

Чернігівський Іван Андрійович

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID 0009-0003-4568-3212

Крючкова Лариса Петрівна

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID 0000-0002-8509-6659

ІНФОРМАЦІЙНІ ВПЛИВИ НА ІНФОКОМУНІКАЦІЙНІ МЕРЕЖІ ІЗ ЗАЛУЧЕННЯМ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. Стрімкий розвиток інформаційних технологій спричинив як позитивні зміни сучасного суспільства, так і нові загрози для інформаційної безпеки. Шкідливі інформаційні впливи на інфокомунікаційні мережі стають усе більш витонченими й небезпечними, оскільки зловмисники постійно вдосконалюють методи впливу, зокрема і шляхом залучення штучного інтелекту (ШІ). Сучасні можливості технологій ШІ, зокрема великих мовних моделей (LLM), таких як GPT-4 і BERT, та генеративних змагальних мереж (GAN), дозволяють вирішувати складні завдання обробки природної мови (NLP). Змінюється сам кіберпростір, за рахунок застосування технологій ШІ як в оборонних, так і наступальних операціях. Актуальність роботи полягає у необхідності виявлення напрямків залучення штучного інтелекту до формування шкідливих інформаційних впливів на інфокомунікаційні мережі та вирішення задач захисту. Метою дослідження є аналіз сучасних інформаційних впливів на інфокомунікаційні мережі із залученням штучного інтелекту для виявлення вразливостей, прийняття обґрунтованих рішень щодо забезпечення інформаційної безпеки та вдосконалення методів захисту. В межах даного дослідження розглянуто можливість залучення ШІ як для атак на інфокомунікаційну мережу, так і для захисту інфокомунікаційних мереж. В результаті виконаного аналізу встановлено, що не всі інформаційні впливи із залученням ШІ несуть однакову шкоду. Серед потенційних загроз та сценаріїв інформаційних впливів із залученням ШІ, які можуть негативно вплинути на інформаційну безпеку інфокомунікаційної мережі, найбільш небезпечними на сьогодні варто вважати ті, які спроможуть фішинг користувачів, генерують людиноподібний текст, правдоподібні голосові повідомлення та відео (deepfake), для більш переконливого введення в оману користувачів. Перспективним напрямом є активне залучення ШІ до захисту інфокомунікаційних мереж для протистояння кібератакам, особливо тим, які виконуються із залученням ШІ, забезпечуючи виявлення, аналіз та реагування на кіберзагрози в реальному часі. Подальші дослідження доцільно зосередити на тестуванні наявних нейромережових моделей для оцінки можливості їх застосування при вирішенні задач захисту інфокомунікаційних мереж.

Ключові слова: інформаційні впливи; кібербезпека; LLM; штучний інтелект; моделі; кіберзагрози; кібератаки; фішинг; інфокомунікаційні мережі

Chernihivskyi Ivan Andriiovych

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID 0009-0003-4568-3212

Kriuchkova Larysa Petrivna

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID 0000-0000-0000-0000

INFORMATION INFLUENCES ON INFOCOMMUNICATION NETWORKS WITH THE INVOLVEMENT OF ARTIFICIAL INTELLIGENCE

Abstract. The rapid development of information technologies has caused both positive changes in modern society and new threats to information security. Harmful information impacts on infocommunication networks are becoming increasingly sophisticated and dangerous, as attackers are constantly improving methods of influence, including by involving artificial intelligence (AI). Modern capabilities of AI technologies, in particular large language models (LLM), such as GPT-4 and BERT, and generative adversarial networks (GAN), allow solving complex problems of natural language processing (NLP). Cyberspace itself is changing, due to the use of AI technologies in both defensive and offensive operations. The relevance of the work lies in the need to identify areas of involving artificial intelligence in the formation of harmful information impacts on infocommunication networks and solving protection problems. The purpose of the study is to analyze modern information impacts on infocommunication networks using artificial intelligence to

identify vulnerabilities, make informed decisions on ensuring information security, and improve protection methods. This study considers the possibility of using AI both for attacks on the infocommunication network and for protecting infocommunication networks. As a result of the analysis, it was found that not all information impacts involving AI are equally harmful. Among the potential threats and scenarios of information impacts involving AI that can negatively affect the information security of the infocommunication network, the most dangerous ones today are those that simplify phishing of users, generate human-like text, plausible voice messages, and video (deepfake) for more convincing deception of users. A promising direction is the active involvement of AI in the protection of infocommunication networks to counter cyberattacks, especially those carried out with the involvement of AI, ensuring the detection, analysis and response to cyberthreats in real time. Further research should focus on testing existing neural network models to assess the possibility of their application in solving problems of protecting infocommunication networks.

Keywords: *information impacts; cybersecurity; LLM; artificial intelligence; models; cyber threats; cyber attacks; phishing; infocommunication networks*

1. Постановка проблеми. Стрімкий розвиток інформаційних технологій спричинив як позитивні зміни сучасного суспільства, так і нові загрози для інформаційної безпеки. Кібератаки на інфокомунікаційні мережі стають усе більш витонченими й небезпечними, оскільки зловмисники постійно вдосконалюють методи впливу, зокрема і шляхом залучення штучного інтелекту (ШІ).

Інформаційний конфлікт є процесом боротьби за інформацію, в рамках якої сторони реалізують найрізноманітніші заходи щодо її «руйнування», «спотворення», «приховування» і «вилучення» при конфліктній взаємодії систем, здійснюваній за допомогою передавання та приймання даних у інфокомунікаційних системах. Взаємний вплив, що призводить до зміни інформаційних станів об'єктів, часто здійснюється безконтактним способом і в залежності від атаки, може навіть іноді не вимагати від користувача ніяких дій.

В реаліях російсько-української війни в умовах високої залежності військового обладнання від безпроводових інфокомунікаційних систем, питання їхньої надійності та безпеки постає на перше місце. Конфліктна взаємодія в інформаційному просторі становить значну загрозу для стабільного функціонування цих систем.

ШІ є важливим інструментом для аналізу та прогнозування наслідків конфліктної взаємодії як для захисту власних, так і для знешкодження ворожих систем. В умовах кібервійни найбільш ймовірними видаються впливи, спрямовані на перехоплення та спотворення даних, перехоплення управління та зміни функціонування, і навіть знищення інфокомунікаційної мережі. Зростаюча складність та глобальність кіберпростору створює серйозні виклики для досягнення необхідної кіберстійкості, особливо для організацій з невеликими ресурсами.

Війна спонукає організації переглядати свої стратегії, балансуючи проблеми безпеки з глобальними операціями, оскільки у воєнні часи значно частіше виконуються цілеспрямовані атаки від АРТ (Advanced Persistent Threat) на вразливі державні та недержавні об'єкти для шпигунства, саботажу і нанесення максимальної шкоди. Цей динамічний ландшафт вимагає адаптивних стратегій, які враховують зміну глобальних ризиків та залежність від ланцюга поставок. Водночас, зростаюча складність кіберзлочинів залишається постійною проблемою. Загальновідомі тактики, вдосконалені штучним інтелектом, програми-вимагачі як послуга (RaaS) та передові методи соціальної інженерії дозволяють кіберзлочинцям випереджати традиційні засоби захисту. Боротьба з цими загрозами, що стрімко розвиваються, вимагає не лише передових технологічних рішень, але й співпраці та обміну знаннями. Організації, які застосовують проактивне управління ризиками, надають пріоритет спільним підходам в різних екосистемах, свідомо використовують ШІ у своїй роботі та інвестують у масштабовані захисні рішення, можуть суттєво зменшити вплив кіберзагроз [1-3].

2. Аналіз останніх досліджень і публікацій.

Кіберзлочинці використовують GenAI для переконливого відтворення стилів спілкування керівників вищої ланки організації. Ці інструменти використовують контекстуальні дані з таких джерел, як соціальні мережі, публічні заяви чи витік документів, що робить спроби соціальної інженерії набагато складнішими для ідентифікації. GenAI також допомагає зловмисникам розробляти надійні атаки соціальної інженерії ширшим колом мов,

що допомагає зловмисникам атакувати більшу кількість людей у більшій кількості країн за меншими витратами [4].

Близько 72% респондентів повідомляють про зростання кіберризиків в організаціях, причому програми-вимагачі залишаються головною проблемою. Майже 47% організацій називають зловмисні дії, що базуються на генеративному штучному інтелекті (GenAI), що в свою чергу дозволяє здійснювати більш складні та масштабовані атаки – своєю основною проблемою. У 2024 році спостерігалось різке зростання кількості фішингових атак та атак соціальної інженерії, причому 42% організацій повідомили про такі інциденти [4].

У працях науковців Hakan T. Otal, M. Abdullah Canbaz [5], Yazi Gholami [6], Luigi Coppolino a, Salvatore D'Antonio a, Giovanni Mazzeo a, Federica Uccello [7] досліджуються великі мовні моделі (LLM – Large Language Model) та їх можливе застосування в інформаційній безпеці, оскільки LLM продовжують розвиватися, вони відіграватимуть дедалі важливішу роль у захисті цифрових систем від нових загроз.

Актуальність роботи полягає у необхідності виявлення напрямків залучення штучного інтелекту до формування шкідливих інформаційних впливів на інфокомунікаційні мережі та вирішення задач захисту.

3. Мета і задачі дослідження

Метою дослідження є аналіз сучасних інформаційних впливів на інфокомунікаційні мережі із залученням ШІ для виявлення вразливостей, прийняття обґрунтованих рішень щодо забезпечення інформаційної безпеки та вдосконалення методів захисту.

Задачі дослідження:

1. Розуміння потенційних загроз та сценаріїв інформаційних впливів із залученням ШІ, які можуть негативно вплинути на інформаційну безпеку інфокомунікаційної мережі.
2. Використання результатів аналізу для прийняття обґрунтованих рішень щодо забезпечення інформаційної безпеки, управління ризиками та вдосконалення методів захисту.

4. Результати дослідження

Цифрова революція останніх десятиліть відкрила епоху безпрецедентної залежності від взаємозалежних інформаційних систем. Безперервне функціонування сучасного суспільства, його критично важливої інфраструктури, від енергомереж та фінансових установ до мереж охорони здоров'я та каналів зв'язку залежить від інформаційних технологій. Однак ця залежність також створила сприятливе підґрунтя для кібератак, що створюють значні ризики для національної, політичної, економічної та суспільної безпеки.

Кібербезпека це динамічна і складна область, що характеризується постійною еволюцією та ескалацією кіберзагроз, швидким розвитком і впровадженням нових технологій та передбачає захист інформаційних систем та мереж від несанкціонованого доступу, використання, зміни чи знищення. Вона необхідна для забезпечення конфіденційності, цілісності та доступності даних та послуг, а також вимагає міждисциплінарного та комплексного підходу, що включає технічні, людські, організаційні та суспільні фактори, і обов'язково враховує технічні, етичні, правові та соціальні наслідки кібердіяльності.

За минулі роки кіберзлочинці перетворилися на високоорганізовані та складні структури, що нагадують традиційні злочинні угруповання з чіткою ієрархією та спеціалізованими ролями. Вони активно співпрацюють та використовують передові методи, такі як експлойти нульового дня, складне шкідливе програмне забезпечення та методи соціальної інженерії, такі як фішинг. Спонсоровані державою суб'єкти та кіберзлочинні угруповання використовують новітні технології для автоматизації та посилення своїх атак, експлуатуючи вразливості організацій-жертв. Фінансові мотиви спонукають їх використовувати криптовалюти для анонімних транзакцій, що сприяє зростанню числа програм-вимагачів та кібершахрайства. Їх глобальне охоплення та здатність атакувати критично важливу інфраструктуру створюють серйозні проблеми, вимагаючи постійного вдосконалення заходів кібербезпеки та міжнародного співробітництва для ефективної

протидії цим загрозам. Традиційно заходи кібербезпеки часто мають реактивний характер і обмежені у своїх можливостях справлятися з ландшафтом загроз, що швидко змінюється.

Варто зазначити, що стрімкий розвиток технологій штучного інтелекту (ШІ), зокрема великих мовних моделей (LLM), таких як GPT-4 і BERT, та генеративних змагальних мереж (GAN), дозволили вирішувати складні завдання обробки природної мови (NLP). Змінюється сам кіберпростір, за рахунок застосування технологій ШІ як в оборонних, так і наступальних операціях. У кібербезпеці впровадження LLM відкриває нові перспективні можливості для вирішення завдань, пов'язаних із зростаючою складністю та масштабом кіберзагроз.

Розвиток засобів та способів шкідливого інформаційного впливу на інфокомунікаційні системи призвів до відповідного розвитку засобів захисту, покликаних виявити, ідентифікувати та нейтралізувати небажані інформаційні впливи на інфокомунікаційну систему. Відомі способи інформаційних впливів на інфокомунікаційну систему із залученням ШІ доцільно поділити на такі групи: *ШІ для атак на інфокомунікаційну мережу*, *ШІ для захисту інфокомунікаційної мережі*.

Залучення ШІ для атак на інфокомунікаційну мережу

На світовому економічному форумі було розглянуто тренди у кібербезпеці на 2025 рік, що відображено на рис. 1,2,3.

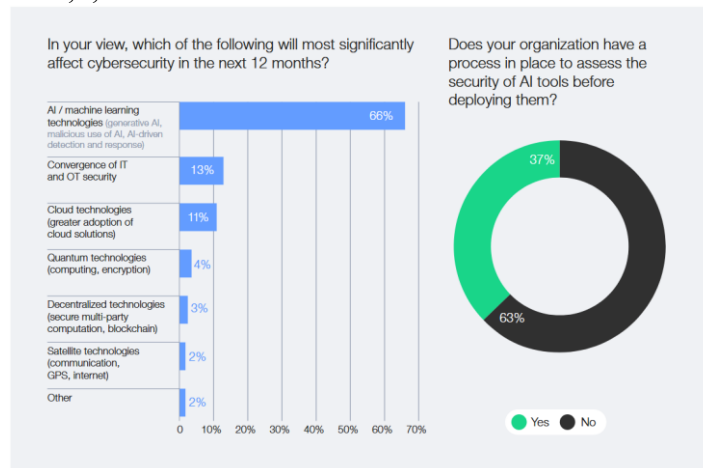


Рис.1. Вразливості що очікувались у 2025 році [4].

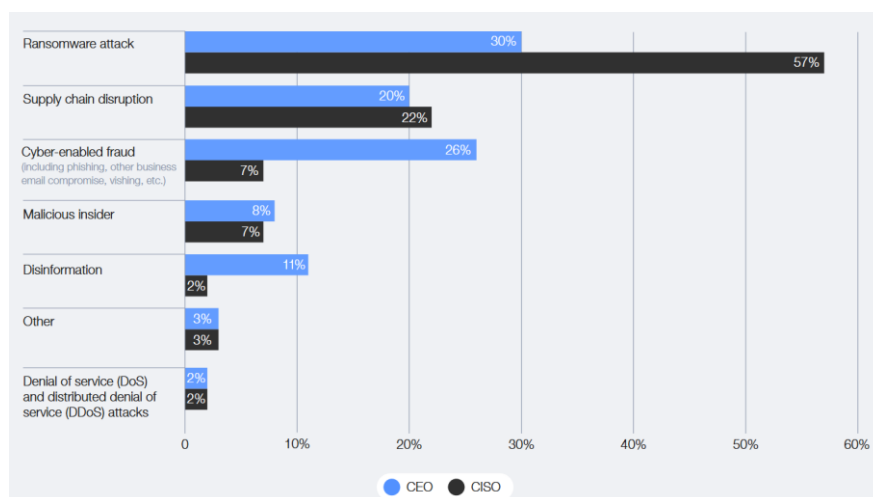


Рис.2. Очікувані кіберризики у 2025 році [4].

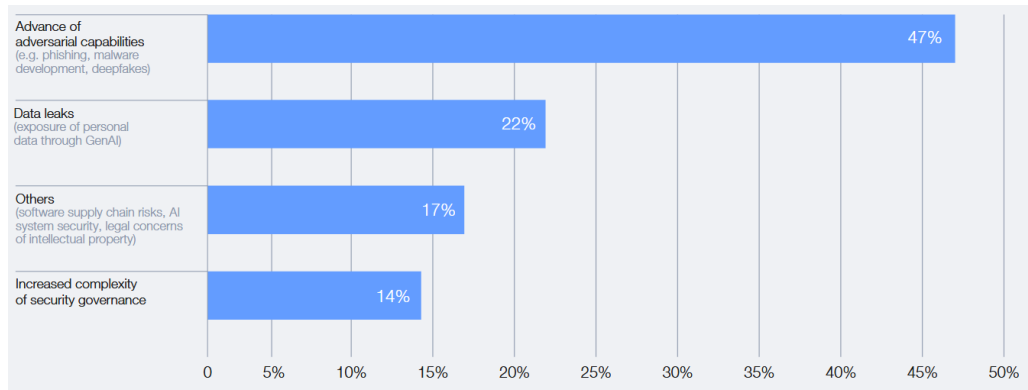


Рис.3. Загрози для кібербезпеки, пов'язані з GenAI [4].

Цей загальносвітовий тренд на використання ШІ у різних сферах у тому числі і для атаки на комп'ютерну інфраструктуру (weaponized AI) також корелює і з українськими даними про інциденти, оскільки однією з головних можливостей проникнути в мережу – використовувати фішингові повідомлення для викрадення облікових даних або для запуску шкідливого коду, а ШІ вже може генерувати людиноподібний текст, правдоподібні голосові повідомлення та відео (deepfake), для більшої переконливості [8].

Наприклад, CERT-UA разом з SSSCIP знайшли 4315 кіберінцидентів у 2024, що близько на 70% більше ніж у попередньому році [9].

За даними CERT-UA, від початку 2025 року основні типи кіберзлочинної активності на українську інфраструктуру включають в себе [10]:

- Кібершпигунство, переважно у державній та військовій сферах;
- Саботаж або Кібертероризм;
- Фінансово вмотивовані злочини, спрямовані здебільшого на викрадення коштів або даних;
- Інші специфічні атаки.

Основні кіберзлочинні угруповання, що втілюють вищевказані атаки [10]:

- UAC-0006;
- UAC-0010 (Gamaredon, Armageddon, Primitive Bear etc.) [11];
- UAC-0050;
- UAC-0184;
- UAC-0200;
- UAC-0218;
- UAC-0219.

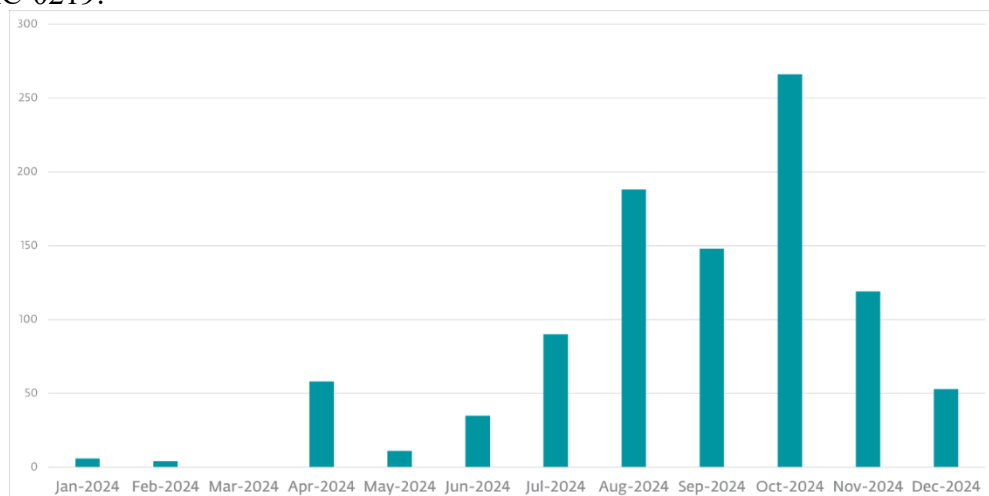


Рис.4. Інтенсивність фішингу Gamaredon проти України за даними компанії ESET [12].

В [7] розглядаються основні методи та прийоми, які застосовують зломисники за допомогою генеративного ШІ для розробки складних кіберзагроз та обходу традиційних механізмів захисту. Створюється нова стратегія приховування атак із залученням ШІ: усувається критичний пропуск у методологіях приховування атак, розробляючи та оцінюючи підхід на основі умовної генеративно-мережевої мережі (CGAN). Метод дозволяє обійти систему виявлення вторгнень (IDS) на основі нейронної мережі (NN), забезпечуючи успішне проведення низки веб та мережевих атак без зміни їх фактичних сигнатур. Зокрема, серед успішно прихованих атак були: атака типу «Людина посередині» (MITM), три варіанти Mirai та SQL-ін'єкція. Автори також перевіряють цей підхід, використовуючи набори даних про атаки нульового дня, щоб оцінити його ефективність проти нових шаблонів атак [7].

Щодо безпеки паролів, у статті [13,14], автори натренували модель GPT-2 на відомих витоках даних та словниках паролів (таких як Rockyou), і це вже перевершує наявні алгоритми для підбору паролів, що продемонстровано на рис 5,6,7. Останній навчається після того, як перший зійдеться і є необхідним для генерації. Декодери трансформаторів в обох моделях є незалежними [10].

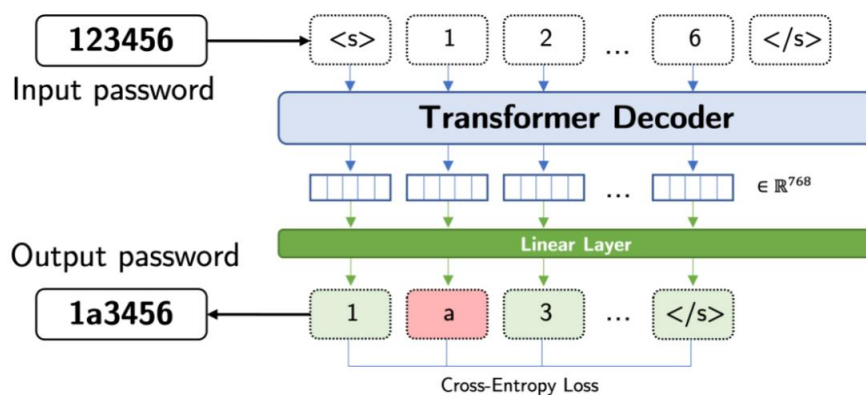


Рис.5. Прогнозування моделлю PassGPT вхідного символу в позиції n з використанням попередніх tokenів. Зелений колір вказує на правильне прогнозування; червоний – на неправильне [13].

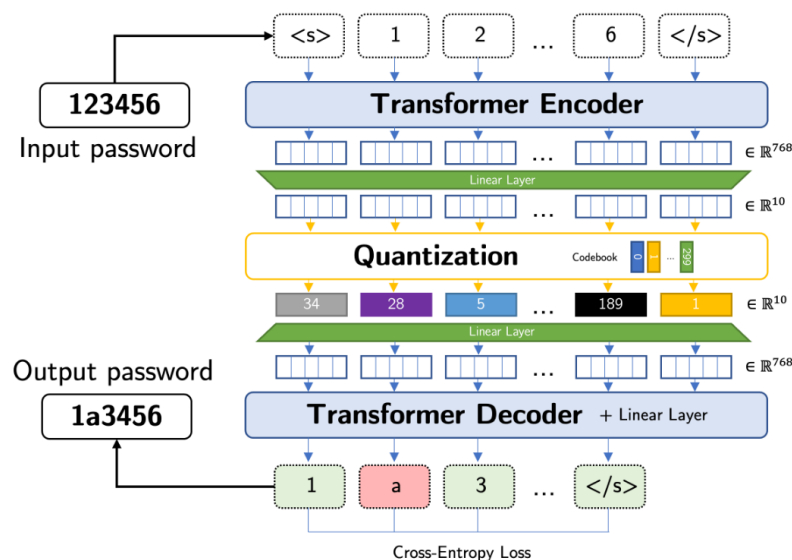


Рис.6. Стискання паролів в квантований латентний простір моделлю PassVQT [13].

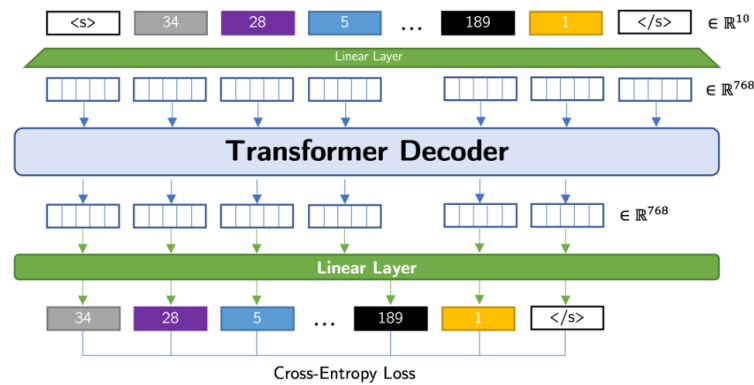


Рис.7. Параметризація умовного розподілу індексів моделлю PassVQT [13].

У [15] розглянуто шкідливі інформаційні впливи на кіберфізичні системи, що використовують методи машинного навчання. Для оцінки ефективності атак з ухиленням (як найбільш поширених) у цьому огляді пропонується класифікація атак, заснована на кількісних показниках, таких як рівень збурень та кількість змінених ознак, що наведена на рисунку 8.

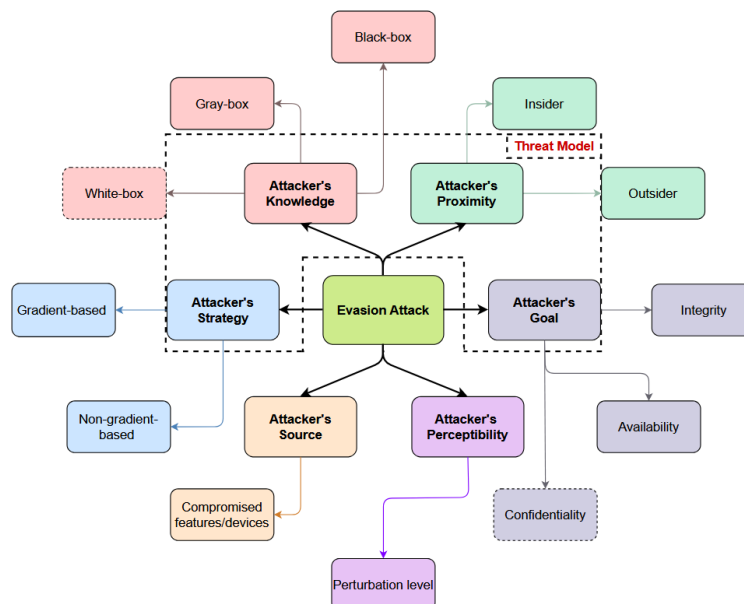


Рис.8. Класифікація атак [15].

Викладене надає можливість зрозуміти, що ІІІ починає відігравати дедалі значущу роль у сучасних кібератаках, зокрема в їх автоматизації. А також, що фішинг із залученням ІІІ та технологій *deepfake* для створення аудіо\відеовізуального контенту, на сьогодні є одним із найбільш небезпечних видів шкідливих інформаційних впливів.

Залучення ІІІ для захисту інфокомунікаційної мережі.

Різноманітність можливих способів інформаційних впливів, створення інформаційних впливів «на замовлення» припускають, що інформаційний вплив спершу буде невідомим системі захисту. Таким чином, у системі захисту обов'язково має бути реалізований механізм аналізу поведінки інфокомунікаційної системи.

Вже сьогодні для рішення задач захисту інформації ІІІ теж може бути дієвим інструментом. Традиційні рішення все ще актуальні для захисту, проте організаціям вже необхідно використовувати та розвивати свої ІІІ для захисту їх мережевої інфраструктури, даних тощо. Перспективним напрямом є активне залучення ІІІ до ідентифікації стану та захисту інфокомунікаційних мереж для протистояння кібератакам, особливо тим, які

виконуються із залученням ШІ, забезпечуючи виявлення, аналіз та реагування на кіберзагрози в реальному часі та в майбутньому стати одним із важливих компонентів захисту цифрових систем від нових і невідомих загроз. Необхідність ідентифікації стану інфокомунікаційної мережі для забезпечення належного захисту розглянуто в [15].

У [17] автори виконали систематичний огляд літератури (SLR), для аналізу останніх досліджень щодо LLM4Security [18,19] та надали комплексне картографування ландшафту, визначили, як LLM впроваджуються для покращення заходів кібербезпеки, що відображено на рисунку 9. Окрім того, автори детально розглянули використання LLM у різних сферах безпеки, виділивши шість основних сфер, що відповідають темам зібраних статей: безпека програмного забезпечення та систем, безпека мережі, безпека інформації та контенту, безпека апаратного забезпечення та безпека блокчейну, загалом 185 статей. На рисунку 10 зображено розподіл LLM у цих шести сферах [17].

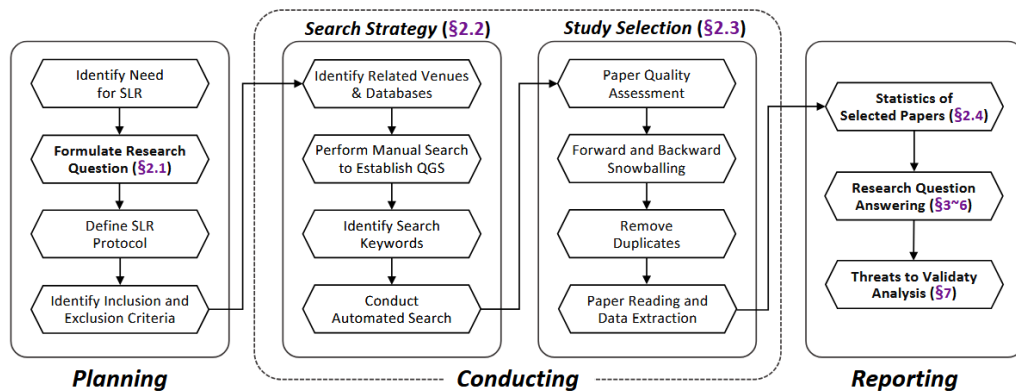


Рис.9. Схема систематичного огляду літератури для LLM4Security [17].

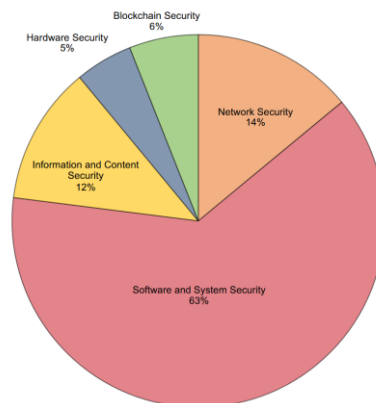


Рис.10. Розподіл використання LLM, орієнтованих на безпеку [17].

У [15] введено класифікацію захисту в машинному навчанні, засновану на чотирьох перспективах, що демонструють методи захисту від вхідних даних моделей до їх вихідних даних, яка подана на рисунку 11.

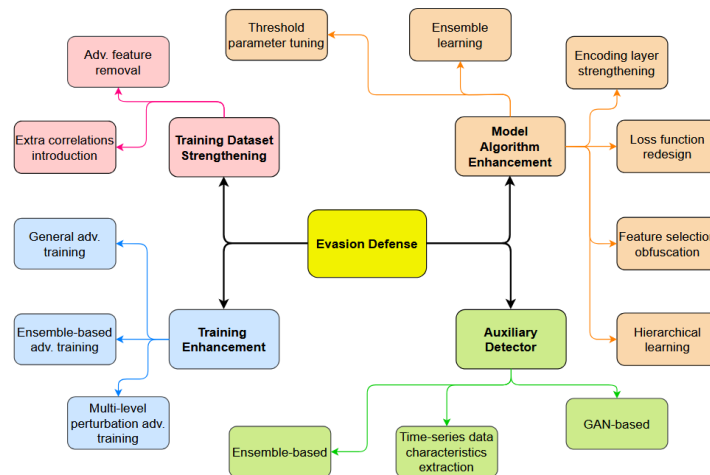


Рис.11. Класифікація захисту [15].

Зазначене додатково підтверджує доцільність використання ШІ для ефективного забезпечення інформаційної безпеки корпоративної інфокомунікаційної мережі та вдосконалення методів захисту.

Застосування ШІ для захисту інфокомунікаційних мереж відкриває нові можливості для виявлення, аналізу та реагування на кіберзагрози в реальному часі.

Перспективним напрямком є впровадження превентивних заходів, системи ШІ можуть використовувати прогностичні моделі для ідентифікації потенційних загроз ще до їх виникнення, аналізуючи поточні тенденції в поведінці мережі або користувачів.

5. Висновки та перспективи подальших досліджень

В результаті виконаного аналізу встановлено, що ШІ все більш широко залучається до формування шкідливих інформаційних впливів на інформаційні системи. При цьому не всі інформаційні впливи із залученням ШІ несуть однакову шкоду.

Серед потенційних загроз та сценаріїв інформаційних впливів із залученням ШІ, які можуть негативно вплинути на інформаційну безпеку інфокомунікаційної мережі, найбільш небезпечними на сьогодні варто вважати ті, які спрощують фішинг користувачів, генерують людиноподібний текст, правдоподібні голосові повідомлення та відео (deepfake), для більш переконливого введення в оману користувачів.

Перспективним напрямом є активне залучення ШІ до захисту інфокомунікаційних мереж для протистояння кібератакам, особливо тим, які виконуються із залученням ШІ, забезпечуючи виявлення, аналіз та реагування на кіберзагрози в реальному часі.

Подальші дослідження доцільно зосередити на тестуванні наявних нейромережевих моделей для оцінки можливості їх застосування при вирішенні задач захисту інфокомунікаційних мереж.

Список використаної літератури

1. Скіцько, О., Складанний, П., Ширшов, Р., Гуменюк, М., & Ворохоб, М. (2023). Загрози та ризики використання штучного інтелекту. Кібербезпека: освіта, наука, техніка, 2(22), 6–18. <https://doi.org/10.28925/2663-4023.2023.22.618>
2. Довженко, Н., Мазур, Н., Костюк, Ю., & Рзаєва, С. (2024). Інтеграція IoT та штучного інтелекту в інтелектуальні транспортні системи. Кібербезпека: освіта, наука, техніка, 2(26), 430–444. <https://doi.org/10.28925/2663-4023.2024.26.708>
3. Олійник, Я., Платоненко, А., Черевик, В., Ворохоб, М., & Шевчук, Ю. (2025). Методи захисту інформації в технологіях IoT. Кібербезпека: освіта, наука, техніка, 3(27), 100–108. <https://doi.org/10.28925/2663-4023.2025.27.705>
4. World Economic Forum: Global Cybersecurity Outlook 2025. URL: https://reports.weforum.org/docs/WEF_Global_Cybersecurity_Outlook_2025.pdf.

5. Otal H. T., Canbaz M. A. LLM HoneyPot: Leveraging Large Language Models as Advanced Interactive HoneyPot Systems. URL: <https://doi.org/10.1109/CNS62487.2024.10735607> (date of access: 09.09.2025).
6. Yazı Gholami. Large Language Models (LLMs) for Cybersecurity: A Systematic Review. World Journal of Advanced Engineering Technology and Sciences. 2024. Vol. 13, no. 1. P. 057–069. URL: <https://doi.org/10.30574/wjaets.2024.13.1.0395> (date of access: 09.09.2025).
7. The good, the bad, and the algorithm: The impact of generative AI on cybersecurity / L. Coppolino et al. Neurocomputing. 2025. Vol. 623. P. 129406. URL: <https://doi.org/10.1016/j.neucom.2025.129406> (date of access: 09.09.2025).
8. 2025 Phishing Statistics: (Updated August 2025) - Keepnet. Keepnet Labs. URL: <https://keepnetlabs.com/blog/top-phishing-statistics-and-trends-you-must-know> (date of access: 09.09.2025).
9. CERT-UA recorded 4,315 cyber incidents in 2024. State Service of Special Communications and Information Protection of Ukraine. URL: <https://cip.gov.ua/en/news/cert-ua-minulogo-roku-opracuyvala-4315-kiberincidentiv> (date of access: 09.09.2025).
10. Огляд кіберзагроз та стратегій захисту в 2025 році: досвід CERT-UA. Державна служба спеціального зв'язку та захисту інформації України. URL: <https://cip.gov.ua/ua/faqs/cyber-threat-overview-and-defense-strategies-in-2025-cert-ua-s-experience> (date of access: 09.09.2025).
11. Gamaredon Group, IRON TILDEN, Primitive Bear, ACTINIUM, Armageddon, Shuckworm, DEV-0157, Aqua Blizzard, Group G0047 | MITRE ATT&CK®. MITRE ATT&CK®. URL: <https://attack.mitre.org/groups/G0047/> (date of access: 09.09.2025).
12. ESET Research: Russia's Gamaredon APT group unleashed spearphishing campaigns against Ukraine with an evolved toolset | | ESET. ESET Newsroom. URL: <https://www.eset.com/us/about/newsroom/research/eset-research-russias-gamaredon-apt-group-unleashed-spearphishing-campaigns-against-ukraine-with-an-evolved-toolset/> (date of access: 09.09.2025).
13. Rando J., Perez-Cruz F., Hitaj B. PassGPT: Password Modeling and (Guided) Generation with Large Language Models. URL: <https://doi.org/10.48550/arXiv.2306.01545>.
14. GitHub - javirandor/passgpt. GitHub. URL: <https://github.com/javirandor/passgpt> (date of access: 09.09.2025).
15. S. Wang, R. K. L. Ko, G. Bai, N. Dong, T. Choi and Y. Zhang, "Evasion Attack and Defense on Machine Learning Models in Cyber-Physical Systems: A Survey," in IEEE Communications Surveys & Tutorials, vol. 26, no. 2, pp. 930-966, Secondquarter 2024, URL: doi.org/10.1109/COMST.2023.3344808 (date of access: 09.09.2025).
16. Chernihivskyi I., Kriuchkova L. Systematic approach to solving the task of protecting information in the infocommunication network from the influence of computer viruses. Cybersecurity: Education, Science, Technique. 2025. P. 572–590. URL: <https://doi.org/10.28925/2663-4023.2025.27.781> (date of access: 09.09.2025).
17. Large Language Models for Cyber Security: A Systematic Literature Review / H. Xu et al. URL: <https://doi.org/10.48550/arXiv.2405.04760> (date of access: 09.09.2025).
18. GitHub - tmylla/Awesome-LLM4Cybersecurity: An overview of LLMs for cybersecurity. GitHub. URL: <https://github.com/tmylla/Awesome-LLM4Cybersecurity> (date of access: 09.09.2025).
19. GitHub - liu673/Awesome-LLM4Security: This project aims to consolidate and share high-quality resources and tools across the cybersecurity domain. GitHub. URL: <https://github.com/liu673/Awesome-LLM4Security> (date of access: 09.09.2025).