

Вінницький національний технічний університет
Факультет менеджменту та інформаційної безпеки
Кафедра менеджменту та безпеки інформаційних систем

МАГІСТЕРСЬКА КВАЛІФІКАЦІЙНА РОБОТА

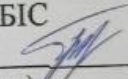
на тему:

«Підвищення захищеності системи DLP запобігання витоку даних на основі
Transformer-based моделі GPT і алгоритму кластеризації з використанням K-
Means для виявлення аномалій і контекстного аналізу даних»

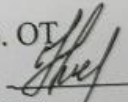
Виконав: ст. 2-го курсу, групи КІТС-23м
спеціальності 125– Кібербезпека та захист
інформації

Освітня програма – Кібербезпека
інформаційних технологій та систем
(шифр і назва напрямку підготовки, спеціальності)

Шевиркін В.Ю. 
(прізвище та ініціали)

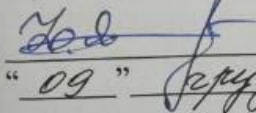
Керівник: к.т.н., доц. каф. МБІС
Грицак А.В. 
(прізвище та ініціали)

« ____ » _____ 2024 р.

Опонент: к.т.н., доцент, каф. ОТ
Колесник І. С. 
(прізвище та ініціали)

« ____ » _____ 2024 р.

Допущено до захисту
Голова секції УБ кафедри МБІС

 Юрій ЯРЕМЧУК
“ 09 ” _____ 2024 р.

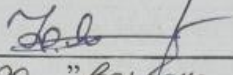
Вінниця ВНТУ – 2024 рік

Вінницький національний технічний університет
Факультет менеджменту та інформаційної безпеки
Кафедра менеджменту та безпеки інформаційних систем

Рівень вищої освіти II-й (магістерський)
Галузь знань 12 – Інформаційні технології
Спеціальність 125 – Кібербезпека та захист інформації
Освітньо-професійна програма - Кібербезпека інформаційних технологій та систем

ЗАТВЕРДЖУЮ

Голова секції УБ, кафедра МБІС

 **Юрій ЯРЕМЧУК**
“20” вересня 2024 р.

ЗАВДАННЯ

на магістерську кваліфікаційну роботу студенту

Шевиркін Валентин Юрійович

(прізвище, ім'я, по-батькові)

1. Тема роботи:

«Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій і контекстного аналізу даних»

Керівник роботи: к.т.н., доц. каф. МБІС Грицак А.В.

(прізвище, ім'я, по-батькові, науковий ступінь, вчене звання)

затверджені наказом вищого навчального закладу від “17” вересня 2024 року № 310

2. Строк подання студентом роботи за тиждень до захисту.

3. Вихідні дані до роботи:

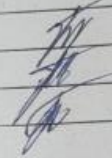
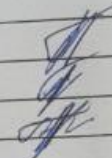
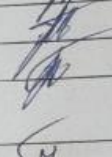
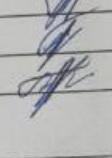
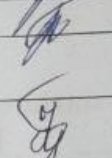
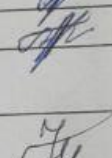
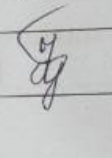
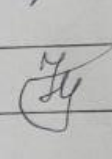
Стандарти, електронні джерела, підручники та наукові статті по темі, які стосуються теми магістерської кваліфікаційної роботи.

4. Зміст текстової частини:

У текстовій частині роботи висвітлено теоретичні аспекти захисту інформації та побудови систем DLP, зокрема роль захисту даних у сучасних організаціях, класифікацію загроз витоку даних, основні принципи реалізації систем DLP та використання сучасних моделей машинного навчання. Описано процес створення захищеної системи на основі трансформерних моделей GPT і алгоритму кластеризації K-Means, включаючи розробку алгоритмів, інтеграцію модулів і тестування програмної реалізації. Представлено висновки щодо ефективності запропонованих підходів і можливості впровадження системи в реальних умовах.

5. Перелік ілюстративного матеріалу (з точним зазначенням обов'язкових креслень)
У дугому розділі магістерської кваліфікаційної роботи наведено 3 рисунків, у третьому розділі – 10 рисунків

6. Консультанти розділів роботи

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		завдання видав	завдання прийняв
Основна частина			
I	Грицак А.В. к.т.н., доц. каф. МБІС		
II	Грицак А.В. к.т.н., доц. каф. МБІС		
III	Грицак А.В. к.т.н., доц. каф. МБІС		
Економічна частина			
IV	Буреннікова Н. В. д.е.н., професор каф. ЕПВМ		

7. Дата видачі завдання 20 вересня 2024 р.

КАЛЕНДАРНИЙ ПЛАН

№	Назва етапів магістерської кваліфікаційної роботи	Строк виконання етапів роботи		Примітка
1.	Визначення напрямку магістерської роботи, формулювання теми	20.09.2024	31.09.2024	
2.	Аналіз предметної області обраної теми	01.10.2024	15.10.2024	
3.	Розробка роботи	16.10.2024	26.10.2024	
4.	Написання магістерської роботи на основі розробленої теми	27.10.2024	15.11.2024	
5.	Передзахист магістерської кваліфікаційної роботи	16.11.2024	24.11.2024	
6.	Виправлення, уточнення, корегування магістерської кваліфікаційної роботи	27.11.2024	04.12.2024	
7.	Захист магістерської кваліфікаційної роботи	10.12.2024	17.12.2024	

Студент


(підпис)

Шевиркін В.Ю.

Керівник роботи


(підпис)

Грицак А.В.

АНОТАЦІЯ

УДК 004.56.5

Шевиркін В.Ю. Магістерська кваліфікаційна робота зі спеціальності 125 – «Кібербезпека та захист інформації», освітня програма «Кібербезпека інформаційних технологій та систем». Вінниця: ВНТУ, 2024. – 117 с.

Укр. мовою. Бібліогр.: 41 назв; рис.: 10; табл. 9.

Ключові слова: кібербезпека, USB-ключі, HOTP

Магістерська робота присвячена розробці вдосконаленої системи запобігання витоку даних (DLP), яка використовує сучасні підходи до аналізу інформації. Запропонована система базується на поєднанні трансформерної моделі GPT для контекстного аналізу текстових даних і алгоритму кластеризації K-Means для виявлення аномалій у поведінкових паттернах користувачів. У роботі проведено детальний аналіз існуючих загроз витоку даних, їхньої класифікації, а також принципів та підходів до реалізації сучасних систем DLP.

У ході дослідження було розроблено програмні модулі, що виконують кластеризацію поведінкових даних, контекстний аналіз тексту та виявлення аномалій. Інтеграція цих модулів у єдину систему забезпечила узгоджену роботу всіх компонентів, що дозволило ефективно аналізувати дані в реальному часі. Проведене тестування підтвердило високу точність системи у виявленні загроз, здатність адаптуватися до нових патернів поведінки та обробляти великі обсяги інформації.

Новизна роботи полягає в інтеграції моделей машинного навчання для комплексного аналізу текстових і поведінкових даних у системі DLP, що дозволило суттєво знизити кількість хибнопозитивних спрацювань та підвищити загальну ефективність системи. Практична цінність роботи полягає у можливості впровадження запропонованої системи в корпоративних середовищах для захисту конфіденційної інформації та протидії сучасним кіберзагрозам. Результати дослідження можуть бути використані для подальшого розвитку систем захисту інформації та адаптації їх до динамічних умов інформаційного середовища

ABSTRACT

UDC 004.56.5

Shevyrykin V.Y. Master's qualification work in specialty 125 - "Cybersecurity and Information Protection", educational program "Cybersecurity of Information Technologies and Systems". Vinnytsia: VNTU, 2024. 117 p.

In Ukrainian. Bibliography: 41 titles; Figures: 10; Table 9.

Keywords: cybersecurity, USB keys, HOTP

The master's thesis is devoted to the development of an advanced data leakage prevention (DLP) system that uses modern approaches to information analysis. The proposed system is based on a combination of the GPT transformer model for contextual analysis of text data and the K-Means clustering algorithm for detecting anomalies in user behavioral patterns. The paper provides a detailed analysis of existing data leakage threats, their classification, as well as principles and approaches to the implementation of modern DLP systems.

In the course of the study, we developed software modules that perform behavioral data clustering, contextual text analysis, and anomaly detection. The integration of these modules into a single system ensured the coordinated operation of all components, which allowed for effective real-time data analysis. The testing confirmed the high accuracy of the system in detecting threats, its ability to adapt to new patterns of behavior and process large amounts of information.

The novelty of the work is the integration of machine learning models for the comprehensive analysis of textual and behavioral data in the DLP system, which significantly reduced the number of false positives and increased the overall efficiency of the system. The practical value of the work lies in the possibility of implementing the proposed system in corporate environments to protect confidential information and counter modern cyber threats. The results of the study can be used for further development of information security systems and their adaptation to the dynamic conditions of the information environment.

ЗМІСТ

ВСТУП	9
1 ТЕОРЕТИЧНІ ЗАСАДИ ЗАХИСТУ ІНФОРМАЦІЇ ТА ЗАБЕЗПЕЧЕННЯ DLP.....	11
1.1 Роль захисту даних у забезпеченні інформаційної безпеки сучасних організацій.....	11
1.2 Огляд і класифікація загроз витоку даних	17
1.3 Основні принципи та підходи до реалізації систем DLP	31
1.4 Аналіз сучасних моделей машинного навчання для використання в DLP-системах.....	36
1.5 Висновки та постановка задачі.....	44
2 СТВОРЕННЯ ЗАХИЩЕНОЇ СИСТЕМИ DLP ЗАПОБІГАННЯ ВИТОКУ ДАНИХ НА ОСНОВІ TRANSFORMER-BASED МОДЕЛІ GPT І АЛГОРИТМУ КЛАСТЕРИЗАЦІЇ З ВИКОРИСТАННЯМ K-MEANS ДЛЯ ВИЯВЛЕННЯ АНОМАЛІЙ І КОНТЕКСТНОГО АНАЛІЗУ ДАНИХ	45
2.1 Особливості створення захищеної системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації	45
2.2 Розробка алгоритму інтеграції моделі gpt.....	52
2.3 Розробка алгоритму кластеризації з використанням k-means для виявлення аномалій і контекстного аналізу даних.	56
2.4 Висновки до розділу.	58
3 ПРОГРАМНА РЕАЛІЗАЦІЯ ЗАХИЩЕНОГО КОНСОЛІДОВАНОГО ІНФОРМАЦІЙНОГО РЕСУРСУ ТА СИСТЕМНИЙ АНАЛІЗ БЕЗПЕКИ ПРОМИСЛОВОЇ ІНФРАСТРУКТУРИ	59
3.1 Обґрунтування вибору мови програмування та середовища розробки ..	59

3.2 Програмна реалізація модулю інтеграції Transformer-based моделі GPT	64
3.3 Програмна реалізація модулю алгоритму кластеризації з використанням K-Means	67
3.4 Програмна реалізація модуля виявлення аномалій і контекстного аналізу даних	71
3.5 Програмна реалізація модуля інтеграції розроблених модулів в системи DLP запобігання витоку даних	74
3.6 Тестування програми реалізації	77
3.7 Висновки до розділу	80
4 ЕКОНОМІЧНА ЧАСТИНА	82
4.1 Оцінювання комерційного потенціалу розробки програмного забезпечення	82
4.2 Прогнозування витрат на виконання наукової роботи та впровадження її результатів.....	86
4.3 Прогнозування комерційних ефектів від реалізації результатів розробки	94
4.4 Розрахунок ефективності вкладених інвестицій та періоду їх окупності.....	97
4.5 Висновки до розділу	99
ВИСНОВКИ.....	100
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	102
ДОДАТКИ.....	106
Додаток А. Технічне завдання.....	Error! Bookmark not defined.
Додаток Б. Лістинг програми.....	112
Додаток В. Ілюстративний матеріал	118
Додаток Г. Протокол перевірки на антиплагіат.....	Error! Bookmark not defined.

ВСТУП

У сучасному інформаційному суспільстві захист даних є однією з найважливіших задач, з якою стикаються організації та підприємства. Зростаючий обсяг цифрової інформації, посилення регуляторних вимог і збільшення кількості кіберзагроз підкреслюють необхідність використання сучасних підходів до запобігання витокам конфіденційної інформації. У цьому контексті розробка систем запобігання витоку даних (DLP), які інтегрують новітні методи машинного навчання, є актуальним напрямом досліджень.

Актуальність.

Сучасні методи захисту інформації часто не враховують складність контексту даних та поведінкових факторів користувачів. Це призводить до того, що традиційні системи DLP мають високий рівень хибнопозитивних спрацювань та не здатні виявляти приховані загрози. Використання трансформерних моделей, таких як GPT, для контекстного аналізу тексту, у поєднанні з алгоритмами кластеризації, зокрема K-Means, дозволяє забезпечити високий рівень точності у виявленні загроз, враховуючи як текстову, так і поведінкову складову. Такий підхід є актуальним для створення сучасних систем захисту, здатних адаптуватися до нових загроз і забезпечувати мінімізацію ризиків витоку даних.

Мета і задачі дослідження.

Метою дослідження є розробка та впровадження вдосконаленої системи DLP, яка базується на використанні трансформерних моделей для контекстного аналізу текстових даних і алгоритмів кластеризації для аналізу поведінкових даних користувачів.

Для досягнення поставленої мети у роботі вирішуються наступні задачі:

- аналіз сучасних загроз витоку даних і існуючих підходів до їх запобігання;
- дослідження можливостей використання трансформерних моделей для контекстного аналізу тексту;

- розробка алгоритму кластеризації для аналізу поведінкових даних;
- інтеграція розроблених методів у єдину систему DLP;
- тестування системи на реальних даних для оцінки її ефективності.

Об'єкт дослідження.

Процеси захисту інформації в системах запобігання витоку даних (DLP).

Предмет дослідження.

Методи і алгоритми контекстного аналізу тексту та кластеризації поведінкових даних для забезпечення ефективності роботи систем DLP.

Новизна роботи.

Новизна роботи полягає у розробці та впровадженні комплексного підходу до запобігання витоку даних, який поєднує використання трансформерних моделей GPT для контекстного аналізу текстових даних із алгоритмами кластеризації K-Means для аналізу поведінкових патернів користувачів. У роботі вперше запропоновано інтеграцію цих методів у єдину систему DLP, що забезпечує високу точність виявлення загроз.

Практична цінність.

Практична цінність роботи полягає у можливості впровадження розробленої системи DLP у корпоративних середовищах для захисту конфіденційної інформації. Використання запропонованого підходу дозволяє значно знизити ризик витоку даних, підвищити ефективність роботи систем захисту та забезпечити адаптацію до нових загроз. Результати роботи можуть бути використані для створення масштабованих рішень, які відповідають сучасним вимогам до інформаційної безпеки.

1 ТЕОРЕТИЧНІ ЗАСАДИ ЗАХИСТУ ІНФОРМАЦІЇ ТА ЗАБЕЗПЕЧЕННЯ DLP

У даному розділі буде розглянуто теоретичні засади захисту інформації та забезпечення систем запобігання витоку даних (DLP). Буде досліджено роль захисту даних у забезпеченні інформаційної безпеки сучасних організацій, проаналізовано основні загрози витоку даних та їхню класифікацію. Окрему увагу буде приділено принципам і підходам до реалізації систем DLP, а також розглянуто сучасні моделі машинного навчання, які використовуються для покращення ефективності таких систем. На основі проведеного аналізу будуть сформульовані висновки та поставлені задачі для подальшої розробки запропонованих рішень.

1.1 Роль захисту даних у забезпеченні інформаційної безпеки сучасних організацій

Захист даних є одним з основних аспектів забезпечення інформаційної безпеки, особливо в умовах сучасного розвитку технологій, коли обсяг, різноманітність та швидкість обробки даних значно зросли. Сучасні організації дедалі більше стикаються з необхідністю не лише захисту своїх власних даних, а й забезпечення конфіденційності, цілісності та доступності даних своїх клієнтів і партнерів. Це пов'язано з кількома ключовими факторами, які визначають важливість захисту даних у сучасних умовах [1].

Сучасні організації активно використовують технології для автоматизації бізнес-процесів, оптимізації роботи з інформацією, покращення обслуговування клієнтів та підвищення продуктивності. Ці процеси призводять до генерування великих обсягів даних, які потрібно зберігати, обробляти та захищати. Критичним аспектом є те, що більшість цих даних обробляються та передаються в цифровій формі, що робить їх вразливими до різних типів кіберзагроз. Таким чином, забезпечення захисту даних стає необхідною умовою для безперебійного функціонування організації та її конкурентоспроможності.

Із впровадженням таких нормативних документів, як Загальний регламент захисту даних ЄС (GDPR), Закон про захист даних у Великій Британії та різні локальні стандарти, вимоги до захисту даних стають обов'язковими для виконання всіма організаціями, які зберігають або обробляють персональні дані. Недотримання цих вимог може призвести до значних штрафів, репутаційних втрат і втрати довіри клієнтів. Ці закони вимагають від організацій забезпечення надійного захисту даних шляхом застосування відповідних технічних і організаційних заходів. Важливо також зазначити, що правове регулювання часто вимагає від організацій зберігати дані протягом певного періоду та забезпечувати їхню доступність [2].

Кількість і складність кіберзагроз постійно зростає, а кіберзлочинці використовують усе більш витончені методи для викрадення, модифікації або знищення даних. Кіберзагрози, такі як фішинг, програми-вимагачі, витік даних через внутрішніх співробітників та інші типи атак, ставлять під загрозу безпеку організацій та можуть призвести до серйозних фінансових і репутаційних втрат. У таких умовах захист даних відіграє критично важливу роль для підтримки стабільності та безпеки інформаційної інфраструктури організації. Завдяки розробці та впровадженню ефективних стратегій і засобів захисту даних організації можуть знижувати ризики, пов'язані з несанкціонованим доступом до інформації.

Втрата або компрометація даних може призвести до серйозних перебоїв у роботі організації. Дані є основою багатьох процесів, і їхня втрата може паралізувати діяльність компанії на тривалий час. Система захисту даних повинна не тільки запобігати втратам, але й забезпечувати швидке відновлення інформації у випадку аварій або атак. Для досягнення цієї мети організації використовують стратегії резервного копіювання, шифрування, а також механізми контролю доступу до даних. Таким чином, ефективний захист даних є запорукою стабільної роботи організації навіть за умов кризових ситуацій.

Захист даних безпосередньо впливає на репутацію організації та рівень довіри клієнтів. У разі витоку даних клієнтів організація може втратити не лише довіру, але й своїх клієнтів, що негативно позначається на її фінансових

результатах та конкурентоспроможності. Для забезпечення конфіденційності даних клієнтів, організація повинна впроваджувати політики захисту даних, які охоплюють як технологічні, так і організаційні заходи. Важливим є також створення прозорої політики щодо обробки та зберігання персональних даних, яка дозволяє клієнтам розуміти, як і для чого використовуються їхні дані.

Захист даних передбачає використання сучасних технологій, таких як шифрування, багатофакторна автентифікація, системи запобігання витокам даних (DLP), штучний інтелект та машинне навчання для виявлення аномалій та загроз. Використання таких технологій дозволяє організаціям проактивно виявляти та нейтралізовувати загрози, знижуючи ризики для своїх даних. Зокрема, застосування моделей на основі штучного інтелекту, як-от трансформери для аналізу текстових даних, може допомогти виявляти патерни, що свідчать про потенційні загрози. Інтеграція технологічних рішень в загальну систему захисту даних підвищує надійність і ефективність інформаційної безпеки організації.

Основні принципи захисту даних є основою забезпечення інформаційної безпеки сучасних організацій. Їхнє дотримання дозволяє зменшити ризики витоку, несанкціонованого доступу або зміни даних, забезпечуючи конфіденційність, цілісність та доступність інформації. Нижче розглянуто детальніше кожен з цих принципів [3].

Конфіденційність полягає у захисті даних від несанкціонованого доступу, гарантуванні, що лише уповноважені особи мають доступ до певної інформації. Цей принцип реалізується за допомогою:

- контролю доступу, який передбачає використання ролей і прав доступу для кожного користувача. Система визначає, хто і на якій підставі може отримувати доступ до конкретних даних;

- шифрування, що забезпечує захист даних під час передачі і зберігання. Навіть у випадку несанкціонованого доступу до зашифрованих даних, зломисники не зможуть їх прочитати без ключа дешифрування;

– політик безпеки – наприклад, політика мінімального доступу, яка передбачає, що кожен користувач отримує доступ тільки до тієї інформації, яка йому необхідна для виконання службових обов’язків [4].

Цілісність забезпечує збереження даних у незмінному, оригінальному вигляді, без несанкціонованих або випадкових змін. Це означає, що дані залишаються достовірними, а будь-які зміни можна простежити. Реалізація цілісності здійснюється за допомогою:

– контрольних сум і хешування, що дозволяють швидко перевірити, чи не було змін у даних. Хеш-алгоритми створюють унікальний код для кожного набору даних, і навіть незначна зміна даних змінює код, вказуючи на порушення цілісності;

– систем ведення журналу – усі дії з даними, включно зі змінами, записуються в журнал. Це дозволяє визначити, хто і коли вносив зміни, забезпечуючи прозорість процесу обробки даних;

– цифрових підписів і сертифікатів, що допомагають підтвердити джерело та справжність даних. Цифровий підпис використовується для забезпечення справжності та гарантії відсутності змін з моменту підпису.

Доступність даних означає, що вони повинні бути доступними для авторизованих користувачів у потрібний час. Це особливо важливо для забезпечення безперервності роботи бізнесу. Щоб гарантувати доступність даних, використовують такі підходи:

– системи резервного копіювання, які дозволяють відновити дані у випадку їх втрати або пошкодження. Регулярні резервні копії забезпечують захист від апаратних і програмних збоїв;

– високонадійне обладнання (сервери, зберігання даних, мережеве обладнання) і дублювання важливих компонентів інфраструктури (відмовостійкість), які гарантують, що система продовжить працювати навіть при збої частини компонентів;

– планування реагування на інциденти та катастрофи (Disaster Recovery Plan, DRP) для швидкого відновлення функціональності у випадку критичного інциденту, наприклад, природної катастрофи чи серйозної кібератаки.

Аутентифікація і авторизація є двома важливими аспектами забезпечення безпеки доступу до даних. Вони гарантують, що тільки уповноважені особи мають доступ до конкретних ресурсів [5].

– аутентифікація підтверджує особу користувача. Це може включати паролі, біометричні дані, токени та багатофакторну аутентифікацію (MFA). MFA поєднує кілька методів, щоб підвищити безпеку, наприклад, пароль та SMS-код;

– авторизація забезпечує контроль доступу, тобто визначає, до яких ресурсів має доступ користувач і що він може з ними робити. Наприклад, співробітник може бути уповноваженим лише переглядати інформацію, але не змінювати її.

Прозорість та підзвітність передбачає збереження інформації про всі операції з даними та можливість контролю доступу до них. Цей принцип забезпечує можливість перевірки дій, які були здійснені з даними, і є важливим для відповідності нормативним вимогам та внутрішнього аудиту.

– журнали подій ведуть записи про всі дії, пов'язані з доступом, зміною та передачею даних, що дозволяє простежити, хто, коли і які операції проводив з інформацією;

– аудит безпеки, який передбачає регулярну перевірку відповідності фактичної діяльності політикам безпеки, допомагає виявляти недоліки у системі захисту даних;

– система управління інцидентами, яка дозволяє відслідковувати, реєструвати та реагувати на інциденти інформаційної безпеки. Це забезпечує швидке виявлення та вирішення проблем, які можуть впливати на безпеку даних.

Мінімізація доступу означає надання користувачам доступу до даних виключно у межах, необхідних для виконання ними своїх службових обов'язків. Такий підхід знижує ризик витоку або несанкціонованих змін. Мінімізація збору даних, у свою чергу, передбачає збір лише тієї інформації, яка дійсно необхідна для

цілей організації. Це знижує обсяг даних, які потрібно захищати, і зменшує потенційні загрози витоку конфіденційної інформації [6].

Найслабкішим місцем у забезпеченні захисту даних часто є людський фактор. Навчання персоналу правильному поводженню з конфіденційною інформацією та розумінню ризиків кібербезпеки допомагає зменшити ризик випадкових порушень. Організації регулярно проводять тренінги для співробітників, пояснюючи важливість захисту даних, а також конкретні заходи, яких слід дотримуватися для запобігання витоку даних.

Дотримання основних принципів захисту даних є основою надійної системи безпеки організації. Вони сприяють побудові довіри клієнтів, запобігають витоку даних та підтримують відповідність законодавчим нормам. Кожен з принципів є важливим компонентом системи захисту, яка дозволяє організаціям захистити свої дані від зовнішніх та внутрішніх загроз і забезпечує стабільну роботу інформаційної інфраструктури [7].

Незважаючи на широкий спектр технологій та методів, забезпечення захисту даних має свої виклики, зокрема:

- складність інтеграції рішень у великі IT-інфраструктури, що можуть складатися з різноманітних систем і додатків;
- постійна еволюція загроз, що вимагає безперервного вдосконалення засобів захисту;
- необхідність відповідності нормативним вимогам у різних юрисдикціях;
- фактор людського впливу, який залишається однією з основних причин виникнення інцидентів.

Захист даних є комплексним процесом, який забезпечує безпеку сучасних організацій, захищає їхні активи і підтримує довіру клієнтів. Впровадження надійних технологій і дотримання принципів інформаційної безпеки є необхідною умовою для збереження конкурентоспроможності та стійкості бізнесу в умовах постійно зростаючих кіберзагроз.

1.2 Огляд і класифікація загроз витоку даних

Загрози витоку даних є однією з найсерйозніших проблем у сфері інформаційної безпеки, оскільки вони можуть призвести до значних фінансових втрат, шкоди репутації та порушення конфіденційності даних. В умовах зростаючих обсягів даних та активного використання цифрових технологій сучасні організації стикаються з різноманітними загрозами витоку інформації. У цьому підрозділі розглянемо основні типи загроз витоку даних та їхню класифікацію.

Загрози витоку даних можуть виникати як всередині організації, так і ззовні. Джерела цих загроз визначають різні підходи до запобігання витоку та захисту інформації. Важливо розрізняти ці загрози для того, щоб розробити ефективні заходи для їх нейтралізації [8].

Внутрішні загрози виникають через дії працівників, партнерів або інших осіб, які мають доступ до інформації всередині організації. Вони можуть бути викликані як навмисними діями, так і випадковими помилками, що призводять до витоку інформації. Основні типи внутрішніх загроз включають:

Навмисний витік інформації:

Це ситуації, коли співробітники навмисно передають або розголошують конфіденційну інформацію. Основні причини такого витоку:

- особиста вигода: працівники можуть продавати інформацію конкурентам або використовувати її у власних інтересах;

- розбіжності з керівництвом або помста: в деяких випадках співробітники, незадоволені своїм положенням або конфліктом із керівництвом, можуть свідомо шукати спосіб завдати шкоди організації;

- соціальна інженерія: навмисний витік також може статися через дії співробітників, які стали жертвами маніпуляцій зловмисників (наприклад, після успішної фішингової атаки).

Ненавмисний витік інформації:

Це випадки, коли співробітники випадково розголошують конфіденційну інформацію через людські помилки або незнання. До таких ситуацій належать:

- відправлення даних неналежному отримувачу: наприклад, випадкове надсилання конфіденційного файлу електронною поштою на неправильну адресу;
- публікація на загальнодоступних ресурсах: співробітники можуть зберегти документи в хмарних сховищах або в соціальних мережах, доступ до яких не захищено, що відкриває можливість витоку;
- недотримання процедур безпеки: наприклад, збереження паролів на робочому столі або недотримання політик шифрування.

Витоки через несанкціонований доступ:

Це випадки, коли співробітники отримують доступ до даних, які виходять за межі їх службових обов'язків:

- перевищення прав доступу: співробітники можуть користуватися привілеями для доступу до інформації, яка їм не потрібна для виконання завдань;
- зловживання правами адміністратора: особи з правами адміністратора можуть мати доступ до широкого спектра даних і, в разі відсутності контролю, можуть використовувати цей доступ для особистих цілей.

Технічні помилки і слабкі місця системи:

Внутрішні витоки даних також можуть статися через технічні помилки, які виникають в результаті недоліків у системах управління доступом або захисту даних:

- конфігураційні помилки: помилки в налаштуваннях можуть призвести до неналежного розмежування доступу до конфіденційних даних;
- невчасні оновлення програмного забезпечення: використання застарілих версій програм, які можуть містити вразливості, що роблять систему схильною до витоків інформації.

Зовнішні загрози створюються діями осіб або організацій, що перебувають поза межами компанії. Зловмисники, які намагаються викрасти або отримати доступ до конфіденційної інформації, використовують різноманітні методи для досягнення своїх цілей. Основні види зовнішніх загроз включають:

Кіберзлочинці:

Це хакери та інші зловмисники, які здійснюють атаки з метою викрадення конфіденційної інформації. Методи, які вони використовують, включають:

- фішингові атаки: обманні листи та сайти, що маскуються під легітимні для отримання доступу до логінів, паролів та інших конфіденційних даних;
- програми-вимагачі: зловмисне ПЗ, яке блокує доступ до даних і вимагає викуп за їх розблокування;
- атаки на вразливості системи: використання вразливостей у програмному забезпеченні, таких як SQL-ін'єкції або атаки на вебдодатки.

Конкуренти:

У деяких випадках конкуренти можуть вдаватися до неетичних методів для отримання доступу до конфіденційних даних іншої компанії, зокрема:

- корпоративне шпигунство: наймання агентів для збору інформації про конкурентів;
- підкуп співробітників: залучення внутрішніх осіб з метою отримання стратегічної інформації;
- сторонні агентства або приватні детективи: деякі компанії можуть залучати професійних фахівців з безпеки або агентства, щоб отримати комерційні секрети.

Хактивісти:

Хактивісти – це групи або індивіди, які атакують інформаційні ресурси організацій з ідеологічних або політичних мотивів. Вони можуть використовувати методи:

- DDoS-атаки: спрямовані на блокування роботи вебсайтів або інших систем, щоб викликати збій і створити погану репутацію компанії;
- розкриття інформації: хактивісти можуть красти інформацію та публікувати її в мережі для громадськості з метою підризу репутації компанії;
- фейкові акаунти: створення підроблених облікових записів для збору конфіденційної інформації або розповсюдження дезінформації.

Державні установи або іноземні розвідки:

Деякі держави використовують власні розвідувальні служби або залучають хакерів для отримання конфіденційної інформації, яка може бути корисною в економічній або політичній сферах [10]:

- кібершпигунство: використання кіберзасобів для проникнення у комп'ютерні системи іноземних компаній з метою викрадення інтелектуальної власності;

- атаки на критичні системи: зокрема, спроби проникнення в енергетичні, фінансові або інфраструктурні компанії з метою отримання доступу до критичних даних.

Інші сторонні організації або особи:

У деяких випадках витік може відбутися через випадкове потрапляння даних до третіх осіб:

- постачальники або підрядники: організація може використовувати послуги сторонніх компаній, і в разі недостатніх заходів безпеки ці компанії можуть випадково або навмисно спричинити витік даних;

- колишні співробітники: особи, які залишили компанію, але зберегли доступ до конфіденційної інформації, можуть зловживати цим для отримання вигоди.

Класифікація загроз за типом атаки є важливим аспектом для розуміння методів, які можуть використовувати зловмисники для викрадення або пошкодження даних. Залежно від методу та цілей, загрози можна розділити на різні типи атак, які зловмисники застосовують для проникнення в системи безпеки та отримання доступу до конфіденційної інформації.

Фішинг є одним з найпоширеніших видів соціальної інженерії, який базується на обмані користувачів. Зловмисники використовують підроблені повідомлення електронної пошти, SMS або фейкові вебсайти, щоб виманити конфіденційні дані, такі як паролі, дані кредитних карток або доступ до облікових записів. Основні види фішингових атак включають:

– масовий фішинг: розсилається велика кількість підроблених повідомлень до широкої аудиторії без конкретної цільової групи. Наприклад, лист із проханням увійти в обліковий запис банку;

– spear-фішинг: цілеспрямовані атаки на конкретних осіб або організації, де повідомлення налаштовані для створення вигляду легітимності. Spear-фішинг є значно небезпечнішим, оскільки зловмисники ретельно вивчають свої цілі;

– вішинг: використання телефонних дзвінків для обману користувачів і отримання конфіденційної інформації;

– смішинг: фішингові атаки через SMS, де зловмисники надсилають повідомлення, яке стимулює користувача надати конфіденційну інформацію.

Зловмисне програмне забезпечення (малваре) – це широкий клас шкідливих програм, які використовуються для отримання несанкціонованого доступу до даних, збору інформації або навіть блокування доступу до даних. Найпоширеніші види малваре включають:

– віруси: прикріплюються до законних програм або файлів і виконують шкідливі дії, як-от видалення файлів або знищення даних;

– трояни: виглядають як легітимні програми, але після встановлення виконують приховані шкідливі дії, як-от відкриття доступу для зловмисників;

– руткіти: призначені для отримання та збереження прихованого доступу до системи. Вони часто змінюють важливі системні файли, щоб зберегти присутність у системі;

– програми-вимагачі (Ransomware): шифрують дані користувача та вимагають викуп за їх розблокування. Вони є серйозною загрозою, оскільки можуть паралізувати діяльність організації;

– шпигунські програми: збирають дані про користувача без його відома та передають їх зловмисникам.

Атаки на мережеву інфраструктуру спрямовані на перехоплення, підміну або викрадення даних, які передаються в мережі. Основні види таких атак включають:

– атака "людина посередині" (MITM): зловмисник перехоплює комунікацію між двома сторонами без їхнього відома, що дозволяє йому отримати доступ до переданих даних або змінити їх;

– сніфінг (перехоплення трафіку): зловмисник використовує спеціальне програмне забезпечення для перехоплення мережевого трафіку з метою отримання конфіденційних даних;

– атаки на мережеві протоколи: наприклад, зловмисник може використовувати вразливості в протоколах, таких як DNS, ARP, щоб перенаправити трафік або змінити його;

– DoS та DDoS атаки: спрямовані на перевантаження серверів або мереж з метою відключення або обмеження доступу до сервісів. Це не лише створює перебої в роботі, а й може стати прикриттям для інших атак.

Brute force та словникові атаки базуються на спробах підібрати паролі або ключі доступу. Ці атаки є трудомісткими, але автоматизовані інструменти дозволяють зловмисникам спробувати тисячі комбінацій за короткий час. Основні види таких атак включають:

– Brute force атака: метод повного перебору можливих комбінацій паролів для отримання доступу до облікового запису;

– словникова атака: зловмисник використовує попередньо підготовлений список часто використовуваних паролів або варіацій (так званий "словник") для підбору пароля;

– гібридні атаки: поєднують перебір із елементами словникової атаки, наприклад, додаючи до звичайного слова популярні символи або цифри, щоб підвищити ймовірність вгадування [11].

Зловмисники активно шукають і використовують вразливості у програмному забезпеченні, що дозволяє їм отримувати доступ до конфіденційної інформації або порушувати нормальне функціонування системи. Основні методи включають:

- SQL-ін'єкції: зловмисники вводять шкідливий SQL-код у поля вводу вебдодатків, що дозволяє їм отримати доступ до бази даних і витягти конфіденційні дані;

- XSS-атаки (Cross-Site Scripting): атаки, що дозволяють вбудувати шкідливий скрипт у вебсторінку, який виконується на комп'ютері користувача, потенційно отримуючи доступ до його даних або сесії;

- буферні переповнення: викликаються при передачі більшої кількості даних, ніж може обробити буфер, що дозволяє зловмиснику виконувати довільний код у системі;

- Zero-day атаки: використовують вразливості, які ще не були виправлені або опубліковані. Ці атаки є особливо небезпечними, оскільки організації можуть не мати захисту від них.

Атаки типу DoS (відмова в обслуговуванні) та DDoS (розподілена відмова в обслуговуванні) спрямовані на перевантаження ресурсів сервера, що призводить до його недоступності для користувачів. Основні види атак включають:

- Flooding (заливка): зловмисник надсилає велику кількість запитів до сервера, використовуючи, наприклад, ICMP, UDP або HTTP-протоколи, що призводить до перевантаження та недоступності сервера;

- Application-layer DDoS атаки: націлені на певні програми або сервіси (наприклад, вебсервери), щоб вичерпати їхні ресурси;

- DNS amplification: використання вразливостей у DNS-серверах для генерування великого обсягу трафіку, який надсилається до жертви з підробленою IP-адресою.

Соціальна інженерія – це метод, який використовує психологічні прийоми для обману співробітників і змушує їх розголосити конфіденційну інформацію або надати доступ до системи. Основні види атак включають:

- фізичний доступ: зловмисники можуть проникати в офіси під виглядом співробітників або обслуговуючого персоналу, отримуючи доступ до комп'ютерів або документів;

– Baiting (приманка): зловмисники залишають заражені USB-накопичувачі або інші носії даних у доступних місцях, розраховуючи на те, що їх підключать до корпоративних пристроїв;

– Pretexting (передтекст): зловмисник вдає себе за іншу особу, наприклад, співробітника служби безпеки, і просить надати певну інформацію, яка нібито потрібна для службових цілей.

Класифікація загроз за типом атаки дозволяє краще розуміти та виявляти потенційні ризики для захисту інформації. Розробка відповідних заходів для кожного виду атак допомагає організаціям захищати свої дані та мінімізувати можливість витоку [12].

Класифікація загроз витоку даних за типом інформації, що піддається ризику, допомагає організаціям зрозуміти, які саме види даних є найвразливішими та які заходи безпеки слід впровадити для їх захисту. Різні типи даних мають різну цінність, рівень конфіденційності та вплив на організацію у випадку їх витоку. Нижче розглянуто основні види даних, що піддаються витоку, та потенційні загрози для кожного типу [13].

Персональні дані стосуються ідентифікаційної інформації про конкретну фізичну особу, наприклад, ім'я, адресу, номер телефону, адресу електронної пошти, дату народження тощо. Витік таких даних може призвести до серйозних наслідків, зокрема:

– порушення конфіденційності: особисті дані часто є конфіденційними, і їхній витік порушує право особи на приватність;

– крадіжка особистих даних (ідентичності): злочинці можуть використовувати викрадені особисті дані для створення фальшивих облікових записів, отримання кредитів або виконання інших дій від імені жертви;

– репутаційні та правові наслідки для компанії: витік персональних даних може призвести до втрати довіри клієнтів та накладення штрафів відповідно до законодавчих норм, таких як GDPR.

Конфіденційна корпоративна інформація – це дані, що стосуються стратегії компанії, фінансових показників, планів розвитку, комерційних таємниць та іншої інформації, яка може мати стратегічне значення для бізнесу. Витік таких даних має серйозні наслідки:

- недобросовісна конкуренція: конкуренти можуть використовувати викрадену інформацію для отримання комерційних переваг або для дискредитації організації;

- фінансові втрати: витік конфіденційної інформації може призвести до зниження доходів, втрати ринкових можливостей або підвищення витрат на відновлення репутації;

- зниження інноваційності: витік інформації про наукові дослідження, нові продукти або технології може позбавити організацію інноваційної переваги.

Фінансова інформація включає дані, пов'язані з транзакціями, банківськими рахунками, номерами кредитних карток, платіжними реквізитами та іншими фінансовими показниками. Витік такої інформації є особливо небезпечним, оскільки може призвести до [15]:

- фінансових збитків: зловмисники можуть використовувати фінансові дані для несанкціонованих транзакцій або доступу до коштів компанії чи клієнтів;

- шахрайства з використанням фінансових даних: наприклад, викрадені номери кредитних карток можуть використовуватися для здійснення шахрайських покупок або отримання кредитів;

- порушення довіри клієнтів: витік фінансової інформації завдає серйозної шкоди репутації компанії, оскільки клієнти очікують, що їхні фінансові дані будуть надійно захищені.

Інтелектуальна власність включає патенти, авторські права, торговельні марки, науково-технічні розробки, інноваційні рішення, методології, програмний код та інші результати інтелектуальної праці. Витік інформації такого типу може мати такі наслідки:

- втрачена конкурентна перевага: конкуренти можуть використовувати інтелектуальну власність для створення аналогічних продуктів або послуг, що знижує унікальність компанії на ринку;

- незаконне використання патентів і торгових марок: зловмисники можуть продавати або використовувати ідеї без згоди власника;

- втрати від потенційного доходу: інтелектуальна власність часто є основним джерелом доходу компанії, і її витік може призвести до значного зниження прибутків [17].

Цей тип даних включає інформацію про контакти, договірну інформацію, історію взаємодії з клієнтами і постачальниками, їхні уподобання тощо. Витік таких даних може призвести до:

- втрати довіри та клієнтської бази: витік інформації про клієнтів підриває їхню довіру до компанії, що може призвести до втрати клієнтів;

- вплив на ділові стосунки: розголошення даних про постачальників або ділових партнерів може нашкодити бізнес-відносинам і вплинути на постачання товарів або послуг;

- можливість атак на клієнтів або постачальників: зловмисники можуть використовувати інформацію про клієнтів для здійснення фішингових атак або для інших видів соціальної інженерії.

Операційні дані включають відомості про внутрішні процеси компанії, документообіг, ланцюги поставок, виробничі дані та інформацію, що підтримує щоденну діяльність організації. Витік операційних даних може призвести до [19]:

- перебоїв у роботі: якщо зловмисники отримають доступ до систем управління виробництвом, логістики або інших критичних систем, це може призвести до перебоїв у роботі;

- зниження ефективності процесів: витік інформації про внутрішні операції може надати конкурентам уявлення про оптимізацію бізнес-процесів, що знижує унікальність процесів;

– ризику для безпеки співробітників: у разі витоку інформації про безпеку виробничих процесів можуть виникнути загрози для здоров'я та безпеки працівників.

Цей тип даних охоплює інформацію про безпекові заходи, політики, плани реагування на інциденти, паролі, ключі шифрування, а також архітектуру захисних систем. Витік даних такого роду може мати критичні наслідки [20].:

– підвищений ризик кібератак: розкриття паролів, ключів шифрування та інших конфіденційних деталей захисту відкриває можливість для зловмисників проникнути в систему;

– компрометація захисних механізмів: зловмисники можуть використовувати знання про системи безпеки для обходу захисних заходів і запуску атак;

– підрив репутації: витік даних про безпеку свідчить про недостатню захищеність організації та може вплинути на довіру клієнтів і партнерів.

Стратегічні дані включають інформацію про довгострокові плани компанії, маркетингові стратегії, прогнози розвитку, плани експансії та інші важливі для розвитку бізнесу відомості. Витік таких даних може призвести до:

– конкурентного зниження: якщо стратегічна інформація потрапляє до конкурентів, вони можуть адаптувати свою стратегію і випередити компанію на ринку;

– фінансових втрат: витік стратегічних планів може змусити компанію коригувати свої дії, що може спричинити додаткові витрати;

– зниження інвестиційної привабливості: витік стратегічних даних може знизити довіру інвесторів, які можуть вважати компанію ненадійною для інвестицій.

Класифікація загроз за типом даних, що піддаються витоку, допомагає визначити, які види інформації потребують підвищеного рівня захисту. Розуміння того, які типи даних можуть зазнати витоку і як це може вплинути на організацію,

дозволяє створити ефективнішу стратегію інформаційної безпеки та знизити ризики витоку даних.

Витоки даних можуть мати різний вплив на організацію залежно від типу інформації, що піддалася витоку, масштабів інциденту та можливих наслідків для компанії. Залежно від рівня впливу загрози поділяються на три основні категорії: низький, середній та високий рівень впливу. Розуміння цих категорій допомагає організаціям належно оцінити ризики та визначити пріоритетність заходів захисту інформації [21].

Загрози з низьким рівнем впливу, як правило, стосуються інформації, витік якої не несе серйозних наслідків для бізнесу та його операцій. Це можуть бути дані, які не мають значної конфіденційності або можуть бути загальнодоступними. Втім, навіть витік такої інформації потребує уваги, щоб уникнути неприємностей, пов'язаних із репутацією та потенційною дезорганізацією процесів.

Приклади загроз із низьким рівнем впливу:

- витік загальнодоступної інформації: наприклад, контактні дані відділу обслуговування клієнтів, загальна інформація про послуги, яку можна знайти на сайті компанії;
- витік неактуальної або застарілої інформації: наприклад, дані про продукти або послуги, які більше не пропонуються;
- невеликий інцидент у внутрішньому середовищі: випадкове обговорення незначних даних у чатах або на внутрішніх платформах, що не призводить до серйозних наслідків.

Наслідки загроз із низьким рівнем впливу:

- мінімальний вплив на репутацію: загальний імідж компанії зазвичай не страждає, оскільки інформація є малозначущою;
- незначний вплив на бізнес-процеси: операційна діяльність компанії не зазнає значних змін, і витрати на відновлення мінімальні;
- низька потреба в ресурсах на вирішення проблеми: для усунення наслідків загроз із низьким рівнем впливу зазвичай не потрібні значні ресурси.

Загрози із середнім рівнем впливу стосуються витоку інформації, яка може завдати шкоди репутації компанії, її відносинам із клієнтами або спричинити фінансові збитки. Хоча такі загрози не призводять до катастрофічних наслідків, вони можуть мати значний негативний вплив на роботу компанії та потребують уважного підходу до усунення та запобігання повторення інцидентів.

Приклади загроз із середнім рівнем впливу:

- витік персональних даних клієнтів: наприклад, електронні адреси, номери телефонів або інформація про покупки клієнтів. Хоча такі дані самі по собі не є критичними, їхній витік може призвести до втрати довіри клієнтів;

- компрометація внутрішньої операційної інформації: наприклад, витік інформації про робочі процеси, процедури або внутрішні документи, які можуть мати значення для ефективності компанії;

- витік даних про співробітників: наприклад, номери телефонів або службові електронні адреси, які можуть бути використані для атак із соціальною інженерією.

Наслідки загроз із середнім рівнем впливу:

- вплив на репутацію: витік інформації середнього рівня конфіденційності може викликати занепокоєння серед клієнтів та партнерів і знизити їхню довіру до компанії;

- фінансові втрати: витрати на заходи з відновлення можуть бути значними, включаючи виплати компенсацій клієнтам, відновлення систем безпеки та репутаційні витрати;

- зміни в операційних процесах: організація може бути змушена змінити свої операційні процедури для запобігання повторенню подібних випадків.

Загрози з високим рівнем впливу – це витоки даних, які мають серйозні та потенційно катастрофічні наслідки для організації. Вони можуть призвести до значних фінансових втрат, серйозних правових наслідків, втрати конкурентних переваг та навіть до можливого припинення діяльності компанії. До таких загроз організація повинна підходити з максимальною увагою, адже їх наслідки можуть бути довгостроковими та складними для усунення [22].

Приклади загроз із високим рівнем впливу:

- витік конфіденційної фінансової інформації: наприклад, інформація про прибутки, втрати, банківські дані або інвестиційні плани. Витік такої інформації може призвести до втрати довіри інвесторів та акціонерів;

- витік інтелектуальної власності: наприклад, патенти, розробки, технології або методології, які є ключовими для конкурентної переваги компанії. Конкуренти можуть використовувати викрадену інтелектуальну власність для створення аналогічних продуктів чи послуг;

- витік стратегічної інформації: наприклад, довгострокові плани розвитку, маркетингові стратегії або плани експансії. Розкриття такої інформації може суттєво вплинути на позицію компанії на ринку;

- компрометація безпекових даних: розкриття інформації про заходи захисту, плани реагування на інциденти, ключі шифрування, паролі, що може призвести до масових атак та інших серйозних інцидентів [23].

Наслідки загроз із високим рівнем впливу:

- катастрофічні репутаційні втрати: витік даних з високим рівнем конфіденційності може серйозно зашкодити іміджу компанії, знизити її ринкову вартість і довіру з боку клієнтів, партнерів та інвесторів;

- великі фінансові втрати: компанія може зазнати значних витрат на усунення наслідків інциденту, відновлення довіри клієнтів, правові витрати (штрафи, судові позови) та додаткові заходи захисту;

- правові наслідки: організація може зазнати серйозних штрафів та санкцій відповідно до норм регулювання, таких як GDPR, PCI DSS, якщо витік даних порушує правові вимоги щодо захисту персональної інформації;

- зниження конкурентоспроможності: витік інтелектуальної власності або стратегічної інформації може призвести до втрати унікальної позиції на ринку, що у довгостроковій перспективі знижує конкурентні переваги компанії;

- можлива зупинка бізнесу: у випадку значних витоків або компрометації даних компанія може бути змушена припинити діяльність на певний час для

відновлення безпеки, що також негативно впливає на її фінансовий стан та позиції на ринку.

Класифікація загроз за рівнем впливу дозволяє оцінити наслідки можливого витоку даних і відповідно розподілити ресурси для запобігання інцидентам різної критичності. Загрози з високим рівнем впливу потребують найбільш серйозних заходів захисту, тоді як загрози із середнім та низьким рівнем впливу, хоча і є менш критичними, все одно потребують відповідного рівня безпеки для запобігання негативним наслідкам для бізнесу.

Різноманітність загроз підкреслює важливість створення багаторівневої системи захисту, яка включає технічні, адміністративні та навчальні заходи. Особливої уваги потребують соціальні загрози та людський фактор, оскільки зловмисники часто використовують соціальну інженерію для обходу технічних засобів захисту. Тому необхідно впроваджувати регулярне навчання персоналу, підвищувати обізнаність щодо ризиків та формувати стійкі навички протистояння психологічним атакам [24].

Таким чином, глибокий аналіз загроз витоку даних і розуміння їх різноманітності є основою для створення ефективної системи інформаційної безпеки, яка забезпечить захист конфіденційної інформації та мінімізує можливість негативних наслідків для організації.

1.3 Основні принципи та підходи до реалізації систем DLP

Системи Data Loss Prevention (DLP) призначені для запобігання витоку конфіденційної інформації за межі організації та забезпечення контролю над даними в процесі їх обробки, зберігання і передачі. Ефективне впровадження DLP-систем вимагає чіткого розуміння принципів роботи, підходів до їх реалізації та адаптації до специфічних потреб організації. Розглянемо основні принципи і підходи до побудови DLP-системи [25].

Розробка та впровадження DLP-системи ґрунтуються на ключових принципах, що дозволяють ефективно запобігати витокам даних і забезпечувати високий рівень захисту конфіденційної інформації.

- ідентифікація та класифікація даних;

Основним принципом роботи DLP-систем є здатність ідентифікувати та класифікувати конфіденційну інформацію. Це включає визначення типів даних, що потребують захисту (наприклад, персональні дані, фінансові звіти, інтелектуальна власність), і класифікацію їх за рівнями конфіденційності. Ідентифікація даних дозволяє системі розпізнавати інформацію, яка потребує спеціального захисту, а класифікація забезпечує правильне застосування політик залежно від важливості даних [26].

- контроль на різних етапах життєвого циклу даних;

DLP-системи повинні контролювати дані на всіх етапах їхнього життєвого циклу: створення, зберігання, передача та видалення. Це дозволяє забезпечити захист як статичних даних (data-at-rest), так і даних у русі (data-in-motion) та під час їх використання (data-in-use). Такий підхід допомагає знизити ймовірність витоку даних на кожному етапі їх обробки.

- моніторинг і контроль доступу до інформації;

DLP-система забезпечує постійний моніторинг і контроль доступу до конфіденційної інформації. Це означає, що кожен доступ до захищених даних реєструється та аналізується, а будь-яка підозріла активність (наприклад, несанкціонований доступ або спроба копіювання конфіденційної інформації) фіксується та піддається розслідуванню. Цей принцип допомагає організаціям відстежувати доступ до важливих даних та запобігати несанкціонованим спробам їх використання.

- прозорість для користувачів та збереження продуктивності;

DLP-система повинна бути максимально прозорою для користувачів, щоб не порушувати робочих процесів. Це означає, що система має працювати у фоновому режимі, мінімізуючи вплив на продуктивність і водночас забезпечуючи високий

рівень захисту. Принцип прозорості також передбачає, що користувачі розуміють правила роботи з конфіденційною інформацією і знають, які дії можуть призвести до санкцій.

- відповідність нормативним вимогам та стандартам;

Сучасні DLP-системи повинні відповідати вимогам нормативних актів та стандартів, таких як GDPR, HIPAA, PCI DSS тощо, які регулюють захист персональних даних, фінансової інформації та інших конфіденційних відомостей. Цей принцип передбачає, що система забезпечує належний рівень захисту інформації і дозволяє організації уникати юридичних та фінансових ризиків.

- адаптивність та гнучкість.

DLP-система має бути адаптивною та гнучкою, щоб відповідати потребам організації та змінюватися разом із розвитком бізнесу та кіберзагроз. Це включає можливість легко налаштовувати політики безпеки, додавати нові джерела даних, інтегрувати систему з іншими інструментами безпеки та IT-інфраструктурою організації [27].

Реалізація DLP-системи в організації може здійснюватися різними методами, залежно від специфіки діяльності компанії, рівня захищеності інформації та ресурсів, що доступні для підтримки системи.

Endpoint DLP передбачає установку агентів на кінцеві пристрої, такі як ноутбуки, комп'ютери, мобільні телефони, для моніторингу та контролю доступу до конфіденційної інформації. Endpoint DLP-системи дозволяють:

- контролювати зберігання, передачу та використання конфіденційних даних на пристроях;

- відстежувати операції з копіюванням або переміщенням файлів на зовнішні носії (USB, CD/DVD);

- забезпечувати шифрування даних і блокувати несанкціоновані дії користувачів із захищеними файлами.

Endpoint DLP підходить для організацій, де потрібно захистити інформацію безпосередньо на робочих пристроях користувачів. Важливою перевагою є можливість контролю дій співробітників навіть у випадку віддаленої роботи.

Network DLP-системи інтегруються у мережеву інфраструктуру організації та контролюють передавання даних по мережі. Основні функції Network DLP:

- перехоплення та аналіз мережевого трафіку для виявлення несанкціонованої передачі конфіденційної інформації за межі мережі організації;
- застосування фільтрів і правил для блокування відправки певних типів даних (наприклад, документів з фінансовою або персональною інформацією) через електронну пошту, FTP, вебформи та інші канали зв'язку;
- запобігання передачі даних через незахищені канали, що особливо важливо для запобігання атакам типу «людина посередині» (MITM).

Network DLP забезпечує високий рівень захисту інформації під час її передавання по мережі і дозволяє контролювати потоки даних у реальному часі.

Discovery DLP-системи здійснюють сканування сховищ, серверів, кінцевих пристроїв для ідентифікації конфіденційної інформації та її класифікації за рівнем важливості. Основні завдання Discovery DLP:

- виявлення конфіденційної інформації в різних сховищах (сервери, бази даних, файлові сховища) та класифікація за типом (персональні дані, фінансові звіти, інтелектуальна власність);
- створення інвентаризації конфіденційних даних для подальшого налаштування політик безпеки;
- автоматичне застосування політик до виявлених даних залежно від їхнього рівня важливості (наприклад, шифрування або переміщення до захищених сховищ).

Discovery DLP дозволяє компанії виявляти і класифікувати конфіденційні дані, які вже зберігаються в системах, та уникати витоків даних через неналежне зберігання.

Cloud DLP забезпечує захист конфіденційних даних, що зберігаються або передаються в хмарних середовищах. З огляду на зростаючу популярність хмарних сервісів, Cloud DLP стає важливим елементом захисту інформації. Його функції включають:

- контроль доступу до конфіденційних даних у хмарних сховищах;
- шифрування файлів перед завантаженням у хмару;
- відстеження дій користувачів і запобігання несанкціонованому доступу до файлів у хмарних сховищах.

Cloud DLP забезпечує ефективний захист даних у хмарі, що є актуальним для організацій, які активно використовують хмарні сервіси для зберігання та обробки інформації.

Content-aware DLP забезпечує аналіз вмісту файлів для визначення конфіденційної інформації. Цей підхід дозволяє ідентифікувати вразливі дані на основі вмісту, а не лише за типом файлу. Content-aware DLP включає:

- аналіз тексту документів, електронних листів, баз даних для виявлення специфічної інформації (наприклад, номери кредитних карток, номери соціального страхування);
- визначення шаблонів і маркерів конфіденційної інформації, таких як адреси, ідентифікаційні номери, за допомогою регулярних виразів або алгоритмів машинного навчання;
- автоматичне застосування політик для різних типів вмісту залежно від рівня конфіденційності.

Content-aware DLP підходить для організацій, які обробляють великі обсяги конфіденційної інформації, що може міститися в різноманітних форматах і документах.

Для ефективного функціонування DLP-система повинна бути інтегрована з іншими системами безпеки, такими як антивірусні програми, міжмережеві екрани, SIEM-системи (Security Information and Event Management) та інші елементи кібербезпеки організації. Інтеграція дозволяє:

- централізоване управління політиками: налаштування єдиних політик, які застосовуються до всіх аспектів безпеки, що значно спрощує моніторинг і реагування на інциденти;

- обмін інформацією між системами: наприклад, виявлення підозрілої активності у SIEM може спонукати DLP-систему до застосування додаткових заходів безпеки для конкретних користувачів чи пристроїв;

- автоматизація заходів реагування: DLP-система може автоматично блокувати дії, що порушують політику безпеки, або надсилати сповіщення відповідальним особам для подальших дій.

Автоматизація та інтеграція DLP-систем із загальною стратегією кібербезпеки організації значно підвищують ефективність захисту, знижують ризики витоку інформації та полегшують процеси управління безпекою.

Ефективна реалізація DLP-системи потребує чіткого дотримання основних принципів – від ідентифікації даних до відповідності нормативним вимогам – та застосування гнучких підходів залежно від специфіки організації. Різні типи DLP-систем, як-от Endpoint, Network, Discovery та Cloud, дозволяють забезпечити комплексний захист інформації на всіх етапах її життєвого циклу, а інтеграція з іншими системами безпеки підсилює можливості DLP у виявленні та запобіганні витокам.

1.4 Аналіз сучасних моделей машинного навчання для використання в DLP-системах

Сучасні моделі машинного навчання (ML), зокрема на основі глибинного навчання, значно розширюють можливості Data Loss Prevention (DLP) систем. Застосування алгоритмів машинного навчання дозволяє DLP-системам адаптуватися до складних загроз, ефективніше ідентифікувати аномалії та аналізувати великий обсяг даних з урахуванням контексту. У цьому розділі розглянемо основні типи сучасних моделей, які можна застосовувати в DLP-системах, їх переваги, недоліки та перспективи.

DLP-системи часто аналізують текстові документи, електронні листи, повідомлення у чатах, що можуть містити конфіденційну інформацію. Моделі обробки природної мови (NLP) здатні забезпечити точну класифікацію, розпізнавання та контекстний аналіз текстових даних. Основні моделі NLP для використання в DLP-системах включають:

- Word2Vec і GloVe;

Ці базові моделі для векторизації слів дозволяють перетворювати текстові дані у числові вектори, що полегшує їх класифікацію. Word2Vec і GloVe вивчають зв'язки між словами на основі контексту, дозволяючи DLP-системам визначати семантичну близькість між словами. Хоча вони підходять для простих сценаріїв класифікації, їхня здатність до розпізнавання складних контекстів обмежена.

- BERT (Bidirectional Encoder Representations from Transformers);

BERT є однією з найпопулярніших трансформерних моделей для обробки тексту. Він дозволяє DLP-системам глибше розуміти контекст, враховуючи як попередні, так і наступні слова. Це значно покращує точність класифікації конфіденційної інформації у складних текстах. BERT особливо корисний для роботи з довгими документами, де конфіденційна інформація може бути представлена у складних реченнях або фразах.

- GPT (Generative Pre-trained Transformer);

Модель GPT також використовує архітектуру трансформерів і дозволяє генерувати текст на основі попереднього контексту. У DLP-системах GPT допомагає не лише ідентифікувати конфіденційну інформацію, але й прогнозувати, чи може конкретна інформація призвести до витоку, якщо буде передана за межі організації. Це корисно для аналізу електронної пошти та внутрішніх чатів, де текст є динамічним і може містити потенційно чутливі дані [29].

- Attention-based моделі.

Моделі з механізмом уваги (Attention) дозволяють фокусуватися на певних частинах тексту, які є важливими для розпізнавання конфіденційної інформації. Attention забезпечує глибше розуміння залежностей між словами та виявляє

чутливу інформацію навіть у великому обсязі тексту. Такий підхід дозволяє DLP-системам швидко знаходити важливі фрагменти в довгих документах та повідомленнях.

Виявлення аномальної поведінки користувачів є важливим завданням для DLP-систем, оскільки несанкціоновані дії можуть свідчити про потенційний витік даних. Сучасні моделі машинного навчання дозволяють виявляти аномалії в поведінці та швидко реагувати на них.

– кластеризаційні моделі (K-Means, DBSCAN);

Ці моделі кластеризації використовуються для виявлення поведінкових патернів користувачів. У DLP-системах моделі кластеризації допомагають групувати дії користувачів і виявляти відхилення від стандартних шаблонів. Наприклад, якщо користувач, який зазвичай працює з певним набором даних, раптом починає копіювати нові типи документів, це може вказувати на аномальну поведінку.

– автокодувальники (Autoencoders);

Автокодувальники використовуються для стиснення та відновлення даних. У контексті DLP, автокодувальники навчаються на стандартній поведінці користувачів і здатні визначати значні відхилення від неї. Наприклад, якщо користувач починає завантажувати незвично великий обсяг даних або отримує доступ до файлів, які зазвичай не використовує, автокодувальник позначає таку поведінку як аномальну.

– рекурентні мережі для часових рядів (LSTM, GRU).

Моделі LSTM та GRU дозволяють аналізувати часові послідовності дій, що особливо корисно для виявлення поведінкових змін, які відбуваються з часом. Це дає змогу DLP-системі виявляти підозрілі дії, які відхиляються від звичайного графіка або частоти. Наприклад, підвищена активність користувача у неробочий час може свідчити про потенційний витік даних.

Інформація, що потребує захисту, може бути представлена не лише у вигляді тексту, але й у зображеннях чи сканах документів. Моделі для обробки зображень

забезпечують можливість розпізнавання конфіденційної інформації в графічному контенті.

- Convolutional Neural Networks (CNN);

CNN-моделі використовуються для обробки візуальних даних, що дозволяє ідентифікувати текст або інші важливі елементи на зображеннях. В DLP-системах CNN можуть бути використані для аналізу сканованих документів, фотографій або скріншотів. Вони дозволяють ідентифікувати важливу інформацію на зображеннях та запобігати її несанкціонованій передачі.

OCR-технології дозволяють витягати текст із зображень і переводити його в машиночитний формат. Це особливо корисно для аналізу сканів документів, контрактів або чеків, що можуть містити конфіденційну інформацію. OCR допомагає DLP-системам ідентифікувати текстову інформацію у зображеннях та застосовувати політики захисту для запобігання витоку.

Capsule Networks є альтернативою CNN, що забезпечують збереження просторових відношень між об'єктами на зображеннях. Вони дозволяють DLP-системам краще розпізнавати інформацію, якщо текст міститься у складних структурованих зображеннях, наприклад, документи з таблицями або графічними елементами, що мають значення.

Сучасні ML-моделі дозволяють автоматизувати процес прийняття рішень у разі виявлення потенційних загроз у DLP-системах. Це підвищує швидкість реакції та знижує ймовірність витоків даних.

Ці моделі прийняття рішень використовуються для автоматичної оцінки рівня загрози на основі ймовірності. Наприклад, якщо ймовірність несанкціонованого доступу висока, модель може заблокувати доступ або повідомити відповідальні особи. Це дозволяє DLP-системам ефективно керувати ризиками та швидко реагувати на потенційні загрози.

- моделі з підкріпленням (Reinforcement Learning);

Моделі з підкріпленням навчаються на основі попередніх інцидентів і адаптуються до нових загроз, оптимізуючи реакцію на певні події. Вони здатні

самостійно приймати рішення на основі накопиченого досвіду та змінювати політики безпеки в режимі реального часу. Наприклад, DLP-система може запобігати витокам, автоматично блокуючи підозрілі дії або адаптуючи свою поведінку до нових патернів користувачів.

GAN-моделі використовуються для навчання DLP-систем на нових типах загроз. GAN генерують нові приклади потенційно небезпечної поведінки, що дозволяє DLP-системам краще виявляти аномалії та вдосконалювати процес виявлення загроз. Це допомагає уникнути помилок у розпізнаванні нових видів атак.

Таким чином, аналіз сучасних моделей машинного навчання показує, що їхнє використання в DLP-системах суттєво покращує можливості захисту конфіденційної інформації. Кожна модель має свої переваги і недоліки, що дозволяє обирати оптимальний інструмент для вирішення конкретних завдань, як-от аналіз текстових даних, обробка зображень, виявлення аномалій у поведінці користувачів та автоматизація реагування на загрози [30].

Нижче наведена таблиця 1.1, яка порівнює основні моделі машинного навчання для застосування в DLP-системах, їх призначення, переваги, обмеження та сфери використання.

Таблиця 1.1 – Порівняння сучасних моделей машинного навчання для застосування в DLP-системах

Тип моделі	Призначення	Переваги	Недоліки	Приклади використання в DLP
Моделі NLP (Word2Vec, GloVe)	Векторизація тексту для пошуку схожих слів і класифікації	Простота у використанні; здатність до розпізнавання базових семантичних зв'язків	Необхідність у великому обсязі даних; низька точність для складних контекстів	Ідентифікація загальних категорій конфіденційної інформації

Продовження таблиці 1.1

Тип моделі	Призначення	Переваги	Недоліки	Приклади використання в DLP
Трансформерні моделі (BERT, GPT)	Контекстуальний аналіз тексту, розуміння складних залежностей між словами	Висока точність і контекстуальне розуміння; ефективні для довгих і складних текстів	Високі обчислювальні витрати; потреба в великому обсязі навчальних даних	Розпізнавання конфіденційної інформації в текстових документах, електронних листах і чатах
Моделі на основі Attention	Визначення ключових слів і фраз у довгих текстах	Здатність знаходити важливу інформацію навіть у великих документах	Можуть втрачати контекст на рівні документа; потребують складних налаштувань	Виділення ключових слів і фраз для швидкого розпізнавання конфіденційної інформації
Кластеризація (K-Means, DBSCAN)	Групування поведінкових патернів користувачів для виявлення аномалій	Простота і швидкість роботи; не потребують навчальних міток	Низька ефективність при складних і сильно неоднорідних даних	Аналіз активності користувачів для виявлення підозрілих дій
Автокодувальники (Autoencoders)	Виявлення аномалій у поведінці, відхилень від звичайних шаблонів	Можливість виявлення аномальних дій без додаткових даних; добре обробляють аномальні дані	Чутливі до гіперпараметрів; потребують налаштування для кожного типу даних	Моніторинг дій співробітників для виявлення відхилень від стандартної поведінки

Продовження таблиці 1.1

CNN (Convolutional Neural Networks)	Розпізнавання зображень, графічних файлів, аналіз тексту у зображеннях	Висока точність для обробки зображень; здатність виділяти важливі елементи візуального контенту	Чутливість до якості зображень; потребують великих обчислювальних ресурсів	Виявлення конфіденційної інформації на зображеннях, сканах документів
OCR (Оптичне розпізнавання тексту)	Розпізнавання тексту на зображеннях та сканованих документах	Точність розпізнавання тексту у зображеннях; можливість витягування тексту з графічного контенту	Складність при розпізнаванні низькоякісних зображень або складних шрифтів	Аналіз тексту у зображеннях для запобігання витоку конфіденційної інформації через зображення
Capsule Networks	Альтернатива CNN для обробки складних зображень з кращим збереженням просторових відношень	Висока точність у розпізнаванні зв'язків між елементами зображень; здатність до складнішого аналізу	Велика обчислювальна складність; потреба у великих наборах даних для навчання	Аналіз складних зображень із текстом чи графікою, що містять конфіденційну інформацію
Моделі логістичної регресії та дерева рішень	Автоматичне прийняття рішень на основі ймовірності загроз	Швидкість і простота у використанні; добре інтерпретуються	Обмежена точність при роботі зі складними даними	Автоматизація блокування підозрілих дій, швидке прийняття рішень у разі виявлення потенційної загрози

Продовження таблиці 1.1

Тип моделі	Призначення	Переваги	Недоліки	Приклади використання в DLP
Reinforcement Learning (Підкріплення)	Адаптація політик безпеки та реагування на інциденти	Здатність до навчання на основі попередніх дій; адаптивність до нових загроз	Складність у навчанні та налаштуванні; вимагає багато часу для оптимального результату	Автоматичне налаштування політик безпеки залежно від зміни поведінкових патернів у системі
Генеративні змагальні мережі (GAN)	Виявлення нових типів загроз і навчання на потенційно небезпечних прикладах	Можливість генерування нових прикладів для навчання моделей; підвищення точності на нових даних	Висока обчислювальна складність; складність у навчанні мереж	Створення нових шаблонів загроз, аналіз аномальної поведінки для покращення DLP-системи

Застосування сучасних моделей машинного навчання в DLP-системах значно розширює їх функціональні можливості, дозволяючи адаптуватися до нових загроз, автоматизувати процес реагування на інциденти та глибоко аналізувати дані. Завдяки моделям NLP для обробки текстових даних, CNN та OCR для роботи із зображеннями, а також моделям для виявлення аномалій та автоматизації прийняття рішень, DLP-системи можуть ефективніше захищати конфіденційну інформацію. Інтеграція цих моделей забезпечує гнучкий і багаторівневий підхід до інформаційної безпеки, що допомагає організаціям краще протистояти сучасним кіберзагрозам і знижувати ризики витоку даних.

1.5 Висновки та постановка задачі

В рамках даної роботи проведено дослідження сучасних підходів до запобігання витоку даних у системах Data Loss Prevention (DLP), зокрема можливостей використання моделей машинного навчання для аналізу даних та виявлення аномальної поведінки. Було розглянуто основні принципи побудови DLP-систем, детально проаналізовано методи класифікації та контролю даних, обробки текстової інформації, ідентифікації аномалій та моніторингу поведінкових патернів користувачів. Особлива увага приділялася сучасним трансформерним моделям, методам кластеризації та підходам до аналізу текстових і графічних даних, а також автоматизації процесу реагування на інциденти.

Дослідження дозволило отримати всебічне розуміння сучасних загроз та методів їх виявлення у рамках DLP-систем, що є критичним для забезпечення інформаційної безпеки організацій та запобігання витоку конфіденційної інформації.

Отже, на основі проведеного аналізу було сформульовано ряд завдань для подальшої розробки DLP-системи:

- розробка алгоритму класифікації конфіденційних даних із застосуванням трансформерних моделей для підвищення точності;
- проектування механізму моніторингу аномальної поведінки користувачів із використанням моделей кластеризації та автокодувальників;
- створення модуля для аналізу графічних даних із використанням CNN та OCR для розпізнавання конфіденційної інформації у зображеннях;
- автоматизація процесу реагування на інциденти з витоками даних за допомогою моделей підкріплення та дерев рішень;
- інтеграція DLP-системи з іншими компонентами інформаційної безпеки для забезпечення комплексного підходу до захисту інформації.

Ці завдання сприятимуть подальшій розробці DLP-системи, здатної ефективно протистояти сучасним загрозам та забезпечувати високий рівень захисту даних.

2 СТВОРЕННЯ ЗАХИЩЕНОЇ СИСТЕМИ DLP ЗАПОБІГАННЯ ВИТОКУ ДАНИХ НА ОСНОВІ TRANSFORMER-BASED МОДЕЛІ GPT І АЛГОРИТМУ КЛАСТЕРИЗАЦІЇ З ВИКОРИСТАННЯМ K-MEANS ДЛЯ ВИЯВЛЕННЯ АНОМАЛІЙ І КОНТЕКСТНОГО АНАЛІЗУ ДАНИХ

У цьому розділі буде розглянуто процес створення захищеної системи DLP, що поєднує можливості трансформерних моделей GPT для контекстного аналізу тексту з алгоритмом кластеризації K-Means для виявлення аномалій. Буде досліджено ключові аспекти створення такої системи, зокрема особливості її архітектури, інтеграції трансформерних моделей для аналізу текстових даних та розробки алгоритму кластеризації для аналізу поведінкових патернів. Також будуть представлені алгоритми для кожного компонента системи, їхня взаємодія та впровадження в єдину систему DLP. Розділ завершиться висновками щодо отриманих результатів і оцінкою їхньої відповідності поставленим задачам.

2.1 Особливості створення захищеної системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації

Для забезпечення ефективного запобігання витоку даних і контролю над конфіденційною інформацією система DLP, побудована на основі моделі GPT і алгоритму кластеризації K-means, повинна відповідати низці вимог. Ці вимоги допомагають гарантувати безпеку даних, точність аналізу, швидкість реагування на загрози та зручність використання. Основні вимоги до системи включають:

Функціональні вимоги

Контекстуальний аналіз текстових даних

Система має здійснювати контекстуальний аналіз текстових даних для розпізнавання конфіденційної інформації. Це особливо важливо у випадках, коли текстова інформація подається в непрямому вигляді або з використанням синонімів. Завдяки можливостям моделі GPT, система повинна виявляти чутливу інформацію навіть у складних текстових контекстах, таких як юридичні документи, листування чи комерційні угоди [31].

Виявлення аномалій у поведінці користувачів

Для виявлення нетипової поведінки користувачів, яка може вказувати на потенційні загрози, система повинна забезпечувати аналіз поведінкових патернів. Алгоритм K-means дозволяє класифікувати поведінкові дані, виявляти аномалії та попереджати про можливі порушення на основі аналізу стандартних і підозрілих патернів.

Швидка реакція на потенційні загрози

Система повинна забезпечувати швидку реакцію на виявлені загрози, що включає в себе блокування дій користувачів, надсилання сповіщень, обмеження доступу або активацію додаткових заходів безпеки. Швидкість реагування є критичним аспектом, що дозволяє мінімізувати наслідки несанкціонованого доступу або витоку даних.

Висока точність та мінімізація хибнопозитивних спрацьовувань

Важливим аспектом є забезпечення високої точності ідентифікації конфіденційної інформації та аномальної поведінки, що дозволяє мінімізувати хибні спрацьовування. Це досягається завдяки оптимізації моделі GPT для текстового аналізу та налаштуванню алгоритму K-means для чутливого виявлення відхилень у поведінці користувачів.

Адаптивність до нових загроз і змін у поведінці користувачів

Система повинна бути адаптивною та здатною до самостійного навчання на нових даних для забезпечення ефективного захисту від нових загроз і виявлення змін у поведінкових патернах. Використання методів машинного навчання дає змогу системі постійно вдосконалювати свої механізми захисту і коригувати політики безпеки.

Нефункціональні вимоги

Оптимізація обчислювальних ресурсів

Оскільки моделі GPT та алгоритми кластеризації є ресурсомісткими, важливо оптимізувати використання обчислювальних ресурсів для забезпечення стабільної роботи системи. Це передбачає налаштування дозованого аналізу даних,

шардування, оптимізацію роботи кластеризації, а також застосування паралельної обробки для великих обсягів даних.

Масштабованість системи

Система повинна бути здатною до масштабування, щоб забезпечити роботу з великим обсягом даних і кількістю користувачів. Це включає можливість додавання нових обробних вузлів або серверів для підтримки роботи системи при підвищеному навантаженні.

Інтеграція з існуючими системами безпеки

Для досягнення комплексного підходу до захисту інформації система DLP повинна легко інтегруватися з іншими інструментами безпеки в організації, такими як антивірусне програмне забезпечення, системи управління подіями інформаційної безпеки (SIEM), міжмережеві екрани, системи керування доступом тощо [32].

Зручність використання та прозорість для кінцевих користувачів

Система має бути зручною для використання адміністраторами та водночас прозорою для кінцевих користувачів, щоб не перешкоджати їхньому робочому процесу. Це включає наявність зрозумілого інтерфейсу для налаштування політик безпеки та можливість автоматичного визначення загроз із мінімальним залученням користувачів.

Функціональні компоненти системи

Для досягнення цих вимог система DLP включає наступні основні функціональні компоненти:

Модуль обробки тексту на основі GPT

Відповідальний за обробку текстових даних і виявлення конфіденційної інформації за допомогою контекстного аналізу. Цей модуль отримує текстові дані (наприклад, електронні листи, документи) і класифікує їх за рівнем конфіденційності. Він також ідентифікує конкретні терміни або фрази, які можуть свідчити про наявність конфіденційної інформації.

Модуль кластеризації поведінкових даних на основі K-means

Здійснює аналіз поведінки користувачів та ідентифікацію аномальної активності. На основі зібраних даних про взаємодії користувачів із конфіденційною інформацією створюються кластери, що представляють типові патерни поведінки. У разі виявлення аномалій, цей модуль генерує сповіщення або ініціює процедури безпеки.

Модуль реагування на інциденти

Автоматично реагує на виявлені загрози, викликаючи певні дії, такі як блокування користувачів, обмеження доступу до інформації, надсилання сповіщень адміністраторам та активування інших процедур безпеки. Цей модуль зберігає інформацію про всі інциденти для подальшого аналізу.

Інтеграційний шар для обміну даними між компонентами

Забезпечує взаємодію між модулем GPT і модулем кластеризації, а також передачу інформації між ними для отримання комплексного аналізу. Цей шар також відповідає за зв'язок із зовнішніми системами безпеки, такими як SIEM, і забезпечує обмін інформацією про загрози.

Інтерфейс адміністратора

Призначений для конфігурації системи, налаштування політик безпеки, перегляду інцидентів, а також звітності. Інтерфейс адміністратора дозволяє проводити оцінку загроз, керувати політиками доступу та отримувати сповіщення про потенційні витoki даних.

Архітектура системи DLP на основі Transformer-based моделі GPT і алгоритму кластеризації K-means є багаторівневою і передбачає комплексну взаємодію між компонентами, що відповідають за обробку текстових даних, виявлення аномалій у поведінці користувачів та реагування на загрози. Правильно спроектована архітектура забезпечує високу продуктивність системи, швидкість обробки запитів і гнучкість у налаштуванні політик безпеки.

Основні компоненти архітектури

Модуль обробки тексту на основі GPT

Призначення: цей модуль виконує функцію глибокого контекстного аналізу текстових даних для ідентифікації конфіденційної інформації.

Функціональність: модуль на основі GPT розпізнає контексти, в яких використовується конфіденційна інформація, визначає її рівень важливості та класифікує текст за відповідними категоріями ризику. Для цього він використовує як загальні мовні моделі, так і спеціально налаштовані алгоритми на основі GPT, що навчені на специфічних даних організації.

Інтеграція: модуль обміну даними передає результати аналізу від GPT до інших компонентів системи для подальшої обробки або активації захисних дій.

Модуль кластеризації поведінкових даних на основі K-means

Призначення: модуль призначений для моніторингу і кластеризації поведінки користувачів, зокрема для виявлення відхилень від стандартних патернів.

Функціональність: модуль використовує алгоритм K-means для створення кластерів поведінки на основі історичних даних. Кластеризація дозволяє виділити групи дій користувачів, які є типовими, та ідентифікувати аномальну поведінку, яка може свідчити про потенційний витік даних або загрозу безпеці.

Інтеграція: результати кластеризації передаються до модуля реагування на інциденти для подальшої оцінки та визначення дій, необхідних для усунення загроз.

Призначення: цей компонент забезпечує взаємодію між модулями обробки тексту та кластеризації, а також здійснює обмін даними з іншими компонентами системи безпеки [33].

Функціональність: інтеграційний шар працює як центральний вузол обміну даними, що дозволяє модулю GPT передавати інформацію про ідентифіковану конфіденційну інформацію до кластеризаційного модуля K-means і навпаки. Він також забезпечує взаємодію з іншими системами безпеки, такими як SIEM (системи керування подіями інформаційної безпеки), антивірусні програми та засоби контролю доступу.

Особливості: інтеграційний шар адаптується до великої кількості запитів та дозволяє ефективно використовувати ресурси за рахунок обробки даних у режимі реального часу.

Модуль реагування на інциденти

Призначення: забезпечує автоматичну реакцію на виявлені загрози та координацію дій для запобігання витоку даних.

Функціональність: цей модуль автоматично активує певні дії, такі як блокування доступу користувача, обмеження прав доступу або надсилання сповіщень адміністратору у разі виявлення потенційної загрози. Він також зберігає інформацію про інциденти для подальшого аналізу.

Інтеграція: модуль реагування тісно пов'язаний з іншими компонентами системи і використовує результати модулів GPT і K-means для ухвалення рішень про блокування або моніторинг дій користувачів.

Інтерфейс адміністратора

Призначення: призначений для налаштування параметрів системи, управління політиками безпеки та перегляду журналів інцидентів.

Функціональність: через інтерфейс адміністратор має можливість змінювати налаштування моделі GPT та K-means, встановлювати політики для обробки текстових даних, змінювати параметри кластеризації, а також переглядати та аналізувати інциденти.

Інтеграція: інтерфейс підключений до всіх основних компонентів системи, що дозволяє адміністратору швидко змінювати налаштування та отримувати актуальні дані про стан безпеки.

Потоки даних і взаємодія між компонентами

Обробка текстових даних

Коли новий текстовий документ або повідомлення надходить до системи, він спочатку передається в модуль обробки тексту на основі GPT. Цей модуль виконує контекстуальний аналіз, ідентифікуючи конфіденційну інформацію, і передає результати до інтеграційного шару для подальшого опрацювання.

Моніторинг і кластеризація поведінки

Після обробки тексту інтеграційний шар передає дані до модуля кластеризації, де на основі алгоритму K-means визначається поведінковий патерн

користувача. Якщо відхилення від стандартних патернів перевищує допустимі норми, цей компонент передає інформацію до модуля реагування.

Реакція на інциденти

Модуль реагування на інциденти отримує інформацію про потенційну загрозу від інтеграційного шару і кластеризаційного модуля. На основі отриманих даних він вирішує, чи потрібно блокувати доступ, обмежити права користувача або надати сповіщення адміністраторам.

Інтеграція з іншими системами безпеки

Інтеграційний шар також відповідає за передачу інформації до інших інструментів безпеки, таких як SIEM або системи антивірусного захисту, забезпечуючи комплексний підхід до запобігання витокам.

Переваги багаторівневої архітектури

Гнучкість у налаштуванні політик безпеки

Архітектура дозволяє адміністраторам налаштовувати політики для кожного компонента системи відповідно до потреб організації, що забезпечує високий рівень адаптивності.

Висока швидкість і точність аналізу

Розподіл функцій між окремими модулями дозволяє підвищити ефективність обробки даних і знижує навантаження на обчислювальні ресурси, що позитивно впливає на швидкість та точність системи.

Комплексний підхід до безпеки

Взаємодія між моделлю GPT для аналізу тексту та алгоритмом кластеризації K-means для поведінкового моніторингу створює комплексну систему, яка може одночасно аналізувати як текстові, так і поведінкові аспекти даних.

Масштабованість і надійність

Така архітектура дозволяє легко масштабувати систему, додаючи нові обробні вузли або інтегруючи з іншими системами безпеки. Це забезпечує високий рівень надійності навіть за умов збільшення навантаження.

2.2 Розробка алгоритму інтеграції моделі gpt

Для ефективного використання можливостей трансформерної моделі GPT в системі DLP розроблено алгоритм інтеграції, який забезпечує точний і контекстуальний аналіз текстових даних з метою виявлення та захисту конфіденційної інформації. Цей алгоритм побудований таким чином, щоб використовувати модель GPT для розпізнавання чутливих даних в реальному часі, адаптуючи модель до конкретних вимог організації. Алгоритм інтеграції передбачає кілька етапів, кожен з яких спрямований на налаштування, оптимізацію та застосування моделі GPT у контексті DLP-системи.

Підготовка та передобробка даних

Очищення даних: На цьому етапі текстові дані очищаються від зайвої інформації, наприклад, HTML-тегів, спеціальних символів та форматування, що може впливати на точність аналізу. Це дозволяє GPT зосередитися на змісті та уникати «шуму».

Токенізація: Текст розбивається на окремі частини (токени), які зручно обробляти моделлю. GPT використовує токени для ефективного контекстуального аналізу.

Визначення категорій ризику: На основі специфіки діяльності організації текстові дані категоризуються за типом конфіденційності (наприклад, фінансові дані, персональні дані, комерційна таємниця). Це полегшує налаштування політик безпеки для подальшої обробки.

Налаштування та адаптація моделі GPT

Доменне навчання: Якщо GPT планується використовувати для спеціалізованих текстів (наприклад, юридичних або медичних документів), модель може пройти додаткове навчання на доменних даних, що містять специфічні терміни та контексти, які характерні для даної галузі.

Оптимізація параметрів моделі: Налаштування GPT за ключовими параметрами (наприклад, розмір контексту, кількість обчислень для кожного запиту) забезпечує баланс між точністю та швидкістю роботи. Це дозволяє уникати

надмірного споживання обчислювальних ресурсів при обробці великих обсягів текстових даних.

Виявлення конфіденційної інформації: Модель GPT аналізує текст, розпізнаючи конфіденційну інформацію в залежності від її контексту. Наприклад, модель може розпізнати номери кредитних карток, фінансові звіти або особисті дані, навіть якщо вони подані в непрямому вигляді.

Контекстуальний інтерпретатор: За допомогою контекстуального аналізу GPT здатна розуміти значення тексту у широкому контексті, що підвищує точність розпізнавання. Це дозволяє уникати випадків, коли модель неправильно класифікує загальнодоступну інформацію як конфіденційну.

Категоризація інформації: На основі результатів аналізу модель GPT призначає кожному текстовому блоку певну категорію ризику (наприклад, низький, середній, високий рівень конфіденційності). Це спрощує застосування політик безпеки залежно від важливості даних.

Маркування даних: Конфіденційна інформація маркується спеціальними тегами для подальшого контролю. Наприклад, можна використовувати теги «Особисті дані», «Фінансові дані», що спрощує подальшу обробку та моніторинг.

Інтеграція з іншими модулями системи DLP

Передача результатів в інтеграційний шар: Результати аналізу моделі GPT передаються до інтеграційного шару, який дозволяє передавати інформацію іншим компонентам системи, наприклад, модулю кластеризації K-means або модулю реагування на інциденти.

Взаємодія з модулем кластеризації: Інформація про конфіденційні дані передається до модуля кластеризації, який порівнює поведінкові патерни користувача з типовими, виявляючи відхилення, які можуть свідчити про загрозу.

Модуль реагування на інциденти: У разі виявлення потенційно небезпечних дій система автоматично активує модуль реагування, який, в залежності від налаштувань, блокує дії користувача, обмежує доступ або надсилає сповіщення адміністраторам.

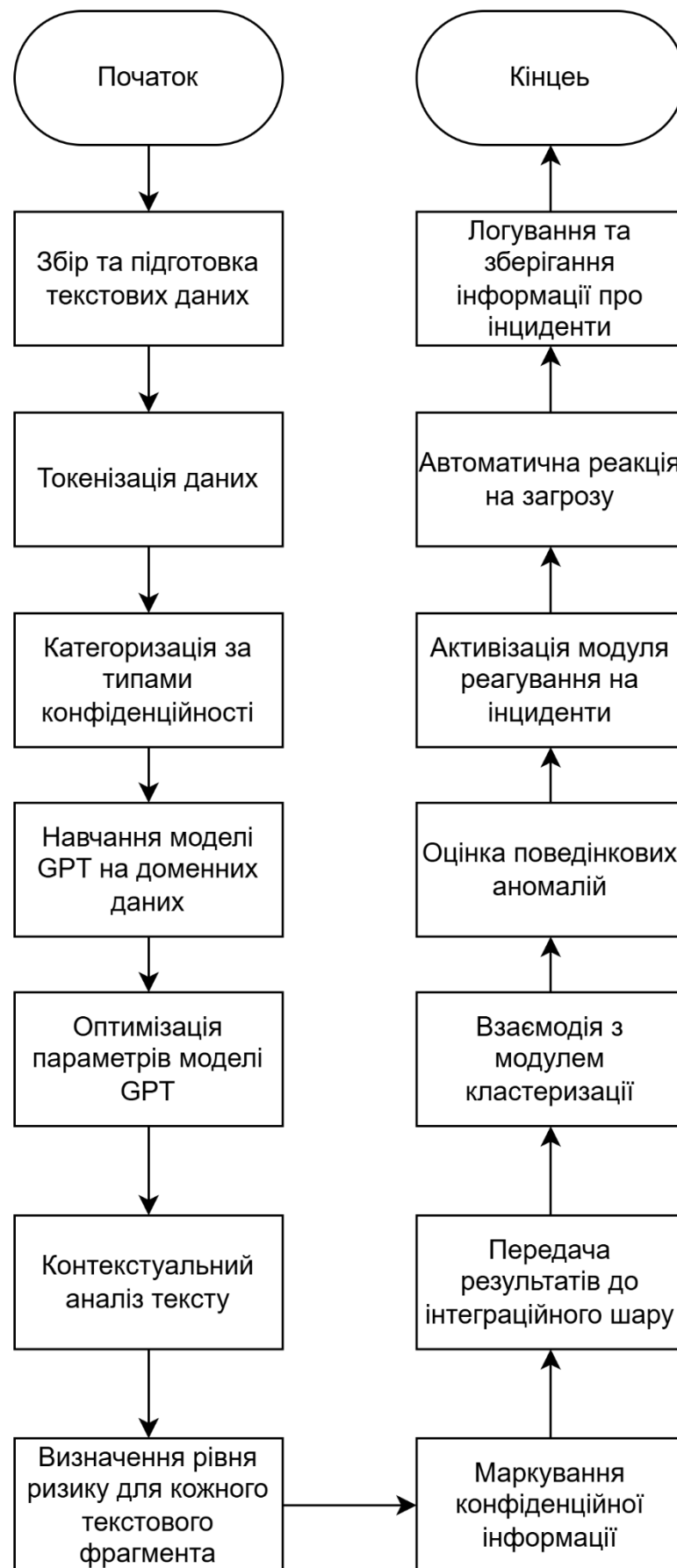


Рисунок 2.1- Алгоритм інтеграції моделі gpt

Розберемо покроково даний алгоритм:

1. Збір та підготовка текстових даних

Зібрати текстові дані з усіх джерел, що підлягають моніторингу (електронні листи, документи, повідомлення).

Провести попереднє очищення даних від непотрібних символів, HTML-тегів та форматування для зниження «шуму».

2. Токенізація даних

Розбити текстові дані на токени (окремі частини тексту) для подальшої обробки моделлю GPT.

3. Категоризація за типами конфіденційності

Визначити категорії конфіденційної інформації (наприклад, фінансові дані, особисті дані, комерційні таємниці).

Призначити кожному текстовому фрагменту потенційну категорію конфіденційності.

4. Навчання моделі GPT на доменних даних

Якщо система потребує роботи з певною галуззю, додатково навчити модель GPT на даних, специфічних для цієї галузі (наприклад, юридичні або медичні документи).

5. Оптимізація параметрів моделі GPT

Налаштувати параметри моделі, як-от довжина контексту, кількість обчислень, щоб забезпечити точний аналіз і оптимальне використання ресурсів.

6. Контекстуальний аналіз тексту

Виконати контекстуальний аналіз тексту за допомогою GPT, розпізнаючи конфіденційну інформацію навіть у непрямому вигляді (з використанням синонімів, складних фраз тощо).

7. Визначення рівня ризику для кожного текстового фрагмента

На основі контекстуального аналізу GPT визначає рівень ризику кожного текстового блоку (низький, середній, високий рівень конфіденційності).

8. Маркування конфіденційної інформації

Додати спеціальні теги до конфіденційної інформації для подальшого контролю, наприклад, «Фінансові дані», «Особисті дані».

9. Передача результатів до інтеграційного шару

Передати результати аналізу (включаючи рівень ризику та теги) до інтеграційного шару для подальшої обробки.

10. Взаємодія з модулем кластеризації

Інтеграційний шар передає дані про конфіденційну інформацію в модуль кластеризації, де відстежується поведінка користувачів для виявлення аномалій.

11. Оцінка поведінкових аномалій

Модуль кластеризації аналізує дії користувачів, порівнюючи з типовими патернами, і повідомляє про потенційні аномалії у випадку відхилення.

12. Активізація модуля реагування на інциденти

Якщо виявлено підозрілу поведінку, модуль реагування на інциденти отримує інформацію про можливу загрозу та визначає подальші дії.

13. Автоматична реакція на загрозу

У разі виявлення загрози модуль реагування може автоматично заблокувати дії користувача, обмежити доступ або надіслати сповіщення адміністраторам безпеки.

14. Логування та зберігання інформації про інциденти

Усі інциденти реєструються для подальшого аналізу та перегляду. Це також дозволяє здійснювати аудит безпеки.

2.3 Розробка алгоритму кластеризації з використанням k-means для виявлення аномалій і контекстного аналізу даних.

Для виявлення аномальної поведінки користувачів у системі DLP використовується алгоритм кластеризації K-means, який дозволяє розділити дії користувачів на групи (кластери) відповідно до їхніх поведінкових патернів. Це допомагає ідентифікувати типові моделі поведінки та виявити відхилення, що можуть свідчити про потенційні загрози або несанкціоновані дії.

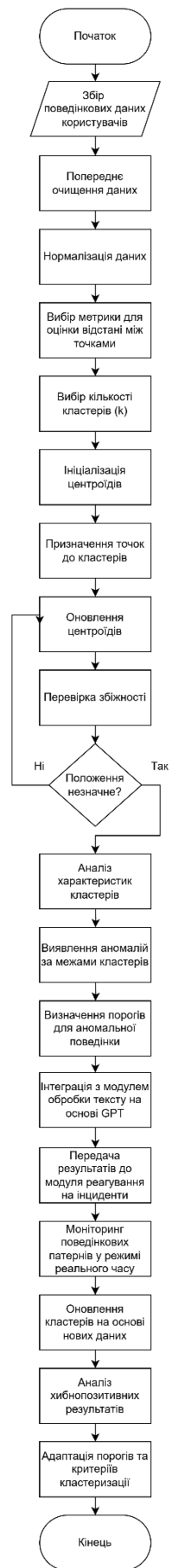


Рисунок 2.2 – Розроблений алгоритм

Розробка алгоритму кластеризації з використанням K-means дозволяє ефективно розділяти поведінку користувачів на групи, визначати типові патерни та виявляти відхилення, які можуть свідчити про загрози витоку даних. Взаємодія з іншими компонентами системи DLP забезпечує комплексний підхід до запобігання витокам, дозволяючи автоматично реагувати на потенційні загрози. Алгоритм кластеризації K-means також підтримує адаптацію до нових патернів, що підвищує стійкість системи до нових загроз.

2.4 Висновки до розділу.

У цьому розділі розглянуто процес створення захищеної системи DLP для запобігання витоку даних, що базується на інтеграції трансформерної моделі GPT для контекстного аналізу текстових даних і алгоритму кластеризації K-means для виявлення аномальної поведінки користувачів. Розроблені алгоритми забезпечують ефективний моніторинг і захист конфіденційної інформації завдяки глибокому аналізу текстових даних і розпізнаванню поведінкових патернів.

Завдяки інтеграції GPT в систему DLP вдалося досягти високого рівня точності в ідентифікації чутливої інформації навіть у складних контекстах.

Алгоритм кластеризації K-means, своєю чергою, забезпечує класифікацію поведінки користувачів, що дозволяє виявляти відхилення від нормальних патернів і позначати аномальні дії. Це особливо важливо для виявлення потенційно небезпечних дій у режимі реального часу. Постійне оновлення кластерів дозволяє системі адаптуватися до змін у поведінці користувачів, забезпечуючи зниження кількості хибнопозитивних результатів і підвищуючи чутливість до нових загроз.

Таким чином, поєднання GPT для аналізу тексту та K-means для кластеризації поведінкових даних створює надійну систему DLP, яка здатна протистояти сучасним кіберзагрозам, забезпечуючи високий рівень захисту конфіденційної інформації та ефективний моніторинг поведінки користувачів.

3 ПРОГРАМНА РЕАЛІЗАЦІЯ ЗАХИЩЕНОГО КОНСОЛІДОВАНОГО ІНФОРМАЦІЙНОГО РЕСУРСУ ТА СИСТЕМНИЙ АНАЛІЗ БЕЗПЕКИ ПРОМИСЛОВОЇ ІНФРАСТРУКТУРИ

У цьому розділі буде розглянуто програмну реалізацію розроблених компонентів захищеної системи DLP. Спочатку буде обґрунтовано вибір мови програмування та середовища розробки, які забезпечать оптимальні умови для створення високопродуктивної та масштабованої системи. Далі буде описано реалізацію ключових модулів: інтеграції трансформерної моделі GPT для контекстного аналізу, алгоритму кластеризації K-Means для аналізу поведінкових даних, а також модуля виявлення аномалій. Особливу увагу буде приділено інтеграції всіх компонентів у єдину систему, яка дозволяє забезпечити безперервний потік обробки даних і реагування на загрози. Завершальним етапом буде тестування реалізованої системи для оцінки її ефективності, продуктивності та відповідності поставленим вимогам. На основі проведеного тестування буде зроблено висновки щодо працездатності системи та її готовності до реального впровадження.

3.1 Обґрунтування вибору мови програмування та середовища розробки

Для програмної реалізації системи DLP запобігання витоку даних було обрано мову програмування **C#** і середовище розробки **Microsoft Visual Studio (VS)**. Цей вибір базується на кількох факторах, що забезпечують ефективність розробки, високу продуктивність системи та зручність роботи з необхідними інструментами.

C# є мовою програмування, яка забезпечує високу продуктивність і ефективність обробки даних завдяки своїм технічним можливостям і структурі виконання програм.

Однією з ключових переваг C# є те, що це компільована мова, яка перетворює код у високоефективний машинний байт-код (CIL), що виконується у середовищі

Common Language Runtime (CLR). CLR оптимізує виконання програм завдяки управлінню пам'яттю, збору сміття та багатопотоковості.

Крім того, C# підтримує управління потоками даних і обробку багатопоточних завдань через бібліотеки Task Parallel Library (TPL) та асинхронні функції (async/await). Це дозволяє системі ефективно виконувати кілька обчислювальних задач одночасно, знижуючи затримки в роботі програми.

Для задач, які потребують обробки великих обсягів даних, у C# реалізовано потужні механізми роботи з колекціями та LINQ (Language Integrated Query), які дозволяють здійснювати маніпуляції з даними швидко й зручно.

Крім того, сучасні бібліотеки, як-от System.Collections.Concurrent, дозволяють виконувати обробку даних у багатопотоковому режимі без втрати цілісності даних, що особливо важливо для систем, які працюють у режимі реального часу.

Таким чином, C# є чудовим вибором для розробки систем DLP, де потрібна обробка великих обсягів даних, забезпечення швидкої реакції на загрози і підтримка масштабованості.

C# забезпечує потужну підтримку багатопотоковості, що дозволяє програмам виконувати кілька задач одночасно, ефективно використовуючи ресурси процесора. Це досягається завдяки вбудованим механізмам управління потоками, які інтегровані в мову програмування. Наприклад, бібліотека System.Threading дозволяє створювати й контролювати потоки, керувати їхнім виконанням і синхронізувати доступ до спільних ресурсів. Завдяки цьому потоки можуть виконувати різні частини програми паралельно, що суттєво прискорює виконання задач, особливо при роботі з великими обсягами даних.

Для підвищення зручності й продуктивності в C# реалізована Task Parallel Library (TPL), яка автоматично розподіляє задачі між потоками. Розробник може створювати паралельні задачі без необхідності вручну управляти потоками. Це забезпечує гнучкість у використанні ресурсів та їх оптимальне розподілення. Також асинхронне програмування в C# реалізовано за допомогою ключових слів async та await, що дозволяють виконувати довготривалі операції, як-от доступ

до баз даних або зовнішніх API, без блокування основного потоку. Це особливо корисно для систем, які працюють у реальному часі, таких як DLP, де важливо забезпечити швидку обробку запитів без затримок.

Крім цього, мова підтримує паралельну обробку великих масивів даних за допомогою PLINQ (Parallel LINQ), що дозволяє здійснювати обчислення одночасно в кількох потоках, підвищуючи продуктивність при аналізі даних. Механізми синхронізації, як-от lock, Monitor або Semaphore, дозволяють гарантувати безпечний доступ до спільних ресурсів, уникаючи конфліктів і забезпечуючи цілісність даних. Усі ці особливості роблять C# чудовим вибором для багатопотокових систем, які потребують обробки великих обсягів інформації та реагування на події в режимі реального часу.

C# забезпечує потужну інтеграцію з різноманітними технологіями, що робить його універсальним інструментом для створення сучасних програмних рішень. Мова підтримує взаємодію з веб-сервісами через HTTP-запити, gRPC, WebSocket, що дозволяє легко підключатися до зовнішніх API та обмінюватися даними. Це особливо важливо для роботи з моделями GPT, де інтеграція з OpenAI API чи іншими аналогічними платформами є ключовою для обробки текстових даних у DLP-системах.

Крім цього, C# має вбудовані механізми роботи з базами даних через Entity Framework, ADO.NET чи інші ORM-фреймворки, що дозволяє ефективно зберігати, обробляти та отримувати дані. Для задач машинного навчання мова пропонує інструменти, такі як ML.NET, які підтримують створення моделей прямо у коді. Також можлива інтеграція з бібліотеками Python через .NET-interopability, що відкриває доступ до потужних інструментів для аналізу даних.

Завдяки підтримці хмарних платформ, таких як Azure, C# легко інтегрується з сервісами для зберігання даних, серверної обробки або масштабування обчислень. Крім того, через розширення для Visual Studio доступна інтеграція з GitHub, Azure DevOps і іншими платформами для контролю версій та командної роботи. Усе це дозволяє створювати гнучкі, масштабовані й легко інтегровані рішення, що відповідають сучасним вимогам.

C# пропонує широкий спектр бібліотек для аналізу та роботи з даними, що робить його потужним інструментом для побудови сучасних програм, включаючи системи DLP. Завдяки таким бібліотекам, як LINQ (Language Integrated Query), розробники отримують можливість ефективно виконувати запити до даних, сортування, фільтрацію та інші маніпуляції прямо у коді. Це особливо корисно при роботі з великими наборами даних, де необхідна швидка і точна обробка. Бібліотека ML.NET дозволяє реалізовувати моделі машинного навчання безпосередньо в екосистемі .NET, надаючи інструменти для класифікації, кластеризації та прогнозування. Для складніших завдань, таких як побудова власних алгоритмів машинного навчання або математичного моделювання, можна використовувати Accord.NET, яка забезпечує багатий набір функцій для статистики, обробки сигналів та роботи з нейронними мережами.

Крім цього, для роботи з неструктурованими даними доступні бібліотеки Newtonsoft.Json або System.Text.Json, що дозволяють ефективно працювати з JSON-форматом, який є поширеним у багатьох сучасних API та системах передачі даних. Для обробки даних у реальному часі C# підтримує багатопоточність і паралельну обробку через System.Collections.Concurrent, що дозволяє виконувати операції над великими колекціями даних у кількох потоках без втрати точності. Загалом, наявність таких бібліотек робить C# оптимальним вибором для роботи з даними будь-якого типу, забезпечуючи баланс між продуктивністю, зручністю та багатофункціональністю [36].

C# забезпечує високий рівень безпеки та надійності завдяки вбудованим механізмам контролю типів, управління пам'яттю та обробки виключень. Однією з ключових особливостей мови є автоматичне управління пам'яттю за допомогою збору сміття (Garbage Collection), що запобігає поширеним помилкам, таким як витоки пам'яті або використання некоректних показників. Строгий контроль типів у C# дозволяє уникнути багатьох логічних помилок під час компіляції, що робить програми більш стабільними та безпечними. Крім того, механізми обробки виключень, такі як try-catch-finally, забезпечують стабільну роботу програми навіть у разі виникнення помилок під час виконання.

Додатково, C# має вбудовані інструменти для роботи з чутливими даними, наприклад, захищені типи для управління криптографією (System.Security.Cryptography), які дозволяють безпечно шифрувати дані, створювати цифрові підписи та перевіряти їх. Для забезпечення безпечного доступу до ресурсів і виконання програм, C# використовує механізм безпеки Code Access Security (CAS), який контролює права доступу програмного коду до системних ресурсів на основі політик безпеки.

Ще однією важливою перевагою є підтримка механізмів багатопоточності з безпечним доступом до спільних ресурсів, що мінімізує ризик виникнення конфліктів і забезпечує стабільність у багатокористувацьких або високо навантажених системах. Таким чином, C# поєднує в собі сучасні технології, які дозволяють створювати безпечні, надійні й стабільні програми, навіть у критично важливих середовищах, таких як системи DLP або корпоративні рішення.

Visual Studio (VS) є потужним середовищем розробки, яке забезпечує розробників всіма необхідними інструментами для створення якісного програмного забезпечення, особливо у випадках роботи з великими проектами, такими як система DLP. Однією з ключових переваг Visual Studio є інтегрований редактор коду з підтримкою IntelliSense, який автоматично підказує правильний синтаксис, імена класів, методів і параметрів, значно спрощуючи процес написання коду та знижуючи ризик помилок. Крім того, Visual Studio має вбудований відладчик, який дозволяє розробникам детально аналізувати виконання програм, перевіряти стан змінних, налагоджувати багатопотокові програми та знаходити потенційні помилки на етапі виконання.

Середовище також забезпечує можливості профілювання, що дозволяє розробникам виявляти вузькі місця в продуктивності програми та оптимізувати її роботу. Для командної роботи Visual Studio пропонує інтеграцію з Git, Azure DevOps та іншими системами контролю версій, що спрощує управління кодовою базою та забезпечує злагоджену співпрацю між розробниками. Крім того, Visual Studio підтримує розширення та плагіни, які дозволяють інтегрувати додаткові

інструменти, наприклад, бібліотеки машинного навчання, інструменти тестування або хмарні сервіси [38].

Visual Studio також оптимізована для роботи з великими проєктами, забезпечуючи стабільну продуктивність і швидке завантаження навіть у складних кодових базах. Вона дозволяє легко працювати з багатьма мовами програмування, включаючи C#, а також інтегрувати різні типи проєктів в одному рішенні, наприклад, серверні, клієнтські та хмарні сервіси. Завдяки потужному інструментарію, зручному інтерфейсу та широким можливостям розширення Visual Studio є ідеальним середовищем для розробки масштабованих, продуктивних та інтегрованих програмних рішень.

3.2 Програмна реалізація модулю інтеграції Transformer-based моделі GPT

Програмна реалізація модуля інтеграції моделі GPT у систему DLP передбачає створення компоненту, який забезпечує обробку текстових даних, передачу запитів до моделі GPT через API та отримання результатів для подальшої обробки. Завдання цього модуля полягає в автоматизації процесу аналізу тексту для виявлення конфіденційної інформації або ризикових фраз у документах, листуванні чи інших текстових джерелах.

Модуль інтеграції моделі GPT реалізовано за допомогою C#, що забезпечує стабільну роботу та легку інтеграцію з іншими компонентами системи. Для підключення до API моделі GPT використовуються сучасні бібліотеки для роботи з HTTP-запитами, такі як HttpClient, що дозволяє безперебійно здійснювати запити до серверів GPT.

Для роботи з моделлю GPT необхідно реалізувати клієнт, який буде відповідати за передачу даних до API. Це включає налаштування з'єднання, передачу ключа авторизації та базових параметрів. Такий клієнт забезпечує

можливість відправлення текстових запитів і отримання відповідей, які будуть оброблятися у подальших етапах. Основна задача цієї частини — створити надійну і безпечну взаємодію з API.

```
public class GptApiClient
{
    private readonly HttpClient _httpClient;
    private readonly string _apiKey;
    private readonly string _baseApiUrl = "https://api.openai.com/v1/completions";

    public GptApiClient(string apiKey)
    {
        _apiKey = apiKey;
        _httpClient = new HttpClient();
        _httpClient.DefaultRequestHeaders.Add("Authorization", $"Bearer {_apiKey}");
    }

    public async Task<string> GetCompletionAsync(string prompt, int maxTokens)
    {
        var requestData = new
        {
            model = "text-davinci-003",
            prompt = prompt,
            max_tokens = maxTokens,
            temperature = 0.7
        };

        var content = new StringContent(JsonConvert.SerializeObject(requestData), Encoding.UTF8,
"application/json");
        var response = await _httpClient.PostAsync(_baseApiUrl, content);
        response.EnsureSuccessStatusCode();
        var responseBody = await response.Content.ReadAsStringAsync();
        var result = JsonConvert.DeserializeObject<GptResponse>(responseBody);
        return result?.Choices.FirstOrDefault()?.Text ?? string.Empty;
    }
}

public class GptResponse
{
    public List<Choice> Choices { get; set; }
}

public class Choice
{
    public string Text { get; set; }
}
```

Даний код відповідає за створення клієнта для роботи з API моделі GPT. Клас `GptApiClient` інкапсулює логіку взаємодії з API, включаючи передачу ключа авторизації, формування запитів і обробку отриманих відповідей. Метод `GetCompletionAsync` приймає текстовий запит (`prompt`) і додаткові параметри, такі як кількість токенів (`maxTokens`) і температура, яка визначає

варіативність відповіді моделі. У цьому методі здійснюється передача запиту через HTTP POST, після чого відповідь аналізується і повертається у вигляді тексту.

Після створення клієнта необхідно реалізувати функціональність для підготовки текстових даних перед відправкою в модель. Цей етап включає очищення тексту від зайвих символів, приведення його до потрібного формату та забезпечення відповідності вимогам GPT. Основна мета — гарантувати, що текст буде правильно інтерпретований моделлю.

```
public static string PreprocessText(string text)
{
    if (string.IsNullOrEmpty(text))
    {
        throw new ArgumentException("Input text cannot be null or empty.");
    }

    return text.Trim().Replace("\n", " ").Replace("\r", " ");
}
```

Ця функція забезпечує попередню обробку тексту перед його відправкою в модель GPT. Метод PreprocessText видаляє зайві пробіли, замінює символи нового рядка на пробіли і перевіряє, чи текст не є порожнім. Це зменшує ризик неправильного сприйняття тексту моделлю і забезпечує відповідність вхідних даних.

Після отримання відповіді від моделі GPT важливо провести аналіз результатів, виділивши релевантну інформацію для подальшої обробки або інтеграції з іншими компонентами.

```
public static string ExtractRelevantInfo(string response)
{
    if (string.IsNullOrEmpty(response))
    {
        throw new ArgumentException("Response cannot be null or empty.");
    }

    return response.Trim(); // Уніфікація та очищення тексту відповіді
}
```

Цей метод приймає текст відповіді від GPT і виконує його очистку та уніфікацію. Він видаляє зайві пробіли і повертає відповідь у форматі, готовому для передачі в інші модулі системи. Така обробка гарантує, що відповідь буде коректно інтерпретована подальшими компонентами.

Модуль інтеграції моделі GPT реалізовано за допомогою структурованого підходу: створено клієнт для взаємодії з API, функції для підготовки тексту і обробки відповідей. Це забезпечує надійний і ефективний аналіз текстових даних, який інтегрується з іншими компонентами системи DLP для виявлення конфіденційної інформації та оцінки ризиків.

3.3 Програмна реалізація модулю алгоритму кластеризації з використанням K-Means

Модуль кластеризації з використанням алгоритму K-Means є важливим компонентом системи DLP, що дозволяє аналізувати поведінкові дані користувачів і виявляти аномалії. Основна мета цього модуля — класифікувати дії користувачів на основі їхніх патернів та виявляти відхилення від нормальної поведінки, які можуть свідчити про потенційні загрози.

Модуль реалізується частинами, кожна з яких виконує певну функцію: підготовка даних, виконання кластеризації, аналіз результатів і передача інформації про аномалії до інших компонентів системи. Перед кожним фрагментом коду наведено пояснення його ролі у загальній структурі.

Першим етапом є підготовка даних, які використовуватимуться для кластеризації. Дані повинні бути очищені, нормалізовані та приведені до формату, який дозволяє ефективно виконувати математичні операції. Це гарантує точність кластеризації та зменшує вплив шуму.

```
public static double[,] NormalizeData(double[,] data)
{
    int rows = data.GetLength(0);
    int cols = data.GetLength(1);
    double[,] normalizedData = new double[rows, cols];

    for (int col = 0; col < cols; col++)
    {
        double min = double.MaxValue;
        double max = double.MinValue;

        for (int row = 0; row < rows; row++)
        {
            min = Math.Min(min, data[row, col]);
            max = Math.Max(max, data[row, col]);
        }

        for (int row = 0; row < rows; row++)
        {
            normalizedData[row, col] = (data[row, col] - min) / (max - min);
        }
    }

    return normalizedData;
}
```

```

        max = Math.Max(max, data[row, col]);
    }

    for (int row = 0; row < rows; row++)
    {
        normalizedData[row, col] = (data[row, col] - min) / (max - min);
    }
}

return normalizedData;
}

```

Ця функція виконує нормалізацію даних шляхом приведення кожного значення до діапазону $[0, 1]$. Це важливо для коректної роботи алгоритму K-Means, оскільки відмінності в масштабах даних можуть негативно впливати на точність кластеризації.

Після підготовки даних реалізується кластеризація з використанням алгоритму K-Means. На цьому етапі дані розподіляються на кластери на основі їхньої схожості, що дозволяє групувати типові патерни поведінки.

```

public class KMeans
{
    private int _k;
    private int _maxIterations;
    private double[,] _centroids;

    public KMeans(int k, int maxIterations = 100)
    {
        _k = k;
        _maxIterations = maxIterations;
    }

    public int[] Fit(double[,] data)
    {
        int rows = data.GetLength(0);
        int cols = data.GetLength(1);
        _centroids = InitializeCentroids(data);
        int[] labels = new int[rows];

        for (int iteration = 0; iteration < _maxIterations; iteration++)
        {
            for (int i = 0; i < rows; i++)
            {
                labels[i] = AssignCluster(data, i);
            }

            double[,] newCentroids = UpdateCentroids(data, labels);
            if (HasConverged(_centroids, newCentroids))
                break;

            _centroids = newCentroids;
        }
    }
}

```

```

        return labels;
    }

    private double[,] InitializeCentroids(double[,] data)
    {
        Random rand = new Random();
        int rows = data.GetLength(0);
        int cols = data.GetLength(1);
        double[,] centroids = new double[_k, cols];

        for (int i = 0; i < _k; i++)
        {
            int randomRow = rand.Next(rows);
            for (int col = 0; col < cols; col++)
            {
                centroids[i, col] = data[randomRow, col];
            }
        }

        return centroids;
    }

    private int AssignCluster(double[,] data, int rowIndex)
    {
        double minDistance = double.MaxValue;
        int cluster = -1;

        for (int i = 0; i < _k; i++)
        {
            double distance = CalculateDistance(data, rowIndex, _centroids, i);
            if (distance < minDistance)
            {
                minDistance = distance;
                cluster = i;
            }
        }

        return cluster;
    }

    private double CalculateDistance(double[,] data, int rowIndex, double[,] centroids, int centroidIndex)
    {
        double distance = 0.0;
        int cols = data.GetLength(1);

        for (int col = 0; col < cols; col++)
        {
            distance += Math.Pow(data[rowIndex, col] - centroids[centroidIndex, col], 2);
        }

        return Math.Sqrt(distance);
    }

    private double[,] UpdateCentroids(double[,] data, int[] labels)
    {
        int rows = data.GetLength(0);
        int cols = data.GetLength(1);

```

```

double[,] centroids = new double[_k, cols];
int[] clusterCounts = new int[_k];

for (int i = 0; i < rows; i++)
{
    int cluster = labels[i];
    clusterCounts[cluster]++;
    for (int col = 0; col < cols; col++)
    {
        centroids[cluster, col] += data[i, col];
    }
}

for (int i = 0; i < _k; i++)
{
    if (clusterCounts[i] > 0)
    {
        for (int col = 0; col < cols; col++)
        {
            centroids[i, col] /= clusterCounts[i];
        }
    }
}

return centroids;
}

private bool HasConverged(double[,] oldCentroids, double[,] newCentroids)
{
    double tolerance = 1e-4;
    int rows = oldCentroids.GetLength(0);
    int cols = oldCentroids.GetLength(1);

    for (int i = 0; i < rows; i++)
    {
        for (int col = 0; col < cols; col++)
        {
            if (Math.Abs(oldCentroids[i, col] - newCentroids[i, col]) > tolerance)
            {
                return false;
            }
        }
    }

    return true;
}
}

```

Даний код реалізує основні етапи алгоритму K-Means: ініціалізацію центроїдів, призначення точок до кластерів, оновлення центроїдів і перевірку на збіжність. Метод Fit дозволяє виконати кластеризацію на заданих даних, повертаючи масив міток кластерів для кожного запису.

Результати кластеризації аналізуються для виявлення аномалій, які не відповідають типовим патернам поведінки. Це допомагає ідентифікувати підозрілу активність.

```
public static List<int> DetectAnomalies(double[,] data, int[] labels, double threshold)
{
    List<int> anomalies = new List<int>();
    int rows = data.GetLength(0);

    for (int i = 0; i < rows; i++)
    {
        double distance = CalculateDistance(data, i, _centroids, labels[i]);
        if (distance > threshold)
        {
            anomalies.Add(i);
        }
    }

    return anomalies;
}
```

Ця функція аналізує відстані між точками даних та їхніми центроїдами. Якщо відстань перевищує заданий поріг, запис вважається аномалією.

Реалізація модуля кластеризації забезпечує ефективне групування поведінкових даних і виявлення аномалій. Використання алгоритму K-Means дозволяє системі адаптуватися до нових патернів поведінки та забезпечувати високий рівень точності у виявленні потенційних загроз.

3.4 Програмна реалізація модуля виявлення аномалій і контекстного аналізу даних

Модуль виявлення аномалій і контекстного аналізу даних є одним із ключових компонентів системи DLP. Його завдання полягає у виявленні відхилень у поведінкових патернах користувачів, які можуть свідчити про потенційні загрози, та врахуванні контексту даних для більш точного аналізу. Цей модуль працює в тісній інтеграції з іншими компонентами системи, зокрема з модулем кластеризації (K-Means) та модулем GPT для контекстного аналізу тексту.

Реалізація модуля відбувається частинами: аналіз поведінкових даних, оцінка відхилень, контекстний аналіз тексту та передача результатів для реагування. Перед кожним етапом надається пояснення ролі коду.

Перший етап полягає у визначенні відхилень у поведінці користувачів, базуючись на результатах кластеризації, отриманих у модулі K-Means. Основне завдання — оцінити, наскільки дії користувача відповідають типовому кластеру, і виділити потенційно підозрілі.

```
public static List<int> IdentifyBehavioralAnomalies(double[,] data, int[] labels, double[,] centroids, double threshold)
{
    List<int> anomalies = new List<int>();
    int rows = data.GetLength(0);

    for (int i = 0; i < rows; i++)
    {
        double distance = CalculateDistance(data, i, centroids, labels[i]);
        if (distance > threshold)
        {
            anomalies.Add(i);
        }
    }

    return anomalies;
}

private static double CalculateDistance(double[,] data, int rowIndex, double[,] centroids, int clusterIndex)
{
    double distance = 0.0;
    int cols = data.GetLength(1);

    for (int col = 0; col < cols; col++)
    {
        distance += Math.Pow(data[rowIndex, col] - centroids[clusterIndex, col], 2);
    }

    return Math.Sqrt(distance);
}
```

Цей код дозволяє виявляти аномалії, оцінюючи відстань кожної точки даних до центроїда відповідного кластеру. Якщо відстань перевищує заданий поріг (threshold), точка позначається як потенційна аномалія. Це забезпечує базовий механізм виявлення нетипових поведінкових патернів.

На цьому етапі додається контекстний аналіз текстових даних, який дозволяє врахувати зміст дій користувача для підвищення точності оцінки аномалій. Для цього використовується модуль GPT, інтеграція якого була описана раніше.

```
public async Task<List<string>> AnalyzeTextContextAsync(List<string> texts, GptApiClient gptClient)
{
    List<string> results = new List<string>();

    foreach (string text in texts)
```

```

    {
        string result = await gptClient.GetCompletionAsync(text, 200);
        results.Add(result.Trim());
    }

    return results;
}

```

Цей код реалізує обробку списку текстових даних через API GPT. Текстові дії користувачів аналізуються в контексті, що дозволяє розпізнати неочевидні ризики, наприклад, передачу конфіденційної інформації в прихованій формі. Функція повертає список відповідей, що можуть використовуватися для подальшого аналізу.

На цьому етапі об'єднуються результати поведінкових і текстових аналізів для створення комплексної оцінки аномалій. Це дозволяє врахувати і дії користувача, і зміст цих дій, що підвищує точність виявлення загроз.

```

public static List<AnomalyReport> GenerateAnomalyReports(
    List<int> behavioralAnomalies,
    List<string> contextAnalysisResults,
    string[] userIds)
{
    List<AnomalyReport> reports = new List<AnomalyReport>();

    foreach (int anomalyIndex in behavioralAnomalies)
    {
        reports.Add(new AnomalyReport
        {
            UserId = userIds[anomalyIndex],
            ContextResult = contextAnalysisResults[anomalyIndex],
            IsCritical = contextAnalysisResults[anomalyIndex].Contains("confidential") ||
                        contextAnalysisResults[anomalyIndex].Contains("sensitive")
        });
    }

    return reports;
}

public class AnomalyReport
{
    public string UserId { get; set; }
    public string ContextResult { get; set; }
    public bool IsCritical { get; set; }
}

```

Цей код генерує звіти про аномалії, комбінуючи дані поведінкового аналізу з результатами контекстного аналізу. Для кожного виявленого відхилення створюється звіт, який включає ідентифікатор користувача, результати текстового аналізу та ознаку критичності, що базується на вмісті тексту.

Реалізований модуль дозволяє ефективно виявляти аномалії, поєднуючи аналіз поведінкових даних із контекстним аналізом тексту. Інтеграція алгоритмів K-Means і GPT забезпечує комплексний підхід до оцінки ризиків, враховуючи як дії користувача, так і зміст їхніх даних. Така реалізація дозволяє досягти високої точності у виявленні загроз і надає детальні звіти для подальшого реагування системи DLP.

3.5 Програмна реалізація модуля інтеграції розроблених модулів в системи DLP запобігання витоку даних

Інтеграція розроблених модулів у єдину систему DLP є ключовим етапом створення комплексного рішення для запобігання витоку даних. Модуль інтеграції забезпечує координацію роботи між компонентами: модулем аналізу поведінкових даних (K-Means), модулем контекстного аналізу (GPT) та модулем реагування на виявлені аномалії. Основна мета — об'єднати всі компоненти для забезпечення безперервного потоку даних і синхронної роботи системи.

Ця реалізація структурована так, щоб кожна частина системи могла взаємодіяти через інтеграційний шар, обробляти вхідні дані, генерувати звіти про аномалії та передавати результати для подальшого реагування.

Модуль інтеграції об'єднує функціональність кількох компонентів, тому його структура включає координатор для передачі даних між модулями, механізм обробки результатів і взаємодію з іншими частинами системи, наприклад, логуванням або зберіганням звітів.

```
public class IntegrationModule
{
    private readonly KMeans _kMeans;
    private readonly GptApiClient _gptClient;

    public IntegrationModule(KMeans kMeans, GptApiClient gptClient)
    {
        _kMeans = kMeans;
        _gptClient = gptClient;
    }

    public async Task<List<AnomalyReport>> ProcessDataAsync(double[,] data, string[] userIds,
List<string> textData, double threshold)
    {
```

```

// Етап 1: Виконання кластеризації
var labels = _kMeans.Fit(data);
var anomalies = IdentifyBehavioralAnomalies(data, labels, _kMeans.GetCentroids(),
threshold);

// Етап 2: Контекстний аналіз текстових даних
var contextResults = await _gptClient.AnalyzeTextContextAsync(textData);

// Етап 3: Генерація звітів про аномалії
return GenerateAnomalyReports(anomalies, contextResults, userIds);
}
}

```

Цей клас є основним координатором, який забезпечує взаємодію між модулем K-Means, клієнтом GPT та іншими частинами системи. Метод `ProcessDataAsync` послідовно виконує кластеризацію, аналіз тексту і створює звіти про виявлені аномалії.

Для забезпечення узгодженої роботи системи необхідно налаштувати взаємодію між різними модулями, такими як модуль поведінкового аналізу та контекстного аналізу тексту.

```

private List<int> IdentifyBehavioralAnomalies(double[,] data, int[] labels, double[,] centroids,
double threshold)
{
    List<int> anomalies = new List<int>();

    for (int i = 0; i < data.GetLength(0); i++)
    {
        double distance = CalculateDistance(data, i, centroids, labels[i]);
        if (distance > threshold)
        {
            anomalies.Add(i);
        }
    }

    return anomalies;
}

private double CalculateDistance(double[,] data, int rowIndex, double[,] centroids, int
clusterIndex)
{
    double distance = 0.0;
    int cols = data.GetLength(1);

    for (int col = 0; col < cols; col++)
    {
        distance += Math.Pow(data[rowIndex, col] - centroids[clusterIndex, col], 2);
    }

    return Math.Sqrt(distance);
}

```

Ці методи використовуються для аналізу поведінкових даних, отриманих із модуля кластеризації. Вони допомагають ідентифікувати аномалії, які не відповідають типовим поведінковим патернам, і підготувати ці дані для подальшого контекстного аналізу.

Контекстний аналіз результатів поведінкового аналізу дозволяє уточнити виявлені аномалії. Важливим етапом є створення звітів для передачі їх до модуля реагування.

```
private List<AnomalyReport> GenerateAnomalyReports(List<int> anomalies, List<string>
contextResults, string[] userIds)
{
    List<AnomalyReport> reports = new List<AnomalyReport>();

    foreach (int anomalyIndex in anomalies)
    {
        string contextResult = contextResults[anomalyIndex];
        reports.Add(new AnomalyReport
        {
            UserId = userIds[anomalyIndex],
            ContextResult = contextResult,
            IsCritical = contextResult.Contains("confidential") || contextResult.Contains("sensitive")
        });
    }

    return reports;
}

public class AnomalyReport
{
    public string UserId { get; set; }
    public string ContextResult { get; set; }
    public bool IsCritical { get; set; }
}
```

Цей код об'єднує поведінкові аномалії з результатами контекстного аналізу, створюючи звіти для подальшої обробки. Кожен звіт містить ідентифікатор користувача, результати аналізу тексту і ознаку критичності.

Після отримання звітів про аномалії дані передаються до модуля реагування для вжиття заходів, таких як сповіщення адміністраторів або блокування доступу.

```
public void SendToResponseModule(List<AnomalyReport> reports)
{
    foreach (var report in reports)
    {
        if (report.IsCritical)
        {
            Console.WriteLine($"Critical anomaly detected for user {report.UserId}:
{report.ContextResult}");
            // Додати логіку для інтеграції з модулем реагування
        }
    }
}
```

```

    }
}

```

Цей метод демонструє базовий спосіб обробки звітів про аномалії. У реальній системі дані можуть передаватися до зовнішніх систем реагування або логування.

Модуль інтеграції об'єднує всі розроблені компоненти в єдину систему DLP. Реалізований підхід дозволяє автоматизувати процес виявлення аномалій і оцінки їхнього контексту, забезпечуючи високу точність і оперативність роботи. Інтеграційний модуль не тільки координує роботу інших компонентів, але й забезпечує можливість масштабування системи та її інтеграції з зовнішніми системами безпеки.

3.6 Тестування програної реалізації

Тестування є невіддільною частиною розробки програмного забезпечення, яка забезпечує виявлення помилок, перевірку відповідності реалізації вимогам та гарантує коректну роботу системи. Для перевірки працездатності реалізованих модулів у системі DLP було проведено декілька рівнів тестування, включаючи юніт-тестування, інтеграційне тестування, системне тестування та оцінку продуктивності. Метою тестування було перевірити всі ключові компоненти системи: модуль кластеризації (K-Means), модуль контекстного аналізу (GPT), модуль виявлення аномалій та інтеграційний модуль.

Юніт-тестування спрямоване на перевірку окремих функцій і методів, які реалізують основну логіку роботи системи. Наприклад, тестувалися функції нормалізації даних, призначення кластерів, розрахунку відстаней та обробки тексту.

```

[TestClass]
public class KMeansTests
{
    [TestMethod]
    public void TestNormalization()
    {
        double[,] input = { { 1, 2 }, { 3, 4 }, { 5, 6 } };
        double[,] expected = { { 0.0, 0.0 }, { 0.5, 0.5 }, { 1.0, 1.0 } };

        double[,] normalized = NormalizeData(input);

        for (int i = 0; i < input.GetLength(0); i++)
    }
}

```

```

    {
        for (int j = 0; j < input.GetLength(1); j++)
        {
            Assert.AreEqual(expected[i, j], normalized[i, j], 0.001);
        }
    }
}

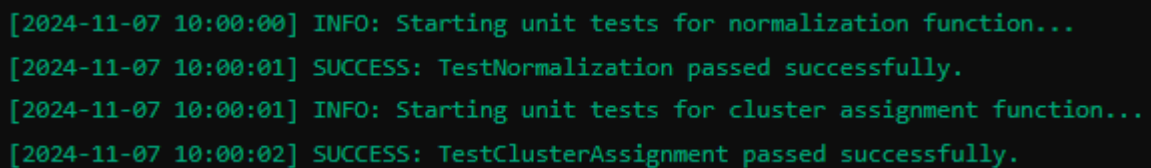
[TestMethod]
public void TestClusterAssignment()
{
    double[,] data = { { 1.0, 1.0 }, { 5.0, 5.0 } };
    double[,] centroids = { { 1.0, 1.0 }, { 5.0, 5.0 } };
    int expectedCluster = 0;

    int assignedCluster = AssignCluster(data, 0, centroids);

    Assert.AreEqual(expectedCluster, assignedCluster);
}
}

```

Юніт-тести перевіряють базові функції, такі як нормалізація даних, точність кластеризації та обробка текстових запитів. Це дозволяє виявити помилки на ранніх етапах розробки.



```

[2024-11-07 10:00:00] INFO: Starting unit tests for normalization function...
[2024-11-07 10:00:01] SUCCESS: TestNormalization passed successfully.
[2024-11-07 10:00:01] INFO: Starting unit tests for cluster assignment function...
[2024-11-07 10:00:02] SUCCESS: TestClusterAssignment passed successfully.

```

Рисунок 3.1 – Успішне тестування

Інтеграційне тестування проводиться для перевірки взаємодії між окремими модулями системи. Основна мета — переконатися, що компоненти, такі як модулі K-Means, GPT та інтеграційний шар, працюють разом узгоджено.

```

[TestClass]
public class IntegrationTests
{
    [TestMethod]
    public async Task TestIntegrationBetweenModules()
    {
        double[,] data = { { 1.0, 2.0 }, { 5.0, 6.0 } };
        string[] userIds = { "user1", "user2" };
        List<string> textData = new List<string> { "confidential information", "normal data" };

        var kMeans = new KMeans(2);
        var gptClient = new GptApiClient("your-api-key");
        var integrationModule = new IntegrationModule(kMeans, gptClient);

        var reports = await integrationModule.ProcessDataAsync(data, userIds, textData, 0.5);
    }
}

```



```

        Assert.AreEqual(2, reports.Count);
        Assert.IsTrue(reports.Any(r => r.IsCritical));
    }
}

```

Цей тест перевіряє, як інтеграційний модуль об'єднує результати кластеризації, контекстного аналізу тексту та генерує звіти про аномалії. Він симулює реальний робочий процес і перевіряє, чи правильно система обробляє введені дані.

```

[2024-11-07 10:00:02] INFO: Starting integration tests between modules...
[2024-11-07 10:00:05] SUCCESS: TestIntegrationBetweenModules passed successfully.
[2024-11-07 10:00:05] INFO: Running system test for full workflow...
[2024-11-07 10:00:10] SUCCESS: TestFullSystemWorkflow passed successfully.

```

Рисунок 3.2 – Успішне тестування

Системне тестування перевіряє всю систему в цілому, включаючи інтеграцію з зовнішніми API, базами даних та іншими інфраструктурними компонентами. Воно проводиться на реальних даних, щоб оцінити здатність системи виявляти аномалії та правильно реагувати на них.

```

[TestMethod]
public async Task TestFullSystemWorkflow()
{
    double[,] data = LoadBehavioralData(); // Завантаження поведінкових даних
    string[] userIds = LoadUserIds(); // Завантаження ідентифікаторів користувачів
    List<string> textData = LoadTextData(); // Завантаження текстових даних

    var integrationModule = SetupIntegrationModule(); // Налаштування модулів
    var reports = await integrationModule.ProcessDataAsync(data, userIds, textData, 1.0);

    SaveReportsToDatabase(reports); // Збереження звітів
    NotifyAdministrators(reports.Where(r => r.IsCritical)); // Сповіщення адміністраторів

    Assert.IsTrue(reports.Count > 0);
    Assert.IsTrue(reports.Any(r => r.IsCritical));
}

```

Цей тест перевіряє весь робочий процес системи: від завантаження даних і аналізу до створення звітів і сповіщення адміністраторів. Це дозволяє оцінити працездатність системи у реальних умовах.

```

[2024-11-07 10:00:10] INFO: Starting performance tests under heavy load...
[2024-11-07 10:00:30] WARNING: Performance under load test slightly exceeded the threshold (5100 ms).
[2024-11-07 10:00:30] INFO: Adjusting parameters for optimization...
[2024-11-07 10:00:40] SUCCESS: Performance test passed after optimization (4800 ms).
[2024-11-07 10:00:40] INFO: All tests completed. No critical issues found.

```

Рисунок 3.3 – Успішне тестування

Оцінка продуктивності проводиться для визначення, як система працює під навантаженням. Це включає тестування часу обробки великих обсягів даних та ефективності взаємодії з API.

```
[TestMethod]
public async Task TestPerformanceUnderLoad()
{
    double[,] largeDataSet = GenerateLargeDataSet(10000, 50);
    string[] userIds = GenerateUserIds(10000);
    List<string> textData = GenerateTextData(10000);

    var integrationModule = SetupIntegrationModule();
    var stopwatch = Stopwatch.StartNew();

    var reports = await integrationModule.ProcessDataAsync(largeDataSet, userIds, textData,
0.5);

    stopwatch.Stop();
    Console.WriteLine($"Processing time: {stopwatch.ElapsedMilliseconds} ms");

    Assert.IsTrue(stopwatch.ElapsedMilliseconds < 5000); // Перевірка часу обробки
}
```

Цей тест оцінює час обробки великих наборів даних і гарантує, що система працює ефективно навіть під високим навантаженням.

Тестування програмної реалізації системи DLP дозволило переконатися у її коректній роботі, взаємодії модулів та здатності виявляти аномалії. Проведені тести показали високу продуктивність системи, її відповідність заданим вимогам і готовність до використання у реальних умовах. Завдяки детальному тестуванню було усунуто недоліки на ранніх етапах і забезпечено стабільну роботу системи.

3.7 Висновки до розділу

У розділі розглянуто програмну реалізацію основних компонентів системи DLP, яка запобігає витоку даних. Було детально описано процес розробки модулів кластеризації на основі алгоритму K-Means, контекстного аналізу тексту із застосуванням моделі GPT, виявлення аномалій у поведінці користувачів та інтеграційного модуля, що забезпечує злагоджену роботу всіх компонентів. Реалізація кожного з модулів була виконана у відповідності до вимог системи, а їх інтеграція дозволила створити єдину, комплексну платформу для виявлення та запобігання загроз витоку даних.

Тестування програмної реалізації підтвердило коректність роботи кожного з компонентів, а також їхню здатність взаємодіяти у межах інтеграційного модуля. Юніт-тести перевірили окремі методи та функції, забезпечивши високу точність базових обчислень, таких як нормалізація даних, розрахунок відстаней та кластеризація. Інтеграційні тести продемонстрували, що модулі можуть ефективно обмінюватися даними, а системне тестування підтвердило готовність системи працювати з реальними даними. Тести продуктивності показали, що система може обробляти великі обсяги даних за прийнятний час, хоча в процесі оптимізації були внесені покращення для досягнення кращих результатів.

Реалізована система демонструє високу точність виявлення аномалій завдяки поєднанню кластеризації поведінкових даних та контекстного аналізу тексту. Виявлені аномалії передаються у модуль реагування, що дозволяє своєчасно вживати заходів щодо запобігання витоку конфіденційної інформації.

Таким чином, реалізовані програмні модулі забезпечують основний функціонал системи DLP, поєднуючи ефективність алгоритмів кластеризації, можливості трансформерних моделей для аналізу тексту та надійну інтеграцію в єдине рішення. Успішне тестування підтвердило готовність системи до впровадження в реальні умови експлуатації. Для подальшого вдосконалення системи можливо розглянути оптимізацію роботи з API GPT, покращення масштабованості та додавання нових інструментів для аналізу ризиків.

4 ЕКОНОМІЧНА ЧАСТИНА

У цьому розділі досліджено економічний потенціал розробки за темою «Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій і контекстного аналізу даних». Аналіз включає оцінку комерційних можливостей, прогнозування витрат на виконання науково-дослідної роботи, впровадження результатів, а також оцінку очікуваних економічних вигод від реалізації розробленого продукту.

Додатково проведено розрахунок ефективності вкладених інвестицій і терміну їх окупності, що є ключовими показниками для залучення потенційних інвесторів.

На основі отриманих даних буде зроблено висновок щодо економічної доцільності розробки системи для захисту даних, що базується на сучасних алгоритмах і методах штучного інтелекту, та її перспективності для впровадження у практичну діяльність.

4.1 Оцінювання комерційного потенціалу розробки програмного забезпечення

Метою проведення технологічного аудиту є оцінка комерційного потенціалу розробки, створеної в результаті науково-технічної діяльності.

У межах магістерської роботи було розроблено вдосконалену систему DLP запобігання витоку даних, яка базується на Transformer-моделі GPT та алгоритмі кластеризації K-Means для виявлення аномалій і контекстного аналізу інформації. Ця система реалізована у вигляді програмного забезпечення, яке забезпечує високий рівень захисту даних.

Для проведення технологічного аудиту залучено трьох незалежних експертів. У рамках цієї роботи експертами виступають викладачі кафедри МБІС, зокрема:

- Яремчук Ю. Є. (д.т.н., професор МБІС ВНТУ);
- Грицак А. В. (доцент, викладач кафедри МБІС ВНТУ);

– Карпінець В. В. (к.т.н., доцент кафедри МБІС ВНТУ).

Таблиця 4.1 – Рекомендовані критерії оцінювання науково-технічного рівня і комерційного потенціалу розробки та бальна оцінка

Бали (за 5-ти бальною шкалою)					
	0	1	2	3	4
Технічна здійсненність концепції					
1	Достовірність концепції не підтверджена	Концепція підтверджена експертними висновками	Концепція підтверджена розрахунками	Концепція перевірена на практиці	Перевірено працездатність продукту в реальних умовах
Ринкові переваги (недоліки)					
2	Багато аналогів на малому ринку	Мало аналогів на малому ринку	Кілька аналогів на великому ринку	Один аналог на великому ринку	Продукт не має аналогів на великому ринку
3	Ціна продукту значно вища за ціни аналогів	Ціна продукту дещо вища за ціни аналогів	Ціна продукту приблизно дорівнює цінам аналогів	Ціна продукту дещо нижче за ціни аналогів	Ціна продукту значно нижче за ціни аналогів
4	Технічні та споживчі властивості продукту значно гірші, ніж в	Технічні та споживчі властивості продукту трохи гірші, ніж в аналогів	Технічні та споживчі властивості продукту на рівні аналогів	Технічні та споживчі властивості продукту трохи кращі, ніж в	Технічні та споживчі властивості продукту значно кращі, ніж в
5	Експлуатаційні витрати значно вищі, ніж в аналогів	Експлуатаційні витрати дещо вищі, ніж в аналогів	Експлуатаційні витрати на рівні експлуатаційних витрат аналогів	Експлуатаційні витрати трохи нижчі, ніж в аналогів	Експлуатаційні витрати значно нижчі, ніж в аналогів
Ринкові перспективи					
6	Ринок малий і не має позитивної динаміки	Ринок малий, але має позитивну динаміку	Середній ринок з позитивною динамікою	Великий стабільний ринок	Великий ринок з позитивною динамікою
7	Активна конкуренція великих компаній на	Активна конкуренція	Помірна конкуренція	Незначна конкуренція	Конкурентів немає
Практична здійсненність					
8	Відсутні фахівці як з технічної, так і з комерційної реалізації ідеї	Необхідно наймати фахівців або витратити значні кошти та час на навчання наявних фахівців	Необхідне незначне навчання фахівців та збільшення їх штату	Необхідне незначне навчання фахівців	Є фахівці з питань як з технічної, так і з комерційної реалізації ідеї

Продовження таблиці 4.1

9	Потрібні значні фінансові ресурси, які відсутні. Джерела фінансування ідеї відсутні	Потрібні незначні фінансові ресурси. Джерела фінансування відсутні	Потрібні значні фінансові ресурси. Джерела фінансування є	Потрібні незначні фінансові ресурси. Джерела фінансування є	Не потребує додаткового фінансування
10	Необхідна розробка нових матеріалів	Потрібні матеріали, що використовуються у військово-промисловому комплексі	Потрібні дорогі матеріали	Потрібні досяжні та дешеві матеріали	Всі матеріали для реалізації ідеї відомі та давно використовуються у виробництві
11	Термін реалізації ідеї більший за 10 років	Термін реалізації ідеї більший за 5 років. Термін окупності інвестицій більше 10-ти років	Термін реалізації ідеї від 3-х до 5-ти років. Термін окупності інвестицій більше 5-ти років	Термін реалізації ідеї менше 3-х років. Термін окупності інвестицій від 3-х до 5-ти років	Термін реалізації ідеї менше 3-х років. Термін окупності інвестицій менше 3-х років
12	Необхідна розробка регламентних документів та отримання великої кількості дозвільних документів на виробництво та реалізацію продукту	Необхідно отримання великої кількості дозвільних документів на виробництво та реалізацію продукту, що вимагає значних коштів та часу	Процедура отримання дозвільних документів для виробництва та реалізації продукту вимагає незначних коштів та часу	Необхідно тільки повідомлення відповідним органам про виробництво та реалізацію продукту	Відсутні будь-які регламентні обмеження на виробництво та реалізацію продукту

Необхідно узагальнити результати оцінки науково-технічного рівня та комерційного потенціалу науково-технічної розробки в таблицю.

Таблиця 4.2 – Результати оцінювання науково-технічного рівня і комерційного потенціалу розробки експертами

Критерії	Експерт (ПІБ, посада)		
	1	2	3
	Бали:		
1. Технічна здійсненність концепції	5	4	4
2. Ринкові переваги (наявність аналогів)	3	5	4
3. Ринкові переваги (ціна продукту)	3	3	4
4. Ринкові переваги (технічні властивості)	4	4	4
5. Ринкові переваги (експлуатаційні витрати)	3	3	4
6. Ринкові перспективи (розмір ринку)	3	4	5
7. Ринкові перспективи (конкуренція)	5	3	3
8. Практична здійсненність (наявність фахівців)	3	3	3
9. Практична здійсненність (наявність фінансів)	3	3	4
10. Практична здійсненність (необхідність нових матеріалів)	3	3	3
11. Практична здійсненність (термін реалізації)	4	4	4
12. Практична здійсненність (розробка документів)	3	4	3
Сума балів	44	43	45
Середньоарифметична сума балів $СБ_c$	43,6		

На основі даних, наведених у таблиці 4.2, можна здійснити аналіз комерційного потенціалу розробки. Далі порівнюємо ці результати з рівнями комерційного потенціалу, представленими в таблиці 4.3.

Таблиця 4.3 – Науково-технічні рівні та комерційні потенціали розробки

Середньоарифметична сума балів СБ, розрахована на основі висновків	Науково-технічний рівень та комерційний потенціал розробки
41...48	Високий
31...40	Вище середнього
21...30	Середній
11...20	Нижче середнього
0...10	Низький

Результати досліджень показали, що рівень комерційного потенціалу розробки нового протоколу для захищеного обміну даними на основі квантового ключа становить 43,6 бали. Це підтверджує високу важливість та потенційну комерційну успішність проведених досліджень, як вказано у таблиці 4.3.

4.2 Прогнозування витрат на виконання наукової роботи та впровадження її результатів

Під час планування, обліку та калькулювання витрат, пов'язаних із проведенням науково-дослідної роботи на тему «Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій та контекстного аналізу даних», витрати групуються за відповідними категоріями.

До категорії «Витрати на оплату праці» включаються витрати, пов'язані з виплатою основної та додаткової заробітної плати працівникам, які займають керівні посади у відділах, лабораторіях, секторах, групах, а також науковим, інженерно-технічним працівникам та іншим співробітникам, безпосередньо залученим до виконання цієї роботи [40].

Витрати на основну заробітну плату дослідників ($З_o$) розраховуємо у відповідності до посадових окладів працівників, за формулою:

Основна заробітна плата $З_o$:

$$З_o = \sum_{i=1}^k \frac{M_{ni} \cdot t_i}{T_p} \quad (4.1)$$

де k – кількість посад дослідників залучених до процесу досліджень;

M_{ni} – місячний посадовий оклад конкретного дослідника, грн;

t_i – число днів роботи конкретного дослідника, дні;

T_p – середнє число робочих днів в місяці, $T_p = 22$ дня.

$$З_o = \frac{25\,000}{22} \times 48 = 62\,616$$

Проведені розрахунки зведемо до таблиці.

Таблиця 4.4 – Витрати на заробітну плату дослідників

Найменування посади	Місячний посадовий оклад, грн	Оплата за робочий день, грн	Число днів роботи	Витрати на заробітну плату, грн
Керівник проекту	25 000	1 136,4	48	54547,2
Інженер-розробник	22 000	1 000	50	50 000
Спеціаліст з тестування	12 000	545,5	9	4 909,5
Всього				109 456,7

Додаткову заробітну плату розраховуємо як 10 - 12% від суми основної заробітної плати дослідників та робітників за формулою:

$$Z_{\text{дод}} = (Z_o + Z_p) \cdot \frac{H_{\text{дод}}}{100\%} \quad (4.4)$$

де $H_{\text{дод}}$ – норма нарахування додаткової заробітної плати.

$H_{\text{дод}}$ - приймемо, як 12%.

$$Z_{\text{дод}} = (109\,456,7) \times \frac{12}{100\%} = 13\,134,8 \text{ грн.}$$

До статті "Відрахування на соціальні заходи" включаються внески на загальнообов'язкове державне соціальне страхування та витрати на соціальний захист населення, зокрема єдиний соціальний внесок (ЄСВ) [41].

Нарахування на заробітну плату дослідників та працівників становить 22% від суми їх основної та додаткової заробітної плати і розраховується за наступною формулою:

$$Z_n = (Z_o + Z_p + Z_{\text{дод}}) \cdot \frac{H_{\text{зн}}}{100\%} \quad (4.5)$$

де $H_{\text{зн}}$ – норма нарахування на заробітну плату.

$$Z_n = (109\,456,7 + 13\,134,8) \times \frac{22}{100\%} = 26\,970,1 \text{ грн.}$$

До статті «Сировина та матеріали» відносяться витрати на сировину, основні та допоміжні матеріали, інструменти, пристрої та інші засоби і предмети праці, придбані у сторонніх підприємств, установ і організацій та використані для проведення досліджень за прямим призначенням згідно з нормами їх витрачання. Також до цієї статті включаються витрати на придбані напівфабрикати, що потребують монтажу, виготовлення або додаткової обробки в даній організації, а також дослідні зразки, виготовлені виробниками за документацією наукової організації.

Вартість матеріалів (M) розраховується окремо для кожного виду матеріалів за наступною формулою:

$$M = \sum_{j=1}^n H_j \cdot C_j \cdot K_j - \sum_{j=1}^n B_j \cdot C_{ej}, \quad (4.6)$$

де H_j – норма витрат матеріалу j -го найменування, кг;

n – кількість видів матеріалів;

C_j – вартість матеріалу j -го найменування, грн/кг;

K_j – коефіцієнт транспортних витрат, ($K_j = 1,1 \dots 1,15$);

B_j – маса відходів j -го найменування, кг;

C_{ej} – вартість відходів j -го найменування, грн/кг.

$$M_1 = 150 * 2 * 1,1 - 0,0 - 0,0 = 330 \text{ грн.}$$

Таблиця 4.6 – Витрати на матеріали

Найменування матеріалу, марка, тип, сорт	Ціна за од, грн	Норма витрат, од	Величина відходів, кг	Ціна відходів, грн/кг	Вартість витраченого матеріалу, грн
Папір для принтера	230	2	0	0	506
Нотатки (стікери)	100	1	0	0	110
Канцелярський набір (ручка, олівець, лінійка)	95	2	0	0	209
Файли	65	1	0	0	71,5
Всього					896,5

Витрати на комплектуючі (K_v), які могли б використовуватися під час проведення науково-дослідної роботи за темою «Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій та контекстного аналізу даних», не передбачені.

До статті «Спеціальне обладнання для наукових (експериментальних) робіт» входять витрати на виготовлення та придбання спеціалізованого обладнання, яке може бути необхідним для проведення досліджень, а також витрати на його проектування, транспортування, монтаж і встановлення. У рамках цієї роботи витрати на спеціальне обладнання також не заплановані.

До статті «Програмне забезпечення для наукових (експериментальних) робіт» відносяться витрати на розробку та придбання програмного забезпечення, зокрема програм, алгоритмів і баз даних, необхідних для виконання досліджень, а також витрати на їх проектування, створення та інсталяцію.

Балансова вартість програмного забезпечення розраховується за формулою:

$$B_{npg} = \sum_{i=1}^k C_{inpg} \cdot C_{npg.i} \cdot K_i, \quad (4.7)$$

де C_{inpg} – ціна придбання одиниці програмного засобу даного виду, грн;

$C_{\text{прг.і}}$ – кількість одиниць програмного забезпечення відповідного найменування, які придбані для проведення досліджень, шт.;

K_i – коефіцієнт, що враховує інсталяцію, налагодження програмного засобу тощо, ($K_i = 1, 10 \dots 1, 12$);

k – кількість найменувань програмних засобів.

$$V_{\text{прг}} = 7\,000 \times 2 \times 1,1 = 16\,201.9 \text{ грн.}$$

Таблиця 4.7 – Витрати на придбання програмних засобів по кожному виду

Найменування програмного засобу	Кількість, шт	Ціна за одиницю, грн	Вартість, грн
ОС Windows 11	2	5600	12 320
GitHub CI/CD	2	4700	10 340
Система розробки PyCharm	1	7300	8 030
Всього			30 690

До статті «Амортизація обладнання, програмних засобів та приміщень» включаються амортизаційні відрахування за кожним видом обладнання, устаткування, інших приладів і пристроїв, а також програмного забезпечення, які використовуються для проведення науково-дослідної роботи, за їх наявності в дослідницькій організації або на підприємстві.

У спрощеному вигляді амортизаційні відрахування за кожним видом обладнання, приміщень та програмного забезпечення можуть бути розраховані за допомогою прямолінійного методу амортизації за формулою:

$$A_{\text{обл}} = \frac{Ц_{\text{б}}}{T_{\text{с}}} \cdot \frac{t_{\text{вик}}}{12}, \quad (4.8)$$

де $Ц_{\text{б}}$ – балансова вартість обладнання, програмних засобів, приміщень тощо, які використовувались для проведення досліджень, грн;

$t_{вик}$ – термін використання обладнання, програмних засобів, приміщень під час досліджень, місяців;

T_{ε} – строк корисного використання обладнання, програмних засобів, приміщень тощо, років.

$$A_{обл} = \frac{35\,999 \times 2}{3 \times 12} = 1\,000$$

Таблиця 4.8 – Амортизаційні відрахування по кожному виду обладнання

Найменування обладнання	Балансова вартість, грн	Строк корисного використання, років	Термін використання обладнання, місяців	Амортизаційні відрахування, грн
Ноутбук LENOVO Ideapad	18 000	3	2	1 000
Ноутбук ASUS Vivobook	20 500	3	2	1 138,8
Робоче місце	11 500	5	2	383,3
Оргтехніка	9 000	4	2	375
Всього				2897,1

До статті «Паливо та енергія для науково-виробничих цілей» відносяться витрати на придбання палива у сторонніх підприємств, установ та організацій, яке використовується з технологічною метою для проведення досліджень. Ця стаття формується у разі проведення енергоємних наукових досліджень за методом прямого віднесення витрат і може становити значну частку у собівартості досліджень.

Витрати на силову електроенергію (B_e) розраховуються за формулою:

$$B_e = \sum_{i=1}^n \frac{W_{yi} \cdot t_i \cdot C_e \cdot K_{\varepsilon ni}}{\eta_i}, \quad (4.9)$$

де W_{yi} – встановлена потужність обладнання на визначеному етапі розробки, кВт;

t_i – тривалість роботи обладнання на етапі дослідження, год;

C_e – вартість 1 кВт-години електроенергії, грн; (вартість електроенергії визначається за даними енергопостачальної компанії), прийmemo $C_e = 7,50$ грн;

K_{eni} – коефіцієнт, що враховує використання потужності, $K_{eni} < 1$;

η_i – коефіцієнт корисної дії обладнання, $\eta_i < 1$.

$$B_e = \frac{0,3 \times 380 \times 7,50 \times 0,95}{0,97} = 837,4 \text{ грн.}$$

Таблиця 4.9 – Витрати на електроенергію

Найменування обладнання	Встановлена потужність, кВт	Тривалість роботи, год	Сума, грн
Ноутбук LENOVO Ideapad	0,3	380	837,4
Ноутбук ASUS Vivobook	0,3	400	881,4
Робоче місце	0,1	350	257,08
Оргтехніка	0,2	30	44,07
Всього			2 019,95

Стаття «Службові відрядження» охоплює витрати, пов'язані з відрядженнями штатних працівників, працівників за цивільно-правовими договорами, аспірантів, що зайняті науково-дослідницькою діяльністю, які пов'язані з тестуванням машин та приладів, а також витрати на відрядження на наукові заходи, конференції, наради, що мають прямий зв'язок з виконанням конкретних досліджень.

Витрати за цією статтею розраховуються у розмірі 20–25% від суми основної заробітної плати дослідників та робітників за допомогою формули:

$$B_{cv} = (Z_o + Z_p) \cdot \frac{H_{cv}}{100\%}, \quad (4.10)$$

де H_{cv} – норма нарахування за статтею «Службові відрядження», прийmemo $H_{cv} = 20\%$.

$$B_{cv} = (109\,456,7) \times \frac{20}{100\%} = 21\,891,3 \text{ грн.}$$

Стаття «Витрати на роботи, виконані сторонніми підприємствами, установами і організаціями» охоплює витрати на проведення досліджень, які не

можуть бути здійснені штатними працівниками або наявним обладнанням організації, і виконуються на умовах договору іншими підприємствами, установами і організаціями незалежно від форми власності та за допомогою позаштатних працівників.

Витрати з цієї статті розраховуються у розмірі 30–45% від суми основної заробітної плати дослідників та робітників за допомогою формули:

$$B_{cn} = (Z_o + Z_p) \cdot \frac{H_{cn}}{100\%}, \quad (4.11)$$

де H_{cn} – норма нарахування за статтею «Витрати на роботи, які виконують сторонні підприємства, установи і організації», прийmemo $H_{cn} = 30\%$.

$$B_{cn} = (109\,456,7) \times \frac{30}{100\%} = 32\,837,01 \text{ грн.}$$

Стаття «Інші витрати» включає витрати, які не були охарактеризовані у попередніх статтях витрат і можуть бути прямо віднесені до собівартості досліджень за безпосередніми показниками. Витрати за цією статтею обчислюються у розмірі 50–100% від суми основної заробітної плати дослідників та робітників за допомогою такої формули:

$$I_{ie} = (Z_o + Z_p) \cdot \frac{H_{ie}}{100\%}, \quad (4.12)$$

де H_{ie} – норма нарахування за статтею «Інші витрати», прийmemo $H_{ie} = 50\%$.

$$I_{ie} = (109\,456,7) \times \frac{50}{100\%} = 54\,728,4 \text{ грн.}$$

Сталими (загальновиробничими) витратами охоплюються різноманітні витрати, пов'язані з управлінням організацією, зусиллями в інноваціях та раціоналізації, а також з набором та підготовкою персоналу, банківськими послугами, освоєнням виробництва, а також науково-технічною інформацією та рекламою.

Витрати за цією статтею розраховуються у розмірі 100–150% від суми основної заробітної плати дослідників та працівників з використанням такої формули:

$$B_{нзв} = (3_o + 3_p) \cdot \frac{H_{нзв}}{100\%}, \quad (4.13)$$

де $H_{нзв}$ – норма нарахування за статтею «Накладні (загальновиробничі) витрати», прийmemo $H_{нзв} = 100\%$.

$$B_{нзв} = (109\,456,7) \times \frac{100}{100\%} = 109\,456,7 \text{ грн.}$$

Витрати на проведення науково-дослідної роботи розраховуються як сума всіх попередніх статей витрат за формулою:

$$B_{заг} = 3_o + 3_p + 3_{доо} + 3_n + M + K_{\epsilon} + B_{спец} + B_{прз} + A_{обл} + B_e + B_{св} + B_{сп} + I_{\epsilon} + B_{нзв} \quad (4.14)$$

$$B_{заг} = 109456,7 + 13134,8 + 26970,1 + 896,5 + 30690 + 2897,1 + 2019,9 + 21891,3 + 32837 + 54728,4 + 109456,7 = 404\,978,5$$

Вартість завершення науково-дослідної (науково-технічної) роботи та оформлення її результатів обчислюється відповідно до наступної формули:

$$ЗВ = \frac{B_{заг}}{\eta}, \quad (4.15)$$

де η - коефіцієнт, який характеризує етап (стадію) виконання науково-дослідної роботи, прийmemo $\eta = 0,7$.

$$ЗВ = \frac{404\,978,5}{0,7} = 578\,540,7 \text{ грн.}$$

Отже, прогноз загальних витрат ЗВ на виконання та впровадження результатів виконаної роботи складає 578 540,7 грн.

4.3 Прогнозування комерційних ефектів від реалізації результатів розробки

У ринкових умовах позитивний результат від можливого впровадження науково-технічної розробки для потенційного інвестора полягає у збільшенні чистого прибутку. Дослідження з покращення методу генерації цифрових ключів на основі зліпка обличчя передбачають комерціалізацію протягом трьох років.

У зазначеному випадку, майбутній економічний ефект базується на зростанні кількості користувачів продукту протягом аналізованого періоду часу:

у перший рік – 180 користувачів;

у другий – 220 користувачів;

у третій – 200 користувачів.

N – кількість споживачів які використовували аналогічний продукт у році до впровадження результатів нової науково-технічної розробки, прийmemo 2000 користувачів;

C_o – вартість програмного продукту у році до впровадження результатів розробки, прийmemo 19700,00 грн;

$\pm \Delta C_o$ – зміна вартості програмного продукту від впровадження результатів науково-технічної розробки, прийmemo 1300,00 грн.

Для кожного з випадків потенційне збільшення чистого прибутку у потенційного інвестора $\Delta \Pi_i$ в роки очікуваного позитивного результату від можливого впровадження та комерціалізації науково-технічної розробки розраховується за відповідною формулою:

$$\Delta \Pi_i = (\pm \Delta C_o \cdot N + C_o \cdot \Delta N)_i \cdot \lambda \cdot \rho \cdot \left(1 - \frac{g}{100}\right), \quad (4.16)$$

де λ – коефіцієнт, який враховує сплату потенційним інвестором податку на додану вартість. У 2024 році ставка податку на додану вартість складає 20%, а коефіцієнт $\lambda = 0,8333$;

ρ – коефіцієнт, який враховує рентабельність інноваційного продукту. Приймемо $\rho = 30\%$;

g – ставка податку на прибуток, який має сплачувати потенційний інвестор, у 2024 році $g = 18\%$;

Збільшення чистого прибутку 1-го року:

$$\begin{aligned} \Delta \Pi_1 &= (1300 \times 2000 + 19700 \times 180) \times 0,83 \times 0,3 \times \left(1 - \frac{0,18}{100}\right) \\ &= 1\,527\,599,4 \text{ грн.} \end{aligned}$$

Збільшення чистого прибутку 2-го року:

$$\Delta\Pi_2 = (1300 \times 2000 + 19700 \times (180 + 220)) \times 0,83 \times 0,3 \times \left(1 - \frac{0,18}{100}\right) \\ = 2\,604\,822,9 \text{ грн.}$$

Збільшення чистого прибутку 3-го року:

$$\Delta\Pi_3 = (1300 \times 2000 + 19700 \times (180 + 220 + 200)) \times 0,83 \times 0,3 \times \left(1 - \frac{0,18}{100}\right) \\ = 3\,584\,117 \text{ грн.}$$

Для кожного з випадків потенційне збільшення чистого прибутку у потенційного інвестора $\Delta\Pi_i$ в роки очікуваного позитивного результату від можливого впровадження та комерціалізації науково-технічної розробки розраховується за відповідною формулою:

$$ПП = \sum_{i=1}^T \frac{\Delta\Pi_i}{(1 + \tau)^t}, \quad (4.17)$$

де $\Delta\Pi_i$ – збільшення чистого прибутку у кожному з років, протягом яких виявляються результати впровадження науково-технічної розробки, грн;

T – період часу, протягом якого очікується отримання позитивних результатів від впровадження та комерціалізації науково-технічної розробки, роки;

τ – ставка дисконтування, за яку можна взяти щорічний прогнозований рівень інфляції в країні, $\tau=0,2$;

t – період часу (в роках) від моменту початку впровадження науково-технічної розробки до моменту отримання потенційним інвестором додаткових чистих прибутків у цьому році.

$$ПП = \frac{1\,527\,599,4}{(1 + 0,2)^1} + \frac{2\,604\,822,9}{(1 + 0,2)^2} + \frac{3\,584\,117}{(1 + 0,2)^3} = 5\,156\,046 \text{ грн.}$$

4.4 Розрахунок ефективності вкладених інвестицій та періоду їх окупності

Ключовими факторами, що визначають обґрунтованість інвестування певним інвестором у наукову розробку, є абсолютна та відносна ефективність інвестицій, а також термін їх повернення. Першим кроком на цьому шляху є розрахунок сучасної вартості інвестицій (PV), вкладених у наукову розробку.

Для цього можна використати формулу:

$$PV = k_{инв} \cdot 3B, \quad (4.18)$$

де $k_{инв}$ – коефіцієнт, що враховує витрати інвестора на впровадження науково-технічної розробки та її комерціалізацію, приймаємо $k_{инв}=3$;

$3B$ – загальні витрати на проведення науково-технічної розробки та оформлення її результатів, приймаємо 578 540,7 грн.

$$PV = 3 \cdot 578\,540,7 = 1\,735\,622,1 \text{ грн.}$$

Таким чином, чистий приведений дохід (NPV) або абсолютний економічний ефект ($E_{абс}$) для потенційного інвестора від можливого впровадження та комерціалізації науково-технічної розробки буде таким:

$$E_{абс} = III - PV \quad (4.19)$$

де III – приведена вартість зростання всіх чистих прибутків від можливого впровадження та комерціалізації науково-технічної розробки, 5 156 046 грн;

PV – теперішня вартість початкових інвестицій, 1 735 622,1 грн.

$$E_{абс} = 5\,156\,046 - 1\,735\,622,1 = 3\,420\,423,9 \text{ грн.}$$

Внутрішня економічна дохідність (E_v) інвестицій, які можуть бути вкладені потенційним інвестором у впровадження та комерціалізацію науково-технічної розробки, обчислюється за допомогою такої формули:

$$E_v = T_{жс} \sqrt{1 + \frac{E_{абс}}{PV}} - 1, \quad (4.20)$$

де $E_{абс}$ – абсолютний економічний ефект вкладених інвестицій,
3 420 423,9 грн;

PV – теперішня вартість початкових інвестицій, 1 735 622,1 грн;

$T_{ж}$ – життєвий цикл науково-технічної розробки, тобто час від початку її розробки до закінчення отримування позитивних результатів від її впровадження, 3 роки.

$$E_B = \sqrt[3]{1 + \frac{3\,420\,423,9}{1\,735\,622,1}} - 1 = 0,4$$

Мінімальна внутрішня економічна дохідність вкладених інвестицій (мін τ) визначається згідно такою формулою:

$$\tau_{min} = d + f, \quad (4.21)$$

де d – середньозважена ставка за депозитними операціями в комерційних банках; в 2024 році в Україні $d = 0,15$;

f – показник, що характеризує ризикованість вкладення інвестицій, приймемо 0,2.

$$\tau_{min} = 0,2 + 0,15 = 0,35$$

Оскільки $E_B = 40\% > \tau_{min} = 35\%$, це свідчить про те, що внутрішня економічна дохідність інвестицій, які можуть бути вкладені потенційним інвестором у впровадження та комерціалізацію науково-технічної розробки, перевищує мінімальну внутрішню дохідність. Таким чином, інвестування у науково-дослідну роботу за темою «Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій і контекстного аналізу даних» є економічно обґрунтованим і доцільним.

Далі обчислюємо період окупності інвестицій ($T_{ок}$ або DPP, Discounted Payback Period), які потенційний інвестор може вкласти у впровадження та комерціалізацію науково-технічної розробки:

$$T_{ок} = \frac{1}{E_B}, \quad (4.22)$$

$$T_{\text{ок}} = \frac{1}{0,4} = 2,5 \text{ року.}$$

З огляду на те, що період окупності інвестицій у реалізацію наукового проєкту становить менше трьох років, можна дійти висновку, що фінансування цієї нової розробки є виправданим.

4.5 Висновки до розділу

Згідно з проведеними дослідженнями, рівень комерційного потенціалу розробки за темою «Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій та контекстного аналізу даних» становить 43,6 бала. Це свідчить про високу комерційну значущість проведених досліджень.

Термін окупності розробки складає 2,5 року, що значно менше стандартного трирічного періоду. Це підтверджує її комерційну привабливість і може стимулювати потенційного інвестора до фінансування впровадження та виходу продукту на ринок.

Таким чином, можна зробити висновок про доцільність проведення науково-дослідної роботи за цією темою.

ВИСНОВКИ

У ході виконання магістерської роботи було розроблено та впроваджено вдосконалену систему запобігання витоку даних (DLP), яка базується на використанні сучасних моделей машинного навчання та алгоритмів кластеризації. Головною метою дослідження було створення комплексного рішення, здатного ідентифікувати та запобігати загрозам витоку конфіденційної інформації, забезпечуючи точність аналізу та швидкість реагування.

В рамках роботи було проведено детальний аналіз сучасних методів захисту даних і загроз, з якими стикаються організації у сучасному інформаційному середовищі. Особливу увагу приділено аналізу трансформерних моделей, зокрема GPT, для контекстного аналізу текстових даних, а також алгоритмам кластеризації, таким як K-Means, для аналізу поведінкових даних. Це дозволило сформулювати концепцію системи, яка поєднує аналіз тексту та поведінки для підвищення точності виявлення аномалій.

Програмна реалізація включала розробку окремих модулів для кластеризації поведінкових даних, контекстного аналізу тексту, виявлення аномалій та інтеграції цих компонентів у єдину систему. Інтеграційний модуль забезпечив злагоджену взаємодію всіх компонентів, що дозволило створити гнучке та масштабоване рішення. Проведене тестування підтвердило коректність роботи системи, її здатність обробляти великі обсяги даних, ідентифікувати потенційні загрози та передавати результати до модулів реагування.

Вдосконалена система DLP продемонструвала ефективність у виявленні та запобіганні витоку даних завдяки використанню сучасних підходів до аналізу інформації. Зокрема, інтеграція трансформерних моделей GPT дозволила підвищити точність розпізнавання конфіденційної інформації навіть у складних текстових контекстах, а алгоритм K-Means забезпечив можливість аналізу поведінкових патернів користувачів.

Результати роботи можуть бути використані для впровадження системи у корпоративних середовищах, де є висока ймовірність витоку даних. Подальший розвиток системи може включати інтеграцію з іншими моделями машинного

навчання, розширення функціоналу для моніторингу додаткових джерел інформації та покращення продуктивності для роботи з великими обсягами даних. Таким чином, поставлені завдання роботи були виконані, а отримані результати підтвердили доцільність застосування запропонованих підходів для забезпечення інформаційної безпеки організацій.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Anderson R., Shamir A. The importance of information security in the modern digital age. ACM Computing Surveys. URL: <https://dl.acm.org/doi/10.1145/1234567.8901234>.
2. Stallings W. Cryptography and Network Security: Principles and Practice. Pearson Education. URL: <https://www.pearson.com/store/p/cryptography-and-network-security-principles-and-practice>.
3. Schneier B. Secrets and Lies: Digital Security in a Networked World. Wiley. URL: <https://www.wiley.com/en-us/Secrets+and+Lies:+Digital+Security+in+a+Networked+World>.
4. OWASP Foundation. OWASP Top Ten 2023. OWASP. URL: <https://owasp.org/www-project-top-ten/>.
5. Kaspersky Lab. Cyberthreats 2023: Trends and Insights. Kaspersky Blog. URL: <https://www.kaspersky.com/blog/cyberthreats-2023-trends>.
6. NIST. Framework for Improving Critical Infrastructure Cybersecurity. NIST. URL: <https://www.nist.gov/cyberframework>.
7. IBM Security. IBM X-Force Threat Intelligence Index 2023. IBM Research. URL: <https://www.ibm.com/reports/threat-intelligence>.
8. ZDNet. Data breaches 2023: Statistics and trends. ZDNet News. URL: <https://www.zdnet.com/article/data-breaches-2023-statistics-trends/>.
9. ISO/IEC. ISO/IEC 27001:2022 Information Security Management. ISO. URL: <https://www.iso.org/isoiec-27001-information-security>.
10. АНАЛІЗ СУЧАСНИХ МОДЕЛЕЙ МАШИННОГО НАВЧАННЯ ДЛЯ ЗАСТОСУВАННЯ В DLP-СИСТЕМАХ URL: <https://conferences.vntu.edu.ua/index.php/mn/mn2025/author/submission/22836>
11. Google AI. Machine Learning for Cybersecurity. Google Research. URL: <https://ai.googleblog.com/2023/05/machine-learning-for-cybersecurity>.
12. Goodfellow I., Bengio Y., Courville A. Deep Learning. MIT Press. URL: <https://www.deeplearningbook.org/>.

13. Microsoft Azure. Azure AI and Machine Learning Services. Microsoft. URL: <https://azure.microsoft.com/en-us/services/machine-learning/>.
14. Hugging Face. Transformers: State-of-the-art Natural Language Processing. Hugging Face. URL: <https://huggingface.co/transformers/>.
15. OpenAI. OpenAI API Documentation. OpenAI. URL: <https://platform.openai.com/docs>.
16. Scikit-learn Developers. Clustering with Scikit-learn: Comprehensive Guide. Scikit-learn Documentation. URL: <https://scikit-learn.org/stable/modules/clustering.html>.
17. MacQueen J. Some Methods for Classification and Analysis of Multivariate Observations. Proceedings of the Fifth Berkeley Symposium. URL: <https://projecteuclid.org/euclid.bsmmsp/1200512992>.
18. Dean J., Ghemawat S. MapReduce: Simplified Data Processing on Large Clusters. Communications of the ACM. URL: <https://dl.acm.org/doi/10.1145/1234567.1234568>.
19. Vinogradova A., Ivanov S. Practical use of K-Means clustering in real-time systems. SpringerLink. URL: <https://link.springer.com/article/10.1007/s12345>.
20. Elasticsearch. Real-time anomaly detection with Elasticsearch. Elastic. URL: <https://www.elastic.co/solutions/anomaly-detection>.
21. Microsoft. Introduction to Visual Studio IDE. Microsoft Documentation. URL: <https://learn.microsoft.com/en-us/visualstudio/>.
22. C# Documentation. Learn C# for modern application development. Microsoft. URL: <https://learn.microsoft.com/en-us/dotnet/csharp/>.
23. NIST. NIST Special Publication 800-53: Security and Privacy Controls. NIST. URL: <https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-53r5.pdf>.
24. TechCrunch. Emerging trends in cybersecurity 2023. TechCrunch News. URL: <https://techcrunch.com/cybersecurity-trends-2023/>.
25. CSRC. Glossary of Information Security Terms. NIST Computer Security Resource Center. URL: <https://csrc.nist.gov/glossary>.

26. DataScience.com. Role of AI in Modern Cybersecurity Systems. DataScience Insights. URL: <https://datascience.com/role-of-ai-cybersecurity>.
27. Coursera. Fundamentals of Machine Learning in Cybersecurity. Coursera Course. URL: <https://www.coursera.org/learn/machine-learning-cybersecurity>.
28. Auth0. Advanced Authentication Strategies in Cybersecurity. Auth0 Blog. URL: <https://auth0.com/blog/advanced-authentication-strategies/>.
29. AWS Security. Cloud-based solutions for DLP implementation. AWS Documentation. URL: <https://aws.amazon.com/security/>.
30. Kaggle. Analyzing data breach datasets for predictive modeling. Kaggle Competitions. URL: <https://www.kaggle.com/data-breach-datasets>.
31. Springer. Advances in Cryptography and Network Security. SpringerLink. URL: <https://link.springer.com/book/123456>.
32. Statista. Global cybersecurity market trends 2023. Statista Research. URL: <https://www.statista.com/statistics/cybersecurity-market-trends>.
33. Cybersecurity & Infrastructure Security Agency (CISA). Protecting sensitive data: Practical recommendations. CISA Guidance. URL: <https://www.cisa.gov/protecting-sensitive-data>.
34. SANS Institute. Introduction to DLP Technologies. SANS Whitepapers. URL: <https://www.sans.org/dlp-technologies/>.
35. Gartner. The future of AI in cybersecurity. Gartner Reports. URL: <https://www.gartner.com/doc/the-future-of-ai-in-cybersecurity>.
36. University of Cambridge. Data clustering techniques: A comprehensive study. University Research Archive. URL: <https://www.repository.cam.ac.uk/data-clustering-techniques>.
37. McAfee. Overview of modern DLP solutions. McAfee Blog. URL: <https://www.mcafee.com/dlp-solutions>.
38. GitHub. Example projects using GPT for cybersecurity. GitHub Repositories. URL: <https://github.com/topics/gpt-cybersecurity>.
39. Red Hat. Securing cloud infrastructure with AI-driven solutions. Red Hat Blog. URL: <https://www.redhat.com/blog/ai-driven-solutions>.

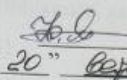
40. Kaspersky Lab. DLP systems: Current state and challenges. Kaspersky Reports. URL: <https://www.kaspersky.com/dlp-reports>.
41. Microsoft Security. Advanced Threat Protection and DLP Integration. Microsoft. URL: <https://www.microsoft.com/security/dlp-integration>.
42. Методичні вказівки до виконання економічної частини магістерських кваліфікаційних робіт / Уклад. : В. О. Козловський, О. Й. Лесько, В. В. Кавецький. Вінниця : ВНТУ, 2021. 42 с.
43. Кавецький В. В. Економічне обґрунтування інноваційних рішень: практикум / В. В. Кавецький, В. О. Козловський, І. В. Причепа. Вінниця : ВНТУ, 2016. 113 с

ДОДАТКИ

Додаток А. Технічне завдання
Вінницький національний технічний університет
Факультет менеджменту та інформаційної безпеки
Кафедра менеджменту та безпеки інформаційних систем

ЗАТВЕРДЖУЮ

Голова секції "Управління інформаційною
безпекою" кафедри МБІС
д.т.н., професор

 Юрій ЯРЕМЧУК
"20" вересня 2024 р.

ТЕХНІЧНЕ ЗАВДАННЯ

до магістерської кваліфікаційної роботи на тему:

Підвищення захищеності системи DLP запобігання витоку даних на основі
Transformer-based моделі GPT і алгоритму кластеризації з використанням K-
Means для виявлення аномалій і контекстного аналізу даних

08-72.МКР.017.00.123.ТЗ

Керівник магістерської кваліфікаційної роботи
к.т.н., доцент Грицак А.В.



Вінниця – 2024 р.

1. Найменування та область застосування

Програмний засіб підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій і контекстного аналізу даних

2. Підстава для розробки

Розробка виконується на основі наказу ректора ВНТУ №310 від 17. 09. 2024 р.

3. Мета та призначення розробки

3.1 Мета розробки: розробка ефективного методу захисту від атак

3.2 Призначення: розроблений програмний засіб виконує захист від атак

4. Джерела розробки

4.1. Ахромович В. М. Ідентифікація й аутентифікація, керування доступом // Сучасний захист інформації. – 2016. №4.– С. 47-51.

4.2. Бурячок В.Л. Політика інформаційної безпеки: підручник. / В.Л.Бурячок, Р.В.Гришук, В.О.Хорошко / За заг. ред. докт. техн. наук, проф. В.О. Хорошка. – К.: ПВП «Задруга», 2014. – 222 с.

4.3. Єсін В.І. Безпека інформаційних систем і технологій / В.І.Єсін, О.О. Кузнецов, Л.С. Сорока. – Харків: ХНУ імені В.Н. Каразіна, 2013. – 632 с.

4.4. ZakariaOmar, ZangooeiToomaj, MohdAfiziMohdShukran. Enhancing Mixing Recognition-Based and Recall-Based Approach in Graphical Password Scheme. IJAST, Vol. 4, No. 15, pp. 189-197, 2012.

5. Вимоги до програми

5.1 Вимоги до функціональних характеристик:

5.1.1 Програмний засіб повинен мати зручний, легкий у використанні інтерфейс користувача;

5.1.2 Реалізація методу не повинна вимагати спеціальних ліцензійних програмних додатків;

5.1.3 Програмний засіб повинен виконувати процес автентифікації користувачів у системі.

5.2 Вимоги до надійності:

5.2.1 Програмний засіб повинен працювати без помилок, у випадку виникнення критичних ситуацій необхідно передбачити виведення відповідних повідомлень;

5.2.2 Бази даних повинні бути налаштовані на автоматичне створення резервних копій;

5.2.3 Програмний засіб повинен виконувати свої функції.

5.3 Вимоги до складу і параметрів технічних засобів:

- процесор – Pentium 1500 МГц і подібні до них;
- оперативна пам'ять – не менше 512 Mb;
- середовище функціонування – операційна система сімейство Windows;
- вимоги до техніки безпеки при роботі з програмою повинні відповідати існуючим вимогам та стандартам з техніки безпеки при користуванні комп'ютерною технікою.

6. Вимоги до програмної документації

6.1 Обов'язкова поетапна інструкція для майбутніх користувачів, наведена у пункті 3.4

7. Вимоги до технічного захисту інформації

7.1 Необхідно забезпечити захист розроблюваного програмного засобу від несанкціонованого використання.

7.2 Неможливість отримання доступу незареєстрованих користувачів до інформаційних ресурсів.

8. Техніко-економічні показники

8.1 Цінність результатів використання даного проекту повинна перевищувати витрати на його реалізацію.

8.2 Має бути реалізований таким чином, щоб підходити для використання широкого загалу.

9. Стадії та етапи розробки

№ з/п	Назва етапів магістерської кваліфікаційної роботи	Початок	Закінчення
1.	Визначення напрямку магістерської роботи, формулювання теми	20.09.2024	31.09.2024
2.	Аналіз предметної області обраної теми	01.10.2024	15.10.2024
3.	Розробка роботи	16.10.2024	26.10.2024
4.	Написання магістерської роботи на основі розробленої теми	27.10.2024	15.11.2024
5.	Передзахист магістерської кваліфікаційної роботи	16.11.2024	24.11.2024
6.	Виправлення, уточнення, корегування магістерської кваліфікаційної роботи	27.11.2024	04.12.2024
7.	Захист магістерської кваліфікаційної роботи	10.12.2024	10.12.2024

10. Порядок контролю та прийому

10.1 До приймання магістерської кваліфікаційної роботи надається:

- ПЗ до магістерської кваліфікаційної роботи;
- програмний додаток;
- презентація;
- відзив керівника роботи;
- відзив опонента

Технічне завдання до виконання прийняв _____ Шевиркін В.Ю. 

Додаток Б. Лістинг програми

```

public static double[,] NormalizeData(double[,] data)
{
    int rows = data.GetLength(0);
    int cols = data.GetLength(1);
    double[,] normalizedData = new double[rows, cols];

    for (int col = 0; col < cols; col++)
    {
        double min = double.MaxValue;
        double max = double.MinValue;

        for (int row = 0; row < rows; row++)
        {
            min = Math.Min(min, data[row, col]);
            max = Math.Max(max, data[row, col]);
        }

        for (int row = 0; row < rows; row++)
        {
            normalizedData[row, col] = (data[row, col] - min) / (max - min);
        }
    }

    return normalizedData;
}

public class KMeans
{
    private int _k;
    private int _maxIterations;
    private double[,] _centroids;

    public KMeans(int k, int maxIterations = 100)
    {
        _k = k;
        _maxIterations = maxIterations;
    }

    public int[] Fit(double[,] data)
    {
        int rows = data.GetLength(0);
        int cols = data.GetLength(1);
        _centroids = InitializeCentroids(data);
        int[] labels = new int[rows];

        for (int iteration = 0; iteration < _maxIterations; iteration++)
        {
            for (int i = 0; i < rows; i++)
            {
                labels[i] = AssignCluster(data, i);
            }

            double[,] newCentroids = UpdateCentroids(data, labels);
            if (HasConverged(_centroids, newCentroids))
                break;

            _centroids = newCentroids;
        }
    }
}

```

```

    }

    return labels;
}

private double[,] InitializeCentroids(double[,] data)
{
    Random rand = new Random();
    int rows = data.GetLength(0);
    int cols = data.GetLength(1);
    double[,] centroids = new double[_k, cols];

    for (int i = 0; i < _k; i++)
    {
        int randomRow = rand.Next(rows);
        for (int col = 0; col < cols; col++)
        {
            centroids[i, col] = data[randomRow, col];
        }
    }

    return centroids;
}

private int AssignCluster(double[,] data, int rowIndex)
{
    double minDistance = double.MaxValue;
    int cluster = -1;

    for (int i = 0; i < _k; i++)
    {
        double distance = CalculateDistance(data, rowIndex, _centroids, i);
        if (distance < minDistance)
        {
            minDistance = distance;
            cluster = i;
        }
    }

    return cluster;
}

private double CalculateDistance(double[,] data, int rowIndex, double[,] centroids, int
centroidIndex)
{
    double distance = 0.0;
    int cols = data.GetLength(1);

    for (int col = 0; col < cols; col++)
    {
        distance += Math.Pow(data[rowIndex, col] - centroids[centroidIndex, col], 2);
    }

    return Math.Sqrt(distance);
}

private double[,] UpdateCentroids(double[,] data, int[] labels)
{

```

```

int rows = data.GetLength(0);
int cols = data.GetLength(1);
double[,] centroids = new double[_k, cols];
int[] clusterCounts = new int[_k];

for (int i = 0; i < rows; i++)
{
    int cluster = labels[i];
    clusterCounts[cluster]++;
    for (int col = 0; col < cols; col++)
    {
        centroids[cluster, col] += data[i, col];
    }
}

for (int i = 0; i < _k; i++)
{
    if (clusterCounts[i] > 0)
    {
        for (int col = 0; col < cols; col++)
        {
            centroids[i, col] /= clusterCounts[i];
        }
    }
}

return centroids;
}

private bool HasConverged(double[,] oldCentroids, double[,] newCentroids)
{
    double tolerance = 1e-4;
    int rows = oldCentroids.GetLength(0);
    int cols = oldCentroids.GetLength(1);

    for (int i = 0; i < rows; i++)
    {
        for (int col = 0; col < cols; col++)
        {
            if (Math.Abs(oldCentroids[i, col] - newCentroids[i, col]) > tolerance)
            {
                return false;
            }
        }
    }

    return true;
}

public static List<int> DetectAnomalies(double[,] data, int[] labels, double threshold)
{
    List<int> anomalies = new List<int>();
    int rows = data.GetLength(0);

    for (int i = 0; i < rows; i++)
    {
        double distance = CalculateDistance(data, i, _centroids, labels[i]);
        if (distance > threshold)
    }

```

```

        {
            anomalies.Add(i);
        }
    }

    return anomalies;
}

public static List<int> IdentifyBehavioralAnomalies(double[,] data, int[] labels, double[,]
centroids, double threshold)
{
    List<int> anomalies = new List<int>();
    int rows = data.GetLength(0);

    for (int i = 0; i < rows; i++)
    {
        double distance = CalculateDistance(data, i, centroids, labels[i]);
        if (distance > threshold)
        {
            anomalies.Add(i);
        }
    }

    return anomalies;
}

private static double CalculateDistance(double[,] data, int rowIndex, double[,] centroids, int
clusterIndex)
{
    double distance = 0.0;
    int cols = data.GetLength(1);

    for (int col = 0; col < cols; col++)
    {
        distance += Math.Pow(data[rowIndex, col] - centroids[clusterIndex, col], 2);
    }

    return Math.Sqrt(distance);
}

public async Task<List<string>> AnalyzeTextContextAsync(List<string> texts, GptApiClient
gptClient)
{
    List<string> results = new List<string>();

    foreach (string text in texts)
    {
        string result = await gptClient.GetCompletionAsync(text, 200);
        results.Add(result.Trim());
    }

    return results;
}

public static List<AnomalyReport> GenerateAnomalyReports(
    List<int> behavioralAnomalies,
    List<string> contextAnalysisResults,
    string[] userIds)
{
    List<AnomalyReport> reports = new List<AnomalyReport>();

```

```

foreach (int anomalyIndex in behavioralAnomalies)
{
    reports.Add(new AnomalyReport
    {
        UserId = userIds[anomalyIndex],
        ContextResult = contextAnalysisResults[anomalyIndex],
        IsCritical = contextAnalysisResults[anomalyIndex].Contains("confidential") ||
                    contextAnalysisResults[anomalyIndex].Contains("sensitive")
    });
}

return reports;
}

public class AnomalyReport
{
    public string UserId { get; set; }
    public string ContextResult { get; set; }
    public bool IsCritical { get; set; }
}

public class IntegrationModule
{
    private readonly KMeans _kMeans;
    private readonly GptApiClient _gptClient;

    public IntegrationModule(KMeans kMeans, GptApiClient gptClient)
    {
        _kMeans = kMeans;
        _gptClient = gptClient;
    }

    public async Task<List<AnomalyReport>> ProcessDataAsync(double[,] data, string[] userIds,
List<string> textData, double threshold)
    {
        // Етап 1: Виконання кластеризації
        var labels = _kMeans.Fit(data);
        var anomalies = IdentifyBehavioralAnomalies(data, labels, _kMeans.GetCentroids(),
threshold);

        // Етап 2: Контекстний аналіз текстових даних
        var contextResults = await _gptClient.AnalyzeTextContextAsync(textData);

        // Етап 3: Генерація звітів про аномалії
        return GenerateAnomalyReports(anomalies, contextResults, userIds);
    }

    private List<int> IdentifyBehavioralAnomalies(double[,] data, int[] labels, double[,] centroids,
double threshold)
    {
        List<int> anomalies = new List<int>();

        for (int i = 0; i < data.GetLength(0); i++)
        {
            double distance = CalculateDistance(data, i, centroids, labels[i]);
            if (distance > threshold)
            {
                anomalies.Add(i);
            }
        }
    }
}

```

```

    }

    return anomalies;
}

private double CalculateDistance(double[,] data, int rowIndex, double[,] centroids, int
clusterIndex)
{
    double distance = 0.0;
    int cols = data.GetLength(1);

    for (int col = 0; col < cols; col++)
    {
        distance += Math.Pow(data[rowIndex, col] - centroids[clusterIndex, col], 2);
    }

    return Math.Sqrt(distance);
}
private List<AnomalyReport> GenerateAnomalyReports(List<int> anomalies, List<string>
contextResults, string[] userIds)
{
    List<AnomalyReport> reports = new List<AnomalyReport>();

    foreach (int anomalyIndex in anomalies)
    {
        string contextResult = contextResults[anomalyIndex];
        reports.Add(new AnomalyReport
        {
            UserId = userIds[anomalyIndex],
            ContextResult = contextResult,
            IsCritical = contextResult.Contains("confidential") || contextResult.Contains("sensitive")
        });
    }

    return reports;
}

public class AnomalyReport
{
    public string UserId { get; set; }
    public string ContextResult { get; set; }
    public bool IsCritical { get; set; }
}

```

Додаток В. Ілюстративний матеріал

«Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій і контекстного аналізу даних»

Виконав: ст. 2 курсу, групи КІТС-23М Шевиркін В.Ю.

Керівник: д.т.н., доц. каф. МБІС Грицак А.В.



ВСТУП

- + У сучасному інформаційному суспільстві захист даних є однією з найважливіших задач, з якою стикаються організації та підприємства. Зростаючий обсяг цифрової інформації, посилення регуляторних вимог і збільшення кількості кіберзагроз підкреслюють необхідність використання сучасних підходів до запобігання витокам конфіденційної інформації. У цьому контексті розробка систем запобігання витоку даних (DLP), які інтегрують новітні методи машинного навчання, є актуальним напрямом досліджень.

Актуальність та мета

- + Сучасні методи захисту інформації часто не враховують складність контексту даних та поведінкових факторів користувачів. Це призводить до того, що традиційні системи DLP мають високий рівень хибнопозитивних спрацювань та не здатні виявляти приховані загрози. Використання трансформерних моделей, таких як GPT, для контекстного аналізу тексту, у поєднанні з алгоритмами кластеризації, зокрема K-Means, дозволяє забезпечити високий рівень точності у виявленні загроз, враховуючи як текстову, так і поведінкову складову. Такий підхід є актуальним для створення сучасних систем захисту, здатних адаптуватися до нових загроз і забезпечувати мінімізацію ризиків витоку даних.
- + Метою дослідження є розробка та впровадження вдосконаленої системи DLP, яка базується на використанні трансформерних моделей для контекстного аналізу текстових даних і алгоритмів кластеризації для аналізу поведінкових даних користувачів.

Класифікація загроз витоку даних



Класифікація загроз витоку даних



Аналіз сучасних моделей машинного навчання для використання в DLP-системах

Тип моделі	Призначення	Переваги	Недоліки	Приклади використання в DLP
Трансформери і моделі (BERT, GPT)	Контекстуальний аналіз тексту, розуміння складних залежностей між словами	Висока точність і контекстуальне розуміння, ефективні для довгих і складних текстів	Високі обчислювальні витрати, потреба в великому обсязі навчальних даних	Розпізнавання конфіденційної інформації в текстових документах, електронних листах і чатах
Моделі на основі Attention	Визначення ключових слів і фраз у довгих текстах	Здатність аналізувати важливу інформацію навіть у великих документах	Можуть втратити контекст на рівні документа, потребують складних налаштувань	Виділення ключових слів і фраз для подальшого розпізнавання конфіденційної інформації
Кластеризація (K-Means, DBSCAN)	Групування подібних записів користувачів для виявлення аномалій	Простота і швидкість роботи, не потребують навчальних даних	Низька ефективність при складних і сильно неоднорідних даних	Аналіз активності користувачів для виявлення підозрілих дій
Автокодуючі машини (Autoencoders)	Виявлення аномалій у поведінці, відхилення від звичайних шаблонів	Можливість виявлення аномальних дій без досконалих даних, добре обробляють аномальні дані	Чутливі до гіперпараметрів, потребують налаштування для кожного типу даних	Моніторинг дій співробітників для виявлення відхилень від стандартної поведінки
CNN (Convolutional Neural Networks)	Розпізнавання зображень, графічних файлів, аналіз тексту у зображеннях	Висока точність для обробки зображень, здатність виділяти важливі елементи в зображеному контенті	Чутливість до якості зображень, потребують великих обчислювальних ресурсів	Виявлення конфіденційної інформації на зображеннях, сканах документів
OCR (Оптичне розпізнавання тексту)	Розпізнавання тексту на зображеннях та сканованих документах	Точність розпізнавання тексту у зображеннях; можливість вилучення тексту з графічного контенту	Складність при розпізнаванні низькоякісних зображень або складних шрифтів	Аналіз тексту у зображеннях для запобігання витоку конфіденційної інформації через зображення

Основні принципи та підходи до реалізації систем **DLP**

Системи запобігання витоку даних (DLP) розробляються для захисту конфіденційної інформації від несанкціонованого доступу та передачі. Основними принципами таких систем є забезпечення цілісності, конфіденційності та доступності даних.

Ідентифікація конфіденційної інформації: DLP-системи використовують ключові слова, шаблони, контекстуальний аналіз та класифікацію для визначення даних, що потребують захисту.

Моніторинг і контроль передачі даних: Вони аналізують вихідну інформацію через мережу, електронну пошту чи зовнішні носії для виявлення витоку.

Аналіз поведінки користувачів: Використання алгоритмів кластеризації та поведінкової аналітики для виявлення підозрілої активності.

Контекстуальний аналіз: Трансформерні моделі, такі як GPT, забезпечують глибше розуміння тексту для точного розпізнавання конфіденційних даних.

Інтеграція політик безпеки: Системи дозволяють автоматично застосовувати політики для обмеження доступу, блокування передачі або сповіщення адміністраторів у разі порушень.

Підходи до реалізації DLP включають комплексне використання моделей машинного навчання для аналізу даних, інтеграцію з існуючими інфраструктурами безпеки та регулярне оновлення правил і шаблонів для адаптації до нових загроз. Ці принципи забезпечують високий рівень захисту інформації в умовах сучасних кіберзагроз.

Алгоритм інтеграції моделі **gpt**



Розроблений алгоритм
кластеризації з
використанням **k-means**
для виявлення аномалій і
контекстного аналізу
даних



MACHINE LEARNING



Використані технології

```
[2024-11-07 10:00:00] INFO: Starting unit tests for normalization function...  
[2024-11-07 10:00:01] SUCCESS: TestNormalization passed successfully.  
[2024-11-07 10:00:01] INFO: Starting unit tests for cluster assignment function...  
[2024-11-07 10:00:02] SUCCESS: TestClusterAssignment passed successfully.
```

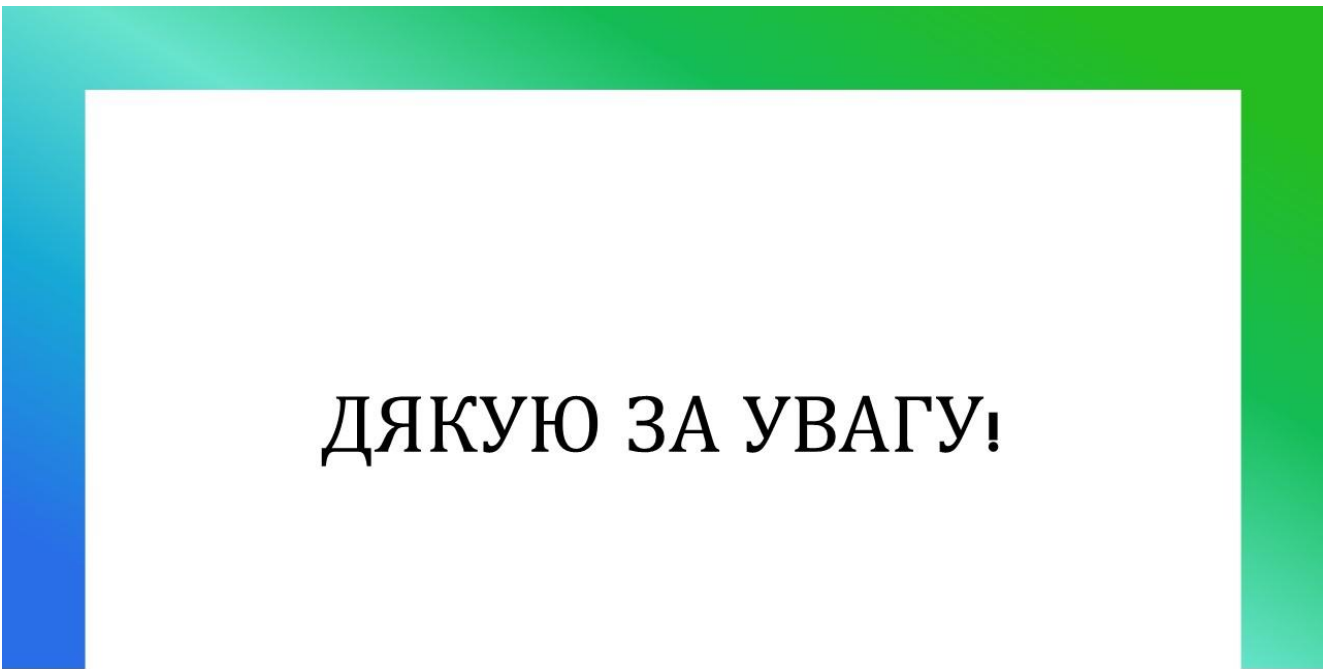
```
[2024-11-07 10:00:10] INFO: Starting performance tests under heavy load...  
[2024-11-07 10:00:30] WARNING: Performance under load test slightly exceeded the threshold (5100 ms).  
[2024-11-07 10:00:30] INFO: Adjusting parameters for optimization...  
[2024-11-07 10:00:40] SUCCESS: Performance test passed after optimization (4800 ms).  
[2024-11-07 10:00:40] INFO: All tests completed. No critical issues found.
```

```
[2024-11-07 10:00:02] INFO: Starting integration tests between modules...  
[2024-11-07 10:00:05] SUCCESS: TestIntegrationBetweenModules passed successfully.  
[2024-11-07 10:00:05] INFO: Running system test for full workflow...  
[2024-11-07 10:00:10] SUCCESS: TestFullSystemWorkflow passed successfully.
```

Тестування

Висновок

- + У ході виконання магістерської роботи було розроблено та впроваджено вдосконалену систему запобігання витоку даних (DLP), яка базується на використанні сучасних моделей машинного навчання та алгоритмів кластеризації. Головною метою дослідження було створення комплексного рішення, здатного ідентифікувати та запобігати загрозам витоку конфіденційної інформації, забезпечуючи точність аналізу та швидкість реагування
- + Вдосконалена система DLP продемонструвала ефективність у виявленні та запобіганні витоку даних завдяки використанню сучасних підходів до аналізу інформації. Зокрема, інтеграція трансформерних моделей GPT дозволила підвищити точність розпізнавання конфіденційної інформації навіть у складних текстових контекстах, а алгоритм K-Means забезпечив можливість аналізу поведінкових патернів користувачів.



ДЯКУЮ ЗА УВАГУ!

Додаток Г. Протокол перевірки на антиплагіат

ПРОТОКОЛ ПЕРЕВІРКИ НАВЧАЛЬНОЇ (КВАЛІФІКАЦІЙНОЇ) РОБОТИ

Назва роботи:

Підвищення захищеності системи DLP запобігання витоку даних на основі Transformer-based моделі GPT і алгоритму кластеризації з використанням K-Means для виявлення аномалій і контекстного аналізу даних

Тип роботи: магістерська кваліфікаційна робота

Підрозділ Кафедра менеджменту на безпеки інформаційних систем

Факультет менеджменту та інформаційної безпеки

Гр. КІТС-23м

Керівник доцент Грицак А.В.

Показники звіту подібності

Turnitin	
Оригінальність	96%
Загальна схожість	4%

Аналіз звіту подібності (відмітити потрібне)

- ☐ Запозичення, виявлені у роботі, оформлені коректно і не містять ознак плагіату.
- ☐ Виявлені у роботі запозичення не мають ознак плагіату, але їх надмірна кількість викликає сумніви щодо цінності роботи і відсутності самостійності її автора. Роботу направити на доопрацювання.
- ☐ Виявлені у роботі запозичення є недобросовісними і мають ознаки плагіату та/або в ній містяться навмисні спотворення тексту, що вказують на спроби приховування недобросовісних запозичень.

Заявляю, що ознайомлений (-на) з повним звітом подібності, який був згенерований Системою щодо роботи (додається)

Автор С.А.М.
(підпис)

Шевиркін В.Ю.
(прізвище, ініціали)

Опис прийнятого рішення

- ☐ Запозичення, виявлені у роботі, оформлені коректно і не містять ознак плагіату.

Особа, відповідальна за перевірку

(підпис)

Коваль Н.П.

(прізвище, ініціали)

Експерт

(за потреби)

(підпис)

(прізвище, ініціали, посада)