

ГЕНЕРАТИВНІ МОВНІ МОДЕЛІ В АНАЛІЗІ ШКІДЛИВОГО КОДУ

GENERATIVE LANGUAGE MODELS IN MALWARE ANALYSIS

Стрімке розширення цифрових сервісів та зростання кількості користувачів в онлайн-просторі призводять до появи дедалі складніших кіберзагроз, які безпосередньо впливають на безпеку як організацій, так і окремих людей. Генеративні мовні моделі (LLM) відіграють дедалі важливішу роль у сфері кібербезпеки, зокрема в аналізі шкідливого програмного забезпечення. Зростання масштабів кіберзагроз та ускладнення технік обходу традиційних систем захисту призвели до того, що malware став одним із ключових векторів атак у 2023–2025 рр. За даними звіту CrowdStrike та Mandiant кількість нових зразків шкідливого ПЗ щорічно збільшується на 30–40%, а значна частина з них створюється або модифікується із використанням автоматизованих інструментів і генеративних алгоритмів, що суттєво ускладнює їхню ідентифікацію класичними методами.

Сучасний malware часто приховує свій справжній функціонал за допомогою обфускації, шифрування або спеціальних антианалізних механізмів. Це знижує ефективність традиційних інструментів і створює потребу в інтелектуальних підходах, здатних інтерпретувати логіку коду на рівні його намірів, а не лише поверхневої структури. Генеративні моделі, побудовані на архітектурі Transformer, здатні аналізувати програмний код подібно до експерта: виявляти приховані зв'язки між функціями, визначати контекст використання API та оцінювати потенційно небезпечну поведінку.

Особливу практичну цінність становить здатність LLM пояснювати результати аналізу природною мовою. Це зменшує навантаження на аналітиків SOC-центрів і допомагає швидше ухвалювати рішення у разі інцидентів. Додаткове навчання моделей на наборах даних та інтеграція з доменною базою знань покращують точність виявлення загроз і дозволяють адаптувати систему до нових шкідливих сімейств.

Разом із тим впровадження LLM у процес аналізу шкідливого коду супроводжується певними викликами. Зокрема, моделі можуть генерувати хибні висновки або небажаний код, якщо не дотримано політик безпеки. Не менш важливими є питання обробки конфіденційних даних, вразливості до некоректних промптів та потреба у значних обчислювальних ресурсах.

Попри це, експериментальні дослідження демонструють суттєве зростання ефективності систем кіберзахисту у разі поєднання класичних підходів і можливостей LLM. Інтелектуальний аналіз дозволяє зменшити кількість хибнопозитивних спрацювань, скоротити час обробки інцидентів та підвищити якість прогнозування ризиків. Використання генеративних моделей формує основу для нового покоління адаптивних систем безпеки, здатних оперативно реагувати на швидко мінливі загрози сучасного кіберпростору.

Література

1. CrowdStrike. (2024). 2024 Global Threat Report. CrowdStrike Holdings, Inc. <https://www.crowdstrike.com/global-threat-report/> (date of access: 03.12.2025).
2. Mandiant. (2024). M-Trends 2024 Special Report. Mandiant, Google Cloud. <https://www.mandiant.com/resources/m-trends> (date of access: 03.12.2025).