

УДК 004.8+004.49

д-р філософії Власенко О. В. ORCID: 0000-0001-6671-870X (ВІТІ ім. Героїв Крут)

ІНТЕГРАЦІЯ ГЕНЕРАТИВНОГО ШТУЧНОГО ІНТЕЛЕКТУ В АРХІТЕКТУРУ SIEM-СИСТЕМИ

В умовах ескалації гібридних кіберзагроз та інтенсифікації кібернетичного протиборства, проблематика забезпечення стійкості і захищеності об'єктів критичної інфраструктури набуває особливої актуальності та стратегічного значення. Найвищого ступеня критичності це питання досягає у військово-оборонній сфері, де будь-яка компрометація інформаційних активів або порушення штатного функціонування інформаційно-комунікаційних систем здатні спричинити незворотні деструктивні наслідки для обороноздатності держави.

У роботі досліджено обмеження традиційних SIEM-систем, що базуються на детермінованих правилах і демонструють низьку ефективність у детектуванні невідомих загроз, високий рівень хибних спрацьовувань та проблеми масштабування. Здійснено формалізацію класичної архітектури SIEM-системи для ідентифікації її функціональних компонентів. На основі цього аналізу запропоновано удосконалену концептуальну модель SIEM-системи нового покоління, посилену можливостями генеративного штучного інтелекту. Ключовими елементами запропонованої моделі є модуль проактивного моделювання загроз, який використовує генеративний штучний інтелект для генерації гіпотез про невідомі вектори атак, та модуль генеративної верифікації і пріоритетизації, що застосовує глибокий семантичний аналіз для автоматичного валідування та контекстуального збагачення кіберінцидентів.

Запропонована архітектура забезпечує концептуальний перехід від реактивного моніторингу до проактивного, семантично-орієнтованого аналізу, підвищуючи ефективність виявлення кіберінцидентів, знижуючи навантаження на аналітиків та прискорюючи час реагування на критичні кіберінциденти, що є надзвичайно важливим для захисту стратегічних об'єктів. Перспективи подальших досліджень включають практичну реалізацію та валідацію запропонованих модулів, а також розробку методики оцінки їхньої обчислювальної ефективності та масштабованості.

Ключові слова: кібербезпека, штучний інтелект, SIEM-система, кіберінцидент, кіберзахист, інформаційно-комунікаційна система.

O. Vlasenko. Integration of generative artificial intelligence into the architecture of the SIEM system

In the context of escalating hybrid cyber threats and the intensification of cyber conflict, the problem of ensuring the resilience and security of critical infrastructure facilities acquires particular urgency and strategic importance. This issue reaches its highest degree of criticality in the military and defense sector, where any compromise of information assets or disruption of the standard functioning of information and communication systems can lead to irreversible destructive consequences for the state's defense capability.

This paper investigates the limitations of traditional SIEM systems, which are based on deterministic rules and demonstrate low efficiency in detecting unknown threats, a high rate of false positives, and scalability issues. A formalization of the classical SIEM architecture was carried out to identify its functional components. Based on this analysis, an improved conceptual model for a next-generation SIEM system is proposed, enhanced with the capabilities of Generative Artificial Intelligence (GenAI). The key elements of the proposed model are a proactive threat modeling module, which uses GenAI to generate hypotheses about unknown attack vectors, and a generative verification and prioritization module, which applies deep semantic analysis for the automated validation and contextual enrichment of cyber incidents.

The proposed architecture provides a conceptual shift from reactive monitoring to proactive, semantically-oriented analysis, significantly increasing threat detection effectiveness, reducing the workload on analysts, and accelerating response times to critical cyber incidents, which is extremely important for the protection of strategic assets. Future research directions include the practical implementation and validation of the proposed modules, as well as the development of methodologies for evaluating their computational efficiency and scalability.

Keywords: cybersecurity, artificial intelligence, SIEM system, cyber incident, cyber defense, information and communication system.

Актуальність та постановка завдання в загальному вигляді. Динамічний розвиток інформаційних технологій та їх глибока інтеграція в усі сфери життєдіяльності супроводжуються пропорційним зростанням кількості та рівня складності кібератак. Існуючі системи кіберзахисту, що здебільшого базуються на реактивних та сигнатурних методах, дедалі частіше виявляються неспроможними ефективно протидіяти адаптивним,

багатовекторним та раніше невідомим кібератакам. Це зумовлює необхідність у проведенні наукових досліджень, спрямованих на вдосконалення теоретико-методологічних засад кібербезпеки та розробку інноваційних методів і моделей захисту інформаційно-комунікаційних систем (ІКС).

Особливої важливості зазначена проблематика набуває у військовій сфері, що характеризується дедалі глибшою інтеграцією ІКС в усі ланки управління військами. Функціональне призначення таких систем полягає у забезпеченні процесів оперативного управління військами та озброєнням, веденні розвідки, підвищенні ситуаційної обізнаності та автоматизації циклів прийняття рішень. Зважаючи на те, що стабільне та захищене функціонування ІКС військового призначення безпосередньо корелює з ефективністю виконання бойових завдань та є критичним фактором забезпечення обороноздатності держави, завдання щодо їх надійного кіберзахисту стає одним із пріоритетних імперативів національної безпеки.

Для вирішення завдань комплексного захисту та забезпечення належного рівня кіберстійкості ІКС військового призначення створюються спеціалізовані організаційно-технічні структури — центри моніторингу та реагування на кіберінциденти (Security Operations Center, SOC). Такі центри виконують функції безперервного моніторингу стану безпеки підконтрольних інформаційних ресурсів, проактивного виявлення загроз, а також оперативного реагування на кіберінциденти. Організаційно-технологічним ядром, що забезпечує агрегацію, нормалізацію, кореляцію та інтелектуальний аналіз потоків подій безпеки з різноманітних джерел в рамках SOC, виступають системи класу Security Information and Event Management (SIEM-системи) [1].

SIEM-системи, виступаючи технологічним ядром SOC, дозволяють ефективно вирішувати задачі виявлення кіберінцидентів, що базуються на попередньо визначених правилах кореляції та відомих сигнатурах атак. Це значно підвищує рівень ситуаційної обізнаності та скорочує середній час реагування. Водночас, такий підхід, детермінований наявністю експертних знань для формування правил, демонструє іманентні обмеження при зіткненні з раніше невідомими загрозами, складними багатоетапними атаками та векторами, що використовують адаптивні тактики ухилення. Крім того, нездатність класичних систем до глибокого семантичного аналізу контексту подій призводить до генерації значного обсягу хибних спрацьовувань (False Positives), що перевантажує роботу аналітиків. Зазначені недоліки зумовлюють гостру науково-практичну потребу в удосконаленні аналітичних можливостей SIEM-систем шляхом їх інтеграції з проактивними модулями, здатними до когнітивної обробки даних та автономного навчання.

Сучасний етап розвитку інформаційних технологій характеризується прискореною еволюцією методологій штучного інтелекту (ШІ), що охоплюють широкий спектр: від класичного машинного навчання до складних глибоких нейронних мереж. Одним із найбільш динамічних елементів цього розвитку є генеративний ШІ, зокрема технології великих мовних моделей (LLM). Визначний потенціал генеративного ШІ у задачах аналізу, структурування та інтерпретації великомасштабних, неструктурованих масивів даних стимулює наукову спільноту до активного дослідження перспектив його інтеграції у прикладні галузі, серед яких одним із найактуальніших напрямків є посилення аналітичних функцій систем кібербезпеки.

Інтеграція модулів на основі генеративного ШІ в архітектуру SIEM-систем відкриває новий напрям розвитку інтелектуальних засобів кіберзахисту. Такий підхід дає змогу усунути традиційні обмеження SIEM-систем шляхом автоматичного формування нових правил кореляції, динамічного оновлення баз знань, зниження кількості хибних спрацьовувань та підвищення ефективності виявлення складних, багатовекторних атак. Використання генеративного ШІ у SIEM-системах може стати основою для формування самонавчальних систем кіберзахисту, здатних адаптуватися до змінного кіберсередовища й забезпечувати проактивне виявлення загроз у режимі реального часу.

Аналіз попередніх досліджень. Оскільки інтеграція генеративного ШІ в системи кіберзахисту є відносно новим напрямом, для визначення актуальних тенденцій та наукових підходів до його реалізації доцільним є проведення аналізу попередніх досліджень у цій галузі. Сучасна наукова спільнота демонструє підвищений інтерес до застосування генеративного ШІ у сфері кібербезпеки, зокрема в контексті автоматизації процесів виявлення кіберінцидентів, аналізу загроз та підвищення адаптивності захисних систем. Огляд основних публікацій дозволить виявити основні наукові підходи, існуючі моделі й інструменти, що базуються на генеративних методах, а також визначити наявні обмеження, які потребують подальшого вдосконалення в контексті інтеграції генеративного ШІ в архітектуру SIEM-систем.

У [2] наведено аналіз трансформаційного впливу генеративного ШІ, зокрема Великих Мовних Моделей (LLMs) та Генеративно-Змагальних Мереж (GANs), на сфери оцінки вразливостей та управління ризиками у кібербезпеці. Суть роботи полягає у розробці комплексної таксономії, яка класифікує подвійне використання генеративного ШІ (наступальне та оборонне) на усіх етапах життєвого циклу оцінки вразливостей та управління ризиками, включаючи виявлення вразливостей, моделювання загроз, оцінку ризиків та автоматизоване усунення.

Дослідження підкреслює, що генеративний ШІ здатний автоматизувати сортування логів, скорочувати час реагування та підвищувати точність виявлення складних кіберзагроз. На основі проведеного систематичного огляду зроблено висновки, що генеративний ШІ може допомогти у вирішенні багатьох традиційних проблем циклу оцінки вразливостей та управління ризиками. Однак, його повна інтеграція стримується ключовими викликами, серед яких вразливість самих моделей ШІ до кібератак (наприклад, data poisoning та prompt injection), проблеми зрозумілості (explainability) рішень у критично важливих доменах та необхідність забезпечення надійності (trustworthiness) згенерованого коду. Для подолання цих обмежень запропоновано пріоритетні напрямки майбутніх досліджень, що сфокусовані на розробці систем безперервного навчання, які адаптуються до загроз у реальному часі; еволюції процесу зрозумілості до інтерпретованих за задумом архітектур; створенні гібридних фреймворків "людина-ШІ" для нагляду та зворотного зв'язку, а також запровадженні стандартизованих рамок оцінювання ШІ-інструментів для підтвердження їхньої стійкості, надійності та відповідності етичним нормам.

У статті [3] автори досліджують ключову та трансформаційну роль генеративного ШІ у сучасній кібербезпеці, наголошуючи на його потенціалі посилити виявлення та реагування на кіберінциденти і удосконалити процеси прийняття рішень. Публікація ґрунтується на систематичному огляді літератури, переважна більшість якої опублікована у 2023 році, що підкреслює релевантність і сучасність даних, враховуючи новизну генеративного ШІ у дослідницькій арені.

Необхідність інтеграції генеративного ШІ обумовлена тим, що традиційні підходи часто поступаються темпам еволюції кіберзагроз та є реактивними, залишаючи організації вразливими до нових методів та моделей кібератак. Технології генеративного ШІ, використовуючи принципи глибокого навчання та нейронних мереж, позиціонується як проактивний елемент, що здатен автономно генерувати рішення та виявляти аномалії в мережевому трафіку, поведінці користувачів та системних логах у реальному часі.

Дослідження визначає конкретні технології застосування генеративного ШІ, що включають: підвищення ефективності аналізу кіберінцидентів та захисту від шкідливого програмного забезпечення шляхом вивчення патернів; посилення автентифікації через біометричні та поведінкові патерни користувачів; швидке та точне сканування програмних кодів для виявлення прихованих вразливостей; а також генерацію надійних паролів і симуляцію кібератак для виявлення слабких місць. Автори роблять висновки, що генеративний ШІ забезпечує постійне вдосконалення безпеки, оскільки інтелектуальні системи адаптуються до нових технік експлуатації та можуть пріоритезувати ризики під час

одночасних кібератак. Водночас, наголошується на значних обмеженнях, включно з неспроможністю ІІІ імітувати людське мислення, його залежністю від великих та точних наборів даних і вразливістю самих систем на базі генеративного ІІІ до зловмисного використання. Кінцевим висновком є те, що найбільш ефективна стратегія кібербезпеки полягає у об'єднанні традиційних методів із можливостями генеративного ІІІ.

Комплексний систематичний оглядом публікацій, присвяченим дослідженню ролі Генеративного ІІІ у підвищенні ефективності розвідки кіберзагроз (Cyber Threat Intelligence, СТІ) висвітлено у публікації [4]. Дослідження використовує методологію PRISMA та охоплює 330 релевантних публікацій за період з 2018 по 2024 роки з метою надання наскрізного (end-to-end) огляду застосувань, типів моделей, викликів та обмежень генеративного ІІІ у сфері СТІ. Автори систематизують генеративний ІІІ на дві основні категорії: Великі Мовні Моделі (LLMs) та Глибокі Генеративні Моделі (DGMs). У роботі детально проаналізовано інтеграцію генеративного ІІІ з різнорідними джерелами СТІ, такими як текстові дані, мережевий трафік, зразки шкідливого програмного забезпечення, дані Dark Web та платформи розвідки загроз. Публікація висвітлює ключові аспекти застосування генеративного ІІІ, які охоплюють виявлення шкідливого ПЗ, аналіз мережевих аномалій, детекцію фішингу, атрибуцію суб'єктів загроз та захист від соціальної інженерії.

Зроблено висновок, що технології генеративного ІІІ значно удосконалили сферу СТІ, пропонуючи адаптивні та проактивні рішення, що перевершують традиційні методи та моделі. Водночас, у статті виділено критичні виклики, включно з вразливістю до змагальних атак, питаннями конфіденційності даних (зокрема, через deepfakes та prompt injection), необхідністю інтерпретованості та прозорості моделей, а також проблемою адаптації до загроз, що постійно еволюціонують.

Публікація [5] присвячена аналізу потенціалу, реальних застосувань та обмежень генеративного ІІІ у сфері кібербезпеки. Публікація розглядає генеративний ІІІ як ефективний інструмент для протидії зростаючій частоті, витонченості та впливу кібератак, дозволяючи системам перейти від традиційного захисного підходу (countermeasures once an attack has been made) до наступального (проактивного), здатного передбачати ймовірну тактику зловмисників. Генеративний ІІІ надає можливість автоматизовано вирішувати загальні ситуації з загрозами, дозволяючи фахівцям концентруватися на більш критичних аспектах безпеки, що вимагають людського втручання. У статті систематизовано застосування генеративного ІІІ, включно з генеруванням безпечного коду, скануванням уразливостей із фільтрацією помилкових спрацювань на основі контексту, симуляцією реалістичних атак (Red Team simulations), створенням динамічних приманок (Honeypots) та генерацією запитів для полювання на загрози (threat-hunting queries).

Серед комерційних рішень, що використовують GAI, аналізуються Microsoft Security Copilot, Google Cloud Security AI Workbench та SentinelOne Purple AI. Водночас, автори публікації підкреслюють, що застосування генеративного ІІІ несе значні ризики, зокрема: схильність надавати невірні/неетичні результати, легкість експлуатації зловмисниками, проблеми інтерпретованості та зрозумілості ("black box" nature), контекстуальні обмеження та відсутність стандартизованих метрик для емпіричної оцінки систем на базі генеративного ІІІ. Зроблено висновок про необхідність розробки етичних настанов та надійних запобіжних заходів для запобігання зловмисному використанню генеративного ІІІ.

Аналіз публікацій, присвячених інтеграції генеративного ІІІ в процеси кіберзахисту, підтверджує високу актуальність та доцільність цього напрямку. Проте, переважна більшість робіт обмежується декларативним описом загальних перспектив та потенційних переваг, не пропонуючи конкретних прикладних рішень чи деталізованих моделей для впровадження. Зокрема, питання методологічних та архітектурних засад інтеграції генеративних компонентів із SIEM-системами для посилення їхніх аналітичних можливостей залишається практично нерозкритим. Ця недостатня дослідженість проблеми на практичному рівні формує чітку

наукову прогалину та підтверджує доцільність подальшої роботи над розробкою відповідних методів та моделей.

Метою статті є розробка удосконаленої моделі архітектури SIEM-системи з інтегрованими модулями на основі генеративного штучного інтелекту для підвищення ефективності виявлення кіберінцидентів в інформаційно-комунікаційних системах військового призначення.

Виклад основного матеріалу. SIEM-система виконує роль центральної аналітичної платформи у SOC, забезпечуючи консолідацію та інтелектуальну обробку потоків подій (event streams) з усього спектру контрольованих інформаційних ресурсів (мережеве обладнання, сервери, кінцеві точки, засоби захисту). Її фундаментальна цінність полягає у трансформації величезних масивів "сирих" даних, які самі по собі мають низьку інформативність, у верифіковані, пріоритезовані та збагачені контекстом кіберінциденти, що є прямими директивами для реагування аналітиками. Таким чином, SIEM-система реалізує повний цикл операційного управління подіями безпеки – від первинного збору до генерації сповіщення. Ключові функціональні модулі та процеси, що забезпечують цей цикл, деталізовано наведено в таблиці 1 [1; 6–8].

Таблиця 1

Декомпозиція функціональних компонентів SIEM-системи

Функціональний компонент	Деталізований опис призначення та механізмів
Збір та Агрегація даних	Прийом та консолідація потоків даних з гетерогенних джерел (кінцеві точки, сервери, мережеві пристрої, СЗІ). Використовуються різні протоколи (Syslog, SNMP, NetFlow) та механізми (агенти, API-інтеграція)
Парсинг та Нормалізація	Перетворення логів з різноманітних, пропріетарних форматів у єдину канонічну модель даних (Common Event Format). Цей етап є критично важливим, оскільки уможливорює подальший спільний аналіз та кореляцію подій з різних джерел
Кореляція подій	Основний аналітичний механізм класичних SIEM-систем. Виконує детермінований аналіз нормалізованих подій у реальному часі на основі попередньо заданої бази правил (correlation rules) для виявлення складних, розподілених у часі та просторі патернів атак
Збагачення даних	Процес автоматичного додавання додаткового контексту до подій або кіберінцидентів. Може включати дані геолокації, інформацію про користувача з Active Directory, репутаційні дані з Threat Intelligence каналів
Генерація сповіщень	Формування та пріоритезація сповіщення для аналітика SOC при спрацюванні кореляційного правила. Ефективність цього модуля визначається співвідношенням True Positives до False Positives
Довготривале зберігання та Індексация	Забезпечення захищеного та відмовостійкого довготривалого зберігання (часто з використанням WORM-технологій) як "сирих", так і нормалізованих даних для потреб ретроспективного аналізу (Threat Hunting), розслідувань (Forensics) та аудиту відповідності (Compliance)
Візуалізація та Звітність	Надання інструментів для оперативного моніторингу та формування аналітичної і статистичної звітності (reports) для технічних спеціалістів та керівництва, а також для підтвердження відповідності регуляторним вимогам

На основі ґрунтовного аналізу провідних наукових публікацій [8–11], присвячених архітектурі та функціонуванню SIEM-систем, було визначено ключові функціональних компонентів: збір та агрегація, парсинг та нормалізація, кореляція подій, збагачення даних, генерація сповіщень, довготривале зберігання та візуалізація/звітність.

Для формалізації взаємозв'язків між цими компонентами та опису системи на концептуальному рівні, запропоновано впорядковану модель (1):

$$S = \langle C_{coll}, P_{norm}, E_{corr}, D_{enrich}, A_{alert}, S_{store}, V_{report} \rangle, \quad (1)$$

де кожен елемент є підсистемою (або множиною функцій), що виконує специфічне завдання:

C_{coll} (Збір та Агрегація) – це вхідна функція системи. Вона приймає множину потоків сирих даних (D_{raw}) з множини джерел (I) і передає їх на наступний етап (2):

$$C_{coll} : I \rightarrow D_{raw}; \quad (2)$$

P_{norm} (Парсинг та Нормалізація) – функція, що перетворює хаотичні сирі дані (D_{raw}) у структурований, єдиний формат (таксономію) системи (D_{norm}), дозволяє нормалізувати дані (3):

$$P_{norm} : D_{raw} \rightarrow D_{norm}; \quad (3)$$

E_{corr} (Кореляція подій) – це множина правил та алгоритмів (R), які застосовуються до потоку нормалізованих подій (D_{norm}) для виявлення значущих патернів загроз (4)

$$(P_{threat}).E_{corr} : D_{norm} \times R \rightarrow P_{threat}; \quad (4)$$

D_{enrich} (Збагачення даних) – функція, що додає контекст (C_{tx}) до виявлених патернів або подій. Вона використовує зовнішні та внутрішні (списки активів, дані користувачів) джерела (5):

$$D_{enrich} : (P_{threat} \cup D_{norm}) \times C_{tx} \rightarrow P_{enriched}; \quad (5)$$

A_{alert} (Генерація сповіщень) – підсистема, що приймає збагачені патерни загроз $P_{enriched}$ і перетворює їх на конкретні, пріоритезовані кіберінциденти (алерти, A) для аналітиків (6):

$$A_{alert} : P_{enriched} \rightarrow A; \quad (6)$$

S_{store} (Зберігання та Індексация) – Підсистема, що відповідає за довготривале, захищене від втручання зберігання всіх даних (як D_{raw} , так і D_{norm} , і A) та їх ефективну індексацію для швидкого пошуку.

V_{report} (Візуалізація та Звітність) – інтерфейс представлення інформації. Це компонент, який надає інструменти (дашборди, звіти, інструменти пошуку) для взаємодії аналітика з кіберінцидентами (A) та збереженими даними (S_{store}).

Модель репрезентує архітектуру SIEM-системи як логічно послідовний механізм обробки даних про події безпеки, де кожен елемент є функціональною підсистемою, що перетворює дані, забезпечуючи перехід від логів до пріоритезованих кіберінцидентів та аналітичних звітів.

Таблиця 2

Аналіз іманентних обмежень класичних SIEM-систем

Категорія обмеження	Сутність проблеми	Негативний наслідок для SOC
Детермінована логіка аналізу	Ефективність виявлення інцидентів прямо детермінована повнотою та актуальністю бази кореляційних правил, що задаються експертами	"Сліпота" до невідомих загроз. Низька або нульова ефективність проти атак Zero-Day, нетипових векторів та складних багатоетапних загроз (APT), які не мають відомих сигнатур
Високий рівень хибних спрацювань	Відсутність глибокого семантичного аналізу та розуміння контексту призводить до генерації значної кількості сповіщень	Перевантаження аналітиків. Розсіювання уваги, збільшення часу на "ручну" верифікацію та тріаж, що критично підвищує

Категорія обмеження	Сутність проблеми	Негативний наслідок для SOC
	(алертів), що не є реальними інцидентами (False Positives)	середній час реагування (MTTR) на справжні загрози
Реактивний характер	Системи переважно реагують на події, що вже відбулися та відповідають заданим правилам, замість проактивного виявлення підготовки до атаки чи аномальної активності	Втрата критичного часу; виявлення інциденту вже після настання деструктивних наслідків або компрометації системи
Проблема масштабування аналізу	Експоненційне зростання обсягів даних (Big Data) у сучасних ІКС призводить до того, що аналіз на основі правил стає обчислювально складним та нездатним охопити всі нелінійні залежності	Пропуск складних інцидентів. Нездатність виявити слабкі, приховані сигнали (weak signals) або складні, розподілені атаки, що "губляться" у загальному шумі даних

Після деталізації функціональних компонентів SIEM-систем стає можливим проведення їх критичного аналізу в контексті протидії сучасним багатовекторним та адаптивним кіберзагрозам. Попри свою фундаментальну роль у SOC, традиційні SIEM-системи мають низку іманентних архітектурних та методологічних обмежень, які представлено у таблиці 2.

Аналіз актуальних викликів, притаманних традиційним архітектурам SIEM-систем, що охоплює детерміновану логіку аналізу, високий рівень хибних спрацьовувань (False Positives) та проблеми масштабування при роботі з великими даними, чітко вказує на потребу в парадигмальних змінах. Існуючі підходи, засновані на статичних кореляційних правилах, а також у лінійному аналізі, демонструють "сліпоту" до невідомих загроз, атак класу Zero-Day та комплексних багатоетапних загроз (APT), які не мають відомих сигнатур.

Крім того, відсутність глибокого семантичного розуміння контексту подій призводить до перевантаження аналітиків SOC нерелевантними оповіщеннями, істотно збільшуючи середній час реагування. З огляду на ці обмеження, запропоновано інтеграцію модулів, що використовують можливості генеративного ШІ, для трансформації ключових функціональних компонентів SIEM-системи [4; 5].

Зокрема, для подолання детермінованості та реактивності розроблено *модуль проактивного моделювання загроз*, який інтегрується у компонент кореляції подій (E_{corr}), перетворюючи його на гібридний кореляційний механізм (E_{corr}^*).

Модуль проактивного моделювання загроз являє собою компонент, що застосовує генеративні моделі для емуляції когнітивних процесів та тактик потенційного зловмисника (adversarial thinking). На відміну від традиційних систем, що оперують на основі спрацювання статичних правил, даний модуль виконує наступні функції:

генерація гіпотетичних векторів атак. Модуль здійснює синтез та валідацію гіпотез щодо можливих векторів компрометації. Цей процес базується на аналізі "слабких сигналів" (weak signals) – даних низької інтенсивності або непрямих індикаторів, які зазвичай фільтруються або ігноруються стандартними системами моніторингу;

симуляційне моделювання загроз. Компонент виконує симуляції за сценаріями (scenario-based simulation) для моделювання невідомих та раніше не ідентифікованих загроз, включно з атаками нульового дня. Це дозволяє створювати синтетичні набори даних, що репрезентують патерни новітніх загроз, які ще не внесені до баз сигнатур;

метод детекції аномалій. Виявлення аномальної активності відбувається не на основі порівняння з встановленою базовою лінією "нормальної" поведінки (як у системах UEBA), а шляхом кореляції спостережуваних подій зі згенерованими модулем проактивними патернами атак.

Застосування *модуля проактивного моделювання загроз* спрямоване на подолання двох ключових обмежень існуючих підходів до моніторингу безпеки:

подолання детермінованої логіки аналізу. Модуль долає обмеження систем, заснованих на правилах (rule-based systems), які демонструють низьку ефективність при зіткненні з невідомими загрозами. Він забезпечує перехід від статичного детермінізму до динамічної генерації гіпотез, що є критично важливим для ідентифікації та протидії складним багатоетапним атакам та експлойтам нульового дня;

трансформація реактивного підходу. Компонент змінює парадигму кібербезпеки, переводячи її з режиму реактивного реагування на вже dokonані кіберінциденти (reactive incident response) до проактивного прогнозування, моделювання та превентивного виявлення загроз (proactive threat modeling and forecasting).

Також для мінімізації хибних спрацьовувань та оптимізації процесу аналізу великих обсягів даних введено новий компонент – *модуль генеративної верифікації та пріоритетизації* (M_{gvp}). Цей модуль, розташований після збагачення даних (D_{enrich}) та перед генерацією сповіщень (A_{alert}), використовує генеративний ШІ для глибокого семантичного аналізу, автоматичної верифікації кіберінцидентів та створення контекстуально збагачених резюме, тим самим знижуючи навантаження на аналітиків та прискорюючи реагування.

Модуль генеративної верифікації та пріоритетизації функціонує як автоматизований аналітичний компонент, що еквівалентний функціям аналітика SOC першого рівня (Tier 1 Analyst), реалізований на базі генеративного ШІ. Усі потенційні кіберінциденти, згенеровані гібридним кореляційним механізмом (E_{corr}^*), першочергово надходять до цього модуля для автоматизованої обробки.

Основні функції *модуля генеративної верифікації та пріоритетизації*:

Семантичний аналіз подій. Модуль виконує поглиблений семантичний розбір сутностей та контексту подій безпеки, переходячи від простого зіставлення полів (field matching) до розуміння логіки та намірів, що стоять за послідовністю дій.

Автоматизована верифікація. Модуль миттєво ініціює процес розслідування (automated investigation) для кожного згенерованого звіту про кіберінцидент. Це включає автоматичний збір додаткових релевантних даних із сховища (S_{store}), збагачення подій контекстуальною інформацією та винесення вердикту щодо класифікації кіберінциденту.

Генерація природномовного резюме. Замість надання аналітику необробленого списку логів, модуль синтезує стислий, когнітивно засвоювальний опис кіберінциденту природною мовою. Це резюме містить пояснення виявленої небезпеки, її потенційного впливу та перелік вже виконаних кроків верифікації.

Виявлення прихованих зв'язків. Модуль здатен ідентифікувати латентні кореляції та недосконалі сигнали, розподілені у великих масивах даних, які не були виявлені на початковому етапі кореляції через їхню низьку амплітуду або неявний характер.

Впровадження *модуля генеративної верифікації та пріоритетизації* безпосередньо вирішує наступні критичні проблеми:

Мітигація високого рівня хибних спрацьовувань. Модуль виконує превентивну фільтрацію та відсіювання некоректних повідомлень до їх ескалації на рівень людського аналізу. Це призводить до значного зниження когнітивного навантаження на персонал SOC та скорочення середнього часу реагування (Mean Time to Respond).

Розв'язання проблеми масштабування аналізу. Забезпечує здатність до осмислення (comprehension) та аналізу великої кількості необроблених подій у режимі реального часу. Це дозволяє виявляти складні, приховані патерни атак, які в традиційних системах "губляться у шумі" (lost in the noise) через обмежені можливості обробки та кореляції даних.

На підставі проведеного аналізу фундаментальних обмежень, притаманних традиційним SIEM-системам, зокрема їхньої детермінованої логіки, реактивного характеру та нездатності ефективно обробляти великі обсяги даних, у даній роботі запропоновано удосконалену архітектурну модель SIEM-системи (7):

$$S_{GenAI} = \langle C_{coll}, P_{norm}, E_{corr}^*, D_{enrich}, M_{gvp}, A_{alert}, S_{store}, V_{report} \rangle, \quad (7)$$

де E_{corr}^* (Гібридний кореляційний механізм) – це розширення E_{corr} . $E_{corr}^* = E_{corr} \cup M_{pmv}$, включення модулю проактивного моделювання загроз.

M_{gvp} (Модуль генеративної верифікації та пріоритетизації) – це новий, вбудований компонент, що виконує функції семантичного аналізу, верифікації та пріоритетизація перед тим, як кіберінцидент потрапить до A_{alert} .

Удосконалена модель S_{GenAI} інтегрує спеціалізовані модулі на базі генеративного штучного інтелекту. Впровадження модуля проактивного моделювання загроз, інтегрованого у гібридний кореляційний механізм, та модуля генеративної верифікації і пріоритетизації як нового компоненту архітектури, безпосередньо спрямоване на подолання вищезазначених недоліків.

Таким чином, запропонована модель забезпечує удосконалений механізм від пасивного моніторингу до проактивного, семантично-обґрунтованого аналізу загроз, що є критично важливим для виявлення невідомих кібератак та зниження навантаження на аналітиків SOC.

Висновки. У статті було формалізовано архітектуру традиційної SIEM-системи, що дозволило чітко ідентифікувати її фундаментальні обмеження, такі як детермінованість аналізу, реактивний характер та високий рівень хибних спрацьовувань.

Для подолання цих недоліків, що є критичними в умовах сучасних кіберзагроз, було розроблено та запропоновано удосконалену модель архітектури. Ключовою інновацією запропонованої архітектури є інтеграція двох спеціалізованих модулів на базі генеративного штучного інтелекту: Модуля проактивного моделювання загроз, що трансформує компонент кореляції у гібридний прогностичний механізм, та модуля генеративної верифікації та пріоритетизації, який впроваджує семантичний аналіз для автоматичної верифікації кіберінцидентів.

Таким чином, розроблена модель забезпечує парадигмальний перехід від реактивного моніторингу до проактивного, контекстуально-обґрунтованого аналізу загроз, що підвищує ефективність виявлення кіберінцидентів та знижує навантаження на аналітиків SOC.

Ключовим напрямом подальших досліджень є практична реалізація запропонованих модулів на базі генеративного ШІ. Це завдання передбачає проведення поглибленого дослідження з метою вибору та адаптації найбільш ефективних архітектур генеративних моделей для вирішення специфічних завдань проактивного моделювання загроз та семантичної верифікації кіберінцидентів у режимі реального часу.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Miloslavskaya N. Analysis of SIEM Systems and Their Usage in Security Operations and Security Intelligence Centers. *Advances in Intelligent Systems and Computing*. P. 282–288. URL: https://doi.org/10.1007/978-3-319-63940-6_40.
2. Seyedeh Leili Mirtaheri, Narges Movahed, Reza Shahbazian, Valerio Pascucci, Andrea Pugliese. Cybersecurity in the age of generative AI: A systematic taxonomy of AI-powered vulnerability assessment and risk management. *Future Generation Computer Systems*. 2025. Vol. 175. URL: <https://doi.org/10.1016/j.future.2025.108107>.
3. Curran K., Curran E., Killen J., Duffy C. The role of generative AI in cyber security. *Metaverse*. 2024. No. 5 (2): 2796. URL: <https://doi.org/10.54517/m2796>.
4. Balasubramanian P., Liyana S., Sankaran H. et al. Generative AI for cyber threat intelligence: applications, challenges, and analysis of real-world case studies. *Artif Intell Rev* 58. 2025. URL: <https://doi.org/10.1007/s10462-025-11338-z>.
5. Sai S., Yashvardhan U., Chamola V., Sikdar B. Generative AI for cyber security: analyzing the potential of ChatGPT, DALL-E and other models for enhancing the security space. *IEEE Access PP* (99). 2024. P. 53499–53516. URL: <https://doi.org/10.1109/ACCESS.2024.3385107>.

6. T. Laue, C. Kleiner, K.-O. Detken and T. Klecker. A SIEM Architecture for Multidimensional Anomaly Detection. 2021. *11th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications (IDAACS)*, Cracow, Poland. 2021. Pp. 136–142. URL: <https://doi.org/10.1109/IDAACS53288.2021.9660903>.
7. Sheeraz M., Paracha M. A., Haque M. U., Durad M. H., Mohsin S. M., Band S. S., Mosavi A. Effective security monitoring using efficient SIEM architecture. *Hum.-Centric Comput. Inf. Sci*, 13, P. 1–18. URL: <https://doi.org/10.22967/HCIS.2023.13.02>.
8. González-Granadillo, Gustavo, Susana González-Zarzosa, and Rodrigo Diaz. Security Information and Event Management (SIEM): Analysis, Trends, and Usage in Critical Infrastructures. *Sensors* 2021. No. 14: 4759. URL: <https://doi.org/10.3390/s21144759>.
9. І. Субач, В. Кубрак, А. Микитюк. Архітектура та функціональна модель перспективної проактивної інтелектуальної системи SIEM-системи для кіберзахисту об'єктів критичної інфраструктури. *Information Technology and Security*. 2019. № 7 (2). Pp. 208–215. URL: <https://doi.org/10.20535/2411-1031.2019.7.2.190570>.
10. Субач І., Власенко О. Архітектура інтелектуальної SIEM-системи для виявлення кіберінцидентів у базах даних інформаційно-телекомунікаційних системах військового призначення. *Збірник наукових праць ВІТІ*. 2023. № 4. С. 82–92. URL: <https://doi.org/10.58254/viti.4.2023.07.82>.
11. SIEM Architecture Explained: Components, Design & Best Practices. *PuppyGraph Cybersecurity*. URL: <https://www.puppygraph.com/blog/siem-architecture>.