

A Comparative Analysis of LLMs in Mapping Malware Behaviors to MITRE ATT&CK Techniques from Textual Threat Intelligence Reports

Thursday 18 September 2025 15:00 (15 minutes)

Cyber Threat Intelligence (CTI) Reports are valuable sources of information for understanding adversarial behaviors and malware functionalities. But their lack of consistency and structure can represent a challenge for security analysts in interpreting, correlating, and applying them effectively. Structuring the data in a common format, such as the MITRE ATT&CK framework, is crucial for integrating CTI into detection and response processes.

This article analyzes the extent to which Large Language Models (LLMs) - GPT (OpenAI), Claude (Anthropic), and Gemini (Google) - can extract and map malware descriptions from natural language CTI reports to specific MITRE ATT&CK techniques. To achieve this, a set of publicly available CTI reports was used that already contained verified MITRE ATT&CK techniques labels. This served as ground truth for evaluating the outputs of each model.

Although issues were observed in the model's execution, such as technique confusion and context loss, the results suggest a strong potential in the use of LLMs for mapping threat intelligence. Their ability to reduce manual effort and improve consistency could address a major gap in today's cyber threat analysis workflow.

Author: RESUL, Ebru

Co-authors: TURCANU, Dinu (Technical University of Moldova); Prof. RUGHINIȘ, Răzvan Victor (National University of Science and Technology POLITEHNICA Bucharest)

Presenter: RESUL, Ebru

Session Classification: Security & Resilience in Cyber-Physical Systems

Track Classification: Doctoral Symposium