intelligence” OR “CTI” OR “Threat Intelligence”)
- AND (“Machine Learning” OR “Deep Learning” OR “Artificial Intelligence” OR “LLM” OR “Large Language Model”)
- AND (“Threat Detection” OR “Attack detection” OR “Threat Prediction” OR “Forecasting” OR “Attribution” OR “Recognition” OR “Early Warning”)
- AND (“Cyber Threat” OR “Cyber Attack” OR “APT”)

This refined query design ensured broader, more precise coverage than prior reviews. Initial retrieval yielded 153 studies, from which duplicates and non-relevant titles were excluded. The final sample comprised 61 peer-reviewed studies published between 2019 and 2025.

Workflow of Systematic Literature Review

A total of 294 research publications were identified after the initial database keyword searches, of which 28 were duplicates. 87 papers were found to still need to be included after their titles and keywords were carefully compared with the predetermined inclusion/exclusion criteria. Furthermore, the full-text analysis excluded 100 papers; as a result, only 79 documents remain in the corpus for review, of which 61 scientific research papers were judged appropriate for inclusion in our study after a thorough review of these 79 articles and application of the specified inclusion/exclusion criteria. The SLR workflow is shown in Figures 1 and 2, from the total number of studies evaluated to the stage of accepted studies. A thorough quality assessment was conducted to ensure the robustness of the extracted data. This evaluation assesses the accuracy of the information in these studies, accounting for the completeness and relevance of the data.

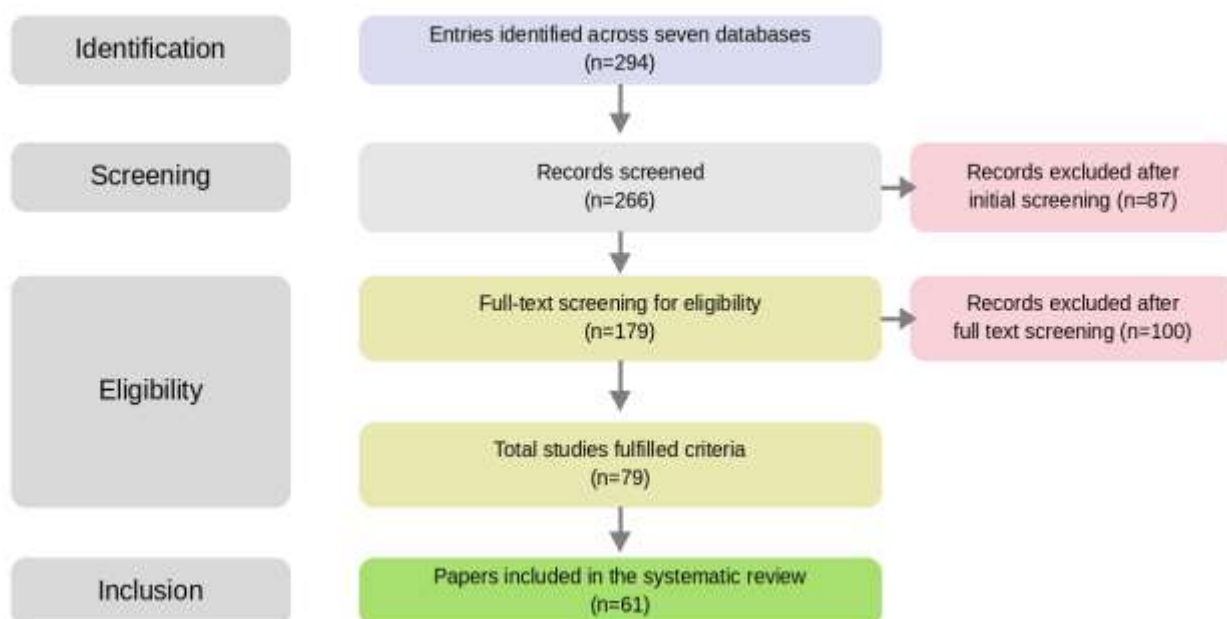


Figure 1. PRISMA flow diagram of study selection

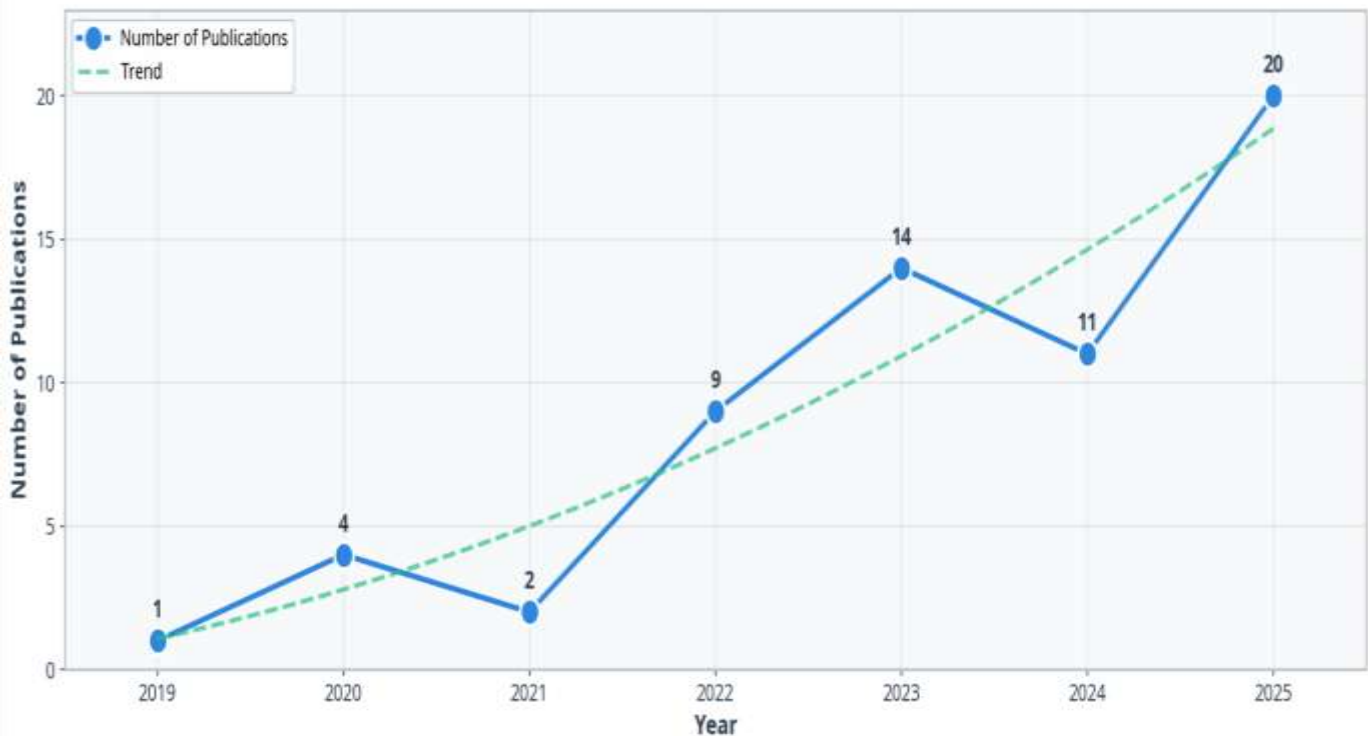


Figure 2. Distribution of research papers included in a systematic review.

The temporal distribution of the 61 included studies is summarized in Figure 2. Publication counts rose from a single study in 2019 to 4 in 2020, 2 in 2021, 9 in 2022, 14 in 2023, 11 in 2024, and 20 in 2025, indicating accelerating research interest in AI and LLM-based CTI from 2022 onward, coinciding with the wider availability of general-purpose large language models.

Quality Assessment Criteria

In addition to the topical inclusion and exclusion criteria applied during title, abstract, and full-text screening, each of the 79 full-text candidates was assessed against the following quality criteria, adapted from Kitchenham et al. [7]:

- QA1: The study clearly states its research objective(s) and the CTI task(s) addressed (e.g., detection, attribution, forecasting, or automated reporting).
- QA2: The methodology, technique, or model architecture is described in sufficient detail to assess or reproduce the approach.
- QA3: The data source(s) used for training, testing, or validation are identified and adequately described.
- QA4: The study reports at least one quantitative or qualitative evaluation of the proposed method's performance.
- QA5: The study discusses limitations, threats to validity, or the scope boundaries of the proposed approach.
- QA6: The publication venue is a peer-reviewed journal, a peer-reviewed conference proceeding, or a preprint from a recognized research group with subsequent citation uptake.

Studies were retained only where they satisfied at least five of the six of these criteria; this quality screening, applied jointly with the topical exclusion criteria described above, accounts for the reduction from 79 full-text candidates to the 61 studies retained in the final corpus.

Study Identification and Eligibility Criteria

A heat map of the keyword distribution across the reviewed articles is shown in Figure 3. By grouping keywords into color-coded clusters and displaying their frequencies and connections, this visualization highlights key concepts essential to this field of study.

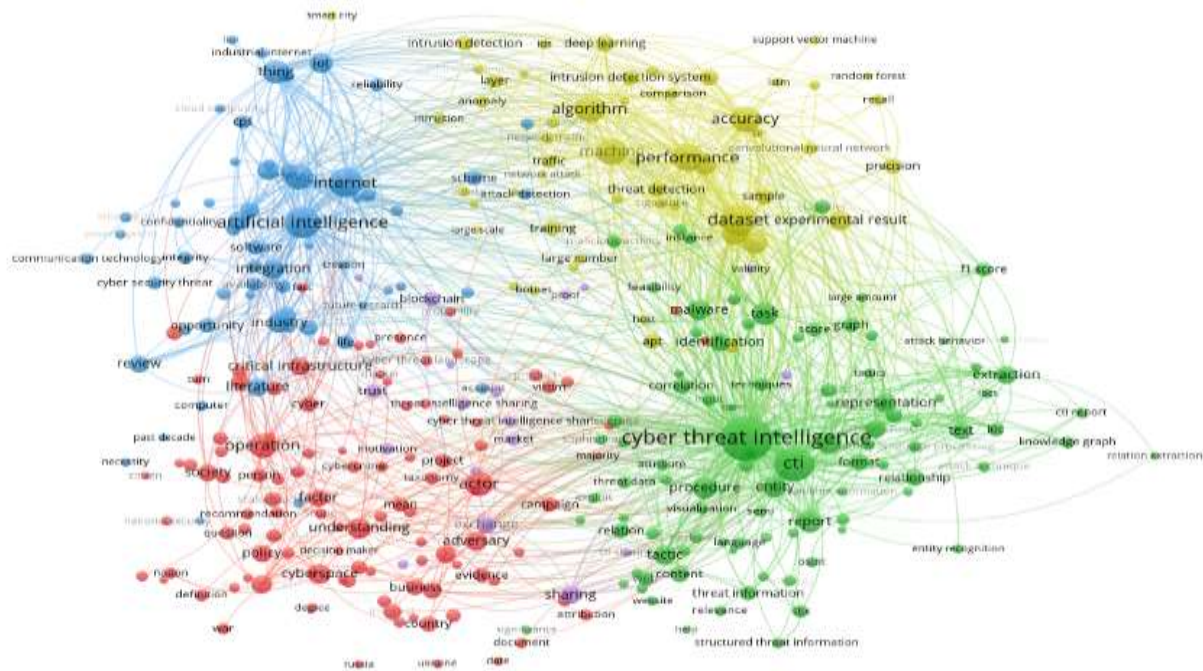


Figure 3. A heat map illustrates the frequency of keywords found in journal articles on the topic of CTI

It is worth noting that the retrieved literature contained some redundancy and irrelevance. To identify the most pertinent publications for our review, inclusion and exclusion criteria were rigorously applied. Publications were chosen for this review based on the following inclusion criteria: 1) Research studies used rule-based, match-based, text analysis, ML, DL, and LLM techniques to detect, recognize, identify, and forecast cyber threats and CTI forecasting. 2) Publications that were released between 2019 and 2025. 3) Extensive research papers or those published in prestigious conferences, journals, and preprints from internationally recognized universities. The following criteria were used to exclude publications from our review: 1) works published prior to 2019; 2) publications with insufficient content on modeling CTI. We identified and screened 61 documents using these criteria.

Motivating Factors for CTI Advancement

Advanced Persistent Threats (APTs) represent highly sophisticated attack campaigns in which adversaries maintain long-term, covert access to targeted networks to conduct reconnaissance, exfiltrate sensitive data, or disrupt sensitive operations. Unlike conventional attacks, APTs are resource-intensive and carefully planned, often initiated through phishing and social engineering techniques and sustained by malware such as trojans, spyware, or ransomware. Additional tactics, such as man-in-the-middle attacks and zero-day exploits, further enhance their stealth and effectiveness. The diversity, persistence, and evolving complexity of APT techniques significantly challenge traditional security defenses and underscore the need for intelligent, adaptive detection mechanisms that can identify threats before substantial damage occurs [52].

In response, CTI has emerged as a proactive approach to anticipating and mitigating cyber threats; however, conventional CTI methods remain constrained by their reliance on static Indicators of Compromise (IoCs) and signature-based detection. These approaches are largely reactive and struggle to adapt to the scale, speed, and variability of modern APT campaigns [40][7]. Consequently, recent research advocates integrating AI, particularly DL and LLMs, into CTI frameworks. AI-driven approaches enable scalable analysis of large volumes of unstructured threat data, improve detection accuracy through advanced pattern recognition, and support predictive and automated intelligence generation, positioning CTI to more effectively counter

sophisticated and evolving cyber threats [34][50][13]. Despite their promise, challenges remain in real-time intelligence integration, reducing reliance on historical data and human intervention, and addressing data privacy concerns [12]. These limitations motivate the development of advanced AI and LLM-driven CTI methodologies to strengthen the predictive, adaptive, and proactive capabilities of modern cybersecurity systems.

Related Work

In recent decades, there have been few dedicated high-quality systematic literature reviews (SLRs) on CTI, most of which remain narrowly focused on traditional techniques and frameworks. For example, the SLR by Rahman [40] offered a detailed and structured overview of CTI extraction from unstructured text, outlining various extraction goals, data sources, and the use of natural language processing (NLP) and ML. While this provides a solid foundation for researchers, it primarily focuses on conventional approaches and doesn't address recent developments in intelligent CTI systems. Similarly, Saeed et al. [7] examined CTI from an organizational perspective and introduced a multilayered framework comprising detection models, knowledge bases, and visualization tools to enhance cybersecurity resilience. However, it stopped short of exploring how modern AI, especially LLMs, can transform CTI practices. More recent systematic reviews continue to exhibit comparable limitations. For example, SLR by Santos et al. [54] provided a comprehensive assessment of CTI technologies, strategies, and collaboration mechanisms, emphasizing their effectiveness in mitigating contemporary cyber threats. Despite its methodological rigor and adherence to PRISMA guidelines, the review predominantly focused on established ML models, platforms, and intelligence-sharing practices, offering only limited discussion on higher-level semantic reasoning, autonomous intelligence generation, or cross-source contextual synthesis enabled by LLMs. Furthermore, the systematic multi-vocal literature review by Cascavilla et al. [55] presented an extensive taxonomy of cybercrime threat intelligence across surface, deep, and dark web sources, integrating both academic and grey literature. While this work significantly advances understanding of cybercrime intelligence indicators and risk assessment techniques, its emphasis remains on threat detection, categorization, and investigative support rather than on intelligent CTI automation or predictive threat reasoning. This pattern is further supported by survey-based reviews. Whitman [53], in their review, offered a thorough analysis of CTI mining methods, sources, and challenges, emphasizing the importance of IoC and TTP extraction for preventive cybersecurity. Nevertheless, as in previous SLRs, the analysis relies primarily on traditional NLP and ML techniques and does not systematically investigate how LLMs support end-to-end intelligent CTI workflows, multi-stage threat inference, or contextual understanding. All of these studies show that current CTI SLRs and reviews remain largely based on incremental ML improvements and conventional intelligence pipelines.

When combined, these findings indicate common limitations in systematic literature reviews of current CTI. Although earlier research provided valuable insights into aspects of CTI, such as cybercrime analysis, threat detection models, indicator extraction, and intelligence-sharing mechanisms, it primarily treated CTI as a set of technical tasks rather than as an integrated intelligence system [40][7][54][55]. As a result, the SLR's does not adequately address higher-level intelligence capabilities, such as semantic reasoning, cross-source correlation, temporal inference, and automated decision support [53][55]. More significantly, existing CTI SLRs do not thoroughly examine how LLMs can be integrated throughout the CTI lifecycle, despite the rapid development of AI and the recent emergence of LLMs.

Table 1 presents a comparative analysis of prior CTI systematic literature reviews to contextualize these limitations and clarify the contribution of the current study. This comparison shows that no existing survey provides a thorough, methodical analysis of LLM-enabled CTI frameworks that covers intelligence gathering, analysis, reasoning, prediction, detection, and mitigation. By offering a cohesive and forward-thinking synthesis of traditional, AI-enhanced, and LLM-based CTI approaches, the current SLR fills this gap.

Table 1 Comparative Analysis of Current CTI Systematic Literature Reviews

Study	Scope of Review	CTI Lifecycle Scope	AI-Driven Intelligence Techniques	LLMs Integration	Identified Limitations
Rahman et al. [44]	CTI extraction from unstructured text	Collection Processing	NLP, ML	✗	Focused on IoCs/TTPs; no intelligent reasoning
Saeed et al. [8]	Organizational CTI	Analysis, Visualization	ML-based detection	✗	Lacks AI-driven automation
Santos et al. [59]	CTI technologies & collaboration	Collection, Sharing, Analysis	ML	✗	Limited semantic reasoning
Cascavilla et al. [60]	Cybercrime CTI	Collection, Classification	ML	✗	Detection-centric
Whitman [58]	CTI mining survey	Extraction, Analysis	NLP, ML	✗	No end-to-end intelligence
This study	AI and LLM-driven CTI	Full lifecycle	ML/DL	✓	Addresses prior gaps

LITERATURE REVIEW

The literature review of main studies concentrated on classifying research topics associated with conventional CTI, AI in CTI, and LLM-based CTI, which are covered in greater detail within the particular focus areas of individual publications. While Section 3.1 explores conventional CTI methods, Section 3.2 discusses AI-based CTI techniques that improve threat detection and present novel models and characteristics for threat identification and recognition, Section 3.3 focuses on LLMs that process unstructured CTI reports for threat detection, type classification, and practical mitigation steps.

Conventional Methods in CTI

Digital forensics and information security laws are facing challenges due to the growing volume of data, driven by the widespread use of security devices and the increasing complexity of information technology. Key findings from the reviewed studies are highlighted in Table 2, which summarizes research papers on conventional methods in CTI.

Table 2 Summary of studies related to the Conventional methods in CTI

No.	Method	Technique	Data Source	Key Contribution	Limitations	CTI Task	Evaluation
1.	DFR [10]	Digital Forensic Readiness model	Audit Database Log	Combines CTI with forensic readiness to reduce incident response time and cost	False positives; lacks zero-day detection; not scalable to enterprise environments	Incident Response Readiness	Validated on audit log dataset; qualitative improvement in response readiness; lacks precision/recall metrics
2.	Decentralized CTI [17]	Blockchain-based CTI-sharing network	Microsoft Azure data from ISPs and routers	Decentralized protection mechanism using blockchain	Limited validators; OSI layer 3-only inspection	Threat Intelligence Sharing	Prototype tested on Ethereum network logs; demonstrates feasibility; lacks empirical benchmarks
3.	TIAM [9]	Text analysis & threat	APT reports from Indian Infosec	Classifies and scores threats	Sparse data; lacks real-time	Threat Recognition	Manual validation on APT reports;

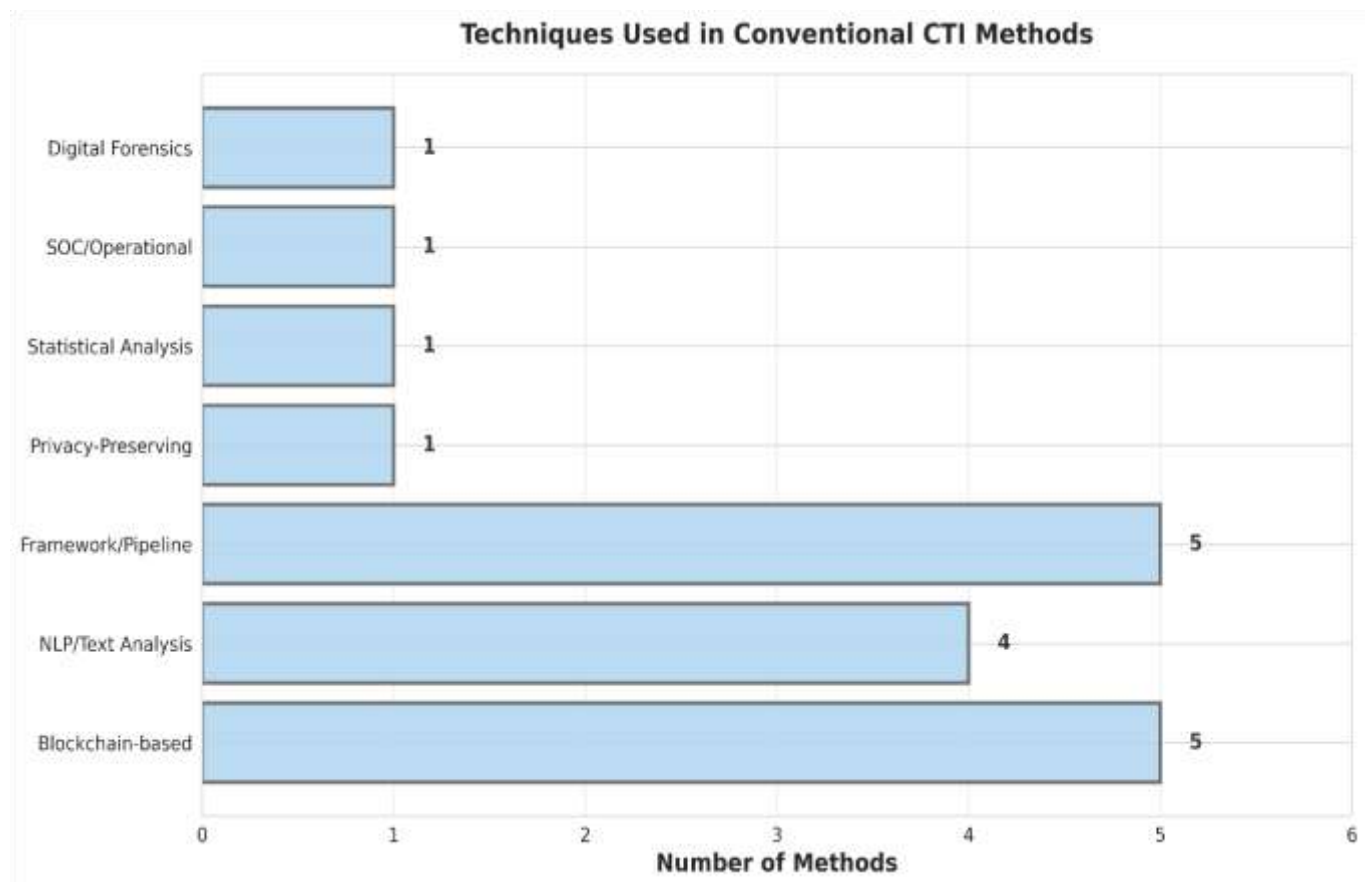


No.	Method	Technique	Data Source	Key Contribution	Limitations	CTI Task	Evaluation
		scoring with vector space model	Consortium	based on textual characteristics	analysis; manual feature engineering	IOC Extraction	limited quantitative metrics; relies on rule-based parsing
4.	Eight-step CTI [13]	Unified analytical CTI pipeline	Pegasus Spyware & SolarWinds reports	Improves CTI data handling across diverse sources	Requires structured STIX/TAXII; lacks real-time capability	Threat Intelligence Lifecycle	Validated against Pegasus and SolarWinds datasets; qualitative insights; lacks automation metrics
5.	AttacKG [14]	Template-based graph alignment	CTI reports from Cisco Talos & Microsoft Security	Constructs structured attack graphs from unstructured CTI	Templates rigid; limited adaptability to evolving threats	Threat Behavior Extraction	Extracted attack graphs from 1,515 reports; precision and recall reported around 84% and 79%
6.	SECDFAN [15]	Multi-layered NLP-driven system	Surface web, deep web, and darknet forums	Transforms raw forum posts into shareable CTI formats	High dependence on NLP; vulnerable to threat dynamics	CTI Generation from Open Source	Forum data processed; empirical tests show 72% accuracy in entity recognition; subjective scoring
7.	Cyber Defenders [16]	SOC lifecycle framework	SOC operational data	Categorizes vulnerabilities and recommends risk management improvements	Assumes uniform threat landscape; lacks specialization	SOC Optimization / False Positive Reduction	Tested on SOC data; qualitative results reported; lacks standardized accuracy metrics
8.	BLOCIS [18]	Blockchain smart contracts	Simulated Sybil attacks	Improves CTI reliability through blockchain validation	Smart contract flaws; blockchain bloat risks	Secure Threat Sharing / Data Integrity	Blockchain simulations conducted; effectiveness in Sybil resistance shown; no detection metrics
9.	CyRiPred [12]	Statistical cyber risk forecasting	National Vulnerability Database	Identifies future threats from known vulnerability patterns	Lacks continual learning; can't detect zero-day threats	Vulnerability Forecasting	Evaluated using NVD data; statistical accuracy metrics not reported; retrospective trend analysis
10.	ABC [11]	Probabilistic blockchain model	Simulation-based CTI feeds	Adds PoQ validation to CTI collection and storage	Relies on validator reputation; bias in scoring	Secure CTI Evaluation & Storage	Simulated CTI feeds; PoQ scoring accuracy not quantitatively validated
11.	SCERM [46]	STIX-based CTI graph framework	Schema test, HAILATAXII, IBM X-Force Exchange	Scores and completes CTI data across STIX phases	Assumes trustworthy sources; does not verify CTI accuracy	CTI Assessment & Refinement	Applied to STIX datasets; generates completeness scores; lacks ground truth validation
12.	BlockIntelChain [56]	Blockchain-based CTI sharing framework	CTI sharing networks (ENISA guidelines)	Enables secure, decentralized threat intelligence sharing with superior	Blockchain scalability concerns; consensus overhead; network	Threat Intelligence Sharing	Performance comparison with centralized approaches; ENISA compliance

No.	Method	Technique	Data Source	Key Contribution	Limitations	CTI Task	Evaluation
				performance over centralized techniques; enhances trust and transparency	latency		validation
13.	CTI-Risk Management Integration [57]	Integration of CTI and risk management frameworks	Critical infrastructure threat data	Proposes integrated approach combining CTI with risk management for comprehensive critical infrastructure protection	Requires organizational alignment; complex implementation across heterogeneous systems; integration overhead	Risk Management / Critical Infrastructure Protection	Framework validation on critical infrastructure environments
14.	Cloud-based CTI Sharing [58]	CTI information-sharing procedures in cloud ecosystems	Cloud environment threat data	Examines CTI and information-sharing procedures for proactive threat identification in cloud-native environments	Cloud-specific threat landscape variability; multi-tenancy isolation challenges; cloud provider dependency	Threat Identification / Cloud Security	Proactive threat detection validation
15.	MANTRA Network [59]	Privacy-preserving CTI sharing network	Multi-organizational threat intelligence feeds	Proposes MANTRA framework enabling privacy-preserving CTI sharing while maintaining collaborative benefits and threat visibility	Privacy-utility trade-off; requires trusted intermediaries; overhead in privacy mechanisms	Privacy-Preserving Threat Sharing	Privacy preservation validation; framework effectiveness assessment
16.	Global CTI Framework [60]	Standardized global framework for CTI sharing	International CTI networks	Argues for establishing trusted global CTI sharing frameworks to enable collective cyber resilience and international cooperation	Requires international coordination; standardization challenges across jurisdictions; policy alignment difficulties	International Threat Intelligence Coordination	Framework feasibility assessment; policy recommendations; strategic analysis
17.	Blockchain Incentive Mechanism [61]	Tripartite incentive mechanism for blockchain CTI sharing	Blockchain-based CTI networks	Develops incentive mechanisms to encourage participation and contribution in blockchain-based CTI sharing systems	Incentive design complexity; blockchain transaction overhead; economic sustainability concerns	Collaborative Threat Sharing / Incentivization	Mechanism effectiveness validation; participation rate analysis
18.	STIX/TAXII Optimization [62]	Optimization of threat intelligence sharing across platforms	STIX/TAXII-formatted threat data	Addresses optimization of STIX/TAXII-based threat intelligence sharing across multiple security	Platform heterogeneity challenges; standardization gaps; interoperability limitations	Multi-Platform Threat Intelligence Sharing	Cross-platform sharing efficiency metrics; integration validation

No.	Method	Technique	Data Source	Key Contribution	Limitations	CTI Task	Evaluation
				platforms and tools			

Figure 4 presents a visual analysis of 18 conventional CTI methods, as shown in Table 2. Analysis show that CTI techniques primarily rely on framework-based pipelines, blockchain solutions, and NLP-driven text analysis, with a focus on IOC extraction and threat intelligence sharing rather than full lifecycle support. Although these methods enable the exchange of structured information, they have several significant drawbacks, including limited scalability, high integration complexity, limited evaluation rigor, and inadequate real-time and zero-day detection capabilities. Furthermore, obstacles such as blockchain overhead, data quality issues, and trade-offs between privacy and utility further limit their efficacy, underscoring the need for more flexible, scalable, and intelligence-driven CTI methods.



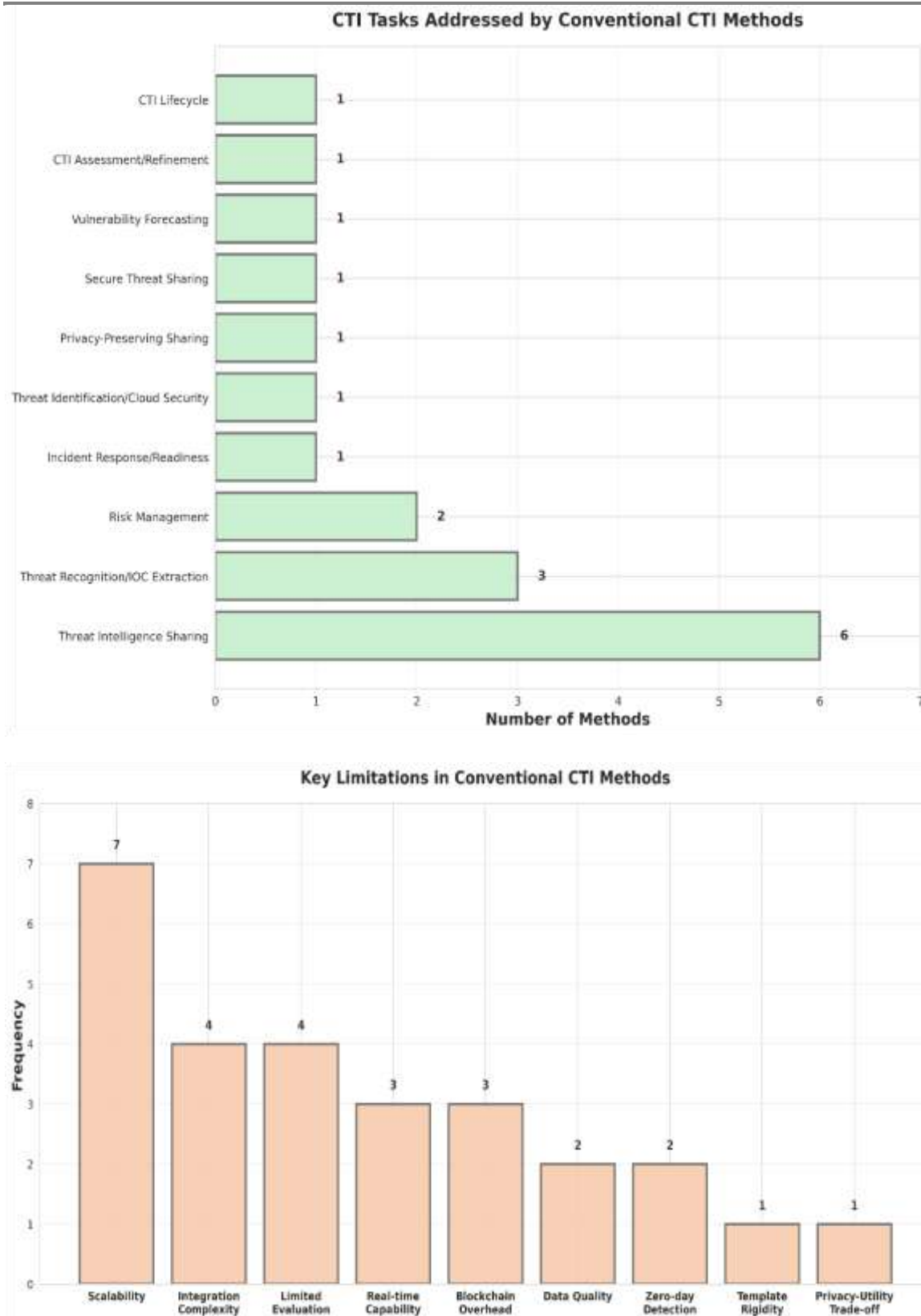


Figure 4. Visual analysis of 18 conventional CTI methods

Findings and Reflections on Conventional CTI Methods

This section provides a categorized analysis of conventional methods applied in CTI, offering critical insights summarized in Table 2. Over the past few years, there's been a clear shift in how CTI is being approached: from static, isolated tools toward more intelligent, automated, and interconnected systems. While the field has made remarkable technical strides, this section unpacks some of the less visible tensions and trade-offs that emerge

across recent approaches. In addition to summarizing each method, we explored where these tools succeed, where they fall short, and why that matters.

One recurring pattern across several conventional models is a strong emphasis on automation, particularly for extracting or organizing threat data. Tools such as TIAM [9] and AttackKG [14] illustrate how structured analysis and NLP can support faster threat recognition. AttackKG, in particular, performs well in identifying techniques across real CTI reports, achieving over 80% accuracy. Despite this, success often depends on the quality and format of the underlying data. Many real-world reports are messy, incomplete, or inconsistent, which can pose challenges even for the most sophisticated parsing models. TIAM's reliance on manual feature design highlights the ongoing gap between experimental performance and practical application. In a different context, blockchain has gained attention as a trust mechanism for CTI sharing, as seen in BLOCIS [18] and ABC [11]. In theory, these systems create secure, tamper-resistant ecosystems for the exchange of threat data. However, in practice, the results remain largely theoretical or simulation-based. While smart contracts and validation protocols add layers of integrity, the actual intelligence value of the shared CTI isn't automatically improved. These systems excel in protecting provenance but are less effective in enhancing the analytical richness of the data itself.

A different category of work, such as the Eight-step CTI process [13] or Cyber Defenders [16], focuses on CTI process management across the full threat intelligence lifecycle. These models are thorough and address genuine operational concerns, particularly in Security Operations Center (SOC) environments. Nonetheless, they often fall short in terms of empirical benchmarking. Their effectiveness is mostly demonstrated through specific use cases from organizations. Without broader validation, it is difficult to determine the applicability of these methods across sectors with varying infrastructures or threat landscapes. An especially interesting, though less discussed, area is the use of community-driven or open-source platforms for CTI extraction. SECDFAN [15], for example, draws on online forums and uses semantic layers to structure informal conversations into actionable insights. This method reflects a growing recognition that valuable CTI isn't always published in official channels. Still, with informal sources comes unpredictability. The tool may identify relevant patterns, but it also risks absorbing noise, bias, or speculative claims, especially in high-traffic or unmoderated environments.

Models such as CyRiPred [12] aim to provide a forward-looking perspective on CTI by predicting cyber risks based on known vulnerabilities. This effort is commendable, as the model generates useful categories for assessing risks. In contrast, reliance on historical Common Vulnerabilities and Exposures (CVE) data and the lack of continuous learning hinder the model's ability to adapt effectively to emerging attack vectors, particularly zero-day exploits. The reviewed methods indicate a positive shift in CTI towards more adaptable and data-driven approaches. There are still challenges to overcome, as many traditional tools rely on idealized conditions and emphasize design over measurable outcomes. Recent advances in conventional CTI have increasingly focused on blockchain-based architectures to address persistent challenges in establishing trust and ensuring secure information exchange. BlockIntelChain [56] presents a comprehensive blockchain-based framework that adheres to the ENISA Threat Intelligence Sharing Guidelines (2024) and demonstrates that deep learning algorithms operating on blockchain-stored intelligence achieve superior performance compared with traditional centralized techniques. However, the framework relies on structured, preformatted threat data and lacks the capability to ingest unstructured CTI reports. The system cannot automatically extract threat entities from narrative text; manual data preparation is required before storage in the blockchain.

The integration of CTI with broader organizational risk management has received substantial attention. [57] propose a multilayered architecture that integrates CTI with risk assessment methodologies to protect critical infrastructure. Their work highlights the necessity of contextualizing threat intelligence within organizational risk profiles rather than treating CTI as an isolated security function. The limitation lies in the static nature of risk categorization, which depends on predefined taxonomies and manual threat classification. The framework lacks the semantic understanding necessary to automatically interpret emerging threat descriptions or to correlate different intelligence sources without human intervention.

Cloud computing presents unique challenges, as discussed in [58]. Their research emphasizes that cloud-specific CTI procedures enable organizations to identify and contain threats before they materialize, accounting for multi-tenancy, virtualization, and dynamic resource provisioning. The study's primary gap is the absence of automated intelligence extraction from cloud security logs and alerts. The proposed procedures require analysts to manually

interpret and categorize threat indicators, creating bottlenecks in high-volume cloud environments where real-time processing is essential.

Privacy preservation has emerged as a critical concern. The MANTRA framework [59] introduces a conceptual architecture for privacy-preserving CTI sharing that enables collective threat visibility without compromising sensitive operational data, addressing regulatory concerns under frameworks such as GDPR. While the framework addresses privacy constraints, it does not incorporate mechanisms for automated anonymization or summarization of threat intelligence. The lack of intelligent text processing means that sensitive information must be manually redacted, limiting scalability and introducing potential for human error in privacy-critical operations.

At the strategic level, authors [60] argue that establishing a trusted global framework for CTI sharing is essential to collective cyber resilience and defense, examine STIX/TAXII standards, and propose improvements to international cooperation. The work remains largely conceptual and policy-oriented, failing to address the technical challenge of cross-lingual threat intelligence processing. International CTI sharing requires the ability to translate, normalize, and semantically align threat reports across multiple languages, a capability that current STIX/TAXII implementations do not natively support.

Zhou et al. [61] discuss a tripartite incentive mechanism for blockchain-based CTI sharing that simulates strategic interactions among producers, consumers, and platform operators using a game-theoretic approach. Although the model assumes that CTI quality can be measured objectively, it provides no automated method for evaluating the timeliness, relevance, or accuracy of shared intelligence. The incentive system may reward quantity over actionable value if threat-report quality is not assessed using natural language understanding.

Additionally, [62] bridges the gap between CTI's theoretical interoperability and its operational reality by addressing the practical challenge of optimizing STIX/TAXII sharing across heterogeneous security platforms. Semantic interoperability is not addressed by the optimization; instead, it concentrates on transport and format compatibility. Therefore, without intelligent mapping capabilities, the same threat may be repeated or misclassified across platforms, and different organizations may use different terminology to describe the same threat.

In summary, Figure 5 depicts the overall lifecycle of traditional CTI techniques [55] applied to different downstream tasks that generate data from multiple sources. Every stage of this process involves a variety of approaches and strategies, and every CTI technique incorporates these choices.

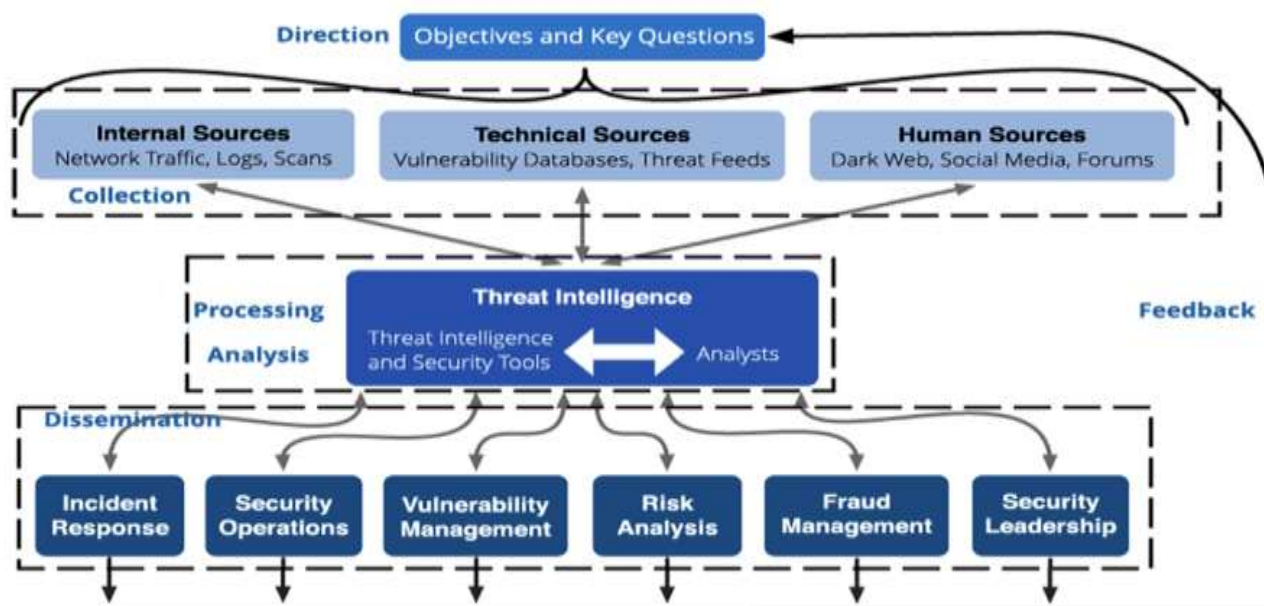


Figure 5. Lifecycle of conventional Cyber Threat Intelligence (CTI) methods

AI Methods in CTI

CTI frameworks can process massive volumes of data and spot patterns of malicious activity by leveraging AI, providing insights for preventive defense tactics. However, existing CTI methods frequently fail to identify novel or previously unidentified threats, underscoring a major shortcoming in the rapidly evolving cybersecurity environment. Research results on AI techniques used in CTI are compiled in the Table 3.

Table 3 Summary of the AI methods in CTI

Table: Comparative Review of CTI-Based Methods and Techniques

No.	Method	Contribution	Technique	Data Source	Limitations	CTI Task	Evaluation
1	TIMiner [27]	CNN-based system for categorizing CTI across domains; introduces Threat-Index metric.	Word2Vec, CNN, BiLSTM	Security blogs, social media, vendor bulletins, hacking forums	Dependent on social media data; lacks evaluation against APTs	IOC Identification	No direct APT benchmarking; qualitative results
2	CTI-Twitter [28]	Utilizes Twitter to classify threat intelligence using NLP techniques.	TF-IDF, BERT, LDA	Twitter, NVD, CERT alerts	Short text limits clustering; lacks temporal relevance filtering	Threat Extraction & Relevance Classification	F1-score on tweet classification task
3	HinCTI [30]	Heterogeneous information network with GCN and MIIS for CTI modeling.	GCN	IBM X-Force, VirusTotal	Limited relation types and static analysis	Threat Type Classification	Accuracy metrics on 11K+ nodes dataset
4	INTIME [26]	Lifecycle model for CTI processing using ML.	SVM, CNN	Social, deep, and dark web	Not robust to zero-days; relies on social feeds	Content Relevance Scoring	Precision/recall on tweet classification
5	CTI View [33]	Automated extraction of APT-related entities from text.	BERT + GRU	English APT threat corpora	Low relation extraction accuracy	Entity Extraction	Accuracy: 72% on APT entity dataset
6	Malware Forensics [24]	Introduces a tri-layer CTI framework separating threat feed collection, semi-supervised analysis, and threat mitigation.	Ensembled XGBoost and Random Forest	Sandbox environment, Malimg, and Microsoft's Malware Dataset	Misclassifications with zero-day attacks; complex, layered processing pipeline	Threat Classification, CTI Report Filtering	Classification accuracy over malware datasets; lack of benchmark comparison for zero-day handling
7	THREATKG [25]	Creates an OSCTI knowledge graph by extracting and	ML classifiers, BiLSTM for NER, PCNN-ATT for RE	OSCTI reports from diverse sources	OSCTI sources may be unreliable; schema rigidity and absence of	Threat Entity and Relation Extraction, Threat Knowledge	Effectiveness in extracting knowledge entities; evaluated



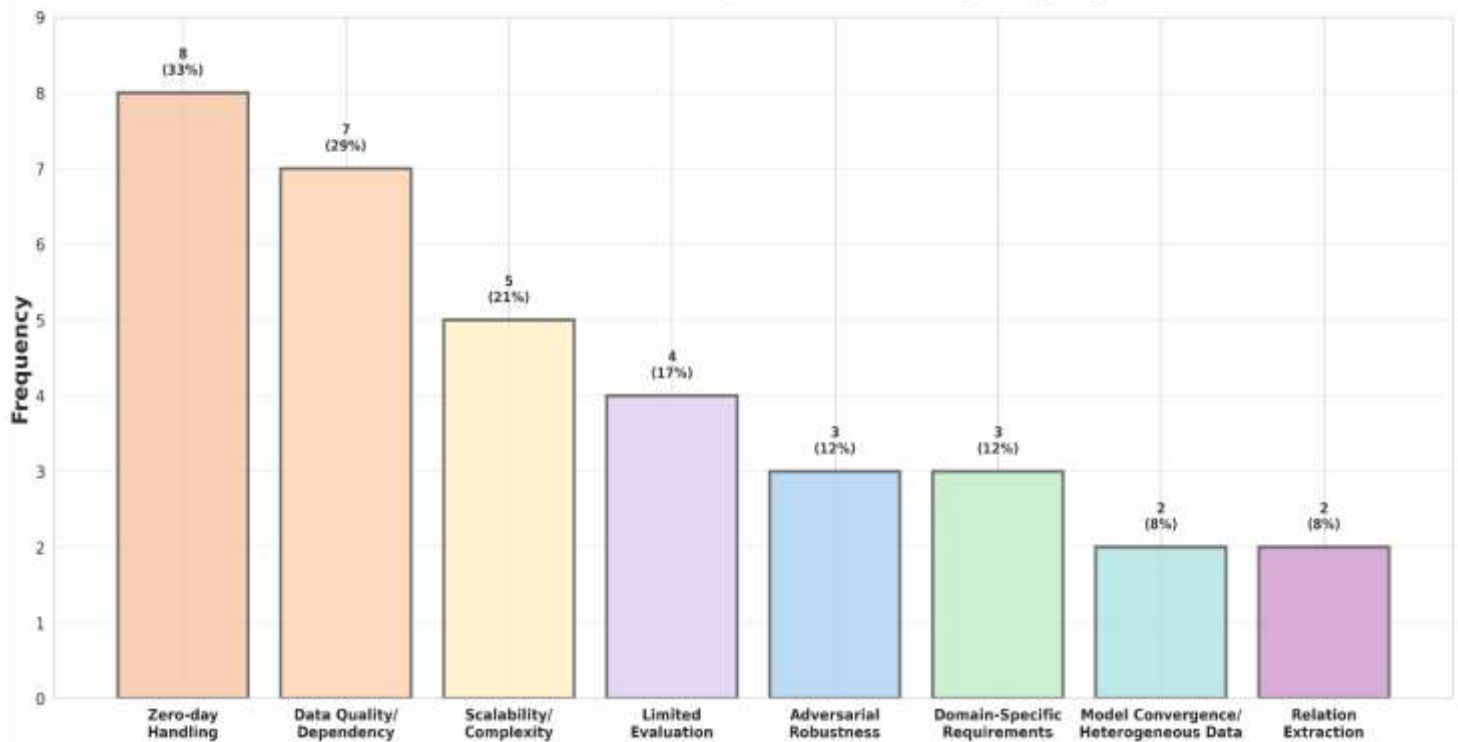
		organizing threat entities and relationships using ML and NLP.			end-to-end gradient flow	Structuring	using custom knowledge graph completeness metrics
8	TriCTI [29]	Develops an automated system to extract and classify actionable CTI from reports with campaign stage mapping.	Trigger-enhanced neural classification based on weighted trigger vectors	DS-1 and DS-2 with 3,468 IOCs and labeled sentences	Ambiguities in campaign stage assignment; regular expression misalignment issues	IOC Detection, Campaign Stage Classification	Improved accuracy over BERT; tested on DS-1 and DS-2 labeled CTI datasets
9	Botnet Detection [31]	Combines Explainable AI and OSINT to improve DGA-based botnet detection and CTI transparency.	Random Forest with OSINT feature engineering	Alexa Top DNS data and labeled malicious botnet traffic	High feature computation cost; reduced generalization across diverse DGA variants	Malicious Traffic Detection, Botnet Behavior Analysis	Comparison with character-based DGA classifiers; improvement in accuracy using 7 key features
10	BFLS [32]	Proposes a decentralized, privacy-preserving CTI model sharing system using federated learning and blockchain.	Blockchain-based KNN with federated aggregation and smart contract updates	ISCX-IDS-2012 and CIC-DDoS-2019 with 13 DDoS attack types	Lack of resilience to colluding malicious nodes; dependency on over 50% committee trust	Threat Model Sharing, Privacy-Preserving CTI Exchange	Evaluated using distributed threat model performance; focus on CTI validation accuracy and decentralization effectiveness
11	CTI Attribution [20]	Development of a CTA framework for detailed and accurate profiling of cyber-threat actors to identify real individuals or organizations behind attacks.	Attack2vect based on Word2Vec; Random Forest for CTA	CTI reports and malware samples from multiple data sources	Limited CTA datasets; unstructured formats complicate feature extraction	Attribution	Not specified; assumed internal validation on collected samples
12	Social Media-CNN [21]	A cyber threat index framework using CNN-based	CNN-based anomaly detection	Twitter feed	Susceptible to manipulation and false content in social media; vulnerable to	Threat Detection & Indexing	Qualitative case study and anomaly detection accuracy

		anomaly detection on social media content.			information operations		
13	Intrusion Detection [22]	Privacy-preserving CTI scheme using Federated Learning to identify intrusions.	LSTM with Federated Learning	NF-UNSW-NB15-v2 and NF-BoT-IoT-v2	Processing inefficiency for heterogeneous datasets; domain adaptation required	Intrusion Detection	Detection rate and performance on benchmark intrusion datasets
14	Few-shot CTI [23]	Few-shot learning CTI classifier for emerging threats with limited training data.	Few-shot and transfer learning based on BERT	Microsoft Exchange Server data breach of 2021	Limited to BERT base model; scalability not demonstrated	Classification & Forecasting	Few-shot classification performance on novel CTI inputs
15	CTIMD [19]	Malware detection via API call sequence and threat behavior relationships.	TextCNN and BiLSTM	Datacon 2019 malware dataset	Heuristic limitations; poor performance on obfuscated or unknown malware	Malware Detection	Benchmark comparison with static malware detection models
16	Deep Maxout [47]	Blockchain and DL framework for alert segregation and intrusion mitigation.	Entropy-based K-means, DMN for prediction, Blockchain integration	Mendeley dataset	Weak against zero-day attacks; works only on structured data	Situational Awareness & Alert Prediction	Performance measured via alert prediction accuracy
17	CyberEntRel [48]	Joint extraction of cyber entities and their relations from text using DL.	RoBERTa + BiGRU + CRF	1,098 CTI reports from Microsoft, Cisco, McAfee, and Kaspersky	Precision and recall drop for overlapping triples; no coverage of zero-day threats	Entity & Relation Extraction	Precision, recall, and F1 score on annotated CTI datasets
18	ML Threat Classification [63]	Machine learning algorithms for threat classification from underground sources.	Machine Learning algorithms	Dark web forums and marketplaces	High dependency on dark web data quality; noisy and multilingual data; lack of ground-truth labels; rapid threat-actor platform migration	Threat Classification / Dark Web Monitoring	Classification accuracy on dark web datasets
19	BERT-GRU-BiLSTM-CRF [64]	Hybrid transformer-based architecture for threat entity extraction.	BERT + GRU + BiLSTM + CRF	CTI data and threat reports	High computational complexity; fixed context window requiring document	Threat Entity Extraction / Named Entity Recognition	Entity extraction accuracy metrics; real-time deployment feasibility

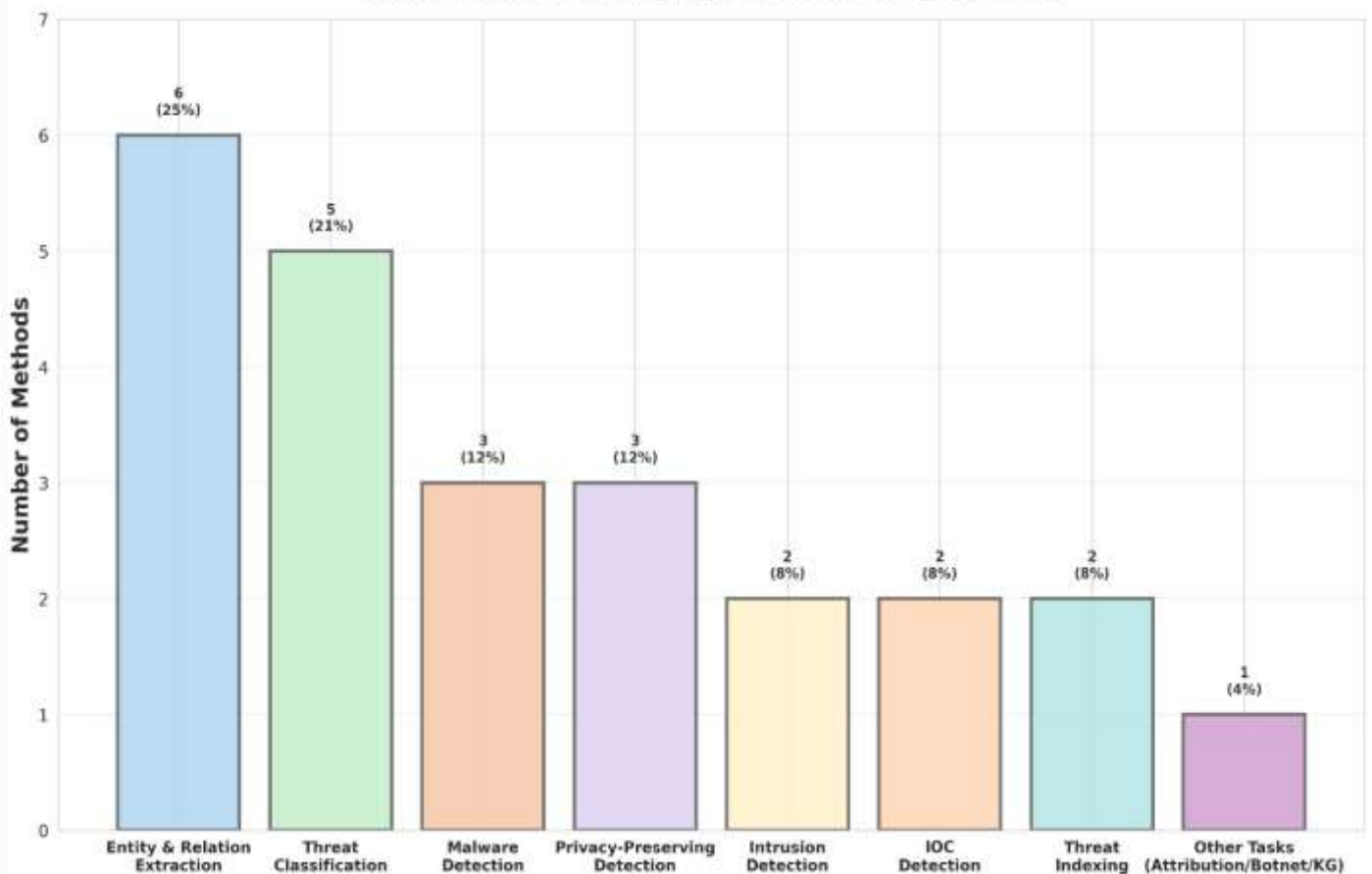
					segmentation; potential threat narrative fragmentation		assessment
20	Federated Learning IoT [65]	Privacy-preserving federated learning framework.	Federated Learning	Distributed IoT device threat data	Model convergence challenges with heterogeneous data distributions; communication overhead; adversarial participant risks	Privacy-Preserving Threat Detection / IoT Security	Privacy preservation validation; IoT environment performance metrics
21	Federated Learning Analysis [66]	Systematic analysis of FL approaches for cybersecurity.	Federated Learning implementations	Security-related federated learning studies	Quality verification challenges; lack of model-update integrity mechanisms; non-IID data aggregation issues	Privacy-Preserving Threat Detection / Distributed Systems	Systematic literature review; best-practice identification
22	ML Security Analysis [67]	Comprehensive review of ML threat models, attacks, and defenses.	ML security approaches	ML security literature and threat models	ML systems vulnerable to poisoning and evasion attacks; robustness-accuracy trade-off; threat actors can craft evasion-specific attacks	ML System Security / Adversarial Robustness	Threat model analysis; defense mechanism evaluation
23	Neural Network Malware Detection [68]	Verification methods for neural network malware classifiers.	Neural network verification methods	Malware samples and adversarial perturbations	Tension between model expressiveness and verifiability; highly accurate models are difficult to formally verify; adversarial evasion risks remain	Malware Detection / Adversarial Robustness	Verification method effectiveness; adversarial robustness testing

Figure 6 shows a shift toward automated knowledge extraction and pattern recognition, with AI/ML-based CTI techniques mainly concentrating on entity and relation extraction, threat classification, and malware detection. Real-time network and IoT data are still underutilized in these methods, which primarily rely on unstructured textual sources like CTI reports, documents, and social media, augmented by malware datasets and security databases. Notwithstanding their analytical benefits, AI/ML techniques continue to face challenges, including limited zero-day handling, reliance on data quality, scalability and complexity issues, and limited evaluation procedures. The need for more thorough, robust, and operationally grounded evaluation frameworks in AI-driven CTI research is highlighted by the fact that evaluation is primarily dominated by standard accuracy-based metrics and benchmark datasets, with relatively little emphasis on robustness testing and qualitative analysis.

Common Limitations in AI/ML CTI Methods (24 Papers)



CTI Tasks Addressed by AI/ML Methods (24 Papers)



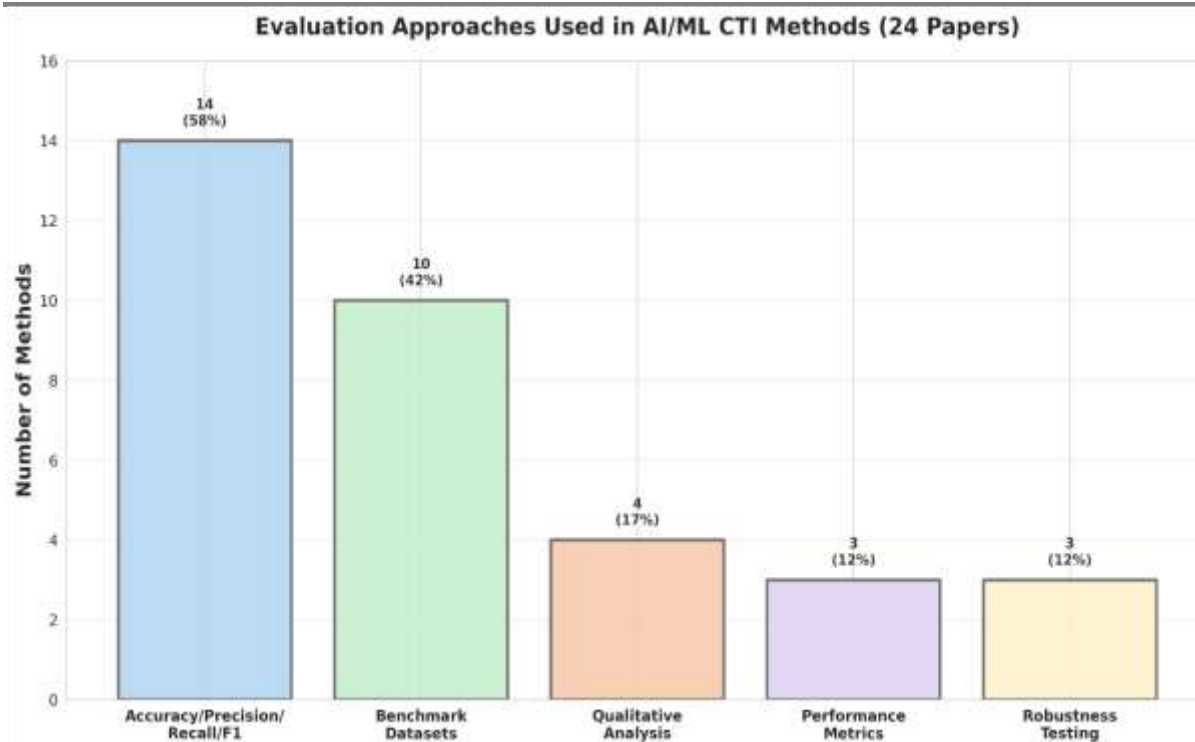


Figure 6. Analysis of ML/AI-based methods in CTI

Findings and Reflections on AI Methods in CTI

This section provides ML/AI-driven methods applied in CTI, offering critical insights and findings summarized in Table 3. AI has rapidly become integral to the evolution of CTI, addressing the increasing complexity and volume of cyber threats. However, despite considerable advancements, critical limitations remain in areas such as generalization, real-time adaptability, and zero-day attacks detection.

A key finding across the reviewed studies is the use of ML and NLP for IOCs and actionable intelligence from unstructured or noisy data. For example, TIMiner [27] and CTI-Twitter [28] demonstrate the utility of convolutional neural networks (CNNs) and supervised classification for deriving IOCs and categorizing domain-specific threats from social media content. Although these models achieved strong accuracy scores, their effectiveness is limited by the volatility and bias present in social platforms. These tools depend on static classifiers and specific taxonomies, which limit their ability to adapt to new attack vectors such as zero-day attacks or evolving threat semantics.

To address data heterogeneity and the lack of structure in CTI sources, studies such as HinCTI [30] and THREATKG [25] employ graph-based approaches using Heterogeneous Information Networks (HINs) and Knowledge Graphs. These models excel at capturing relationships among threat entities, such as malware, IP addresses, and domains, thereby improving situational awareness. Despite this, they are constrained by a rigid structure imposed by predefined schemas and have limited capacity to model newly emerging relationships or unknown threat types. This underlines the necessity for schema-flexible architectures and continual learning capabilities.

In contrast to other findings, deep learning models, such as BERT and BiLSTM, are widely used to enhance entity extraction and classification, as demonstrated in TriCTI [29] and CTI-View [33]. These systems enhance traditional CTI methods by using contextual embeddings and sequence modeling, enabling more precise classification of threats and their campaign stages. A limitation of this study is the continued scarcity of data, particularly in low-resource areas and when addressing rare threat types. Few-shot CTI learning methods, such as that described in [23], aim to address this problem through techniques such as data augmentation and transfer learning. However, reliance on synthetic or GPT-generated samples raises concerns about data quality, the risk of overfitting, and the applicability of the results to real-world scenarios.

Federated learning BFLS [32] and privacy-preserving models [22] offer promising directions for secure, collaborative intelligence sharing. These approaches allow organizations to train shared models without exposing sensitive internal data. Nonetheless, the reliability of these systems is still questionable, particularly against adversarial nodes or poisoning attacks. Their implementation demands a high level of trust in consensus mechanisms, which may not be practical across all domains. Explainability is another crucial area that needs attention. Models such as Botnet-XAI [31] address this issue by combining Explainable AI with Open Source Intelligence (OSINT), thereby enhancing transparency in CTI decisions. While these approaches help build trust among analysts, there remains a significant challenge in balancing explainability and model complexity. This is particularly important in real-time decision-making environments where latency is a critical factor.

Dynamic malware detection methods, such as CTIMD [19] and hybrid DL-blockchain models such as Deep Maxout [47], demonstrate progress in identifying malicious behaviors through runtime analysis and distributed alert systems. Unlike traditional static signature-based detection, these models offer enhanced capabilities. Because they rely on structured data formats, they often exhibit limitations and struggle to generalize to unknown malware behaviors, particularly zero-day exploits.

With [63] promoting ML as a strategic tool for identifying and categorizing new threats from dark web sources, ML algorithms for CTI data classification have gained considerable traction. According to their research, early warning of emerging cyber campaigns can be obtained by automatically classifying threat intelligence from underground forums and marketplaces. However, the method relies heavily on the representativeness and quality of dark web data, which are inherently noisy, multilingual, and prone to rapid changes as threat actors switch between platforms. The efficacy of supervised learning is further limited by the absence of ground-truth labels for novel threats.

In [48], whose work on joint extraction of cyber entities and relations using DL has attracted substantial scholarly attention, has significantly advanced DL approaches for named entity recognition and relation extraction in CTI. Threat actors, malware families, vulnerabilities, attack methods, and their semantic relationships can all be identified simultaneously through their methodology. Although the model performs well on benchmark datasets, it needs a large amount of cybersecurity-specific annotated training data. Therefore, dependencies in large CTI reports pose another challenge for the joint extraction architecture, which may overlook relationships spanning several paragraphs.

Transformer-based architectures have been tailored for CTI analysis [64], who created a threat entity extraction model called BERT-GRU-BiLSTM-CRF. This hybrid architecture combines the structured prediction of conditional random fields, the sequential modeling capabilities of recurrent networks, and the contextual understanding of BERT. Yet, the computational complexity of this multi-layered architecture makes deployment difficult in real-time operational settings. Additionally, the model inherits BERT's fixed context window limitation, necessitating document segmentation to break threat narratives into segments.

Using Federated Learning (FL) [66] frameworks, privacy-preserving collaborative threat detection in CTIs has emerged as a promising approach. Recent research has offered systematic analyses of FL approaches for cybersecurity and advanced AI frameworks for cyberthreat detection in IoT-assisted smart cities [65]. These studies demonstrate how FL can overcome organizational reluctance to share proprietary threat intelligence by enabling distributed training of CTI models without centralizing sensitive security data. However, real-time scalability is limited by communication overhead; heterogeneous data distributions among participating organizations hinder model convergence; and existing implementations lack reliable mechanisms to verify the integrity of model updates, creating potential attack vectors for malicious participants—all of which pose serious obstacles to practical deployment.

These vulnerability concerns extend beyond FL architectures to ML-based CTI systems more broadly. Recent research examining threat models, attacks, and defenses for machine learning systems highlights that CTI tools built on ML are themselves susceptible to adversarial manipulation, including training data poisoning and detection evasion [67]. The implications are significant: sophisticated threat actors aware of ML-based defenses may deliberately craft attacks designed to evade or manipulate these systems, creating a fundamental tension between model accuracy and adversarial robustness. This challenge is further demonstrated in malware detection

contexts, where neural classifiers remain vulnerable to adversarial perturbations that cause misclassification of malicious samples as benign [68]. Despite proposed verification methods, the fundamental trade-off between model expressiveness and formal verifiability remains unresolved, leaving operational ML-based CTI deployments without guarantees against adversarial evasion—a critical gap that motivates the exploration of alternative approaches such as LLM-based reasoning [67] [68].

Beyond FL architectures, ML-based CTI systems are generally vulnerable. CTI tools based on machine learning are themselves vulnerable to adversarial manipulation, including training data poisoning and detection evasion, according to recent research on threat models, attacks, and defenses for machine learning systems [67]. There is a fundamental conflict between model accuracy and adversarial robustness, as sophisticated threat actors aware of ML-based defenses may deliberately create attacks to evade or manipulate these systems. Neural classifiers remain susceptible to adversarial alarms that misclassify malicious samples as benign, further illustrating this difficulty in malware-detection contexts [68]. The fundamental trade-off between model expressiveness and formal verifiability remains unresolved despite proposed verification techniques, leaving operational ML-based CTI deployments without guarantees against adversarial evasion. This crucial gap motivates the investigation of alternative strategies, such as LLM-based reasoning [67].

DDoS Attacks through AI-Driven CTI

Attacks known as distributed denial-of-service (DDoS) pose a serious risk to network infrastructure and are frequently exacerbated by the proliferation of unreliable Internet of Things (IoT) devices. Many IoT devices are easily compromised and recruited into botnets such as the Mirai botnet, to launch massive attacks because they lack security measures in place, such as default login credentials and open ports. In order to counter these threats, CTI's advanced platforms automatically extract IOCs, such as malicious IP addresses, from public sources, such as social media and the dark web, using AI and NLP models [75]. Researchers in [17] suggest using blockchain technology to distribute CTI to Internet service providers (ISPs) and end users without a central authority, making this intelligence actionable against DDoS. When an intrusion detection system (such as Snort) detects a DDoS signature, it initiates a blockchain transaction that instantly propagates the threat data to other nodes, automatically updating firewalls to contain the attack and stop the source of the malicious traffic.

The IoT-CTIS system, in this study [76], which employs ML and a three-layer architecture, identifies malware in IoT devices. The Edge layer platform handles preliminary processing, the Cloud layer handles final processing and predictive analysis, and the Internet of Things layer gathers device data. The model integrates autoencoders to lower data complexity and NLP for IOC extraction, utilizing both supervised and unsupervised learning. Outperforming models like SVM and CNN, threats like DDoS botnets were identified with high precision and recall. By lowering false positives and negatives, this combination of methods improves IoT network security.

While ML/AI has significantly enhanced the automation and complexity of CTI tasks, such as data fusion, threat actor attribution, anomaly detection, and IOC extraction, several gaps remain. One major gap is the limited ability to manage evolving, dynamic, or zero-day threats. Additionally, there is a heavy reliance on labeled or structured data, which is often scarce. This issue is compounded by a lack of contextual awareness and the difficulty in applying insights across different platforms. Furthermore, high-performance models face explainability challenges, making their decision-making processes difficult to understand. Finally, there are governance and trust concerns related to decentralized or federated sharing of CTI.

LLM Methods in CTI

In the rapidly evolving field of CTI, Large Language Models (LLMs) offer significant potential to transform how threat data is processed, analyzed, and operationalized. Their ability to interpret vast volumes of unstructured information, extract meaningful insights, and support the generation of coherent intelligence reports positions them as a powerful tool for enhancing CTI workflows. LLMs can assist in automating narrative synthesis, enriching threat context, and accelerating the production of timely, actionable intelligence. As interest in leveraging LLMs for CTI continues to grow, Table 4 summarizes current research efforts exploring their integration and application in this domain.

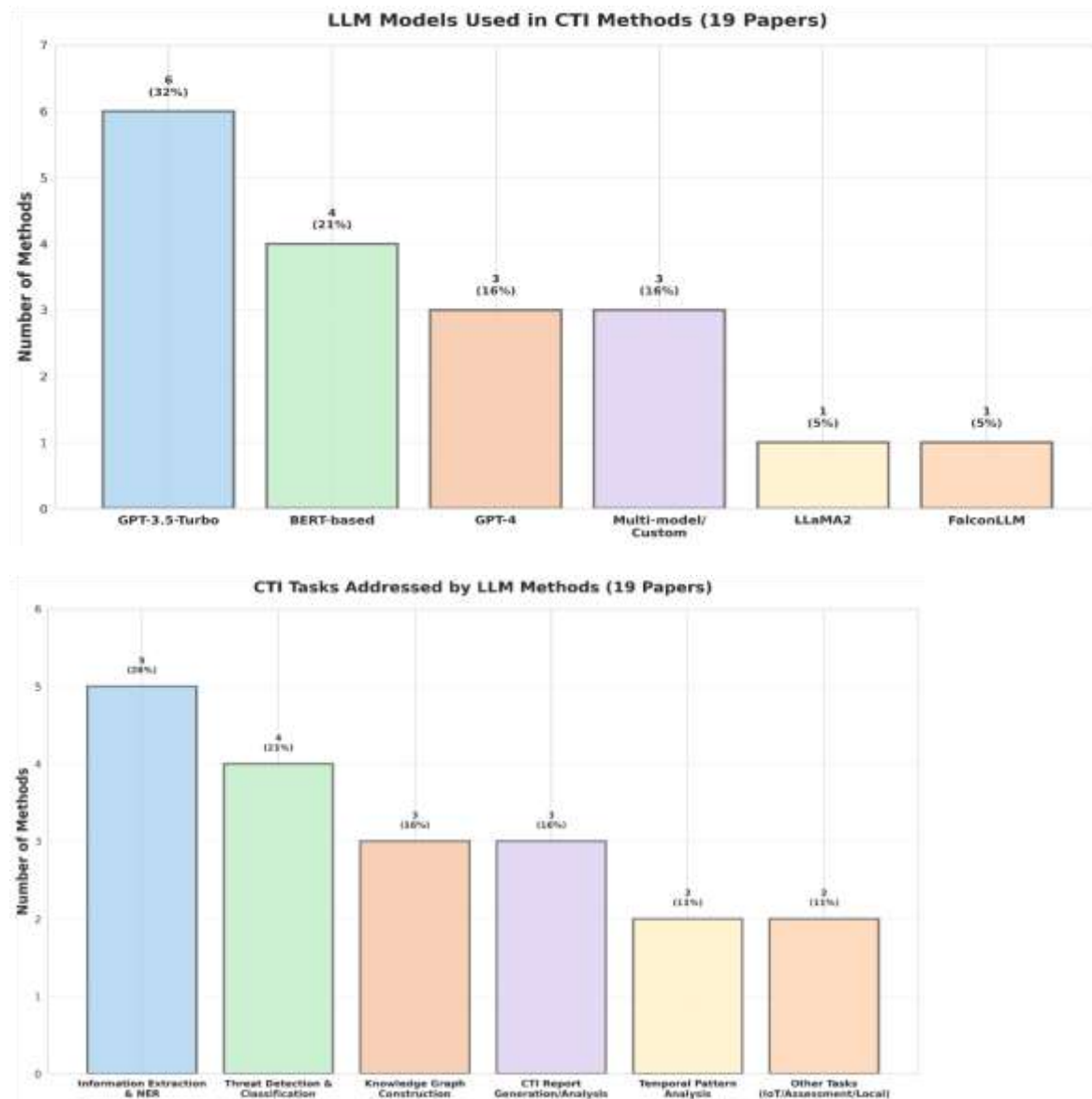
Table 4 summarizes the research work related to LLMs in CTI

No.	Method	Contribution	Technique	Data Source	Limitations	CTI Task	Evaluation
1.	AGIR [34]	Template and LLM-based framework to auto-generate CTI reports from STIX-encoded JSON	Template processing + ChatGPT	Custom STIX graphs	Depends on structured input; lacks robust evaluation; LLM used only for fluency	CTI Report Generation	Qualitative performance; lacks systematic validation
2.	SecurityLLM [35]	FLLE encoding + FalconLLM for incident response with privacy-preserving CTI classification	BERT + FalconLLM	EdgeIoTset dataset	Limited generalization to novel threats; model training constraints	Threat Detection and Incident Response	Multi-class classification performance; lacks adaptability validation
3.	Plan-act-report [36]	Prompt engineering framework for multi-stage threat analysis and decision-making	GPT-3.5-Turbo with prompt chaining	NMap scan in Docker	Simplified environment; no multi-hop logic; limited scope	Threat Analysis and Action Planning	Proof-of-concept demonstration only
4.	LLM-TIKG [37]	Automated threat extraction using few-shot GPT and LLaMA2 for knowledge graph assembly	GPT + LoRA-tuned LLaMA2	LLM-TIKG dataset from 6 sources	Only tested on LLaMA2; limited evaluation framework	Threat Knowledge Extraction and Structuring	Few-shot NER + relation extraction with partial manual correction
5.	aCTIon [38]	Zero-shot prompting with self-verification for CTI report classification and annotation	GPT-3.5-Turbo	204 CTI reports (Feb 2022–2023)	Low precision and recall; token length limits	Information Extraction and Annotation	Zero-shot LLM on benchmarked CTI dataset
6.	CRUSH [39]	LLM for constructing enriched event knowledge graphs (EKGs) for threat detection	GPT-4	Labeled scripts + blog URLs from Microsoft team	Lacks SME input; TIG complexity not fully addressed	Threat Detection and Mapping	Experimental, case-based validation
7.	Threat Behavior [44]	Evaluates GPT-3.5/4 in generating and answering MCQs for threat knowledge assessment	Few-shot prompting with GPT-3.5 & GPT-4	MITRE ATT&CK	Single dataset and small sample; SME bottleneck	Threat Knowledge Assessment	SME-rated MCQ performance; limited generalizability
8.	ChronoCTI [40]	LLM for identifying temporal patterns in	GPT-3.5	713 CTI reports with 94 curated for time relations	Lower recall; dataset constraints	Temporal Threat Pattern Extraction	Temporal relationship F1 score; limited test size

		threat campaigns					
9.	SecurityBERT [41]	15-layer BERT with traffic encoding for IoT threat classification	PPFLE + BBPE + BERT	EdgeIoTset	Only handles classification; lacks interpretability	IoT Threat Classification	Outperforms DL baselines; lacks interpretability test
10.	LocalIntel [42]	Development of a retrieval-enhanced QA framework for contextualized local CTI.	GPT-3.5-turbo	326 organization-specific local knowledge databases.	Dependency on vector database; only used for formulating defensive strategies in text.	Threat analysis, recommendation generation	Demonstrated response quality via internal tests; no large-scale benchmark used.
11.	Chatbots [43]	Evaluation of pre-trained LLMs for attack detection and NER.	ChatGPT, GPT4all, Dolly, Alpaca, Falcon	IoCs from Twitter OSINT; 31,281 constructed and annotated questions.	Limited domain-specific NER accuracy; lack of fine-tuning.	Threat entity detection, attack identification	Manual annotation accuracy and performance metrics compared across LLMs.
12.	Virtual Assistant [45]	Security mechanisms against manipulation in LLM-integrated virtual assistants.	Azure GPT-3.5 Turbo	Simulated Xbox customer support agent scenario.	Text-only assistant; lacks integration with real mitigation systems.	Threat detection support, defense advisory	Scenario-based evaluation; lacks quantitative benchmarks.
13.	Vulcan [49]	NER and RE models for recognizing cyber threats from unstructured CTI.	BERT (MLM + NSP)	ThreatPost and Malwarebytes; ~540,000 articles with 6,791 CTI instances and 4,323 relations.	Fixed data schema limits adaptability to emerging threats.	Threat entity extraction, relationship mapping	Evaluated on annotated dataset; precision and recall reported.
14.	LLM Cybersecurity Applications [69]	Comprehensive systematic review of LLM applications	Cybersecurity and LLM literature	Systematic literature review covering LLM applications in threat detection, vulnerability assessment, and CTI analysis; highly influential (198 citations)	Rapidly evolving field; incomplete coverage of emerging LLM techniques; evaluation standardization challenges	Threat Detection / Vulnerability Assessment / CTI Analysis	Systematic literature review; comprehensive coverage
15.	Transformers for Intrusion Detection [71]	LLM-based and Transformer intrusion detection systems	Network traffic and intrusion detection datasets	Reviews Transformers and LLMs in cyber-threat detection; demonstrates LLM	Context window limitations for large network flows; computational overhead for	Intrusion Detection / Threat Detection	Detection accuracy metrics; real-time performance evaluation

				effectiveness for real-time intrusion detection	real-time detection; training data requirements		
16.	ChronoCTI [70]	Temporal attack pattern extraction using MITRE ATT&CK	CTI reports and MITRE ATT&CK framework	Develops ChronoCTI for mining temporal attack patterns; enables understanding of attack progression and evolution over time	Temporal data sparsity in CTI reports; MITRE ATT&CK mapping gaps; attack pattern generalization challenges	Attack Pattern Analysis / Temporal Intelligence	Pattern extraction accuracy; MITRE ATT&CK alignment validation
17.	CTINexus [71]	Automatic CTI knowledge graph construction using LLMs	CTI reports and threat intelligence documents	LLM-powered framework for automated CTI knowledge extraction and cyber security knowledge graph (CSKG) construction; enables semantic threat relationship modeling	Knowledge graph completeness challenges; entity disambiguation difficulties; relationship extraction accuracy limitations	Knowledge Graph Construction / Information Extraction	CSKG quality metrics; entity and relation accuracy assessment
18.	GNN with Domain Knowledge [72]	Graph Neural Networks embedded with domain knowledge	CTI knowledge graphs and threat relationships	Combines GNNs with domain knowledge for enhanced CTI knowledge graph construction and threat relationship modeling; improves semantic understanding	Domain knowledge encoding complexity; graph scalability with large threat networks; knowledge graph maintenance overhead	Knowledge Graph Construction / Threat Relationship Modeling	Knowledge graph quality metrics; relationship extraction accuracy
19.	Multi-Agent LLM System[73]	Proposes a multi-agent architecture of large language models to automate CTI analysis, reducing human workload in threat report interpretation and intelligence extraction	Multi-agent coordination of LLMs with specialized roles to perform CTI analysis collaboratively	Unstructured CTI text sources (cyber threat reports, advisories, security blogs)	Primarily textual analysis, no real-time detection, limited empirical evaluation beyond qualitative benchmarks	CTI report analysis, extraction, and summarization	Qualitative evaluation of system workflow demonstrating automation potential; no standardized benchmark metrics reported

Analysis of LLM-based CTI methods illustrated in Figure 7 shows that they primarily use GPT-family and BERT-based models, with a growing but still restricted use of multimodal and open-source LLMs. These methods, which reflect a move toward contextual reasoning and narrative intelligence, primarily address high-level semantic CTI tasks, including information extraction, threat detection and classification, knowledge graph construction, and automated CTI report generation. While operational performance comparisons and expert-driven assessments are still less common, evaluation practices are primarily focused on conventional precision, recall, and F1-based metrics, supplemented by benchmark datasets and qualitative case studies. The need for more reliable, scalable, and operationally validated LLM-centric CTI frameworks is highlighted by the significant limitations of LLM-based CTI approaches, despite their potential. These limitations include limited generalization, evaluation gaps, context and token constraints, data quality dependency, and issues with scalability, real-time capability, and interpretability.



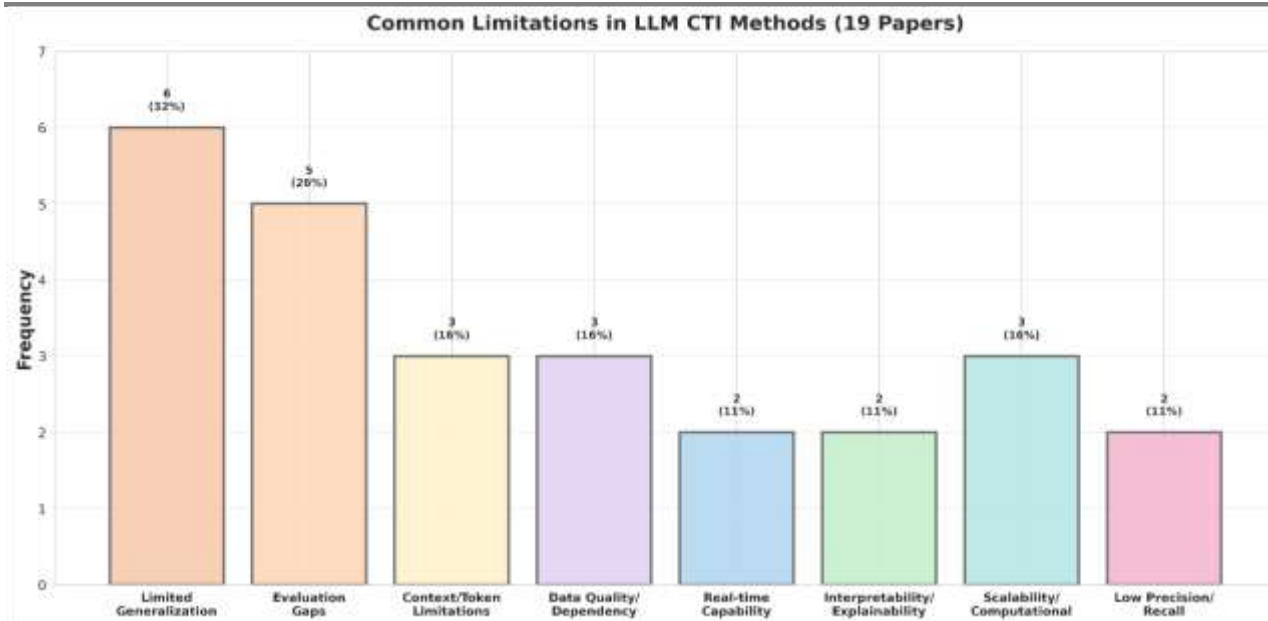


Figure 7. LLM-Based CTI Methods analysis

Findings and Reflections on LLM Methods in CTI

Beyond early investigations, the use of LLMs in CTI has advanced rapidly, with current studies tackling key issues related to deployment, accuracy, and operational integration. An authoritative analysis of LLMs across cybersecurity domains is provided by a thorough systematic review published in [70], which records more than 180 studies covering threat intelligence extraction, malware analysis, vulnerability detection, and phishing detection. This work demonstrates that although LLMs show impressive zero-shot and few-shot learning capabilities for threat classification, serious concerns about hallucination risks—where models produce plausible but factually incorrect intelligence—remain. For mission-critical CTI applications, the review highlights the necessity of grounding mechanisms and human-in-the-loop validation.

Another study reveals that LLM-enabled systems such as AGIR [34] and aCTIon [38] demonstrate an early maturation of natural language generation in CTI. AGIR introduces a dual-pipeline that leverages structured STIX graphs for template-driven report construction, later refined by ChatGPT. While this approach bridges structured threat data with narrative intelligence, its reliance on predefined templates limits its adaptability to emerging threat behaviors. Similarly, aCTIon's two-stage LLM pipeline effectively mitigates hallucination and token limitations, offering reliable extraction from unstructured reports. Its reliance on fixed input lengths and zero-shot prompting, however, limits its scalability across various CTI formats and domains.

SecurityLLM [35] and SecurityBERT [41] exemplify efforts to merge LLMs with classic classification models. By incorporating novel encodings (e.g., FLLE and BBPE), these systems achieve higher accuracy in identifying cyber threats, particularly within IoT environments. SecurityBERT, in particular, achieves superior performance with lightweight deployment, proving its suitability for edge-based CTI. Yet, both models exhibit brittleness against novel threats, largely due to their dependence on fixed pretraining corpora and limited real-time adaptability.

A consistent theme across Plan-Act-Report [36], ChronoCTI [40], and CRUSH [39] is the use of LLMs for procedural and temporal reasoning. Plan-Act-Report deploys GPT-3.5 to simulate multi-stage threat campaigns using prompt chaining, reflecting early explorations of LLMs in autonomous adversarial modeling. ChronoCTI's extraction of attack timelines adds a novel temporal layer to CTI, although its limited dataset and narrow recall metrics indicate challenges in generalizing to broader threat contexts. CRUSH offers a unique attempt to build and enrich Threat Intelligence Graphs (TIGs) with LLMs. It reduces reliance on SMEs but remains experimental in balancing automation and knowledge validation.

Entity-centric methods such as LLM-TIKG [37], CyberEntRel [48], and Vulcan [49] showcase the growing

sophistication of LLMs in knowledge extraction from noisy CTI sources. Vulcan, in particular, moves beyond traditional IoC detection by uncovering semantic threat relationships through fine-tuned NER and RE models. All three models still encounter limitations in moving beyond static entity categories, highlighting a significant weakness in LLMs' ability to adaptively capture the dynamic nature of attacker tactics.

On the operational front, LOCALINTEL [42] and Threat Behavior [44] target SOC-driven applications and adversarial behavior modeling, respectively. LOCALINTEL advances retrieval-augmented LLM use for localized CTI generation, yet its dependency on vector databases and constraint to textual outputs limits its utility in active defense orchestration. Meanwhile, the MCQ-based threat behavior assessment in [44] is an ambitious step toward benchmarking LLMs' domain knowledge. Nevertheless, small-scale evaluations and format limitations raise concerns about its validity across broader real-world applications.

On the other hand, studies evaluating generalized LLM capabilities such as Chatbots [43] and Virtual Assistant [45] reveal persistent gaps in domain-specific fine-tuning. These models exhibit high performance in simple classification tasks but struggle with precise Named Entity Recognition (NER) and system integration, reflecting the ongoing tension between LLM generalization and cybersecurity-specific precision.

Furthermore, surveys on transformer-based models for intrusion detection systems have systematically examined the architectural underpinnings of LLM-based threat detection. This thorough analysis in [71] shows that self-attention mechanisms enable superior pattern recognition compared to traditional recurrent architectures, particularly for identifying sophisticated attack patterns spanning extended temporal windows. Nevertheless, as threat actors develop evasion strategies to evade attention-based classifiers, adversarial robustness remains an open research question, and the computational overhead of transformer inference poses deployment challenges for real-time network monitoring.

Investigation into continual learning approaches for cybersecurity has been encouraged by the difficulty of maintaining current threat intelligence within LLM systems. This work [72] addresses the fundamental limitation that static model parameters encode knowledge only up to training cutoff dates, rendering models progressively obsolete as new threats emerge. Without catastrophically forgetting previously learned threat patterns, the proposed framework supports incremental knowledge updates. However, continuous learning adds complexity to validation; comprehensive regression testing is necessary to ensure that knowledge updates do not impair performance across recognized threat categories.

Additional research on [73] multi-agent LLM systems for CTI analysis shows that coordinating specialized agents can automate complex CTI workflows and lessen analyst workload. Each agent is responsible for tasks such as intelligence extraction, contextual reasoning, and report synthesis. The study demonstrates that, unlike single-model approaches, distributing CTI analysis across multiple cooperating LLM agents enables a more thorough, organized interpretation of unstructured threat reports. Individual agents or prompts can be improved separately due to this architecture's support for modularity and scalability. The study, however, also highlights issues with agent coordination, growing system complexity, and reliance on qualitative assessment, all of which could prevent immediate deployment in real-time or mission-critical CTI environments.

The reviewed literature shows that while LLMs significantly improve the automation of CTI tasks, they typically operate within isolated, task-specific silos. Most models concentrate on a single aspect—such as extraction, generation, or classification—without addressing the comprehensive needs of end-to-end threat intelligence. Additionally, a common limitation across studies is reliance on synthetic or narrowly defined datasets, which undermines the ecological validity of the results. Only a few models exhibit resilience to adversarial manipulation, adapt to evolving threats, or handle the complexities of multilingual and cross-platform intelligence.

Comparative Evaluation Across CTI Paradigms

Sections 3.1 to 3.3 report evaluation results using inconsistent combinations of metrics, which limits direct comparison across conventional, AI/ML-based, and LLM-based CTI methods. Table 5 consolidates the performance indicators reported across Tables 2 to 4 against six common criteria: detection accuracy,

precision/recall, false-positive rate, response time and real-time capability, and computational cost.

Table 5 Comparative summary of performance indicators reported across conventional, AI/ML-based, and LLM-based CTI methods

CTI Paradigm	Detection Accuracy	Precision / Recall	False-Positive Rate	Response Time / Real-Time Capability	Computational Cost
Conventional Methods (Table 2)	Rarely reported as a standalone metric; where present, ranges from qualitative validation (e.g., BLOCIS [18], ABC [11]) to ~72% entity-recognition accuracy (SECDFAN [15]).	Reported in a minority of studies; AttackKG reports precision and recall of approximately 84% and 79% for attack-technique extraction [14].	Not reported in any of the 18 reviewed conventional studies.	Not benchmarked; several methods explicitly note the absence of real-time capability (Eight-step CTI pipeline [13]; CyRiPred [12]).	Not quantified; discussed only qualitatively in terms of blockchain overhead or integration complexity (BLOCIS [18]; BlockIntelChain [56]).
AI/ML-Based Methods (Table 3)	More consistently reported, ranging from ~72% entity-extraction accuracy (CTI-View [33]) to high classification accuracy on benchmark datasets (HinCTI [30]).	Commonly reported for entity-relation-extraction tasks; CyberEntRel reports precision, recall, and F1-score on annotated CTI corpora [48].	Occasionally reported for intrusion-detection variants (Federated Learning IDS [22]) but absent from most entity-extraction and classification studies.	Rarely benchmarked; real-time performance is claimed but not measured for models such as BERT-GRU-BiLSTM-CRF, which the authors note carries high computational complexity [64].	Discussed qualitatively as a deployment barrier for transformer-based models [64]; not reported in FLOPs, latency, or hardware terms.
LLM-Based Methods (Table 4)	Task-dependent and inconsistently reported; SecurityBERT outperforms deep-learning baselines on IoT threat classification [41], while aCTIon reports comparatively low precision and recall on zero-shot extraction [38].	Reported for extraction-oriented tasks (Vulcan [49]; CyberEntRel-style models); temporal-reasoning studies such as ChronoCTI instead report F1-score for relation extraction [40].	Not reported in any of the 19 reviewed LLM-based studies.	Identified as an open challenge rather than measured; transformer inference overhead is flagged as a barrier to real-time deployment [71].	Noted qualitatively as a barrier to large-scale adoption (training and inference cost of large-scale LLMs) but not quantified in any reviewed study.

Table 5 makes explicit a pattern that recurs throughout Sections 3.1 to 3.3: reporting of detection accuracy, and to a lesser extent precision and recall, becomes more consistent moving from conventional to AI/ML-based and

LLM-based methods, but false-positive rate, response time, and computational cost remain almost universally unreported across all three paradigms. This gap prevents like-for-like comparison of conventional, AI/ML-based, and LLM-based CTI methods on operational grounds and motivates the standardized evaluation framework proposed in Section 5.

Enhanced LLM-Based CTI Framework

Building upon the insights gained from conventional, machine learning-based, and LLM-enhanced approaches to CTI this study introduces an enhanced, integrated framework that strategically leverages the strengths of LLMs. While traditional and ML-based CTI methods have significantly enhanced foundational threat detection capabilities, they often fall short in adaptability, scalability, and semantic understanding. In contrast, LLMs demonstrate transformative potential by effectively processing unstructured data, enabling contextual reasoning, and generating structured intelligence outputs. Nevertheless, the operational success of LLMs in CTI depends on domain-specific fine-tuning, robust performance evaluation, and seamless integration into threat analysis workflows. To address these needs, we propose a tailored LLM-based framework that enables automated threat forecasting, strategic mitigation planning, and head-to-head performance assessment against general-purpose LLMs. As depicted in Figure 8, the enhanced model is embedded across the CTI lifecycle, serving multiple functions, including data ingestion and processing, analysis, and the development of mitigation strategies, thereby streamlining and improving the end-to-end threat response process.

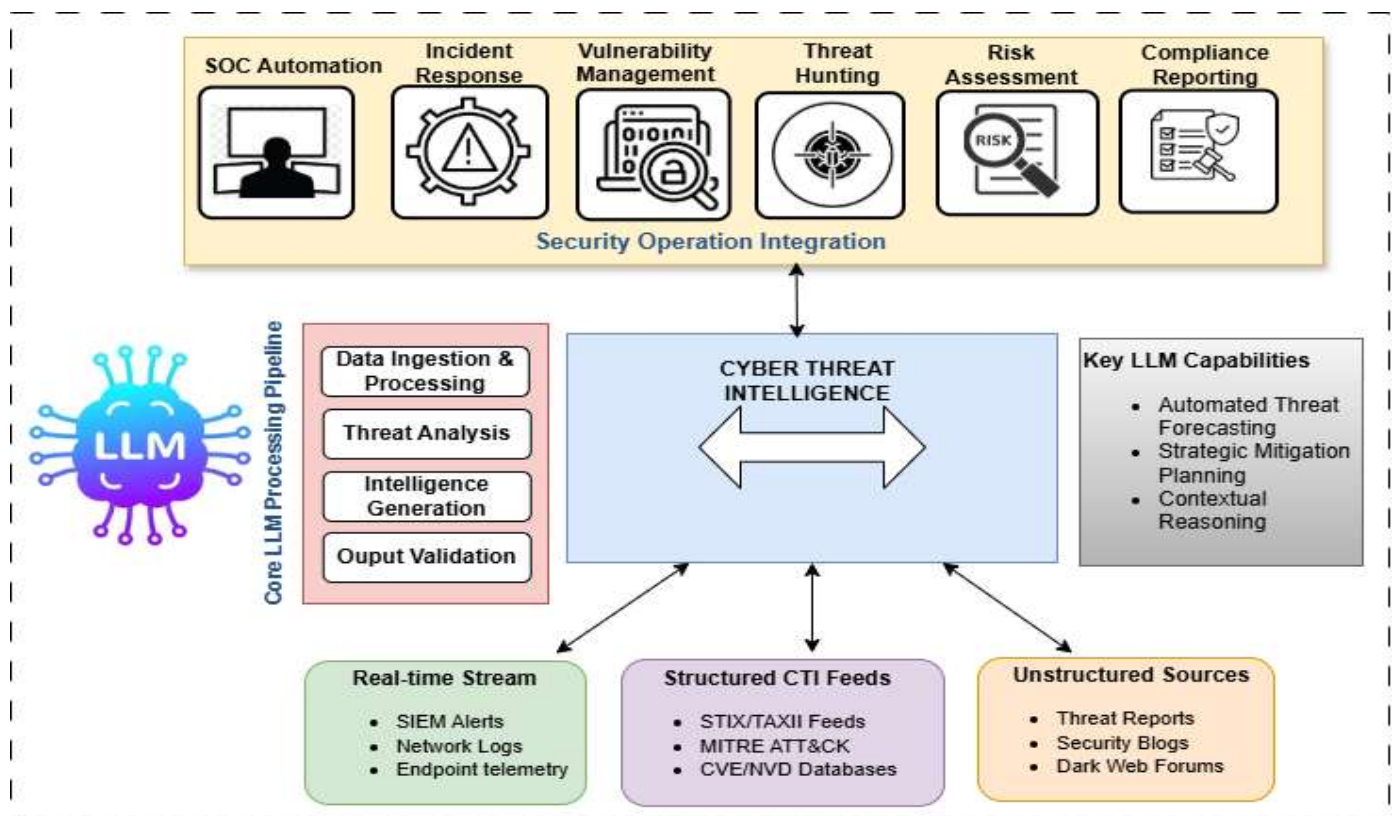


Figure 8. LLM-driven Cyber Threat Intelligence architecture integrating multi-source heterogeneous data

Challenges And Future Research Directions

Even though they are fundamental, conventional CTI techniques have serious drawbacks when confronting contemporary, sophisticated cyber threats. These techniques are ineffective against APT malware and zero-day attacks because they are frequently manual and rely on signature-based detection. They struggle to keep up with the quick evolution of threats, and their static nature leads to a high volume of false positives and negatives. Their capacity to examine the enormous volumes of unstructured data, which contain numerous threat indicators, is constrained by their reliance on structured data and predetermined correlation rules. Additionally, traditional CTI is frequently compartmentalized and lacks the more comprehensive context that results from cross-organizational data sharing, which is frequently impeded by privacy concerns. Consequently, conventional CTI's

future depends on how well it integrates with more intelligent and dynamic systems. These techniques need to advance toward more automated data gathering and analysis that incorporate information from a broader range of sources, including the deep and dark web.

Though it has greatly improved threat detection and analysis, AI/ML-based CTI is not without its challenges. The reliance on substantial amounts of high-quality, labeled data for training—which is frequently lacking in the cybersecurity field—is a major drawback. Adversarial attacks, in which malicious actors alter input data to avoid detection or contaminate the model itself, can affect these models. Many AI/ML models operate as "black boxes," lacking transparency and making it difficult for security analysts to understand their logic and trust their results. Furthermore, the high computational cost and complexity of some models can be a barrier to their adoption in resource-constrained environments. Because cyber threats are dynamic, models can quickly become outdated, necessitating ongoing retraining and adaptation. Future developments in AI/ML-based CTI will focus on addressing interpretability, model robustness, and data dependency. Building successful models with limited labeled data will require developing few-shot learning, transfer learning, and unsupervised learning strategies. More robust and resilient models that can withstand attacks will result from research into adversarial machine learning. A major area of attention will be explainable AI (XAI), which aims to develop more transparent and comprehensible models that can offer convincing explanations for their choices.

LLMs have demonstrated significant promise in CTI, especially in producing human-readable reports and analyzing unstructured text. They do, however, have some serious drawbacks. In a security setting, LLMs' propensity for "hallucinations," which produce believable but false information, can have detrimental effects. Their expertise is restricted to the data they were trained on, which may not contain the most recent threats and may be out of date. Additionally, LLMs are susceptible to prompt injection attacks, which exploit malicious input to bypass safety filters and produce dangerous content. One major obstacle to the widespread adoption of large-scale LLMs may be the high cost of training and operating them. The potential for abuse in producing phishing emails or other malicious content, as well as the privacy implications of training models on sensitive data, are additional ethical issues surrounding the use of LLMs. In the future, LLM-based CTI will focus on improving its security, accuracy, and reliability. To improve LLMs' understanding and reduce hallucinations, methods for fine-tuning them on domain-specific cybersecurity data will be developed. LLMs will receive more up-to-date and contextually relevant information when integrated with external knowledge bases and real-time threat intelligence feeds. It is also essential to conduct research on strengthening LLMs against adversarial attacks. They will become more widely available as smaller, more specialized, and more affordable LLMs are developed. To fully utilize LLMs in cybersecurity while reducing risks, it will be crucial to establish clear ethical guidelines and best practices for their responsible development and deployment.

Toward a standardized evaluation framework: A recurring limitation identified across Tables 2 through 4 is the absence of a common evaluation protocol for CTI systems: studies report inconsistent combinations of accuracy, precision, recall, and F1-score, rarely disclose false-positive rates or response latency, and evaluate on incompatible, often privately collected datasets. Future research should adopt a standardized evaluation framework for LLM-based CTI that specifies common CTI datasets, task definitions, and metrics for detection, attribution, forecasting, summarization, and automated reporting, enabling like-for-like comparison across conventional, AI/ML-based, and LLM-based approaches. Operational metrics such as detection accuracy, precision, recall, false-positive rate, response time, and computational cost should be reported alongside task-specific measures so that both predictive performance and deployment feasibility can be assessed jointly.

Hybrid architectures: Given the hallucination and reliability concerns documented throughout Section 3.3, hybrid architectures that combine LLMs with traditional threat-intelligence platforms, knowledge graphs, rule-based systems, and machine-learning classifiers represent a promising direction for future work. Coupling LLM-based reasoning with structured, verifiable knowledge sources, such as STIX/TAXII feeds, MITRE ATT&CK-aligned knowledge graphs, or ensemble ML classifiers, can constrain LLM outputs to validated intelligence, reduce hallucination, and preserve the contextual reasoning and unstructured-text processing capabilities that make LLMs valuable for CTI. Systematically evaluating these hybrid designs against the standardized framework proposed above would help determine which combinations of components yield the greatest gains in reliability without sacrificing the scalability benefits of LLM-based automation.

REFERENCES

1. Tounsi, W., & Rais, H. (2018). A survey on technical threat intelligence in the age of sophisticated cyber attacks. *Computers & security*, 72, 212-233.
2. Sharif, M. H. U., & Mohammed, M. A. (2022). A literature review of financial losses statistics for cyber security and future trend. *World Journal of Advanced Research and Reviews*, 15(1), 138-156.
3. Ainslie, S., Thompson, D., Maynard, S., & Ahmad, A. (2023). Cyber-Threat Intelligence for Security Decision-Making: A Review and Research Agenda for Practice. *Computers & Security*, 103352.
4. Kotsias, J., Ahmad, A., & Scheepers, R. (2023). Adopting and integrating cyber-threat intelligence in a commercial organisation. *European Journal of Information Systems*, 32(1), 35-51.
5. PwC. (2024). PwC's 24th Annual Global CEO Survey: CEOs on their tech concerns., Report by PwC annual survey of CEO on IT or technology concerns. . <https://www.pwc.com/gx/en/issues/c-suite-insights/ceo-survey.html>
6. Lin, P. C., Hsu, W. H., Lin, Y. D., Hwang, R. H., Wu, H. K., Lai, Y. C., & Chen, C. K. (2023). Correlation of cyber threat intelligence with sightings for intelligence assessment and augmentation. *Computer Networks*, 228, 109736.
7. Saeed, S., Suayyid, S. A., Al-Ghamdi, M. S., Al-Muhaisen, H., & Almuhaideb, A. M. (2023). A systematic literature review on cyber threat intelligence for organizational cybersecurity resilience. *Sensors*, 23(16), 7273.
8. Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., ... & Moher, D. (2021). The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *International journal of surgery*, 88, 105906.
9. Zhang, S., Chen, P., Bai, G., Wang, S., Zhang, M., Li, S., & Zhao, C. (2022). An automatic assessment method of cyber threat intelligence combined with ATT&CK matrix. *Wireless Communications and Mobile Computing*, 2022.
10. Serketzis, N., Katos, V., Ilioudis, C., Baltatzis, D., & Pangalos, G. (2019). Improving forensic triage efficiency through cyber threat intelligence. *Future Internet*, 11(7), 162.
11. Chatziamanetoglou, D., & Rantos, K. (2023). Blockchain-Based Cyber Threat Intelligence Sharing Using Proof-of-Quality Consensus. *Security and Communication Networks*, 2023.
12. Kia, A. N., Murphy, F., Sheehan, B., & Shannon, D. (2024). A cyber risk prediction model using common vulnerabilities and exposures. *Expert Systems with Applications*, 237, 121599.
13. Borges Amaro, L. J., Percilio Azevedo, B. W., Lopes de Mendonca, F. L., Giozza, W. F., Albuquerque, R. D. O., & Garcia Villalba, L. J. (2022). Methodological framework to collect, process, analyze and visualize cyber threat intelligence data. *Applied Sciences*, 12(3), 1205.
14. Li, Z., Zeng, J., Chen, Y., & Liang, Z. (2022, September). AttackKG: Constructing technique knowledge graph from cyber threat intelligence reports. In *European Symposium on Research in Computer Security* (pp. 589-609). Cham: Springer International Publishing.
15. Sakellariou, G., Fouliras, P., & Mavridis, I. (2023). SECDFAN: A Cyber Threat Intelligence System for Discussion Forums Utilization. *Eng*, 4(1), 615-634.
16. Sacher-Boldewin, D., & Leverett, E. (2022). The Intelligent Process Lifecycle of Active Cyber Defenders. *Digital Threats: Research and Practice (DTRAP)*, 3(3), 1-17.
17. Mendez Mena, D., & Yang, B. (2021). Decentralized actionable cyber threat intelligence for networks and the internet of things. *IoT*, 2(1), 1-16. <https://doi.org/10.3390/iot2010001>
18. Gong, S., & Lee, C. (2020). Blocis: blockchain-based cyber threat intelligence sharing framework for sybil-resistance. *Electronics*, 9(3), 521.
19. Chen, T., Zeng, H., Lv, M., & Zhu, T. (2024). CTIMD: Cyber threat intelligence enhanced malware detection using API call sequences with parameters. *Computers & Security*, 136, 103518.
20. Irshad, E., & Siddiqui, A. B. (2023). Cyber threat attribution using unstructured reports in cyber threat intelligence. *Egyptian Informatics Journal*, 24(1), 43-59.
21. Sufi, F. (2023). A New Social Media-Driven Cyber Threat Intelligence. *Electronics*, 12(5), 1242.
22. Sarhan, M., Layeghy, S., Moustafa, N., & Portmann, M. (2023). Cyber threat intelligence sharing scheme based on federated learning for network intrusion detection. *Journal of Network and Systems Management*, 31(1), 3.
23. Bayer, M., Frey, T., & Reuter, C. (2023). Multi-level fine-tuning, data augmentation, and few-shot learning

- for specialized cyber threat intelligence. *Computers & Security*, 134, 103430.
24. Keim, Y., & Mohapatra, A. K. (2022). Cyber threat intelligence framework using advanced malware forensics. *International Journal of Information Technology*, 14(1), 521-530.
25. Gao, P., Liu, X., Choi, E., Ma, S., Yang, X., Ji, Z., ... & Song, D. (2022). ThreatKG: A Threat Knowledge Graph for Automated Open-Source Cyber Threat Intelligence Gathering and Management. *arXiv preprint arXiv:2212.10388*.
26. Koloveas, P., Chantzios, T., Alevizopoulou, S., Skiadopoulos, S., & Tryfonopoulos, C. (2021). intime: A machine learning-based framework for gathering and leveraging web data to cyber-threat intelligence. *Electronics*, 10(7), 818.
27. Zhao, J., Yan, Q., Li, J., Shao, M., He, Z., & Li, B. (2020). TIMiner: Automatically extracting and analyzing categorized cyber threat intelligence from social data. *Computers & Security*, 95, 101867.
28. Kristiansen, L. M., Agarwal, V., Franke, K., & Shah, R. S. (2020, December). CTI-Twitter: gathering cyber threat intelligence from twitter using integrated supervised and unsupervised learning. In *2020 IEEE International Conference on Big Data (Big Data)* (pp. 2299-2308). IEEE.
29. Liu, J., Yan, J., Jiang, J., He, Y., Wang, X., Jiang, Z., ... & Li, N. (2022). TriCTI: an actionable cyber threat intelligence discovery system via trigger-enhanced neural network. *Cybersecurity*, 5(1), 8.
30. Gao, Y., Li, X., Peng, H., Fang, B., & Philip, S. Y. (2020). Hincti: A cyber threat intelligence modeling and identification system based on heterogeneous information network. *IEEE Transactions on Knowledge and Data Engineering*, 34(2), 708-722.
31. Suryotrisongko, H., Musashi, Y., Tsuneda, A., & Sugitani, K. (2022). Robust botnet DGA detection: Blending XAI and OSINT for cyber threat intelligence sharing. *IEEE Access*, 10, 34613-34624.
32. Jiang, T., Shen, G., Guo, C., Cui, Y., & Xie, B. (2023). BFLS: Blockchain and Federated Learning for sharing threat detection models as Cyber Threat Intelligence. *Computer Networks*, 224, 109604.
33. Zhou, Y., Tang, Y., Yi, M., Xi, C., & Lu, H. (2022). CTI view: APT threat intelligence analysis system. *Security and Communication Networks*, 2022, 1-15.
34. Perrina, F., Marchiori, F., Conti, M., & Verde, N. V. (2023, December). AGIR: Automating Cyber Threat Intelligence Reporting with Natural Language Generation. In *2023 IEEE International Conference on Big Data (BigData)* (pp. 3053-3062). IEEE.
35. Ferrag, M. A., Ndhlovu, M., Tihanyi, N., Cordeiro, L. C., Debbah, M., & Lestable, T. (2023). Revolutionizing Cyber Threat Detection with Large Language Models. *arXiv preprint arXiv:2306.14263*.
36. Moskal, S., Laney, S., Hemberg, E., & O'Reilly, U. M. (2023). LLMs Killed the Script Kiddie: How Agents Supported by Large Language Models Change the Landscape of Network Threat Testing. *arXiv preprint arXiv:2310.06936*.
37. Hu, Y., Zou, F., Han, J., Sun, X., & Wang, Y. (2024). Llm-Tikg: Threat Intelligence Knowledge Graph Construction Utilizing Large Language Model. *Computers & Security*, 145, 103999. <https://doi.org/10.1016/j.cose.2024.103999>.
38. Siracusano, G., Sanvito, D., Gonzalez, R., Srinivasan, M., Kamatchi, S., Takahashi, W., ... & Bifulco, R. (2023). Time for aCTIon: Automated Analysis of Cyber Threat Intelligence in the Wild. *arXiv preprint arXiv:2307.10214*.
39. Sewak, M., Emani, V., & Naresh, A. (2023). CRUSH: Cybersecurity Research using Universal LLMs and Semantic Hypernetworks.
40. Rahman, M. R., Wroblewski, B., Matthews, Q., Morgan, B., Menzies, T., & Williams, L. (2024). Mining Temporal Attack Patterns from Cyberthreat Intelligence Reports. *arXiv preprint arXiv:2401.01883*.
41. Ferrag, M. A., Ndhlovu, M., Tihanyi, N., Cordeiro, L. C., Debbah, M., Lestable, T., & Thandi, N. S. (2024). Revolutionizing Cyber Threat Detection with Large Language Models: A privacy-preserving BERT-based Lightweight Model for IoT/IIoT Devices. *IEEE Access*.
42. Mitra, S., Neupane, S., Chakraborty, T., Mittal, S., Piplai, A., Gaur, M., & Rahimi, S. (2024). LOCALINTEL: Generating Organizational Threat Intelligence from Global and Local Cyber Knowledge. *arXiv preprint arXiv:2401.10036*.
43. Shafee, S., Bessani, A., & Ferreira, P. M. (2024). Evaluation of LLM Chatbots for OSINT-based Cyberthreat Awareness. *arXiv preprint arXiv:2401.15127*.
44. Garza, E., Hemberg, E., Moskal, S., & O'Reilly, U. M. (2023). Assessing Large Language Model's knowledge of threat behavior in MITRE ATT&CK.
45. Chan, C. F., Yip, D. W., & Esmradi, A. (2024). Detection and Defense Against Prominent Attacks on

- Preconditioned LLM-Integrated Virtual Assistants. arXiv preprint arXiv:2401.00994.
46. Iqbal, Z., & Anwar, Z. (2020). SCERM—A novel framework for automated management of cyber threat response activities. *Future Generation Computer Systems*, 108, 687-708.
 47. Mohan, J. S., Thirunavukkarasu, M., Kumaran, N., & Thamaraiselvi, D. (2024). Deep Learning with Blockchain Based Cyber Security Threat Intelligence and Situational Awareness System for Intrusion Alert Prediction. *Sustainable Computing: Informatics and Systems*, 100955.
 48. Ahmed, K., Khurshid, S. K., & Hina, S. (2024). CyberEntRel: Joint extraction of cyber entities and relations using deep learning. *Computers & Security*, 136, 103579.
 49. Jo, H., Lee, Y., & Shin, S. (2022). Vulcan: Automatic extraction and analysis of cyber threat intelligence from unstructured text. *Computers & Security*, 120, 102763.
 50. Zacharis, A., Katos, V., & Patsakis, C. (2024). Integrating AI-driven threat intelligence and forecasting in the cybersecurity exercise content generation lifecycle. *International Journal of Information Security*, 23, 2691–2710. <https://doi.org/10.1007/s10207-024-00860-w>
 51. B. D. Le, G. Wang, M. Nasim, and M. A. Babar, "Gathering Cyber Threat Intelligence from Twitter Using Novelty Classification," in *Proc. 2019 Int. Conf. Cyberworlds (CW)*, 2019, pp. 316–323. [Online]. Available: <https://doi.org/10.1109/CW.2019.00058>
 52. Kaspersky, "What is an Advanced Persistent Threat (APT)?," Kaspersky, 2023. [Online]. Available: <https://www.kaspersky.com/resource-center/threats/advanced-persistent-threats>. [Accessed: Jul. 5, 2025].
 53. C. Whitman, "Leveraging cyber threat intelligence mining for enhanced proactive cybersecurity: A comprehensive review and future directions," *International Journal of Cyber Threat Intelligence and Secure Networking*, pp. 14–19, Dec. 15, 2024.
 54. P. Santos, R. Abreu, M. J. C. S. Reis, C. Serôdio, and F. Branco, "A systematic review of cyber threat intelligence: The effectiveness of technologies, strategies, and collaborations in combating modern threats," *Sensors*, vol. 25, no. 14, Art. no. 4272, Jul. 2025, doi: 10.3390/s25144272.
 55. G. Cascavilla, S. Gupta, F. Massacci, and J. Camacho, "Cybercrime threat intelligence: A systematic multi-vocal literature review," *Computers & Security*, vol. 105, Art. no. 102442, Jun. 2021, doi: **10.1016/j.cose.2021.102442**.
 56. A. Tolah, "BlockIntelChain: A blockchain-based cyber threat intelligence sharing framework," *Scientific Reports*, vol. 15, Art. no. 29152, 2025. [Online]. Available:
 57. H. El Amin, A. E. Samhat, M. Chamoun, and L. Oueidat, "An integrated approach to cyber risk management with cyber threat intelligence framework to secure critical infrastructure," *Journal of Cybersecurity and Privacy*, vol. 4, no. 2, Art. no. 18, 2024. [Online]. Available:
 58. A. Mahida and A. Tyagi, "Cyber threat intelligence and information sharing in cloud ecosystems," *Proceedings on Engineering*, vol. 7, no. 1, pp. 1–12, 2025. [Online]. Available:
 59. P. Fuxen, M. Hachani, R. Hackenberg, and M. Ross, "MANTRA: Towards a conceptual framework for elevating cybersecurity applications through privacy-preserving cyber threat intelligence sharing," in *Proc. 15th Int. Conf. Cloud Computing, GRIDs, and Virtualization (CLOUD COMPUTING 2024)*, 2024, pp. 44–50.
 60. D. M. Janosek, "Toward a global framework for cyber threat intelligence sharing," *Cyber Defense Review*, vol. 10, no. 2, 2025. [Online]. Available:
 61. Y. Zhou et al., "A blockchain based efficient incentive mechanism in cyber threat intelligence sharing," *Journal of Information Security and Applications*, vol. 82, Art. no. 103769, 2025. [Online]. Available:
 62. J. Komarthy, "Optimizing threat intelligence sharing across multiple security platforms," *Emerging Frontiers Library for The American Journal of Engineering and Technology*, vol. 3, no. 1, Art. no. 552, 2025.
 63. Varbanov, V., et al. (2024). Advancing cyber threat intelligence through machine learning algorithms. *ACM International Conference on Computing Frontiers*, Art. no. 3660879. <https://dl.acm.org/doi/fullHtml/10.1145/3660853.3660879>
 64. Demirolo, D., et al. (2025). A novel approach for cyber threat analysis systems using BERT model from cyber threat intelligence data. *Symmetry*, 17(4), 587. <https://www.mdpi.com/2073-8994/17/4/587>
 65. Ragab, M., et al. (2025). Advanced artificial intelligence with federated learning framework for privacy-preserving cyberthreat detection in IoT-assisted sustainable smart cities. *Scientific Reports*, 15, Art. no. 88843. <https://www.nature.com/articles/s41598-025-88843-2>
 66. Hamad, N. A., et al. (2025). Systematic analysis of federated learning approaches for cybersecurity. *IEEE*

- Access, 13, 1–15. <https://ieeexplore.ieee.org/document/11017512/>
67. Koball, C., et al. (2024). Machine learning security: Threat model, attacks, and defenses. *Computer*, 57(10), 42–51. <https://www.computer.org/csdl/magazine/co/2024/10/10687326/>
68. Robinette, P. K., et al. (2024). Neural network malware detection verification for feature robustness. arXiv preprint, arXiv:2404.05703. <https://arxiv.org/abs/2404.05703>
69. Mouiche and S. Saad, "Entity and relation extractions for threat intelligence knowledge graphs," *Computers & Security*, vol. 148, Art. no. 104255, 2025. <https://www.sciencedirect.com/science/article/pii/S0167404824004255>
70. M. Motlagh, A. Hajizadeh, and M. Raahemi. (2024). Large language models in cybersecurity: State-of-the-art and future directions. <https://arxiv.org/abs/2402.00891>
71. H. Kheddar, Y. Himeur, S. Al-Maadeed, A. Amira, and F. Bensaali. (2024). Transformers and large language models for efficient intrusion detection systems: A comprehensive survey. arXiv preprint, arXiv:2408.07583.
72. G. Liu, K. Lu, and S. Pi, "Graph neural networks embedded with domain knowledge for cyber threat intelligence entity and relationship mining," *PeerJ Computer Science*, vol. 11, p. e2769, 2025, doi: 10.7717/peerj-cs.2769.
73. Y. Hmimou, M. Tabaa, A. Khiat, and Z. Hidila, "A multi-agent system for cybersecurity threat detection and correlation using large language models," *IEEE Access*, vol. 13, pp. 150199-150214, 2025, doi: 10.1109/ACCESS.2025.3602681.
74. E. M. Hutchins, M. J. Cloppert, and R. M. Amin, "Intelligence-driven computer network defense informed by analysis of adversary campaigns and intrusion kill chains," *Lockheed Martin Corporation*, White Paper, 2011.
75. P. Balasubramanian, S. Nazari, D. K. Khlogh, A. Mahmoodi, J. Seby, and P. Kostakos, "A cognitive platform for collecting cyber threat intelligence and real-time detection using cloud computing," *Decision Analytics Journal*, vol. 14, p. 100545, 2025, doi: 10.1016/j.dajour.2025.100545.
76. P. Xiao, "Malware cyber threat intelligence system for Internet of Things (IoT) using machine learning," *Journal of Cyber Security and Mobility*, vol. 13, no. 1, pp. 53–90, Dec. 2023, doi: 10.13052/jcsm2245-1439.1313.