

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ

НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ЗАХИСТУ ІНФОРМАЦІЇ
КАФЕДРА УПРАВЛІННЯ ІНФОРМАЦІЙНОЮ ТА КІБЕРНЕТИЧНОЮ
БЕЗПЕКОЮ

КВАЛІФІКАЦІЙНА РОБОТА

на тему: “ СПОСОБИ ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В
УПРАВЛІННІ КІБЕРНЕТИЧНОЮ БЕЗПЕКОЮ ”

на здобуття освітнього ступеня магістра
зі спеціальності 125 Кібербезпека
освітньо-професійної програми Управління інформаційною безпекою

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

Дмитро СОКУРЕНКО
(підпис) Ім'я, ПРІЗВИЩЕ здобувача

Виконав: Здобувач вищої освіти гр. УБДМ 61
Дмитро СОКУРЕНКО

Керівник: К.Т.Н. Юрій ЯКИМЕНКО

Рецензент: д.т.н., професор Галина ГАЙДУР

ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ

Навчально-науковий інститут захисту інформації

Кафедра Управління інформаційною та кібернетичною безпекою

Ступінь вищої освіти магістр

Спеціальність 125 Кібербезпека

Освітньо-професійна програма Управління інформаційною безпекою

ЗАТВЕРДЖУЮ

Завідувач кафедри УІКБ

_____ Світлана ЛЕГОМІНОВА

“ _____ ” _____ 2023 р.

ЗАВДАННЯ НА КВАЛІФІКАЦІЙНУ РОБОТУ

студенту Сокуренку Дмитру Олександровичу

(прізвище, ім'я, по батькові здобувача)

1. Тема кваліфікаційної роботи: “Способи застосування штучного інтелекту в управлінні кібернетичною безпекою”

керівник кваліфікаційної роботи Юрій ЯКИМЕНКО, к.т.н., доцент

(Ім'я, ПРІЗВИЩЕ, науковий ступінь, вчене звання)

затверджені наказом Державного університету інформаційно-комунікаційних технологій від “19” жовтня 2023 р. №145

2. Строк подання кваліфікаційної роботи “18” грудня 2023 р

3. Вихідні дані до кваліфікаційної роботи: *управління кібернетичною безпекою, штучний інтелект, планування кібербезпеки, нормативні документи, стандарти, міжнародні стандарти кібербезпеки, наукова та технічна література*

4. Перелік питань, які потрібно розробити:

1. Розглянути теоретичні засади кібернетичної безпеки.
2. Проаналізувати методологічні підходи до застосування штучного інтелекту в управлінні кібернетичною безпекою.
3. Провести дослідження і розробити рекомендації щодо вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою.

5. Перелік ілюстративного матеріалу: *презентація*

6. Дата видачі завдання “02” жовтня 2023 р.

КАЛЕНДАРНИЙ ПЛАН

№з/п	Назва етапів кваліфікаційної роботи	Строк виконання етапів роботи	Примітка
1.	Визначення об'єкту, предмету, мети та завдань дослідження.	10.10.2023	
2.	Збір та аналіз літератури.	23.10.2023	
3.	Аналіз основних засад кібернетичної безпеки	27.10.2023	
4.	Дослідження методологічних підходів до застосування штучного інтелекту в управлінні кібернетичною безпекою.	10.11.2023	
5.	Розроблення рекомендацій щодо вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою	15.11.2023	
6.	Формулювання висновків за результатами проведеного дослідження.	22.11.2023	
7.	Оформлення роботи.	04.12.2023	
8.	Оформлення презентації.	14.12.2023	
9.	Отримання рецензії на роботу.	18.12.2023	
10.	Захист в ДЕК.	.01.2024	

Здобувач вищої освіти

(*niɔnuc*)

Дмитро СОКУРЕНКО

(Ім'я, ПРІЗВИЩЕ)

Керівник
кваліфікаційної роботи

(niɒnuc)

Юрій ЯКИМЕНКО

(Ім'я, ПРІЗВИЩЕ)

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

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

Направляється здобувач Сокурєнко Д.О., до захисту кваліфікаційної роботи
(прізвище та ініціали)

за спеціальністю 125 Кібербезпека
(код, найменування спеціальності)

Освітньо-професійної програми Управління інформаційною безпекою
(назва)

на тему: “Способи застосування штучного інтелекту в управлінні кібернетичною безпекою ”

Кваліфікаційна робота і рецензія додаються.

Директор ННІЗІ _____
(підпис)

Віталій САВЧЕНКО
(Ім'я, ПРІЗВИЩЕ)

Висновок керівника кваліфікаційної роботи

Здобувач **СОКУРЕНКО Дмитро** у кваліфікаційній роботі дослідив сучасний стан управління кібернетичною безпекою та способи застосування штучного інтелекту в управлінні кібербезпекою. Дослідження розглядає не лише технічні аспекти використання штучного інтелекту, але й враховує етичні проблеми, пов'язані з його застосуванням у сфері кібербезпеки. Розроблена методика та програма автоматизованого виявлення загроз в соціальних мережах дозволяє ефективно реагувати на них керівниками структурних підрозділів кібербезпеки. **СОКУРЕНКО Дмитро** показав розуміння проблеми дослідження та бачення основних теоретичних та практичних напрямів її вирішення, довів уміння самостійного застосування методів наукового дослідження, проявив себе як організований, відповідальний виконавець. Результати дослідження апробовані на науково-практичній конференції.

Це дозволяє оцінити виконану кваліфікаційну роботу здобувача **СОКУРЕНКА Дмитра** на оцінку “відмінно” та присвоїти йому кваліфікацію “Магістр кібербезпеки за освітньо-професійною програмою “Управління інформаційною безпекою””.

Керівник кваліфікаційної роботи _____ Юрій ЯКИМЕНКО
(підпис) (Ім'я, ПРІЗВИЩЕ)

“ ____ ” _____ 2024 року

Висновок кафедри про кваліфікаційну роботу

Кваліфікаційна робота розглянута . Здобувач Сокурєнко Д.О.. допускається до захисту роботи в екзаменаційній комісії

Завідувач кафедрою
Управління інформаційною
та кібернетичною безпекою

_____ Світлана ЛЕГОМІНОВА
(підпис) (Ім'я, ПРІЗВИЩЕ)

ВІДГУК РЕЦЕНЗЕНТА
на кваліфікаційну магістерську роботу

здобувача вищої освіти Сокурєнка Дмитра Олександровича
на тему “ СПОСОБИ ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В
УПРАВЛІННІ КІБЕРНЕТИЧНОЮ БЕЗПЕКОЮ ”

Актуальність. В сучасному цифровому світі зростаюча залежність від штучного інтелекту супроводжується постійною загрозою кібератак та інших форм кіберзлочинності. Разом з тим, штучний інтелект стає важливим інструментом у боротьбі з кіберзагрозами та способом вдосконалення систем управління кібернетичною безпекою організацій. Тому тема дослідження є надзвичайно актуальною.

Позитивні сторони. В роботі правильно визначено, що ефективним способом при організації заходів та створенні системи управління кібербезпекою є застосування штучного інтелекту. Обґрунтовані переваги для організацій і підприємств створення власних моделей виявлення загроз та небезпек з використанням мови програмування Python та її стандартних бібліотек. Запропонований алгоритм створення моделі виявлення загроз, який складає основу нової методики, дозволяє автоматизувати процес управління кібербезпекою організації на початковому етапі. Також використання моделі організаціями і підприємствами середнього бізнесу дозволить оперативно реагувати на нові види загроз.

Недоліки. Разом з тим, доцільно було б застосувати при проведенні дослідження і інші бібліотеки Python для порівняння результатів, але це не зменшує важливості даної роботи та її науковості.

Висновок: Кваліфікаційна робота виконана на належному науково-методичному рівні і заслуговує оцінки «відмінно» а її автор Сокурєнко Дмитро Олександрович - присвоєння кваліфікації «Магістр кібербезпеки» за освітньо-професійною програмою «Управління інформаційною безпекою».

Рецензент: завідувач кафедри
Інформаційної та кібернетичної
безпеки,
д.т.н., професор

підпис

Галина ГАЙДУР

(Ім'я, ПРИЗВИЩЕ)

РЕФЕРАТ

Текстова частина кваліфікаційної роботи на здобуття освітнього ступеня магістра: 87 стор., 19 рис., 1 табл., 67 джерел. 3 додатки.

Метою роботи є дослідження та аналіз способів застосування штучного інтелекту в управлінні кібернетичною безпекою.

Об'єктом дослідження є управління кібернетичною безпекою.

Предмет дослідження – способи і методи застосування штучного інтелекту в управлінні кібернетичною безпекою.

Методи дослідження. Для вирішення визначених наукових завдань в роботі використані наступні методи: системний аналіз, синтез, машинне навчання для виявлення загроз, моделювання, аналіз великих обсягів даних, нейронні мережі для прогнозування атак, методи теорії управління.

Короткий зміст роботи. Здійснено аналіз наукової літератури та публікацій, пов'язаних із сучасним станом та тенденціями розвитку кібернетичної безпеки.

У першому розділі були розглянуті ключові аспекти, що становлять теоретичну базу кібернетичної безпеки та роль штучного інтелекту в управлінні кібербезпекою. Встановлено, що сучасні загрози, від клонованих ідентичностей до діпфейкових кампаній, стає все важче виявити та зупинити, загрози постійно видозмінюються та удосконалюються, що потребує негайної перебудови стратегії кібербезпеки організації. Визначено, що ефективним способом виявлення загроз є застосування штучного інтелекту (ШІ) при організації заходів та створенні системи управління кібербезпекою.

У другому розділі на основі проаналізованих сучасних тенденцій розвитку штучного інтелекту в кібернетичній безпеці встановлено, що майбутнє кібербезпеки буде забезпечене інтеграцією штучного інтелекту в кібербезпеку. Зазначено, що впровадження штучного інтелекту в кібербезпеку вимагає глибокого розуміння та врахування етичних та конфіденційних аспектів, а також постійного спілкування та співпраці з експертами з кібербезпеки. Впровадження передових технологій ШІ в управління кібернетичною безпекою має потенціал

значно поліпшити виявлення та запобігання кіберзагрозам, забезпечуючи високий рівень безпеки в цифровому середовищі.

Обґрунтовані переваги для організацій і підприємств створення власних моделей виявлення з використанням мови програмування Python та її стандартних бібліотек. Запропонований алгоритм створення моделі виявлення, який складає основу нової методики.

Окрему увагу приділено етичним питанням застосування штучного інтелекту у кібербезпеці. Виявлено, що разом із великим потенціалом штучного інтелекту для підвищення рівня безпеки, існують серйозні етичні виклики, такі як конфіденційність даних, автономність систем та можливість виникнення системних помилок.

У третьому розділі на основі запропонованої методики у другому розділі за вказаним алгоритмом вперше створений зразок моделі раннього виявлення загроз за допомогою штучного інтелекту шляхом моніторингу соціальних мереж. Створена на основі стандартних бібліотек Python, за результатами експерименту модель показала досить добрі результати.

Використання моделі організаціями і підприємствами середнього бізнесу дозволить оперативно реагувати на нові види загроз, що змінюються з часом шляхом відкритого доступу до коду та постійного навчання моделі.

Галузь застосування. Розроблені рекомендації фахівцям управління кібербезпекою, керівникам структурних підрозділів інформаційної безпеки щодо удосконалення застосування штучного інтелекту в управлінні кібербезпекою. Разом із запропонованою методикою створення моделей виявлення загроз кібербезпеці рекомендації дозволяють ефективно будувати та управляти кібернетичною безпекою організацій, підприємств.

КЛЮЧОВІ СЛОВА: ІНФОРМАЦІЙНА БЕЗПЕКА, ШТУСНИЙ ІНТЕЛЕКТ, КІБЕРНЕТИЧНА БЕЗПЕКА, УПРАВЛІННЯ ІНФОРМАЦІЙНОЮ БЕЗПЕКОЮ.

ABSTRACT

The text part of the qualification work for obtaining a master's degree: 87 pages, 19 figures, 1 tables, 67 sources, 3 appendix.

The purpose of the work is research and analysis of ways to use artificial intelligence in cyber security management.

Object of research is cyber security management.

Subject of research is ways and methods of using artificial intelligence in cyber security management.

Research methods. The following methods are used to solve the identified scientific tasks in the work: system analysis, synthesis, machine learning for threat detection, modeling, analysis of large volumes of data, neural networks for predicting attacks, control theory methods.

Brief content of research. An analysis of scientific literature and publications related to the current state and trends in the development of cyber security was carried out.

In the first chapter, the key aspects that constitute the theoretical basis of cyber security and the role of artificial intelligence in cyber security management were considered. Today's threats, from cloned identities to deepfake campaigns, have been found to be increasingly difficult to detect and stop, and threats are constantly evolving and evolving, requiring an immediate overhaul of an organization's cybersecurity strategy. It was determined that an effective way to detect threats is the use of artificial intelligence (AI) when organizing events and creating a cyber security management system.

In the second chapter, based on the analyzed current trends in the development of artificial intelligence in cyber security, it is established that the future of cyber security will be ensured by the integration of artificial intelligence in cyber security. It is noted that the implementation of artificial intelligence in cyber security requires a deep understanding and consideration of ethical and confidential aspects, as well as constant communication and cooperation with cyber security experts. The

implementation of advanced AI technologies in cyber security management has the potential to significantly improve the detection and prevention of cyber threats, ensuring a high level of security in the digital environment.

Reasonable benefits for organizations and enterprises to create their own detection models using the Python programming language and its standard libraries. The proposed algorithm for creating a detection model, which forms the basis of the new technique.

Special attention is paid to the ethical issues of using artificial intelligence in cyber security. It has been found that along with the great potential of artificial intelligence to improve security, there are serious ethical challenges, such as data privacy, system autonomy, and the possibility of system errors.

In the third chapter, on the basis of the proposed methodology in the second chapter, according to the specified algorithm, a sample model of early detection of threats using artificial intelligence by monitoring social networks is created for the first time. Created on the basis of standard Python libraries, according to the results of the experiment, the model showed quite good results.

The use of the model by medium-sized business organizations and enterprises will allow prompt response to new types of threats that change over time through open access to the code and constant training of the model. ***Field of research.*** Recommendations have been developed for cyber security management specialists, heads of structural divisions of information security on improving the use of artificial intelligence in cyber security management. Together with the proposed method of creating models for detecting threats to cyber security, the recommendations allow you to effectively build and manage the cyber security of organizations and enterprises.

KEYWORDS: INFORMATION SECURITY, ARTIFICIAL INTELLIGENCE, CYBERNETIC SECURITY, INFORMATION SECURITY MANAGEMENT.

ЗМІСТ

ВСТУП	11
РОЗДІЛ 1. ОГЛЯД ТЕОРЕТИЧНИХ ЗАСАД КІБЕРНЕТИЧНОЇ БЕЗПЕКИ.....	14
1.1 Основні загрози та виклики для кібернетичної безпеки.....	14
1.2 Роль штучного інтелекту в контексті управління кібернетичною безпекою	31
РОЗДІЛ 2. МЕТОДОЛОГІЧНІ ПІДХОДИ ДО ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В УПРАВЛІННІ КІБЕРНЕТИЧНОЮ БЕЗПЕКОЮ	40
2.1 Аналіз сучасних тенденцій використання штучного інтелекту в кібернетичній безпеці	41
2.2 Застосування машинного навчання для виявлення кібератак.....	53
2.3 Етичні проблеми застосування штучного інтелекту у кібербезпеці....	60
РОЗДІЛ 3. ПРАКТИЧНЕ ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В УПРАВЛІННІ КІБЕРНЕТИЧНОЮ БЕЗПЕКОЮ ...	64
3.1 Створення моделі виявлення кіберзагроз з використанням штучного інтелекту	64
3.2 Рекомендації щодо вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою.....	71
ВИСНОВКИ	76
ПЕРЕЛІК ПОСИЛАНЬ	78
ДОДАТКИ	87
ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ	90

ВСТУП

Актуальність теми. Штучний інтелект (ШІ) відіграє важливу роль в сфері кібербезпеки, допомагаючи виявляти загрози, виявляти вразливості та реагувати на кібератаки.

В сучасному цифровому світі, де зростаюча залежність від інформаційних технологій супроводжується постійною загрозою кібератак та інших форм кіберзлочинності, питання забезпечення кібернетичної безпеки стає надзвичайно актуальним та важливим. Щоденно мільйони організацій та індивідів стикаються із загрозами, що походять із великого різноманіття джерел, включаючи атаки на мережеві інфраструктури, фішингові атаки, витoki конфіденційної інформації, і багато інших.

Штучний інтелект стає важливим інструментом у боротьбі з кіберзагрозами та вдосконалення систем управління кібернетичною безпекою. Використання алгоритмів машинного навчання, глибокого навчання та аналізу великих обсягів даних дозволяє ефективно виявляти, прогнозувати та запобігати кібератакам. Застосування ШІ також дозволяє підвищити швидкість реакції на інциденти та створювати адаптивні системи, які самостійно коригуються відповідно до характеру кіберзагроз, що постійно змінюється.

Магістерська робота присвячена вивченню та аналізу способів застосування штучного інтелекту в управлінні кібернетичною безпекою. Дослідження розглядає не лише технічні аспекти використання ШІ, але й враховує етичні проблеми, пов'язані з його застосуванням у сфері кібербезпеки. Робота також ставить за мету визначення оптимальних стратегій та підходів до імплементації та вдосконалення систем управління кібернетичною безпекою за допомогою штучного інтелекту.

Обґрунтування вибору даної теми базується на необхідності впровадження передових технологій для ефективної боротьби із зростаючими загрозами кіберпростору. Розвиток ШІ відкриває нові можливості для розбудови інтелектуальних та адаптивних систем управління кібербезпекою, які мають

потенціал підвищити рівень захисту від кіберзагроз та забезпечити стійкість у глобальному цифровому середовищі.

Робота спрямована на глибоке вивчення проблем кібербезпеки та визначення оптимальних шляхів використання штучного інтелекту для створення надійних та ефективних систем управління кібернетичною безпекою в сучасному інформаційному середовищі.

Об'єктом дослідження є управління кібернетичною безпекою.

Предмет дослідження – способи і методи застосування штучного інтелекту в управлінні кібернетичною безпекою.

Мета роботи полягає у проведенні дослідження та аналізу способів застосування штучного інтелекту в управлінні кібернетичною безпекою.

Для досягнення цієї мети в роботі виконано наступні завдання:

1. Розглянуті теоретичні засади кібернетичної безпеки.
2. Проаналізовані методологічні підходи до застосування штучного інтелекту в управлінні кібернетичною безпекою.
3. Визначені етичні аспекти застосування штучного інтелекту в кібербезпеці.
4. Проведено практичне дослідження і розроблені рекомендації щодо вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою.

Методи дослідження. Для вирішення визначених наукових завдань в роботі використані наступні методи: системний аналіз, синтез, машинне навчання для виявлення загроз, моделювання, аналіз великих обсягів даних, нейронні мережі для прогнозування атак, методи теорії управління.

Наукова новизна одержаних результатів.

1. В магістерській роботі вперше запропонована нова методика створення моделей виявлення загроз кібернетичній безпеці, які відрізняється тим, що використовує при створенні стандартні бібліотеки мови програмування Python моніторингу соціальних мереж з метою виявлення потенційних загроз та випереджального реагування на них.

2. Дістали подальшого розвитку способи застосування штучного інтелекту в управлінні кібербезпекою, які полягають в поєднанні традиційних методів кібербезпеки та методики розроблення моделей на основі нейронних мереж у відповідності із новітніми вимогами NIST щодо застосування штучного інтелекту в управлінні кібербезпекою .

Практичне значення одержаних результатів. Розроблені рекомендації фахівцям управління кібербезпекою, керівникам структурних підрозділів інформаційної безпеки щодо удосконалення застосування штучного інтелекту в управлінні кібербезпекою. Разом із запропонованою методикою створення моделей виявлення загроз кібербезпеці рекомендації дозволяють ефективно будувати та управляти кібернетичною безпекою організацій, підприємств.

Апробація результатів кваліфікаційної роботи. Результати дослідження було оприлюднено на Всеукраїнській науково-практичній Інтернет-конференції «Стратегії кіберстійкості: управління ризиками та безперервність бізнесу» (м. Київ, 23 лютого 2023 року), яка проведена в Навчально-науковому інституті захисту інформації ДУІКТ.

.

РОЗДІЛ 1

ОГЛЯД ТЕОРЕТИЧНИХ ЗАСАД КІБЕРНЕТИЧНОЇ БЕЗПЕКИ

1.1 Основні загрози та виклики для кібернетичної безпеки

Кібербезпека – це стан або процес захисту та відновлення комп'ютерних систем, мереж, пристроїв і програм від будь-якого типу кібератак.

Як визначено законодавчими документами [1], під кібербезпекою розуміють захищеність життєво важливих інтересів людини, громадянина, суспільства і держави під час використання кіберпростору, при якій забезпечуються сталий розвиток інформаційного суспільства та цифрового комунікативного середовища, своєчасне виявлення, запобігання і нейтралізація реальних і потенційних загроз національній безпеці України у кіберпросторі.

Термін «кібербезпека» відноситься до сукупності технологій, процесів і практик для захисту мереж, пристроїв, програмного забезпечення та даних від атак, пошкоджень або несанкціонованого доступу [2]. Кібербезпека передбачає політику, процедури та технічні механізми для захисту, виявлення, виправлення та захисту від пошкодження, несанкціонованого використання, модифікації або експлуатації інформаційно-комунікаційних систем та інформації, яку вони містять.

Сучасний світ все більше залежить від технологій, і ця залежність збережеться, коли ми представимо наступне покоління нових технологій, які матимуть доступ до наших підключених пристроїв через Bluetooth і Wi-Fi. Зі зростанням залежності від технологій захист конфіденційної інформації став критичним пріоритетом. Кібербезпека передбачає різні заходи та практики, які захищають комп'ютерні системи та мережі від несанкціонованого доступу, пошкодження чи крадіжки. Це передбачає необхідність впровадження надійних протоколів безпеки, складних методів шифрування та проактивних контрзаходів.

Важливість кібербезпеки очевидна через збільшення кількості кібератак,

які руйнують бізнес і впливають на людей у всьому світі. Кіберзагрози можуть спричинити витік даних, фінансові втрати та репутацію.

Кібератаки стають дедалі складнішою небезпекою для ваших конфіденційних даних, оскільки зловмисники використовують нові методи на основі соціальної інженерії та штучного інтелекту (ШІ), щоб обійти традиційні засоби контролю безпеки даних.

Станом на жовтень 2023 року в результаті витоку даних виявлено понад 600 мільйонів записів. Наслідки великих витоків даних, які стосуються секторів і організацій, таких як Healthcare, Twitter і MOVEit (клієнти), означають щось більше, ніж просто необхідність змінити пароль [3]. Це означає, що окремі люди та групи орієнтуються на технологію, яка, по суті, визначає та підтримує вас у сучасному світі. Вони націлені на системи, які містять ваші особисті дані. Іншими словами, зловмисники з усього світу націлені на вас. Ось чому захист наших інформаційних систем є обов'язковим.

Кібербезпека важлива буквально для всіх сфер життєдіяльності людини, суспільства та держави. Вона допомагає захистити від кібератак організації та окремих осіб, підприємства і банки, критичну інфраструктуру економіки держави.

Кібербезпека важлива для студентів, оскільки вони часто націлені на кібератаки. У нещодавньому випадку група студентів коледжу в Сполучених Штатах стала мішенню хакерів, які отримали доступ до їхньої особистої інформації, включаючи номери соціального страхування та дані кредитних карток. Потім хакери використали цю інформацію, щоб шахрайським шляхом стягнути тисячі доларів з кредитних карток студентів. Студенти залишилися з величезними боргами і були змушені витратити місяці на відновлення свого кредиту. Цей кейс підкреслює важливість кібербезпеки для студентів, які часто стають жертвами кіберзлочинів.

Якщо особиста інформація студента буде викрадена під час кібератаки, вона може бути використана для крадіжки особистих даних. Це може зіпсувати кредит студента, ускладнивши студенту отримання кредиту на навчання або

автомобіль. У крайніх випадках крадіжка особистих даних може призвести навіть до тюремного ув'язнення.

Кібербезпека може допомогти запобігти витоку даних, крадіжці особистих даних та інших видів кіберзлочинності. Важливість кібербезпеки для бізнесу та організацій можна побачити у випадку цільового витоку даних. У цьому випадку хакери змогли отримати доступ до даних клієнтів цілі, включаючи інформацію про кредитні та дебетові картки. Це призвело до того, що Target довелося виплатити мільйони доларів збитків і втратити довіру клієнтів. Витік даних Target є лише одним із прикладів того, наскільки кібербезпека важлива для бізнесу та організацій.

Реальним прикладом важливості кібербезпеки для банківського сектору є витік даних JPMorgan Chase у 2014 році. Під час цього злomu хакери отримали доступ до імен, адрес, номерів телефонів та адрес електронної пошти 76 мільйонів домогосподарств та 7 мільйонів малих підприємств. Хакери також отримали доступ до інформації про облікові записи, включаючи номери рахунків і баланси, 83 мільйонів клієнтів JPMorgan Chase [4].

Цей злом підкреслює важливість кібербезпеки для банківського сектору, оскільки хакери змогли отримати доступ до великої кількості конфіденційних даних клієнтів. Якби ці дані потрапили в чужі руки, вони могли б бути використані для крадіжки особистих даних, шахрайства або інших зловмисних цілей.

Іншим прикладом витоку даних може бути атака програм-вимагачів WannaCry, націлена на підприємства та організації по всьому світу. Ця атака призвела до втрати даних і грошей для багатьох організацій, а деякі навіть були змушені закритися [5].

Останніми роками сталося кілька резонансних кібератак, які мали руйнівний вплив на бізнес і приватних осіб. Це крадіжка номерів соціального страхування, реквізитів банківських рахунків, даних кредитних карток та витоку конфіденційних даних. Основна причина полягає в тому, що більшість людей зберігають свої дані в хмарних сховищах, таких як Dropbox або Google Drive. Ці

атаки підкреслили важливість наявності надійних заходів кібербезпеки [6].

Таким чином, загрози кібербезпеці продовжують зростати та розвиватися за частотою, вектором та складністю. На рисунку 1.1 показані основні види загроз.

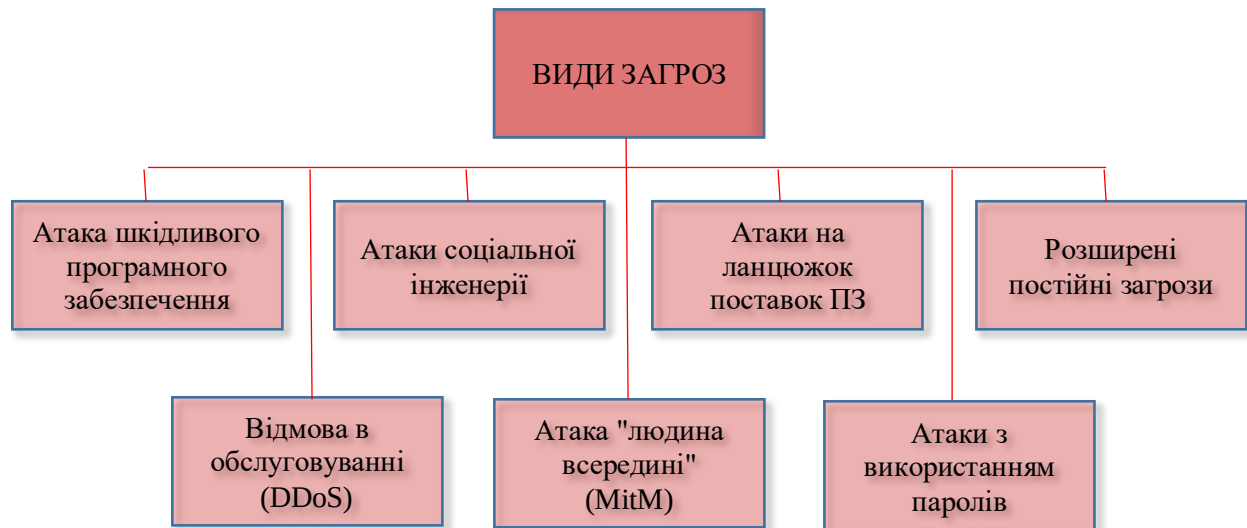


Рис. 1.1. Основні види загроз

1. Атака шкідливого програмного забезпечення

В атаках використовується безліч методів проникнення шкідливого програмного забезпечення на пристрій користувача, найчастіше це соціальна інженерія. Користувачам може бути запропоновано виконати певну дію, наприклад клацнути посилання або відкрити вкладення. В інших випадках шкідливе програмне забезпечення використовує уразливості в браузері або операційних системах для самовстановлення без відома чи згоди користувача [7].

Після встановлення зловмисного програмного забезпечення воно може відстежувати дії користувача, надсилати конфіденційні дані зловмиснику, допомагати зловмиснику проникати в інші цілі в мережі та навіть змушувати пристрій користувача брати участь у ботнеті, який зловмисник використовує для зловмисних намірів.

Атаки шкідливого програмного забезпечення включають (рис. 1.2):

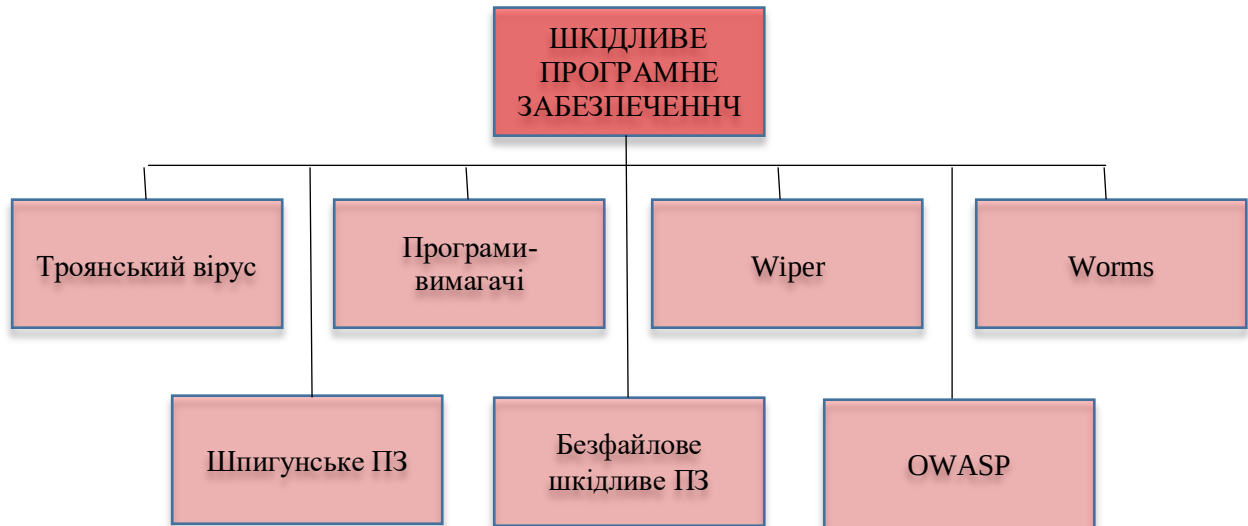


Рис. 1.2. Види атаки шкідливого програмного забезпечення

Троянський вірус – обманює користувача, змушуючи його думати, що це нешкідливий файл. Троян може запустити атаку на систему і може встановити бекдор, яким можуть скористатися зловмисники.

Програми-вимагачі – перешкоджають доступу до даних жертви та погрожують видалити або опублікувати їх, якщо не буде сплачено викуп. Дізнайтеся більше в нашому посібнику із запобігання програмам-вимагачам.

Шкідливе програмне забезпечення Wiper – має намір знищити дані або системи, перезаписуючи цільові файли або знищуючи цілу файлову систему. Склоочисники зазвичай призначені для надсилання політичного повідомлення або приховування хакерської діяльності після витоку даних.

Worms – це шкідливе програмне забезпечення призначене для використання бекдорів і вразливостей для отримання несанкціонованого доступу до операційних систем. Після інсталяції хробак може виконувати різні атаки, включаючи розподілену відмову в обслуговуванні (DDoS).

Шпигунське програмне забезпечення – це шкідливе програмне забезпечення дозволяє зловмисникам отримувати несанкціонований доступ до даних, включаючи конфіденційну інформацію, таку як платіжні реквізити та облікові дані. Шпигунське програмне забезпечення може вражати мобільні

телефони, настільні програми та настільні браузері.

Безфайлове шкідливе ПЗ – цей тип шкідливого ПЗ не вимагає встановлення програмного забезпечення на операційну систему. Це робить рідні файли, такі як PowerShell і WMI, доступними для редагування, щоб увімкнути шкідливі функції, завдяки чому вони розпізнаються як законні та їх важко виявити.

Маніпулювання додатками або веб-сайтами – OWASP описує 10 основних ризиків для безпеки додатків, починаючи від порушеного контролю доступу та неправильної конфігурації безпеки до ін'єкційних атак та криптографічних збоїв. Після того, як вектор встановлено шляхом придбання сервісних облікових записів, запускаються нові атаки на шкідливе програмне забезпечення, облікові дані або APT [8].

2. Атаки соціальної інженерії працюють, психологічно маніпулюючи користувачами, змушуючи їх виконувати дії, бажані зловмиснику, або розголошуючи конфіденційну інформацію.

Атаки соціальної інженерії включають (рис.1.3):



Рис.1.3. Атаки соціальної інженерії

Фішинг – зловмисники надсилають шахрайську кореспонденцію, яка, здається, надходить із законних джерел, зазвичай електронною поштою. Електронний лист може спонукати користувача виконати важливу дію або натиснути посилання на шкідливий веб-сайт, що призведе до того, що він передасть конфіденційну інформацію зловмиснику або піддасть себе шкідливим завантаженням. Фішингові електронні листи можуть містити вкладення електронної пошти, заражене зловмисним програмним забезпеченням.

Цільовий фішинг – варіант фішингу, при якому зловмисники спеціально націлені на осіб, які мають привілеї безпеки або вплив, наприклад, на системних адміністраторів або керівників вищої ланки.

Шкідлива реклама – контрольована хакерами інтернет-реклама, яка містить шкідливий код, який заражає комп'ютер користувача при натисканні або навіть просто перегляді реклами. Шкідлива реклама була виявлена в багатьох провідних інтернет-виданнях.

Drive-by-завантаження – зловмисники можуть зламувати веб-сайти та вставляти шкідливі скрипти в PHP або HTTP-код на сторінці. Коли користувачі відвідують сторінку, шкідливе програмне забезпечення встановлюється безпосередньо на їхній комп'ютер; Або скрипт зловмисника перенаправляє користувачів на шкідливий сайт, який виконує завантаження. Випадкові завантаження залежать від вразливостей у браузерях або операційних системах. Дізнайтеся більше в посібнику із завантаження автомобілів.

Захисне програмне забезпечення Scareware – вдає, що сканує наявність шкідливих програм, а потім регулярно показує користувачеві фальшиві попередження та виявлення. Зловмисники можуть попросити користувача заплатити за видалення підроблених загроз зі свого комп'ютера або за реєстрацію програмного забезпечення. Користувачі, які дотримуються вимог, передають свої фінансові дані зловмиснику.

Приманка – виникає, коли зловмисник обманом змушує ціль використовувати шкідливий пристрій, розміщуючи заражений шкідливим

програмним забезпеченням фізичний пристрій, наприклад USB, де ціль може його знайти. Після того, як ціль вставляє пристрій у свій комп'ютер, вона ненавмисно встановлює шкідливе програмне забезпечення.

Вішинг – атаки голосового фішингу (вішингу) використовують методи соціальної інженерії, щоб змусити жертв розголошувати фінансову або особисту інформацію по телефону.

Китобійний промисел – ця фішингова атака націлена на високопоставлених співробітників (китів), таких як головний виконавчий директор (CEO) або фінансовий директор (CFO). Зловмисник намагається обманом змусити ціль розкрити конфіденційну інформацію.

Претекстинг – виникає, коли зловмисник бреше цілі, щоб отримати доступ до привілейованих даних. Шахрайство з претекстом може включати зловмисника, який вдає, що підтверджує особу цілі, запитуючи фінансові або особисті дані.

Scareware – зловмисник обманом змушує жертву подумати, що вона ненавмисно завантажила незаконний вміст або що її комп'ютер заражений шкідливим програмним забезпеченням. Далі зловмисник пропонує жертві рішення для вирішення фальшивої проблеми, обманом змушуючи жертву завантажити та встановити зловмисне програмне забезпечення.

Крадіжка переадресації – зловмисники використовують соціальних інженерів, щоб обманом змусити кур'єра або кур'єрську компанію вирушити до неправильного місця висадки або посадки, перехопивши транзакцію.

Медова пастка – соціальний інженер припускає фальшиву особистість як привабливу особу для взаємодії з цільовою аудиторією в Інтернеті. Соціальний інженер імітує стосунки в Інтернеті та збирає конфіденційну інформацію через ці стосунки.

Затримка хвоста або контрабанда – виникає, коли зловмисник проникає в будівлю, що охороняється, слідуючи за уповноваженим персоналом. Як правило, персонал, який має законний доступ, припускає, що людині позаду дозволено увійти, тримаючи двері відчиненими для них.

Фармінг – схема онлайн-шахрайства, під час якої кіберзлочинець встановлює шкідливий код на сервер або комп'ютер. Код автоматично спрямовує користувачів на підроблений веб-сайт, де користувачів обманом змушують надати особисті дані.

3. Атаки на ланцюжок поставок програмного забезпечення

Атака на ланцюжок поставок програмного забезпечення – це кібератака на організацію, яка націлена на слабкі ланки в її надійному оновленні програмного забезпечення та ланцюжку поставок. Ланцюг поставок - це мережа всіх осіб, організацій, ресурсів, видів діяльності та технологій, які беруть участь у створенні та продажу продукту. Атака на ланцюжок поставок програмного забезпечення використовує довіру, яку організації мають до своїх сторонніх постачальників, особливо в оновленнях і виправленнях.

Це особливо актуально для інструментів мережевого моніторингу, промислових систем управління, «розумних» машин та інших мережевих систем з сервісними обліковими записами. Атака може бути здійснена в багатьох місцях проти життєвого циклу програмного забезпечення безперервної інтеграції та безперервної доставки (CI/CD) або навіть проти сторонніх бібліотек і компонентів, як це видно через Apache і Spring.

Типи атак на ланцюжок поставок програмного забезпечення:

- компрометація інструментів створення програмного забезпечення або інфраструктури розробки/тестування;
- компрометація пристроїв або облікових записів, що належать привілейованим стороннім постачальникам;
- шкідливі додатки, підписані сертифікатами підпису викраденого коду або ідентифікаторами розробників;
- шкідливий код, розгорнутий на апаратних або мікропрограмних компонентах;
- зловмисне програмне забезпечення, попередньо встановлене на таких пристроях, як камери, USB-накопичувачі та мобільні телефони.

4. Розширені постійні загрози (APT). Коли особа або група осіб отримує

несанкціонований доступ до мережі та залишається непоміченою протягом тривалого періоду часу, зловмисники можуть викрасти конфіденційні дані, навмисно уникаючи виявлення співробітниками служби безпеки організації. АРТ вимагають досвідчених зловмисників і вимагають великих зусиль, тому вони зазвичай запускаються проти національних держав, великих корпорацій або інших дуже цінних цілей [8].

Загальні показники наявності АРТ включають (рис. 1.4):

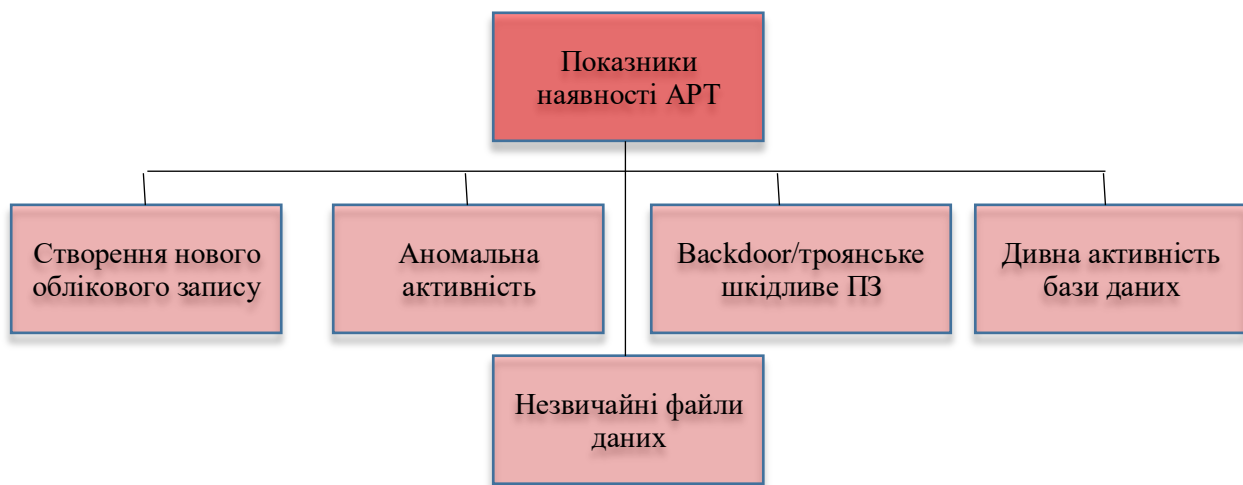


Рис. 1.4. Показники наявності АРТ

Створення нового облікового запису – у Persistent походить від зловмисника, який створює ідентифікаційні дані або облікові дані в мережі з підвищеними привілеями.

Аномальна активність – законні облікові записи користувачів зазвичай працюють за шаблонами. Аномальна активність у цих облікових записах може вказувати на АРТ, зокрема позначити застарілий обліковий запис, який було створено, а потім деякий час не використовувався, який раптово став активним.

Backdoor/троянське шкідливе програмне забезпечення – широке використання цього методу дозволяє АРТ підтримувати довгостроковий доступ.

Дивна активність бази даних – наприклад, раптове збільшення операцій з базами даних з величезними обсягами даних.

Незвичайні файли даних – наявність цих файлів може вказувати на те, що

дані були об'єднані у файли, щоб допомогти в процесі викрадення.

5. Розподілена відмова в обслуговуванні (DDoS)

Мета атаки типу «відмова в обслуговуванні» (DoS) полягає в тому, щоб перевантажити ресурси цільової системи та змусити її припинити функціонування, заборонивши доступ своїм користувачам. Розподілена відмова в обслуговуванні (DDoS) – це варіант DoS, при якому зловмисники компрометують велику кількість комп'ютерів або інших пристроїв і використовують їх для скоординованої атаки на цільову систему.

DDoS-атаки часто використовуються в комплексі з іншими кіберзагрозами. Ці атаки можуть запустити відмову в обслуговуванні, щоб привернути увагу співробітників служби безпеки та створити плутанину, тоді як вони здійснюють більш тонкі атаки, спрямовані на крадіжку даних або заподіяння іншої шкоди.

Методи DDoS-атак включають (рис. 1.5):

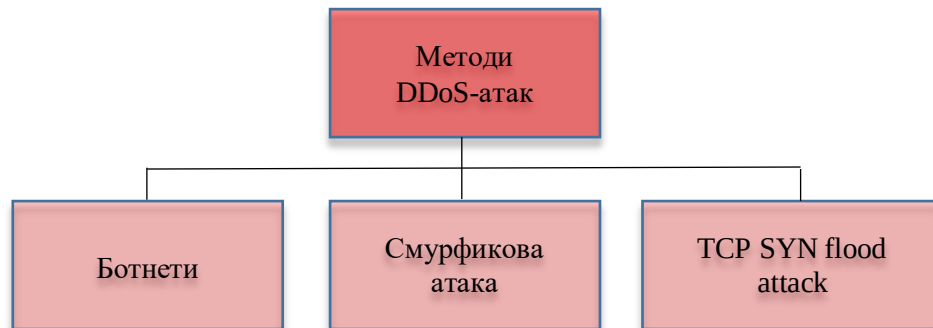


Рис. 1.5. Методи DDoS-атак

Ботнети – системи під контролем хакерів, які були заражені шкідливим програмним забезпеченням. Зловмисники використовують цих ботів для здійснення DDoS-атак. Великі ботнети можуть включати мільйони пристроїв і запускати атаки в руйнівних масштабах.

Смурфикова атака – надсилає echo-запити Internet Control Message Protocol (ICMP) на IP-адресу жертви. Запити ICMP генеруються з «піддроблених» IP-адрес. Зловмисники автоматизують цей процес і виконують його в масштабі, щоб

перевантажити цільову систему.

TCP SYN flood attack – атаки заповнюють цільову систему запитами на підключення. Коли цільова система намагається завершити з'єднання, пристрій зловмисника не відповідає, змушуючи цільову систему перевищити час. Це швидко заповнює чергу підключення, не дозволяючи законним користувачам підключатися.

6. Атака "людина посередині" (MitM - англ. Man in the middle)

Коли користувачі або пристрої отримують доступ до віддаленої системи через Інтернет, вони припускають, що спілкуються безпосередньо з сервером цільової системи. Під час атаки MitM зловмисники порушують це припущення, опиняючись між користувачем і цільовим сервером. Після того, як зловмисник перехопив повідомлення, він може скомпрометувати облікові дані користувача, викрасти конфіденційні дані та повернути користувачеві різні відповіді.

Атаки MitM включають (рис. 1.6):

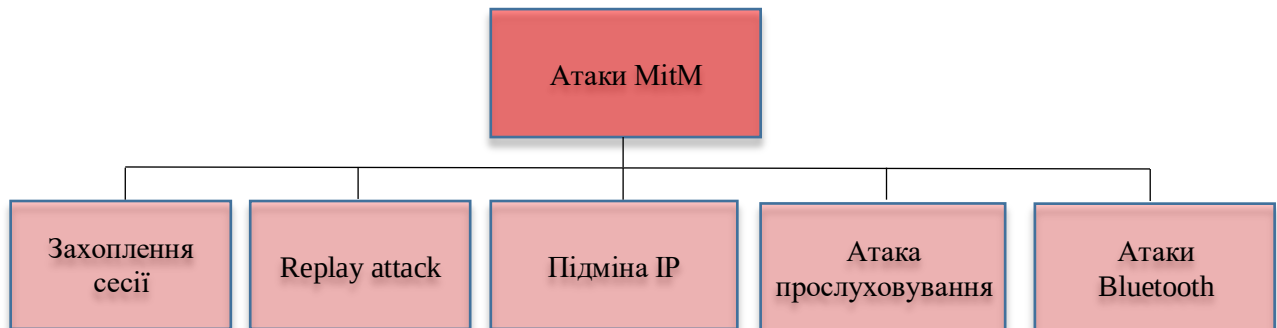


Рис. 1.6. Складові атаки "людина посередині" (MitM)

Захоплення сесії – зловмисник перехоплює сеанс між мережевим сервером і клієнтом. Атакуючий комп'ютер підмінює IP-адресу клієнта своєю IP-адресою. Сервер вважає, що листується з клієнтом, і продовжує сесію.

Replay attack – кіберзлочинець підслуховує мережевий зв'язок і відтворює повідомлення пізніше, видаючи себе за користувача. Атаки Replay були значною мірою пом'якшені додаванням часових позначок до мережевих комунікацій.

Підміна IP – зловмисник переконує систему, що вона листується з

довіреною відомою сутністю. Таким чином, система надає зловмиснику доступ. Зловмисник підробляє свій пакет з IP-адресою джерела довіреного хоста, а не з власною IP-адресою.

Атака прослуховування – зловмисники використовують незахищений мережевий зв'язок для доступу до інформації, що передається між клієнтом і сервером. Ці атаки важко виявити, оскільки мережеві передачі, схоже, працюють нормально.

Атаки Bluetooth – Оскільки Bluetooth часто відкритий у безладному режимі, існує багато атак, особливо на телефони, які скидають картки контактів та інше шкідливе програмне забезпечення через відкриті та отримані з'єднання Bluetooth. Зазвичай ця компрометація кінцевої точки є засобом для досягнення мети, від збору облікових даних до особистої інформації.

7. Атаки на паролі

Хакер може отримати доступ до інформації про пароль людини, «пронюхавши» з'єднання з мережею, використовуючи соціальну інженерію, вгадування або отримавши доступ до бази паролів. Зловмисник може «вгадати» пароль випадковим або систематичним способом.

До атак з використанням паролів належать (рис. 1.7):

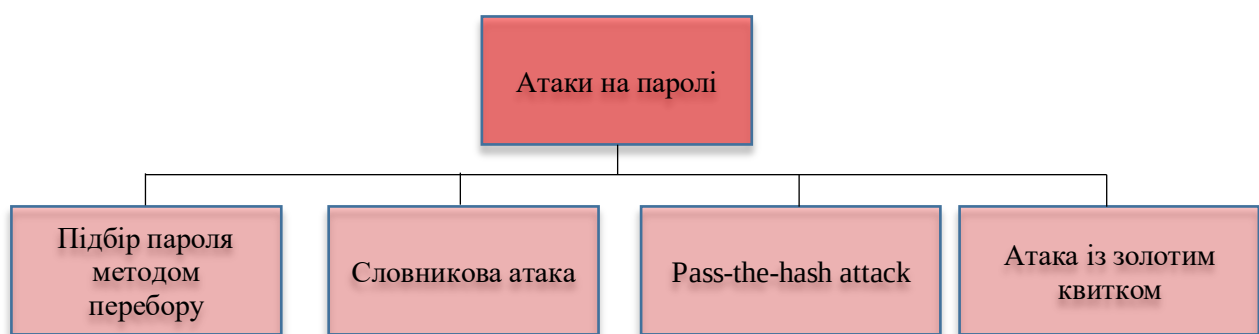


Рис. 1.7. Різновид атак на паролі

Підбір пароля методом перебору – зловмисник використовує програмне забезпечення, щоб спробувати багато різних паролів, сподіваючись вгадати правильний. Програмне забезпечення може використовувати певну логіку для

спроби паролів, пов'язаних з ім'ям людини, її роботою, сім'єю тощо.

Словникова атака – словник поширених паролів використовується для отримання доступу до комп'ютера та мережі жертви. Один із методів полягає в тому, щоб скопіювати зашифрований файл із паролями, застосувати те саме шифрування до словника паролів, які регулярно використовуються, і порівняти отримані результати.

Pass-the-hash attack – зловмисник використовує протокол аутентифікації під час сеансу та перехоплює хеш пароля (на відміну від символів пароля безпосередньо), а потім передає його для аутентифікації та бічного доступу до інших мережових систем. У цих типах атак зловмиснику не потрібно розшифровувати хеш, щоб отримати пароль із відкритим текстом.

Атака із золотим квитком починається так само, як і атака за допомогою хеш-атаки, коли в системі Kerberos (Windows AD) зловмисник використовує вкрадений хеш пароля для доступу до центру розподілу ключів, щоб підробити хеш квитка на видачу квитків (TGT). В атаках Мімікаца часто використовується цей вектор атаки.

При визначенні кіберзагроз важливо розуміти, хто є суб'єктом загрози, а також його тактику, методи та процедури. До поширених джерел кіберзагроз належать:

- *спонсоровані державою* – кібератаки з боку країн можуть порушити зв'язок, військову діяльність чи інші послуги, якими громадяни користуються щодня;
- *терористи* – терористи можуть атакувати урядові або військові об'єкти, але іноді також можуть атакувати цивільні веб-сайти, щоб порушити роботу та завдати довготривалої шкоди;
- *промислові шпигуни* – організована злочинність і міжнародні корпоративні шпигуни здійснюють промислове шпигунство і крадіжку грошей. Їх основний мотив – фінансовий;
- *організовані злочинні угруповання* – злочинні угруповання, які проникають у системи з метою отримання грошової вигоди. Організовані злочинні групи

використовують фішинг, спам і шкідливе програмне забезпечення для крадіжки особистих даних та шахрайства в Інтернеті. Існують організовані злочинні групи, які існують для того, щоб продавати хакерські послуги іншим, підтримуючи навіть підтримку та послуги як для спекулянтів, так і для промислових шпигунів;

- *хакери* – існує велика глобальна популяція хакерів, починаючи від початківців «дітей зі сценарієм» або тих, хто використовує готові набори інструментів для захисту від загроз, і закінчуючи досвідченими операторами, які можуть розробляти нові типи загроз і уникати організаційного захисту;

- *хактивісти* – це хакери, які проникають або руйнують системи з політичних чи ідеологічних причин, а не заради фінансової вигоди;

- *зловмисний інсайдери* становлять дуже серйозну загрозу, оскільки вони мають існуючий доступ до корпоративних систем і знання цільових систем і конфіденційних даних. Внутрішні загрози можуть бути руйнівними, і їх дуже важко виявити.

Кібершпигунство – це форма кібератаки, яка викрадає секретні або конфіденційні інтелектуальні дані з метою отримання переваги над конкурентною компанією чи державною установою.

З розвитком технологій у 2023 році зростають і загрози та проблеми, з якими стикаються команди безпеки. Одної з таких важливих проблем є застосування штучного інтелекту у кібербезпеці. Штучний інтелект – палиця з двома кінцями. Він покращує рішення безпеки, але в той же час використовується зловмисниками для обходу цих рішень. Однією з причин цього є зростаюча доступність штучного інтелекту. У минулому розробка моделей машинного навчання була можлива лише за наявності доступу до значних бюджетів та ресурсів. Однак зараз моделі можна розробляти і на персональних ноутбуках. Ця доступність робить штучний інтелект інструментом, який розширився від великих цифрових гонок озброєнь до повсякденних атак. У той час як команди безпеки використовують штучний інтелект, щоб спробувати виявити підозрілу поведінку, злочинці використовують його для створення ботів, які видають себе за користувачів-людей, і для динамічної зміни

характеристик і поведінки шкідливого програмного забезпечення [10].

Особливу загрозу сьогодні хакерські атаки складають для на Інтернету речей (IoT) та транспортних засобів. IoT – це не тільки розумний дім, це поняття включає і віддалений контроль за допомогою сенсорів за технологічними процесами виробництва з передачею даних мережевою інфраструктурою [11].

Дані на цих пристроях можуть надавати конфіденційну інформацію злочинцям. Ця інформація включає приватні розмови, конфіденційні зображення, інформацію про відстеження та доступ до будь-яких облікових записів, які використовуються з пристроями. Ці пристрої можуть бути легко використані зловмисниками для шантажу або особистої вигоди. Наприклад, зловживання фінансовою інформацією або продаж інформації на чорному ринку.

Обсяг даних, що містяться в сучасному транспортному засобі, величезний. Навіть автомобілі, які не є автономними, завантажені різноманітними розумними датчиками. Сюди входять пристрої GPS, вбудовані комунікаційні платформи, камери та контролери штучного інтелекту. Будинки, робочі місця та громади багатьох людей переповнені подібними розумними пристроями. Наприклад, персональні помічники, вбудовані в колонки, є розумними пристроями.

Зокрема, у випадку з транспортними засобами загроза заподіяння шкоди здоров'ю також дуже реальна. Коли транспортні засоби частково або повністю контролюються комп'ютерами, зловмисники мають можливість зламати транспортні засоби так само, як і будь-який інший пристрій. Це може дозволити їм використовувати транспортні засоби як зброю проти інших або як засіб заподіяння шкоди водієві чи пасажиром.

Навіть якщо люди не повністю освоїли розумні технології, майже у кожного є якийсь мобільний пристрій. Поширені смартфони, ноутбуки та планшети. Ці пристрої часто є багатоцільовими, використовуються як для роботи, так і для особистої діяльності, і користувачі можуть підключати пристрої до кількох мереж протягом дня.

Така велика кількість і широке використання роблять мобільні пристрої привабливою мішенню для зловмисників. Таргетинг не є чимось новим, але

справжня проблема пов'язана з тим, що команди безпеки не мають повного контролю над пристроями. Правила використання власного пристрою (BYOD) є поширеними, але вони часто не включають внутрішній контроль або керування.

Часто служби безпеки можуть контролювати лише те, що відбувається з цими пристроями в периметрі мережі. Можливо, пристрої застаріли, вже заражені шкідливим програмним забезпеченням або мають недостатній захист. Єдиний спосіб, яким служби безпеки можуть заблокувати ці загрози, – це відмовитися від підключення, що непрактично.

У зв'язку з тим, що компанії щодня переходять на хмарні ресурси, багато середовищ стають все складнішими. Це особливо актуально у випадку гібридних та мультихмарних середовищ, які потребують ретельного моніторингу та інтеграції [12-15].

З кожним хмарним сервісом і ресурсом, які включені в середовище, збільшується кількість кінцевих точок і ймовірність неправильної конфігурації. Крім того, оскільки ресурси знаходяться в хмарі, більшість, якщо не всі кінцеві точки, орієнтовані на Інтернет, що надає зловмисникам доступ у глобальному масштабі.

Щоб убезпечити ці середовища, командам з кібербезпеки потрібні передові, централізовані інструменти та часто більше ресурсів. Це включає ресурси для захисту та моніторингу 24/7, оскільки ресурси працюють і потенційно вразливі, навіть коли робочий день закінчився.

Російсько-українська війна та нова геополітична ситуація підняли ставки спонсорованих державою атак на західні країни та організації. У міру того, як все більше світу переходить у цифрову сферу, кількість великомасштабних атак, спонсорованих державою, зростає. Мережі хакерів тепер можуть бути використані та куплені протиборчими національними державами та групами інтересів, щоб завдати шкоди урядовим та організаційним системам.

Для деяких з цих атак результати очевидні. Наприклад, було виявлено численні атаки, пов'язані з фальсифікацією виборів. Інші, однак, можуть залишитися непоміченими, мовчки збираючи конфіденційну інформацію, таку

як військові стратегії або бізнес-розвідка. У будь-якому випадку, ресурси, що фінансують ці атаки, дозволяють злочинцям використовувати передові та розподілені стратегії, які важко виявити та запобігти.

Кібербезпека важлива, оскільки захищає всі категорії даних від крадіжки та пошкодження: конфіденційні дані, особиста ідентифікаційна інформація, захищена медична інформація, особиста інформація, інтелектуальна власність, дані, а також урядові та галузеві інформаційні системи. Без програми кібербезпеки жодна організація не може захистити себе від кампаній з витоку даних, що робить її мішенню для кіберзлочинців.

Підсумовуючи, кібербезпека має вирішальне значення для захисту конфіденційної інформації, зменшення ризику витоку даних, фінансових втрат і шкоди репутації. Це має основне значення для безпечної та надійної навігації серед загроз.

1.2 Роль штучного інтелекту в контексті управління кібернетичною безпекою

Історично склалося так, що кібербезпека була сферою, в якій домінували ресурсомісткі зусилля. Моніторинг, пошук загроз, реагування на інциденти та інші обов'язки часто виконуються вручну та забирають багато часу, що може затримати заходи з виправлення ситуації, збільшити вразливість та підвищити вразливість до кіберзловмисників.

За останні кілька років рішення штучного інтелекту швидко розвинулися до такої міри, що вони можуть принести значну користь операціям кіберзахисту в широкому спектрі організацій і місій. Автоматизуючи ключові елементи основних функцій, що вимагають великої кількості праці, штучний інтелект може перетворити робочі кіберпроцеси на оптимізовані, автономні, безперервні процеси, які прискорюють виправлення та максимізують захист.

Бізнес-лідери більше не можуть покладатися виключно на готові рішення кібербезпеки, такі як антивірусне програмне забезпечення та брандмауери,

кіберзлочинці стають розумнішими, а їхня тактика стає більш стійкою до звичайного кіберзахисту. Важливо охопити всі сфери кібербезпеки, щоб залишатися добре захищеним.

Штучний інтелект (англ. Artificial Intelligence, AI) та машинне навчання (англ. Machine learning, ML) все частіше використовуються в кібербезпеці для виявлення загроз, реагування та автоматизації завдань безпеки. Однак такі проблеми, як змагальне машинне навчання, упередження в алгоритмах штучного інтелекту та потенціал атак, керованих штучним інтелектом, повинні бути ретельно розглянуті та пом'якшені [16-17].

Штучний інтелект став невід'ємною частиною управління кібербезпекою. Рішення на основі штучного інтелекту можуть допомогти командам безпеки підвищити швидкість, точність і продуктивність, виявляючи загрози безпеці та реагуючи на них у режимі реального часу. AI також може допомогти виявити вразливі місця та захистити від кіберзлочинців.

Рішення на основі штучного інтелекту в кібербезпеці відрізняються від традиційних підходів за кількома параметрами.

Традиційні підходи до кібербезпеки значною мірою покладалися на сигнатурні системи виявлення, які були ефективними лише проти відомих загроз. Це означало, що нові та невідомі загрози могли залишитися непоміченими.

Навпаки, рішення на основі штучного інтелекту використовують алгоритми машинного навчання, які можуть виявляти та реагувати як на відомі, так і на невідомі загрози в режимі реального часу.

Алгоритми машинного навчання навчаються на основі величезних обсягів даних, включаючи історичні дані про загрози та дані з мережі та кінцевих точок, щоб виявляти закономірності, які людям важко побачити. Це дозволяє рішенням на основі штучного інтелекту виявляти загрози та реагувати на них у режимі реального часу, без необхідності втручання людини.

Наприклад, алгоритми машинного навчання можуть аналізувати моделі мережевого трафіку, щоб виявити аномальну поведінку, яка може вказувати на

кібератаку, а потім попереджати персонал служби безпеки або навіть вживати автоматизованих заходів для пом'якшення загрози [17-18].

Ще одна відмінність рішень на основі штучного інтелекту від традиційних підходів полягає в тому, що вони призначені для постійного навчання та адаптації.

У міру появи нових загроз алгоритми машинного навчання можна навчати на нових даних, щоб покращити їхню здатність виявляти та реагувати на ці загрози. Це означає, що рішення на основі штучного інтелекту можуть йти в ногу з мінливим ландшафтом загроз і з часом забезпечувати більш ефективний захист кібербезпеки.

Використання штучного інтелекту в кібербезпеці є серйозним зрушенням у підході організацій до кібербезпеки. Рішення на основі штучного інтелекту можуть забезпечити більш ефективний захист як від відомих, так і від невідомих загроз – використовуючи алгоритми машинного навчання для виявлення загроз і реагування на них у режимі реального часу. Це допомагає організаціям краще захищати свої конфіденційні дані та критично важливі системи.

Штучний інтелект використовується в кібербезпеці для виявлення та реагування на кіберзагрози в режимі реального часу. Алгоритми штучного інтелекту можуть аналізувати великі обсяги даних і виявляти закономірності, які свідчать про кіберзагрозу.

Шкідливе програмне забезпечення є серйозною загрозою для кібербезпеки. Традиційне антивірусне програмне забезпечення покладається на виявлення сигнатур для виявлення відомих варіантів шкідливого програмного забезпечення.

Виявлення на основі сигнатур — це метод, який порівнює файл із базою даних відомих сигнатур зловмисного програмного забезпечення та виявляє збіг. Ця техніка ефективна лише проти відомих варіантів зловмисного програмного забезпечення, і її можна легко обійти за допомогою шкідливого програмного забезпечення, яке було модифіковано, щоб уникнути виявлення.

Рішення на основі штучного інтелекту використовують алгоритми

машинного навчання для виявлення та реагування як на відомі, так і на невідомі загрози зловмисного програмного забезпечення. Алгоритми машинного навчання можуть аналізувати великі обсяги даних, щоб виявити закономірності та аномалії, які людям важко виявити. Аналізуючи поведінку шкідливого програмного забезпечення, штучний інтелект може виявляти нові та невідомі варіанти шкідливого програмного забезпечення, які можуть бути пропущені традиційним антивірусним програмним забезпеченням.

Рішення для виявлення шкідливого ПЗ на основі штучного інтелекту можна навчати, використовуючи як марковані, так і немарковані дані [19].

Позначені дані – це дані, позначені певними атрибутами, наприклад, чи є файл шкідливим або безпечним. З іншого боку, немарковані дані не позначаються тегамі і можуть бути використані для навчання алгоритмів машинного навчання для виявлення закономірностей і аномалій у даних.

Рішення для виявлення шкідливого програмного забезпечення на основі штучного інтелекту можуть використовувати різні методи для виявлення шкідливого програмного забезпечення, такі як статичний аналіз і динамічний аналіз.

Статичний аналіз передбачає аналіз характеристик файлу, таких як його розмір, структура та код, для виявлення закономірностей та аномалій. Динамічний аналіз передбачає аналіз поведінки файлу під час його виконання для виявлення закономірностей та аномалій.

Рішення на основі штучного інтелекту забезпечують більш просунутий та ефективний підхід до виявлення шкідливого програмного забезпечення, ніж традиційне антивірусне програмне забезпечення. Вони можуть виявляти нові та невідомі варіанти шкідливого програмного забезпечення, які можуть бути пропущені традиційним антивірусним програмним забезпеченням.

Традиційні підходи до виявлення фішингу зазвичай покладаються на фільтрацію на основі правил або чорний список для виявлення та блокування відомих фішингових електронних листів. Ці підходи мають обмеження, оскільки вони ефективні лише проти відомих атак і можуть пропустити нові атаки або

атаки, що розвиваються.

Рішення для виявлення фішингу на основі штучного інтелекту використовують алгоритми машинного навчання для аналізу вмісту та структури електронних листів, щоб виявити потенційні фішингові атаки. Ці алгоритми можуть навчатися на величезних обсягах даних, щоб виявляти закономірності та аномалії, які вказують на фішингову атаку.

Рішення на основі штучного інтелекту також можуть аналізувати поведінку користувачів під час взаємодії з електронними листами, щоб виявити потенційні фішингові атаки. Наприклад, якщо користувач натискає на підозріле посилання або вводить особисту інформацію у відповідь на фішинговий електронний лист, рішення на основі штучного інтелекту можуть позначити цю активність і попередити команди безпеки.

Традиційний аналіз журналів безпеки спирається на системи, засновані на правилах, які обмежені у своїй здатності виявляти нові загрози.

Аналіз журналів безпеки на основі штучного інтелекту використовує алгоритми машинного навчання, які можуть аналізувати великі обсяги даних журналу безпеки в режимі реального часу.

Алгоритми штучного інтелекту можуть виявляти закономірності та аномалії, які можуть вказувати на порушення безпеки, навіть за відсутності відомої сигнатури загрози. Таким чином, організації можуть швидко виявляти потенційні інциденти безпеки та реагувати на них, зменшуючи ризик витоку даних та інших інцидентів безпеки.

Аналіз журналу безпеки на основі штучного інтелекту також може допомогти організаціям виявити потенційні внутрішні загрози. Аналізуючи поведінку користувачів у кількох системах і програмах, алгоритми штучного інтелекту можуть виявляти аномальну поведінку, яка може вказувати на внутрішні загрози, такі як несанкціонований доступ або незвичайна передача даних. Після цього організації можуть вжити заходів, щоб запобігти витоку даних та іншим інцидентам безпеки до того, як вони відбудуться [20].

Аналіз журналів безпеки на основі штучного інтелекту надає організаціям

потужний інструмент для виявлення потенційних загроз і вжиття заходів для їх пом'якшення.

Алгоритми штучного інтелекту можна навчити відстежувати мережі на предмет підозрілої активності, виявляти незвичайні моделі трафіку та виявляти пристрої, які не мають права перебувати в мережі.

Штучний інтелект може підвищити безпеку мережі за допомогою виявлення аномалій. Це передбачає аналіз мережевого трафіку для виявлення закономірностей, які виходять за рамки норми. Аналізуючи історичні дані про трафік, алгоритми штучного інтелекту можуть дізнатися, що є нормальним для певної мережі, і визначити трафік, який є аномальним або підозрілим. Це може включати незвичайне використання портів, незвичне використання протоколу або трафік із підозрілих IP-адрес.

Штучний інтелект також може підвищити безпеку мережі, контролюючи пристрої в мережі. Алгоритми штучного інтелекту можна навчити виявляти пристрої, які не мають права перебувати в мережі, і попереджати команди безпеки про потенційні загрози [20-21].

Наприклад, якщо в мережі виявлено новий пристрій, який не був авторизований IT-відділом, система штучного інтелекту може позначити його як потенційну загрозу безпеці. Штучний інтелект також можна використовувати для моніторингу поведінки пристроїв у мережі, наприклад, незвичайних моделей активності, для виявлення потенційних загроз.

Однією з ключових переваг рішень безпеки кінцевих точок на основі штучного інтелекту є їхня здатність адаптуватися та розвиватися з часом. У міру того, як кіберзагрози розвиваються і стають все більш складними, алгоритми штучного інтелекту можуть навчатися на нових даних і виявляти нові закономірності, які вказують на потенційні загрози. Це означає, що рішення для захисту кінцевих точок на основі штучного інтелекту можуть забезпечити кращий захист від нових і невідомих загроз, ніж традиційне антивірусне програмне забезпечення.

Рішення для захисту кінцевих точок на основі штучного інтелекту

забезпечують захист у реальному часі. Алгоритми штучного інтелекту можуть аналізувати поведінку кінцевих точок у режимі реального часу та попереджати команди безпеки про потенційні загрози. Це означає, що команди безпеки можуть швидше реагувати на загрози та запобігати їх завданню шкоди.

Кібердодатки штучного інтелекту пропонують значні переваги для уряду та бізнес-лідерів, відповідальних за захист людей, систем, організацій та громад від сучасних кіберсупротивників. Діючи як мультиплікатор сили для досвідчених кіберпрофесіоналів.

Негайні та довгострокові переваги інтеграції штучного інтелекту в екосистему кібербезпеки організації полягають у тому, що він:

- покращує захист і виправлення завдяки здатності виявляти тонкі атаки, підвищує безпеку та покращує реагування на інциденти;
- прискорює час циклу виявлення та реагування, швидко кількісно оцінюючи ризики та прискорюючи прийняття рішень аналітиками;
- зміцнює захист репутації бренду та довіри до систем безпеки та протоколів організації;
- підвищує задоволеність працівників, оскільки фахівці з кібербезпеки можуть зосередитися на завданнях вищого рівня, а не на трудомістких ручних діях.

Майбутнє кібербезпеки буде забезпечене, якщо буде інтегроване зі AI. Штучний інтелект відповідає потребам сучасних підвищених вимог безпеки. Для урядових організацій, які потребують найвищого рівня захисту кібербезпеки, зокрема в оборонних установах та органах національної безпеки, штучний інтелект розширює можливості.

Завдяки автоматизації штучний інтелект забезпечує конкурентну перевагу влюбій галузі людської діяльності, в тому числі і у кібербезпеці. Застосування функцій штучного інтелекту, інтеграція можливостей штучного інтелекту в ручні та напівручні процеси дозволить мінімізувати людські помилки та невідповідності.

Для кіберпрофесіоналів будуть затребувані нові набори навичок.

Організації найматимуть експертів зі штучного інтелекту за їхній досвід застосування технологій штучного інтелекту та машинного навчання для кібербезпеки, а не просто шукатимуть традиційні набори кібернавичок.

За останні 4 роки угоди у сфері кібербезпеки та штучного інтелекту зросли більш ніж удвічі, і здатність штучного інтелекту підтримувати кібербезпеку з часом лише зростатиме. Керівники з кібербезпеки повинні вивчити широту варіантів використання штучного інтелекту та потенційні застосування у місії безпеки.

Функції AI протягом життєвого циклу кібербезпеки включають моніторинг величезних масивів даних для виявлення нюансів кібератак, кількісну оцінку ризиків, пов'язаних із відомими вразливими місцями, і прийняття рішень на основі даних під час пошуків загроз [22].

IBM Security надає трансформаційні рішення на основі штучного інтелекту, які оптимізують час аналітиків за рахунок прискорення виявлення загроз, прискорення реагування та захисту ідентифікаційних даних користувачів і наборів даних, водночас зберігаючи команди з кібербезпеки в курсі та контролюючи. IBM Security QRadar Suite – одне з таких рішень, яке охоплює технології виявлення, дослідження та реагування на загрози та побудоване на відкритій основі, яка може задовольнити вимоги кіберзахисту (рис. 1.8).

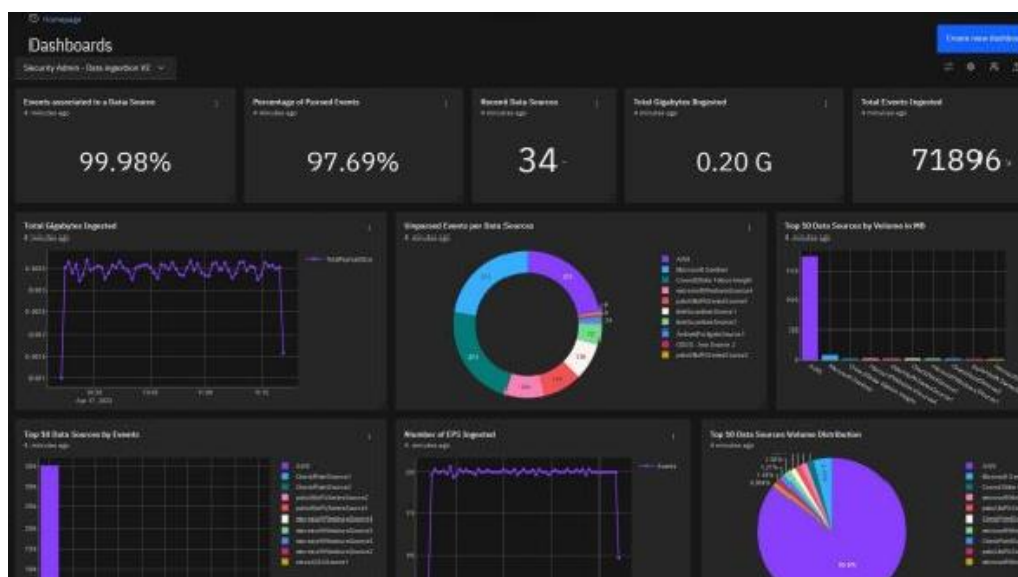


Рис. 1.8. Інтерфейс журналу статистики IBM Security QRadar

ВИСНОВОК до розділу 1. Стан кібербезпеки у 2023 році характеризується динамічним ландшафтом загроз, що розвивається, проблемами дотримання нормативних вимог і конфіденційності, а також зростаючим використанням штучного інтелекту та машинного навчання, серед інших тенденцій. Організації та окремі особи повинні залишатися пильними, проактивними та адаптивними, щоб ефективно реагувати на ці тенденції та пов'язані з ними виклики та захищатися від кіберзагроз. Бути в курсі новітніх технологій, вимог відповідності та найкращих практик, а також створювати кваліфіковану робочу силу з кібербезпеки, має вирішальне значення для ефективної кібербезпеки в нинішніх умовах.

Спеціалізовані заходи безпеки та найкращі практики допомагають організаціям встановити надійні позиції безпеки та захистити свої хмарні середовища від загроз і вразливостей. Однак важливо зазначити, що хмарна безпека є постійним зусиллям, і організації повинні постійно адаптувати свої методи безпеки для протидії загрозам, що розвиваються, і змінам у хмарних середовищах.

Сьогодні існує постійне занепокоєння щодо нестачі навичок у сфері кібербезпеки. Експертів з кібербезпеки просто не вистачає, щоб заповнити всі необхідні посади. У міру того, як створюється все більше компаній та інші оновлюють свої існуючі стратегії безпеки, це число зростає.

Сучасні загрози, від клонованих ідентичностей до дідфейкових кампаній, стає все важче виявити та зупинити. Навички безпеки, необхідні для боротьби з цими загрозами, виходять далеко за рамки простого розуміння того, як впровадити інструменти або налаштувати шифрування. Ці загрози вимагають різноманітних знань про широкий спектр технологій, конфігурацій і середовищ. Щоб отримати ці навички, організації повинні найняти експертів високого рівня або виділити ресурси на навчання власних.

РОЗДІЛ 2

МЕТОДОЛОГІЧНІ ПІДХОДИ ДО ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В УПРАВЛІННІ КІБЕРНЕТИЧНОЮ БЕЗПЕКОЮ

Управління кібербезпекою відноситься до стратегічних зусиль організації щодо захисту інформаційних ресурсів. Він зосереджений на тому, як компанії використовують свої активи безпеки, включаючи програмне забезпечення та рішення IT-безпеки, для захисту бізнес-систем.

Ці ресурси стають все більш вразливими до внутрішніх і зовнішніх загроз безпеці, таких як промислове шпигунство, крадіжки, шахрайство та саботаж. Управління кібербезпекою має використовувати різноманітні адміністративні, юридичні, технологічні, процедурні та службові практики, щоб зменшити ризики організацій і, особливо важливо на сьогоднішній день – застосування штучного інтелекту для підвищення ефективності кожної процедури та практик.

Хоча загальноприйнята структура кібербезпеки не встановлена, існують деякі керівні принципи, запобіжні заходи та технології, які багато організацій вирішили прийняти (рис. 2.1): Open Web Application Security Project (OWASP); Програма Національного інституту стандартів і технологій (NIST); Міжнародна організація зі стандартизації (ISO) серії 27000.

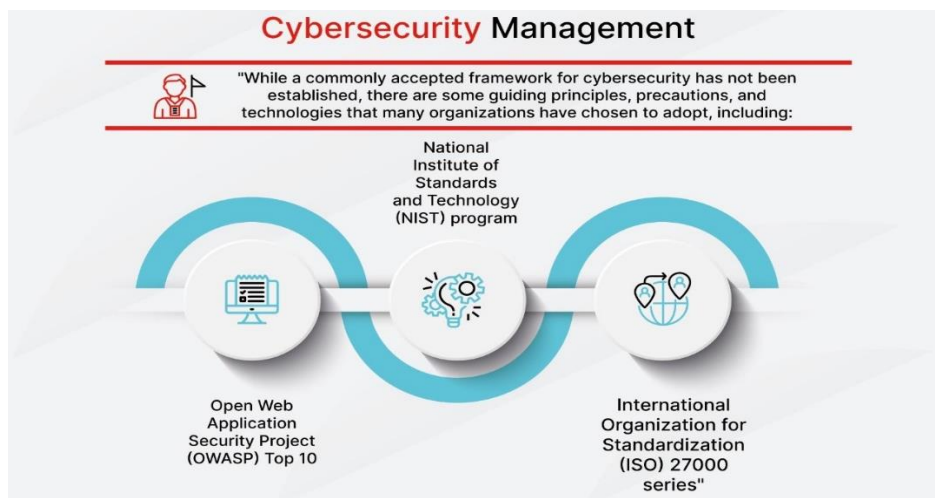


Рис. 2.1. Міжнародні організації, які відпрацювали принципи керівництва кібербезпекою

Вони служать де-факто основою для управління кібербезпекою, а також окреслюють методи та стандарти захисту цифрових активів.

Традиційна кібербезпека також спиралася на ручний аналіз. Аналітики безпеки вручну досліджуватимуть сповіщення та журнали безпеки, шукаючи закономірності або індикатори порушення безпеки. Цей процес був трудомістким і часто покладався на досвід аналітика безпеки для виявлення загроз.

Системи, засновані на правилах, працювали, встановлюючи правила або політики, які визначали прийнятну поведінку в мережі. Якщо дорожній рух порушує ці правила, він активує тривогу. Хоча системи, засновані на правилах, можуть бути ефективними в певних ситуаціях, вони часто були негнучкими і не могли адаптуватися до нових і нових загроз.

Традиційний підхід до кібербезпеки до впровадження штучного інтелекту був значною мірою реактивним, покладаючись на ручний аналіз, сигнатурні системи виявлення та системи, засновані на правилах. Цей підхід часто був неефективним проти нових і невідомих загроз, і він міг генерувати велику кількість помилкових спрацьовувань, що могло призвести до виснаження ресурсів.

2.1 Аналіз сучасних тенденцій у використанні штучного інтелекту в кібернетичній безпеці

Штучний інтелект (ШІ або англ. AI) – це потужна технологія, яка допомагає командам з кібербезпеки автоматизувати повторювані завдання, прискорити виявлення загроз і реагування на них, а також підвищити точність своїх дій для посилення захисту від різних проблем безпеки та кібератак.

Для здійснення аналізу проведено дослідження [21] з використанням методу таксономії, яке виявило сучасні тенденції використання штучного інтелекту у кібербезпеці.

Враховуючи ефективність та сучасний стан розвитку, AI сьогодні почав використовуватися у всіх сферах людської діяльності, в тому числі і у інформаційній і кібербезпеці. Основні напрямки його застосування показані на рис. 2.2.

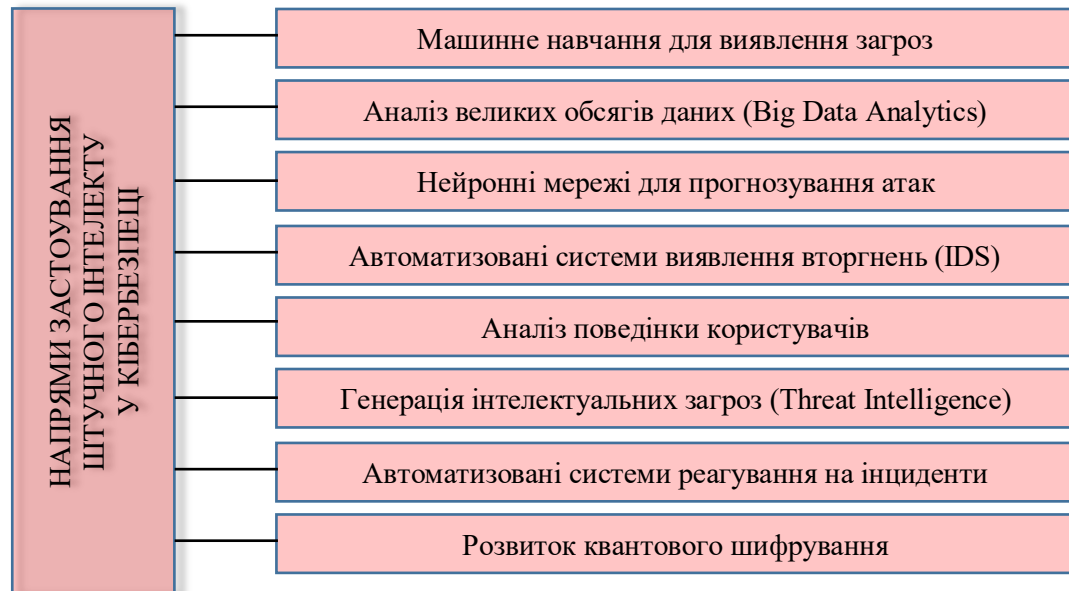


Рис. 2.2. Основні напрямки застосування штучного інтелекту в кібербезпеці

Машинне навчання для виявлення загроз застосовується для :

- використання алгоритмів класифікації, наприклад, дерева рішень, метод опорних векторів для виявлення вразливостей та загроз (Класифікація);.
- групування схожих даних для виявлення несподіваних аномалій та вразливостей (Кластеризація).

При аналіз великих обсягів даних (Big Data Analytics) AI використовує технології обробки великих обсягів даних для виявлення патернів та аномалій, що вказують на можливі атаки.

Для прогнозування можливих атак на основі аналізу змін у мережевому трафіку та системних журналах використовуються глибокі нейронні мережі.

AI використовується для розробки та підтримки систем IDS, які здатні ефективно виявляти та відповідати на вторгнення в реальному часі.

Для виявлення аномальної поведінки користувачів, що може свідчити про

компрометацію облікового запису, підготовку терористичного акту на об'єкти критичної інфраструктури використовуються алгоритми машинного навчання (семантичний та сентиментний аналіз).

Також AI використовується для автоматизації процесу збору, аналізу інформації про загрози для покращення захисту мереж та систем, для автоматизації процесу виявлення, аналізу та відповіді на кібератаки в реальному часі.

Особливо новим є дослідження застосування квантового шифрування для забезпечення безпеки зв'язку та обміну інформацією.

Дослідники повідомляють про постійну еволюцію кіберзагроз, пов'язаних з національними державами та злочинними угрупованнями, а також про зростаючу складність кібератак, які знаходять нові та інвазивні способи націлитися навіть на найрозумніші цілі [22]. Ця еволюція призводить до збільшення кількості, масштабу та впливу кібератак, а також вимагає впровадження кібербезпеки на основі аналітики для забезпечення динамічного захисту від кібератак, що розвиваються, та управління великими даними. Консультативні організації, такі як Національний інститут стандартів і технологій (NIST), також заохочують використання більш проактивних та адаптивних підходів, переходячи до оцінки в режимі реального часу, безперервного моніторингу та аналізу на основі даних для виявлення, захисту, виявлення, реагування та каталогізації кібератак для запобігання майбутнім інцидентам безпеки [23].

Функція ідентифікації за фреймворком NIST забезпечує основу для інших функцій кібербезпеки, точно визначаючи критичні функції та ризики, пов'язані з системами, людьми, активами та даними. Це допомагає розвинути розуміння поточного стану кібербезпеки, виявити прогалини та розробити відповідну стратегію управління ризиками для досягнення бажаної безпеки на основі власних потреб, ризиків та бюджету організації. Початком ідентифікації є управління активами.

Управління активами - це процес виявлення та відстеження інформації,

людей, обладнання, систем та будівель, які допомагають організації досягати своїх цілей і є пропорційними відносній важливості активу цим цілям та стратегіям ризику. Вона включає в себе виявлення, інвентаризацію, управління та відстеження активів для їх захисту. Управління активами кібербезпеки стає все більш складним, оскільки організації мають більше платформ, ніж будь-коли раніше, від систем операційних технологій та Інтернету речей (IoT) до локальних і хмарних сервісів. Таке поширення нових типів активів і можливість віддаленої роботи призвели до появи високорозподілених активів, якими важко управляти та інвентаризувати. Тому система управління активами на основі штучного інтелекту може вирішити багато з цих проблем, надаючи нові рівні інтелекту людській команді в наступних випадках використання.

Виявлення та управління активами мають вирішальне значення для забезпечення повної видимості та контролю над усіма активами в розширеній мережі. Штучний інтелект може допомогти з безперервним і автоматичним виявленням усіх пристроїв, додатків і користувачів, а також їх класифікацією та критичністю для операцій. Завдяки точній та актуальній інвентаризації активи можна відстежувати та аналізувати для оцінки ризиків за відомими векторами атак, а монітори відповідності можуть виявляти шахрайські активи та несанкціоноване використання.

Дослідники розробили різні підходи до класифікації активів за допомогою алгоритмів машинного навчання. Наприклад, у дослідженні [24] використовувалася кластеризація K-середніх для класифікації активів відповідно до їхніх вимог до кібербезпеки на основі їхньої безпеки, функціональності та цілісності на атомній електростанції. У дослідженні [25] запропоновано випадковий класифікатор машинного навчання на основі дерева рішень для класифікації операційної системи та ідентифікації вразливих пристроїв у мережі. Кілька досліджень [26-27] зосереджені на ідентифікації та класифікації пристроїв IoT на основі їх характеристик мережевого трафіку. У дослідженні [28] запропоновано рішення проблеми класифікації в швидко мінливому, гетерогенному і динамічному середовищі за допомогою методу контрольованого

машинного навчання, здатного призначати пристрої IoT заздалегідь визначеним класам на основі значень їх потоків трафіку.

Особливе значення має валідність контролю інформаційної та кібербезпеки підприємств та організацій. Автоматизація валідації контролю безпеки забезпечить моніторинг безпеки в режимі реального часу в мінливому середовищі та ландшафті загроз. Дослідники працюють над впровадженням методів штучного інтелекту для остаточної оцінки загальної безпеки системи з використанням даних прогнозу мережі, структури кібербезпеки і шляхом кореляції загроз, вразливостей та заходів безпеки [29-30].

Аналіз впливу на бізнес є найважливішим методом визначення критичних функцій і додатків у бізнес-середовищі шляхом оцінки впливу інцидентів кібербезпеки на бізнес. Методи штучного інтелекту можна використовувати для автоматизації аналізу впливу на бізнес, оцінюючи економічні ризики на основі відомого вектора атаки або розраховуючи здійсненність загрози разом із ймовірністю подій безпеки з високим впливом на критично важливі сфери бізнесу. Дослідники вимірюють економічний ризик кібербезпеки в різних компаніях, використовуючи моделювання різних відомих профілів атак [31], моделювання рідкісних подій [32] або пов'язуючи бізнес-мету з можливостями зловмисника, щоб керувати сценарним аналізом [33] для визначення його впливу на бізнес-активи.

Управління включає політику, процедури та процеси для розуміння екологічних та операційних вимог, а також моніторинг нормативних вимог організації. Це допомагає визначити обов'язки організації та надає керівництву інформацію про кіберризики. Штучний інтелект можна використовувати для забезпечення дотримання політики або автоматизації пошуку ключових індикаторів ризику. Незважаючи на те, що існують деякі дослідження щодо дотримання політики, у цьому дослідженні не було знайдено жодної дослідницької статті щодо вимірювання показників ризику в режимі реального часу.

Дотримання політики має вирішальне значення для організацій, щоб

забезпечити дотримання ними нормативних актів та належне управління ризиками. Штучний інтелект використовується в автоматизованому застосуванні політик у традиційних мережах, відмінних від SDN, за допомогою контролера та проксі-серверів [34-35]. Контролер – це централізований сервер керування, який використовується для керування програмно визначеними проміжними коробками для традиційних маршрутизаторів і проксі-сервером політики, який ідентифікує трафік, який підпадає під дію політик, і допомагає йому в застосуванні політик.

Оцінка ризиків – це процес виявлення, оцінки та пріоритизації ризиків кібербезпеки, пов'язаних з операціями, операційними активами та особами в даний час або в найближчому майбутньому. Це вимагає ретельного аналізу інформації про загрози, вразливості та атаки, щоб визначити, якою мірою події кібербезпеки можуть негативно вплинути на організацію, і ймовірність того, що такі події відбудуться. Процес ручної оцінки ризиків є складним, дорогим і трудомістким через велику кількість факторів ризику і вимагає активної участі людини на кожному етапі. Процес оцінки ризиків на основі штучного інтелекту вирішує ці проблеми, підтримуючи команду управління ризиками в наступних випадках використання.

Автоматизована оцінка вразливостей – це процес систематичного аналізу слабких місць безпеки в системі за допомогою автоматизованих інструментів для ідентифікації, класифікації, дослідження та пріоритезації вразливостей. Ці автоматизовані інструменти покладаються на репозиторії вразливостей, оголошення про вразливості постачальників, системи управління активами та канали розвідки загроз для виявлення, класифікації та оцінки серйозності, а також надання рекомендацій щодо виправлення.

Автоматизоване виявлення вразливостей - це важливий крок для виявлення вразливостей програм, серверів або інших систем і активів організації. Дослідники працювали над виявленням вразливостей програмного забезпечення, перевіряючи вихідний код за допомогою глибокого навчання та трансферного навчання [36-38]. У цих дослідженнях використовуються методи

інтелектуального аналізу тексту, щоб подавати моделі виявлення вразливостей на основі машинного навчання в унісон із системою рекомендацій, щоб допомогти програмістам писати безпечний код.

Дослідники також працювали над виявленням нових вразливостей програмного забезпечення [39] та кібернетичної інфраструктури [40] за допомогою репозиторіїв вразливостей або соціальних мереж відповідно. У праці дослідники [41] запропонували нову схему виявлення вразливостей на системному та мережевому рівнях шляхом моделювання поведінки кіберфізичних систем (CPS)/IoT під час атаки на системному та мережевому рівнях, а потім використання машинного навчання для виявлення будь-якого потенційного простору атаки.

Дослідники використовують фаззинг на основі штучного інтелекту для виявлення вразливостей у програмних та апаратних інтерфейсах і додатках. Це робиться шляхом введення помилкових, несподіваних або випадково згенерованих даних у програму або інтерфейс, а потім моніторингу таких подій, як збої, невдалі твердження коду, недокументовані стрибки або процедури налагодження, а також можливі витoki пам'яті. Методи штучного інтелекту використовуються для розробки автоматизованої системи для визначення потенційних варіантів атаки, генерації вхідних даних, генерації ймовірних тестових випадків та аналізу збоїв, як показано на рис. Науковці використовують міркування та обробку природної мови для методів генерації тексту, щоб збільшити охоплення коду з більш унікальними шляхами виконання, як один з основних кроків розумної системи фаззингу. Генерація тест-кейсів є однією з найбільш широко вивчених областей фаззингу на основі штучного інтелекту для веб-браузерів, компіляторів, кіберфізичних систем (CPS), програмних бібліотек і простих комп'ютерних програм [42-45].

Автоматизоване тестування на проникнення – це ефективна та розумна спроба визначити початок атаки шляхом використання відомих вразливостей або вразливостей нульового дня, щоб зрозуміти, що зловмисник може отримати від поточних середовищ. Дослідники працюють над розробкою автономного

тестування на проникнення з використанням навчання з підкріпленням для великих мереж [46-47] та алгоритмів керування мікромережами [48].

Класифікація вразливостей є важливим кроком, який сприяє глибшому розумінню інформації, пов'язаної з безпекою, для прискорення оцінки вразливостей. Дослідники працюють над системами автоматичної класифікації вразливостей і для маркування описів вразливостей у звітах про вразливості. Дослідники у своїй роботі [49] запропонували підхід до узагальнення щоденних опублікованих вразливостей і класифікації їх відповідно до визначеної моделі таксономії для галузі. У дослідженні [50,51] науковці розглянули використання інтелектуального аналізу тексту для класифікації вразливостей за допомогою опису, наведеного в списку Common Vulnerabilities and Exposures (CVE).

Моделювання шляху атаки – це проактивний підхід до зниження ризиків, який підтримує команди безпеки, картографуючи вразливі маршрути в мережі для оцінки ризиків, виявлення вразливостей і вжиття контрзаходів для захисту ключових активів. Дослідники використовували методи штучного інтелекту для моделювання шляху з використанням сповіщень про вторгнення [52] або описів вразливостей [53,54]. Деякі дослідники використовували всі кібердані, включаючи оповіщення, вразливості, журнали та мережевий трафік, щоб імітувати дії зловмисника/захисника та вживати превентивних заходів у режимі реального часу [55-57].

Автоматизований аналіз ризиків та оцінка впливу зміцнюють команду з управління ризиками завдяки розумному використанню даних про ризики, доступних всередині компанії та ззовні, для оцінки ризиків та пов'язаних з ними показників у режимі реального часу. Штучний інтелект є каталізатором для прискорення прогресу в управлінні ризиками шляхом автоматизації розрахунку оцінки ризику, висновків про ймовірність інциденту безпеки, визначення ключових показників ризику вразливостей, а також оцінки ризиків та аналізу рішень.

Прогностична розвідка – це розвідувальна інформація, яка є дієвою та актуальною в певному контексті та може бути використана для передбачення

атак. Інструменти прогнозування вторгнень допомагають забезпечити активний захист від майбутніх атак, заздалегідь передбачаючи тип, інтенсивність і ціль вторгнення. Дослідники використовують підходи глибокого навчання [58-60] для прогнозування сповіщень від шкідливих джерел або на задану ціль, використовуючи послідовність попередніх сповіщень, історичні повідомлення електронної пошти зі спамом та мережу трафік дані.

Прогнозування зловмисного програмного забезпечення включає методи прогнозування та блокування шкідливих файлів до того, як вони повністю виконають своє корисне навантаження, щоб запобігти атакам зловмисного програмного забезпечення, а не виправити їх. У цьому напрямку в дослідженні [61] розроблена модель прогнозування шкідливого програмного забезпечення, заснована на моделі рекурентних нейронних мереж (RNNs) для прогнозування шкідливої поведінки з використанням даних про машинну активність.

Вважається, що прогнозування атак має чудовий потенціал для проактивного підвищення кіберстійкості. Дослідники представляють схеми прогнозування атак, використовуючи різні типи даних, отриманих з новинних сайтів і веб-сайтів, форумів даркнету, національних баз даних вразливостей, звітів про інциденти і загальних баз даних вразливостей і вразливостей [62-64].

Стратегія управління ризиками допомагає приймати рішення про операційні ризики, встановлюючи пріоритети, толерантність до ризику та обмеження. Вона повинна забезпечити встановлені та задокументовані прийнятні рівні ризику а також розумні терміни вирішення та інвестиції. Штучний інтелект має потенціал трансформувати цю категорію, автоматизувавши такі види діяльності.

Планування ризиків кібербезпеки передбачає впровадження бажаного портфеля контрзаходів у межах заздалегідь визначеного бюджету. Офіційні системи підтримки прийняття рішень та моделювання атак можуть допомогти фахівцям з планування безпеки проводити економічні порівняння з вартістю контрзаходів та наявними бюджетами ризиків. Проблема прийняття рішень при плануванні ризиків кібербезпеки є важливою через чутливість плану ризиків до

ставлення особи, яка приймає рішення, до ризику та його взаємодії з наявним бюджетом. Таким чином, є дуже важливим впровадження системи підтримки прийняття рішень для оцінки невизначеного ризику, з яким стикається організація під час кібератаки, з урахуванням невизначених рівнів загроз, витрат на контрзаходи та впливу на її активи. У цьому напрямку використовують генетичний алгоритм для пошуку найкращої комбінації контрзаходів для блокування або пом'якшення атак безпеки, що дозволяє бізнесу визначити бажаний компроміс між інвестиційними витратами та результируючим ризиком.

Функція захисту допомагає планувати та впроваджувати відповідні засоби контролю, щоб обмежити або стримати вплив потенційної події кібербезпеки. Це включає низку технічних і процедурних заходів контролю для проактивного захисту від внутрішніх і зовнішніх кіберзагроз. Штучний інтелект може підвищити стійкість системи шляхом автентифікації користувачів, пристроїв та інших активів, моніторингу поведінки користувачів, автоматизованого контролю доступу, адаптивного навчання, запобігання витоку даних і моніторингу цілісності, автоматизованого захисту інформації та процесів, а також надання захисних рішень для проактивного захисту системи.

Управління ідентифікацією, автентифікація та контроль доступу відповідають за обмеження доступу до активів і пов'язаних з ними об'єктів авторизованим користувачам, процесам або пристроям, а також до авторизованої діяльності. Штучний інтелект можна використовувати для управління та захисту фізичного та віддаленого доступу за допомогою інтелектуальної автентифікації користувачів, інтелектуальної автентифікації пристроїв, автоматизованого контролю доступу за допомогою авторизацій та прав доступу для запобігання несанкціонованому доступу та його наслідкам.

Штучний інтелект може покращити автентифікацію користувачів за допомогою фізичної біометрії, поведінкової біометрії або багатофакторної автентифікації захищаючи імена користувачів, паролі і навіть одноразові текстові токени.

Фізична біометрія відноситься до вроджених фізичних характеристик

користувачів, які можуть бути використані для ідентифікації, таких як відбитки пальців, райдужна оболонка ока та біосигнали (рис. 2.3).

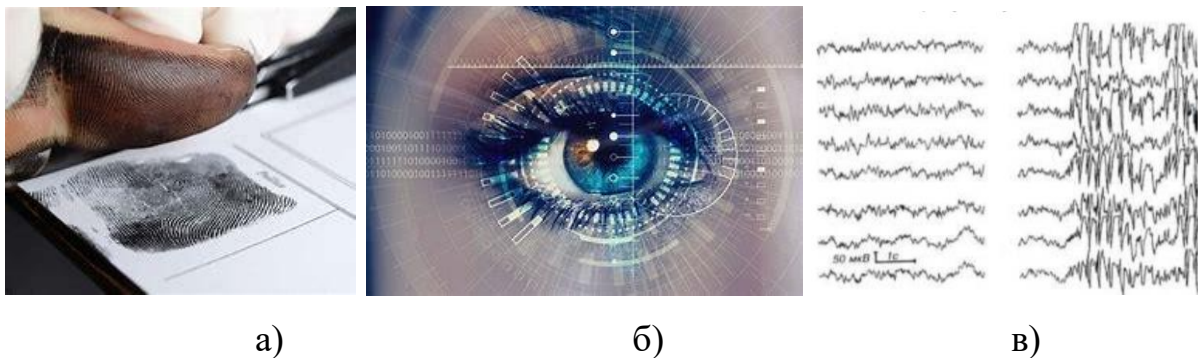


Рис. 2.3. Приклади фізичної біометрії (а – відбитки пальців; б - райдужна оболонка; в – біосигнали)

У 2020 році дослідники представили систему біометричної автентифікації людини на основі PPG (фотоплетизмографії) з використанням глибокого навчання [65].

На противагу фізичній, поведінкова біометрія пов'язана з унікальними ідентифікованими та вимірюваними моделями людської діяльності, які можуть забезпечити безперервну та зручну безпеку. Існує кілька поведінкових біометричних показників, включаючи поведінку користувача і процесу. Використання моделей поведінки, пов'язаних із взаємодією користувача з власним пристроєм, є основною основою систем безперервної автентифікації. У цьому напрямку дослідники запропонували використовувати мобільні функції та дані про використання, наприклад, акселерометр, гіроскоп, магнітометр та статистику взаємодії з різними програмами, щоб визначити, чи є поточний користувач тим самим, що й той, хто пройшов попередню автентифікацію.

У 2019 році розробили та впровадили модель безперервної автентифікації, заснований на поведінкових профілях користувачів відповідно до їхньої взаємодії з різними офісними пристроями для розумних офісів, використовуючи парадигму хмарних обчислень та алгоритм випадкового лісу (RF) (рис. 2.4) [66].

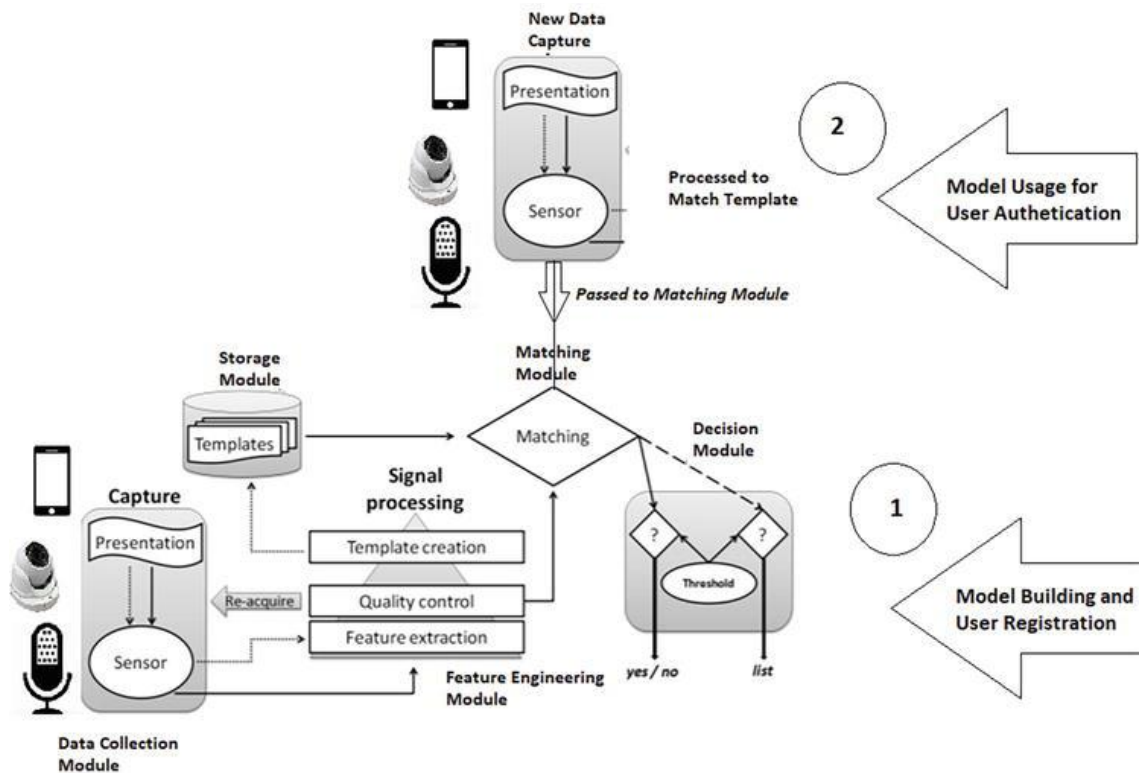


Рис. 2.4. Модель поведінкової автентифікації [66]

Крім того, ці рішення поступово використовуються в управлінні ідентифікацією, що ще більше підвищує інтерес до них. Автентифікація – це ненав'язливий, прозорий і безперервний метод автентифікації в мобільних пристроях, який фіксує інформацію, необхідну для перевірки автентичності користувача під час ходьби людини.

Багатофакторна автентифікація – це багаторівневий підхід до захисту даних і програм, який вимагає комбінації двох або більше облікових даних для перевірки особи або входу користувачів. Дослідники [44, 57] використовували динаміку натискання клавіш як другий рівень безпеки для автентифікації веб-користувачів і автентифікації пристроїв.

Інтелектуальний захист електронної пошти – це клас програмних рішень для запобігання складним кібератакам, орієнтованим на електронну пошту. Традиційно спам-розсилка використовувалася як тактика для просування товарів і послуг шляхом розсилки небажаних електронних листів у масові списки. Однак сьогодні він активно використовується для поширення шкідливого програмного

забезпечення, крадіжки облікових даних автентифікації або фінансових махінацій. Методи штучного інтелекту використовуються для автоматизованого захисту від шкідливих електронних листів зі спамом. Для виявлення спаму в режимі реального часу за допомогою динамічних вхідних даних електронної пошти, включаючи загальні, представлення, теми, вкладення та пов'язані з контентом функції, використовують методи контрольованої класифікації і глибокого навчання.

Методи штучного інтелекту використовуються для динамічного планування резервного копіювання та оптимізованого планування резервного копіювання.

2.2 Застосування машинного навчання для виявлення кібератак

Дуже важливим і ефективним у кібербезпеці є організація та здійснення превентивних заходів шляхом розвідувальних дій у кіберпросторі.

Функція виявлення дозволяє своєчасно виявляти події кібербезпеки, розробляючи та впроваджуючи відповідні заходи для ідентифікації їх виникнення. Ця функція має вирішальне значення для безпеки, оскільки швидке виявлення зведе до мінімуму збої в захисті інформаційної системи. Вона включає заходи щодо своєчасного виявлення вторгнень та аномалій, а також оцінку впливу, впровадження безперервного моніторингу безпеки для перевірки ефективності захисних заходів та відповідне обслуговування процесів виявлення для забезпечення обізнаності про кіберподії. Штучний інтелект може підвищити швидкість виявлення, відстежуючи внутрішні та зовнішні джерела інформації та швидко співвідносячи цю інформацію для виявлення незвичайних дій, щоб мінімізувати наслідки.

Одним із відомих інструментів, які використовують організації для виявлення кіберзагроз є системи UEBA - Аналітика поведінки користувачів і сутностей, які зосереджені на моніторингу та аналізі поведінки користувачів в організації з метою виявлення зловмисної або несанкціонованої активності (рис.

2.5).



Рис. 2.5. Зміст роботи систем UEBA

UEBA можна використовувати для доповнення або заміни традиційних методів безпеки, таких як системи виявлення вторгнень (IDS), надаючи більш повне уявлення про діяльність в організації. Системи UEBA використовують алгоритми машинного навчання для аналізу даних з різних джерел, а потім переходять до використання цих даних для створення моделей нормальної поведінки користувачів, які можуть бути використані для виявлення аномалій і шкідливої активності. UEBA також може використовуватися для контролю за дотриманням внутрішніх політик і зовнішніх правил.

Переваги використання UEBA включають покращене виявлення складних атак, зменшення помилкових спрацьовувань і можливість проактивного виявлення ризиків. Зловмисні інсайдери – визначаючи базовий рівень поведінки користувачів, UEBA може виявляти аномальну активність і допомагати в інтерпретації намірів. Наприклад, користувач може мати справжні привілеї доступу, але не потребувати доступу до конфіденційних даних у певний час або в певному місці.

Скомпрометовані інсайдери – користувачі з привілеями доступу можуть бути скомпрометовані через зловмисне програмне забезпечення або спроби фішингу, що дозволяє використовувати їхні облікові дані для ініціювання атаки. Зловмисники часто змінюють облікові дані, IP-адреси або пристрої, потрапивши в систему. Порівнюючи поведінку пристроїв і користувачів з базовими показниками, UEBA може ідентифікувати ці атаки так, як це не можуть зробити традиційні інструменти безпеки, такі як брандмауери та антивіруси.

Викрадення даних – такі інструменти, як запобігання втраті даних (DLP), які використовують машинне навчання, словникові моделі та моделі поведінки для збору всіх доказів, пов'язаних із викраденням конфіденційних даних, можуть швидко розслідувати та попереджати про аномальну активність. Це включає завантаження даних, віддалені входи, дії з базами даних, доступ до хмари та доступ до спільного доступу до файлів.

Бічний рух – зловмисники часто обходять мережу, використовуючи різноманітні IP-адреси, облікові дані та машини в пошуках ключових активів і даних. Інструменти UEBA виявляють цей рух, збагачуючи дані контекстом, що дозволяє їм розрізняти сервери, користувачів, сервісні облікові записи, персонал відділу кадрів, фінансовий персонал і керівників і визначати, чи поведуться вони підозріло.

Алгоритми навчання, які використовує UEBA, призначені для виявлення закономірностей. Такими закономірностями можуть бути незвично високий рівень активності під час зазвичай неактивних періодів або доступ до конфіденційних даних, до яких користувач зазвичай не має доступу. Прикладом незвичайної активності може бути конкретний користувач, який регулярно завантажує, скажімо, від 10 до 20 МБ файлів щодня, але раптом починає завантажувати гігабайти файлів. Показник ризику користувача або пристрою для події визначається та зшивається з пов'язаними подіями в часову шкалу, щоб оцінити, чи становлять ці події загрозу для організації. Пов'язуючи воєдино поведінку, визначену як аномальну, аналітики можуть відстежувати всі кроки, зроблені зловмисником, і, таким чином, швидко визначати загрозу. У цьому

сценарії система виявить аномалію та попередить вас про потенційну загрозу безпеці. Це дозволяє організаціям швидко виявляти та реагувати на потенційні загрози, які в іншому випадку могли б залишитися непоміченими.

Системи UEBA відстежують не тільки активність користувачів. Вони відстежують активність пристроїв, додатків, серверів і даних. Замість того, щоб аналізувати лише дані про поведінку користувачів, ця технологія поєднує дані про поведінку користувачів із даними про поведінку сутностей.

На відміну від UEBA, системи SIEM зазвичай використовуються для виявлення проблем із відповідністю та використовують статистичні моделі для виявлення підозрілих моделей поведінки, пов'язаних із кібератаками. Продукти SIEM можуть перевантажувати команди інформаційної безпеки великим обсягом сповіщень, багато з яких можуть бути помилковими спрацьовуваннями. SIEM є хорошим інструментом управління безпекою, але він менш складний, коли справа доходить до виявлення та реагування на більш просунуті загрози. Як наслідок, досвідчені суб'єкти загроз не беруть участь у одноразових загрозах, а виконують довгострокові атаки, які можуть залишатися непоміченими протягом кількох тижнів або навіть місяців за допомогою традиційних інструментів управління загрозами.

Інструменти UEBA та SIEM можна поєднувати для кращого захисту підприємств від широкого спектру загроз. Оскільки UEBA, як більш інтелектуальніша, менше зосереджується на системних подіях, а більше на конкретних діях користувачів або організацій, вона створює профіль співробітників і організацій на основі моделей використання та попереджає користувачів, коли спостерігається незвичайна або підозріла поведінка.

Крім того, для більшого захисту інформаційної безпеки доцільне поєднання UEBA та Security, Orchestration, Automation, and Response (SOAR). Таке поєднання дозволяє автоматизувати процеси виявленні активності загроз, пов'язані з її ідентифікацією та аналізом, підвищити ефективність команд безпеки.

Можливості UEBA також збігаються з інструментами XDR. Системи

розширеного виявлення та реагування (XDR) – це останнє покоління систем виявлення та реагування на кінцеві точки (EDR) та систем виявлення та реагування мережі (NDR). EDR фокусується на комп'ютерах кінцевих користувачів і серверах організації, тоді як NDR контролює мережеві передачі. XDR - це нова технологія, яка зближується з функціями UBA і UEBA, а також з системами SOAR і SIEM.

XDR об'єднує всі можливості EDR, UEBA, NTA, антивірусів нового покоління та інших інструментів в єдине рішення для забезпечення найкращої безпеки. Для того, щоб компанії могли виявляти, відстежувати, досліджувати та реагувати на потенційні загрози безпеці, технологія забезпечує всебічну видимість і надійну поведінкову аналітику по всій інфраструктурі, включаючи мережі, хмари та кінцеві точки.

Організації можуть впроваджувати UEBA різними способами, залежно від своїх потреб та ресурсів. Одним із підходів є використання платформи UEBA, яка надає готове рішення з попередньо налаштованими правилами та моделями. Це може бути корисно для організацій, які хочуть швидко розпочати роботу з UEBA і не мають ресурсів для створення власного рішення.

Іншим варіантом є створення кастомного рішення UEBA за допомогою інструментів з відкритим вихідним кодом. Цей підхід вимагає більше технічних знань, але дає організаціям більше гнучкості, щоб адаптувати рішення до своїх конкретних потреб.

Одним із важливих недоліків UEBA, які стосуються невеликих компаній є її велика вартість. Тому варіантом для них є розроблення власних систем виявлення загроз на основі машинного навчання.

Застосування машинного навчання для виявлення кібератак включає в себе розробку моделей, які можуть вчитися на основі історичних даних та виявляти патерни атак у нових даних. Методика застосування машинного навчання для виявлення кібератак наступна (рис. 2.6):

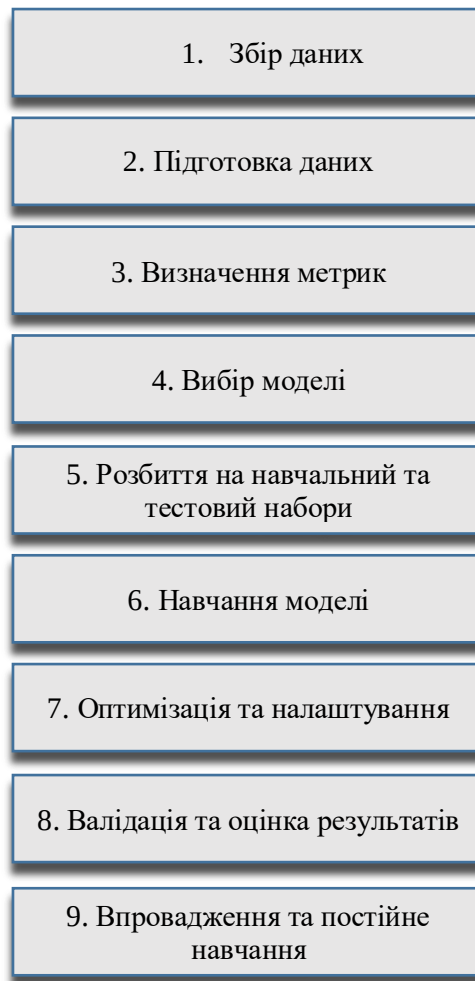


Рис. 2.6. Алгоритм методики застосування машинного навчання для виявлення кіберзагроз

1. Збір різноманітних даних, таких як журнали подій, дані мережевого трафіку, системні журнали, інформація про доступ та інші релевантні дані. Для автоматизації процесу та наступної інтеграції результатів збору даних в загальну модель виявлення загроз є розроблені на мові програмування Python стандартні бібліотеки:

- для роботи з мережевим трафіком *Scapy* для захоплення та аналізу мережевого трафіку. Це може бути корисно для виявлення незвичайної активності або атак в мережі;
- для роботи з журналами подій (Event Logs) - *win32evtlog* в операційних системах *Windows* або *syslog* для операційних систем на базі *UNIX*. Це може бути важливо для виявлення підозрілих подій;
- збір інформації про систему використовуйте бібліотеки, такі як *psutil*

для отримання інформації про процеси та ресурси системи, це може допомогти виявити незвичайну або підозрілу активність (лістинг програми для збору даних з використанням бібліотеки «*psutil*» в ДОДАТКУ А);

- взаємодія з API для отримання інформації з різних джерел можна використовувати бібліотеку *requests* для взаємодії з веб-сервісами та отримання актуальної інформації;
- для зберігання та обробки отриманих даних в базах даних бібліотеки *sqlite3*, *psycopg2* або інші;
- для автоматизації задач моніторингу бібліотеки *schedule*, для автоматизації виконання завдань збору даних та аналізу визначеними інтервалами часу;
- для збору даних з веб-сайтів - бібліотеки, такі як *requests*, *BeautifulSoup* або *Scrapy* (це може бути корисно для виявлення загроз, які можуть бути пов'язані з онлайн-активністю);
- для аналізу та витягування інформації з логів застосунків - регулярні вирази та бібліотеки, такі як *re*.

Python має широкую екосистему бібліотек і інструментів, що полегшує розробку різноманітних інструментів для збору та обробки даних для виявлення загроз.

2. Обробка та очищення даних, видалення аномалій, кодування категоріальних змінних та інші операції для підготовки даних до використання в моделях машинного навчання. Приклад простого Python-скрипта, який використовує бібліотеки *pandas* та *scikit-learn* для обробки та очищення даних, видалення аномалій і кодування категоріальних змінних в ДОДАТКУ Б із простим штучно згенерованим набором даних.

3. Визначення метрик, які будуть використовуватися для оцінки ефективності моделі (наприклад, точність, чутливість, специфічність).

4. Вибір відповідної моделі машинного навчання для конкретного завдання виявлення кібератак (наприклад, дерева рішень, метод опорних векторів, глибокі нейронні мережі тощо). Простий приклад використання

бібліотеки TensorFlow та Keras для створення Глибокої Нейронної Мережі (ГНМ) з одним прихованим шаром у ДОДАТКУ В. Це тільки початок, і реальні системи для виявлення кіберзагроз можуть бути набагато складнішими та потужнішими.

5. Розділення даних на навчальний та тестовий набори для тренування та оцінки ефективності моделі.

6. Використання навчального набору для тренування моделі, щоб вона могла виявляти патерни та аномалії, характерні для кібератак.

7. Оптимізація параметрів моделі та налаштування для досягнення максимальної ефективності (за визначеними в п.3 метриками - точність, чутливість, специфічність).

8. Використання тестового набору для валідації моделі та оцінки її результатів. Аналіз отриманих метрик для забезпечення адекватної ефективності та виявлення можливих покращень.

9. Впровадження та постійне навчання. Інтеграція моделі в систему виявлення кібератак для реального використання. Забезпечення постійного навчання моделі на нових даних для адаптації до нових типів атак та підтримки актуальності.

При цьому необхідно пам'ятати, що застосування машинного навчання для виявлення кібератак – це динамічний процес, і важливо постійно вдосконалювати моделі та методи для ефективного протистояння зростаючим кількісно та якісно кіберзагрозам.

2.3. Етичні проблеми застосування штучного інтелекту у кібербезпеці

Деякі з етичних проблем включають відсутність прозорості та пояснення, надмірну залежність від ШІ, упередженість та дискримінацію, вразливість до атаки, відсутність нагляду з боку людини, висока вартість і занепокоєння щодо конфіденційності. ШІ може ідентифікувати вразливі місця системи, які часто

вислизають від експертів-людей, і використовувати їх для атаки на певну ціль, що може завдати серйозної шкоди нашим інформаційним суспільствам 1. Потрібний етичний аналіз щоб уникнути декваліфікації людей, масового нагляду та ризиків через відсутність контролю систем штучного інтелекту. Розглядаючи роль штучного інтелекту в кібербезпеці на системному рівні, є три сфери великого впливу: надійність системи, стійкість системи, реакція системи. У всіх трьох випадках потенціал штучного інтелекту пов'язаний із серйозними етичними ризиками [67].

Оскільки організації впроваджують системи безпеки на основі штучного інтелекту, існує ризик того, що вони можуть стати занадто залежними від цих систем, що призведе до помилкового відчуття безпеки. Це може призвести до того, що вони знехтують іншими важливими заходами безпеки, такими як навчання співробітників і планування реагування на інциденти.

Системи штучного інтелекту настільки хороші, наскільки хороші дані, на яких вони навчаються; якщо дані упереджені, система штучного інтелекту також може бути упередженою. Це може призвести до несправедливого ставлення до певних осіб або груп і може посилити існуючу соціальну нерівність.

Впровадження систем безпеки на основі штучного інтелекту може бути дорогим, особливо для малого та середнього бізнесу, який може не мати ресурсів для інвестування в таку технологію.

Системи штучного інтелекту, які використовуються в цілях безпеки, можуть збирати та обробляти великі обсяги даних, що може викликати занепокоєння щодо конфіденційності. Ці дані можуть містити особисту інформацію і, якщо їх не захистити належним чином, можуть призвести до витоку даних, крадіжки особистих даних або фінансового шахрайства.

Хоча штучний інтелект і машинне навчання можуть покращити можливості кібербезпеки, існують також занепокоєння щодо їх потенційного зловмисного використання та етичних наслідків:

- моделі штучного інтелекту здатні генерувати відповіді, подібні до людських, які можуть бути використані для поширення дезінформації.

Зловмисники можуть використовувати цю модель для генерування неправдивої інформації, обману окремих осіб та/або маніпулювання громадською думкою;

- моделі штучного інтелекту можуть використовуватися зловмисниками для покращення своїх кампаній з фішингу та соціальної інженерії. Імітуючи людську розмову, зловмисники можуть намагатися обманом змусити користувачів розкрити конфіденційну інформацію, таку як паролі або фінансові дані;

- моделі штучного інтелекту можуть відображати упередження, присутні в даних, на яких вони були навчені. Якщо навчальні дані містять упереджену або несправедливу інформацію, існує ризик того, що модель може ненавмисно генерувати упереджені або дискримінаційні відповіді. Наприклад, у 2018 році повідомлялося, що Amazon розробила інструмент рекрутингу на основі штучного інтелекту для автоматизації процесу найму. Однак алгоритм показав упереджене ставлення до кандидаток. Система була навчена на історичних даних резюме, які переважно складалися з претендентів чоловічої статі. В результаті алгоритм навчився віддавати перевагу кандидатам-чоловікам і знизив рейтинг резюме, що містять терміни, пов'язані з жінками.

Під час взаємодії з моделями штучного інтелекту користувачі можуть надавати особисту або конфіденційну інформацію. Дуже важливо переконатися, що вжито відповідних заходів для захисту конфіденційності користувачів і безпечного поводження з будь-якими даними, якими діляться під час розмов.

Таким чином, неправильне використання штучного інтелекту може призвести до значних ризиків. Ці ризики включають потенційні порушення прав людини, такі як дискримінація та порушення конфіденційності, а також можливість зловмисного використання штучного інтелекту для кібератак, фінансового шахрайства, створення та поширення дезінформації. Крім того, серед ризиків також є надмірна залежність від штучного інтелекту, відсутність прозорості та людського нагляду, а також потенційне переміщення робочих місць. Дуже важливо мати правила, етичні принципи та заходи безпеки, щоб запобігти неправомірному використанню штучного інтелекту та зменшити ризики, які воно може спричинити.

Висновки до розділу 2. У даному розділі проведено аналіз сучасних тенденцій використання штучного інтелекту в кібернетичній безпеці, що виявило широкий спектр можливостей та перспектив розвитку цього напрямку. Зазначено, що штучний інтелект вже активно застосовується для виявлення кібератак за допомогою методів машинного навчання.

Аналіз сучасних тенденцій підтвердив, що застосування машинного навчання є важливим інструментом для виявлення та реагування на кіберзагрози. Моделі машинного навчання, зокрема, дозволяють виявляти аномалії та підозрілу активність на основі аналізу великих обсягів даних. Використання цих методів може сприяти ефективнішому виявленню загроз та скороченню часу реакції на інциденти.

Окрему увагу приділено етичним питанням застосування штучного інтелекту у кібербезпеці. Виявлено, що разом із великим потенціалом штучного інтелекту для підвищення рівня безпеки, існують серйозні етичні виклики, такі як конфіденційність даних, автономність систем та можливість виникнення системних помилок.

Запропонований алгоритм створення моделі виявлення загроз на основі штучного інтелекту з використанням стандартних бібліотек Python, який буде доцільно використовувати для відпрацювання малими і середніми організаціями при створенні власних систем забезпечення інформаційної безпеки.

У цілому, розділ надає обґрунтовану та систематизовану інформацію щодо методологічних підходів до використання штучного інтелекту в управлінні кібернетичною безпекою. Застосування штучного інтелекту в цьому контексті має потенціал значно поліпшити ефективність систем управління кібербезпекою та сприяти розвитку більш безпечного цифрового середовища.

РОЗДІЛ 3

ПРАКТИЧНЕ ЗАСТОСУВАННЯ ШТУЧНОГО ІНТЕЛЕКТУ В УПРАВЛІННІ КІБЕРНЕТИЧНОЮ БЕЗПЕКОЮ

3.1. Створення моделі виявлення кіберзагроз з використанням штучного інтелекту

Створення повноцінної моделі для виявлення кіберзагроз потребує великої кількості даних та реального збору інформації про кіберзагрози.

Як були визначено у розділі 2.2, для створення власної моделі виявлення загроз з використанням штучного інтелекту для організацій і підприємств малого і середнього бізнесу потрібні знання основ інформаційної та кібербезпеки у поєднанні з основами програмування. За методикою, розробленою у п.2.2 визначена послідовність, яка включає:

- збір та розуміння даних;
- очистка та перетворення даних;
- визначення метрик та критеріїв успішності;
- вибір моделі та архітектури;
- тренування моделі;
- оптимізація та налаштування;
- валідація та оцінка;
- впровадження та моніторинг і постійне навчання і вдосконалення.

Одною з важливих переваг розроблення власної моделі виявлення загроз – це можливість постійного навчання та вдосконалення.

Нижче наведений простий приклад створення моделі на Python, який використовує штучні дані для бінарної класифікації (0 - нормальна активність, 1 - підозріла активність) за допомогою бібліотек TensorFlow та Keras (рис. 3.1) для виявлення загроз на основі контролю текстових повідомлень в соціальних мережах. Такий підхід дозволяє застосувати розвідочний аналіз та виявити

майбутні потенційні загрози.

```
import numpy as np
import tensorflow as tf
from tensorflow import keras
from tensorflow.keras import layers
from sklearn.model_selection import train_test_split
from sklearn.preprocessing import StandardScaler
from sklearn.metrics import confusion_matrix, classification_report

# Згенеруємо штучні дані для прикладу
data_size = 1000
feature_size = 10

X = np.random.rand(data_size, feature_size)
y = np.random.randint(2, size=(data_size,))

# Розділяємо дані на тренувальний та тестовий набори
X_train, X_test, y_train, y_test = train_test_split(X, y, test_size=0.2)

# Стандартизуємо числові дані
scaler = StandardScaler()
X_train = scaler.fit_transform(X_train)
X_test = scaler.transform(X_test)

# Створення моделі
model = keras.Sequential([
    layers.Dense(32, activation='relu', input_shape=(feature_size,)),
    layers.Dense(1, activation='sigmoid')
])

model.compile(optimizer='adam', loss='binary_crossentropy', metrics=['a

# Тренування моделі
model.fit(X_train, y_train, epochs=10, batch_size=32, validation_split=

# Оцінка моделі на тестових даних
y_pred = (model.predict(X_test) > 0.5).astype("int32")

# Оцінка результатів
print("Confusion Matrix:")
print(confusion_matrix(y_test, y_pred))

print("\nClassification Report:")
print(classification_report(y_test, y_pred))
```

Рис.3.1. Лістинг програми виявлення загроз

Цей приклад використовує один прихований шар у створеній моделі та використовує функцію активації ReLU для прихованого шару та сигмоїду для виходу. Кількість шарів моделі, функції активації у подальшому можна змінювати, пристосовуючи модель до потреб завдання виявлення загроз. Модель компілюється за допомогою оптимізатора Adam і функції втрати binary_crossentropy.

Основними джерелами збору даних для моніторингу загроз можуть бути:

- соціальні мережі (Twitter, Facebook, LinkedIn, Reddit, Instagram та інші), які містять велику кількість інформації про злочинну діяльність, наміри злочинців, їх плани, способи атак і т.д.;
- спеціалізовані форуми і блоги, де часто обговорюються питання кібербезпеки та злочинної діяльності;
- веб-сайти, пов'язані з кібербезпекою, такі як Stack Exchange, Cybersecurity Forum, Dark Reading і т.д., які можуть містити відгуки користувачів, коментарі, поради та новини про кібербезпеку та критичну інфраструктуру;
- портали новин, де можуть з'являтися новини про кібератаки на ОКІ та інші подібні події;
- відкриті джерела, бази даних з публічної інформації про кібератаки на об'єкти КІП;
- системи моніторингу, що відслідковують кібератаки на об'єкти критичної інфраструктури та інші подібні події;
- документи та звіти, що стосуються кібербезпеки, наприклад, звіти про витоки даних можуть містити важливі дані для створення моделі;
- інші джерела даних можуть включати датасети, які вже були створені та опубліковані для використання в машинному навчанні.

Для автоматизованого збору даних з веб-сайтів соціальних мереж використовують інструменти веб-скрапінгу. Вони можуть бути корисним для автоматизації збору великого об'єму даних з відгуками користувачів, розміщених в Інтернеті. Основний процес веб-скрапінгу включає наступні кроки:

- визначення джерела даних, з якого потрібно отримати інформацію;
- встановлення з'єднання з веб-сайтом та отримання відповіді зі сторони сервера;
- аналіз відповіді та виділення необхідних даних, використовуючи інструменти веб-скрапінгу;

- збереження даних у файлі або базі даних.

Для цього потрібно розробити окремий блок для збору даних. Збір даних з соціальних мереж може бути складним завданням через обмеження доступу та політику конфіденційності. Більшість соціальних мереж надають API для збору обмеженого обсягу даних. Важливо дотримуватися правил використання API та збирати дані відповідно до політики конфіденційності.

Нижче наведений простий приклад використання Python для збору даних з Twitter API за допомогою бібліотеки Tweepy. Цей приклад демонструє, як отримати твіти з вказаного користувача (рис. 3.2).

```
import tweepy

# Замініть наступні ключі на власні, які ви отримаєте після створення додатку
consumer_key = 'your_consumer_key'
consumer_secret = 'your_consumer_secret'
access_token = 'your_access_token'
access_token_secret = 'your_access_token_secret'

# Аутентифікація за допомогою ключів
auth = tweepy.OAuthHandler(consumer_key, consumer_secret)
auth.set_access_token(access_token, access_token_secret)
api = tweepy.API(auth)

# Збір твітів з вказаного користувача
def collect_tweets(username, count=10):
    tweets = api.user_timeline(screen_name=username, count=count, tweet_mode='extended')
    data = []
    for tweet in tweets:
        data.append({
            'created_at': tweet.created_at,
            'text': tweet.full_text,
            'retweet_count': tweet.retweet_count,
            'favorite_count': tweet.favorite_count,
            'user_mentions': [mention['screen_name'] for mention in tweet.user_mentions],
            'hashtags': [hashtag['text'] for hashtag in tweet.entities.get('hashtags', [])]
        })
    return data

# Виклик функції для збору твітів
username = 'example_username'
tweet_data = collect_tweets(username, count=5)

# Виведення результатів
for tweet in tweet_data:
    print(f"Created At: {tweet['created_at']}")
    print(f"Text: {tweet['text']}")
    print(f"Retweet Count: {tweet['retweet_count']}")
    print(f"Favorite Count: {tweet['favorite_count']}")
    print(f>User Mentions: {tweet['user_mentions']}")
    print(f>Hashtags: {tweet['hashtags']}")
    print("\n" + "-"*30 + "\n")
```

Рис.3.2. Лістинг програми для збору даних за допомогою бібліотеки Tweepy

Для збору даних з Facebook на Python використовують бібліотеку PyFacebook, яка надає доступ до Facebook Graph API. Для використання Facebook Graph API потрібно мати обліковий запис розробника Facebook та створити свій додаток, щоб отримати ключ доступу.

Однак слід зауважити, що ця конкретна модель, побудована на штучно згенерованих даних, не буде досить для ефективного виявлення реальних кіберзагроз. У реальних умовах необхідно мати доступ до великого обсягу реальних даних про кіберзагрози та використовувати більш складні та потужні архітектури мереж для досягнення ефективності. Також важливо регулярно оновлювати та навчати модель на нових даних для підтримки актуальності та вдосконалення її ефективності.

Зібрані дані необхідно перевірити та переконатися, що вони є якісними і репрезентативними. Недостовірні дані можуть призвести до некоректних висновків та помилкових рішень. Для цього необхідно застосувати метод експертних оцінок.

При обробленні даних їх перетворюють у формат, зрозумілий для нейронної мережі: прибирають зайві символи, здійснюють токенізацію, лематизацію та інші методи, що зменшують кількість варіантів тексту і полегшують подальшу обробку. Вислови, які можуть вказувати на кіберзагрози, можуть містити такі елементи:

- використання загрозливих слів, наприклад: "знищити", "зламати", "вбити", "завдати шкоди", "взламати", "захопити", "знищити систему", "нанести удар", "зробити безлад".
- використання емоційно заряджених слів або фраз, таких як: "мститися", "ненавидіти", "злість", "розлючення", "обурення", "ворожість", "напруження", "зневага", "відчай", "тривога", "страх", "паніка".
- використання технічних термінів, пов'язаних з кібербезпекою та кібератаками, таких як: "ботнет", "DDoS-атака", "вірус", "шпигунський програмне забезпечення", "троян", "злочинна група", "витівка".
- використання скорочень або кодових слів, що можуть вказувати на

кібератаку, наприклад: "APT" (Advanced Persistent Threat), "Zer0day" (використання вразливості, про яку ніхто не знає), "Bot" (зловмисне програмне забезпечення, що контролюється здалеку) та ін.

Зазвичай текстові дані перетворюють на числові вектори за допомогою алгоритмів векторизації. Для автоматичного перетворення використовують готову модель Word2Vec бібліотеки Gensim, яку необхідно інсталиувати у Python. Приклад векторизації тексту Word2Vec із застосуванням моделі на рис. 3.3.

```
[-0.00709772 -0.00197749 -0.00095769 -0.00504181 -0.00385755 -0.00306635
 0.00358217 0.00183558 -0.00706854 0.00070746 0.00811502 0.00062238
 -0.00503692 0.00324048 -0.00312645 -0.00412871 0.00365291 0.00715654
 -0.00091911 0.00236634 -0.00703393 0.00876777 0.00568507 -0.00531039
 -0.00543225 -0.00756868 -0.00026711 -0.00868249 -0.00246009 0.00621088
 0.00669893 -0.00202668 0.00367284 -0.00387379 -0.00399847 -0.00077337
 0.00558054 -0.00756769 0.00888656 0.00313928 0.00404492 -0.0041762
 -0.00512051 -0.00332049 -0.0017409 0.00495611 -0.0077603 0.00765923
 0.00124323 0.0009654 -0.00873457 0.00045491 -0.00571809 -0.00074356
 -0.00114072 0.00214731 0.00274531 -0.00167239 -0.00550216 -0.00075626
 -0.00074471 -0.00230423 0.00747439 0.00816348 0.00012605 -0.0002532
 -0.00773769 -0.00824024 -0.00552806 -0.00543652 -0.00102763 -0.00490736
 -0.00779188 0.00810343 -0.00853763 -0.00040155 0.00222529 0.00032456]
```

Рис. 3.2. Результат векторизації тексту

Навчити створену модель штучного інтелекту потрібно на зібраних та розмічених даних, використовуючи вибраний алгоритм та перевірити якість моделі на тестових даних, які не використовувалися при навчанні.

Маючи доступ до коду моделі, що є значною перевагою власноруч розроблених моделей штучного інтелекту, варіюючи кількістю шарів та нейронів у кожному шарі, різними функціями активації, кількістю епох, можна дивитись значної якості моделі.

Під час експерименту в процесі навчання було виявлено, що збільшення кількості шарів і нейронів приводить до перенавчання моделі та невірною результату. Збільшення кількості епох приводить до збільшення втрат (рис. 3.3), що не відповідає меті моделі і знижує точність цільового значення.

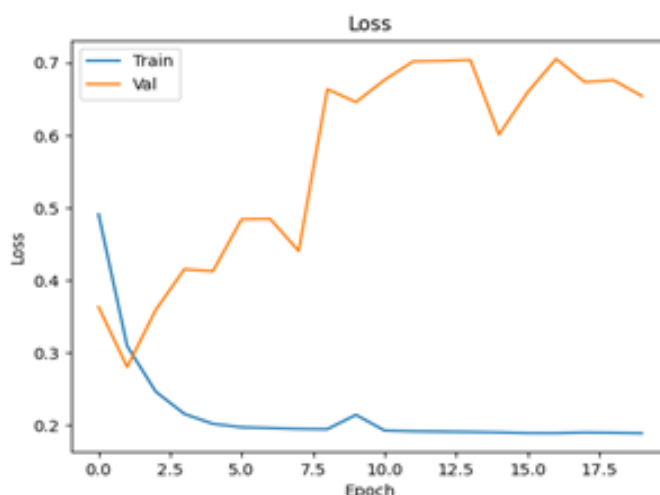


Рис.3.3. Величина втрат за 17 епох

Тому оптимальна кількість епох навчання складає не більше 10, як видно на графіках точності моделі та втрат протягом епох (рис. 3.4).

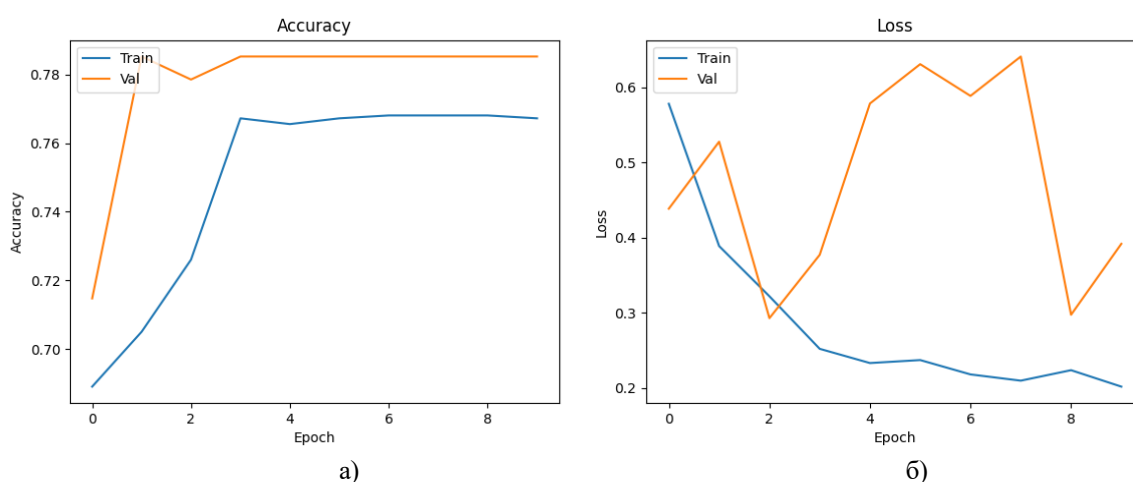


Рис. 3.4. Підвищення точності моделі (а) протягом навчання та зменшення втрат (б)

Після завершення 10 епох виведена фінальна точність (Test Accuracy) моделі на тестовому наборі даних, яка становить 0.7485 (74.85%), оптимальне значення втрат у цьому випадку складає 0.3522. Вивід результату остаточного тестування моделі після тренування зі спливаючим підсумком показує, що модель успішно навчилася на тренувальних даних та має певний рівень універсальності та здатності до визначення відношення до кібератак на об'єкти критичної інформаційної інфраструктури при оцінюванні величини емоційності

повідомлень на невідомих даних.

3.2. Рекомендації щодо вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою

Вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою може значно підвищити ефективність та реакцію на загрози в цифровому середовищі.

Основні рекомендації згруповані за напрямками в таблиці 3.1

Таблиця 3.1.

Рекомендації щодо удосконалення застосування штучного інтелекту в управлінні кібербезпекою

Напрями	Зміст та результат
Розвиток систем аналізу потоків даних	Розроблення систем, які можуть аналізувати великі потоки даних в реальному часі. Використання технологій аналізу потоків даних дозволяє виявляти загрози миттєво, що є критичним для кібербезпеки.
Використання машинного навчання та глибокого навчання	Вдосконалення моделі машинного навчання та глибокого навчання для виявлення нових типів загроз та аномальної активності. Використання адаптивних моделей дозволяє боротися з еволюцією кіберзагроз.
Створення моделей поведінкового аналізу	Розроблення моделі аналізу поведінки користувачів та систем, щоб вчасно виявляти аномальні або підозрілі зміни. Підходи, які враховують звичайну поведінку, можуть виявляти відхилення в реальному часі.
Автоматизована реакція на загрози	Розгляд можливості автоматизації реакції на виявлені загрози. Використання ШІ для автоматизованого виправлення та відновлення може значно скоротити час реакції.

Продовження таблиці 3.1

Напрями	Зміст та результат
Комбінація різноманітних джерел даних	Інтеграція та аналіз інформації з різних джерел (логи, мережевий трафік, поведінка користувачів і т.д.). Комбінація даних може дати повніший образ загроз та подій.
Впровадження технологій кібераналізу та threat intelligence	Використання технологій кібераналізу та врахування Threat Intelligence для збагачення аналізу та виявлення загроз
Навчання на мікрорівнях та оновлення в реальному часі	Забезпечення можливості моделям ШІ навчатися на мікрорівнях і оновлюватися в реальному часі для адаптації до нових видів атак.
Забезпечення етичного використання та захисту конфіденційності	Розвиток та використання ШІ з дотриманням етичних стандартів, а також з урахуванням питань конфіденційності та захисту приватності.
Залучення експертів з кібербезпеки	Співпраця з експертами з кібербезпеки для постійного вдосконалення стратегій та використання новітніх технологій.

Ці рекомендації є загальними та можуть варіюватися залежно від конкретного контексту та потреб організації.

У порівнянні з традиційним програмним забезпеченням, штучний інтелект створює низку різних ризиків. Системи штучного інтелекту навчаються на даних, які можуть змінюватися з часом, іноді суттєво та несподівано, впливаючи на системи таким чином, що їх може бути важко зрозуміти. Ці системи також мають «соціально-технічну» природу, тобто на них впливає соціальна динаміка та поведінка людини. Ризики, пов'язані зі штучним інтелектом, можуть виникати через складну взаємодію цих технічних і соціальних факторів, впливаючи на життя людей у різних ситуаціях, починаючи від їхнього досвіду роботи з онлайн-чат-ботами і закінчуючи результатами заявок на роботу та кредити.

Тому, 30 березня NIST запустив Ресурсний центр з питань штучного інтелекту, який сприятиме впровадженню та міжнародному узгодженню з AI

RMF. У співпраці з приватним і державним секторами NIST розробив структуру для кращого управління ризиками для окремих осіб, організацій і суспільства, пов'язаними зі штучним інтелектом (ШІ). Структура управління ризиками NIST AI (AI RMF) (рис. 3.5) призначена для добровільного використання та покращення здатності враховувати міркування надійності при проектуванні, розробці, використанні та оцінці продуктів, послуг та систем штучного інтелекту.

Оприлюднена 26 січня 2023 року, Фреймворк був розроблений на основі консенсусу, відкритим, прозорим і спільним процесом, який включав запит на інформацію, кілька чернеток версій для публічних коментарів, численні семінари та інші можливості для надання внеску. Він призначений для розвитку, узгодження та підтримки зусиль інших з управління ризиками штучного інтелекту.



Рис.3.5. Структура AI RMF

NIST також опублікував супутній посібник NIST AI RMF Playbook разом із дорожньою картою AI RMF, AI RMF Crosswalk та різними перспективами. Крім того, NIST робить доступним відеопояснення про AI RMF.

AI RMF слідує вказівкам Конгресу для NIST щодо розробки фреймворку та був створений у тісній співпраці з приватним і державним секторами. Він

призначений для адаптації до ландшафту штучного інтелекту в міру того, як технології продовжують розвиватися, і для використання організаціями різного ступеня та можливостей, щоб суспільство могло отримати вигоду від технологій штучного інтелекту, а також бути захищеним від його потенційної шкоди.

ВИСНОВОК до 3 розділу. У розділі досліджено ключові аспекти створення моделей штучного інтелекту для управління та забезпечення кібербезпеки. Вивчено базові етапи розробки та використання моделей, такі як збір та обробка даних, вибір алгоритмів, навчання та тестування.

Розглянуті практичні аспекти включають в себе використання машинного навчання для виявлення аномалій та ідентифікації потенційно підозрілої активності, а також автоматизовану обробку даних для виявлення загроз та реакцію на них. Проаналізовано важливість регулярного навчання моделей на нових даних та їхню здатність адаптуватися до змінюючогося кіберпростору.

Рекомендації щодо удосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою включають в себе впровадження систем аналізу потоків даних для реального виявлення загроз, використання технологій обробки мови та аналізу тексту для розширеного аналізу інформації, а також постійне оновлення моделей та врахування Threat Intelligence для забезпечення найвищого рівня безпеки з урахуванням нових напрацювань NIST.

Зазначено, що впровадження штучного інтелекту в кібербезпеку вимагає глибокого розуміння та врахування етичних та конфіденційних аспектів, а також постійного спілкування та співпраці з експертами з кібербезпеки. Впровадження передових технологій ШІ в управління кібернетичною безпекою має потенціал значно поліпшити виявлення та запобігання кіберзагрозам, забезпечуючи високий рівень безпеки в цифровому середовищі.

На основі запропонованої методики у другому розділі за вказаним алгоритмом створений зразок моделі раннього виявлення загроз за допомогою штучного інтелекту шляхом моніторингу соціальних мереж. Створена на основі

стандартних бібліотек Python, за результатами експерименту модель показала досить добрі результати.

Використання моделі організаціями і підприємствами середнього бізнесу дозволить оперативно реагувати на нові види загроз, що змінюються з часом шляхом відкритого доступу до коду та постійного навчання моделі.

ВИСНОВКИ

У відповідності із визначеними завданнями у роботі розглянуті теоретичні засади кібернетичної безпеки та проаналізовані методологічні підходи до застосування штучного інтелекту в управлінні кібернетичною безпекою. Проведено дослідження і розроблені рекомендації щодо вдосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою.

При розгляді теоретичних засад та аналізі підходів виявлено, що сучасні кіберзагрози, від клонованих ідентичностей до діпфейкових кампаній, стає все важче виявити та зупинити. Разом з тим, проаналізовані види кіберзагроз дозволили визначити основні напрямки зосередження зусиль щодо удосконалення кібербезпеки.

За результатами аналізу сучасних тенденцій використання штучного інтелекту в кібернетичній безпеці встановлено, що майбутнє кібербезпеки буде забезпечене, якщо буде інтегроване із AI. Зазначено, що впровадження штучного інтелекту в кібербезпеку вимагає глибокого розуміння та врахування етичних та конфіденційних аспектів, а також постійного спілкування та співпраці з експертами з кібербезпеки. Впровадження передових технологій ШІ в управління кібернетичною безпекою має потенціал значно поліпшити виявлення та запобігання кіберзагрозам, забезпечуючи високий рівень безпеки в цифровому середовищі.

Моделі машинного навчання, зокрема, дозволяють виявляти аномалії та підозрілу активність на основі аналізу великих обсягів даних. Використання цих методів може сприяти ефективнішому виявленню загроз та скороченню часу реакції на інциденти.

Окрему увагу приділено етичним питанням застосування штучного інтелекту у кібербезпеці. Виявлено, що разом із великим потенціалом штучного інтелекту для підвищення рівня безпеки, існують серйозні етичні виклики, такі як конфіденційність даних, автономність систем та можливість виникнення системних помилок.

Запропонований новий алгоритм створення моделі виявлення загроз на

основі штучного інтелекту з використанням стандартних бібліотек Python, який складає основу нової методики, буде доцільно використовувати для відпрацювання організаціями малого і середнього бізнесу при створенні власних систем забезпечення інформаційної безпеки.

Відпрацьовані рекомендації щодо удосконалення застосування штучного інтелекту в управлінні кібернетичною безпекою включають в себе впровадження систем аналізу потоків даних для реального виявлення загроз, використання технологій обробки мови та аналізу тексту для розширеного аналізу інформації, а також постійне оновлення моделей та врахування Threat Intelligence для забезпечення найвищого рівня безпеки з урахуванням нових напрацювань NIST.

На основі запропонованої методики у другому розділі за вказаним алгоритмом створений зразок моделі раннього виявлення загроз за допомогою штучного інтелекту шляхом моніторингу соціальних мереж. Створена на основі стандартних бібліотек Python, за результатами експерименту модель показала досить добрі результати. Використання моделі організаціями і підприємствами середнього бізнесу дозволить оперативно реагувати на нові види загроз, що змінюються з часом шляхом відкритого доступу до коду та постійного навчання моделі.

Визначено, що для ефективної безпеки управлінцям потрібні навички, необхідні для боротьби з цими загрозами, які виходять далеко за рамки простого розуміння того, як впровадити інструменти або налаштувати шифрування. Ці загрози вимагають різноманітних знань про широкий спектр технологій, конфігурацій і середовищ. Щоб отримати ці навички, організації повинні найняти експертів високого рівня або виділити ресурси на навчання власних фахівців управління кібербезпекою.

ПЕРЕЛІК ПОСИЛАНЬ

1. Про основні засади забезпечення кібербезпеки України: Закон України. (Відомості Верховної Ради (ВВР), 2017, № 45, ст.403). URL: <https://zakon.rada.gov.ua/go/2163-19> (дата звернення 12.11.2023 р.)
2. Bhardwaj A. et al. Secure framework against cyber attacks on cyber-physical robotic systems. J. Electron. Imaging. – 2022. – Volume 31. – №. 6. – p. 061802. URL: <https://doi.org/10.1016/j.inffus.2023.101804>
3. Wasyihun Sema Admass, Yirga Yayeh Munaye, Abebe Abeshu Diro, Cyber security: State of the art, challenges and future directions, Cyber Security and Applications, Volume 2, 2024, 100031, URL: <https://doi.org/10.1016/j.csa.2023.100031>.
4. Gunduz M. Z., Das R. Cyber-security on smart grid: Threats and potential solutions //Computer networks. – 2020. – Vol. 169. – p. 107094. URL: <https://doi.org/10.1016/j.comnet.2019.107094>
5. F. Ullah, H. Naeem, S. Jabbar, S. Khalid, M.A. Latif, F. Al-Turjman, L. Mostarda. Cyber security threats detection in internet of things using deep learning approach. IEEE Access, 7 (2019), pp. 124379-124389, URL: <https://doi.org/10.1109/ACCESS.2019.2937347>
6. B. Guembe, A. Azeta, S. Misra, V.C. Osamor, L. Fernandez-Sanz, V. Pospelova. The emerging threat of ai-driven cyber attacks: a review. Applied Artificial Intelligence, 36, Taylor & Francis (2022), URL: <https://doi.org/10.1080/08839514.2022.2037254>
7. Кавак Г. та ін. Моделювання кібербезпеки: сучасний стан та майбутні напрямки //Журнал кібербезпеки. – 2021. – Vol. 7 – No 1. – p. tyab005. URL: <https://doi.org/10.1093/cybsec/tyab005>
8. H.J. Hejase, H. Kazan, I. Moukadem. Advanced persistent threats (APT): an awareness review. J. Econ. Econ. Educ. Res., 21 (6) (2020), pp. 1-8, URL: <https://doi.org/10.13140/RG.2.2.31300.65927>

9. J. Kaur, K.R. Ramkumar. The recent trends in cyber security: a review. J. King Saud Univ.- Comput. Inform. Sci., 34 (8) (2022), pp. 5766-5781, URL: <https://doi.org/10.1016/j.jksuci.2021.01.018>
10. Z. Zhang, H. Ning, F. Shi, F. Farha, Y. Xu, J. Xu, F. Zhang, K.K.R. Choo. Artificial intelligence in cyber security: research advances, challenges, and opportunities. Artif. Intell. Rev., 55 (2) (2022), pp. 1029-1053, URL: <https://doi.org/10.1007/s10462-021-09976-0>
11. K. Kimani, V. Oduol, K. Langat. Cyber security challenges for IoT-based smart grid networks. Int. J. Crit. Infrastruct. Prot., 25 (2019), pp. 36-49, URL: <https://doi.org/10.1016/j.ijcip.2019.01.001>
12. M. Humayun, M. Niazi, N. Jhanjhi, M. Alshayeb, S. Mahmood. Cyber security threats and vulnerabilities: a systematic mapping study. Arab. J. Sci. Eng., 45 (4) (2020), pp. 3171-3189, URL: <https://doi.org/10.1007/s13369-019-04319-2>
13. A.M. Tonge. Cyber security: challenges for society- literature review. IOSR J. Comput. Eng., 12 (2) (2013), pp. 67-75, URL: <https://doi.org/10.9790/0661-1226775>
14. A. Arabo. Cyber security challenges within the connected home ecosystem futures. Procedia Comput. Sci., 61 (0) (2015), pp. 227-232, URL: <https://doi.org/10.1016/j.procs.2015.09.201>
15. G.D. Rodosek, M. Golling . Cyber security: challenges and application areas. Lect. Note. Logist. (2013), pp. 179-197, URL: https://doi.org/10.1007/978-3-642-32021-7_11
16. Y.-L. Cheng, C.-Y. Lee, Y.-L. Huang, C.A. Buckner, R.M. Lafrenie, J.A. Dénommée, J.M. Caswell, D.A. Want, G.G. Gan, Y.C. Leong, P.C. Bee, E. Chin, A.K.H. Teh, S. Picco, L. Villegas, F. Tonelli, M. Merlo, J. Rigau, D. Diaz, ..., R.H.J. Mathijssen. Smart health and cybersecurity in the era of artificial intelligence. Intech, 11 (2016), p. 13. URL: <https://www.intechopen.com/books/advanced-biometric-technologies/liveness-detection-in-biometrics> (дата звернення 24.11.2023 р.)
17. A. Yazdinejad, M. Kazemi, R.M. Parizi, A. Dehghantanha, H. Karimipour. An ensemble deep learning model for cyber threat hunting in industrial internet of

things. Digit. Commun. Netw., 9 (1) (2022), pp. 101-110, URL: <https://doi.org/10.1016/j.dcan.2022.09.008>

18. K. Thakur, M. Qiu, K. Gai, M.L. Ali. An investigation on cyber security threats and security models. Proceedings - 2nd IEEE International Conference on Cyber Security and Cloud Computing, CSCloud 2015 - IEEE International Symposium of Smart Cloud, IEEE SSC 2015 (2016), pp. 307-311, URL: <https://doi.org/10.1109/CSCloud.2015.71>

19. N.N. Abbas, T. Ahmed, S.H.U. Shah, M. Omar, H.W. Park. Investigating the applications of artificial intelligence in cyber security. Scientometrics, 121 (2) (2019), pp. 1189-1211, URL: <https://doi.org/10.1007/s11192-019-03222-9>

20. H. HaddadPajouh, A. Dehghantanha, R. Khayami, K.K.R. Choo. A deep Recurrent Neural Network based approach for Internet of Things malware threat hunting. Futu. Gener. Comput. Syst., 85 (2018), pp. 88-96, URL: <https://doi.org/10.1016/j.future.2018.03.007>

21. Ramanpreet Kaur, Dušan Gabrijelčič, Tomaž Klobučar. Artificial intelligence for cybersecurity: Literature review and future research directions. Information Fusion. Volume 97. 2023. 101804. <https://doi.org/10.1016/j.inffus.2023.101804>.

22. Chithaluru P. et al. Computational Intelligence Inspired Adaptive Opportunistic Clustering Approach for Industrial IoT Networks . In IEEE Internet of Things Journal. № 9— 2023. . pp. 7884 - 7892 URL: <https://doi.org/10.1109/JIOT.2022.3231605>

23. Р. Чіталуру, А.Т. Фаді, М. Кумар, Т. Стефан. Обчислювальний інтелект надихнув адаптивний опортуністичний підхід кластеризації для промислових мереж IoT. IEEE Internet Things J (2023). URL: <https://doi.org/10.1109/JIOT.2022.3231605>

24. В.Г. Промислов, К.В. Семенов, А.С. Шумов. Кластерний метод класифікації кібербезпеки активів. IFAC-PapersOnLine, 52 (13) (2019), с. 928-933. URL: <https://doi.org/10.1016/j.ifacol.2019.11.313>

25. Міллар К. та ін. Класифікація операційних систем: мінімалістичний підхід. In 2020 Міжнародна конференція з машинного навчання та кібернетики (ICMLC). – IEEE,. – С. 143-150. URL: <https://doi.org/10.1109/ICMLC51923.2020.9469571>

26. А. Аксой, М.Г. Гюнес. Автоматична ідентифікація пристроїв IoT за допомогою мережевого трафіку. Міжнародна конференція IEEE з комунікацій (ICC) (2019), с. 1-7. URL: <https://icc2019.ieee-icc.org/program/best-paper-awards.html>

27. А. Сіванатан, Х.Х. Гарахейлі, Ф. Лой, А. Редфорд, К. Відженаке, А. Вішванатх, В. Сівараман. Класифікація пристроїв IoT в розумних середовищах за характеристиками мережевого трафіку. IEEE Trans. Mobile Comput., 18 (8) (2018), pp. 1745-1759. URL: <https://doi.org/10.1109/TMC.2018.2866249>

28. І. Квітич, Д. Перакович, М. Періша, Н. Гупта. Ансамблевий підхід до машинного навчання для класифікації IoT-пристроїв у розумному домі. Міжнар. Кібернетика, 12 (11) (2021), с. 3179-3202. URL: <https://doi.org/10.1007/s13042-020-01241-0>

29. Zhan Z., Xu M., Xu S. A characterization of cybersecurity posture from network telescope data. Trusted Systems: 6th International Conference, INTRUST 2014, Beijing, China, December 16-17, 2014, Revised Selected Papers 6. – Springer International Publishing, 2015. – С. 105-126. URL: https://doi.org/10.1007/978-3-319-27998-5_7

30. Gourisetti S. N. G. et al. Multi-scenario use case based demonstration of buildings cybersecurity framework webtool. In 2017 IEEE Symposium Series on Computational Intelligence (SSCI). – IEEE, 2017. – С. 1-8. URL: <https://ieeexplore.ieee.org/xpl/conhome/8267146/proceeding>

31. Narasimhan V. L. et al. Using deep learning for assessing cybersecurity economic risks in virtual power plants. Proceeding in 2021 7th International Conference on Electrical Energy Systems (ICEES). – IEEE, 2021. – pp. 530-537. URL: <https://ieeexplore.ieee.org/xpl/conhome/9383648/proceeding>

32. Nguyen H. H., Nicol D. M. Estimating loss due to cyber-attack in the presence of uncertainty. In 2020 IEEE 19th International Conference on Trust, Security and Privacy in Computing and Communications (TrustCom). – IEEE, 2020. – pp. 361-369. URL: <https://doi.org/10.1109/TrustCom50675.2020.00057>

33. Ponsard C., Ramon V., Touzani M. Improving Cyber Security Risk Assessment by Combined Use of AI and Infrastructure Models. In the 14th International iStar Workshop. – 2021. – pp. 63-69.

34. Odegbile O., Chen S., Wang Y. Dependable Policy Enforcement in Traditional Non-SDN Networks. Processing 2019 IEEE 39th International Conference on Distributed Computing Systems (ICDCS). – IEEE, 2019. – pp. 545-554. URL: <https://doi.org/10.1109/ICDCS.2019.00061>

35. Odegbile O. et al. Policy enforcement in traditional non-SDN networks . Journal of Parallel and Distributed Computing. – 2023. – №. 177. – pp. 39-52. URL: <https://dl.acm.org/toc/jpdc/2023/177/C>

36. Nembhard F. D., Carvalho M. M., Eskridge T. C. Towards the application of recommender systems to secure coding //EURASIP Journal on Information Security. – 2019. – Vol. №. 1. – pp. 1-24. URL: <https://doi.org/10.1109/SoutheastCon42311.2019.9020339>

37. Liu S. et al. DeepBalance: Deep-learning and fuzzy oversampling for vulnerability detection. IEEE Transactions on Fuzzy Systems. – 2019. – T. 28. – №. 7. – pp. 1329-1343. URL: <https://doi.org/10.1109/DDCLS.2019.8908835>

38. Jeon S., Kim H. K. AutoVAS: An automated vulnerability analysis system with a deep learning approach. Computers & Security. – 2021. – T. 106. – pp. 102308. URL: <https://doi.org/10.1016/j.cose.2021.102308>

39. Huff P. et al. A recommender system for tracking vulnerabilities. Proceedings of the 16th International Conference on Availability, Reliability and Security. – 2021. – pp. 1-7. URL: <https://doi.org/10.1145/3465481.3470039>

40. Iorga D. et al. Yggdrasil—early detection of cybernetic vulnerabilities from Twitter //2021 23rd International Conference on Control Systems and Computer

Science (CSCS). – IEEE, 2021. – pp. 463-468. URL: <https://doi.org/10.1109/CSCS52396.2021.00082>

41. Saha T. et al. SHARKS: Smart hacking approaches for risk scanning in Internet-of-Things and cyber-physical systems based on machine learning. IEEE Transactions on Emerging Topics in Computing. – 2021. – Vol. 10. – No. 2. – C. 870-885.

42. Cummins C. et al. Compiler fuzzing through deep learning. Proceedings of the 27th ACM SIGSOFT International Symposium on Software Testing and Analysis. – 2018. – pp. 95-105. URL: <https://doi.org/10.1145/3213846.3213848>

43. Xu H. et al. Dsmith: Compiler fuzzing through generative deep learning model with attention. In 2020 International Joint Conference on Neural Networks (IJCNN). – IEEE, 2020. – pp. 1-9. URL: <https://doi.org/10.1186/s12964-019-0473-9>

44. Chen Y. et al. Learning-guided network fuzzing for testing cyber-physical system defences //2019 34th IEEE/ACM International Conference on Automated Software Engineering (ASE). – IEEE, 2019. – C. 962-973. URL: <https://doi.org/10.1109/ASE.2019.00093>

45. She D. et al. Neuzz: Efficient fuzzing with neural program smoothing. In 2019 IEEE Symposium on Security and Privacy (SP). – IEEE, 2019. – C. 803-817. URL: <https://doi.org/10.1109/SP.2019.00052>

46. Zhou S. et al. Autonomous penetration testing based on improved deep q-network. Applied Sciences. – 2021. – T. 11. – No. 19. – C. 8823. URL: <https://doi.org/10.3390/app11198823>

47. Gangupantulu R. et al. Crown jewels analysis using reinforcement learning with attack graphs. In 2021 IEEE Symposium Series on Computational Intelligence (SSCI). – IEEE, 2021. – C. 1-6. URL: <https://doi.org/10.1109/SSCI50451.2021.9659947>

48. Neal C. et al. Reinforcement learning based penetration testing of a microgrid control algorithm . IN 2021 IEEE 11th Annual Computing and Communication Workshop and Conference (CCWC). – IEEE, 2021. – C. 0038-0044. URL: <https://doi.org/10.1109/CCWC51732.2021.9376126>

49. Russo E. R. et al. Summarizing vulnerabilities' descriptions to support experts during vulnerability assessment activities. *Journal of Systems and Software*. – 2019. – T. 156. – C. 84-99. URL: <https://doi.org/10.1016/j.jss.2019.06.001>

50. Aota M. et al. Automation of vulnerability classification from its description using machine learning //2020 IEEE Symposium on Computers and Communications (ISCC). – IEEE, 2020. – C. 1-7. URL: <https://doi.org/10.1109/ISCC50000.2020.9219568>

51. M. Vanamala, X. Yuan and K. Roy, "Topic Modeling And Classification Of Common Vulnerabilities And Exposures Database," 2020 International Conference on Artificial Intelligence, Big Data, Computing and Data Communication Systems (icABCD), Durban, South Africa, 2020, pp. 1-5, doi: URL: <https://doi.org/10.1109/icABCD49160.2020.9183814>.

52. Nadeem A. et al. Alert-driven attack graph generation using s-pdf //IEEE Transactions on Dependable and Secure Computing. – 2021. – T. 19. – №. 2. – C. 731-746. URL: <https://doi.org/10.1109/TDSC.2021.3117348>

53. Binyamini H. et al. A framework for modeling cyber attack techniques from security vulnerability descriptions. *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery & Data Mining*. – 2021. – C. 2574-2583. URL: <https://doi.org/10.1145/3447548.3467159>

54. Falco G. et al. A master attack methodology for an AI-based automated attack planner for smart cities. *IEEE Access*. – 2018. – T. 6. – C. 48360-48373. URL: <https://doi.org/10.1109/ACCESS.2018.2867556>

55. Cam H. Model-Guided Infection Prediction and Active Defense Using Context-Specific Cybersecurity Observations. *MILCOM 2019-2019 IEEE Military Communications Conference (MILCOM)*. – IEEE, 2019. – C. 1-6. URL: <https://doi.org/10.1109/MILCOM47813.2019.9020980>

56. Wollaber A. et al. Proactive cyber situation awareness via high performance computing //2019 IEEE High Performance Extreme Computing Conference (HPEC). – IEEE, 2019. – C. 1-7. URL: <https://doi.org/10.1109/HPEC.2019.8916528>

57. Sancho J. C. et al. New approach for threat classification and security risk estimations based on security event management. *Future Generation Computer Systems*. – 2020. – T. 113. – C. 488-505. URL: <https://doi.org/10.1016/j.future.2020.07.015>

58. M.S. Ansari, V. Bartoš, B. Lee. GRU-based deep learning approach for network intrusion alert prediction. *Future Gener. Comput. Syst.*, Vol. 2 No. 4 (2022). pp. 35-47. URL: <https://doi.org/10.54183/jssr.v2i4.53>

59. L. Wang, R. Jones. Big data analytics in cybersecurity: network data and intrusion prediction. *IEEE 10th Annual Ubiquitous Computing, Electronics & Mobile Communication Conference (UEMCON)* (2019), pp. 0105-0111 URL: <https://doi.org/10.1109/UEMCON47517.2019.8993037>

60. H. Al Najada, I. Mahgoub, I. Mohammed. Cyber intrusion prediction and taxonomy system using deep learning and distributed big data processing. *IEEE symposium series on computational intelligence (SSCI)* (2018), pp. 631-638. URL: <https://doi.org/10.1109/SSCI50451.2018.9659982>

61. M. Rhode, P. Burnap, K. Jones. Early-stage malware prediction using recurrent neural networks. *Comput. Secur.*, 77 (2018), pp. 578-594. URL: <https://doi.org/10.1016/j.cose.2018.05.010>

62. I. Perera, J. Hwang, K. Bayas, B. Dorr, Y. Wilks. Cyberattack prediction through public text analysis and mini-theories. *IEEE International Conference on Big Data (Big Data)* (2018), pp. 3001-3010. URL: <https://doi.org/10.1109/BigData.2018.8622106>

63. E. Marin, M. Almukaynizi, P. Shakarian. Inductive and deductive reasoning to assist in cyber-attack prediction. *10th Annual Computing and Communication Workshop and Conference (CCWC)* (2020), pp. 0262-0268. URL: <https://doi.org/10.1109/CCWC47524.2020.9031154>

64. N. Polatidis, E. Pimenidis, M. Pavlidis, S. Papastergiou, H. Mouratidis. From product recommendation to cyber-attack prediction: generating attack graphs and predicting future attacks. *Evolving Syst.*, 11 (3) (2020), pp. 479-490. URL: <https://doi.org/10.1007/s12530-018-9234-z>

65. A.I. Siam, A. Sedik, W. El-Shafai, A.A. Elazm, N.A. El-Bahnasawy, G.M. El Banby, A.A. Khalaf, F.E. Abd El-Samie. Biosignal classification for human identification based on convolutional neural networks. *Int. J. Commun. Syst.*, 34 (7) (2021), pp. 1-22. URL: <https://doi.org/10.4018/IJFSA.285553>

66. M. Sharma, H. Elmiligi. *Behavioral Biometrics: Past, Present and Future*. Reviewed: 24 January 2022 Published: 06 March 2022. URL: <https://doi.org/10.5772/intechopen.102841>.

67. Таддео, М. Три етичні виклики застосування штучного інтелекту в кібербезпеці. *Minds & Machines* 29, 2019. С. 187–191. <https://doi.org/10.1007/s11023-019-09504-8>

ДОДАТКИ

ДОДАТОК А

Лістинг програми для збору даних з використанням бібліотеки «psutil»

```
import psutil
import time

def collect_system_data():
    # Отримати інформацію про використання ресурсів системи
    cpu_usage_percent = psutil.cpu_percent(interval=1)
    memory_usage = psutil.virtual_memory()
    disk_usage = psutil.disk_usage('/')

    # Отримати інформацію про процеси
    running_processes = psutil.process_iter(['pid', 'name', 'username', 'cpu_percent',
    'memory_percent'])
    processes_info = [{ 'pid': process.info['pid'], 'name': process.info['name'], 'username':
    process.info['username'],
                        'cpu_percent': process.info['cpu_percent'], 'memory_percent':
    process.info['memory_percent']}
                      for process in running_processes]

    return {
        'cpu_usage_percent': cpu_usage_percent,
        'memory_usage': {
            'total': memory_usage.total,
            'available': memory_usage.available,
            'percent': memory_usage.percent
        },
        'disk_usage': {
            'total': disk_usage.total,
            'used': disk_usage.used,
            'free': disk_usage.free,
            'percent': disk_usage.percent
        },
        'processes_info': processes_info
    }

def main():
    while True:
        data = collect_system_data()
        print("System Data:", data)
        time.sleep(5) # Очікування 5 секунд перед наступним збором даних

if __name__ == "__main__":
    main()
```

**Лістинг програми для збору даних з використанням бібліотеки
«pandas» та «sklearn»**

```
import pandas as pd
from sklearn.preprocessing import StandardScaler, LabelEncoder
from sklearn.ensemble import IsolationForest

# Згенеруємо штучні дані
data = {'Feature1': [1.0, 2.0, 3.0, 4.0, 5.0, 100.0],
        'Feature2': ['A', 'B', 'A', 'C', 'B', 'C'],
        'Target': [0, 1, 0, 1, 0, 1]}

df = pd.DataFrame(data)

# Очищення та кодування даних
def process_data(dataframe):
    # Видалення аномалій за допомогою Isolation Forest
    clf = IsolationForest(contamination=0.2, random_state=42)
    dataframe['Anomaly'] = clf.fit_predict(dataframe[['Feature1']])
    dataframe = dataframe[dataframe['Anomaly'] != 1].drop('Anomaly', axis=1)

    # Кодування категоріальних змінних
    le = LabelEncoder()
    dataframe['Feature2'] = le.fit_transform(dataframe['Feature2'])

    return dataframe

# Стандартизація числових змінних
def standardize_numerical(dataframe):
    scaler = StandardScaler()
    dataframe[['Feature1']] = scaler.fit_transform(dataframe[['Feature1']])
    return dataframe

# Застосуємо обробку та очищення даних
df_processed = process_data(df)
df_standardized = standardize_numerical(df_processed)

# Виведемо результат
print("Original Data:")
print(df)

print("\nProcessed and Standardized Data:")
print(df_standardized)
```


ДОДАТОК В

Лістинг моделі з використанням бібліотеки «numpy» і «tensorflow»

```
import numpy as np
import tensorflow as tf
from tensorflow import keras
from tensorflow.keras import layers

# Згенеруємо штучні дані для прикладу
data_size = 1000
feature_size = 10

X_train = np.random.rand(data_size, feature_size)
y_train = np.random.randint(2, size=(data_size,))

# Створення моделі
model = keras.Sequential([
    layers.Dense(32, activation='relu', input_shape=(feature_size,)),
    layers.Dense(1, activation='sigmoid')
])

model.compile(optimizer='adam', loss='binary_crossentropy', metrics=['accuracy'])

# Тренування моделі
model.fit(X_train, y_train, epochs=10, batch_size=32, validation_split=0.2)

# Використання моделі для прогнозування
example_data = np.random.rand(5, feature_size)
predictions = model.predict(example_data)
print("Predictions:", predictions)
```

Цей код створює просту ГНМ для бінарної класифікації (наприклад, виявлення кіберзагроз: 0 - нормальна активність, 1 - підозріла активність). Модель має один прихований шар і використовує функцію активації ReLU для прихованого шару та сигмоїду для виходу, оскільки вирішується задача бінарної класифікації.