

LLM을 이용한 바이러스 코드의 클론 생성 기법

장수혁, 강성혁, 원일용, 김현정*, 유상현**

서울호서전문학교, *건국대학교, **경민대학교

suv240113@naver.com, detole@naver.com, clccclcc@shoseo.ac.kr, *nygirl@konkuk.ac.kr,
**simonyoo@kyungmin.ac.kr

Clone Generation Techniques for Virus Code Using LLM

Jang Su Hyeok, Kang Seong Huck, Won Il Yong, Kim Hyun Jung*, Yoo Sang Hyun**

Seoul Hoseo College, *Konkuk Univ., Kyungmin Univ.**

요약

본 논문은 대규모 언어 모델(LLM)을 활용하여 악성 소프트웨어의 변종을 생성하고, 백신 탐지 시스템을 회피하는 기법을 제안한다. 실험 결과, 생성된 변종은 기능적 유사성을 유지하면서도 구조적으로 변형되어 탐지 가능성을 크게 감소시켰다.

I. 서론

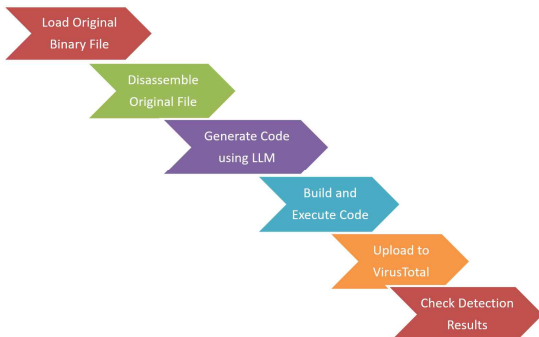
디지털 기술의 비약적 발전으로 사이버 공간의 중요성이 날로 증대되고 있으며, 이에 따라 사이버 공격과 방어는 국가 안보의 핵심 요소로 자리 잡고 있다[1][2][5][6]. 이러한 국가 안보의 핵심 요소를 위협하는 새로운 방법인 대규모 언어 모델(LLM)을 활용한 바이러스 코드 생성 및 분석은 사이버 공격의 진화를 가속화하고 있으며, 이에 대한 효과적인 방어 기술의 필요성이 더욱 부각되고 있다[3][4][7][8].

본 논문에서는 LLM을 활용하여 MS-DOS 환경에서 다양한 변종 바이러스 코드를 어셈블리 코드 변형을 통해 생성하고, 그 빌드 성공 여부를 분석함으로써 악성 코드 생성 기술의 최신 동향과 메커니즘을 심층적으로 탐구한다. 이를 통해 사이버 보안 분야에서 LLM의 잠재적 활용 가능성을 탐색하고, 악성 코드의 진화에 대응하기 위한 방어 전략 수립에 필요한 기초 자료를 제공하고자 한다.

II. 실험 데이터 및 결과

1. 실험 데이터

본 연구에서는 바이러스 코드 생성 알고리즘을 통해 실험 데이터를 생성하였으며, 이 과정을 시각적으로 설명하기 위해 단계별 프로세스를 나타냈다.

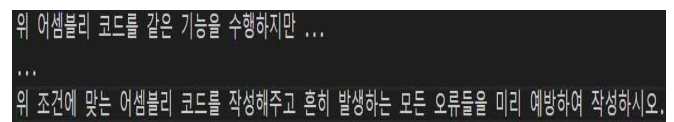


(그림 1) Virus Code Generation Process

바이러스 코드 생성 알고리즘은 원본 파일을 분석하여 변환한 후, LLM을 활용해 변종을 생성하는 방식으로 이루어진다. 생성된 변종은 실행 후 분석 도구에 의해 평가되며, 이를 통해 LLM이 생성한 변종 코드의 탐지 회피 성능을 확인할 수 있다.

2. 생성된 바이러스 코드

본 연구에서는 원본 바이러스 파일을 분석하여 저수준 코드로 변환한 후, 이를 기반으로 새로운 변종 코드를 생성하였다. 생성된 변종 코드는 실행 과정에서 특정 동작을 수행하며, 이후 화면 출력 등과 같은 결과를 나타낸다.



(그림 2) Assembly Code Generation Program

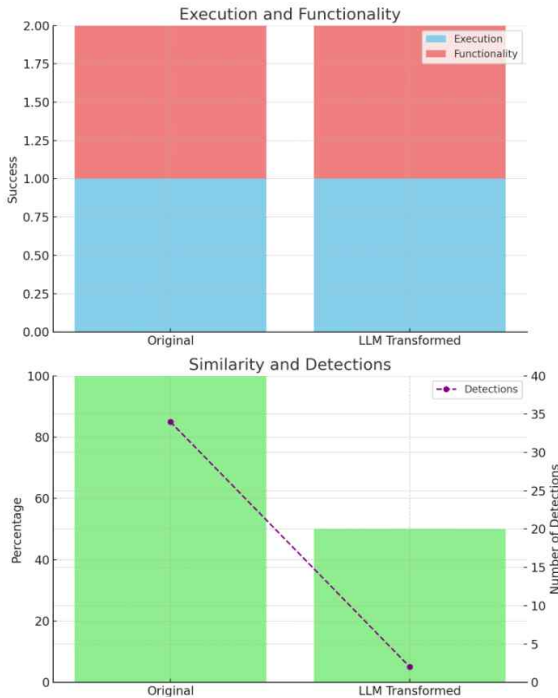


(그림 3) Virus Execution Result

3. 실험 결과

실험 결과, LLM을 사용하여 변환된 바이러스 코드는 원본과 비교하여 기능적 동등성을 유지하면서도 구조적 유사도가 크게 감소하였다. 이러한 변환으로 인해 바이러스 탐지율에서 현저한 차이가 나타났다. 원본 코드는 다수의 백신에 의해 탐지되었으나, 변종 코드는 극히 일부 백신에만 탐지되어 탐지 회피 성능이 약 94% 향상되었다.

(그림 4)는 실험 결과를 시각적으로 나타낸 것으로, 첫 번째 차트는 바이러스 코드의 실행 및 기능성을, 두 번째 차트는 구조적 유사도와 탐지 횟수를 보여준다. 이를 통해, LLM을 통한 코드 변환이 탐지 회피에 효과적임을 확인할 수 있다.



(그림 4) Virus Detection Rate Experiment Result

III. 결론 및 향후 과제

본 연구에서는 대규모 언어 모델(LLM)을 활용하여 MS-DOS 환경에서 바이러스 코드의 클론을 생성하고, 이를 통해 백신 탐지 회피 기법을 제시하였다. 실험 결과, LLM을 통해 생성된 클론 코드가 기능적 동등성을 유지하면서도 구조적 유사성을 감소시켜 백신 탐지율을 현저히 낮출 수 있음을 확인하였다. 이를 기반으로 LLM을 활용한 코드 변환이 바이러스 코드의 탐지 회피에 효과적임을 보여주었다.

향후 연구에서는 Windows 및 Linux와 같은 주요 운영체제에서도 유사한 실험을 진행하여, 더욱 강력하고 실용적인 백신 회피 기술을 개발하는데 주력할 필요가 있다. 이러한 연구는 사이버 보안에서 새로운 위협 모델을 이해하고, 방어 전략을 개선하는 데 중요한 기여를 할 것으로 기대된다.

참 고 문 헌

- [1] 권혁천, 이용준, 박원형, "한국의 사이버공격 비교 분석과 정책적 대응 방안," 융합보안논문지, v.20, no.5, pp.19-26, 2020.
- [2] 박휘락, "북한의 사이버전 위협분석과 대응방안 고찰," 정보보호학회 논문지, v.30, no.5, pp.118-128, 2015.
- [3] Lonergan, Erica D., and Jacquelyn Schneider, "Artificial Intelligence for Cybersecurity: Literature Review and Future Research Directions," Journal of Cybersecurity, vol. 9, no. 1, pp. 1-30, 2023.
- [4] Mohamed Amine Ferrag, et al., "Generative AI and Large Language Models for Cyber Security: All Insights You Need,"

arXiv:2405.12750 [cs.CR], vol. 50, pp. 1-50, 2024.

[5] "AI in Cybersecurity: The Role of Large Language Models," cybersecurity.dev, pp.1-8, 2023.

[6] 윤오준, 배광용, 김재홍, 서형준, 신용태, "사이버공격 대응 분석을 통한 사이버안보 강화 방안 연구," 융합보안논문지, v.15, no.4, pp.65-72, 2015.

[7] Z. Wang et al., "Revolutionizing Cyber Threat Detection with Large Language Models: A privacy-preserving BERT-based Lightweight Model for IoT/IIoT Devices," arXiv preprint arXiv:2306.14263, 2023.

[8] Tom Gleason, "LLM-assisted Malware Review: AI and Humans Join Forces to Combat Malware," Endor Labs, pp.1-10, 2023.