

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ**

**НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ
ІНФОРМАЦІЇ
КАФЕДРА УПРАВЛІННЯ КІБЕРБЕЗПЕКОЮ ТА ЗАХИСТОМ
ІНФОРМАЦІЇ**

КВАЛІФІКАЦІЙНА РОБОТА

**на тему: “МЕТОДИКА ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ У
ЗАБЕЗПЕЧЕННІ КІБЕРБЕЗПЕКИ ОБ’ЄКТІВ КРИТИЧНОЇ
ІНФРАСТРУКТУРИ”**

на здобуття освітнього ступеня бакалавра
зі спеціальності 125 Кібербезпека
освітньої програми Управління інформаційною та кібернетичною безпекою

*Кваліфікаційна робота містить результати власних досліджень.
Використання ідей, результатів і текстів інших авторів мають посилання на
відповідне джерело*

(підпис) Софія ДУДІЙ
Ім'я, ПРІЗВИЩЕ здобувача

Виконала: здобувачка вищої освіти гр. УБД-42

Софія ДУДІЙ
Ім'я, ПРІЗВИЩЕ

Керівник:
к. держ. упр.,
доцент

Тетяна МУЖАНОВА
Ім'я, ПРІЗВИЩЕ

Рецензент:

Ім'я, ПРІЗВИЩЕ

Київ 2025

ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ

Навчально-науковий інститут кібербезпеки та захисту інформації

Кафедра управління кібербезпекою та захистом інформації

Ступінь вищої освіти бакалавр

Спеціальність 125 Кібербезпека

Освітня програма Управління інформаційною та кібернетичною безпекою

ЗАТВЕРДЖУЮ

Завідувач кафедри УКБЗІ

_____ Світлана ЛЕГОМІНОВА
“ ” _____ 2025 р.

ЗАВДАННЯ НА КВАЛІФІКАЦІЙНУ РОБОТУ

Дудій Софії Андріївні

(прізвище, ім'я, по батькові здобувача)

1. Тема кваліфікаційної роботи “Методика застосування штучного інтелекту у забезпеченні кібербезпеки об’єктів критичної інфраструктури”,
керівник кваліфікаційної роботи МУЖАНОВА Тетяна, к. держ. упр., доцент,
(ПРИЗВИЩЕ, Ім'я, науковий ступінь, вчене звання)

затверджені наказом Державного університету інформаційно-комунікаційних технологій від “24” лютого 2025 р. №56.

2. Строк подання кваліфікаційної роботи “20” травня 2025р.
3. Вихідні дані до кваліфікаційної роботи: *кібербезпека об’єктів критичної інфраструктури (ОКІ), штучний інтелект (ШІ) у забезпеченні кібербезпеки ОКІ, методика застосування ШІ у забезпеченні кібербезпеки ОКІ.*
4. Перелік питань, які мають бути розроблені:
 - 4.1. Дослідити теоретичні основи використання штучного інтелекту в галузі кібербезпеки.
 - 4.2. Проаналізувати особливості й методику застосування ШІ в забезпеченні кібербезпеки об’єктів критичної інфраструктури.
 - 4.3. Вивчити засади практичного використання ШІ для посилення кібербезпеки підприємств та ОКІ.
5. Перелік ілюстративного матеріалу: *презентація PowerPoint*
6. Дата видачі завдання “11” березня 2025 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Етапи кваліфікаційної роботи	Термін виконання етапів роботи	Примітка
1.	Визначення об'єкту, предмету, мети та завдань дослідження.	18.03.2025	
2.	Збір та аналіз літератури.	29.03.2025	
3.	Дослідження теоретичних основ використання штучного інтелекту в галузі кібербезпеки.	08.04.2025	
4.	Аналіз особливостей і визначення методики застосування ШІ в забезпеченні кібербезпеки ОКІ.	22.04.2025	
5.	Встановлення засад практичного використання ШІ для посилення кібербезпеки підприємств та ОКІ.	08.05.2025	
6.	Формулювання висновків за результатами проведеного дослідження.	20.05.2025	
7.	Оформлення роботи.	22.05.2025	
8.	Оформлення презентації.	03.06.2025	
9.	Отримання рецензії на роботу.	03.06.2025	
10.	Захист в ЕК.	__ .06.2025	

Здобувачка вищої освіти

(підпис)

Софія ДУДІЙ

(Ім'я, ПРІЗВИЩЕ)

Керівник
кваліфікаційної роботи

(підпис)

Тетяна МУЖАНОВА

(Ім'я, ПРІЗВИЩЕ)

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ КІБЕРБЕЗПЕКИ ТА ЗАХИСТУ
ІНФОРМАЦІЇ**

**ПОДАННЯ
ГОЛОВІ ЕКЗАМЕНАЦІЙНОЇ КОМІСІЇ
ЩОДО ЗАХИСТУ КВАЛІФІКАЦІЙНОЇ РОБОТИ
на здобуття освітнього ступеня бакалавра**

Направляється здобувачка Дудій С.А. до захисту кваліфікаційної роботи
(прізвище та ініціали)

за спеціальністю 125 Кібербезпека
(код, найменування спеціальності)

освітньої програми Управління інформаційною та кібернетичною безпекою
(назва)

на тему: “Методика застосування штучного інтелекту у забезпеченні кібербезпеки об’єктів критичної інфраструктури”

Кваліфікаційна робота і рецензія додаються.

Директор ННІКБЗІ _____
(підпис)

Свєнєнїя ІВАНЧЕНКО
(Ім'я, ПРІЗВИЩЕ)

Висновок керівника кваліфікаційної роботи

Здобувачка ДУДІЙ Софія у кваліфікаційній роботі дослідила теоретичні основи використання штучного інтелекту в галузі кібербезпеки, проаналізувала особливості й методику застосування ШІ в забезпеченні кібербезпеки об’єктів критичної інфраструктури, вивчила засади практичного використання ШІ для посилення кібербезпеки підприємств та ОКІ.

ДУДІЙ Софія показала розуміння проблеми дослідження та бачення основних теоретичних і практичних напрямів її вирішення, довела володіння методами наукового дослідження, проявила себе як організований, відповідальний виконавець. Результати дослідження апробовані на Всеукраїнській науково-практичній конференції “Стратегії кіберстійкості: управління ризиками та безперервність бізнесу” 27 лютого 2025 року.

Все це дозволяє оцінити кваліфікаційну роботу здобувачки ДУДІЙ Софії на оцінку “відмінно” та присвоїти їй кваліфікацію бакалавра з кібербезпеки за освітньою програмою Управління інформаційною та кібернетичною безпекою.

Керівник кваліфікаційної роботи _____
(підпис)

Тетяна МУЖАНОВА
(Ім'я, ПРІЗВИЩЕ)

“___” _____ 2025 року

Висновок кафедри про кваліфікаційну роботу

Кваліфікаційна робота розглянута. Здобувачка Дудій С.А. допускається до захисту даної роботи в Екзаменаційній комісії.

Завідувач кафедри
управління кібербезпекою та
захистом інформації

(підпис)

Світлана ЛЕГОМІНОВА
(Ім'я, ПРІЗВИЩЕ)

ВІДГУК РЕЦЕНЗЕНТА

на кваліфікаційну бакалаврську роботу

здобувачки вищої освіти ДУДІЙ Софії

на тему “Методика застосування штучного інтелекту у забезпеченні кібербезпеки об’єктів критичної інфраструктури”

Актуальність. У сучасних умовах кіберзагрози для об’єктів критичної інфраструктури стають дедалі масштабнішими та складнішими. Традиційні засоби захисту не завжди здатні оперативно виявляти новітні атаки, тому виникає потреба в інноваційних підходах до забезпечення кібербезпеки. Технології ШІ дають змогу автоматизувати виявлення аномалій, передбачати поведінку зловмисників та підвищувати стійкість інформаційних систем. Впровадження ШІ забезпечує високу швидкість обробки даних і масштабованість, здатність до самонавчання й адаптивність до нових загроз, зменшення навантаження на фахівців з кібербезпеки.

Враховуючи ці аспекти, дослідження методики застосування штучного інтелекту у забезпеченні кібербезпеки об’єктів критичної інфраструктури є актуальним науковим завданням.

Позитивні сторони.

1. У роботі досліджено теоретичні та практичні засади використання штучного інтелекту в забезпеченні кібербезпеки підприємств та об’єктів критичної інфраструктури, проаналізовано особливості й методику застосування ШІ для посилення кіберзахисту ОКІ.

2. Кваліфікаційна робота оформлена відповідно до вимог. Виклад матеріалу здійснено відповідно до плану, зроблено логічні висновки. Ключові положення роботи представлено у вигляді рисунків.

3. Здобувачка опрацювала достатню джерельну базу: понад 50 публікацій, в тому числі англomовні статті й нормативні документи Рівненської АЕС.

4. Робота має практичну спрямованість, містить аналіз особливостей і методики застосування ШІ для посилення кіберзахисту ОКІ на прикладі РАЕС.

Недоліки.

Доцільно було б глибше проаналізувати управлінські аспекти застосування ШІ для посилення кіберзахисту об’єктів критичної інфраструктури.

Однак, вищезгадані зауваження не впливають на загальну позитивну оцінку кваліфікаційної роботи.

Висновок: Кваліфікаційна робота виконана на високому науково-методичному рівні і заслуговує оцінки “відмінно”, а здобувачка ДУДІЙ Софія заслуговує присвоєння кваліфікації бакалавра з кібербезпеки за освітньою програмою Управління інформаційною та кібернетичною безпекою.

Рецензент:

підпис

Ім’я, ПРИЗВИЩЕ

РЕФЕРАТ

Кваліфікаційна робота присвячена дослідженню методики застосування штучного інтелекту у забезпеченні кібербезпеки об'єктів критичної інфраструктури. Робота складається зі вступу, трьох розділів, що містять 17 рисунків і 4 таблиці, висновків і списку використаних джерел із 56 найменувань. Загальний обсяг роботи становить 82 аркуша, з яких 5 аркушів займає список використаних джерел.

Метою роботи є дослідження методики застосування штучного інтелекту у забезпеченні кібербезпеки об'єктів критичної інфраструктури.

Об'єктом дослідження є забезпечення кібербезпеки об'єктів критичної інфраструктури.

Предмет дослідження – методика застосування штучного інтелекту у забезпеченні кібербезпеки об'єктів критичної інфраструктури.

Методи дослідження. Для вирішення поставленого наукового завдання використано методи аналізу та синтезу, порівняння, системного підходу, прогнозування, а також аналізу нормативно-правової бази.

Як результат у роботі досліджено теоретичні основи використання штучного інтелекту в галузі кібербезпеки, проаналізовано особливості й методику застосування ШІ в забезпеченні кібербезпеки об'єктів критичної інфраструктури, вивчено засади практичного використання ШІ для посилення кібербезпеки підприємств та ОКІ.

Галузь застосування. Описані підходи можуть бути використані для підвищення рівня кібербезпеки об'єктів критичної інфраструктури України, зокрема в галузі енергетики, шляхом впровадження методів штучного інтелекту.

Ключові слова: КІБЕРБЕЗПЕКА ОБ'ЄКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ (ОКІ), ШТУЧНИЙ ІНТЕЛЕКТ (ШІ) У ЗАБЕЗПЕЧЕННІ КІБЕРБЕЗПЕКИ ОКІ, МЕТОДИКА ЗАСТОСУВАННЯ ШІ У ЗАБЕЗПЕЧЕННІ КІБЕРБЕЗПЕКИ ОКІ.

ABSTRACT

The qualification work is devoted to the study of the methodology of applying artificial intelligence in the cybersecurity of critical infrastructure. The work consists of an introduction, three chapters containing 17 figures and 4 tables, conclusions and a list of references which consists of 56 items. The total volume of the work is 82 pages, of which 5 pages are occupied by the list of references.

The purpose of the study is to investigate the methodology of applying artificial intelligence in cybersecurity of critical infrastructure.

The object of the study is cybersecurity of critical infrastructure.

The subject of the study is the methodology of applying artificial intelligence in cybersecurity of critical infrastructure.

Research methods. In order to solve the mentioned higher scientific task, the methods of analysis and synthesis, comparison, systematic approach, forecasting as well as methods of legislative analysis were used.

As a result, the study explores the theoretical foundations of the AI application in the field of cybersecurity, analyzes the features and methodology of using AI in cybersecurity of critical infrastructure, and studies the principles of practical use of AI to enhance the cybersecurity of enterprises and critical infrastructure.

Field of application. The described approaches can be used to increase the level of cybersecurity of critical infrastructure in Ukraine, in particular in the energy sector, by implementing artificial intelligence methods.

Keywords: CYBERSECURITY OF CRITICAL INFRASTRUCTURE (CI), ARTIFICIAL INTELLIGENCE (AI) IN CYBERSECURITY OF CI, METHODOLOGY OF AI APPLICATION IN CYBERSECURITY OF CI.

ЗМІСТ

ВСТУП.....	9
РОЗДІЛ 1 ТЕОРЕТИЧНІ ОСНОВИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ГАЛУЗІ КІБЕРБЕЗПЕКИ.....	11
1.1 Основні поняття кібербезпеки та штучного інтелекту	11
1.2 Методи та алгоритми штучного інтелекту у виявленні кіберзагроз	15
1.3 Використання машинного навчання для аналізу поведінки користувачів..	20
1.4 Глибоке навчання й аналіз великих даних у кібербезпеці	24
1.5 Вплив технологій ШІ на традиційні системи кіберзахисту	28
Висновки до розділу 1.....	31
РОЗДІЛ 2 ОСОБЛИВОСТІ Й МЕТОДИКА ЗАСТОСУВАННЯ ШІ В ЗАБЕЗПЕЧЕННІ КІБЕРБЕЗПЕКИ ОБ’ЄКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ.....	33
2.1 Політика забезпечення кібербезпеки критичної інфраструктури в Україні	33
2.2 Основні кіберзагрози для об’єктів критичної інфраструктури	36
2.3 Особливості забезпечення кібербезпеки атомних електростанцій в Україні	39
2.4 Аналіз програм кіберзахисту на прикладі Рівненської АЕС.....	42
2.5 Методика використання штучного інтелекту для кіберзахисту ОКІ	46
Висновки до розділу 2.....	50
РОЗДІЛ 3 ПРАКТИЧНЕ ВИКОРИСТАННЯ ШІ ДЛЯ ПОСИЛЕННЯ КІБЕРБЕЗПЕКИ ПІДПРИЄМСТВ ТА ОКІ.....	52
3.1 Сучасні інструменти кіберзахисту на основі ШІ	52
3.2 Автоматизація аналізу загроз за допомогою ШІ.....	56
3.3 Використання ШІ для виявлення аномалій у мережевому трафіку	59
3.4 Захист критичних систем: прогнозування атак і превентивні заходи	63
3.5 Виклики та ризики використання ШІ у кібербезпеці	70
Висновки до розділу 3.....	74
ВИСНОВКИ	76
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	78

ВСТУП

Актуальність теми. У сучасних умовах цифрової трансформації об'єкти критичної інфраструктури (ОКІ) стають дедалі більш вразливими до кіберзагроз. Саме тому впровадження інтелектуальних систем захисту, здатних оперативно виявляти аномальну активність, передбачати потенційні атаки та мінімізувати наслідки інцидентів набуває особливого значення. Штучний інтелект (ШІ), зокрема технології машинного навчання, аналізу поведінки користувачів (UEBA) й автоматизованого реагування, вже демонструє високу ефективність у системах кіберзахисту. Враховуючи важливість енергетичної безпеки, актуальним є дослідження можливостей застосування таких підходів для підвищення кіберстійкості атомних електростанцій, зокрема на прикладі Рівненської АЕС.

З огляду на зазначене, дослідження методики застосування штучного інтелекту у забезпеченні кібербезпеки об'єктів критичної інфраструктури є актуальним науковим завданням.

Мета роботи полягає у дослідженні методики застосування штучного інтелекту у забезпеченні кібербезпеки об'єктів критичної інфраструктури.

Об'єкт дослідження – забезпечення кібербезпеки об'єктів критичної інфраструктури.

Предмет дослідження – методика застосування штучного інтелекту у забезпеченні кібербезпеки об'єктів критичної інфраструктури.

Для досягнення цієї мети в роботі необхідно виконати наступні **завдання**:

1. Дослідити теоретичні основи використання штучного інтелекту в галузі кібербезпеки.
2. Проаналізувати особливості й методику застосування ШІ в забезпеченні кібербезпеки об'єктів критичної інфраструктури.
3. Вивчити засади практичного використання ШІ для посилення кібербезпеки підприємств та ОКІ.

Методи дослідження. Для вирішення означеного вище наукового

завдання в роботі використано методи аналізу та синтезу, порівняння, елементи системного підходу, а також методи аналізу нормативно-правової бази і прогнозування.

Практичне значення одержаних результатів. Проаналізовані підходи до застосування методів штучного інтелекту для підвищення кібербезпеки можуть бути використані як орієнтир для впровадження сучасних рішень у захист об'єктів критичної інфраструктури. Результати дослідження можуть стати основою для удосконалення механізмів управління доступом, виявлення аномальної активності та побудови ефективних систем моніторингу кіберзагроз у високоризикових секторах, зокрема на підприємствах енергетики.

Апробація результатів кваліфікаційної роботи відбулася на Всеукраїнській науково-практичній конференції “Стратегії кіберстійкості: управління ризиками та безперервність бізнесу” 27 лютого 2025 року.

РОЗДІЛ 1 ТЕОРЕТИЧНІ ОСНОВИ ВИКОРИСТАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В ГАЛУЗІ КІБЕРБЕЗПЕКИ

1.1 Основні поняття кібербезпеки та штучного інтелекту

Під інформаційною безпекою розуміється комплекс заходів, спрямованих на захист інформації та її інфраструктури від випадкових або навмисних загроз природного або техногенного походження. Основною метою інформаційної безпеки є запобігання втраті, модифікації чи несанкціонованому використанню даних, що можуть завдати значної шкоди власникам інформації та користувачам [1].

Кібербезпека є важливою складовою інформаційної безпеки та зосереджується на захисті інформаційних систем, мереж і цифрових даних від кібератак. Вона охоплює комплекс технічних, організаційних та правових заходів, спрямованих на запобігання несанкціонованому доступу, витоку інформації, зміні або знищенню даних. Рівень загроз кібербезпеки постійно зростає через поширення хмарних технологій, Інтернету речей (IoT) та дистанційної роботи, що робить кіберзахист однією з ключових пріоритетів держав та бізнесу [2].

Сучасні концепції кібербезпеки базуються на так званій CIA-тріаді, яка включає (Рис. 1.1):

- *Конфіденційність* – обмеження доступу до інформації лише для авторизованих осіб;
- *Цілісність* – забезпечення точності та достовірності даних шляхом запобігання їх несанкціонованій зміні;
- *Доступність* – гарантія того, що легітимні користувачі можуть отримати доступ до даних у потрібний час [3].



Рис. 1.1. Основні принципи інформаційної безпеки.

Окрім цього, у сучасній кібербезпеці широко застосовуються принципи нульової довіри (Zero Trust), які передбачають, що жоден користувач чи пристрій не повинен мати автоматичної довіри, і кожна взаємодія повинна бути перевірена.

Ця концепція є основою сучасної інформаційної безпеки, оскільки традиційні периметрові підходи, що базуються на розмежуванні внутрішньої та зовнішньої мережі, більше не забезпечують достатнього рівня захисту. Zero Trust впроваджує принцип "ніколи не довіряй, завжди перевіряй", що означає, що кожен запит на доступ до системи чи мережі підлягає ретельному аналізу та автентифікації, незалежно від того, де знаходиться користувач – всередині чи поза межами корпоративної інфраструктури.

Основні принципи Zero Trust включають:

- Ідентифікацію та автентифікацію користувачів – всі запити доступу потребують багатофакторної автентифікації (MFA) або біометричних методів підтвердження особи.
- Контроль доступу за мінімальними привілеями (Least Privilege Access) – користувачі та пристрої отримують лише ті права доступу, які необхідні для виконання конкретних завдань, що мінімізує ризики компрометації.

- Безперервний моніторинг активності – система аналізує поведінку користувачів та пристроїв у режимі реального часу, що дає змогу виявляти аномалії та потенційні загрози.

- Мікросегментація – поділ мережі на ізольовані зони, що обмежує поширення шкідливого програмного забезпечення у разі успішної атаки на один із сегментів.

- Шифрування даних – захист інформації як під час її передавання, так і в стані зберігання, що унеможливує її перехоплення [4].

На рис. 1.2 представлено спрощену схему впровадження Zero Trust, де спочатку здійснюється автентифікація користувача, а потім йому надається обмежений доступ на основі принципу мінімальних привілеїв.

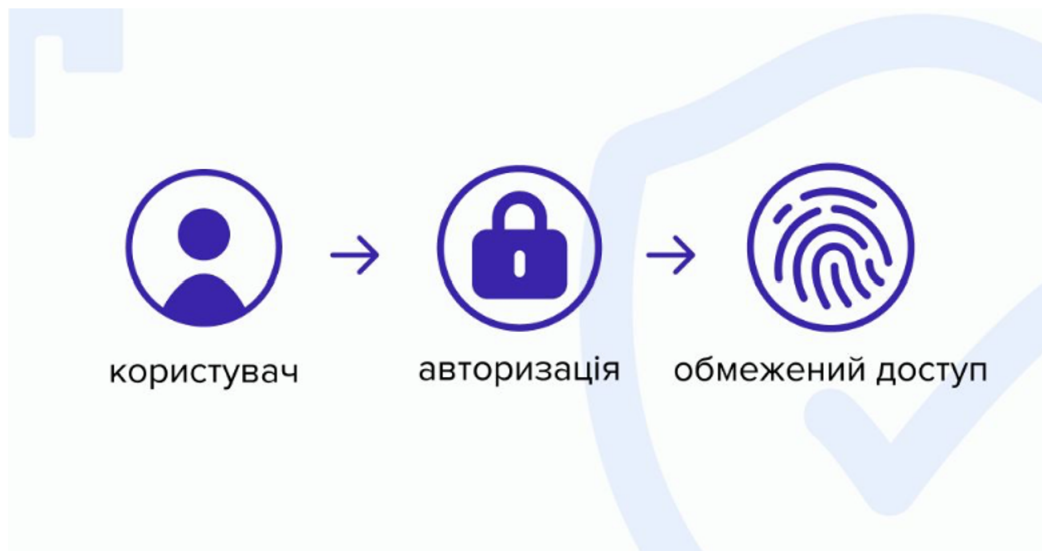


Рис. 1.2. Принцип роботи Zero Trust: автентифікація користувача та обмежений доступ.

Запровадження Zero Trust у сучасних корпоративних і державних системах значно підвищує рівень кібербезпеки, мінімізуючи ризики несанкціонованого доступу та зловживання привілеями. Цей підхід забезпечує гнучкість у захисті інформації, дозволяючи адаптувати механізми контролю доступу відповідно до мінливих загроз і нових технологічних викликів.

Модель Zero Trust була розроблена для протидії зростаючим кіберзагрозам, оскільки традиційні методи безпеки, засновані на периметровому підході, більше не можуть ефективно протидіяти сучасним атакам. Основний

принцип Zero Trust полягає в тому, що безпека має бути інтегрована на кожному рівні IT-інфраструктури, а доступ до будь-яких ресурсів має надаватися лише після ретельної перевірки.

Штучний інтелект у кібербезпеці

Штучний інтелект (ШІ) – це галузь комп'ютерних наук, що розробляє системи, здатні самостійно навчатися, аналізувати дані та приймати рішення без прямого людського втручання. У сфері кібербезпеки ШІ використовується для виявлення загроз, аналізу мережевого трафіку, реагування на атаки та прогнозування можливих векторів зловмисних дій. Основною перевагою ШІ є його здатність виявляти атаки на ранніх етапах, аналізуючи величезні обсяги даних у реальному часі [5].

Основні напрямки використання ШІ в кібербезпеці включають:

- *Автоматизований аналіз кіберзагроз* – використання алгоритмів машинного навчання для аналізу великих обсягів даних і виявлення аномалій;
- *Виявлення вторгнень (Intrusion Detection Systems – IDS)* – моніторинг мережевого трафіку з метою ідентифікації підозрілих активностей;
- *Фільтрація шкідливих повідомлень та фішингових атак* – застосування методів обробки природної мови (NLP) для аналізу текстів електронних листів та виявлення фішингових схем;
- *Автоматизація реагування на загрози* – впровадження SIEM-систем (Security Information and Event Management), які використовують штучний інтелект для швидкого аналізу та запобігання кібератакам [6].

Крім того, одним із перспективних напрямків є використання генеративних змагальних мереж (GANs) у створенні реалістичних сценаріїв атак для тестування захисту інформаційних систем [7].

Останні дослідження у сфері кібербезпеки підтверджують, що штучний інтелект значно підвищує ефективність боротьби з кіберзагрозами. Наприклад, у дослідженні "The Future of Artificial Intelligence in Cybersecurity: A Comprehensive Survey" зазначається, що використання методів машинного навчання значно покращує точність виявлення атак та аналізу загроз. Дослідники відзначають

важливість гібридних систем, що поєднують традиційні методи кібербезпеки з алгоритмами ШІ, що дозволяє підвищити рівень захисту від новітніх загроз [6].

Українські вчені також приділяють увагу проблемам кібербезпеки, зокрема впровадженню інтелектуальних систем для захисту критичної інфраструктури. В їхніх роботах аналізується використання Big Data та методів штучного інтелекту для автоматичного виявлення атак та підвищення рівня кіберзахисту у державних установах та великих підприємствах [8].

Таким чином, інтеграція ШІ у систему кібербезпеки є важливим кроком для створення ефективного та гнучкого захисту інформаційних ресурсів у сучасному цифровому світі.

1.2 Методи та алгоритми штучного інтелекту у виявленні кіберзагроз

Сучасні кіберзагрози стають дедалі складнішими та масовішими, що вимагає застосування інноваційних методів для їхнього виявлення та нейтралізації. Традиційні системи безпеки, які базуються на сигнатурному аналізі та ручному налаштуванні правил, часто виявляються неефективними проти новітніх загроз, таких як атаки нульового дня та складні багатовекторні атаки. Саме тому штучний інтелект (ШІ) став невід'ємним інструментом у сфері кібербезпеки [9].

Методи ШІ у виявленні кіберзагроз поділяються на такі основні категорії (Рис. 1.3):

- *Машинне навчання (ML)* – дозволяє системам самостійно навчатися на основі попередніх атак та аномальної поведінки. Використовується для аналізу великих обсягів даних, розпізнавання патернів, автоматичного створення моделей загроз та їх подальшого виявлення. Основні алгоритми ML включають методи класифікації, регресії та кластеризації. Цей підхід значно покращує можливість проактивного захисту від атак нульового дня та інших складних загроз.

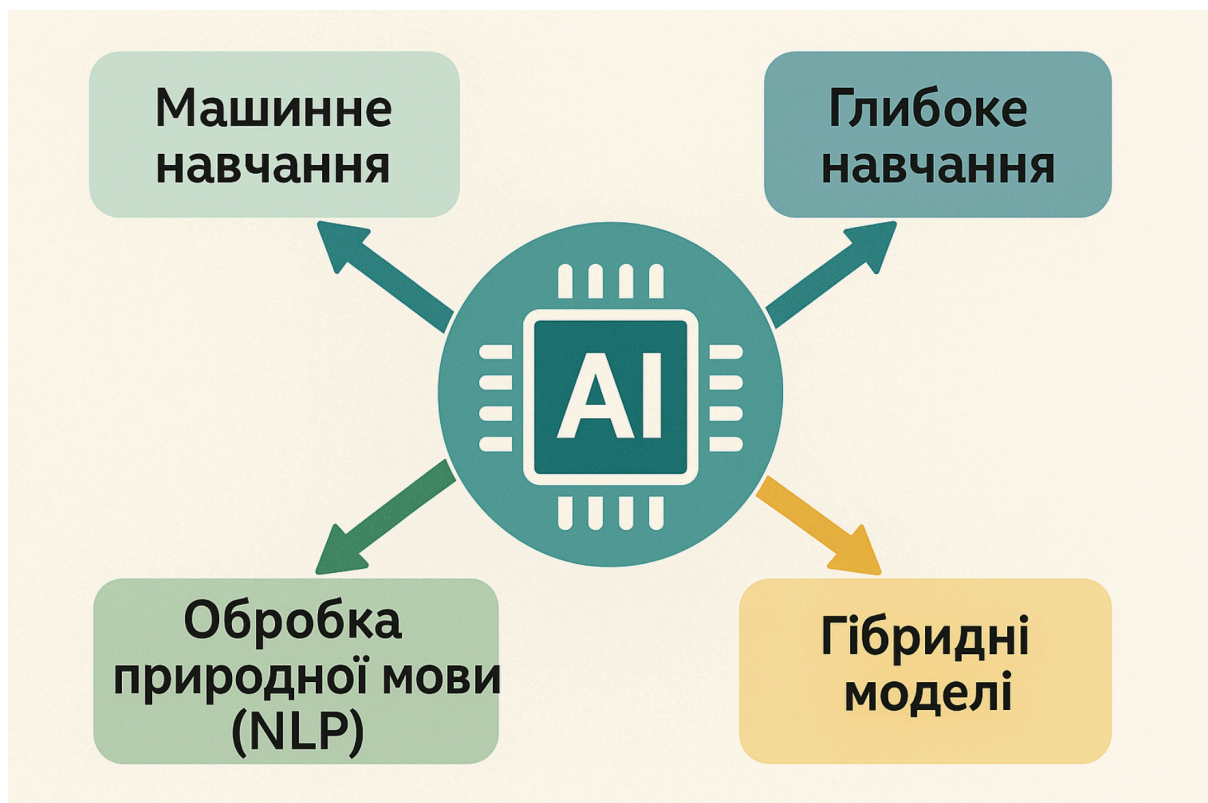


Рис. 1.3. Методи ШІ у виявленні кіберзагроз

- *Глибоке навчання (DL)* – є підвидом машинного навчання, що використовує багатошарові нейронні мережі для аналізу складних загроз. Дозволяє розпізнавати приховані взаємозв’язки між даними, які можуть свідчити про потенційну атаку. Застосовується для аналізу поведінки користувачів, виявлення аномального трафіку та виявлення складних атак, які змінюють свою форму (поліморфні загрози).

- *Обробка природної мови (NLP)* – використовується для аналізу текстової інформації, зокрема електронної пошти, повідомлень у месенджерах та публікацій у соціальних мережах. Застосовується для виявлення фішингових атак, спаму, шкідливих команд або інструкцій у текстових файлах. За допомогою NLP алгоритми можуть розрізняти тональність повідомлень, ідентифікувати підозрілі запити, а також автоматично блокувати небезпечні повідомлення.

- *Гібридні моделі* – поєднують кілька методів ШІ для підвищення ефективності виявлення загроз. Наприклад, комбінація ML та NLP дозволяє не лише аналізувати кіберзагрози на рівні поведінкової активності, а й одночасно виявляти атаки соціальної інженерії. Гібридні системи також можуть інтегрувати

традиційні підходи до кібербезпеки, такі як аналіз сигнатур загроз та поведінкове моделювання, для забезпечення багаторівневого захисту.

Машинне навчання є одним із найбільш ефективних методів у сфері кібербезпеки. Воно дозволяє аналізувати величезні обсяги даних, визначати потенційні загрози та автоматично реагувати на них. Це значно підвищує ефективність систем безпеки, дозволяючи їм адаптуватися до нових загроз у режимі реального часу [10].

Завдяки алгоритмам машинного навчання можливо:

- Виявляти аномалії в поведінці користувачів і пристроїв.
- Визначати нові види атак без потреби у попередньому навчанні на основі сигнатур.
- Аналізувати великі обсяги журналів подій та виявляти приховані закономірності.
- Оптимізувати процеси кіберзахисту, автоматизуючи аналіз загроз та реагування на них [11].

Основні алгоритми машинного навчання у кібербезпеці

У галузі кібербезпеки використовують декілька типів алгоритмів машинного навчання (Рис. 1.4)



Рис. 1.4. Основні алгоритми машинного навчання

Класифікаційні алгоритми (SVM, Random Forest, Decision Trees):

- *SVM (Support Vector Machine)* – використовується для аналізу шкідливих програм та визначення аномальної поведінки користувачів. Він ефективний у розпізнаванні кіберзагроз шляхом пошуку найбільш релевантних особливостей у даних.
- *Random Forest* – працює на основі ансамблю дерев ухвалення рішень, що дозволяє підвищити точність виявлення загроз. Він особливо ефективний у розпізнаванні складних атак завдяки комбінуванню декількох моделей ухвалення рішень.
- *Decision Trees* – допомагає швидко класифікувати загрози на основі певних правил, що спрощує аналіз лог-файлів безпеки.

Методи кластеризації (K-Means, DBSCAN):

- *K-Means* – дозволяє групувати схожі події у кластери, що допомагає ідентифікувати нетипову активність у поведінці користувачів або мережевому трафіку. Його використовують для виявлення нових загроз, які не мають визначених шаблонів.
- *DBSCAN* – ефективний для виявлення аномалій у великих наборах даних, що допомагає ідентифікувати підозрілі патерни мережевої активності. Він особливо корисний при виявленні ботнетів та інших складних атак.

Байєсівські мережі:

- *Наївний байєсівський класифікатор* – широко використовується для фільтрації спаму та виявлення шахрайських повідомлень. Завдяки своїй швидкодії та точності, цей метод допомагає ефективно визначати небезпечні електронні листи та шкідливі вкладення.
- *Bayesian Networks* – допомагає визначати ймовірність виникнення кіберзагроз, аналізуючи взаємозв'язки між різними подіями. Цей метод є особливо корисним для прогнозування потенційних атак та визначення ризиків кібербезпеки.

Алгоритми аномального виявлення (One-Class SVM, Isolation Forest):

- *One-Class SVM* – використовується для виявлення відхилень у поведінці користувачів, що може сигналізувати про можливу атаку. Він створює модель «нормальної» поведінки та ідентифікує всі відхилення як потенційні загрози.
- *Isolation Forest* – дозволяє швидко знаходити аномалії в потоці мережевого трафіку, що є корисним для боротьби з атаками нульового дня. Його принцип роботи ґрунтується на ізоляції рідкісних та унікальних подій у великих наборах даних.

Алгоритми підкріпленого навчання (Reinforcement Learning):

- *Q-Learning* – дозволяє автоматизувати реагування на загрози шляхом навчання системи на попередніх атаках. Він використовується для створення динамічних механізмів кіберзахисту, які адаптуються до нових атак.
- *Deep Q-Networks (DQN)* – поєднує глибоке навчання та навчання з підкріпленням для адаптації до нових видів загроз. Це дозволяє моделі самостійно коригувати власні рішення та визначати найефективніші стратегії для запобігання кібератакам [12].

Застосування методів та алгоритмів штучного інтелекту у виявленні кіберзагроз дозволяє суттєво підвищити ефективність систем кібербезпеки. Машинне навчання, глибоке навчання, обробка природної мови та гібридні моделі забезпечують багаторівневий підхід до виявлення та нейтралізації складних атак. Ці технології дозволяють автоматизувати процеси аналізу загроз, виявляти аномальну поведінку та забезпечувати проактивний захист інформаційних систем навіть від нових та невідомих загроз.

Поєднання різних методів штучного інтелекту дозволяє досягти більш високої точності виявлення та зменшити ймовірність хибнопозитивних спрацьовувань. Таким чином, використання ШІ є важливим елементом сучасної кібербезпеки, що дозволяє значно знизити ризики кіберінцидентів та забезпечити надійний захист інформаційних ресурсів.

1.3 Використання машинного навчання для аналізу поведінки користувачів

Машинне навчання відіграє ключову роль у сучасних системах кібербезпеки, зокрема в аналізі поведінки користувачів. Використання методів машинного навчання дозволяє створювати поведінкові моделі, які можуть ідентифікувати потенційно небезпечні дії, виявляти відхилення від звичайних шаблонів та запобігати кібератакам ще до їхньої реалізації. Такі алгоритми працюють у режимі реального часу, аналізуючи величезні обсяги даних та застосовуючи аналітичні методи для оцінки ризиків.

Основні підходи до аналізу поведінки користувачів у кібербезпеці включають наступні методи (Рис. 1.5):



Рис. 1.5. Підходи до аналізу поведінки користувачів у кібербезпеці

- *Аналіз аномалій* – використання алгоритмів для виявлення відхилень у поведінці користувачів. Наприклад, якщо працівник намагається отримати доступ до конфіденційних даних у неробочий час або з незвичної IP-адреси, система може виявити цей інцидент як потенційну загрозу. Також цей метод дозволяє виявляти незвичні транзакції у фінансових системах, визначати підозрілу активність у корпоративних мережах та блокувати автоматизовані атаки на основі виявлених патернів поведінки.

- *Методи кластеризації* – алгоритми, такі як K-Means та DBSCAN, групують користувачів за схожими поведінковими характеристиками, що

дозволяє визначати потенційно небезпечні дії, які відрізняються від загальноприйнятих моделей. Наприклад, алгоритми кластеризації допомагають ідентифікувати нових користувачів, які демонструють поведінку, характерну для злоумисників або ботнетів, а також виявляти внутрішніх загроз шляхом аналізу відхилень від типової поведінки працівників у системі.

- *Нейромережеві методи* – рекурентні нейронні мережі (RNN) та моделі довготривалої короткочасної пам'яті (LSTM) дозволяють аналізувати поведінкові тенденції користувачів у часі, що особливо ефективно для виявлення складних атак. Наприклад, LSTM-моделі використовуються для аналізу багаторівневих атак, що розгортаються поступово, а також для прогнозування майбутньої поведінки користувачів на основі історичних даних. Завдяки здатності виявляти закономірності у часових рядах, ці алгоритми успішно застосовуються для аналізу логів безпеки та моніторингу поведінки користувачів у корпоративних мережах.

- *Алгоритми підкріпленого навчання* – використовуються для автоматичного виявлення небезпечних поведінкових моделей та створення механізмів реагування на потенційні загрози. Алгоритми підкріпленого навчання (Reinforcement Learning, RL) застосовуються для створення автоматичних механізмів реагування, які можуть адаптуватися до змін у поведінці злоумисників. Наприклад, RL може бути використане для автоматизованого блокування облікових записів при виявленні підозрілої активності або для управління політиками доступу до корпоративних систем залежно від ризик-профілю користувача [9, 10, 12].

Застосування UBA у кібербезпеці

User Behavior Analytics (UBA) є одним із ключових компонентів сучасних систем кібербезпеки. Вона дозволяє аналізувати дії користувачів, виявляти підозрілу поведінку та своєчасно реагувати на потенційні загрози. Це досягається завдяки поєднанню алгоритмів машинного навчання, які можуть виявляти відхилення від нормальної поведінки, передбачати ризики та автоматизувати процес реагування на інциденти безпеки.

- *Виявлення компрометації облікових записів* – якщо система помічає незвичну активність у профілі користувача (наприклад, раптове використання VPN з іншої країни або доступ до ресурсів, які раніше не використовувалися), вона може заблокувати доступ або надіслати сповіщення адміністраторам. Такі методи аналізу можуть допомогти виявити спроби викрадення облікових даних, а також визначити потенційні ризики ще до їхньої реалізації.

- *Запобігання інсайдерським загрозам* – алгоритми UBA аналізують поведінку співробітників і можуть ідентифікувати потенційні загрози, пов'язані з витоком конфіденційної інформації. Якщо співробітник починає завантажувати великі обсяги даних або отримує доступ до критичних систем, якими раніше не користувався, це може бути сигналом внутрішньої загрози. Крім того, методи машинного навчання можуть аналізувати не лише дії співробітників у системі, але й взаємодію з файлами, електронними листами та корпоративними базами даних, визначаючи потенційні ризики витоку інформації.

- *Моніторинг доступу до критичних систем* – машинне навчання дозволяє створювати поведінкові профілі користувачів, що допомагає ідентифікувати спроби аномального доступу до важливих ресурсів організації. Використання машинного навчання для аналізу доступу до корпоративних серверів дозволяє ідентифікувати підозрілу активність, наприклад, спроби обходу багаторівневої автентифікації або використання викрадених облікових даних для доступу до критичних систем. Крім того, алгоритми можуть автоматично виявляти та блокувати підозрілі запити, мінімізуючи ризики внутрішніх загроз [13].

Машинне навчання є важливим інструментом у сфері аналізу поведінки користувачів, що дозволяє значно підвищити рівень кібербезпеки організацій. Використання алгоритмів класифікації, кластеризації та нейронних мереж дозволяє ефективно виявляти аномалії у поведінці користувачів, передбачати потенційні загрози та своєчасно реагувати на них.

Методи аналізу аномалій дозволяють ідентифікувати відхилення від звичайних поведінкових патернів, допомагаючи виявляти випадки компрометації облікових записів або несанкціонованого доступу до критичних

ресурсів. Кластеризація поведінкових моделей сприяє сегментації користувачів та виявленню потенційних загроз на основі нетипових дій. Використання нейромережевих методів, зокрема рекурентних нейронних мереж (RNN) та моделей довготривалої короткочасної пам'яті (LSTM), забезпечує можливість аналізу тенденцій у поведінці користувачів та прогнозування загроз.

User Behavior Analytics (UBA) виступає ключовим компонентом у сучасних системах кібербезпеки, надаючи змогу аналізувати взаємодію користувачів із системами та виявляти внутрішні загрози, що виникають у корпоративному середовищі. Завдяки автоматизованому моніторингу поведінки працівників, системи кіберзахисту можуть своєчасно ідентифікувати потенційно небезпечну активність та запобігати витоку конфіденційної інформації.

Однак, незважаючи на значні переваги, впровадження машинного навчання в аналіз поведінки користувачів супроводжується викликами, такими як необхідність у великих обсягах якісних даних для тренування моделей, ризики атак на алгоритми ШІ та можливість помилкових спрацювань. Вдосконалення технологій машинного навчання, розвиток методів підкріпленого навчання та комбінування різних моделей ШІ сприятиме підвищенню точності систем безпеки та їх здатності адаптуватися до нових загроз [13].

Таким чином, інтеграція машинного навчання у систему кібербезпеки дозволяє не лише ефективно виявляти загрози, але й забезпечувати проактивний захист цифрових ресурсів, роблячи інформаційні системи більш надійними та безпечними.

Технології штучного інтелекту (ШІ) змінюють традиційні підходи до кіберзахисту, підвищуючи ефективність виявлення загроз, автоматизації захисних заходів та адаптації до нових атак. Вони дозволяють системам кібербезпеки не лише реагувати на загрози, але й прогнозувати потенційні атаки, що дає змогу значно підвищити ефективність захисних механізмів. ШІ аналізує великі обсяги даних у реальному часі, виявляючи закономірності, які можуть вказувати на підготовку кібератак або аномальні дії користувачів.

Класичні методи кібербезпеки, що базуються на сигнатурному аналізі та попередньо визначених правилах, поступаються адаптивним рішенням, які навчаються на основі реальних загроз та атак. Завдяки машинному навчанню та глибоким нейронним мережам, сучасні системи можуть миттєво коригувати алгоритми детекції та реагування без необхідності втручання людини.

Зокрема, технології ШІ дозволяють ефективніше працювати із загрозами нульового дня, складними багатовекторними атаками та кіберзлочинністю, що використовує методи соціальної інженерії. Вони також інтегруються у хмарні системи безпеки та підтримують автоматичне оновлення захисних механізмів у відповідь на нові види загроз.

Сучасні дослідження демонструють, що інтеграція ШІ у системи кібербезпеки може зменшити кількість помилкових спрацювань, покращити швидкість виявлення загроз і мінімізувати шкоду від кібератак. Це сприяє переходу від реактивного до проактивного підходу, коли система безпеки не лише виявляє та блокує загрози, а й прогнозує їх та розробляє контрзаходи ще до їхньої реалізації. Завдяки можливостям машинного та глибокого навчання сучасні системи кібербезпеки можуть аналізувати великі масиви даних, передбачати потенційні загрози та самостійно адаптуватися до нових викликів. Інтеграція ШІ дозволяє перейти від реактивного до проактивного підходу в захисті інформаційних систем.

1.4 Глибоке навчання й аналіз великих даних у кібербезпеці

Глибоке навчання (Deep Learning, DL) та аналіз великих даних (Big Data Analytics) є потужними технологіями у сфері кібербезпеки, які забезпечують вирішення низки завдань (Рис. 1.6).



Рис. 1.6. Глибоке навчання й аналіз великих даних у кібербезпеці

Обидві технології дозволяють обробляти великі обсяги інформації, знаходити приховані закономірності та виявляти кіберзагрози в реальному часі. Інтеграція глибокого навчання з аналізом великих даних надає можливість створювати адаптивні системи безпеки, здатні оперативно реагувати на нові види атак.

Глибоке навчання використовується для аналізу складних загроз і атак, які неможливо виявити традиційними методами. Основні напрямки застосування DL у кібербезпеці охоплюють:

- *Виявлення загроз у мережевому трафіку* – згорткові нейронні мережі (CNN) та автоенкодери застосовуються для ідентифікації аномалій у мережевому трафіку та виявлення складних кібератак, наприклад DDoS-атак, шкідливого трафіку та спроб несанкціонованого доступу. Використання автоенкодерів дозволяє знаходити відхилення в потоці трафіку, а CNN ефективно працює з аналізом пакетів даних у мережі, допомагаючи розпізнавати загрози, які змінюють свої характеристики для обходу традиційних засобів захисту.

- *Аналіз шкідливого програмного забезпечення (Malware Analysis)* – глибоке навчання дозволяє аналізувати поведінку програмного забезпечення, ідентифікувати підозрілі патерни та розрізняти легітимні та шкідливі файли з високою точністю. Використання рекурентних нейронних мереж (RNN) та трансформерів допомагає виявляти шкідливі програми шляхом аналізу їхніх динамічних характеристик, наприклад взаємодії із системними файлами, мережевою активністю та доступом до критичних ресурсів.

- *Розпізнавання атак соціальної інженерії* – технології DL допомагають виявляти фішингові атаки та маніпулятивні методи соціальної інженерії шляхом аналізу текстових даних електронних листів та повідомлень у реальному часі. Використання рекурентних нейронних мереж (RNN) та трансформерів (наприклад, BERT) дозволяє моделювати природну мову для розпізнавання прихованих загроз у великих обсягах електронних повідомлень, а також аналізу тональності текстів, що допомагає виявити підозрілі листи та повідомлення.

- *Автоматизація реагування на загрози* – глибоке навчання допомагає автоматизувати класифікацію інцидентів безпеки, прогнозувати розвиток атак та запускати відповідні механізми реагування. Наприклад, використання глибоких підкріплених навчальних моделей (Deep Reinforcement Learning, DRL) дозволяє системам безпеки самостійно адаптувати свої стратегії протидії загрозам, аналізуючи історичні інциденти та ефективність попередніх заходів [14].

Застосування засобів **аналізу великих даних** дозволяє вирішувати такі завдання з кібербезпеки як:

- *Кореляція подій безпеки* – технології Big Data дозволяють поєднувати дані з різних джерел (лог-файли, мережевий трафік, системи моніторингу) для виявлення зв'язків між загрозами та ідентифікації атак ще на ранніх стадіях. Це допомагає запобігати складним багатовекторним атакам. Використання розподілених обчислювальних систем (наприклад, Apache Spark) дозволяє обробляти величезні масиви подій безпеки в реальному часі та корелювати потенційно небезпечні дії, щоб виявити складні схеми атак.

- *Прогнозування загроз* – аналіз великих масивів історичних даних допомагає передбачати потенційні атаки шляхом розпізнавання повторюваних шаблонів поведінки зловмисників та уразливостей у системах. Використання глибоких нейронних мереж, таких як LSTM та GRU, дозволяє створювати прогнозні моделі, що базуються на часових рядах подій кібербезпеки, і визначати ймовірність виникнення майбутніх атак на основі аналізу попередніх інцидентів.

- *Оптимізація захисних механізмів* – великі дані дозволяють адаптувати механізми безпеки, змінювати політики захисту в режимі реального часу та покращувати алгоритми детектування загроз. Використання методів кластерного аналізу та графових нейронних мереж (Graph Neural Networks, GNN) допомагає створювати комплексні профілі загроз, що дозволяє покращити точність визначення аномальної активності та виявлення атак на корпоративні мережі та державні інформаційні системи [9, 15].

Глибоке навчання та аналіз великих даних суттєво змінюють сучасні підходи до кібербезпеки, роблячи системи захисту більш автономними, швидкими й ефективними. Завдяки використанню алгоритмів глибокого навчання, таких як згорткові та рекурентні нейронні мережі, сучасні системи безпеки здатні аналізувати складні патерни кіберзагроз, виявляти аномалії у мережевому трафіку та ідентифікувати новітні види атак, включаючи атаки нульового дня.

Використання методів аналізу великих даних дозволяє обробляти величезні обсяги інформації з різних джерел, у тому числі з лог-файлів, журналів подій та потоків мережевого трафіку. Це забезпечує можливість виявлення загроз на ранніх стадіях та прогнозування потенційних атак на основі історичних даних. Інтеграція Big Data Analytics у системи кібербезпеки також допомагає мінімізувати кількість хибнопозитивних спрацювань та покращити точність виявлення загроз.

Штучний інтелект у кібербезпеці дозволяє переходити від реактивного підходу до проактивного, де система не лише виявляє та блокує загрози, а й передбачає можливі ризики та автоматично адаптується до змін у середовищі загроз. Використання автоматизованих механізмів реагування на основі

глибокого навчання дозволяє значно скоротити час реагування на інциденти безпеки та підвищити ефективність захисту інформаційних ресурсів.

Однак, незважаючи на значні переваги, впровадження ШІ в кібербезпеку супроводжується певними викликами, такими як висока обчислювальна складність, потреба у великих наборах якісних даних для навчання моделей та ризику атак на самі алгоритми ШІ. Проте подальший розвиток технологій, таких як квантові обчислення, федеративне навчання та розширене використання генеративних змагальних мереж (GANs), може зробити системи кібербезпеки ще більш ефективними та захищеними.

Таким чином, поєднання глибокого навчання та аналізу великих даних відкриває нові горизонти для розбудови надійних, адаптивних та самонавчальних систем кібербезпеки, що дозволить ефективно протидіяти сучасним загрозам у цифровому просторі.

1.5 Вплив технологій ШІ на традиційні системи кіберзахисту

Технології ШІ змінюють традиційні підходи до кіберзахисту, підвищуючи ефективність виявлення загроз, автоматизації захисних заходів та адаптації до нових атак. Вони дозволяють системам кібербезпеки не лише реагувати на загрози, але й прогнозувати потенційні атаки, що дає змогу значно підвищити ефективність захисних механізмів. ШІ аналізує великі обсяги даних у реальному часі, виявляючи закономірності, які можуть вказувати на підготовку кібератак або аномальні дії користувачів.

Класичні методи кібербезпеки, що базуються на сигнатурному аналізі та попередньо визначених правилах, поступаються адаптивним рішенням, які навчаються на основі реальних загроз та атак. Завдяки машинному навчанню та глибоким нейронним мережам, сучасні системи можуть миттєво коригувати алгоритми виявлення та реагування без необхідності втручання людини.

Зокрема, технології ШІ дозволяють ефективніше працювати із загрозами нульового дня, складними багатовекторними атаками та кіберзлочинністю, що

використовує методи соціальної інженерії. Вони також інтегруються у хмарні системи безпеки та підтримують автоматичне оновлення захисних механізмів у відповідь на нові види загроз.

Сучасні дослідження демонструють, що інтеграція ШІ у системи кібербезпеки може зменшити кількість помилкових спрацювань, покращити швидкість виявлення загроз і мінімізувати шкоду від кібератак. Це сприяє переходу від реактивного до проактивного підходу, коли система безпеки не лише виявляє та блокує загрози, а й прогнозує їх та розробляє контрзаходи ще до їхньої реалізації. Завдяки можливостям машинного та глибокого навчання сучасні системи кібербезпеки можуть аналізувати великі масиви даних, передбачати потенційні загрози та самостійно адаптуватися до нових викликів. Інтеграція ШІ дозволяє перейти від реактивного до проактивного підходу в захисті інформаційних систем.

Основні напрями впливу ШІ на традиційні системи кіберзахисту:

1. *Автоматизація виявлення загроз* – традиційні системи кібербезпеки значною мірою покладаються на сигнатурний аналіз та попередньо задані правила. ШІ дозволяє автоматично виявляти нові види загроз, використовуючи поведінковий аналіз та алгоритми машинного навчання. Нейронні мережі можуть аналізувати аномальні патерни трафіку та виявляти загрози, які не мають попередньо визначених сигнатур. Такі методи дають змогу швидко виявляти новітні загрози, у тому числі атаки нульового дня, які не можуть бути виявлені традиційними антивірусними базами даних.

2. *Зниження рівня хибнопозитивних спрацювань* – традиційні системи виявлення загроз часто створюють велику кількість хибних сповіщень. Це може перевантажувати аналітиків безпеки та знижувати ефективність реагування на реальні загрози. ШІ дозволяє більш точно аналізувати інциденти безпеки, використовуючи алгоритми класифікації та кластеризації. Завдяки глибокому навчанню моделі можуть ідентифікувати характерні шаблони атак та зменшити кількість неправдивих попереджень, дозволяючи спеціалістам сфокусувати увагу на дійсно критичних загрозах.

3. *Адаптивний кіберзахист* – технології ШІ можуть постійно оновлювати свої моделі на основі нових атак та автоматично змінювати стратегії захисту. Наприклад, алгоритми підкріпленого навчання (Reinforcement Learning) допомагають системам безпеки навчатися на основі історичних даних про атаки та розробляти оптимальні контрзаходи. Це дозволяє кіберзахисту бути більш гнучким і адаптивним до новітніх загроз.

4. *Розширений аналіз шкідливого програмного забезпечення (Malware Analysis)* – традиційні антивірусні рішення використовують бази сигнатур для виявлення шкідливого ПЗ, що робить їх уразливими до атак нульового дня. ШІ може аналізувати поведінку програм у реальному часі, ідентифікуючи шкідливі дії навіть у випадках, коли сигнатурні методи не працюють. Методи поведінкового аналізу на основі глибокого навчання дозволяють розпізнавати приховані загрози, які змінюють свій код для обходу традиційних антивірусних баз даних [5-6].

Технології ШІ відіграють вирішальну роль у трансформації традиційних систем кіберзахисту, забезпечуючи адаптивність, автоматизацію та підвищену ефективність реагування на загрози. Інтеграція алгоритмів машинного навчання дозволяє системам безпеки не лише виявляти та блокувати загрози, але й прогнозувати потенційні атаки, що сприяє проактивному підходу до кібербезпеки.

Використання ШІ виявляється критично важливим у протидії атакам нульового дня, зниженні рівня хибнопозитивних спрацювань та розширенні аналізу поведінки користувачів. Системи, що використовують нейронні мережі та методи глибокого навчання, здатні швидко адаптуватися до нових загроз, самотійно вдосконалюючи алгоритми виявлення та реагування.

Водночас впровадження ШІ у сферу кібербезпеки супроводжується певними викликами, такими як необхідність великих обчислювальних ресурсів, загроза атак на самі алгоритми штучного інтелекту та складність пояснення рішень, прийнятих моделями ШІ. Проте, незважаючи на ці труднощі, розвиток

штучного інтелекту у кібербезпеці відкриває нові можливості для побудови стійких, ефективних та автономних систем захисту інформаційних ресурсів.

Висновки до розділу 1

Дослідження показало, що в умовах динамічного кіберсередовища, постійного ускладнення й розширення масштабів атак з Інтернету кібербезпека потребує застосування інноваційних підходів до їх виявлення та нейтралізації. З цією метою в галузі кібербезпеки широко застосовуються методи штучного інтелекту, які забезпечують автоматизований аналіз кіберзагроз, виявлення вторгнень, фільтрацію шкідливих повідомлень й ідентифікацію фішингових атак; автоматизацію реагування на загрози.

Встановлено, що для виявлення кіберзагроз використовують різні методи ШІ, такі як машинне навчання (ML), глибоке навчання (DL), обробка природної мови (NLP), гібридні моделі, а також алгоритми класифікації, виявлення аномалій, підкріпленого навчання, кластеризації та байєсівські мережі.

Машинне навчання відіграє ключову роль у системах аналізу поведінки користувачів, створюючи поведінкові моделі для ідентифікації потенційно небезпечних дій, виявлення відхилень і запобігання кібератакам. Такі алгоритми працюють у режимі реального часу, аналізуючи величезні обсяги даних і застосовуючи аналітичні методи для оцінювання ризиків.

Значний внесок у підвищення ефективності кіберзахисту вносять технології глибокого навчання (DL), які забезпечують вирішення низки завдань, зокрема виявлення загроз у мережевому трафіку, розпізнавання атак соціальної інженерії, аналіз шкідливого ПЗ, автоматизація реагування на загрози. Технології аналізу великих даних (Big Data Analytics) дозволяють здійснювати кореляцію подій безпеки, прогнозування загроз, оптимізацію захисних механізмів.

Отже, технології ШІ змінюють традиційні підходи до кіберзахисту, підвищуючи їх ефективність і сприяючи переходу від реактивного до проактивного підходу до забезпечення кібербезпеки. ШІ дозволяє системам

кібербезпеки не лише реагувати на загрози, але й прогнозувати і самостійно адаптуватися до потенційних атак, успішно протидіяти загрозам нульового дня, складним багатовекторним атакам і кіберзлочинам із використанням соціальної інженерії. Технології III дозволяють зменшити кількість помилкових спрацювань, покращити швидкість виявлення загроз і мінімізувати шкоду від кібератак.

РОЗДІЛ 2 ОСОБЛИВОСТІ Й МЕТОДИКА ЗАСТОСУВАННЯ ШІ В ЗАБЕЗПЕЧЕННІ КІБЕРБЕЗПЕКИ ОБ'ЄКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ

2.1 Політика забезпечення кібербезпеки критичної інфраструктури в Україні

Україна як держава з розвиненою цифровою інфраструктурою та активною присутністю в міжнародному кіберпросторі неухильно впроваджує стратегії та нормативні акти, що мають на меті захист критичних об'єктів від кіберзагроз. У рамках національної політики кібербезпеки сформовано комплекс правових, організаційних і технологічних заходів, які відповідають міжнародним стандартам.

Ключовими нормативно-правовими актами України у сфері кіберзахисту є:

- *Закони України «Про інформацію», «Про Національну програму інформатизації», «Про захист інформації в інформаційно-комунікаційних системах».* Ці нормативні акти формують основу регулювання безпеки інформації та інформаційно-комунікаційних систем в Україні [16-18].

- *Закон України «Про основні засади забезпечення кібербезпеки України»* (2017). Закон визначає правові та організаційні основи забезпечення кібербезпеки, регулює діяльність суб'єктів кіберзахисту, встановлює принципи взаємодії між державними органами та об'єктами критичної інфраструктури [19].

- *Закон України «Про критичну інфраструктуру»* (2023), в якому дано визначення ключових термінів, окреслено основні засади державної політики у сфері захисту критичної інфраструктури та віднесення об'єктів до критичної інфраструктури, сформовано перелік секторів та правила категоризації об'єктів критичної інфраструктури, визначено режим функціонування і суб'єктів національна система захисту критичної інфраструктури тощо [20].

- *Національна стратегія кібербезпеки України* (затверджена Указом Президента України № 447/2021). Документ визначає пріоритети державної

політики в сфері кібербезпеки на середньострокову перспективу. Основна увага приділяється розвитку національного кіберпростору, удосконаленню системи кіберзахисту, формуванню кіберстійкості, а також інтеграції з безпековими структурами ЄС і НАТО [21].

- *Концепція створення державної системи захисту критичної інфраструктури України*, затверджена розпорядженням Кабінету Міністрів України № 1009-р від 6 грудня 2017 року. Документ визначає категорії критичних об'єктів, підходи до їх ідентифікації, оцінювання ризиків і механізми забезпечення їхньої безпеки [22].

- Постанова Кабінету Міністрів України від 9 жовтня 2020 р. № 1109 “*Деякі питання об'єктів критичної інфраструктури*”, яка затвердила Порядок віднесення об'єктів до критичної інфраструктури, секторальний перелік об'єктів критичної інфраструктури сектору, перелік секторів критичної інфраструктури і методику категоризації об'єктів критичної інфраструктури [23].

- Указ Президента України «Про рішення Ради національної безпеки і оборони України від 17 жовтня 2023 року «*Про організацію захисту та забезпечення безпеки функціонування об'єктів критичної інфраструктури та енергетики України в умовах ведення воєнних дій*» (з грифом ДСК) [24].

Важливу роль у реалізації політики кіберзахисту відіграють такі державні органи:

- *Державна служба спеціального зв'язку та захисту інформації України (ДССЗІ)* — є головним органом у сфері технічного та криптографічного захисту інформації. ДССЗІ відповідає за створення та реалізацію технічної політики у сфері кіберзахисту, розробляє вимоги до безпеки інформаційно-телекомунікаційних систем, координує роботу Національного координаційного центру кібербезпеки при РНБО, а також здійснює оперативне реагування на кіберінциденти державного рівня [25].

- *Служба безпеки України (СБУ)* — забезпечує кібербезпеку з позиції національної безпеки. Основні завдання СБУ в цій сфері включають протидію кіберрозвідці, виявлення та припинення дій іноземних спецслужб у

кіберпросторі, боротьбу з кібертероризмом і захист критичних об'єктів державного значення від диверсій у кіберпросторі. Також СБУ уповноважена розслідувати серйозні інциденти та координацію роботи спеціалізованих підрозділів [26].

- *Міністерство оборони України та Генеральний штаб Збройних Сил України* — відповідають за формування і реалізацію політики у сфері кібероборони. У структурі ЗСУ діють окремі підрозділи кібербезпеки, які здійснюють моніторинг загроз, забезпечують захист військових інформаційних систем та розробляють наступальні і захисні кібероперації в умовах гібридної війни.

- *Кіберполіція Національної поліції України* — займається розслідуванням кіберзлочинів, таких як фішинг, шахрайство, хакерські атаки на бізнес та приватних осіб. Також кіберполіція веде активну просвітницьку діяльність, проводить кампанії щодо запобігання інтернет-злочинам, бере участь у міжнародних операціях із протидії організованим кіберзлочинності [27].

- *Національний координаційний центр кібербезпеки при РНБО України* — об'єднує роботу державних та приватних структур у сфері кіберзахисту. В рамках його діяльності організовуються тренінги, здійснюються перевірки інформаційної безпеки та формуються рекомендації щодо підвищення кіберстійкості об'єктів критичної інфраструктури [28].

У межах реалізації державної політики кіберзахисту важливу роль відіграють ключові суб'єкти національної системи кібербезпеки. Їхня взаємодія забезпечує комплексний підхід до виявлення, запобігання та реагування на кіберзагрози. Структурну модель взаємодії основних органів представлено на рисунку 2.1.



Рисунок 2.1. Суб'єкти національної системи кібербезпеки України
та їх взаємодія

Україна долучена до міжнародних програм у кіберпросторі, зокрема, взаємодію чи з НАТО, Євросоюзом, та між народними організаціями кібербезпеки. Актуальним залишається приведення українського законодавства до європейських норм, як передбачено Угодою про асоціацію. Посилюється співпраця між приватним та державним секторами, коли підприємства критичної інфраструктури працюють з державою для посилення кіберзахисту.

Загалом, стратегія України щодо кібербезпеки зосереджена на створенні всеосяжної, міцної та єдиної системи, що здатна оперативно протидіяти актуальним викликам. Подальше удосконалення цієї політики передбачає використання передових технологій, таких як штучний інтелект, машинне навчання та автоматизовані системи аналізу загроз.

2.2 Основні кіберзагрози для об'єктів критичної інфраструктури

Об'єкти критичної інфраструктури (ОКІ), до яких належать енергетичні, транспортні, телекомунікаційні, водопровідні, медичні й інші стратегічно важливі системи, відіграють ключову роль у забезпеченні стабільної роботи держави та соціуму. В умовах швидкої цифровізації й посилення кіберзалежності

такі об'єкти стають потенційною ціллю для складних кібернападів, які здатні вивести з ладу не лише окремі складові, а й цілу інфраструктуру, що, у свою чергу, може спричинити масові збої, людські втрати або значні економічні збитки.

У таблиці 2.1 представлено основні типи кіберзагроз, характерні для об'єктів критичної інфраструктури, із прикладами та засобами їх виявлення.

Таблиця 2.1.

Класифікація кіберзагроз для об'єктів критичної інфраструктури

Тип загрози	Приклад	Потенційні наслідки	Засоби виявлення та реагування
Атаки типу DDoS	Перевантаження серверів компанії	Виведення з ладу вебпорталів, затримка реакцій	IDS/IPS, обмеження трафіку, балансування
Шкідливі програми (Malware)	Інфікування SCADA-систем троянами	Зміна режимів роботи обладнання, крадіжка даних	Антивірус, sandbox-аналіз, EDR
Атаки нульового дня	Експлуатація невідомих вразливостей	Неконтрольований доступ до систем, зупинка процесів	ШІ-аналіз поведінки, віртуалізація середовища
Фішинг та соціальна інженерія	Шкідливі листи персоналу	Компрометація облікових записів	Фільтрація е-пошти, навчання персоналу
Внутрішні загрози	Несанкціоновані дії персоналу	Витік конфіденційної інформації	UBA-системи, контроль доступу

- *DDoS-атаки (Distributed Denial of Service)*. Ці атаки спираються на великі мережі заражених машин (ботів) для перевантаження цільових систем, спричиняючи їхню непрацездатність. У контексті критичної інфраструктури, це може означати зупинку систем спостереження або керування, що призводить до повного блокування роботи об'єкта.

- *Шкідливе ПЗ і програми-вимагачі.* Віруси, такі як програми-вимагачі (ransomware), здатні перекрити доступ до ключових систем чи зашифрувати важливу інформацію, вимагаючи грошову винагороду. Оскільки об'єкти критичної інфраструктури не можуть припиняти роботу, наслідки таких атак здатні бути надзвичайно серйозними. Показовим прикладом є атака вірусу NotPetya, яка завдала збитків на мільйони гривень в Україні.

- *Цілеспрямовані довготривалі атаки (APT, Advanced Persistent Threats).* Це складні, багаторівневі напади, які втілюються протягом тривалого періоду часу з метою збору даних розвідки або підготовки до руйнівних дій. Вони можуть включати проникнення через фішинг, використання уразливостей нульового дня, встановлення бекдорів. Характерна риса цих атак – їх непомітність, висока технологічність і стратегія тривалого перебування в мережі об'єкта.

- *Фішинг і атаки соціальної інженерії.* Переважно реалізуються через фішингові листи, телефонні розмови або підроблені онлайн-ресурси. Цілі зловмисників полягають у тому, щоб змусити працівника надати доступ чи виконати дії, які ставлять під загрозу безпеку системи. Людський аспект часто стає «слабкою ланкою» в охороні ОКІ.

- *Загрози з боку внутрішніх користувачів (Insider threats).* Співробітники або підрядники можуть ненавмисно чи спеціально порушити правила безпеки. Маючи фізичний або логічний доступ до систем, інсайдери стають особливо небезпечними та важко піддаються виявленню [29].

Крім розглянутих у таблиці 2.1 серйозні наслідки для кібербезпеки ОКІ можуть мати такі вразливості:

- *Вразливості у SCADA та ICS-системах.* Системи контролю та управління виробничими процесами часто використовують застарілі протоколи з недостатнім рівнем захисту. Відсутність оновлень безпеки, слабкі механізми аутентифікації, незахищені з'єднання з мережею – усі ці фактори відкривають двері для зловмисних впливів на ОКІ.

- *Недостатня зрілість процесів реагування на інциденти.* На багатьох об'єктах критичної інфраструктури досі не розроблені протоколи реагування або

немає команд кібербезпеки, що ускладнює можливості оперативно зменшити негативний вплив у разі появи загрози [8].

Отже, об'єкти критичної інфраструктури стають мішенню для різноманітних кібератак, які здатні завдати шкоди не лише економіці, але й мати критичні наслідки для національної безпеки. З огляду на це, надзвичайно важливо розробляти та впроваджувати комплексні системи кіберзахисту з багаторівневим підходом, спрямовані на активне виявлення загроз. Це, зокрема, передбачає використання інструментів штучного інтелекту та машинного навчання, які дозволяють розпізнавати аномалії у режимі реального часу та реагувати на них, перш ніж виникнуть серйозні наслідки.

2.3 Особливості забезпечення кібербезпеки атомних електростанцій в Україні

Відповідно до Закону України «Про критичну інфраструктуру» [20] до ОКІ, які виконують життєво важливі функції та/або послуги, порушення яких призводить до негативних наслідків для національної безпеки України, належать підприємства галузі енергозабезпечення, в тому числі атомні електростанції.

Атомні електростанції (АЕС) — це надзвичайно складні інженерні споруди, що об'єднують енергетичну, фізичну та інформаційну інфраструктуру. Будь-яке втручання в роботу цих систем, зокрема через кіберпростір, може призвести до техногенних катастроф, витоку радіації чи виведення з ладу енергетичних ланцюгів. З огляду на це, питання кібербезпеки для АЕС має особливе значення та регулюється на національному та міжнародному рівнях.

Особливості забезпечення кібербезпеки АЕС зумовлені такими ключовими аспектами:

- *Інформаційно-технологічна спадковість.* Частина інфраструктури атомних електростанцій зводилася десятиліттями, тому і містить застарілі системи управління, що ускладнює їх оновлення відповідно до актуальних вимог безпеки.

Застарілі протоколи передавання даних, відсутність підтримки оновлень програмного забезпечення – усе це формує «слабкі місця» в архітектурі захисту.

- *Розмежування IT та OT (Operational Technology).* На АЕС існує суворе розмежування між інформаційними (офісними) та технологічними (операційними) мережами. Це робиться з метою мінімізації ризику розповсюдження шкідливого програмного забезпечення. Водночас, це ускладнює об'єднання зусиль моніторингу та запровадження уніфікованої системи реагування на інциденти.

- *Високий рівень вимог до сертифікації змін.* Кожна модифікація програмного чи апаратного забезпечення потребує сертифікації, отримання згоди від регуляторних установ, таких як ДІЯРУ (Державна інспекція ядерного регулювання України), плюс проведення випробувань на спеціалізованих стендах. Вказане значно гальмує оперативне впровадження найсвіжіших інструментів для виявлення загроз.

- *Потенційна загроза для безпеки держави.* Зважаючи на стратегічну вагу атомної енергетики, АЕС можуть стати ціллю для спланованих кібератак з боку державних структур чи організованих угруповань. У 2022 році експерти з МАГАТЕ акцентували увагу на тому, що об'єкти атомної енергетики України постійно перебувають під загрозою як фізичних, так і кібернетичних нападів [29].

- *Міжнародні вимоги до кіберзахисту.* Усі АЕС зобов'язані дотримуватися міжнародних стандартів, зокрема рекомендацій МАГАТЕ (наприклад, NSS No. 17 «Computer Security at Nuclear Facilities»), ISO/IEC 27001 [30] та стандартів ENISA [31]. Ці документи встановлюють вимоги до багаторівневого захисту, менеджменту інцидентів, шифрування, ізоляції систем і моніторингу подій.

- *Важливість людського чинника.* Працівники АЕС зобов'язані систематично підвищувати кваліфікацію з кібербезпеки та взаємодії з захищеними системами. Помилки або недбалі дії операторів здатні спричинити критичні проблеми. Отже, першочерговим завданням є розгортання систем

аналізу поведінки користувачів, що виявляють аномалії у поведінці користувачів у реальному часі.

- *Обов'язковість безперервної роботи систем.* Реактори не можна зупинити або перезавантажити без серйозних наслідків, що унеможлиблює багато традиційних процедур безпеки (наприклад, перезавантаження або сканування системи). Тому превентивні заходи, резервування й мережеве сегментування є основними методами захисту.

- *Потреба у взаємодії з державними та міжнародними структурами.* Енергоатом, як оператор АЕС в Україні, співпрацює з ДССЗІ, СБУ, МАГАТЕ та CERT-UA для координації дій під час кіберінцидентів, тестування готовності до загроз та спільної розробки сценаріїв реагування.

Сучасні стратегії кібербезпеки АЕС ґрунтуються на теорії «поглибленого захисту» (Defence-in-Depth). Він передбачає комплекс заходів захисту: фізичних, адміністративних, технічних та інформаційних, що діють узгоджено.

Для кращого розуміння архітектури кіберзахисту атомних електростанцій у межах принципу «поглибленого захисту» наведено схематичне зображення моделі зонування IT- та OT-інфраструктури АЕС з урахуванням основних засобів безпеки (Рис. 2.2) [32].

Сьогодні широкого вжитку у забезпеченні кібербезпеки ОКІ набувають інструменти штучного інтелекту. Їх використовують для виявлення відхилень у даних телеметрії, аналізу журналів безпеки, знаходження взаємозв'язків між подіями, а також для автоматизованого реагування на інциденти. До того ж, одним з багатообіцяючих напрямків у сфері кібербезпеки АЕС є розробка так званих цифрових двійників (digital twins) критичних об'єктів. Ці віртуальні копії дають змогу імітувати наслідки кібератак на фізичні процеси, не наражаючи на небезпеку функціонування справжніх систем. Це відкриває шлях для безпечного випробування сценаріїв атак, навчання персоналу та відпрацювання планів реагування в умовах, які максимально наближені до реальності.

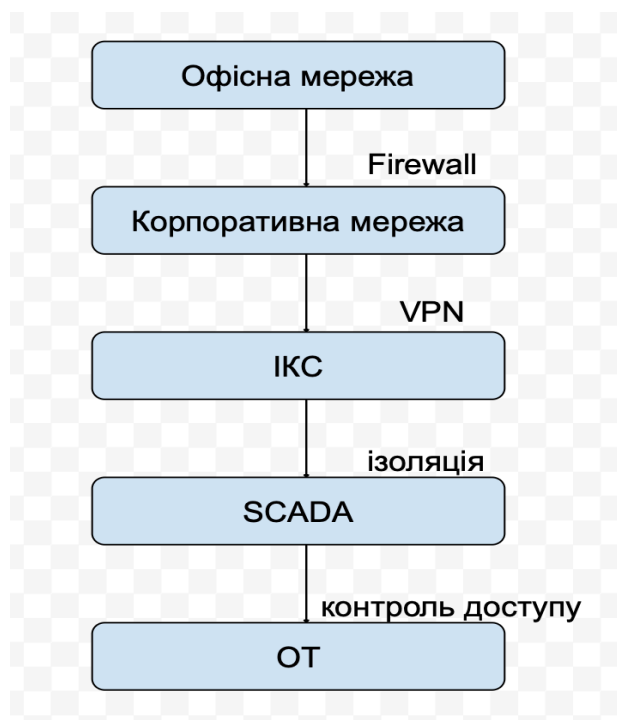


Рисунок 2.2. Модель зонування інформаційної інфраструктури АЕС відповідно до рівнів безпеки

Ще однією перспективною технологією є застосування систем аналізу великих даних для формування поведінкових шаблонів устаткування та персоналу. Виявлення відмінностей у функціонуванні автоматизованих систем або діях користувачів сприяє своєчасному виявленню інцидентів, які непомітні при використанні традиційних методів.

У майбутньому інтеграція систем штучного інтелекту, хмарних технологій (за умови ізолюваного середовища) та автоматизованих систем реагування стане фундаментом для створення стійкої цифрової архітектури безпеки в атомній енергетиці. Завдяки комбінації традиційних та інноваційних підходів АЕС зможуть зберігати високий рівень кіберстійкості навіть в умовах гібридної війни та еволюції кіберзагроз.

2.4 Аналіз програм кіберзахисту на прикладі Рівненської АЕС

Рівненська атомна електростанція (РАЕС), розташована біля міста Вараш у Рівненській області, є однією з найбільших діючих атомних електростанцій

України. РАЕС є структурним підрозділом державного підприємства «НАЕК Енергоатом» і відіграє визначальну роль у забезпеченні енергетичної безпеки держави.

РАЕС використовує у своїй роботі чотири енергоблоки: два з реакторами ВВЕР-440 та два - з ВВЕР-1000 (Рис. 2.4). Загальна проектна потужність електростанції сягає 2880 МВт, що дозволяє їй входити до переліку лідерів з виробництва електроенергії в Україні. Завдяки високим стандартам технологічного процесу, стабільній роботі та ефективній системі управління, РАЕС безперервно задовольняє потреби мільйонів споживачів у електроенергії [33].



Рис. 2.4. Енергоблоки Рівненської атомної електростанції

З огляду на критичний характер об'єкта, кібербезпека на РАЕС є пріоритетом державного значення. Будь-яке втручання у функціонування ІКС атомної станції може призвести до серйозних наслідків не лише для енергетичної системи, а й для національної безпеки. Тому на РАЕС впроваджується комплексна система кіберзахисту, яка охоплює технічні, організаційні, процедурні й аналітичні заходи для забезпечення безпеки інформаційних і технологічних систем.

Рівненська атомна електростанція (РАЕС) – ключовий елемент енергетичної інфраструктури України, що відіграє важливу роль у впровадженні

передових методів кіберзахисту. Враховуючи стратегічну значущість цього об'єкта, на РАЕС впроваджено серію ініціатив та заходів, метою яких є зведення до мінімуму можливості несанкціонованого проникнення в критичні системи управління.

Ключовими елементами програми кіберзахисту РАЕС (Рис. 2.5) є:

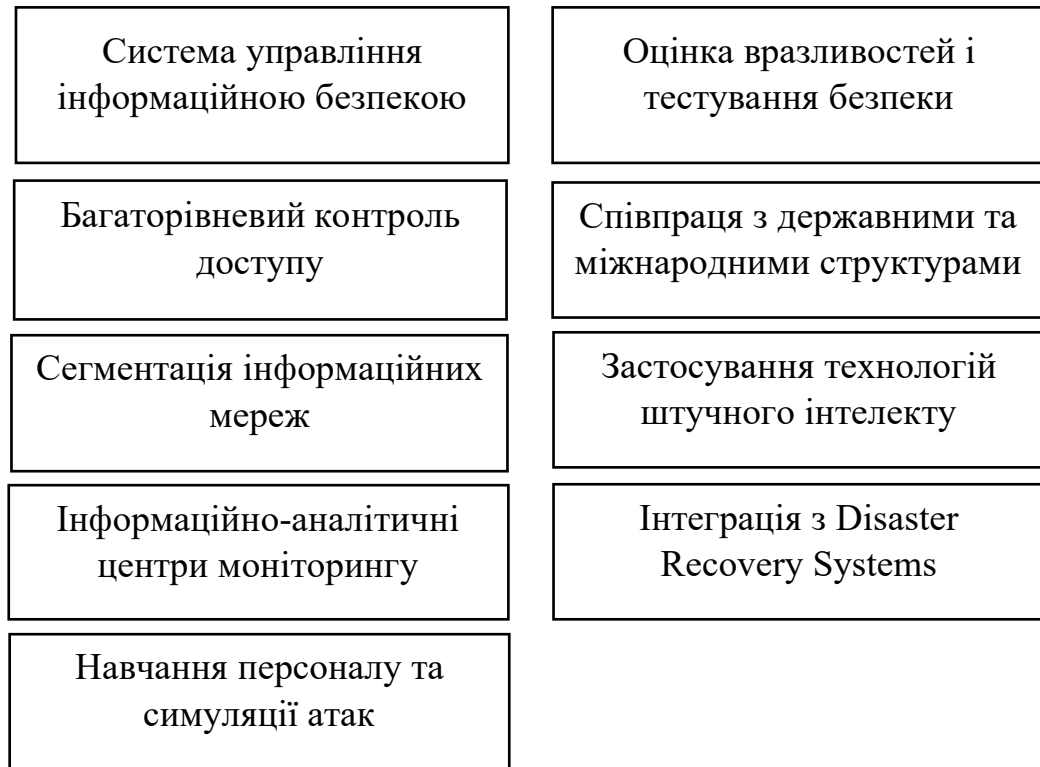


Рис. 2.5. Ключові елементи програми кіберзахисту РАЕС

- *Система управління інформаційною безпекою (СУІБ).* На РАЕС впроваджено політику, яка визначає вимоги до управління безпекою інформаційних ресурсів. Ця політика охоплює оцінювання ризиків, контроль доступу, управління інцидентами, захист від зловмисного ПЗ та аудит безпеки. Вона відповідає міжнародним стандартам ISO/IEC 27001 та регламентам НП 306.2.202-2015 [34].

- *Багаторівневий контроль доступу.* Всі критичні системи РАЕС функціонують із обмеженим доступом на основі принципу найменших привілеїв. Впроваджено механізми багатофакторної автентифікації, електронного журналювання дій користувачів, а також автоматизованих інструментів виявлення аномалій у поведінці персоналу.

- *Сегментація інформаційних мереж.* IT-інфраструктура поділена на логічно ізольовані зони. Операційні системи керування (ОТ) повністю відокремлені від корпоративних мереж. Доступ до ІКС можливий лише через спеціалізовані шлюзи з контролем протоколів та ретельною перевіркою оновлень програмного забезпечення.

- *Інформаційно-аналітичні центри моніторингу.* Для безперервного контролю безпеки застосовуються SIEM-системи, які здійснюють кореляцію подій у реальному часі. Усі події аналізуються за допомогою вбудованих правил та штучного інтелекту для виявлення потенційних інцидентів.

- *Навчання персоналу та симуляції атак.* Персонал РАЕС регулярно проходить тестування знань, тренінги з реагування на кіберінциденти та участь у внутрішніх навчаннях з моделювання кіберзагроз. Це дозволяє сформувати культуру безпеки на всіх рівнях.

- *Оцінка вразливостей і тестування безпеки.* Проводяться регулярні пентести, сканування систем на вразливості, ревізії конфігурацій. Виявлені недоліки усуваються з урахуванням пріоритетів і критичності системи.

- *Співпраця з державними та міжнародними структурами.* РАЕС взаємодіє з ДССЗЗІ, СБУ, CERT-UA, MAGATE та іншими структурами в межах системи обміну інформацією про кіберзагрози, аналізу ризиків та впровадження рекомендацій щодо захисту об'єктів критичної інфраструктури.

- *Застосування технологій штучного інтелекту.* У пілотному режимі впроваджуються модулі поведінкового аналізу, аналізу логів, ідентифікації несанкціонованих дій у системах SCADA. AI-моделі використовуються для побудови профілів нормальної поведінки операторів і автоматичного виявлення відхилень.

- *Інтеграція з Disaster Recovery Systems.* На випадок кіберінцидентів розроблено сценарії автоматичного перемикання на резервні системи, ізоляції пошкоджених компонентів та відновлення працездатності критичних процесів без зупинки енергоблоків [34, 35].

Особливу увагу слід звернути також на впровадження комплексних засобів безпеки, які включають фізичні, технічні, адміністративні та програмні елементи. На РАЕС впроваджено багаторівневий моніторинг доступу до систем управління технологічним процесом, використовуються спеціалізовані шлюзи для контролю обміну даними між ІТ та ОТ-середовищами, проводиться детальна інвентаризація активів та управління життєвим циклом інформаційних ресурсів.

Інноваційним напрямом кіберзахисту є впровадження концепції нульової довіри Zero Trust. Згідно з нею, жоден користувач чи пристрій не отримує автоматичного статусу надійного. Кожен запит проходить ретельну автентифікацію, авторизацію та перевірку на відповідність політикам безпеки. Такий підхід суттєво зменшує ризик компрометації внутрішніх систем, навіть якщо зловмисник зуміє проникнути у корпоративну мережу.

Відтак, РАЕС засвідчує цілісний, системний підхід до кіберзахисту, що об'єднує традиційні засоби інформаційної безпеки з передовими інструментами аналізу та штучного інтелекту. Це сприяє оперативному реагуванню на актуальні загрози у кіберпросторі та гарантує стабільність енергетичної інфраструктури в умовах гібридного протистояння [35].

2.5 Методика використання штучного інтелекту для кіберзахисту ОКІ

У сучасному цифровому світі, де ІКС є визначальними чинниками роботи підприємств та об'єктів критичної інфраструктури, збільшується необхідність впровадження передових методів кіберзахисту. Штучний інтелект відкриває кардинально нові перспективи у цій галузі, трансформуючи усталені підходи до реагування на кіберінциденти й дозволяючи застосовувати випереджаючі стратегії безпеки.

Застосування ШІ для забезпечення захисту ІКС полягає у впровадженні алгоритмів машинного та глибокого навчання в інфраструктуру компанії. Це дозволяє не тільки негайно реагувати на ймовірні загрози, але й досліджувати їхню сутність, траєкторію розповсюдження та передбачати наступні кроки

злочинців. Зокрема, використовуючи поведінковий аналіз, що ґрунтується на аналітиці поведінки користувачів та об'єктів (UEBA), системи здатні виявляти підозрілі дії без необхідності в заздалегідь визначених шаблонах.

Одним з найцікавіших способів є використання алгоритмів глибокого навчання для розбору мережевого трафіку. Ці моделі мають здатність знаходити комплексні багатовекторні атаки, такі як атаки нульового дня чи поліморфне шкідливе програмне забезпечення. Додатково, аналіз великих наборів логів та подій безпеки (Big Data) допомагає виявляти приховані відхилення та тенденції, які важко помітити за допомогою звичайних інструментів моніторингу.

ШІ також активно використовується в системах оркестрації, автоматизації й реагування (Security Orchestration, Automation and Response, SOAR), які відповідають за автоматизоване реагування на інциденти: від виявлення до ізоляції та усунення загроз. Це помітно зменшує навантаження на фахівців безпеки, підвищує оперативність реагування та зводить до мінімуму ризики, пов'язані з людським фактором.

Іншим ключовим моментом виступає гнучке регулювання доступу до об'єктів ІКС. За допомогою парадигми Zero Trust та контекстно-залежного дозволу, системи можуть аналізувати ризики в реальному часі, опираючись на поведінку користувача, місцезнаходження, вид пристрою та інші чинники. Алгоритми ШІ підлаштовують політики доступу у відповідності до змін середовища і ступеня небезпеки.

Необхідно підкреслити, що застосування ШІ сприяє підвищенню адаптивності й масштабованості систем захисту. Замість застарілих правил, що вимагають перманентного оновлення, запроваджуються адаптивні моделі, які мають здатність до самонавчання та самостійної корекції. Це особливо важливо для підприємств енергетичної галузі, зокрема, таких як РАЕС, де безпека має визначальне значення, а атаки здатні призвести до катастрофічних наслідків.

Штучний інтелект відіграє визначальну роль у формуванні надійного захисту ІКС, особливо з огляду на невинне посилення кібернетичних викликів. ІКС охоплюють фізичні пристрої та програмне забезпечення, котрі забезпечують

пересилання, зберігання, опрацювання та захист інформації. У нинішньому кіберпросторі, де напади стають дедалі більш витонченими та замаскованими, ШІ здійснює автоматизований аналіз даних, розпізнавання відхилень від норми, прогнозування інцидентів безпеки та виконання відповідних дій у реальному часі.

Аналіз використання основних підходів і методів ШІ для кіберзахисту ІКС РАЕС дає підстави окреслити типову методику застосування ШІ на об'єктах критичної інфраструктури України (Рис. 2.5).



Рисунок 2.5. Вектори використання ШІ в ІКС

Отже, основними елементами методики застосування ШІ для захисту ІКС ОКІ є:

- *Аналіз поведінки користувачів та систем (UBA/UEBA).* За допомогою машинного навчання створюються поведінкові моделі користувачів і пристроїв. Це дозволяє виявляти відхилення від звичної поведінки, що може свідчити про спробу атаки з боку внутрішнього чи зовнішнього зловмисника.
- *Автоматизоване виявлення вторгнень (IDS/IPS).* ШІ дозволяє ідентифікувати загрози, які не мають чітких сигнатур, завдяки аналізу мережевого трафіку, логів і метаданих. Алгоритми глибокого навчання

розпізнають шаблони, характерні для атак, та попереджають адміністраторів ще до виникнення інциденту.

- *Прогнозування та запобігання атакам.* Використовуючи дані про попередні загрози, системи з ШІ можуть формувати прогностичні моделі, що визначають вірогідність появи нових атак. Це дозволяє організаціям готуватися до загроз заздалегідь.

- *Розширений аналіз логів та великих даних.* Завдяки поєднанню Big Data та ШІ можливо аналізувати мільйони записів логів у режимі реального часу, виявляючи аномалії, які залишаються непоміченими традиційними системами.

- *Реалізація Zero Trust архітектури.* ШІ застосовується для динамічного контролю доступу до ресурсів ІКС на основі контекстної інформації — поведінки, місця розташування, типу пристрою, історії активності користувача тощо.

- *Автоматизація реагування на інциденти.* Інструменти SOAR (Security Orchestration, Automation and Response), що використовують ШІ, дозволяють автоматизувати перевірку, підтвердження та нейтралізацію інцидентів безпеки, значно скорочуючи час реагування.

- *Аналіз вразливостей.* ШІ допомагає виявляти нові вразливості у системах, порівнюючи поведінкові показники програмного забезпечення і систем із нормальними сценаріями експлуатації.

- *Інтелектуальна класифікація трафіку.* Завдяки глибоким нейронним мережам можливо з високою точністю класифікувати типи трафіку, визначаючи легітимну та шкідливу активність навіть при використанні шифрування [34, 35].

У поєднанні зі стандартними методами захисту, штучний інтелект відкриває шлях до створення гнучкої, здатної до самонавчання системи безпеки, що не тільки реагує на атаки, а й передбачає їх виникнення, тим самим значно підвищуючи кіберстійкість ІКС перед лицем сучасних кіберзагроз. Використання цього інструменту набуває ключового значення для об'єктів критичної інфраструктури на зразок РАЕС, де навіть незначні порушення можуть спричинити катастрофічні наслідки.

Підсумовуючи, слід відзначити, що інтеграція ІІІ в інформаційно-комунікаційні системи ОКІ значно підвищує рівень кіберзахисту та дозволяє створювати стійкі до загроз архітектури. З огляду на розвиток технологій та зростання складності атак, впровадження ІІІ є не просто бажаним, а стратегічно необхідним кроком у цифровій трансформації ІКС критичних об'єктів і державних підприємств.

Висновки до розділу 2

У другому розділі розглянуто практичні аспекти застосування методів штучного інтелекту у сфері кібербезпеки об'єктів критичної інфраструктури на прикладі Рівненської атомної електростанції. Аналіз показав, що сучасні кіберзагрози для АЕС характеризуються високим рівнем складності й потенційною небезпекою для національної безпеки, що вимагає особливої уваги до безпеки ІКС та ОТ-систем.

Встановлено, що на загальнодержавному рівні розроблена ґрунтовна нормативно-правова база, яка регламентує вимоги до кіберзахисту, а повноваження з кібербезпеки покладені на відповідальні інституції: ДССЗЗІ, СБУ, Міноборони, кіберполіція тощо на чолі з Національним координаційним центром кібербезпеки при РНБО. Особливості кібербезпеки АЕС включають зонування мереж, розмежування ІТ та ОТ, високу критичність до змін, обмеження на пряме оновлення систем, а також постійний контроль з боку державних та міжнародних регуляторів.

У межах дослідження проаналізовано діючі компоненти системи кіберзахисту РАЕС, серед яких рішення з управління інформацією та подіями безпеки (SIEM), навчання персоналу, політики обмеженого доступу, використання сегментованих мереж та інструментів моніторингу. Підвищенню ефективності цих механізмів забезпечення кібербезпеки АЕС сприяє впровадження методики застосування ІІІ, яка містить такі елементи: аналіз поведінки користувачів та систем, автоматизоване виявлення вторгнень,

прогнозування та запобігання атакам, розширений аналіз логів і великих даних, реалізація архітектури нульової довіри, автоматизація реагування на інциденти, аналіз вразливостей та інтелектуальна класифікація трафіку

Таким чином, результати роботи доводять, що інтеграція ІІІ в системи кіберзахисту не лише підвищує рівень безпеки об'єктів критичної інфраструктури, а й дозволяє реалізувати проактивну модель захисту, яка є ефективною в умовах гібридних загроз і зростаючої цифрової залежності стратегічно важливих підприємств.

РОЗДІЛ 3 ПРАКТИЧНЕ ВИКОРИСТАННЯ ШІ ДЛЯ ПОСИЛЕННЯ КІБЕРБЕЗПЕКИ ПІДПРИЄМСТВ ТА ОКІ

3.1 Сучасні інструменти кіберзахисту на основі ШІ

Стрімкий розвиток ІТ та супутні кіберзагрози підштовхують бізнеси та державні структури до активного пошуку результативних та передових методів кібербезпеки. Виняткове місце серед сучасних технологій належить інструментам, які базуються на використанні штучного інтелекту. Вони дають можливість забезпечувати високий рівень захисту завдяки автоматизації процесів, котрі традиційно потребують значних людських зусиль, а також через здатність до постійного самовдосконалення та пристосування до нових викликів. В умовах безперервного ускладнення кібератак і збільшення вимог до захисту важливої інформації, використання ШІ перетворюється на необхідність, оскільки він спроможний ефективно протистояти не тільки відомим, а й новим загрозам.

Серед основних сучасних інструментів, які використовують ШІ для посилення кібербезпеки, можна виокремити такі (Рис. 3.1):

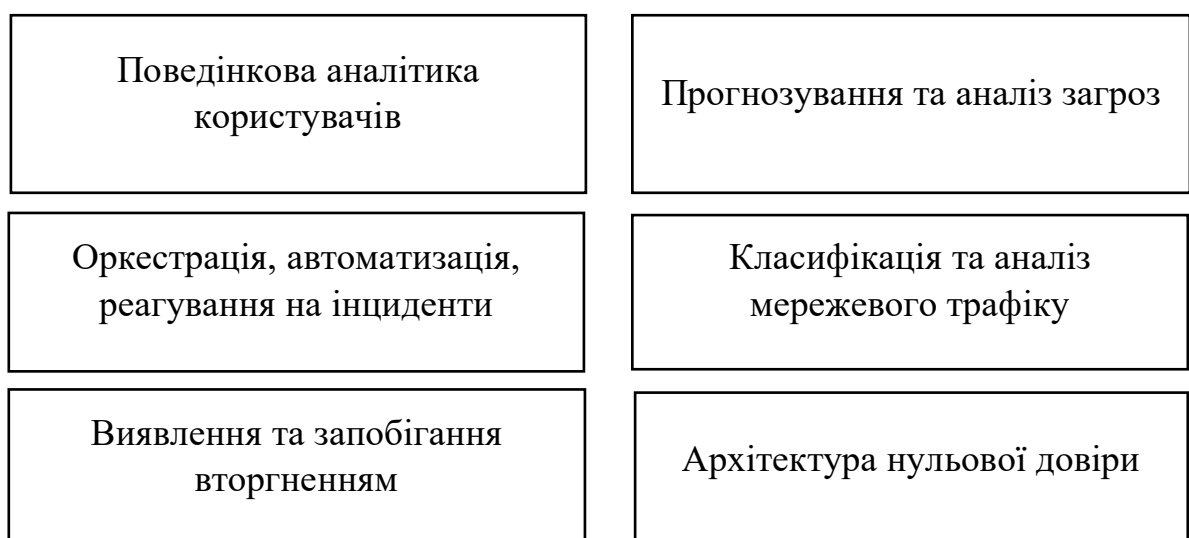


Рис. 3.1. Сучасні інструменти кібербезпеки, які використовують ШІ

Інструменти поведінкової аналітики користувачів (UBA/UEBA)

Інструменти, що спираються на аналіз поведінки користувачів та об'єктів (UEBA – User and Entity Behavior Analytics), дають змогу виявляти нетипові дії, котрі можуть вказувати на кіберзагрозу. Системи, на зразок Splunk UBA, IBM QRadar User Behavior Analytics, Exabeam та Forcepoint, застосовують алгоритми машинного навчання (МН) задля побудови базової поведінкової моделі та швидкого виявлення відмінностей від звичайних шаблонів взаємодії користувачів та пристроїв. Exabeam дає можливість знаходити компрометацію облікових записів, тоді як Forcepoint зосереджується на запобіганні витоків інформації (DLP), досліджуючи поведінку користувачів у реальному часі [36].

Автоматизовані системи реагування на інциденти (SOAR)

Системи автоматизації та оркестрації реагування на інциденти (Security Orchestration, Automation and Response) відіграють критичну роль у сучасному кіберзахисті. Інструменти, такі як Palo Alto Networks Cortex XSOAR, IBM Resilient, Splunk Phantom, Demisto та Swimlane, дозволяють автоматизувати важливі етапи реагування на інциденти, здійснювати аналіз величезних масивів даних про інциденти безпеки та пришвидшувати реакцію на загрози завдяки використанню інтелектуальних скриптів. Swimlane, наприклад, автоматизує реагування через візуалізацію сценаріїв, що відчутно зменшує час реагування на інциденти [37].

Системи виявлення та запобігання вторгненням

Системи виявлення та запобігання вторгненням (IDS/IPS), що впроваджують алгоритми МН і ШІ, гарантують виявлення загроз без потреби у заздалегідь визначених сигнатурах. Прикладами таких систем є Darktrace Enterprise Immune System, Cisco Secure Network Analytics та Vectra Cognito, що застосовують найсучасніші алгоритми глибокого навчання для скрупульозного аналізу мережевого трафіку та виявлення аномалій. Darktrace діє подібно до імунної системи, автоматично вивчаючи типову поведінку мережі для ідентифікації відхилень, а Vectra Cognito миттєво реагує на складні та нещодавно виниклі загрози, що залишаються непомітними для традиційних систем.

Системи прогнозування та аналізу загроз (Threat Intelligence Platforms)

Платформи аналізу загроз, що використовують ІІІ, надають змогу передбачити можливі кібернетичні атаки, аналізуючи величезні обсяги даних. Завдяки їм, як-от Recorded Future, Anomali ThreatStream та ThreatConnect, аналізуються відомості про минулі випадки, актуальні загрози та їх розвиток, що сприяє ефективному прогнозуванню. Recorded Future пропонує точні моделі передбачення, а Anomali ThreatStream об'єднує дані з різних джерел, автоматизуючи аналіз та розповсюдження інформації про небезпеки.

Системи класифікації мережевого трафіку

Класифікація та аналіз трафіку на основі ІІІ забезпечують точне розмежування між дозволеним та зловмисним мережевим потоком. Рішення, як-от Cisco Encrypted Traffic Analytics (ETA) та Awake Security, задіюють глибинні нейронні мережі для аналізу навіть зашифрованого трафіку, що сприяє негайному блокуванню прихованої зловмисної активності та виявленню складних загроз, зокрема АРТ (Advanced Persistent Threats) [6].

В епоху, коли ІІІ все активніше перетворюється на ключовий фактор у сфері кіберзахисту, відкриваючи небачені перспективи для розпізнавання, нейтралізації та протидії атакам, його використання не є простим шляхом. Впровадження систем на основі ІІІ надає небачені можливості автоматизації та покращення захисту, проте супроводжується проблемами, пов'язаними з надійністю, етикою, технічними недоліками й новими кіберзагрозами. Застосування ІІІ в стратегіях безпеки потребує не тільки усвідомлення потенційних плюсів, а й розуміння ризиків, що можуть підірвати стабільність та довіру до даних систем.

Підхід Zero Trust на основі ІІІ

Архітектура Zero Trust, що інтегрує ІІІ, передбачає гнучке управління доступом, ґрунтуючись на контекстних даних про те, як поведуться користувачі, які пристрої вони використовують, звідки підключаються. Алгоритми ІІІ, застосовані в платформах на зразок Microsoft Zero Trust Security та Google BeyondCorp, миттєво змінюють політики безпеки, мінімізуючи загрозу

несанкціонованого проникнення і покращуючи загальну безпеку в кіберпросторі [38].

Для наочного відображення сучасних інструментів на основі ІІІ, які використовуються у сфері кіберзахисту, у таблиці 3.1 наведено їх основні функції та приклади платформ.

Таблиця 3.1.

Інструменти штучного інтелекту для кіберзахисту та їхні основні функції

Інструмент	Основна функція	Приклади платформ
UEBA (поведінкова аналітика)	Виявлення аномальної поведінки користувачів і пристроїв	Splunk UBA, IBM QRadar UBA, Exabeam, Forcepoint
SOAR	Автоматизація реагування на інциденти	Cortex XSOAR, Splunk Phantom, Swimlane
IDS/IPS з підтримкою ІІІ	Виявлення і блокування вторгнень	Darktrace, Cisco Secure Analytics, Vectra Cognito
Threat Intelligence Platforms	Прогнозування загроз, аналіз тенденцій	Recorded Future, Anomali ThreatStream, ThreatConnect
Класифікація трафіку	Аналіз трафіку, виявлення шкідливих даних	Cisco ETA, Awake Security
Zero Trust на основі ІІІ	Динамічне управління доступом	Microsoft Zero Trust, Google BeyondCorp

Отже, використання новітніх засобів кіберзахисту на базі ІІІ дає можливість не тільки швидко реагувати на наявні атаки, але й успішно передбачати майбутні кіберінциденти, підлаштовуючи стратегії безпеки під постійні зміни в ландшафті загроз. Застосування таких прогресивних технологій гарантує не тільки зниження ризиків і запобігання ймовірним інцидентам, а й відчутно скорочує витрати на обслуговування служби безпеки завдяки високому рівню автоматизації процесів. Подальший розвиток та інтеграція ІІІ у сферу

кібербезпеки визначається як основний напрям у створенні потужного та надійного захисту інформаційних активів ОКІ [33].

3.2 Автоматизація аналізу загроз за допомогою ШІ

Сучасна кібербезпека потребує не тільки миттєвого виявлення потенційних загроз, а й злагодженого управління та оперативного реагування на них у реальному часі. Одним із основних інструментів, що сприяють автоматизації аналізу загроз, є SIEM-системи (Security Information and Event Management). Застосовуючи ШІ та МН, ці системи здатні оперувати колосальними обсягами даних, знаходити підозрілі інциденти і швидко реагувати на можливі кібератаки, що відчутно знижує ризики для бізнесу та стратегічно важливих об'єктів.

Основний задум SIEM-систем зводиться до централізованого збирання та аналізу даних з різноманітних джерел (журналів безпеки, мережевого устаткування, серверів, додатків тощо), що дає можливість сформувати цілісне уявлення про стан безпеки організації. Актуальні SIEM-системи, на зразок IBM QRadar, Splunk Enterprise Security, LogRhythm та Microsoft Sentinel, комбінують традиційні функції збирання логів і подій з розширеними можливостями аналітики й автоматизованого реагування.

Застосування ШІ в системах SIEM дозволяє розв'язати ключові проблеми, з якими зустрічаються фахівці з безпеки. Мова йде про велику кількість помилкових спрацьовувань, дефіцит кваліфікованого персоналу та безперервну еволюцію ландшафту кіберзагроз. Засоби ШІ можуть автоматично знаходити приховані тенденції та незвичайну поведінку, що може вказувати на атаку. Це дає змогу негайно інформувати експертів або запускати відповідні сценарії реагування [39].

На рисунку 3.2 наведено узагальнену блок-схему роботи SIEM-системи, яка ілюструє основні етапи обробки та аналізу подій безпеки.

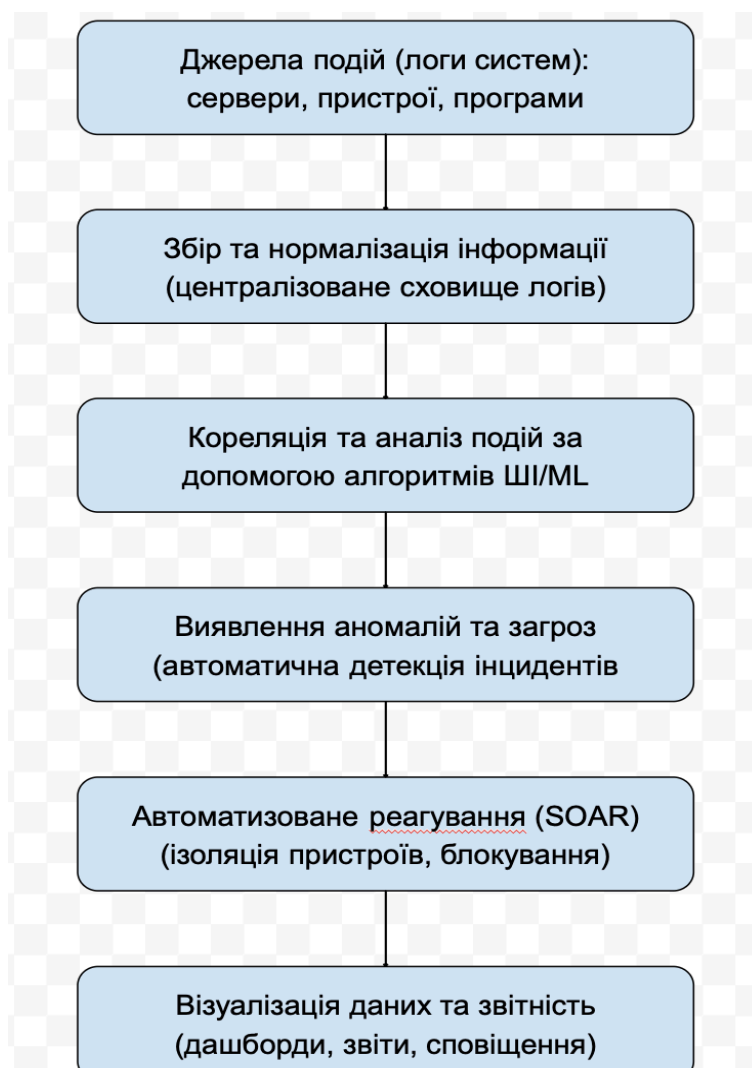


Рис. 3.2. Блок-схема роботи SIEM-системи

До основних переваг автоматизації аналізу загроз за допомогою SIEM-систем відносять:

Автоматизація виявлення загроз SIEM-системах з використанням ШІ. Однією з переваг SIEM-систем є швидкий збір даних з різноманітних джерел – серверів, мережевого обладнання, програмного забезпечення та систем безпеки. Ці дані одразу нормалізуються, тобто приводяться до єдиного вигляду. Це дає змогу сформувати цілісну базу даних, що критично важлива для аналізу потенційних інцидентів. Наприклад, такі системи, як Splunk або IBM QRadar, можуть автоматично обробляти десятки тисяч подій щосекунди – показник, недосяжний при ручній обробці. Завдяки цьому існує можливість негайно виявляти взаємозв'язки між подіями й оперативно реагувати на інциденти, зменшуючи ймовірні збитки.

Використання машинного навчання для виявлення аномалій. Сучасні SIEM-системи широко залучають технології МН для виявлення аномальної активності користувачів та пристроїв. Ці алгоритми тренуються на прикладах звичайного функціонування мережі, ідентифікуючи шаблони допустимої поведінки, після чого автоматично генерують сповіщення про будь-які нетипові або підозрілі дії, що можуть вказувати на наявність загрози. Наприклад, модуль IBM QRadar User Behavior Analytics автоматично формує профілі нормальної поведінки користувачів і в режимі реального часу повідомляє про відхилення, що забезпечує швидке реагування на можливі кіберзагрози.

Зниження рівня хибних спрацювань. Автоматизація розбору великої кількості інцидентів значно скорочує хибні спрацювання, які часто супроводжують ручний аналіз. Застосовуючи методи кластеризації, класифікації та аналізу поведінки, SIEM-системи мають здатність відрізняти реальні ризики від незначних чи помилкових сповіщень. Наприклад, рішення від LogRhythm або Exabeam застосовують алгоритми МН, що підвищує точність виявлення справжніх загроз, дозволяючи фахівцям з безпеки зосереджуватися на найкритичніших інцидентах.

Автоматизація реагування (SOAR-функції). Інтеграція SIEM з платформами SOAR дає змогу суттєво прискорити процес усунення інцидентів. Виявивши шкідливу подію, система автоматично запускає процедури реагування без втручання людини: блокує підозрілі IP-адреси, ізолює заражені комп'ютери або обмежує доступ користувачів. Наприклад, Microsoft Sentinel має вбудовані сценарії автоматичного реагування, які негайно блокують загрозу, що значно скорочує час реакції та зменшує ризики серйозних інцидентів [39-40].

Покращення візуалізації і звітності. SIEM-системи надають наочне графічне уявлення про стан кібербезпеки за допомогою інтерактивних інформаційних панелей (дашбордів). Це дає змогу швидко оцінити загальну картину безпеки, виявити потенційні загрози та оперативно приймати рішення на основі фактів. Наприклад, платформа Splunk Security Intelligence Platform оснащена інтерактивними звітами та візуалізаціями, що відображають важливі

показники і тенденції безпеки. Це дозволяє керівництву оперативно отримувати інформацію про поточний стан і ефективно управляти кіберризиками [41].

Однак варто пам'ятати, що, хоч би якими привабливими були автоматизовані системи, інтеграція SIEM з використанням ШІ вимагає ретельного планування, адаптації та безперервного моніторингу. Основними проблемами є складність початкового налаштування, потреба в точному калібруванні алгоритмів МН, а також гарантування високої якості даних, на основі яких система навчається та функціонує.

Отже, автоматизація аналізу загроз із застосуванням актуальних SIEM-систем, наділених ШІ та МН, становить критичний компонент кіберзахисту сучасних підприємств і критичної інфраструктури. Це дає змогу не тільки суттєво прискорити виявлення загроз та реагування на них, але й обмежити вплив людського чинника на продуктивність безпеки. В подальшому передбачається інтенсивний розвиток цих технологій, що перетворить їх на невід'ємну складову дієвої стратегії кібербезпеки.

3.3 Використання ШІ для виявлення аномалій у мережевому трафіку

У сучасному цифровому світі, де щоденно генерується колосальна кількість мережевого трафіку, вчасне виявлення аномалій є надзвичайно важливим для гарантування кібербезпеки. Застосування ШІ дає змогу суттєво збільшити ефективність моніторингу та захисту ІКС від різних загроз. ШІ відкриває можливості, що виходять за межі простого виявлення відомих загроз, на відміну від традиційних систем. Завдяки аналізу поведінки та алгоритмам МН, він здатний виявляти нові, неklasифіковані аномалії. Це забезпечує проактивний захист, суттєво зменшуючи ймовірність компрометації даних.

Основні напрями застосування ШІ для виявлення аномалій у мережевому трафіку (Рис. 3.3):

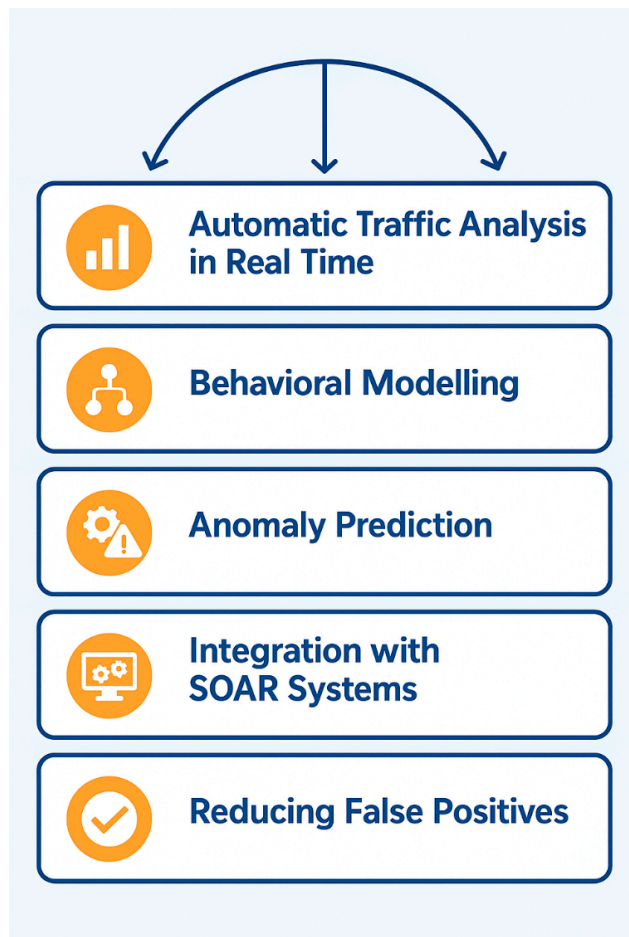


Рис. 3.3. Основні напрями застосування ІІІ для виявлення аномалій у мережевому трафіку

Автоматичний аналіз трафіку в реальному часі. Сучасні алгоритми глибокого навчання мають здатність опрацьовувати надзвичайно великі обсяги даних про пакети, сесії та запити, оперативно виявляючи підозрілу активність, яка потенційно може свідчити про кібератаку.

Поведінкове моделювання. Розробка моделей поведінки користувачів, пристроїв та програм сприяє виявленню аномалій у звичній діяльності, які часто слугують першими сигналами кібернетичних загроз. Ці моделі регулярно оновлюються на основі нових даних, гарантуючи адаптивність системи.

Прогнозування аномалій. Застосування методів прогнозової аналітики дає змогу передбачити ймовірність виникнення аномальної активності, аналізуючи історичні дані та шаблони поведінки. Це дозволяє завчасно уникнути загроз, перш ніж вони стануть реальністю.

Інтеграція з системами автоматизованого реагування (SOAR). Після виявлення аномалій системи здатні автоматично запускати сценарії реагування, наприклад, блокування підозрілих IP-адрес або ізоляція уражених пристроїв, таким чином мінімізуючи потенційний вплив атак.

Зменшення хибних спрацьовувань. Завдяки застосуванню алгоритмів класифікації та кластеризації подій, системи аналізу аномалій на базі ШІ скорочують кількість помилкових сповіщень, що дозволяє аналітикам зосередитися на дійсно важливих загрозах [42].

Перевагами використання ШІ є:

- Значна швидкість обробки і масштабованість: ШІ з легкістю опрацьовує колосальні обсяги інформації в реальному часі, що фізично не під силу людині.
- Здатність до самонавчання та адаптації до нових викликів: МН забезпечує безперервне вдосконалення аналітичних моделей та покращення виявлення загроз.
- Можливість превентивного виявлення багатовекторних атак: алгоритми ШІ здатні розпізнавати складні, багаторівневі атаки, які надзвичайно важко виявити звичними методами.
- Зменшення навантаження на аналітиків SOC: Автоматизація аналізу та реагування звільняє аналітиків від необхідності займатися великою кількістю рутинних сповіщень, даючи змогу зосередитись на більш важливих задачах [43].

Розглянемо найбільш відомі приклади успішного використання ШІ та МН.

Darktrace. Ця система реалізує концепцію «цифрової імунної системи», яка базується на алгоритмах МН. Вона формує поведінкові профілі користувачів, пристроїв та мережевих складових, визначаючи базові сценарії їх функціонування. Завдяки цьому Darktrace здатна автоматично виявляти відхилення від звичної поведінки, які можуть свідчити про спроби несанкціонованого доступу, зараження шкідливим ПЗ або внутрішні загрози. Система функціонує у реальному часі та використовує механізми самонавчання для постійного покращення [44].

Vectra AI. Пропонує інструменти для аналізу даних про поведінку і трафік мережі, що допомагають виявляти комплексні та багатовекторні кіберзагрози. Завдяки алгоритмам глибокого навчання, система Vectra AI аналізує активність користувачів та обладнання, визначаючи аномалії, типові для атак на різних етапах: від розвідки до встановлення командних центрів або нелегального виведення інформації. Це робить рішення особливо ефективним у виявленні атак, що здатні обійти традиційні системи захисту [45].

Cisco Secure Network Analytics (колишня назва Stealthwatch). Пропонує комплексний нагляд за мережевим рухом, враховуючи навіть зашифровані дані, без необхідності їх розшифрування. Застосовуючи МН та евристичні методи, система виявляє аномалії в поведінці мережевих потоків, які можуть свідчити про зловмисну діяльність або витік даних. Cisco Secure Network Analytics також дає змогу фахівцям з безпеки отримувати ґрунтовне розуміння трафіку, що спрощує розслідування інцидентів і підвищує рівень безпеки [46].

У таблиці 3.2 наведено приклади сучасних рішень на основі ШІ, які успішно застосовуються для аналізу та виявлення аномалій у мережевому трафіку.

Таблиця 3.2.

Основні інструменти AI для виявлення аномалій у мережевому трафіку

Інструмент	Основна функція	Ключові можливості
Darktrace	Виявлення відхилень у поведінці	Самонавчання, цифрова імунна система
Vectra AI	Виявлення багатовекторних атак	Поведінковий аналіз, глибоке навчання
Cisco Secure Network Analytics	Моніторинг трафіку, включаючи зашифрований	Виявлення аномалій без дешифрування

Отже, використання ШІ у виявленні аномалій в мережевому трафіку - це одне з найефективніших рішень сучасної кібербезпеки. Завдяки здатності автоматично аналізувати великі обсяги даних в реальному часі, моделювати поведінку та використовувати прогнозну аналітику, системи на базі ШІ здатні

своєчасно виявляти як відомі, так і нові загрози, що є нездійсненним для традиційних методів виявлення.

Алгоритми глибокого навчання, які інтегровані в ці системи, забезпечують високу точність та швидкість виявлення аномалій, допомагаючи мінімізувати ризик компрометації даних і звести до мінімуму хибні спрацювання. Поведінкові моделі користувачів і пристроїв постійно оновлюються, забезпечуючи адаптивність систем до нових типів атак.

Інтеграція з автоматизованими системами реагування (SOAR) дає можливість не тільки виявляти аномалії, а й миттєво реагувати на інциденти без потреби людського втручання. Це суттєво зменшує навантаження на аналітиків безпеки та покращує загальну ефективність діяльності центрів операцій безпеки (SOC).

Використання таких платформ, як Darktrace, Vectra AI та Cisco Secure Network Analytics, демонструє конкретні випадки успішного впровадження ШІ для забезпечення мережевої безпеки організацій та об'єктів критичної інфраструктури.

Відтак, ШІ перетворюється з додаткового інструменту у стратегічному наборі кіберзахисту на вкрай важливий елемент стратегії гарантування стійкості та безпеки ІКС в умовах безперервного зростання кіберзагроз.

3.4 Захист критичних систем: прогнозування атак і превентивні заходи

ШІ у прогнозуванні кібератак. У цифрову епоху критичні інформаційні та технологічні системи (енергетичні мережі, атомні електростанції, транспортні розв'язки, фінансові інституції, медичні заклади й державні інформаційні ресурси) перетворилися на основні складники функціонування сучасного суспільства. Вони гарантують безперебійну роботу економіки, забезпечують національну безпеку та суспільну стабільність. Проте, через власну критичність і технологічну складність, ці системи все частіше привертають увагу

кіберзлочинців, хактивістів і державних актори, які використовують високотехнологічні методи атак.

Характерною рисою сучасних кібератак є їхня комплексна природа, багатогранність і вміння уникати звичних систем захисту. Вони використовують недоліки не лише в програмах, а й в організаційних підходах та людській поведінці. Звичні реагуючі методи кіберзахисту, що зосереджені на ліквідації вже здійснених атак, є недостатньо ефективними. У цьому сенсі актуальним є перехід до випереджальних стратегій, котрі передбачають прогнозування потенційних загроз ще до їх реалізації, та запровадження дієвих превентивних заходів для зменшення ризиків і збереження важливих активів [9].

Основні напрями застосування ШІ для прогнозування кібератак охоплюють (Рис. 3.4):



Рис. 3.4. Основні напрями застосування ШІ для прогнозування кібератак

- *Прогнозна аналітика* (Predictive Analytics). Алгоритми оперують значними обсягами історичних даних про кіберінциденти. До їхнього числа належать часові особливості, способи атак, характерні вектори проникнення та сценарії розгортання подій. Ці моделі дають змогу не тільки розпізнавати

індикатори компрометації, а й прогнозувати сценарії загроз, що можуть виникнути у майбутньому. Прогностична аналітика сприяє організаціям у раціональному використанні ресурсів безпеки, своєчасному виявленню ризиків і розробці превентивних стратегій захисту.

- *Обробка великих даних (Big Data Analytics)*. Платформи ІІІ проводять аналіз структурованих і неструктурованих даних з різних джерел: файли журналів (лог-файли), телеметрія мережевих і кінцевих пристроїв, дані з систем моніторингу трафіку, профілі поведінки користувачів, інформація від Threat Intelligence, OSINT, а також відомості з Dark Web. Це сприяє виявленню складних взаємозв'язків і залежностей, що здатні вказувати на підготовку або вже активну стадію кібератак.

- *Поведінкове моделювання загроз (Behavioral Threat Modeling)*. Системи проводять аналіз типової діяльності користувачів, обладнання та застосунків, формуючи профілі активності. У випадку виявлення відхилень від стандартних зразків, скажімо, підозрілих спроб авторизації або аномальної мережевої взаємодії, система розпізнає це як можливу загрозу. Такий підхід сприяє виявленню новітніх та важко помітних атак, для яких не існує заздалегідь заданих сигнатур і які не реєструються класичними інструментами безпеки.

- *Застосування глибокого навчання (Deep Learning)*. Нейронні мережі, зокрема рекурентні нейронні мережі (RNN) та моделі довготривалої короткочасної пам'яті (LSTM), досліджують часові ряди поведінкових даних. Вони здатні виявляти приховані патерни і закономірності, які можуть вказувати на ранні ознаки кіберзагроз. Такі моделі особливо ефективні для передбачення багатоступеневих атак APT, де зловмисники діють повільно та обережно [47].

Розглянемо найбільш показові приклади застосування ІІІ для прогнозування кібератак.

IBM QRadar Advisor with Watson є рішенням, яке інтегрує когнітивні обчислення, МН та обробку природної мови (NLP) з метою аналізу значних обсягів даних. Система прогнозує можливі кіберзагрози та генерує поради для

фахівців з кібербезпеки, забезпечуючи оперативне виявлення складних атак, а також скорочення часу, необхідного для реагування на них [48].

Платформа *Recorded Future* поєднує ШІ з аналізом загроз, застосовуючи дані розвідки загроз (Threat Intelligence). Вона використовує інформацію з відкритих джерел (OSINT) і мережі Dark Web, здійснює аналіз тенденцій загроз і типових моделей поведінки з метою виявити ознаки підготовки до кібернападів на ранніх етапах. Таким чином організації можуть оперативно реагувати і вживати необхідних заходів [49].

Система *FireEye Helix* автоматично здійснює збір та кореляцію даних про загрози з відомими вразливостями. Її функція полягає у передбаченні ймовірних сценаріїв атак, а також у забезпеченні швидкого реагування. Це досягається за рахунок використання МН, яке аналізує велику кількість джерел інформації [50].

Cortex XDR (Palo Alto Networks) використовує надзвичайні можливості аналізу великих даних, МН і поведінкового аналізу для виявлення й передбачення складних кібератак. Система досліджує мережевий трафік, активність кінцевих пристроїв та хмарні сервіси, гарантуючи проактивний захист організацій [51].

Як свідчить практика прогнозування кібератак, застосування ШІ має низку переваг, серед яких забезпечення проактивного захисту (запобігання кібератакам); висока швидкість аналізу великих обсягів даних в реальному часі; зменшення навантаження на персонал в результаті автоматизації поточних завдань; зменшення хибнопозитивних спрацювань завдяки вдосконаленим алгоритмам класифікації та кластеризації даних; висока адаптивність як наслідок здатності ШІ до самонавчання.

Застосування ШІ у превентивних заходах. Ефективний кіберзахист критичних систем є неможливим без запобіжних кроків, що дають змогу зменшувати ризики задовго до виникнення активних загроз. Впровадження ШІ дозволяє автоматизувати й оптимізувати ці заходи, тим самим підвищуючи опірність систем до потенційних атак.

До найбільш поширених превентивних стратегій із застосуванням ІШ експерти відносять (Рис. 3.5):



Рис. 3.5. Основні превентивні стратегії із застосуванням ІШ

Автоматизоване оновлення й управління захисними політиками. Засоби ІШ ретельно відстежують зміни у ландшафті кіберзагроз, автоматично поновлюють сигнатури атак, правила зіставлення подій та політик безпеки, не вимагаючи втручання людей. Це дозволяє організаціям підтримувати захист на високому рівні та блискавично реагувати на нові загрози, зменшуючи «вікно вразливостей».

Інтелектуальне управління доступом. Замість звичайних правил доступу, ІШ бере до уваги контекст: поведінку користувача, історію доступів, тип пристрою, геолокацію і час активності. Системи використовують поведінкову аналітику для динамічного ухвалення рішень щодо надання чи відмови в доступі. Це мінімізує ризик компрометації облікових записів і реалізує принцип Zero Trust.

Сегментування мережі та мікросегментація. ІШ досліджує потоки трафіку, розпізнаючи шаблони взаємодій між системами, щоб автоматично формувати мережеві сегменти. Мікросегментація зменшує вірогідність подальшого просування зловмисника в ІКС організації-жертви після

проникнення, локалізуючи потенційно загрозливу активність в ізольованих областях інфраструктури.

Прогнозоване управління ресурсами. Алгоритми МН вивчають історію роботи систем, а також сезонні й поведінкові особливості ресурсного використання. Це дає змогу передбачати пікові навантаження й запобігати ситуаціям, які можуть бути використані для атак, наприклад, DDoS, а також забезпечити ефективний розподіл ресурсів і збалансування навантаження.

Автоматизоване навчання персоналу та симуляції атак. Методи ШІ використовуються для відтворення реальних кібернападів і навчання персоналу. Системи автоматично генерують сценарії атак, здійснюють перевірку підготовленості та сприяють вдосконаленню практичних вмінь спеціалістів з безпеки. Цей підхід мінімізує вплив людського фактору на реалізацію кіберризиків.

Інтеграція з аналізом загроз для раннього оповіщення. Системи ШІ поєднують використання зовнішніх даних розвідки загроз, включно з даними з OSINT, Dark Web та комерційних розвідувальних платформ, що дає можливість розпізнавати загальні тенденції, вразливості типу «нульового дня» і сучасні методи атак до їх реалізації [41, 47].

Наведемо кілька прикладів ІТ-продуктів, в яких ШІ застосовується для превентивних заходів.

Система *Darktrace Antigena* є невід'ємною частиною платформи *Darktrace Enterprise Immune System*, що послуговується ШІ для побудови так званої «цифрової імунної системи». *Darktrace Antigena* автоматично проводить аналіз мережевого трафіку, поведінки пристроїв і користувачів в реальному часі. У разі виявлення відхилень, що можуть сигналізувати про кіберзагрозу, система самостійно здійснює обмеження або блокування зловмисної активності. Завдяки цьому організації мають можливість зменшити час реагування на інциденти і запобігти подальшому розвитку атак без безпосереднього залучення аналітиків [52].

Microsoft Defender for Endpoint використовує алгоритми МН і ШІ, що забезпечують безперервний моніторинг кінцевих точок. Система прогнозує ймовірні шкідливі дії, автоматично блокуючи підозрілі процеси та ізолюючи уражені пристрої, перш ніж інцидент вийде з-під контролю. Продукт аналізує поведінкові патерни, ідентифікує нові вектори атак і пропонує дії для усунення вразливостей [53].

Платформа *Cisco SecureX* пропонує цілісне інтегроване середовище для аналізу великих масивів даних і використання ШІ з метою автоматизації реагування. SecureX дає змогу ідентифікувати можливі загрози, автоматично підсилювати політики доступу та блокувати активність, що виходить за межі очікуваних поведінкових шаблонів. Ця система також об'єднує різні джерела інформації та використовує ШІ для оптимізації взаємодії між підрозділами безпеки [54].

CrowdStrike Falcon застосовує аналіз поведінки і технології МН для виявлення складних, багатовекторних атак. Система здійснює аналіз активності на кінцевих точках і в мережевому трафіку, виявляючи навіть ті загрози, які здатні обходити традиційні антивірусні та сигнатурні системи захисту. Falcon не лише дає можливість своєчасно виявляти атаки, але й автоматично блокує підозрілу активність та надає рекомендації для подальшого реагування на інциденти [55].

В умовах постійно наростаючих загроз в кіберпросторі та все складнішого ландшафту атак, захист критичних інформаційних та технологічних систем вимагає випереджаючого, адаптивного реагування. Інтегрування ШІ у стратегії кібербезпеки надає можливість перейти від традиційних реактивних моделей до проактивного, прогнозованого захисту.

Завдяки потенціалу ШІ в обробці великих обсягів даних, побудові поведінкових моделей, прогнозній аналітиці та глибокому навчанню, організації отримують у своє розпорядження інструменти, здатні виявляти складні багатопарові загрози на їхніх ранніх стадіях. ШІ не лише покращує ефективність

виявлення загроз, але й гарантує швидку відповідь, зменшує кількість помилкових спрацьовувань та оптимізує використання людських ресурсів.

Превентивні засоби, автоматизований контроль доступу, сегментування мережі, адаптивні політики реагування та інтеграція з глобальними системами аналізу загроз - це фундамент багаторівневого, надійного захисту. Практичні приклади, як Darktrace Antigena, Microsoft Defender for Endpoint, Cisco SecureX та CrowdStrike Falcon, підтверджують практичну користь від впровадження ШІ для гарантування безпеки важливих систем.

Отже, застосування ШІ не є просто модною тенденцією, а стратегічно важливим кроком для організацій, які ставлять за мету забезпечити стійкість своїх інформаційних ресурсів і захистити критичну інфраструктуру від теперішніх та майбутніх кіберзагроз [10].

3.5 Виклики та ризики використання ШІ у кібербезпеці

В епоху, коли ШІ все активніше перетворюється на вагомий чинник ефективного кіберзахисту, відкриваючи небачені перспективи для розпізнавання, нейтралізації та протидії атакам, його використання не позбавлене викликів і проблем. Застосування ШІ в стратегіях безпеки потребує не тільки усвідомлення потенційних плюсів, а й розуміння ризиків, що можуть підірвати стабільність та довіру до «розумних» систем.

Дослідження показало, що серед основних викликів застосування ШІ в кібербезпеці виділяють: можливість помилкових спрацьовувань, критичну залежність від якості отриманих даних, чутливість моделей ШІ до зовнішніх впливів, а також проблеми забезпечення прозорості процесу прийняття рішень і нормативної відповідності (Рис. 3.6) [56].



Рис. 3.6. Виклики та ризики застосування ІШІ в кібербезпеці

Розглянемо кожен із показаних викликів застосування ІШІ детальніше.

Хибнопозитивні та хибнонегативні результати. Незважаючи на те, що ІШІ суттєво вдосконалює здатність виявляти кібернетичні загрози, він не дає абсолютної гарантії. У результаті хибнопозитивних спрацювань (False Positives), коли система оцінює законну активність як підозрілу, може виникати надмірна кількість сповіщень, що спричиняє перевантаження аналітиків і, потенційно, пропуск важливих інцидентів. Хибнонегативні (False Negatives) спрацювання, коли реальні атаки залишаються не розпізнаними, призводять до зростання ризиків непомітного проникнення злоумисників у систему.

Залежність від якості даних. ІШІ не здатен перевершити якість тих даних, на основі яких відбувається навчання та подальша робота. Ефективність алгоритмів безпосередньо залежить від повноти та репрезентативності даних (дані мають віддзеркалювати всі ймовірні сценарії загроз); їх актуальності (дані мають постійно оновлюватися, використання застарілих даних може перешкодити виявленню новітніх методів атак); відсутність спотворень (упереджені і нерепрезентативні дані можуть спричинити хибні висновки, наприклад недооцінку окремих типів атак).

Вразливість моделей до атак. Попри власні потужні можливості, системи ІІІ самі можуть стати мішенями атак кіберзлочинців, зокрема атака «отруєння даних» (Data Poisoning), коли зловмисники свідомо вводять шкідливі чи викривлені дані у процес навчання моделі; зловмисна атака (Adversarial attack) - зловмисники вносять зміни до вхідних даних, які практично непомітні для людини, внаслідок чого модель починає демонструвати хибні результати; інверсія моделі (Model Inversion) – атаки націлені на відтворення (реконструкцію) навчальних даних з використанням самої моделі, що може призвести до витоку чутливої інформації, зокрема паролів, персональних даних або комерційної таємниці.

Непрозорість і відсутність пояснень (Explainability). Системи ІІІ, особливо ті, що побудовані на основі глибокого навчання, часто функціонують як «чорні скриньки». Це означає, що процес прийняття рішень є складним і непрозорим навіть для розробників і користувачів, що призводить до труднощів при аналізі рішень (без використання спеціалізованих засобів проблематично з'ясувати, чому система оцінила певну подію як потенційну загрозу або санкціонувала певну операцію); виникнення проблем у ході проведення аудиту і встановлення відповідальності (важко пояснити аудитору, начальнику служби безпеки або регулятору рішення системи щодо хибних спрацьовувань або помилкового блокування користувачів); неможливості виправлення: коли модель виносить хибні рішення, для їхнього коригування необхідно повністю або частково переглянути систему.

Етичні та правові виклики. Хоча ІІІ надає численні переваги у сфері кібербезпеки, його застосування створює нові етичні дилеми та правові питання. Насамперед це пов'язано з конфіденційністю даних користувачів. Для навчання та функціонування систем ІІІ необхідні значні масиви інформації, нерідко включаючи особисті або секретні дані. Недотримання правил поведінки з такими даними може призвести до порушення права користувачів на захист конфіденційної інформації, а також вимог нормативних документів, наприклад GDPR.

Можливість систем ІІІ автоматизовано приймати рішення і накладати обмеження, наприклад блокувати облікові записи, обмежувати доступ до певних ресурсів або застосовувати інші санкції без участі людини, може призвести до хибних рішень щодо користувачів як наслідок недосконалості алгоритмів чи моделей. Нормативно-правова невизначеність внаслідок відставання законодавства багатьох держав від технологічного поступу також створює низку проблем, зокрема щодо того, хто відповідальний за недоліки ІІІ: розробник, користувач чи власник системи.

Вартість впровадження та обслуговування. Попри значний потенціал ІІІ в оптимізації процесів кіберзахисту, повномасштабне впровадження таких рішень вимагає значних ресурсів, в тому числі на розробку, придбання або ліцензування рішень на основі ІІІ, придбання обладнання (серверів, обчислювальних потужностей), яке потрібне для опрацювання великих масивів даних; підготовку й навчання персоналу, який має мати відповідну кваліфікацію і систематично покращувати професійний рівень у зв'язку з динамічним розвитком технологій; утримання та модернізація (засоби ІІІ потребують систематичного оновлення вихідної інформації, повторного навчання та забезпечення актуальності); складність інтеграції рішень ІІІ в наявну інфраструктуру, що може зумовлювати необхідність реорганізації поточних процесів, або й формування нових підрозділів [5, 6, 56].

З огляду на наявність перелічених вище викликів, слід наголосити на важливості обдуманого і стратегічного впровадження ІІІ у сферу кібербезпеки, забезпечення постійного контролю та модернізації, а також навчання персоналу ефективній роботі в динамічному і непередбачуваному технологічному просторі.

Саме усвідомлення обмежень і потенційних загроз є фундаментом для розробки зваженої та стійкої стратегії кіберзахисту з максимальним використанням позитивних рис ІІІ і мінімізацією його потенційних негативних наслідків.

Висновки до розділу 3

У розділі 3 досліджено засади практичного використання ШІ в галузі кібербезпеки. Встановлено, що основними інструментами кібербезпеки, які використовують ШІ, є засоби поведінкової аналітики користувачів, оркестрації та автоматизації реагування на інциденти, виявлення та запобігання вторгненням, прогнозування й аналізу загроз, класифікація та аналізу мережевого трафіку, а також архітектура нульової довіри.

Детально розглянуто існуючі інструменти кібербезпеки на основі ШІ від Cisco, Exabeam, IBM, Microsoft, Splunk та інших виробників світового рівня з акцентом на автоматизації аналізу загроз з використанням SIEM-систем, застосуванні ШІ для виявлення аномалій у мережевому трафіку, прогнозування атак і впровадження превентивних заходів для захисту критичних систем.

Дослідження засвідчило, що впровадження ШІ дає організаціям і об'єктам критичної інфраструктури вагомі переваги, серед яких проактивне виявлення загроз, значна швидкість обробки та масштабованість, здатність до самонавчання й висока адаптивність до нових загроз, зменшення робочого навантаження на фахівців з кібербезпеки. Водночас, виявлено проблеми і ризики, пов'язані з використанням ШІ, зокрема, помилкові спрацьовування та пропуски загроз, залежність систем ШІ від якості вихідних даних, чутливість моделей до зовнішніх впливів, непрозорість прийняття рішень, етичні й правові обмеження, а також значні фінансові витрати на впровадження та супровід рішень ШІ.

Підсумовуючи, слід відзначити, що ШІ сьогодні відіграє вирішальну роль у підготовці й реалізації комплексних заходів кіберзахисту, впровадження інноваційних підходів до запобігання кіберзагрозам і, як наслідок забезпечення кіберстійкості підприємств та об'єктів критичної інфраструктури.

Однак, враховуючи наявність багатьох ризиків, пов'язаних із застосуванням ШІ, важливо забезпечити обдумане і стратегічно виважене впровадження ШІ у сферу кібербезпеки, гарантувати постійний контроль і оновлення «розумних» систем, а також проводити систематичне навчання

персоналу ефективній роботі з інструментами ІІІ. За таких умов організації зможуть з максимальною віддачею використати переваги ІІІ і мінімізувати його потенційно негативні наслідки.

ВИСНОВКИ

Дослідження показало, що в умовах динамічного кіберсередовища, постійного ускладнення й розширення масштабів атак з Інтернету кібербезпека потребує впровадження інноваційних підходів, зокрема методів штучного інтелекту, таких як машинне навчання (ML), глибоке навчання (DL), аналіз великих даних (Big Data Analytics), обробка природної мови (NLP) та інші.

Машинне навчання відіграє ключову роль у системах аналізу поведінки користувачів, створюючи поведінкові моделі для ідентифікації потенційно небезпечних дій, виявлення відхилень і запобігання кібератакам. Технології глибокого навчання забезпечують виявлення загроз у мережевому трафіку, розпізнавання атак соціальної інженерії, аналіз шкідливого ПЗ, автоматизацію реагування на загрози. Технології аналізу великих даних дозволяють корелювати події безпеки, прогнозувати загрози, оптимізувати захисні механізми.

Встановлено, що технології ШІ змінюють традиційні підходи до кіберзахисту, підвищуючи їх ефективність і сприяючи переходу від реактивного до проактивного підходу до забезпечення кібербезпеки. ШІ дозволяє системам кібербезпеки не лише реагувати на загрози, але й прогнозувати і самостійно адаптуватися до потенційних атак, успішно протидіяти загрозам нульового дня, складним багатовекторним атакам і кіберзлочинам із використанням соціальної інженерії. Методи ШІ сприяють зменшенню кількості помилкових спрацювань, підвищують швидкість виявлення загроз і мінімізують шкоду від кібератак.

У результаті дослідження практичних аспектів застосування методів ШІ для забезпечення кібербезпеки об'єктів критичної інфраструктури на прикладі Рівненської АЕС встановлено, що об'єкти атомної енергетики України постійно стикаються з кіберзагрозами високого рівня складності, які становлять потенційну небезпеку для національної безпеки. Основними характеристиками забезпечення кібербезпеки АЕС є зонування мереж, розмежування ІТ та ОТ, високу критичність впровадження змін, обмеження на пряме оновлення систем, а також постійний контроль з боку державних та міжнародних регуляторів.

Проаналізовано діючі компоненти системи кіберзахисту РАЕС, серед яких рішення SIEM, навчання персоналу, політики обмеженого доступу, сегментація мереж та інструментів моніторингу. Підвищенню ефективності цих механізмів кібербезпеки АЕС сприяє впровадження методики застосування ІІІ, яка містить такі елементи: аналіз поведінки користувачів і систем, автоматизоване виявлення вторгнень, прогнозування та запобігання атакам, розширений аналіз логів і великих даних, реалізація архітектури нульової довіри, автоматизація реагування на інциденти, аналіз вразливостей та інтелектуальна класифікація трафіку.

Вивчення засад практичного використання ІІІ в галузі кібербезпеки, в тому числі сучасних засобів кіберзахисту на основі ІІІ від Cisco, Exabeam, IBM, Microsoft, Splunk та інших виробників, показало, що основними інструментами кібербезпеки зі ІІІ є засоби поведінкової аналітики користувачів, оркестрації та автоматизації реагування на інциденти, виявлення та запобігання вторгненням, прогнозування й аналізу загроз, класифікації й аналізу мережевого трафіку тощо.

Встановлено, що впровадження ІІІ дає організаціям і ОКІ вагомі переваги, серед яких проактивне виявлення загроз, висока швидкість обробки і масштабованість, здатність до самонавчання й адаптивність до нових загроз, зменшення навантаження на фахівців з кібербезпеки. Водночас, виявлено, що використання ІІІ викликає появу низки проблем, зокрема, пропуски загроз і помилкові спрацювання, залежність систем від якості вихідних даних, чутливість моделей до зовнішніх впливів, непрозорість прийняття рішень, етичні й правові обмеження, значні фінансові витрати на впровадження та супровід рішень ІІІ.

Беручи до уваги наявність багатьох ризиків, пов'язаних із застосуванням ІІІ, важливо забезпечити обдумане і стратегічно виважене впровадження ІІІ у сферу кібербезпеки, гарантувати постійний контроль і оновлення «розумних» систем, а також проводити систематичне навчання персоналу ефективній роботі з інструментами ІІІ. За таких умов організації та підприємства критичної інфраструктури зможуть з максимальною віддачею використати переваги ІІІ і мінімізувати його потенційно негативні наслідки.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Хорошко В.О., Чередниченко В.С., Шелест М.Є. Основи інформаційної безпеки / За ред. проф. В.О. Хорошка. К.: ДУІКТ, 2008. 186 с.
2. Information Security: The Ultimate Guide. *Imperva*. URL: <https://www.imperva.com/learn/data-security/information-security-infosec/>.
3. What is cybersecurity? *IBM*. URL: <https://www.ibm.com/topics/cybersecurity>.
4. What is a Zero Trust network? *Cloudflare*. URL: <https://www.cloudflare.com/learning/security/glossary/what-is-zero-trust/>
5. The Rise of Artificial Intelligence in Cybersecurity: The Benefits and Drawbacks. Cyber Security News, 2024. URL: <https://cybersecuritynews.com/rise-of-ai-in-cybersecurity/>.
6. Tao, Feng, Akhtar, Muhammad, Jiayuan, Zhang. The future of Artificial Intelligence in Cybersecurity: A Comprehensive Survey. EAI Endorsed Transactions on Creative Technologies. 2021. URL: https://www.researchgate.net/publication/353046785_The_future_of_Artificial_Intelligence_in_Cybersecurity_A_Comprehensive_Survey.
7. Striuk, Oleksandr & Kondratenko, Yuriy. Generative Adversarial Networks in Cybersecurity: Analysis and Response. 2023. URL: https://www.researchgate.net/publication/370078494_Generative_Adversarial_Networks_in_Cybersecurity_Analysis_and_Response
8. Машталяр Я., Козачок В., Бржезьська З., Богданов О., Оксанич І., Литвинов В. Дослідження розвитку та інновації кіберзахисту на об'єктах критичної інфраструктури. *Кібербезпека: освіта, наука, техніка*, 2023. № 2(22). С. 156–167. URL: <https://doi.org/10.28925/2663-4023.2023.22.156167>.
9. Big Data and AI in Government Cybersecurity. *National Cybersecurity Center*, 2023. URL: <https://www.ncsc.gov.uk/search?q=big%20data>
10. The NIST Cybersecurity Framework (CSF) 2.0. *National Institute of Standards and Technology*. 2024. URL: <https://nvlpubs.nist.gov/nistpubs/CSWP/NIST.CSWP.29.pdf>.

11. Lev Craig, Linda Tucci. What is machine learning? Guide, definition and examples. *Techtarget*. URL: <https://www.techtarget.com/searchenterpriseai/definition/machine-learning-ML>.
12. Simon Tavasoli. The Machine Learning Algorithms List: Types and Use Cases. *Simplilearn*. URL: <https://www.simplilearn.com/10-algorithms-machine-learning-engineers-need-to-know-article>
13. Shanika Wickramasinghe. Behavioral Analytics in Cybersecurity. *Splunk*. URL: https://www.splunk.com/en_us/blog/learn/behavioral-analytics.html
14. What is Deep Learning? A Tutorial for Beginners. *Datacamp*. URL: <https://surl.li/uescko>
15. Angin, Pelin, Bhargava, Bharat, Ranchal, Rohit. Big Data Analytics for Cyber Security. *Security and Communication Networks*. 2019. 1-2. URL: https://www.researchgate.net/publication/335698795_Big_Data_Analytics_for_Cyber_Security
16. Про інформацію: Закон України від 02.10.1992 № 2657-XII. URL: <https://zakon.rada.gov.ua/laws/show/2657-12>
17. Про Національну програму інформатизації: Закон України від 01.12.2022 №2807-IX. URL: <https://zakon.rada.gov.ua/laws/show/2807-20#Text>
18. Про захист інформації в інформаційно-комунікаційних системах: Закон України від 05.07.1994 № 80/94-ВР. URL: <https://zakon.rada.gov.ua/laws/show/80/94-%D0%B2%D1%80#Text>
19. Про основні засади забезпечення кібербезпеки України: Закон України від 05.10.2017 № 2163-VIII. URL: <https://zakon.rada.gov.ua/laws/show/2163-19#Text>
20. Про критичну інфраструктуру: Закон України від 20.10.2023 № 3430-IX. URL: <https://zakon.rada.gov.ua/laws/show/1882-20>
21. Про рішення Ради національної безпеки і оборони України від 14 травня 2021 року «Про Стратегію кібербезпеки України»: Указ Президента України від 26.08.2021 № 447/2021. URL: <https://zakon.rada.gov.ua/laws/show/447/2021#Text>

22. Про схвалення Концепції створення державної системи захисту критичної інфраструктури: Розпорядження Кабінету Міністрів України від 06.12.2017 № 1009-р. URL: <https://zakon.rada.gov.ua/laws/show/1009-2017-%D1%80#Text>

23. Деякі питання об'єктів критичної інфраструктури: Постанова Кабінету Міністрів України від 09.10.2020 № 1109. URL: <https://zakon.rada.gov.ua/laws/show/1109-2020-%D0%BF>

24. Про рішення Ради національної безпеки і оборони України від 17 жовтня 2023 року «Про організацію захисту та забезпечення безпеки функціонування об'єктів критичної інфраструктури та енергетики України в умовах ведення воєнних дій»: Указ Президента України від 18.10.2023 № 691/2023.

25. Офіційний сайт ДССЗІ України. URL: <https://cip.gov.ua/>.

26. Офіційний сайт Служби безпеки України. URL: <https://ssu.gov.ua/>.

27. Офіційний сайт Кіберполіції України. URL: <https://cyberpolice.gov.ua/>.

28. Офіційний сайт Національного координаційного центру кібербезпеки при РНБО України. URL: <https://cip.gov.ua/ua/natsionalnyi-koordinatsiinyi-tsentr-kiberbezpeky>.

29. IAEA. Nuclear Security Series No. 17. Computer Security at Nuclear Facilities. – International Atomic Energy Agency, Vienna, 2011. URL: <https://www.iaea.org/publications/8906/computer-security-at-nuclear-facilities>.

30. ISO/IEC 27001:2022. Information security, cybersecurity and privacy protection – Information security management systems – Requirements. – International Organization for Standardization, 2022. URL: <https://www.iso.org/standard/27001>.

31. European Union Agency for Cybersecurity (ENISA). Guidelines and standards. URL: <https://www.enisa.europa.eu/>.

32. Внутрішні нормативні документи ВП «Рівненська АЕС» ДП «НАЕК «Енергоатом» щодо вимог до інформаційної безпеки на об'єктах критичної інфраструктури. [Внутрішній документ підприємства. Без публічного доступу].

33. Інформація про потужності ВП «Рівненська АЕС». URL: <https://energoatom.com.ua/about>

34. НП 306.2.202-2015 Вимоги з ядерної та радіаційної безпеки до інформаційних та керуючих систем, важливих для безпеки атомних станцій. URL: https://online.budstandart.com/ua/catalog/doc-page.html?id_doc=68978

35. Внутрішні нормативні документи ВП «Рівненська АЕС» ДП «НАЕК «Енергоатом» щодо організації кіберзахисту. 2023–2024. [Внутрішній документ підприємства. Без публічного доступу].

36. How UEBA Works | User and Entity Behavior Analytics. URL: <https://www.forcepoint.com/cyber-edu/user-and-entity-behavior-analytics-ueba>.

37. Котляров О. Ю., Бортнік Л. Л. Огляд фундаментальної моделі безпечної оркестрації, автоматизації та реагування в контексті кібербезпеки віртуальних мереж. *Computer Systems and Networks* Vol. 7, No. 1, 2025. URL: <https://science.lpnu.ua/sites/default/files/journal-paper/2025/may/38970/250473csn-162-177.pdf>

38. Keith Thomas. Understanding AI risks and how to secure using Zero Trust. *Levelblue*. URL: <https://levelblue.com/blogs/security-essentials/understanding-ai-risks-and-how-to-secure-using-zero-trust>

39. AI SIEM: The 6 Components of AI-Based SIEM. *Stellarcyber*. URL: <https://stellarcyber.ai/learn/ai-driven-siem/>

40. What Is the Role of AI and ML in Modern SIEM Solutions? *Paloaltonetworks*. URL: <https://www.paloaltonetworks.com/cyberpedia/role-of-artificial-intelligence-ai-and-machine-learning-ml-in-siem>

41. Olivia Henderson. Splunk Threat Intelligence Management. *Splunk* URL: https://www.splunk.com/en_us/blog/learn/splunk-threat-intelligence-management.html

42. Aluwala, Aakash. AI-Driven Anomaly Detection in Network Monitoring Techniques and Tools. *Journal of Artificial Intelligence & Cloud Computing*. 2024. 1-6. URL: https://www.researchgate.net/publication/381845747_AI-Driven_Anomaly_Detection_in_Network_Monitoring_Techniques_and_Tools

43. Jon Hencinski. AI-Driven SOC: The Next Evolution in Security Operations. URL: <https://www.prophetsecurity.ai/blog/ai-soc-key-to-solving-persistent-soc-challenges>

44. Darktrace. ActiveAI Security Platform. URL: <https://www.darktrace.com/>
45. Vectra AI. The world leader in AI security. URL: <https://www.vectra.ai/>
46. Cisco Secure Network Analytics. *Cisco*. URL: <https://www.cisco.com/site/us/en/products/security/security-analytics/secure-network-analytics/index.html>
47. AI in Threat Intelligence. URL: <https://www.silobreaker.com/glossary/ai-in-threat-intelligence/>
48. QRadar Advisor with Watson app. *IBM*. URL: <https://www.ibm.com/docs/en/qradar-common?topic=apps-qradar-advisor-watson-app>
49. Recorded Future: Advanced Cyber Threat Intelligence. URL: <https://www.recordedfuture.com/>
50. FireEye Helix: Overview. *FireEye*. URL: <https://fireeye.dev/docs/helix/>
51. Transform Endpoint Security with Cortex XDR. *Palo Alto Networks*. URL: <https://www.paloaltonetworks.com/cortex/cortex-xdr>
52. Darktrace Antigena. URL: <https://www.cyberaiworks.com/datasheets/ds-antigena.pdf>
53. Microsoft Defender for Endpoint. *Microsoft*. URL: <https://www.microsoft.com/en-us/security/business/endpoint-security/microsoft-defender-endpoint>
54. Cisco SecureX – A Simplified Security Experience. *Cisco*. URL: <https://www.cisco.com/site/id/en/products/security/securex-platform.html>
55. The CrowdStrike Falcon® platform. *CrowdStrike*. URL: <https://www.crowdstrike.com/platform/>
56. Gautam, Ashish, Thapaliya, Suman. A Chronology of AI Failures in Safety and Cybersecurity. *NPRC Journal of Multidisciplinary Research*. 2024. №1. P.1-12. URL: https://www.researchgate.net/publication/386039573_A_Chronology_of_AI_Failures_in_Safety_and_Cybersecurity