

DOI 10.28925/2663-4023.2025.30.998УДК 004.8:004.056.53**Шульга Володимир Петрович**

доктор історичних наук, старший дослідник, ректор,
Державний університет інформаційно-комунікаційних технологій, Київ, Україна
ORCID: 0000-0003-4356-7288
v.shulha@duikt.edu.ua

Іванченко Євгенія Вікторівна

доктор технічних наук, професор,
директор навчально-наукового інституту Кібербезпеки та захисту інформації
Державний університет інформаційно-комунікаційних технологій, Київ, Україна
ORCID ID: 0000-0003-3017-5752
evivancenko@gmail.com

Берестяна Тетяна Володимирівна

аспірант
Державний університет інформаційно-комунікаційних технологій, Київ, Україна
ORCID: 0009-0007-1455-234X
t.berestiana@duikt.edu.ua

Шкурченко Олексій Анатолійович

аспірант
Державний університет інформаційно-комунікаційних технологій, Київ, Україна
ORCID: 0009-0006-9534-144X
o.shkurchenko@stud.duikt.edu.ua

МЕТОДИ ТА МОДЕЛІ ПРОТИДІЇ ГРУПОВИМ КІБЕРЗАГРОЗАМ НА ОСНОВІ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація. Розглянуто актуальну проблему протидії груповим кіберзагрозам, які характеризуються високим рівнем організованості, складною структурою та цілеспрямованістю. З розвитком цифрових технологій та зростанням залежності бізнесу і державних структур від інформаційних систем, кіберзагрози, здійснені координаційними групами, такими як АРТ-групи, ботнети та інші об'єднання кіберзлочинців, становлять дедалі більшу небезпеку. Запропоновано дослідження методів протидії груповим кіберзагрозам, що базується на застосуванні технологій штучного інтелекту, зокрема машинного навчання, глибинного навчання та методів обробки великих даних та архітектуру моделі, яка дозволяє виявляти ознаки координованої шкідливої активності, аналізувати поведінкові шаблони атакуювальників та вчасно реагувати на потенційні кіберзагрози. Особливу увагу приділено дослідженню адаптивних моделей, здатних до самонавчання в реальному часі та до ідентифікації нетипових відхилень від нормального функціонування інформаційної системи, які створено з урахуванням ключових особливостей групових атак, включаючи сценарії розподіленої атаки, прихованої комунікації між учасниками атаки, а також використання шифрування та технік ухилення від виявлення. Традиційні підходи до кіберзахисту виявляються малоефективними перед обличчям таких атак, оскільки не враховують динаміку поведінки кіберзагроз, швидкість їх еволюції та адаптації до захисних механізмів. Отримані результати дослідження підтверджують доцільність і ефективність використання інтелектуальних методів для протидії складним груповим кіберзагрозам. Запропоновані метод та моделі можуть бути інтегровані у наявні системи моніторингу кібербезпеки для підвищення їх стійкості до новітніх атак. Також окреслено перспективи подальших досліджень, зокрема щодо вдосконалення методів навчання моделей та їх масштабування для великих корпоративних мереж.



Ключові слова: групові кіберзагрози, штучний інтелект, машинне навчання, глибоке навчання, інформаційна безпека, моделі виявлення, адаптивні системи, кібератаки, обробка великих даних, поведінковий аналіз, кібербезпека.

ВСТУП

Актуальність теми зумовлена стрімким зростанням кількості кіберзагроз, які набувають дедалі складнішого та скоординованого характеру. Групові атаки в кіберпросторі, що здійснюються організованими угрупованнями або через автоматизовані ботнети, становлять серйозну небезпеку для інформаційної безпеки державних структур, бізнесу та критичної інфраструктури. Такі кіберзагрози часто мають високий рівень складності, використовують новітні технології маскування та обходу традиційних засобів кіберзахисту, що значно ускладнює їх виявлення та нейтралізацію. Проблематика групових кіберзагроз полягає у їхній координації, масштабованості та здатності швидко адаптуватися до змін у системах кіберзахисту. Традиційні методи кібербезпеки, зокрема сигнатурні та евристичні аналізи, виявляються недостатніми для своєчасного реагування на атаки, що реалізуються колективно та змінюють свою поведінку в реальному часі. У цьому контексті особливої значущості набуває використання штучного інтелекту (ШІ) у сфері кіберзахисту. Алгоритми машинного навчання, глибокого аналізу даних та нейронні мережі дозволяють здійснювати поведінковий аналіз, прогнозування кіберзагроз та автоматизоване прийняття рішень щодо протидії скоординованим атакам. ШІ здатний виявляти приховані закономірності, оперативно реагувати на нові типи кіберзагроз і забезпечувати динамічну адаптацію систем кіберзахисту.

Метою є дослідження ефективних методів і моделей протидії груповим кіберзагрозам у кіберпросторі з використанням засобів штучного інтелекту. Для досягання мети необхідно вирішити наступні задачі:

- провести аналіз існуючих підходів до виявлення та протидії груповим кіберзагрозам;
- дослідити моделі виявлення скоординованої шкідливої активності на основі алгоритмів ШІ;

Результати цього дослідження сприятимуть підвищенню рівня захищеності інформаційних систем та удосконаленню механізмів реагування на сучасні кіберзагрози.

Аналіз сучасного стану проблеми. Розширена постійна кіберзагроза (APT) - це форма кібератаки або діяльності хакерського угруповання, що застосовує комплексні техніки для здійснення тривалих і скоординованих атак на державні чи приватні організації. Такі атаки зазвичай реалізуються висококваліфікованими та добре організованими групами, які мають значне фінансування та часто діють за підтримки державних структур або в тісному зв'язку з ними. [1]

Основною особливістю APT є їхня здатність тривалий час залишатися непоміченими в інформаційній системі жертви, систематично збираючи конфіденційні дані або спричиняючи порушення у роботі критично важливих сервісів. Зловмисники використовують високорівневі методи соціальної інженерії, спеціалізоване шкідливе програмне забезпечення та інші способи прихованого проникнення й закріплення у мережі.

Однією з ключових тактик APT є так зване «бокове переміщення» (lateral movement) - поступове розширення доступу всередині мережі після початкового проникнення, що дозволяє атакуючим отримувати підвищені привілеї та переміщатися між вузлами системи. Через це APT можуть залишатися в системі жертви протягом місяців або навіть років, перш ніж їх буде виявлено. [1]



Для досягнення своїх цілей АРТ-групи застосовують низку різноманітних векторів атак:

Фішинг - цільові електронні листи, спрямовані на введення в оману ключових працівників організації з метою викрадення даних або встановлення шкідливого ПЗ. Часто доповнюється такими методами, як атаки типу «водопій», компрометація ділової електронної пошти (BEC) або телефонні шахрайства (вішинг). [2]

Експлуатація вразливостей нульового дня (zero-day) - використання нещодавно виявлених і ще не виправлених уразливостей у програмному забезпеченні або апаратних засобах для проникнення в мережу. [2]

Шкідливе програмне забезпечення (malware) - спеціально розроблені засоби, зокрема трояни для віддаленого доступу (RAT), клавіатурні шпигуни тощо, які забезпечують стійку присутність зловмисника в системі та використовують засоби обходу кіберзахисту. [2]

Компрометація постачальників - атака на довірених партнерів або постачальників із метою проникнення у внутрішню інфраструктуру клієнта. Подібні атаки можуть впливати на велику кількість організацій через уразливості в спільному середовищі, зокрема у хмарних сервісах. [2]

Фізичний доступ - у деяких випадках АРТ можуть вдатися до фізичного втручання, наприклад, викрадення пристроїв, несанкціонований доступ до серверних приміщень або пряме втручання у роботу апаратного забезпечення. [2]

АРТ-групи, як правило, комбінують кілька з перелічених методів для досягнення максимальної ефективності та уникнення виявлення. Така багатовекторність та гнучкість роблять АРТ одними з найскладніших для виявлення та протидії сучасних кіберзагроз. [2,3]

До прикладу розглянемо Групу FIN7, яка зазвичай розпочинала свої кібератаки зі спрямування фішингових електронних листів на співробітників цільових компаній. Такі листи містили вкладення, зазвичай у формі на вигляд звичайного документа Microsoft Word, у якому приховувалося шкідливе програмне забезпечення. Зміст повідомлень був стилізований під легітимну ділову кореспонденцію, що мало на меті переконати адресата відкрити файл та несвідомо активувати шкідливий код, який інфікував систему. [4]

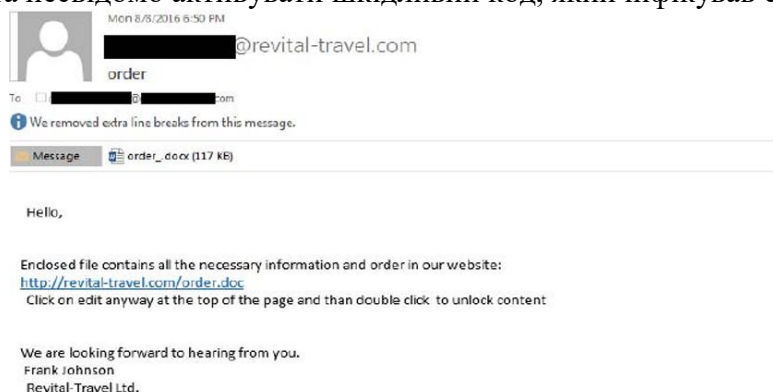


Рис. 1. Фішинговий email FIN7

Залежно від типу установи, сценарії фішингових повідомлень змінювалися. Наприклад, у випадку готельного бізнесу електронний лист міг містити запит щодо бронювання номеру, з деталями у вкладеному файлі. У сфері ресторанного



обслуговування повідомлення могло містити скаргу на якість обслуговування або запит на велике замовлення з детальним описом у вкладенні. [5]

У багатьох випадках FIN7 доповнювала фішингову кампанію телефонним дзвінком співробітнику компанії-жертви. Під час розмови зловмисник повторював тему листа, з метою підвищення довіри до повідомлення та стимулювання відкриття вкладеного файлу. Така комбінація електронної та голосової соціальної інженерії підвищувала ефективність атаки та ймовірність активації шкідливого програмного забезпечення. [6,7,8]

Ботнет - це децентралізована мережа пристроїв, заражених шкідливим програмним забезпеченням, яка перебуває під контролем кіберзлочинців. Така інфраструктура формується без відома користувачів і використовується для реалізації широкого спектра зловмисних дій.

Ботнети зазвичай застосовуються для масової розсилки спаму, встановлення шпигунського програмного забезпечення, а також для викрадення персональних даних та облікових записів. Однією з найпоширеніших кіберзагроз є використання ботнетів для здійснення розподілених атак типу відмови в обслуговуванні (DDoS), які перевантажують сервери надлишковим трафіком і можуть призвести до уповільнення роботи вебресурсу або повної недоступності сервісу. [10]

Інфікування пристроїв ботнет-шкідниками часто відбувається через фішингові листи з шкідливими вкладеннями, завантаження заражених програм або файлів з неперевірених джерел. Зловмисники також активно експлуатують вразливості, спричинені використанням застарілого програмного забезпечення та недостатнім кіберзахистом інтернет-з'єднання. В останні роки спостерігається зростання кількості атак на пристрої Інтернету речей (IoT), включно з камерами відеоспостереження, смарт-телевізорами та навіть підключеними до мережі транспортними засобами. [10]

Даркнети - це інструменти пасивного моніторингу мережі, які базуються на діапазонах IP-адрес, оголошених у системах маршрутизації, однак не асоціюються з реальними сервісами або активними вузлами. Вони постійно фіксують вхідний трафік, який є небажаним за своєю природою, і тому даркнети стають цінним джерелом інформації для досліджень у сфері мережевої безпеки. Відсутність легітимного трафіку дозволяє ефективніше виявляти потенційні кіберзагрози, зокрема сканування внутрішніх сегментів мережі, атаки методом перебору облікових даних тощо. [10]

Ключовим завданням є виявлення та аналіз скоординованих дій у даркнет-трафіку, що може значно знизити обсяг даних, які підлягають ручному аналізу з боку фахівців з кібербезпеки, та надати цілісніше уявлення про характер поточних кібератак у глобальному Інтернеті. З огляду на велику кількість вихідного трафіку, що надходить до даркнетів, ручна обробка такого масиву інформації є малоефективною. Крім того, існує дефіцит структурованої інформації для глибокого аналізу шаблонів цього трафіку. [10,11,12]

У межах даного дослідження розглядається застосування неконтрольованих методів інтелектуального аналізу даних для автоматичного виявлення скоординованих дій серед вихідних IP-адрес, що взаємодіють з даркнет-інфраструктурою. Дослідження базується на двох ключових гіпотезах: (1) IP-адреси, контрольовані одним зловмисним суб'єктом, повинні демонструвати схожий рівень активності, що дозволяє використовувати інтенсивність трафіку як ознаку координації; (2) координовані джерела, ймовірно, надсилають трафік у певному часовому співпадінні, що дає змогу виявляти кореляцію у часі між різними адресами.

Для перевірки цих гіпотез було досліджено повноцінний машинний конвеєр з використанням алгоритмів кластеризації та побудови ознак. У випадку аналізу



інтенсивності трафіку створено набір характеристик, що описують активність IP-адрес, на основі яких здійснено кластеризацію за допомогою неконтрольованих алгоритмів. Для аналізу часової координації застосовано метод word2vec - алгоритм, що використовується в обробці природної мови для виявлення контекстно подібних слів - адаптований для оцінки спільної появи IP-адрес у мережевому трафіку [13]. З метою забезпечення достовірності результатів, використано набір відомих координованих IP-адрес, зокрема з відкритих джерел, пов'язаних із системами мережевого сканування. Результати демонструють високу ефективність використаних ознак у виявленні як інтенсивності, так і часової узгодженості трафіку. Аналіз показав, що IP-адреси, що належать до однієї координативної групи, з високою точністю (89%) класифікуються до одного просторового сегмента.

Утім, аналіз часової стабільності кластерів показав варіативність у складі груп у різні дні, що вказує на змінну поведінку координованих джерел і вимагає застосування моделей до обмежених часових інтервалів для збереження точності.

Загалом дослідження демонструє, що неконтрольовані методи машинного навчання здатні ефективно зменшити обсяг даних для аналізу та підвищити ефективність виявлення скоординованої кіберзлочинної активності у даркнетах, однак їх ефективність тісно пов'язана з правильним вибором часових рамок обробки [14,15,16].

Штучний інтелект став невід'ємною складовою сучасних систем кіберзахисту. Його використання дозволяє підвищити ефективність виявлення кіберзагроз, автоматизувати аналіз інцидентів і зменшити час реагування на атаки. Водночас успішне впровадження AI/ML/DL-рішень вимагає обережного підходу, з урахуванням етичних, технічних та практичних аспектів. У майбутньому розвиток гібридних систем, які поєднують інтелектуальні алгоритми та експертні знання, стане ключовим напрямом підвищення кіберстійкості організацій. Огляд особливостей моделей штучного інтелекту наведено наведено в таблиці №1.

Таблиця 1

Огляд особливостей моделей штучного інтелекту

Моделі та їх особливості	Deep Learning Модель	Random Forest	Isolation Forest
Основне застосування	Комплексне розпізнавання образів	Класифікація та регресія	Виявлення аномалій
Сильні сторони	Висока точність у різних середовищах, багатих на дані	Висока точність і стійкість до переобладнання	Ефективне для виявлення аномалій без навчальних даних
Обмеження	Вимагає значних обчислювальних ресурсів	Може бути вимогливим до обчислювальних ресурсів	Може мати вищі показники помилкових спрацьовувань в деяких контекстах
Вимоги до даних	Підтримуються різні типи	Дані з мітками для навчання	Дані без міток або з мінімальною підготовкою
Обчислювальна інтенсивність	Високий	Від помірної до високої	Помірна
Адаптивність до нових кіберзагроз	Корисно для глибокого аналізу даних	Помірна	Висока
Використання у розслідуванні кіберзлочинів	Комплексне розпізнавання образів	Ефективна у класифікації та визначенні відхилень	Ефективна у виявленні прихованих аномалій



Розширена порівняльна таблиця надає глибокий аналіз кількох моделей штучного інтелекту, висвітлюючи сфери їхнього застосування, основні переваги та обмеження, а також дозволяє оцінити їхню ефективність у контексті розслідування кіберзлочинів.

Моделі з керованим навчанням, зокрема Random Forest, демонструють високу ефективність у задачах класифікації та прогнозування, особливо за наявності великих обсягів розмічених даних. Однак їх ефективність знижується при появі нових, невідомих кіберзагроз через залежність від наперед визначених прикладів. [12, 23]

У свою чергу, некеровані моделі, як-от Isolation Forest, мають здатність виявляти аномалії та нові закономірності без потреби в попередньому маркуванні даних, що надає їм суттєву перевагу для ідентифікації нових типів атак. Проте ці моделі схильні до підвищеного рівня хибнопозитивних результатів.

Моделі глибинного навчання особливо ефективні у роботі зі складними структурами даних, такими як зображення або послідовності подій, забезпечуючи глибокий рівень аналізу. Водночас їх використання вимагає значних обчислювальних ресурсів та великих обсягів даних для навчання.[16]

Традиційні рішення, такі як антивірусне програмне забезпечення, міжмережеві екрани (firewalls), IDS/IPS та системи, що базуються на сигнатурному виявленні, мають низку обмежень. Вони зазвичай розраховані на вже відомі кіберзагрози й працюють на основі попередньо заданих шаблонів. Такий підхід не дозволяє ефективно виявляти нові, цілеспрямовані або багатоетапні атаки (APT), які змінюють свої вектори дій у режимі реального часу.

Однією з головних проблем є низька гнучкість традиційних систем - вони повільно реагують на нові типи атак і не можуть самостійно адаптуватися до нових сценаріїв кіберзагроз. Інша слабка сторона - висока кількість помилкових спрацьовувань (false positives), які ускладнюють роботу аналітиків і знижують ефективність реагування.

Крім того, масштабність та динамічність сучасних цифрових інфраструктур, включно з хмарними обчисленнями, IoT, мобільними пристроями, створює складне середовище, яке класичні інструменти безпеки не можуть охопити повністю. Зловмисники все частіше використовують автоматизовані інструменти, штучний інтелект та скоординовані атаки, які обходять статичний кіберзахист.

У роботі робиться акцент на потребі впровадження інтелектуальних, адаптивних методів кіберзахисту, які включають машинне навчання (ML), аналіз поведінки користувачів та мережевого трафіку (UEBA), а також оркестрацію та автоматизацію відповідей на інциденти (SOAR). Ці підходи дозволяють виявляти аномалії, прогнозувати кіберзагрози, швидко адаптувати кіберзахист до нових умов і зменшувати навантаження на аналітиків безпеки.

Таким чином, ефективна кібербезпека сьогодні вимагає переосмислення стратегій - від реактивних до проактивних, від статичних до адаптивних, від ізольованих систем до комплексних екосистем безпеки, інтегрованих на основі штучного інтелекту та автоматизації.[17, 24]

Вибір алгоритмів штучного інтелекту у сфері кіберзахисту тісно пов'язаний з типом задачі, яку потрібно вирішити. Здебільшого ці задачі поділяються на три основні категорії: класифікація, кластеризація та виявлення аномалій. Кожна з них має свої цілі, характерні підходи та оптимальні алгоритми.

Класифікація (Supervised Learning) використовується тоді, коли доступні мічені (labelled) дані. Метою є розпізнавання кіберзагроз або дій на основі відомих прикладів. Такий підхід ефективний для виявлення шкідливого чи легітимного трафіку, ідентифікації типів атак (наприклад, DoS, фішинг або SQL-ін'єкції), а також виявлення зловмисників



серед користувачів. Найбільш поширені алгоритми в цій категорії – Random Forest, який вирізняється точністю та стійкістю до шумів, Support Vector Machine (SVM) – ефективний у випадках, коли простір ознак складний, Logistic Regression – простий у реалізації та добре підходить для пояснення результатів, а також Neural Networks (MLP), які здатні виявляти глибокі залежності у складних наборах даних, хоча потребують більше ресурсів і часу на навчання.

Кластеризація (Unsupervised Learning) застосовується тоді, коли немає мічених даних, і потрібно групувати об'єкти за схожими характеристиками. Це корисно для виявлення нових типів атак, поведінкової сегментації користувачів та побудови профілів активності. Серед популярних методів – K-Means, який швидкий, але вимагає заздалегідь визначеної кількості кластерів; DBSCAN, який краще працює з даними, що містять викиди; Hierarchical Clustering, який дозволяє бачити ієрархію зв'язків між групами; а також Autoencoders, які дають можливість виявити приховані закономірності у великих наборах даних. [18]

Виявлення аномалій (Anomaly Detection) – це ще один важливий напрям, що фокусується на пошуку відхилень від нормальної поведінки. Такий підхід допомагає знаходити zero-day атаки, виявляти зламані облікові записи, внутрішні кіберзагрози або ботнети. Тут часто застосовуються алгоритми Isolation Forest, який добре масштабується на великих даних і точно ізолює нетипові об'єкти; One-Class SVM, що навчається лише на нормальних даних і фіксує відхилення; Autoencoders, які відстежують розбіжності між вхідними та відновленими даними; а також LOF (Local Outlier Factor), який порівнює локальну щільність об'єктів для виявлення "випадків" з-поза типового контексту. [18]

Вибір **оптимального алгоритму** залежить від кількох факторів: чи доступні мічені дані, наскільки складна задача (наприклад, просте виявлення кіберзагроз чи складний поведінковий аналіз), обсяг даних і ресурси для обробки, а також потреба в інтерпретованості – наприклад, Logistic Regression легше пояснити, ніж Autoencoders. Якщо маєш конкретну задачу в сфері кіберзахисту – можу допомогти підібрати найбільш доречний алгоритм. [18]

Обробка великих масивів даних у сфері кіберзахисту є одним із ключових чинників ефективного кіберзахисту інформаційних систем від сучасних кіберзагроз. Зі стрімким зростанням обсягів інформації, що генерується ІТ-інфраструктурами, виникає необхідність у розробці і впровадженні новітніх технологій для її аналізу в реальному часі. Основна мета таких підходів – своєчасне виявлення кіберзагроз, аналіз аномалій, кореляція подій та прогнозування потенційних атак, що дозволяє мінімізувати наслідки кіберінцидентів і підвищити загальну безпеку організацій.

Сучасні інформаційні системи генерують колосальні обсяги різномірних даних, серед яких основне місце займають журнали системних і мережових подій (логи), телеметрія пристроїв, мережовий трафік, а також сигнали від систем безпеки, таких як IDS/IPS, SIEM, антивірусні та інші. Ці дані мають різний формат, швидкість надходження та ступінь важливості, що зумовлює необхідність використання спеціалізованих платформ і підходів для їхньої збирання, зберігання, обробки та аналізу.

Термін Big Data у контексті кібербезпеки охоплює великі обсяги структурованої та неструктурованої інформації, що генерується з великою швидкістю та у великих обсягах. До таких даних відносяться журнали доступу, мережовий трафік, події безпеки, звіти та алерти SIEM-систем. Для ефективної роботи з цими даними застосовуються технології розподіленої обробки, що дозволяють паралельно аналізувати великі масиви інформації та виявляти у них закономірності.



До ключових платформ для роботи з Big Data належать Hadoop та Apache Spark. Hadoop забезпечує розподілене зберігання і обробку даних на сотнях чи тисячах серверів, дозволяючи масштабуватися горизонтально. Spark, у свою чергу, дозволяє виконувати швидкий аналіз даних у пам'яті, що є критичним для задач, пов'язаних з кібербезпекою, де час реакції має першорядне значення. Для зберігання великих обсягів даних із різною структурою використовують NoSQL-бази даних, наприклад, Cassandra або MongoDB, які здатні швидко індексувати та обробляти інформацію в реальному часі.

Одним із популярних підходів є створення Data Lake - централізованого сховища «сирих» даних, що дозволяє зберігати інформацію в оригінальному вигляді без попередньої обробки. Це дає змогу згодом виконувати глибокий аналіз, застосовувати алгоритми машинного навчання та інші аналітичні методи для виявлення складних шаблонів атак, які складно побачити при традиційних методах. [19]

В умовах сучасного кіберзахисту особливо важлива можливість обробляти події в режимі реального часу. Переважна більшість інцидентів починається з незначних ознак, які потрібно оперативно виявити та заблокувати. Саме тому критичною стає технологія потокової обробки даних, яка дозволяє безперервно збирати, фільтрувати та аналізувати інформацію, що надходить з мережевих сенсорів, пристроїв, логів користувачів і систем безпеки.

Для цих задач застосовуються такі платформи, як Apache Kafka, що служить для надійного збору та передачі повідомлень у потоках, а також Apache Flink та Apache Storm, які виконують обробку даних з мінімальними затримками. Інтеграція поточкових даних з SIEM-системами (наприклад, Splunk, ELK Stack, IBM QRadar) дає змогу виявляти аномалії, корелювати події і реагувати на кіберзагрози миттєво, що суттєво підвищує ефективність кіберзахисту.

Телеметрія - це один із найважливіших джерел даних у кібербезпеці. Вона включає інформацію про стан операційних систем, поведінку користувачів, стан мережевих пристроїв, активність додатків та інші метрики, які надходять з кінцевих точок - серверів, мобільних пристроїв, IoT-пристроїв тощо. Аналіз телеметрії дозволяє виявляти компрометації на ранніх стадіях, відслідковувати аномальну поведінку користувачів та пристроїв, а також формувати контекст для подальшого розслідування інцидентів. [19]

Завдяки поєднанню телеметрії з методами штучного інтелекту та машинного навчання, сучасні системи кіберзахисту можуть автоматично виявляти складні патерни поведінки та потенційні кіберзагрози, мінімізуючи навантаження на аналітиків і підвищуючи швидкість реакції на інциденти. [19]

У сучасних системах кібербезпеки надзвичайно важливо розуміти, наскільки добре вони виконують свої функції. Адже від цього залежить, чи вдасться вчасно виявити та зупинити атаки, чи будуть помічені найдрібніші прояви кіберзагроз, а також як швидко організація зможе відреагувати на небезпеку. Для оцінки ефективності використовують кілька ключових показників, серед яких найважливішими є точність, повнота та швидкість реагування.

Точність показує, наскільки система здатна правильно визначати реальні кіберзагрози, коли вона повідомляє про інцидент. Висока точність означає, що більшість сповіщень про атаки - це справжні кіберзагрози, а не помилкові тривоги. Це дуже важливо, адже у випадку низької точності аналітики отримують багато «хибних спрацьовувань», які займають їхній час і ресурси. Якщо система постійно сигналізує про зкіберзагрозу там, де її немає, користувачі можуть почати ігнорувати ці сповіщення або просто втомляться від постійних попереджень. Це створює ризик, що реальна атака залишиться непоміченою через «шум» у системі. Тому висока точність допомагає підтримувати довіру до системи



і зберігати увагу команди безпеки. Повнота (або чутливість) відображає, яку частку всіх реальних кіберзагроз система змогла знайти і повідомити про них. Висока повнота означає, що жодна або майже жодна атака не залишилася непоміченою. Це особливо важливо в умовах кібербезпеки, адже пропущена кіберзагроза може призвести до серйозних наслідків - від витоку даних до повного злома системи.

Якщо система має низьку повноту, це означає, що багато реальних атак проходять повз кіберзахист і залишаються непоміченими, що може бути критично небезпечно. Однак підвищення повноти часто призводить до зростання кількості помилкових спрацьовувань, бо система намагається виявляти навіть сумнівні активності.

У реальних умовах системи кібербезпеки намагаються знайти баланс між цими двома метриками. Якщо робити систему надто «пильну», вона буде знаходити майже всі кіберзагрози, але разом з тим створюватиме багато помилкових тривог. Якщо ж зосередитися лише на високій точності, можна пропустити деякі важливі атаки.

Для кращого розуміння роботи системи використовують комплексні показники, які враховують одночасно і точність, і повноту. Це дозволяє зробити більш об'єктивну оцінку і вибрати оптимальні параметри системи відповідно до конкретних потреб організації.

Швидкість реагування - це ще одна критично важлива метрика, що визначає, скільки часу проходить від моменту виявлення кіберзагрози до вжиття відповідних заходів. У кібербезпеці це питання життя і смерті, адже атаки можуть розвиватися дуже швидко - за лічені секунди чи хвилини.

Якщо система вчасно і автоматично повідомляє про інцидент і навіть може самостійно застосувати заходи для його блокування (наприклад, ізолювати заражений пристрій або заблокувати підозрілий трафік), це значно знижує ризики і допомагає уникнути масштабних втрат.

Повільне реагування, навпаки, дає зловмисникам час для розвідки, проникнення в систему та поширення атак на інші компоненти мережі. Тому в сучасних рішеннях активно застосовують автоматизацію та штучний інтелект, щоб скоротити час реагування до мінімуму.

Кожна з цих метрик важлива сама по собі, але тільки їхнє поєднання дає повну картину ефективності системи. Висока точність і повнота без швидкої реакції не допоможуть зупинити атаку вчасно. Швидка реакція без точності призведе до хаосу через безліч помилкових тривог, а хороша повнота без точності змусить команду безпеки витрачати багато часу на аналіз незначущих сигналів.

Ідеальна система кіберзахисту - це така, яка швидко і точно виявляє майже всі кіберзагрози, при цьому не перевантажуючи аналітиків помилковими сповіщеннями і оперативно реагуючи на інциденти [20,21,22].

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Використання штучного інтелекту (ШІ), машинного навчання (ML) та глибокого навчання (DL) у сфері кібербезпеки, зокрема для протидії груповим кіберзагрозам, відкриває нові можливості і значно підвищує ефективність кіберзахисту інформаційних систем. Групові атаки, які часто характеризуються складною координацією між великою кількістю зловмисних джерел, ставлять перед традиційними методами численні виклики, пов'язані з масштабом, швидкістю та варіативністю атак. Саме тут ШІ демонструє свої ключові переваги.

По-перше, штучний інтелект дозволяє автоматизувати процеси виявлення та аналізу кіберзагроз, що значно знижує навантаження на аналітиків безпеки. Завдяки алгоритмам



неконтрольованого навчання ШІ може виділяти аномальні патерни поведінки навіть у разі відсутності чітких підказок або позначених даних. Це критично важливо для ідентифікації нових або маловідомих типів атак, які можуть здійснюватись скоординовано кількома джерелами.

По-друге, ШІ має здатність до масштабованості, що робить його незамінним у роботі з великими обсягами даних, характерними для сучасних мереж. Завдяки методам обробки великих даних і потокової аналітики, штучний інтелект здатен оперативно аналізувати інформацію з різних джерел — логів, телеметрії, мережевого трафіку — і миттєво реагувати на кіберзагрози. Така швидкість є вирішальною у протидії розподіленим атакам, які розгортаються вкрай динамічно.

По-третє, ШІ забезпечує гнучкість і адаптивність систем кіберзахисту. Традиційні правила та сигнатури швидко стають неактуальними у відповідь на зміни тактик зловмисників. Моделі машинного навчання можуть постійно оновлюватись, підлаштовуючись під нові кіберзагрози та поведінку атакуючих. Це дозволяє зберігати високу ефективність навіть у швидкоплинних умовах кіберпростору.

Крім того, ШІ сприяє підвищенню точності виявлення кіберзагроз. Алгоритми здатні виділяти ключові ознаки координації між IP-адресами або поведінковими патернами, що робить можливим ідентифікацію скоординованих атак, які інакше могли б залишитись непоміченими. Це значно покращує загальну безпеку та зменшує кількість помилкових спрацьовувань, що допомагає фокусувати ресурси на реальних інцидентах.

Також важливо відзначити, що ШІ може ефективно працювати у комплексі з іншими системами кіберзахисту, інтегруючись із SIEM, SOAR, а також інструментами автоматизованого реагування. Це дозволяє створювати багаторівневі системи, які не лише виявляють кіберзагрози, а й оперативно нейтралізують їх, мінімізуючи шкоду від атак.

Незважаючи на значні досягнення, застосування штучного інтелекту у протидії груповим кіберзагрозам залишається активно розвиваючоюся сферою з великим потенціалом для подальших досліджень. Одним із напрямків є покращення точності моделей за допомогою більш глибокого інтегрування контекстної інформації та багатовимірного аналізу поведінки мережевого трафіку. Наприклад, комбінування часових, просторових та поведінкових характеристик у рамках однієї моделі може допомогти краще виявляти скоординовані атаки і відокремлювати їх від випадкових збігів.

Одним із напрямків є покращення точності моделей за допомогою більш глибокого інтегрування контекстної інформації та багатовимірного аналізу поведінки мережевого трафіку. Наприклад, комбінування часових, просторових та поведінкових характеристик у рамках однієї моделі може допомогти краще виявляти скоординовані атаки і відокремлювати їх від випадкових збігів.

Інший важливий напрямок — розвиток методів самонавчання та адаптивного машинного навчання, які дозволяють системам не лише підлаштовуватися під нові кіберзагрози, але й передбачати потенційні атаки на основі аналізу трендів. Такі системи могли б працювати проактивно, попереджаючи про ризики до їхнього фактичного виникнення.

Додатково, існує потреба у створенні більш прозорих та інтерпретованих моделей штучного інтелекту. Сучасні алгоритми глибокого навчання часто працюють як «чорні скриньки», що ускладнює розуміння причин їхніх рішень. Для кібербезпеки це критично, оскільки аналітики повинні не лише отримувати сигнали про кіберзагрози, а й розуміти їхній контекст для прийняття правильних рішень.

Важливим аспектом також є питання етики, безпеки та кіберзахисту даних, пов'язаних із використанням ШІ. Зловмисники можуть намагатися обходити або вводити



в оману моделі штучного інтелекту, тому розробка стійких до атак ШІ-систем — ще одна актуальна задача.

Не менш перспективними є дослідження у сфері колаборативного машинного навчання (federated learning), що дозволяє різним організаціям об'єднувати досвід і моделі без обміну чутливими даними. Це може сприяти більш ефективному виявленню глобальних кіберзагроз, зокрема координації групових атак, які зазвичай охоплюють багато цілей.

Окрім технічних аспектів, великим викликом залишається інтеграція AI-рішень у реальні бізнес-процеси та системи кібербезпеки, навчання персоналу і підвищення їхньої компетенції в роботі з новими інструментами. Майбутні дослідження мають включати не лише алгоритмічні інновації, а й питання організаційної готовності, регуляторної бази і взаємодії між людьми і машинами.

Отже, штучний інтелект має потенціал докорінно змінити підходи до протидії груповим кіберзагрозам, підвищуючи ефективність, швидкість та точність виявлення атак, а також автоматизуючи багато рутинних процесів. Водночас ця галузь продовжує активно розвиватися, і подальші дослідження спрямовані на подолання існуючих викликів, підвищення прозорості моделей, адаптивності та безпеки AI-систем. Лише комплексний підхід, що поєднує технологічні інновації з організаційними та людськими факторами, дозволить максимально ефективно використовувати штучний інтелект у боротьбі з сучасними кіберзагрозами.

Таким чином, методи та моделі протидії груповим кіберзагрозам на основі штучного інтелекту демонструють високу ефективність у виявленні, аналізі та нейтралізації скоординованих атак, забезпечуючи адаптивність, швидкодію та гнучкість сучасних систем кібербезпеки. Їх подальший розвиток є ключовим чинником для забезпечення сталого цифрового кіберзахисту в умовах зростаючої складності кіберзагроз.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Андреев, О. О. (2021). *Методи виявлення та протидії кібератакам у корпоративних мережах* : монографія. Київ: НТУУ «КПІ».
2. Альянс Лазаря. (2023, 8 березня). *Що таке розширені стійкі кіберзагрози (APT)?* <https://lazarusalliance.com/uk/what-are-advanced-persistent-threats-apt/>
3. Козлов, Д. С., & Терещенко, Л. В. (2022). Інтелектуальні системи захисту інформації: методи машинного навчання. *Інформаційна безпека*, (3), 42–50.
4. Skopik, F., Settanni, G., & Fiedler, R. (2016). A problem shared is a problem halved: A survey on the dimensions of collective cyber defense through security information sharing. *Computers & Security*, 60, 154–176. <https://doi.org/10.1016/j.cose.2016.03.011>
5. Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176. <https://doi.org/10.1109/COMST.2015.2494502>
6. Зверев, І. В. (2022). Штучний інтелект як інструмент у системах інформаційної безпеки. *Сучасні інформаційні технології*, (2), 28–34.
7. Sommer, R., & Paxson, V. (2010). Outside the closed world: On using machine learning for network intrusion detection. *IEEE Symposium on Security and Privacy*, 305–316. <https://doi.org/10.1109/SP.2010.25>
8. Sangkatsanee, P., Wattanapongsakorn, N., & Charnsripinyo, C. (2011). Practical real-time intrusion detection using machine learning approaches. *Computer Communications*, 34(18), 2227–2235. <https://doi.org/10.1016/j.comcom.2011.05.009>
9. Колесник, С. П., Овчарук, В. М., Пархоменко, А. П., та ін. (2020). *Алгоритми штучного інтелекту для систем кіберзахисту* : навч. посіб. Харків: ХНУРЕ.
10. ESET. (n.d.). *Захист від ботнетів. Як не стати частиною ботнет-мережі.* <https://www.eset.com/ua/support/information/entsyklopediya-zahroz/zakhyst-vid-botnetiv/>



11. Berman, D. S., Buczak, A. L., Chavis, J. S., & Corbett, C. L. (2019). A survey of deep learning methods for cyber security. *Information*, 10(4), 122. <https://doi.org/10.3390/info10040122>
12. Tang, T. A., Mhamdi, L., McLernon, D., Zaidi, S. A. R., & Ghogho, M. (2016). Deep learning approach for network intrusion detection in software defined networking. *2016 International Conference on Wireless Networks and Mobile Communications (WINCOM)*, 258–263. IEEE. <https://doi.org/10.1109/WINCOM.2016.7777224>
13. Куліш, В., & Пастух, І. (2023). Сучасні методи виявлення ботнет-мереж у корпоративних системах. *Захист інформації*, (1), 63–70.
14. U.S. Department of Justice. (n.d.). *HOW FIN7 attacked and stole data: Sophisticated social engineering – phishing & calling*. <https://www.justice.gov/archives/opa/press-release/file/1084361/dl?inline=1>
15. Politecnico di Torino. (2021). *Corso di laurea magistrale in ICT for Smart Societies (ICT per la società del futuro)*. <https://webthesis.biblio.polito.it/18007/>
16. Journal of Scientific Papers “Social Development and Security”. (2024). Vol. 14(2). ISSN 2522-9842.
17. Adnovum. (2025, March 14). *Modern cybersecurity strategies: Why traditional solutions fall short*. <https://www.adnovum.com/blog/modern-cybersecurity-strategies-why-traditional-solutions-fall-short>
18. Zscaler. (n.d.). *AI vs. traditional cybersecurity: Which is more effective?* <https://www.zscaler.com/zpedia/ai-vs-traditional-cybersecurity>
19. ResearchGate. (2019, September). *Big data analytics for cyber security*. https://www.researchgate.net/publication/335698795_Big_Data_Analytics_for_Cyber_Security
20. Scarfone, K., & Mell, P. (2007, February). *Guide to intrusion detection and prevention systems (IDPS)* (NIST Special Publication 800-94). National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs/Legacy/SP/nistspecialpublication800-94.pdf>
21. MITRE Corporation. (n.d.). *ATT&CK Matrix for Enterprise*. <https://attack.mitre.org/>
22. Zhurakovskiy, B., Averichev, I., & Shakhmatov, I. (2023, November 21). Using the latest methods of cluster analysis to identify similar profiles in leading social networks. *Information Technology and Implementation (Satellite) Conference Proceedings*. https://ceur-ws.org/Vol-3646/Paper_12.pdf
23. Ponochozny, P. (2024). Low-speed HTTP DDoS attack prevention model for end users. *Cybersecurity: Education, Science, Technique*, 2(26), 291–304. <https://doi.org/10.28925/2663-4023.2024.26.695>
24. Іванченко, Є., Роженко, А., & Берестяна, Т. (2025). Інноваційні підходи до підвищення рівня кібербезпеки корпоративних мереж при використанні хмарних технологій. *Кібербезпека: освіта, наука, техніка*, 4(28), 656–670. <https://doi.org/10.28925/2663-4023.2025.28.858>



Volodymyr Shulga

Doctor of Historical Sciences, Senior Researcher, Rector,
State University of Information and Communication Technologies, Kyiv, Ukraine
ORCID: 0000-0003-4356-7288
v.shulha@duikt.edu.ua

Yevheniia Ivanchenko

Doctor of Technical Sciences, Professor,
Director of the Educational and Scientific Institute of Cybersecurity and Information Protection,
State University of Information and Communication Technologies, Kyiv, Ukraine
ORCID: 0000-0003-3017-5752
evivanchenko@gmail.com

Tatyana Berestyana

Postgraduate student
State University of Information and Communication Technologies, Kyiv, Ukraine
ORCID: 0009-0007-1455-234X
t.berestiana@duikt.edu.ua

Oleksiy Shkurchenko

Postgraduate student
State University of Information and Communication Technologies, Kyiv, Ukraine
ORCID: 0009-0006-9534-144X
o.shkurchenko@stud.duikt.edu.ua

METHODS AND MODELS OF COUNTERING GROUP CYBER THREATS BASED ON ARTIFICIAL INTELLIGENCE

Abstract. The article addresses the pressing issue of countering group threats in the field of cybersecurity, which are characterized by a high level of organization, complex structures, and targeted execution. With the advancement of digital technologies and the growing dependency of businesses and government institutions on information systems, threats carried out by coordinated groups-such as APTs, botnets, and other cybercriminal organizations-pose an increasing danger. Traditional security approaches have proven ineffective against such attacks, as they often fail to consider the dynamic behavior of threats, their rapid evolution, and adaptability to defensive mechanisms. A method for countering group threats based on artificial intelligence technologies is proposed, specifically leveraging machine learning, deep learning, and big data processing techniques. The developed model architecture enables the detection of signs of coordinated malicious activity, analysis of attacker behavioral patterns, and timely response to potential threats. Special attention is given to the development of adaptive models capable of real-time self-learning and identifying atypical deviations from the normal functioning of an information system. The models are designed with key characteristics of group attacks in mind, including scenarios of distributed attacks, hidden communication between attackers, and the use of encryption and evasion techniques. Experimental studies on the effectiveness of the proposed approach were conducted using test datasets and realistic threat scenarios. The results demonstrated high detection accuracy, a reduction in false positives, and improved response times compared to traditional systems. The findings confirm the feasibility and effectiveness of using intelligent methods to counter complex group threats. The proposed method and models can be integrated into existing cybersecurity monitoring systems to enhance their resilience against modern attacks. The article also outlines prospects for further research, particularly in improving model training methods and scaling to large corporate networks.

Keywords: group cyber threats, artificial intelligence, machine learning, deep learning, information security, detection models, adaptive systems, cyberattacks, big data processing, cybersecurity.



REFERENCES

1. Andrieiev, O. O. (2021). *Methods for detecting and countering cyberattacks in corporate networks* : Monograph. Kyiv: National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute".
2. Lazarus Alliance. (2023, March 8). *What are advanced persistent threats (APT)?* <https://lazarusalliance.com/uk/what-are-advanced-persistent-threats-apt/>
3. Kozlov, D. S., & Tereshchenko, L. V. (2022). Intelligent information protection systems: Machine learning methods. *Information Security*, (3), 42–50.
4. Skopik, F., Settanni, G., & Fiedler, R. (2016). A problem shared is a problem halved: A survey on the dimensions of collective cyber defense through security information sharing. *Computers & Security*, 60, 154–176. <https://doi.org/10.1016/j.cose.2016.03.011>
5. Buczak, A. L., & Guven, E. (2016). A survey of data mining and machine learning methods for cyber security intrusion detection. *IEEE Communications Surveys & Tutorials*, 18(2), 1153–1176. <https://doi.org/10.1109/COMST.2015.2494502>
6. Zvieriev, I. V. (2022). Artificial intelligence as a tool in information security systems. *Modern Information Technologies*, (2), 28–34.
7. Sommer, R., & Paxson, V. (2010). Outside the closed world: On using machine learning for network intrusion detection. *IEEE Symposium on Security and Privacy*, 305–316. <https://doi.org/10.1109/SP.2010.25>
8. Sangkatsanee, P., Wattanapongsakorn, N., & Charnsripinyo, C. (2011). Practical real-time intrusion detection using machine learning approaches. *Computer Communications*, 34(18), 2227–2235. <https://doi.org/10.1016/j.comcom.2011.05.009>
9. Kolesnyk, S. P., Ovcharuk, V. M., Parkhomenko, A. P., et al. (2020). *Artificial intelligence algorithms for cybersecurity systems* : Textbook. Kharkiv: Kharkiv National University of Radio Electronics.
10. ESET. (n.d.). *Protection against botnets: How not to become part of a botnet network*. <https://www.eset.com/ua/support/information/entsyklopediya-zahroz/zakhyst-vid-botnetiv/>
11. Berman, D. S., Buczak, A. L., Chavis, J. S., & Corbett, C. L. (2019). A survey of deep learning methods for cyber security. *Information*, 10(4), 122. <https://doi.org/10.3390/info10040122>
12. Tang, T. A., Mhamdi, L., McLernon, D., Zaidi, S. A. R., & Ghogho, M. (2016). Deep learning approach for network intrusion detection in software defined networking. *2016 International Conference on Wireless Networks and Mobile Communications (WINCOM)*, 258–263. IEEE. <https://doi.org/10.1109/WINCOM.2016.7777224>
13. Kulish, V., & Pastukh, I. (2023). Modern methods for detecting botnet networks in corporate systems. *Information Protection*, (1), 63–70.
14. U.S. Department of Justice. (n.d.). *How FIN7 attacked and stole data: Sophisticated social engineering – phishing & calling*. <https://www.justice.gov/archives/opa/press-release/file/1084361/dl?inline=1>
15. Politecnico di Torino. (2021). *Master's degree program in ICT for Smart Societies (ICT for the Society of the Future)*. <https://webthesis.biblio.polito.it/18007/>
16. *Journal of Scientific Papers "Social Development and Security"*. (2024). Vol. 14(2). ISSN 2522-9842.
17. Adnovum. (2025, March 14). *Modern cybersecurity strategies: Why traditional solutions fall short*. <https://www.adnovum.com/blog/modern-cybersecurity-strategies-why-traditional-solutions-fall-short>
18. Zscaler. (n.d.). *AI vs. traditional cybersecurity: Which is more effective?* <https://www.zscaler.com/zpedia/ai-vs-traditional-cybersecurity>
19. ResearchGate. (2019, September). *Big data analytics for cyber security*. https://www.researchgate.net/publication/335698795_Big_Data_Analytics_for_Cyber_Security
20. Scarfone, K., & Mell, P. (2007, February). *Guide to intrusion detection and prevention systems (IDPS)* (NIST Special Publication 800-94). National Institute of Standards and Technology. <https://nvlpubs.nist.gov/nistpubs/Legacy/SP/nistspecialpublication800-94.pdf>
21. MITRE Corporation. (n.d.). *ATT&CK Matrix for Enterprise*. <https://attack.mitre.org/>
22. Zhurakovskiy, B., Averichev, I., & Shakhmatov, I. (2023, November 21). Using the latest methods of cluster analysis to identify similar profiles in leading social networks. *Information Technology and Implementation (Satellite) Conference Proceedings*. https://ceur-ws.org/Vol-3646/Paper_12.pdf
23. Ponochozny, P. (2024). Low-speed HTTP DDoS attack prevention model for end users. *Cybersecurity: Education, Science, Technique*, 2(26), 291–304. <https://doi.org/10.28925/2663-4023.2024.26.695>
24. Ivanchenko, Y., Rozhenko, A., & Berestyana, T. (2025). Innovative approaches to improving the level of cybersecurity of corporate networks using cloud technologies. *Cybersecurity: Education, Science, Technique*, 4(28), 656–670. <https://doi.org/10.28925/2663-4023.2025.28.858>

