

PROMPT ENGINEERING IN CYBERSECURITY – ACHIEVING TECHNOLOGICAL EDGE

Iustin PRIESCU

“Titu Maiorescu” University, Bucharest, Romania
iustin.priescu@prof.utm.ro

Geanina Silvana BANU

“Titu Maiorescu” University, Bucharest, Romania
geanina.banu@prof.utm.ro

Tatiana Corina DOESCU

“Titu Maiorescu” University, Bucharest, Romania
tatiana.doescu@prof.utm.ro

Marius Irinel BANU

“Ferdinand I” Military Technical Academy, Bucharest, Romania
marius.banu@mta.ro

ABSTRACT

Prompt Engineering has emerged as a critical technique for optimizing interactions with Large Language Models (LLMs), enhancing their precision and effectiveness in various applications. This paper explores the role of structured prompting techniques in cybersecurity, focusing on their potential to improve threat detection, security automation, and incident response. By leveraging Chain-of-Thought (CoT) and Multimodal CoT prompting, the study evaluates how well-crafted prompts enhance LLM capabilities in analyzing network traffic, identifying phishing attempts, automating security reporting, and conducting forensic investigations. The findings demonstrate that structured prompts significantly improve LLM performance, enabling more accurate anomaly detection, faster security incident analysis, and enhanced cyber threat intelligence. Additionally, the study outlines a framework for integrating Prompt Engineering into cybersecurity workflows, illustrating its potential for AI-driven security monitoring and defense strategies. These insights contribute to the ongoing advancement of AI applications in cybersecurity, highlighting the transformative role of LLMs and Prompt Engineering in strengthening digital security measures.

KEYWORDS: cybersecurity, Generative Artificial Intelligence, Large Language Models, Prompt Engineering, Threat Detection

1. Introduction

Prompt Engineering is an emerging discipline essential for leveraging the potential of **Large Language Models (LLMs)** in various applications. This practice involves the development and optimization of prompts – natural language inputs – designed to guide

Generative Artificial Intelligence (GAI) models in generating precise and relevant responses (outputs). Unlike traditional programming approaches that require explicit coding, Prompt Engineering offers a high level of accessibility, enabling non-technical users to interact effectively with language

models, specifically LLMs (Patel & Parmar, 2024; Chen et al., 2023).

Prompt Engineering is closely related to **Natural Language Processing (NLP)**, a subfield of Artificial Intelligence (AI) focused on the interaction between computers and human language. NLP enables GAI models to understand, interpret, and generate text in a coherent and meaningful way. It utilizes *machine learning* to allow computers to comprehend and communicate in human language. NLP integrates computational linguistics (i.e., rule-based modeling of human language) with statistical modeling, machine learning, and deep learning (IBM, 2024; Stanford NLP Group, 2025).

In this context, Prompt Engineering represents the art of crafting effective queries to optimize interactions with LLMs, which rely on advanced NLP techniques to comprehend, analyze, and generate the most precise and relevant responses to user inputs.

2. Literature Review, Research Question

Prompt Engineering involves the creation of clear and effective instructions that direct Generative AI (GAI) models to produce the desired outcomes. Unlike traditional programming, where commands are explicit, interacting with LLMs requires careful formulation of prompts (commands/instructions) to achieve optimal performance. This skill is crucial because it allows users to tailor GAI models to specific needs without requiring in-depth technical knowledge. The practice is based on the principle that language models can understand and process a broad spectrum of information when guided correctly (Reynolds & McDonell, 2021).

A deep understanding of how LLMs process and interpret prompts is crucial for achieving optimal results. The systematic and scientific use of prompts can significantly enhance the performance of Generative AI (GAI) models (Shah, 2024).

A key aspect of Prompt Engineering is its flexibility, which enables the application of language models across a wide range of fields, from education and healthcare to marketing and data analysis. This flexibility supports the access to advanced technologies, reducing technological barriers and fostering innovation (Reynolds & McDonell, 2021).

GAI is an advanced branch of Artificial Intelligence (AI) that utilizes machine learning algorithms to create new and original content, such as text, images, sounds, or videos. It is based on analysing and learning patterns from large datasets, generating outputs that mimic the characteristics of the original information without reproducing it exactly. A significant subset of GAI consists of Large Language Models (LLMs), such as GPT (Generative Pre-trained Transformer) and BERT (Bidirectional Encoder Representations from Transformers). These models are pre-trained on vast datasets and can generate coherent text and performing complex natural language processing (NLP) tasks, such as machine translation, summarization, or contextual responses. LLMs rely on deep neural network architectures and employ advanced attention mechanisms to understand and generate complex textual content. Beyond LLMs, GAI also incorporates technologies such as Generative Adversarial Networks (GANs) and Variational Autoencoders (VAEs), extending its applications to diverse domains, including photorealistic image generation, music composition, and even complex scientific simulations (Brown et al., 2020).

Large datasets play an important role in training LLMs. They are extensive and diverse collections of textual information used to train GAI models, such as LLMs. These datasets include texts from various sources, such as scientific articles, web pages, books, conversations, journals, technical documentation, and other types of textual content. In general, they cover a wide range of domains, topics, and writing styles, allowing the model to learn and understand

natural language, recognize semantic patterns, and generate relevant responses (Brown et al., 2020; Devlin et al., 2018).

Common types of datasets used in training LLMs include academic corpora (e.g., PubMed, arXiv) for scientific fields, internet-based data (e.g., Wikipedia, publicly available web pages), conversational corpora (e.g., dialogue datasets and human conversations), literary corpora (e.g., books and classic literature), and source code corpora (e.g., programming code from various platforms) (Brown et al., 2020; Devlin et al., 2018).

The rapid development of LLMs, such as **GPT-3** and **GPT-4**, has revolutionized GAI-assisted content generation in the digital age. These models, trained on massive datasets, can produce coherent, relevant, and creative texts, with applications in automated article writing, code generation, script creation, and even interactive dialogue simulations. As a result, Prompt Engineering has become essential for fully leveraging these capabilities, allowing for customization and control over the generated outputs (Sharma, 2025).

Prompts serve as the interface between users and LLMs, enabling the formulation of explicit requests for content generation while providing an intuitive means of specifying requirements. The **quality and structure of a prompt** directly influence the relevance and accuracy of the response generated by the model. Recent research has shown that techniques such as **structured prompting** can enhance model performance in linguistic prediction tasks (Kirk et al., 2022).

Thus, the structure of a prompt directly impacts the quality and relevance of the generated response. Well-designed prompts are **clear, concise, and provide adequate context**, allowing the model to generate more relevant responses. **Structured prompting**, where users break down complex requests into logical steps, can significantly **improve model performance**, particularly in **linguistic prediction and complex reasoning tasks** (Kirk et al., 2022).

In this context, the paper explores the following research question: To what extent can Prompt Engineering contribute to advancements in cybersecurity?

The rapid evolution of GAI and LLMs has transformed the way humans interact with AI-driven systems. Prompt Engineering has emerged as a crucial discipline that enhances the efficiency and precision of LLM outputs, making AI more accessible and adaptable across various domains. While existing research has extensively explored the applications of Prompt Engineering in content generation, education, and automation, its potential role in cybersecurity remains an underexplored area. Understanding how well-crafted prompts can improve threat detection, vulnerability analysis, and AI-driven security protocols is essential for harnessing the full potential of LLMs in safeguarding digital ecosystems.

In terms of originality and contribution, this paper presents a novel perspective by examining **the intersection of Prompt Engineering and cybersecurity**, an area that has received limited scholarly attention. The research aims to:

1. **Bridge the gap between Prompt Engineering and AI-powered cybersecurity** by evaluating how structured prompts can influence the effectiveness of LLMs in detecting cyber threats, analysing attack patterns, and strengthening automated security defences.

2. **Investigate the impact of structured prompting techniques** – such as breaking down security-related queries into logical steps – on improving the reliability and accuracy of AI-generated cybersecurity insights.

3. **Propose a framework for leveraging Prompt Engineering in cybersecurity applications**, outlining practical use cases where LLMs can assist security analysts, automate forensic investigations, and enhance cyber threat intelligence.

By addressing these aspects, this paper expands the current understanding of Prompt Engineering beyond its traditional applications, offering a pioneering contribution to both AI and cybersecurity research.

3. Material and Methods

This paper examines the role of Prompt Engineering in cybersecurity by applying structured prompting techniques to Large Language Models (LLMs) for threat detection, forensic analysis, and security automation. *The methodology consists of data collection, experimental design, and evaluation.*

Empirical data was sourced from publicly available cybersecurity datasets, threat intelligence reports, and synthetic network logs. The study also leveraged state-of-the-art LLMs (i.e. GPT-4o) to process structured prompts related to anomaly detection, phishing identification, and incident response.

The experimental design involved the application of Multimodal Chain-of-Thought (CoT) Prompting to evaluate LLM performance across various cybersecurity tasks. The structured prompting approach was designed to guide the GAI model in analyzing network logs to detect anomalous behaviors, identifying phishing attempts through contextual email evaluations, automating the analysis of security reports for vulnerability detection, generating cyber threat intelligence summaries from aggregated data, and conducting forensic assessments of security incidents. By integrating structured and multimodal prompts, the study aimed to enhance the accuracy and contextual awareness of GAI-generated cybersecurity insights.

The evaluation process combined quantitative and qualitative assessment methods to measure the effectiveness of Prompt Engineering in cybersecurity applications. The performance of LLMs was assessed using quantitative metrics, such as detection accuracy and response efficiency, to

determine the reliability of AI-generated outputs. Additionally, expert reviews were conducted to qualitatively analyze the clarity, coherence, and applicability of AI-generated security reports. The results indicate that structured prompting significantly improves LLM performance in cybersecurity contexts, leading to enhanced threat detection, forensic analysis, and automated reporting capabilities.

For the purposes of this paper, the scenarios and prompts were designed by the authors.

4. Results and Discussion

Prompt Engineering and LLMs play an increasingly important role in cybersecurity, contributing to threat detection, automation of data analysis processes, and the development of proactive solutions against cyberattacks. The following sections explore various ways in which Prompt Engineering can enhance the performance of cybersecurity systems, along with practical examples and relevant case studies.

4.1. Generative Artificial Intelligence-driven Cyber Threat Detection and Security Automation through Prompt Engineering

Prompt Engineering can be used to design specific queries that help LLMs *analyze network traffic and detect anomalous behaviors*. For instance, a prompt instructing the model to identify suspicious connections and explain the reasoning behind its detection can lead to the rapid identification of malicious activities, such as *Distributed Denial-of-Service (DDoS) attacks or data exfiltration* (Huang et al., 2024; Patel et al., 2025).

We propose the following prompt as an example:

“Analyze the following network traffic logs and identify any anomalous behavior that may indicate a cyberattack. Explain your reasoning step by step.”

A recent case study conducted by Palo Alto Networks demonstrated that

using LLMs to analyze network behavior can significantly reduce attack identification time, thereby improving response speed and network protection (Palo Alto Networks, 2025).

LLMs can be *trained to recognize phishing emails and messages* using Prompt Engineering techniques. By providing examples of both malicious and legitimate messages, models can be trained to detect subtle indicators of phishing attacks, such as suspicious links, urgent requests, or typographical errors (Siemerink, Jansen & Labunets, 2025).

We propose the following prompt as an example:

“Evaluate the following email and determine whether it contains indicators of a phishing attack. Explain the reasoning behind your analysis.”

Various companies have developed AI-powered phishing protection solutions that utilize LLMs to detect phishing attacks (Basit et al., 2021).

Prompt Engineering can be leveraged to *automate security report analysis generated by monitoring systems*. By designing prompts that instruct LLMs to identify critical vulnerabilities or summarize security reports, analysts can save time and quickly access relevant insights (Fieblinger et al. 2024).

We propose the following prompt as an example:

“Analyze the attached security report and identify critical vulnerabilities. Propose immediate actions for remediation.”

Microsoft Defender for Cloud integrates LLM-based AI models that utilize Prompt Engineering to analyze security reports in real-time, enhancing security monitoring processes (Freitas et al., 2024).

Prompt Engineering can be applied to *automatically generate real-time cyber threat reports based on collected data from multiple sources*. LLMs can aggregate and

organize the information, highlight the most critical threats, and propose actionable recommendations (Fieblinger et al. 2024).

We propose the following prompt as an example:

“Generate a security report summarizing identified threats from the past 24 hours and propose mitigation strategies.”

By utilizing Prompt Engineering, LLMs can *assist in analyzing cybersecurity incidents by automatically generating detailed reports* on the nature, cause, and impact of a security breach (Fieblinger et al. 2024).

We propose the following prompt as an example:

“Analyze the following security incident and identify the root causes, the impact on affected systems, and recommended actions to prevent a similar incident in the future.”

A study published in the World Journal of Advanced Research and Reviews highlights unprecedented advancements in automated threat detection and response facilitated by LLMs in cybersecurity applications (Molleti et al., 2024).

4.2. Multimodal CoT Prompting in Cybersecurity

One of the most effective Prompt Engineering techniques in cybersecurity is ***Multimodal CoT Prompting***.

Multimodal CoT Prompting is an advanced Prompt Engineering technique that extends the Chain-of-Thought (CoT) reasoning approach ***by integrating multiple types of data, such as text, images, audio, and video***. This method enables GAI models to process information from various sources and perform ***complex reasoning by combining different modalities*** (Zhang et al., 2024).

Unlike standard CoT Prompting, which relies solely on textual input, Multimodal CoT Prompting enhances ***the interpretation of visual and auditory***

information alongside textual analysis. This improves performance in tasks such as object recognition, medical diagnosis based on imaging, and multimedia data analysis (Zhang et al., 2024).

The effectiveness of this prompting technique can be attributed to the following factors:

• **Integration of Multimodal Information** – Enables GAI models to correlate visual and textual data for better contextual understanding.

• **Step-by-Step Explanation** – Encourages progressive reasoning, reducing the risk of hallucinations and improving the accuracy of responses.

• **Better Generalization** – Adaptable to multiple applications, including medicine, education, and cybersecurity.

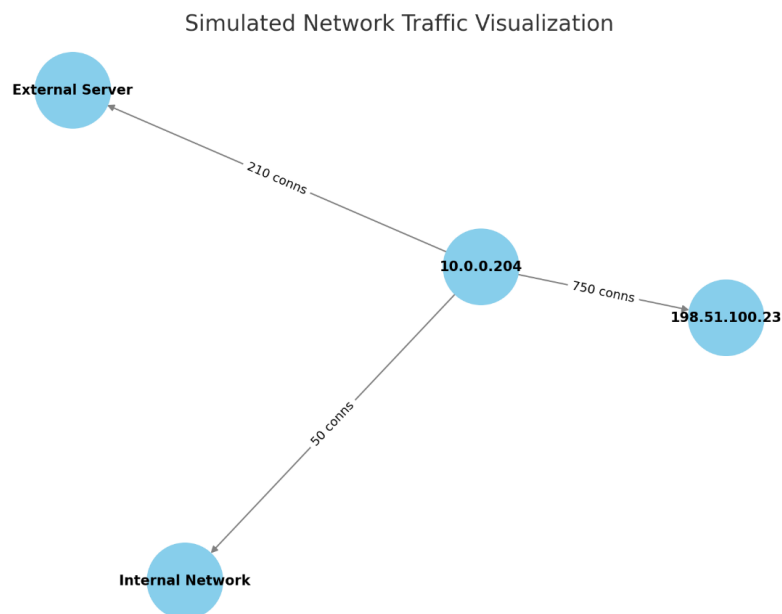
Table no. 1 presents an example of Multimodal CoT Prompting, using ChatGPT-4o. The scenario and the Multimodal CoT prompt were developed by the authors, whereas the output was generated using ChatGPT-4o.

Table no. 1

Practical example of Multimodal CoT Prompting using ChatGPT-4o: in Cybersecurity

Scenario: A cybersecurity analyst uses a GAI model to identify anomalies in an organizations’ network traffic. The available data includes both textual logs (e.g., suspicious IP addresses, unusual data flows) and graphical visualizations of the network traffic (e.g., connection graphs between nodes). Multimodal CoT Prompt: “Analyze the following network traffic data and determine whether there are any anomalous behaviors that could indicate a cyberattack. Use both the textual logs and the network traffic visualization to extract relevant insights. Explain your reasoning step by step. Input: Network Traffic Textual Logs: - 2025-03-12 14:08:21 - IP 10.0.0.204 initiated 750 TCP connections within 15 seconds to 198.51.100.23. - 2025-03-12 14:08:30 - Unusual data transfer detected: 210 GB sent to an unverified external server. - 2025-03-12 14:08:45 - IP 198.51.100.23 flagged in multiple threat intelligence databases. Network Traffic Visualization (Description): • Node 10.0.0.204 shows a sudden spike in outbound connections, primarily directed to a single external node (198.51.100.23). • Data flow volume is significantly higher than normal traffic patterns. • Anomalous activity detected on non-standard ports, deviating from typical usage trends.
--

Network Traffic Visualization:



The diagram represents the suspicious activity, highlighting the **high number of TCP connections** from **10.0.0.204** to **198.51.100.23** and the **large data transfer** to an **external server**.”

Output:

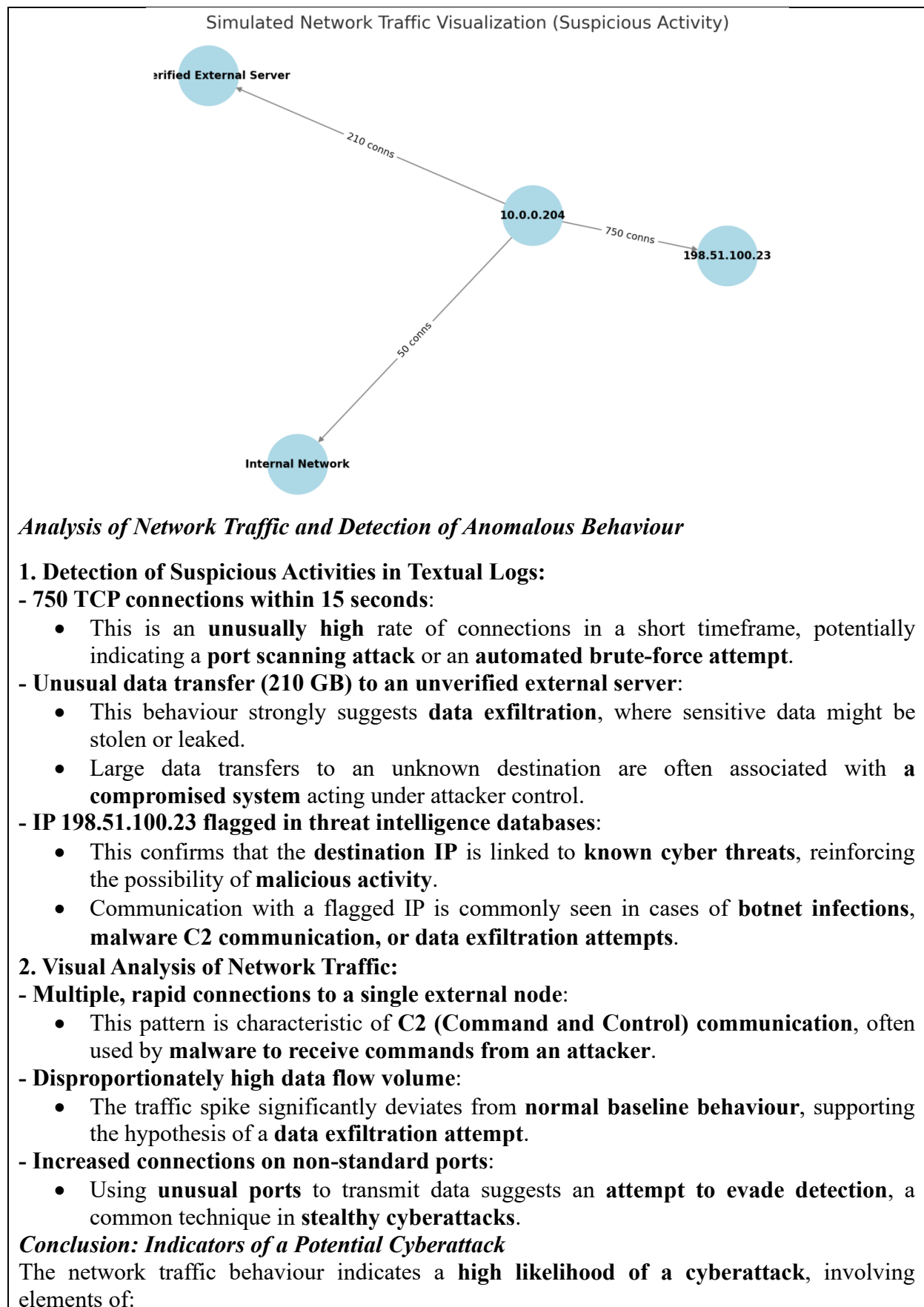
“Network Traffic Analysis Report

Network Traffic Logs

Timestamp	Event	Source IP	Destination IP
12/03/2025 14:08	750 TCP connections in 15 sec	10.0.0.204	198.51.100.23
12/03/2025 14:08	210 GB data transfer	10.0.0.204	Unknown Server
12/03/2025 14:08	IP flagged in threat database	198.51.100.23	N/A

Network Traffic Graph

The diagram represents the suspicious activity, highlighting the high number of TCP connections from 10.0.0.204 to 198.51.100.23 and the large data transfer to an external server.



- **Port scanning or brute-force attempts** to probe vulnerabilities.
- **Data exfiltration**, where large amounts of information are being transferred out of the network.
- **Communication with a malicious server**, suggesting **potential malware infection or compromised credentials**.

Recommended Mitigation Actions:

- **Immediately block IP 198.51.100.23** at the **firewall level** to prevent further malicious communication.
- **Quarantine and investigate IP 10.0.0.204** to determine if it has been compromised (e.g., malware, unauthorized access).
- **Perform a forensic analysis** on outbound traffic to **validate whether sensitive data was exfiltrated**.
- **Monitor all outbound network traffic** for additional anomalies and **implement stricter access controls**.
- **Deploy an Intrusion Detection System (IDS)** to detect and prevent similar attacks in the future.

If you require a more detailed analysis or additional security recommendations, I can provide an incident response strategy to contain and mitigate the impact of the attack.”

4.3. Proposing a Framework for Leveraging Prompt Engineering in Cybersecurity Applications

This paper introduces a structured and actionable framework for integrating Prompt Engineering into cybersecurity applications, with the goal of enhancing threat detection, security automation, and real-time incident response. The proposed framework is designed to optimize the interaction between security analysts and GAI-driven models, ensuring that LLMs and GAI systems can be effectively utilized to support and augment cybersecurity operations. This framework is built upon three core components:

✓ ***Use case-driven prompts*** – The framework advocates for the development of structured, scenario-based prompts that serve as reusable templates within GAI-driven cybersecurity workflows. These contextually tailored prompts enable security analysts to formulate precise and targeted queries, thereby improving the interpretability and accuracy of GAI-generated insights. By standardizing the use of domain-specific prompts, the framework

facilitates consistent and reproducible cybersecurity analyses, reducing the likelihood of ambiguous or misleading outputs. Furthermore, the incorporation of step-by-step logical breakdowns in prompts, such as Chain-of-Thought (CoT) prompting, enhances GAI reasoning capabilities, leading to more reliable threat assessments and forensic investigations.

✓ ***Automation of security intelligence reports*** – The paper illustrates how Prompt Engineering allows LLMs to autonomously generate threat intelligence reports by aggregating data across multiple sources. This highlights the role of Prompt Engineering in automating the generation of cybersecurity intelligence reports, a process that traditionally requires significant manual effort. By leveraging structured prompting techniques, LLMs can autonomously aggregate, process, and synthesize threat intelligence data from multiple sources, including network traffic logs, intrusion detection system alerts, and external cybersecurity databases. This automation significantly reduces the time required for threat reporting, allowing

security teams to focus on decision-making and proactive defense strategies. The structured approach ensures that LLM-generated reports remain accurate, concise, and context-aware, minimizing the risks associated with GAI hallucinations or data misinterpretation.

✓ ***Enhancing real-time security response*** – the detailed example of automated network traffic analysis shows how structured prompts enable faster incident detection and mitigation. Thus, this framework enhances real-time cybersecurity responses through structured Prompt Engineering. The integration of automated network traffic analysis – as demonstrated in the case study on anomaly detection – shows how carefully designed prompts can assist LLMs in rapidly identifying suspicious activities, correlating threat indicators, and recommending mitigation strategies. By enabling automated, GAI-assisted detection and response mechanisms, the framework contributes to faster threat containment and reduced incident response times. Additionally, the incorporation of Multimodal CoT Prompting – which combines textual logs with graph-based network visualizations – further strengthens cyber forensic capabilities, allowing for a more comprehensive assessment of complex security incidents.

This framework establishes a structured methodology for operationalizing Prompt Engineering in cybersecurity, bridging the gap between human expertise and GAI-driven automation, and laying the groundwork for next-generation GAI-powered security solutions.

4.4. Implications and Significance of Findings

The findings highlight the role of Prompt Engineering and LLMs in cybersecurity, demonstrating their capacity to enhance threat detection, security automation, and forensic analysis. Structured prompting enables LLMs to process complex security data with greater

accuracy, efficiency, and contextual awareness, bridging the gap between GAI-driven automation and human expertise.

The study confirms that GAI-assisted anomaly detection improves real-time threat identification. The ability of LLMs to analyze network traffic patterns and detect malicious activity demonstrates their potential to strengthen security infrastructures. The capability of detecting phishing further validates the role of structured prompting in mitigating social engineering threats, where conventional security tools often fall short.

The integration of Multimodal Chain-of-Thought (CoT) prompting enhances LLMs' analytical capabilities by combining text-based inputs with network traffic logs and visual threat intelligence. This approach can improve forensic analysis and automated decision-making for more reliable security insights and proactive cyber defense strategies.

The proposed framework formalizes the role of Prompt Engineering in cybersecurity workflows, ensuring that structured prompts yield interpretable and actionable outputs. The ability to automate security intelligence reporting and anomaly detection reduces manual intervention, accelerates incident response, and enhances threat mitigation strategies. These results suggest that LLM-driven security automation could be integrated into Security Operations Centers (SOCs) and Incident Response systems, strengthening overall cyber resilience.

The findings emphasize the growing role of AI in cybersecurity decision-making. Future research should explore prompt optimization, adaptive AI learning mechanisms, and advanced multimodal threat detection, ensuring that Prompt Engineering continues to evolve as a critical tool in GAI-driven security frameworks.

5. Conclusions

The research highlights the necessity of precisely designed prompts to guide GAI models in interpreting security data, reducing errors, and ensuring coherent threat analysis. The proposed framework provides a structured approach for integrating GAI-assisted cybersecurity automation, reducing reliance on manual analysis while maintaining operational oversight.

Despite these advancements, the adoption of Prompt Engineering in cybersecurity can also present challenges and future research opportunities. The reliability of LLM-generated security insights depends on the quality of input prompts, the robustness of training data,

and the ability to mitigate biases or adversarial manipulation. Future research should explore adaptive and self-learning prompt optimization techniques, allowing GAI models to refine and improve their security responses over time. Additionally, integrating real-world cybersecurity datasets for continuous GAI training and validation will be critical in enhancing the robustness and trustworthiness of AI-powered security solutions. As cyber threats continue to evolve, Prompt Engineering will play a crucial role in shaping next-generation AI-driven cybersecurity frameworks, ensuring greater resilience, automation, and proactive defense against sophisticated cyberattacks.

REFERENCES

- Basit, A., et al. (2021). A comprehensive survey of AI-enabled phishing attacks detection techniques. *Telecommunication Systems*. 10.1007/s11235-020-00733-2.
- Brown, T.B., et al. (2020). *Language Models are Few-Shot Learners*. arXiv:2005.14165 [cs.CL]. Available at: <https://doi.org/10.48550/arXiv.2005.14165>.
- Chen, B., et al. (2023). *Unleashing the potential of prompt engineering in Large Language Models: a comprehensive review*. arXiv:2310.14735 [cs.CL]. Available at: <https://doi.org/10.48550/arXiv.2310.14735>.
- Devlin, J., et al. (2018). BERT: *Pre-training of Deep Bidirectional Transformers for Language Understanding*. arXiv:1810.04805 [cs.CL]. Available at: <https://doi.org/10.48550/arXiv.1810.04805>.
- Fieblinger, R., et al. (2024). *Actionable Cyber Threat Intelligence using Knowledge Graphs and Large Language Models*. arXiv:2407.02528 [cs.CR]. Available at: <https://doi.org/10.48550/arXiv.2407.02528>
- Freitas, S., et al. (2024). *AI-Driven Guided Response for Security Operation Centers with Microsoft Copilot for Security*. arXiv:2407.09017 [cs.LG]. Available at: <https://doi.org/10.48550/arXiv.2407.09017>.
- Huang, Q., et al. (2024). *Selective Prompting Tuning for Personalized Conversations with LLMs*. arXiv:2406.18187 [cs.CL]. Available at: <https://doi.org/10.48550/arXiv.2406.18187>
- IBM (2024). *What is NLP?*. Available at: <https://www.ibm.com/think/topics/ai-agents>, accessed on January 2025.
- Kirk, J.R., et al. (2022). *Improving Language Model Prompting in Support of Semi-autonomous Task Learning*. arXiv:2209.07636 [cs.LG]. Available at: <https://doi.org/10.48550/arXiv.2209.07636>.
- Molleti, R., et al. (2024). Automated threat detection and response using LLM agents. *World Journal of Advanced Research and Reviews*, 2024, 24(02), 079-090. Available at: <https://doi.org/10.30574/wjarr.2024.24.2.3329>.

Palo Alto Networks. (2025). *Cloud Security What Are Large Language Models (LLMs)?* Available at: <https://www.paloaltonetworks.com/cyberpedia/large-language-models-llm>, accessed on February 2025.

Patel, H., & Parmar, S. (2024). *Prompt Engineering for Large Language Model*. DOI:10.13140/RG.2.2.11549.93923.

Patel, V., et al. (2025). Unveiling network intrusions: Leveraging large language models for anomaly detection in cybersecurity. *AIP Conf. Proc.*, Vol. 3255, Nr. 1, 020011 (2025). Available at: <https://doi.org/10.1063/5.0254175>.

Reynolds, L., & McDonell, K. (2021). *Prompt Programming for Large Language Models: Beyond the Few-Shot Paradigm*. arXiv:2102.07350 [cs.CL]. Available at: <https://doi.org/10.48550/arXiv.2102.07350>.

Shah, C. (2024). *From Prompt Engineering to Prompt Science with Human in the Loop*. arXiv:2401.04122 [cs.HC]. Available at: <https://doi.org/10.48550/arXiv.2401.04122>.

Sharma, M. (2025). *What is Prompt Engineering? Almost everything you want to know about one of the hottest jobs in AI*. Techradar. Available at: <https://www.techradar.com/computing/artificial-intelligence/what-is-prompt-engineering-almost-everything-you-want-to-know-about-one-of-the-hottest-job-in-ai>, accessed on February 2025.

Siemerink, A., Jansen, S., & Labunets, K. (2025). *The Dual-Edged Sword of Large Language Models in Phishing*. In: Horn Iwaya, L., Kamm, L., Martucci, L., Pulls, T. (eds) *Secure IT Systems. NordSec 2024. Lecture Notes in Computer Science, Vol 15396*. Springer, Cham. Available at: https://doi.org/10.1007/978-3-031-79007-2_14.

Stanford NLP Group. (n.d.). Available at: <https://nlp.stanford.edu/>, accessed on January 2025.

Zhang, Z., et al. (2024). *Multimodal Chain-of-Thought Reasoning in Language Models*. arXiv:2301.07069 [cs.CL]. Available at: <https://doi.org/10.48550/arXiv.2301.07069>.