

НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ
імені ІГОРЯ СІКОРСЬКОГО»
ФІЗИКО-ТЕХНІЧНИЙ ІНСТИТУТ
Кафедра інформаційної безпеки

До захисту допущено
В.о. завідувача кафедри

_____ Дмитро ЛАНДЕ
(підпис)

« _____ » _____ 2022 р.

Дипломна робота
на здобуття ступеня бакалавра
за освітньо-професійною програмою «Системи, технології та математичні
методи кібербезпеки»
спеціальності: 125 «Кібербезпека»

на тему: «Захист організацій від витоку інформації з використанням систем DLP»

Виконав: здобувач вищої освіти **IV** курсу, групи ФБ-82
(шифр групи)

Кравчук Владислав Сергійович
(прізвище, ім'я, по батькові)


(підпис)

Керівник

к.т.н., доцент Барановський Олексій Миколайович
(посада, науковий ступінь, вчене звання, прізвище, ім'я, по батькові)


(підпис)

Рецензент

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

_____ (підпис)

Засвідчую, що у цій дипломній роботі немає
запозичень з праць інших авторів без
відповідних посилань.

Здобувач вищої освіти


(підпис)

Київ - 2022 року

НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ
імені ІГОРЯ СІКОРСЬКОГО»
ФІЗИКО-ТЕХНІЧНИЙ ІНСТИТУТ
Кафедра інформаційної безпеки

Рівень вищої освіти – перший (бакалаврський)

Спеціальність – 125 «Кібербезпека»

Освітньо-професійна програма «Системи, технології та математичні методи кібербезпеки»

ЗАТВЕРДЖУЮ

В.о. завідувача кафедри

_____ Дмитро ЛАНДЕ
(підпис)

«__» _____ 2022 р.

ЗАВДАННЯ
на дипломну роботу здобувачу вищої освіти

Кравчук Владислав Сергійович
(прізвище, ім'я, по батькові)

1. Тема роботи «Захист організацій від витоку інформації з використанням систем DLP», керівник роботи Барановський Олексій Миколайович к.т.н., доцент, _____ затверджені наказом по університету від «__» _____ 2022 р. № _____
2. Термін подання здобувачем вищої освіти роботи 13 червня 2022 р.
3. Вихідні дані до роботи:
4. Зміст роботи: стани даних, типи та режими роботи DLP, класифікація даних, суб'єкти загроз, огляд статистики збитків, проблематики та рішення, опис ідеї, методи навчання.
5. Перелік ілюстративного матеріалу (із зазначенням плакатів, презентацій тощо) Захист організацій від витоку інформації з використанням систем DLP - презентація
6. Дата видачі завдання 29.03.2022

Календарний план

№ з/п	Назва етапів виконання дипломної роботи	Термін виконання етапів дипломної роботи	Примітка
1	Отримання завдання	29.03.2022	виконано
2	Дослідження напрямку DLP	30.03.2022 – 10.04.2022	виконано
3	Визначення функціоналу рішень	10.04.2022 – 15.04.2022	виконано
4	Дослідження методів класифікації	15.04.2022 – 20.04.2022	виконано
5	Вивчення проблематики та рішень	21.04.2022 – 1.05.2022	виконано
6	Проходження переддипломної практики	02.05.2022 – 29.05.2022	виконано
7	Написання дипломної роботи	11.06.2022 – 13.06.2022	виконано
8	Отримання допуску до захисту	14.06.2022	
9	Захист дипломної роботи	20.06.2022	

Здобувач вищої освіти



(підпис)

Владислав Кравчук
(Власне ім'я, ПРІЗВИЩЕ)

Керівник роботи



(підпис)

Олексій Барановський
(Власне ім'я, ПРІЗВИЩЕ)

РЕФЕРАТ

Обсяг роботи 68 сторінок, 22 ілюстрації, 1 таблиця, 12 джерел літератури.

В цій роботі був розглянутий напрям DLP в контексті захисту організації. Був розкритий загальний підхід до захисту організації на базі стратегії DiD. Проаналізована статистика фінансових збитків підприємств по всьому світу. Розкриті основні поняття напрямку DLP. Досліджені основні проблематики та їх рішення.

Був запропонований підхід для покращення методу класифікації даних завдяки інтеграції системи DLP із технологією машинного навчання. Машинне навчання виступає в якості додаткової перевірки документів із метою правильної класифікації. Як результат – зменшення False-positive спрацювань.

Цей підхід може бути розглянутий та випробуваний зацікавленими сторонами пов'язаними із DLP продуктами з метою покращення класифікації.

Ключові слова: виток даних, DLP, конфіденційний документ, класифікація, False-positive, машинне навчання, організація, рішення, проблематика.

ABSTRACT

The volume of work is 68 pages, 22 illustrations, 1 table, 12 sources of literature.

In this paper, was considered DLP direction in the context of protection of organization. A general approach to protecting the organization based on the DiD strategy was revealed. The statistics of financial losses of enterprises around the world were analyzed. The main concepts of the DLP direction are revealed. The main problems and their solutions are studied.

An approach was proposed to improve the method of data classification by integrating the DLP system with machine learning technology. Machine learning serves as an additional check of documents for the purpose of correct classification. The result is a reduction in False-positive triggers.

This approach can be reviewed and tested by DLP-related stakeholders to improve classification.

Keywords: data breach, DLP, confidential document, classification, False-positive, machine learning, organization, solutions, problems.

**ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ,
СКОРОЧЕНЬ І ТЕРМІНІВ**

DLP – Data Leak Prevention

DiD – Defense in Depth

DPI – Deep Packet Inspection

Labeling – Process/Method of Data Classification

TLS – Transport Layer Security

ACL – Access Control Lists

DAR – Data at Rest

DIM – Data in Motion

DIU – Data in Use

ITAM – IT Asset Management

ML – Machine Learning

NIST – National Institute of Standards and Technologies

ISO – International Organization for Standardization

ROI – Reverse of Investments

PII – Personal Identifiable Information

PHI – Protected Healthcare Information

PCI DSS - Payment Card Industry Data Security Standard

Зміст

Перелік умовних позначень, символів, одиниць, скорочень і термінів	6
Вступ.....	7
1. Загальний підхід до захисту організації.....	11
1.1 Ключові поняття DLP та особливості впровадження.....	16
1.2 Класифікація даних.....	28
1.3 Вибір продукту та партнера.....	31
1.3.1 Вибір DLP системи.....	31
1.3.2 План проведення Пілотного проекту.....	32
1.3.3 Рекомендації щодо планування розгортання	34
1.3.4 Приклад технічного завдання	34
Висновок до розділу 1.....	39
2. Сучасні положення проблематики та _рішення_.....	40
2.1 Огляд статистики збитків.....	40
2.2 Проблема класифікації даних та False-positive спрацювання.....	44
2.3 Порівняння методів класифікації даних машинним навчанням.....	45
Висновок до розділу 2.....	50
3. Експериментальний метод класифікації даних	50
3.1 Інтегрованість машинного навчання з рішенням DLP. Опис ідеї	50
3.2 Реалізація запропонованого методу. Методи навчання ML	53
Висновок до розділу 3.....	58
Висновок.....	59
Перелік джерел посилань	60
Додаток Програмна реалізація.....	62

Вступ

З моменту створення та розвитку інформаційних систем та їх інтеграції в життя суспільства по всьому світу, всі конфіденційні, і не тільки, дані стали зберігатися в цифровому вигляді. Неможливо заперечити той факт, що дані будь-якої чутливості повинні знаходитися в безпеці. Забезпеченням цієї самої безпеки і займаються фахівці з інформаційної безпеки.

В наші часи інформаційна безпека являється невід'ємною частиною будь-якої організації незалежно від її розміру або прибутку. Кожен свідомий керівник має приділяти значну частину уваги та фінансів задля захисту свого бізнесу, адже зловмисники мають можливість завдати шкоди, через яку компанія може понести фінансові та репутаційні збитки або взагалі перестане існувати.

Компанії по всьому світу щорічно втрачають мільйони доларів через вразливості людського фактору та недосконалість інформаційних систем. Тому вкрай важливо впроваджувати заходи по забезпеченню безпеки задля унеможливлення несанкціонованих маніпуляцій над вашою інфраструктурою.

З кожним роком зростає кількість атак на інформаційні системи через їх недосконалість. Нажаль людство не здатне створювати рішення які б забезпечували 100% стійкість до помилок, випадкових або ж навмисних. Відсутність помилок в технологіях можливо лише при відсутності цих технологій, але в наші часи це неможливо.

Захист організації – це постійний безперервний процес. Це комплексний, і, місцями, творчий підхід. Це велика та складна система до якої потрібно підходити правильно та всебічно. Не існує єдиного рішення яке б закривало усі питання безпеки, і це важливо розуміти. Усі заходи стосовно безпеки направлені на зниження ризиків виникнення інцидентів, а не на їх повне позбавлення.

В своїй роботі я маю намір розкрити тему витоку інформації зсередини організацій. DLP - це частина загальної стратегії захисту інформації. Цей напрям

завжди був, є і буде актуальним, оскільки існують не тільки атаки ззовні, але й атаки зсередини, які так само можуть привести до катастрофічних наслідків для компанії.

Вважаю необхідним наголосити на тому, що DLP – це напрям який стосується не лише апаратно-технічного захисту інформації в кіберпросторі. Наявність охоронця та камер відеоспостереження на території будівлі, які слідкують за тим, щоб не було фізичних крадіжок будь-чого, що має цінність – це також DLP. Охоронець, та системи контролю доступу до приміщень/територій та камери – це рішення.

В цій роботі DLP буде розглядатися як технологія, і як рішення і як продукт. Тому що коли ми говоримо про впровадження DLP в архітектуру організації, ми маємо на увазі загально технологію, розвинення напрямку. Коли розмова йде більш предметна і ми розглядаємо конкретні випадки витоків та протидії ним, то DLP розглядатиметься як рішення.

Протягом роботи буде представлений загальний підхід до захисту організації. Буде приведена статистика витоку даних по всьому світу у порівнянні з минулими роками, тим самим наголошено на необхідності звернути увагу на цей напрям. Буде розкритий функціонал рішень DLP. Розказано в яких станах дані розглядають в контексті DLP. Наголошено на тому, Чим і Як дані класифікуються та якими ці дані бувають в контексті їх важливості для підприємства. Розберемо Хто та Що представляє загрозу для організації в розрізі витоку даних. Будуть представлені та розібрані сучасні проблематики, з якими бізнес стикається, та їх вирішення. Також буде представлений квадрант Radicati задля того, щоб сформулювати бачення того, які рішення існують на ринку та які постачальники є лідерами. В якості ілюстрацій будуть продемонстровані порівняльні характеристики продуктів. За допомогою них представники підприємства замовника зможуть підібрати продукт під свої потреби.

Успіх функціонування технічних рішень DLP ґрунтується на двох речах:

- DPI – технологія аналізу мережевих пакетів на предмет вмісту в них даних, в тому числі і конфіденційних;
- Labeling – технологія маркування документів спеціальними мітками, які визначають критичність документа;

Також потрібно наголосити на підвищенні обізнаності в області інформаційної безпеки, яка є складовою частиною комплексу заходів із забезпечення інформаційної безпеки організації та спрямована на формування культурних особливостей кібербезпеки персоналу. Це знання про наявність атак людського фактору з метою отримання зловмисником доступу до даних та їх компрометації, але це не технічне рішення, а питання адміністративного контролю. Це вважається обов'язковим, оскільки більшість (десь 2/3) усіх випадків витоку даних відбувається саме через необачність персоналу, або через вплив зловмисників на персонал використовуючи методи соціальної інженерії тощо. Проведення навчання персоналу допоможе знизити ризики збитків через вплив людського фактору.

Вирішальним критерієм успішності буде налаштування політик, по яким DLP буде спрацьовувати, конкретно під потреби. Рішення є доволі гнучкими, тому потрібно правильно налаштувати систему із урахуванням бізнес-потреб компанії.

Актуальність роботи – з кожним роком кількість витоків даних збільшується. Постійно та примусово збільшується рівень зрілості організацій, через що мотивація впроваджувати рішення захисту стає все більш актуальною.

Мета роботи – дослідження сучасних проблематики та їх рішень, розробка ідеї щодо збільшення ефективності класифікації даних.

Завданням роботи – вивчення напряму DLP, проаналізувати існуючі проблематики та їх рішення, зібрати статистику витоку даних за останні роки,

визначити лідерів на ринку серед постачальників, розробити концепцію ефективної класифікації.

Об’єктом дослідження – є технологія DLP.

Предметом дослідження – є архітектура безпеки організації, існуючі проблематики та рішення.

Методами дослідження – аналіз інформаційних джерел, наукових робіт та літератури стосовно обраного напрямку. Консультації із фахівцями.

Наукова новизна – створення моделі підходу до ефективного маркування даних із використанням машинного навчання інтегрованого із системою DLP. Як результат – зменшення False Positive спрацювань та більш ефективне маркування.

1. Загальний підхід до захисту організації від витоку інформації

Інформація є одним із найбільш цінних і водночас вразливих активів будь-якої компанії, тому кожне підприємство має забезпечити належний рівень інформаційної безпеки. Ефективність бізнесу в багатьох випадках залежить від збереження конфіденційності, цілісності та доступності інформації. В даний час однією з найбільш актуальних загроз у сфері інформаційної безпеки є витік конфіденційних даних. Сьогодні більшість підприємств використовують багаторівневі системи обробки інформації – комп'ютери, хмарні сховища, корпоративні мережі і т. п. Всі ці системи не тільки передають дані, але і є середовищем їх можливого витоку.

Витік інформації - це неконтрольоване розповсюдження інформації, що виходить, за межі організації або кола осіб, котрим ці відомості були довірені, в результаті її розголошення або несанкціонованого доступу до неї.

Витік інформації можна розділити на два типи:

1. Навмисний витік характеризуються тим, що зловмисник, маючи доступ до інформації, що захищається, розуміє протиправність своїх дій і усвідомлює їх можливі негативні наслідки. Основними мотивами для здійснення подібного роду злочинів є в більшості випадків матеріальна нажива або отримання іншої вигоди. Крім того, крадіжка критичних даних хакерами, які обманним або іншим шляхом отримали прямий доступ до них за допомогою співробітника організації, теж відноситься до навмисних витоків, навіть якщо протиправних намірів у співробітника не було.

2. Ненавмисний витік пов'язаний з халатністю або недостатньою обізнаністю користувачів в своїх функціональних обов'язках. Найбільш поширені помилки працівників організації, що ведуть до витоків, представлені на рисунку 1.1.



Рис. 1.1. Поширені помилки користувачів, що ведуть до ненавмисних витоків даних

Будь-який витік інформації несе в собі ті чи інші негативні економічні наслідки для компанії. Особливо значних фінансових збитків компанія може зазнати у результаті витоку інформації, що становить комерційну таємницю.

Комерційна таємниця – інформація про організацію діяльності підприємства, технології розробки продукції, дані про грошові потоки, інтелектуальна власність та інші відомості, володіючи якими отримуються фінансові вигоди.

Причини витоку інформації :

– **Персонал.** Кожен співробітник підприємства є потенційною загрозою для безпеки інформації. Часто люди забирають роботу додому – переміщують робочі файли на свої флеш-носії, передають їх по незахищеним

каналам з'єднання, обговорюють інформацію зі співробітниками конкуруючих компаній.

– **Проблеми підбору кадрів.** Часта зміна персоналу, масштабні зміни в організації роботи компанії, зниження заробітних плат, скорочення співробітників – все це є частиною «плинності» кадрів. Таке явище часто стає причиною витоку секретної інформації. Криза, нестача коштів для видачі зарплат змушують керівництво погіршувати умови роботи персоналу. В результаті підвищується невдоволення працівників, які можуть звільнитися або ж просто почати поширювати секретні дані конкурентам. Проблема зміни персоналу особливо важлива для керівних посад, адже всі керуючі повинні мати доступ до секретної документації.

– **Відрядження.** Робочий процес фірми включає ділові зустрічі, поїздки в інші філії компанії, країни. Співробітники, які часто виїжджають у відрядження, можуть ненавмисно стати основною причиною витоку секретної інформації підприємства. У поїздці такий працівник завжди має при собі особистий або корпоративний ноутбук/смартфон, що обробляє захищені документи. Техніка може бути залишена в громадському місці, зламана або викрадена. Якщо за співробітником ведеться стеження або ж він зустрічається з керівниками конкуруючої компанії, загублений ноутбук може стати головним джерелом розголошення службової інформації.

– **Співпраця з іншими компаніями.** Більшість автоматизованих систем захисту здатні обмежити доступ до службової інформації тільки в рамках однієї будівлі або одного підприємства (якщо кілька філій використовують загальний сервер зберігання даних). У процесі спільного виконання проєкту декількома фірмами, служби безпеки не можуть в повній мірі простежити за тим, як реалізується доступ до службової таємниці кожного з підприємств.

– **Використання складних ІТ-інфраструктур.** Великі корпорації використовують комплексні системи захисту службових відомостей. Автоматизовані системи мають на увазі наявність декількох відділів безпеки і роботу понад п'яти системних адміністраторів, завдання яких полягає тільки в підтримці збереження комерційної таємниці. Складність системи теж є ризиком витоку, адже одночасна робота кількох людей може бути незлагодженою.

– **Несправність техніки.** Наприклад, помилки в роботі програмного забезпечення, збої в роботі серверного обладнання, несправність технічних засобів захисту.

– **Витік з технічних каналів передачі даних.**

Канал витоку даних – це фізичне середовище, всередині якого не контролюється поширення таємної інформації. На будь якому підприємстві, що використовує комп'ютери, серверні стійки, мережі, є канали витоку. З їх допомогою зломисник може отримати доступ до комерційної таємниці.

Технічний канал витоку інформації (ТКВІ) – сукупність джерела небезпечного сигналу, середовища поширення небезпечного сигналу та засобу технічної розвідки (рисунок 1.2)

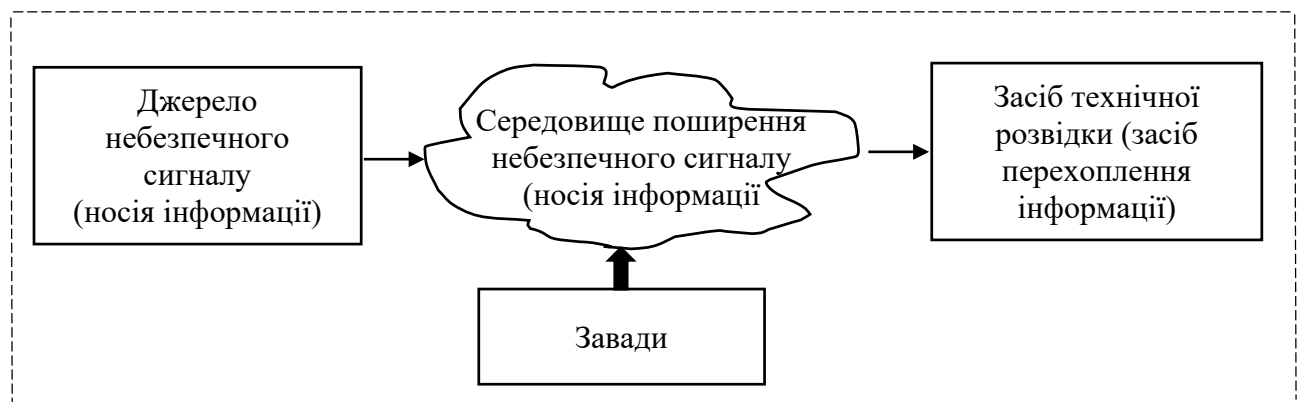


Рисунок 1.2 - Технічний канал витоку інформації

Отже, технічним каналом витоку інформації є фізичний шлях небезпечного сигналу (носія інформації) від джерела небезпечного сигналу до противника.

Існують наступні технічні канали витоку інформації:

1. **Мовний.** Конкуренти часто використовують прослуховувачі та інші закладки, за допомогою яких відбувається крадіжка таємниці.

2. **Віброакустичний.** Цей канал витоку виникає в процесі зіткнення звуку з архітектурними конструкціями (стінами, підлогою, вікнами). Вібраційні хвилі можна зчитати і перевести в мовної текст. За допомогою спрямованих мікрофонів на відстані до 200 метрів від приміщення зловмисник може зчитати розмову, в якій фігурує службова інформація.

3. **Електромагнітний.** В результаті роботи всіх технічних засобів виникає магнітне поле. Між апаратними елементами передаються сигнали, що можна вважати спеціальним обладнанням на великих відстанях і отримати секретні дані.

4. **Візуальний.** Приклад появи візуального каналу крадіжки – це проведення нарад і конференцій з неприкритими вікнами. З сусіднього будинку зловмисник може легко переглянути всю інформацію. Також можливі варіанти використання відеозакладок, які передають картинку того, що відбувається конкурентам.

Для захисту технічних каналів витоку рекомендується використовувати:

- Тепловізор. За допомогою такого девайса можна просканувати всі стіни і частини інтер'єру на наявність закладних пристроїв (жучків, відеокамер).
- Пристрої, що заглушають подачу сигналу радіочастот.
- Засоби захисту архітектурних конструкцій - ущільнювачі для вікон, дверей, підлоги і стелі. Вони ізолюють звук і унеможливають зчитування вібраційних хвиль з поверхні будівлі.

– Пристрої для екранування і зашумлення. Вони використовуються для захисту електромагнітного каналу витоку.

Також слід заземлити всі комунікації, що виходять за межі приміщення і контрольованої зони (труби, кабелі, лінії зв'язку).

1.1. Ключові поняття DLP та особливості впровадження

Розробка та впровадження окремих систем для контролю за обігом конфіденційної інформації та запобігання її витоку – одне із важливих завдань кожного підприємства. Ці системи відрізняються своїми можливостями та функціональністю, але всі вони узагальнюються аббревіатурою **DLP (Data Leak Prevention)**, яка запропонована агентством Forrester в 2005 році. Якщо брати національну термінологію, то це «системи запобігання витоку інформації».

Основне призначення DLP-систем – забезпечувати захист від випадкового або навмисного поширення конфіденційної інформації з боку співробітників, що мають доступ до неї в силу своїх службових обов'язків.

Основними завданнями DLP-систем є:

- формалізація опису даних, що захищаються (налаштування системи);
- розпізнавання даних, що захищаються, у вихідному потоці з внутрішньої інформаційної мережі компанії (розпізнавання дій, спрямованих на переміщення конфіденційної інформації);
- реагування на виявлені спроби витоку даних та формування доказової бази для розслідування інцидентів.

Основні функції DLP-систем:

- контроль передачі інформації через Інтернет з використанням e-mail, http, https, ftp і інших додатків і протоколів;
- контроль збереження інформації на зовнішні носії та мобільні пристрої;
- захист інформації від витоку шляхом контролю друк даних;

- блокування спроб пересилання або збереження конфіденційних даних, інформування адміністраторів ІБ про інциденти, спроби створення копій;
- пошук конфіденційної інформації на робочих станціях і файлових серверах за ключовими словами, мітками документів, атрибутами файлів і цифровими відбитками;
- запобігання витоку інформації шляхом контролю життєвого циклу і руху конфіденційних відомостей.

В результаті використання DLP-систем можуть бути вирішені питання:

- запобігання витокам і несанкціонованій передачі конфіденційної інформації;
- мінімізації ризиків фінансового і репутаційного збитку;
- підвищення дисципліни користувачів;
- розслідування інцидентів та їх наслідків;
- ліквідації загроз безпеки персональних даних, відповідності вимогам щодо захисту персональних даних.

В контексті DLP дані розглядаються в трьох станах:

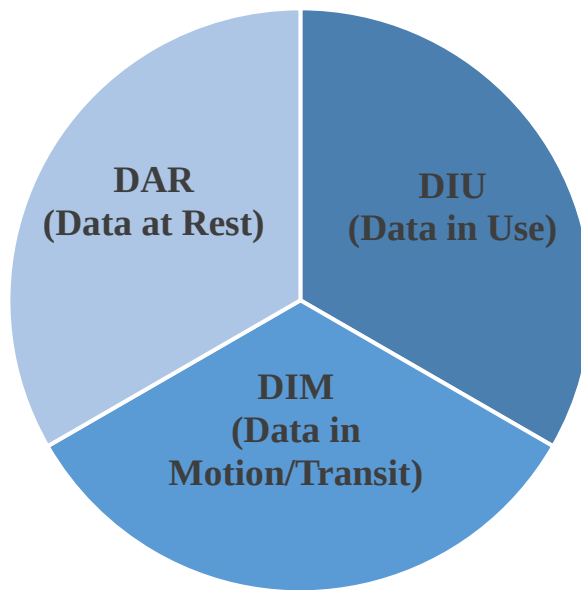


Рисунок 1.1.1 - Стан даних у контексті DLP

– **DAR (Data at Rest)** – дані у стані спокою. Це дані, які фізично зберігаються на носіях інформації таких як жорсткі диски, архівовані бази даних, usb пристрої, хмарні та інші інфраструктурні сховища тощо, і над цими даними не виконуються ніяких маніпуляцій. В цьому стані дані не передаються між мережами та пристроями.

– **DIU (Data in Use)** – це стан даних, при якому вони використовуються в даний момент. В цьому стані дані взаємодіють з людиною або іншим процесом. Наприклад: оновлюються, обробляються, змінюються, стираються тощо.

– **DIM (Data in Motion/Transit)** – дані у стані переміщення. В цьому стані дані переміщуються з одного місця в інше, наприклад між комп'ютерами, віртуальними машинами, між мережами, з кінцевої точки до хмари, на портативний пристрій, через інтернет, електронну пошту або месенджер тощо. По прибуттю в пункт призначення, дані зі стану переміщення (DIM) переходять у стан спокою (DAR).

Останнім часом вимоги до функціональних можливостей DLP-систем постійно зростають, що призводить до перетворення їх в один з найефективніших, комплексних і системних рішень в сфері захисту конфіденційної корпоративної інформації.

Сучасна DLP-система, являє собою розподілений програмно-апаратний комплекс, що складається з декількох модулів, які функціонують на виділених серверах, на робочих місцях співробітників компанії (персональних комп'ютерах, робочих станціях, інших пристроях) і на рівні внутрішньої служби безпеки:

– модулі бази даних, систематизації, аналізу та іншої обробки інформації, що стосується всіх інцидентів, виявлених системою, а також інших даних, відстеження і контроль яких закладені в систему;

– модулі пасивного і (або) активного спостереження, контролю за діями співробітників компанії (користувачів). Перелік дій, що відслідковуються і

контролюються системою, встановлюється заздалегідь і як правило носить обмежений характер. Стандартний перелік зазвичай включає контроль входу і виходу з системи, бездротової передачі даних, підключення зовнішніх пристроїв, на які може бути скопійована інформація, друку документів та інших процесів.

– модулі управління, моніторингу, налаштування системи, аналітичної роботи для потреб служби безпеки.

Кожен з модулів DLP-системи вирішує своє коло завдань. Розглянемо основні функціональні модулі DLP-систем.

Таблиця 1.1.1

Основні функціональні модулі DLP-систем

№	Найменування функціонального модулю DLP-систем	Опис
1	Контроль корпоративної пошти	В базах даних DLP можуть зберігатися все поштові повідомлення контрольованого користувача, незважаючи на те, що користувач може видалити отримане або відправлене повідомлення, а відновлення з резервної копії може займати значний час або ж резервування пошти може не бути зовсім. За адресатам листів може бути побудований відповідний звіт, який дасть розуміння, з ким і в якому обсязі взаємодіє користувач. Також здійснюється аналіз тексту листів на збіг з певними словоформами.
2	Контроль програм обміну повідомленнями	Всі повідомлення месенджерів контрольованого користувача можуть зберігатися в базах даних DLP. Згідно з політиками і правилами, що налаштовані в DLP, відбувається спрацьовування на певний контент, оператор DLP оповіщається про відповідні повідомлення. Незважаючи на те, що обмін між клієнтом і сервером месенджера може відбуватися за допомогою засобів шифрування месенджера, DLP-система отримує незашифровані дані, якщо вона підтримує відповідний месенджер.
3	Контроль друку	Модуль дозволяє контролювати файли, які користувач відправляє на друк, зберігаючи їх у базах даних

		DLP. Також може контролюватися обсяг друку та здійснюватися збір статистики по друку.
4	Контроль зовнішніх накопичувачів	DLP-система фіксує та розпізнає пристрої, які підключаються до робочого місця. Мається можливість задання переліку дозволених пристроїв і шифрування даних на них. У разі відповідного налаштування система здатна зробити копію всієї інформації, що зберігається на підключеному зовнішньому пристрої користувача, або тільки файлів, які зчитуються/записуються на зовнішній пристрій.
5	Контроль пристроїв вводу	Найчастіше це логгер клавіатури і миші, які при відповідних настройках записують в бази даних DLP все натискання на клавіатуру і рухи миші користувача. У деяких DLP-системах при цьому відбувається аналіз словоформ, і при відповідних настройках відбувається спрацьовування на певні слова (поєднання) з оповіщенням оператора DLP.
6	Контроль пристроїв відеовиводу	Дозволяє при відповідних настройках з необхідною періодичністю проводити знімки екрану монітора користувача із записом їх в бази даних DLP для подальшого перегляду та аналізу оператором.
7	Контроль роботи програмного забезпечення (ПЗ) робочого місця користувача	Веде фіксацію початку і закінчення використання ПЗ користувача. Може також контролювати завантаження ресурсів робочого місця користувача, дозволяє оцінити активність та ефективність роботи персоналу.
8	Контроль роботи користувача в браузері (http, https)	Модуль фіксує всі відвідувані сайти (або сторінки) та час, проведений на кожному сайті.
9	Контроль обміну файлами (ftp, sftp)	Виробляє запис копій переданих користувачем файлів з фіксацією адреси одержувача.
10	Аудіоконтроль роботи користувача (при наявності підключеного мікрофона)	Дозволяє вести запис з мікрофону, підключеного до робочого місця користувача за заданим розкладом або в центрі онлайн-контролю.
11	Відеоконтроль роботи користувача (при наявності)	Дозволяє налаштувати відеозапис з камери, підключеної до робочого місця. У більшості випадків використовується спільно з модулем аудіоконтролю.

	підключеної камери)	
--	---------------------	--

До складу більшості DLP-систем входять також центри: адміністратора, звітності, повідомлень, онлайн-контролю.

Існує два режими роботи систем DLP – активний і пасивний. Вибір на користь одного чи іншого залежить від потреб організації.

Система DLP в активному режимі роботи має можливість блокування трафіку, тим самим запобігаючи потенційному витоку інформації. Активний режим незалежний, він сам сканує трафік який проходить крізь нього і у разі виявлення інциденту буде виконана дія (блокування, пропускання, карантин) згідно з встановленими політиками. Головним недоліком активного режиму роботи систем DLP є False-positive спрацювання, через які потенційно можливе призупинення бізнес процесів організації. False-positive спрацювання – це помилкове виявлення порушення там, де його не було, через що абсолютно легальне повідомлення блокується системою автоматично. Сучасні системи DLP здатні до самонавчання з метою зменшення False-positive спрацювань, але для цієї калібровки потрібен час.

Пасивний режим роботи DLP можливості блокування трафіку не має. Зазвичай це окремий пристрій на який виконується зеркалювання трафіку (SPAN) з метою його аналізу на предмет інцидентів. Пасивний режим працює у зв'язці з іншими технологіями, наприклад SIEM, або ж з логами. Адміністратор може виявити конфіденційну інформацію в тілі листа, якщо це e-mail, та донести про інцидент до керівництва, але заборонити відправку листа в цьому режимі роботи він не здатен. У цьому режимі система збирає дані для їх аналітики та у результаті для формування звітів. Після чого ми здатні побачити Хто, Що, Коли, Куди відправив.

присвоєному даним. Адміністратори, відповідальні за напрям DLP, мають змогу створювати, змінювати та видаляти політики в залежності від потреб організації. Стосовно Мережевого DLP адміністратор може налаштувати систему на моніторинг певного протоколу або порту. Це може бути контроль: електронної пошти, відправки файлів, веб-сервісів, месенджерів, чатів, відправки документів на друк та інше. Далі, під час виявлення аномалії, знову ж таки в залежності від налаштованих політик, буде здійснено три можливі дії:

1. Блокування трафіку
2. Пропускання трафіку
3. Відправка до карантину

Мережевий DLP базується на DPI та Labeling. DPI – це розширений та поглиблений метод аналізу мережевого трафіку. Він дозволяє швидко та ефективно відфільтровувати пакети згідно з потребою. Його перевага полягає в тому, що він здатний аналізувати не лише заголовки пакетів, а ще й корисне навантаження пакету на предмет наявності в ньому даних, які згідно з встановленими політиками, там знаходитися не повинні.

DLP Кінцевих Точок. Зазвичай це Агент (програма), яка встановлюється на клієнтських пристрої такі як: ПК, телефони, ноутбуки, сервери, віртуальні машини, IoT та інші. Його задача, так само як і в Мережевому DLP, запобігти витоку чутливих для організації даних, але саме з пристрою, на якому цей клієнт встановлений. Агент захищає дані в усіх трьох станах: DAR, DIU, DIM. Агенти здатні сканувати пристрій на наявність на ньому даних які є конфіденційними (режим Discover). У разі виявлення документу, який містить в собі чутливу інформацію, але при цьому він промаркований як публічний, цей маркер автоматично буде змінений на конфіденційний. Також DLP Кінцевих Точок здатен прослідковувати за діями користувача в системі: що він відкриває, видаляє, змінює, архівує, копіює, робить знімки екрану, переміщує, відправляє.

Головною перевагою DLP Кінцевих Точок над Мережевим DLP є те, що перший буде виконувати свої функції постійно і незалежно від мережі, до якої пристрій приєднаний, в той час як другий буде виконувати свій функціонал лише тоді, коли пристрій під'єднаний безпосередньо до корпоративної мережі організації. Коли співробітник вирішить тричі архівувати та перейменувати файл з метою відправити його туди, куди заборонено, Мережевий DLP не зможе проаналізувати вміст архіву, а DLP Кінцевих Точок, в свою чергу, зрозуміє що саме сталося та запобігне неправомірній передачі даних.

Обидва режими роботи мають свої переваги та недоліки. Кращою практикою для підприємств є використання цих двох типів DLP через те, що система стає складнішою до зламу.

Хмарний DLP. Керівники підприємств по всьому світу поступово переміщують свій бізнес до хмари. Нажаль далеко не всі представники хмарних сервісів пропонують свої послуги із достатнім рівнем безпеки для клієнтів. Тому питання впровадження додаткових технологій для кіберзахисту остається актуальним. Задачею Хмарного DLP є запобігання витоку даних з хмарних сервісів. Принцип роботи такий самий як і у Мережевого DLP.

DLP-система може бути інтегрована в IT-інфраструктуру компанії та поєднана з іншими системами і рішеннями щодо захисту інформації.

Розглянемо можливу архітектуру DLP-системи (рисунок 1.1.3).

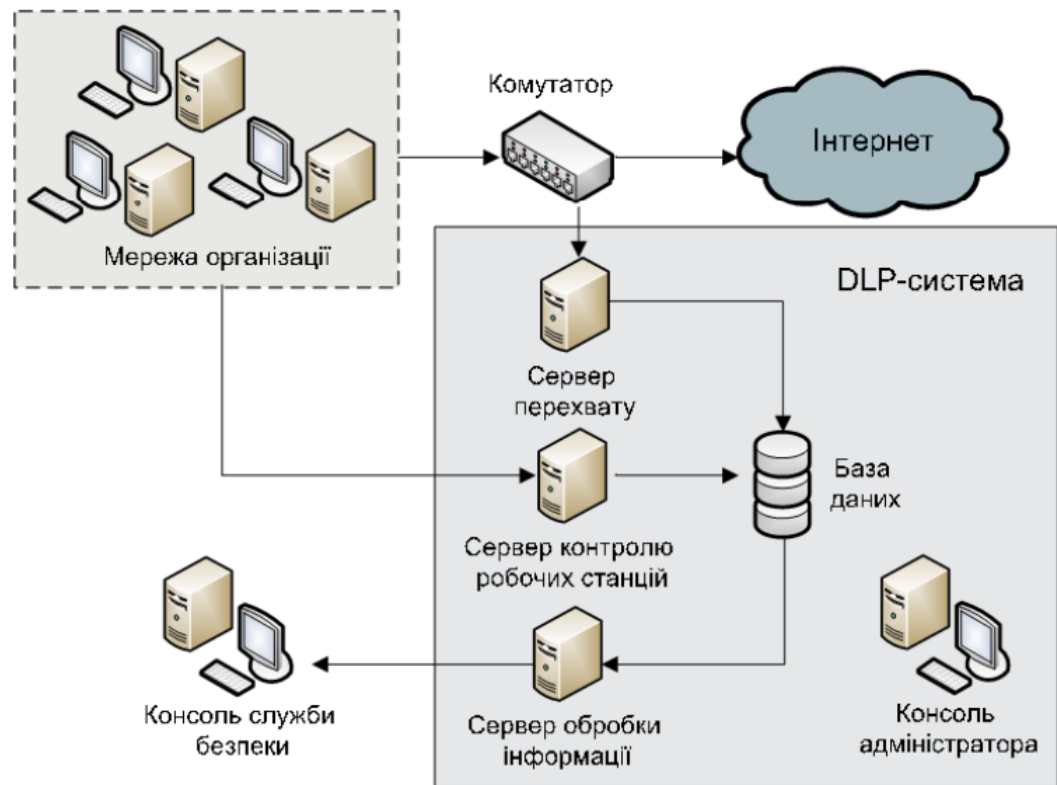


Рис. 1.1.3 Архітектура DLP-системи

Всі інформаційні потоки в компанії перехоплюються сервером перехвату з мережевого адаптера, потім аналізуються і зберігаються в базі.

Інформація, отримана за допомогою сервера перехвату (повідомлення в месенджерах, вкладені файли, електронні листи і т.п.), зберігається на сервері баз даних.

Після збереження в базі даних вся перехоплена інформація надходить на сервер обробки даних, який відповідає за виконання завдань:

- індексування перехопленої інформації;
- повнотекстовий пошук по інформації;

- автоматичний аналіз і відправка на задані адреси електронної пошти повідомлень про факт передачі інформації з порушенням прийнятої в компанії політики безпеки.

Сервер контролю робочих станцій призначений для централізованої установки на комп'ютери програм-агентів, призначених для перехоплення безпосередньо з робочих станцій користувачів трафіку, в тому числі шифрованого, а також даних, що передаються для друку на принтери і зовнішні пристрої. Програми агенти відповідають також за збір статистичної інформації про активність співробітників на робочих місцях. Крім того, сервер контролю робочих станцій здійснює моніторинг стану всіх встановлених в мережі агентів і в разі виявлення збою або примусового відключення агента користувачем будь-якого комп'ютера автоматично здійснює повторну установку.

Консоль адміністратора використовується в DLP-системі для централізованої установки і настройки роботи всіх елементів програми та для віддаленої установки програм-агентів на робочі станції користувачів. Через консоль адміністратора здійснюється настройка параметрів перехоплення і періодичність індексування даних, перегляду статистики по перехопленому трафіку в режимі реального часу, а також налаштовується система повідомлень про збої в роботі системи перехоплення.

Через консоль служби безпеки здійснюється робота з DLP-системою. Відбувається редагування існуючих і створення нових правил безпеки, проводиться оцінка зібраних даних, аналіз звітів, які надають інформацію як про активність окремих співробітників, так і всієї організації в цілому.

Весь нешифрований трафік, який циркулює в мережі компанії, перехоплюється централізовано або за допомогою програм-агентів, встановлених на комп'ютери. Перехоплений трафік передається на сервер перехоплення за допомогою керованого комутатора. На сервері перехоплення

проводиться аналіз отриманого трафіку, в ході якого виділяються необхідні дані (список месенджерів, електронні листи, файли і т.п.) і зберігаються у базах даних DLP-системи.

Програми-агенти дозволяють перехоплювати як нешифровані, так і зашифровані дані. Агенти віддалено встановлюються на комп'ютери сервером контролю робочих станцій. Агенти відстежують дані, що пересилаються в месенджерах і по електронній пошті, всі випадки відправки інформації користувачами на зовнішні пристрої та принтери, та передають все перехоплені дані на сервер обробки. Також за допомогою програм-агентів здійснюється моніторинг діяльності співробітників на робочих місцях і збір статистичних даних для формування звітів.

Після перехоплення вся інформація надходить на сервер обробки даних, де проводиться побудова індексів інформації, що знаходиться в базі, з подальшим збереженням індексів в сховище сервера обробки інформації. Надалі пошук здійснюється по файлах пошукового індексу, тоді як вміст всіх знайдених документів автоматично завантажується з бази даних і відображається в клієнтській консолі.

Результатом впровадження DLP-системи в корпоративну мережу є значне збільшення кількості зареєстрованих випадків порушення політики безпеки співробітниками в частині роботи із конфіденційними даними. Такі дії можуть мати фатальні репутаційні та економічні наслідки. Тому важливість застосування в інфраструктурі мережі такого компонента, як DLP-система, важко взяти під сумнів.

DLP-система – невід'ємна частина системи інформаційної безпеки, неефективна без інших заходів захисту, таких як мережні екрани, антивіруси тощо. Наявність цього елемента захисту інформації та мінімізації ризиків аж ніяк не усуває необхідності в захисті периметра і контролю доступу.

1.2 Класифікація даних

Для того, щоб захистити інформацію від її витоку зсередини організації, або інших несанкціонованих дій над нею, спочатку треба її визначити та промаркувати в залежності від рівня критичності. Саме на підставі цих маркерів будуть застосовуватися певні правила щодо подальших дій.

Boldon James та Titus – є лідерами серед класифікаторів інформації. Свого часу були конкурентами, а наразі належать одній компанії HelpSystems. Вони захищають дані протягом всього їх життєвого циклу. За функціоналом вони майже ідентичні та пропонують дуже широкий спектр можливостей класифікації. Важливо зазначити, що вони відповідають нормативним вимогам таким як CCPA, GDPR, SOX, PCI-DSS, ISO 27001 та HIPAA. Також підтримують інтегрованість з такими технічними партнерами як Microsoft, Forcepoint, Digital Guardian, McAfee та інші. Тисячі компаній по всьому світу впроваджують та довіряють свої дані Boldon James та Titus в різних галузях: фармацевтика, виробництво, професійні та фінансові послуги, уряд та військові сили.

Загалом дані класифікуються за допомогою спеціальних міток, які виставляються класифікатором, та за допомогою аналізу вмісту/контексту файлу (морфологічний аналіз) який виконується автоматично, згідно з встановленими шаблонами.

В першому випадку, кожному документу присвоюється певний гриф який визначає цінність документа. Цих грифів існує 3: **Public**, **Sensitive** та **Confidential**. Маркер **Public** означає, що документ не несе в собі ніякої комерційної таємниці та у разі його витоку організація не зазнає шкоди. Маркер **Sensitive** забороняє документу покидати територію контрольованої корпоративної мережі та означає, що цей документ призначений виключно для внутрішніх процесів компанії. Маркер **Confidential** означає, що документ містить чутливі для організації дані та забороняє його переміщення не тільки за периметр

організації, а навіть між її департаментами. Ці дані мають залишатися в безпеці. Також важливо наголосити, що всі маніпуляції над даними виконуються виключно за правилами та політиками, налаштованими адміністратором класифікатору, та в разі необхідності можуть бути додані, відредаговані або видалені. До конфіденційних даних належать: персональна ідентифікаційна інформація (PII), персональна інформація про стан здоров'я (PHI), дані платіжних карт (PCI), інтелектуальна особистість (IP) тощо.

Classification Scheme



Рисунок 1.2.1 – Схема класифікація для бізнесу та військового секторів

В другому випадку, кожний документ проходить автоматичну перевірку на предмет певних **ключових слів** та/або **фраз**, **регулярних виразів**, певних заданих **шаблонів** та **типів файлів**. В залежності від знайденого, документу буде автоматично присвоєно певний гриф. Це робиться для додаткового рівня захищеності, щоб зменшити ризик випадкової або вимушеної помилки.

Прикладом **ключових слів/виразів** може бути: розробка, таємниця, конфіденційність, персональні дані, банківська карта, переказ, платіж, звіт, паспорт, соціального страхування, гаманець та інші.

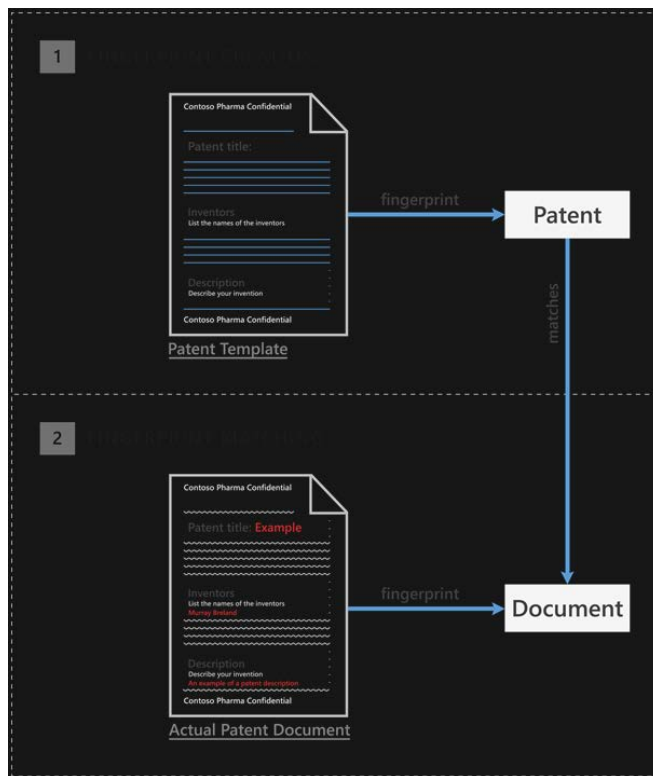


Рисунок 1.2.2 - Приклад співпадіння патентного документа до шаблону

За допомогою **регулярних виразів** можна визначати закономірності послідовностей символів, а саме: номери банківських карт (за кількістю цифр), номер соціального страхування, персональні дані (ПІБ, адреси, номери телефонів та інше), певні дані про розробку та інше.

Стосовно правил та політик. Boldon James та Titus мають дуже гнучке налаштування правил. Адміністратор має змогу створювати, редагувати та видаляти політики, в залежності від потреб організації. Як згадувалося раніше, правила виконуються в залежності від виставленого маркера на документі. Є можливість налаштувати певну політику на певний документ на певного користувача або групу користувачів. Ця гнучкість є дуже практичною та зручною. Можна надати змогу користувачам змінювати грифи на документах, але при цьому ввести додаткові перевірки задля уникнення випадковостей або шахрайства. І що важливо, кожна подібна дія має бути зафіксована та оброблена.

Customer Confidential - 100 Petition

Mortgage Loan Agreement Draft

Sara J. Peterson
Borrower

Co-Borrower

I. TYPE OF MORTGAGE AND TERMS OF LOAN

Mortgage Applied for: ☐ VA ☒ USDA/Rural Housing Service ☐ FHA ☐ Conventional ☐ Other (explain): Agency Case Number: 27036-496 Lender Case Number: 00L27-496

Amount: \$ 146300 Interest Rate: 4.9% No. of Months: 96 Amortization Type: ☒ Fixed Rate ☐ Other (explain): ☐ GPM ☐ ARM (type):

II. PROPERTY INFORMATION AND PURPOSE OF LOAN

Subject Property Address (street, city, state & ZIP): 45 Rachael BLVD, Sylvestre Kansas, 65253 No. of Units: 1

Legal Description of Subject Property (attach description if necessary): Subdivision 4, Section 10, Township 82, West of the 16th Meridian, Dist. 1 Year Built: 1983

Purpose of Loan: ☒ Purchase ☐ Refinance ☐ Construction ☐ Construction-Permanent ☐ Other (explain): Property will be: ☐ Primary Residence ☐ Secondary Residence ☒ Investment

III. BORROWER INFORMATION

Borrower's Name (include Jr. or Sr. if applicable): Sara J. Peterson Co-Borrower's Name (include Jr. or Sr. if applicable):

Social Security Number: 123-12-1234 Home Phone (incl. Area code): 808-488-1212 DOB (mm/dd/yyyy): 12/19/81 Yrs. School: 16

☒ Married ☐ Separated ☐ Divorced ☐ Widowed ☐ Single ☐ Other (explain): Dependents (list last name, date of birth, age): ☐ Married ☐ Separated ☐ Divorced ☐ Widowed ☐ Single ☐ Other (explain): Dependents (list last name, date of birth, age):

Present Address (street, city, state, ZIP): 1 West 45th, LaCrosse, WI 54601 Present Address (street, city, state, ZIP):

Mailing Address, if different from Present Address: Mailing Address, if different from Present Address:

If recording of present address for less than two years, complete the following:

Former Address (street, city, state, ZIP): Former Address (street, city, state, ZIP):

I hereby acknowledge the formal acceptance of this legally binding Mortgage Loan Application document that outlines the type and purpose of the Loan Agreement in the eye of the Legal Entity overseeing this Application and will abide by the Terms and Conditions within.

Метрики Документа

Name	Titus Persistent Metadata Value
TitusGUID	Unique File ID
ML Category	Loan Agreement
GDPR Flag	Yes
LoanID	27038-49b
Legal Hold	LH-4726
PI Type	Name;Address;Location Data
Share Consent	No
Retention	11_9_2021
Data Centre Source	EMEA Zone 4
TVM	0
Classification	Confidential
Sensitivity	Customer Confidential
Customer	New
Visual Markings	Headers Footers
ContainsPersonalData	True
Apply Encryption	None

Рисунок 1.2.3 – Приклад міток у метаданих документа

1.3 Вибір продукту та партнера

1.3.1 Вибір DLP системи

1. Які технології аналізу використовує вендор?
2. Якими каналами DLP система аналізує трафік?
3. Чи може DLP система встановлюватися в розрив та блокувати поширення конфіденційної інформації?
4. Приклади реальних проектів, у яких система встановлювалася на великій кількості робочих місць?
5. Як перевірити зручність роботи із системою, зокрема, її конфігурування та обробку інцидентів.
6. Гнучкість та швидкість роботи інструментів аналізу даних та побудови звітності.
7. Чи вміє система підтримувати політики безпеки в актуальному стані?

8. Чи може система віддавати дані у зовнішні програми, наприклад, SIEM, IRM?
9. Чи надає вендор додаткові послуги із забезпечення правильного з юридичного погляду впровадження DLP системи та розширеної технічної підтримки?
10. Подивіться, що і як каже вендор про себе і своїх конкурентів.

1.3.2 План проведення Пілотного проекту

1. Встановлення модулів та компонентів на віртуальний сервер: база даних Oracle, Enforce Platform управління, сервера детектування (локальна установка);
2. Додавання плагінів, мовних компонентів (локальне налаштування);
3. Налаштування сервера детектування (локальне/віддалене налаштування);
4. Розгортання агентів (тестова група) на робочі станції/сервера (локальне налаштування);
5. Налаштування Політик та правил реагування (локальне/віддалене налаштування);
6. Тестова експлуатація (локальна/віддалена);
7. Налагодження рішення (локальне/віддалене);
8. Тестова експлуатація, після налагодження (локальна/віддалена);
9. Формування наочної звітності (локальна/віддалена);
10. Підбиття підсумків та результати пілотного впровадження (локальна).

Терміни реалізації Пілотного проекту: 3-4 тижні.

Тестовані сценарії:

1. Моніторинг поштових повідомлень на наявність: певних вкладень, певних адресатів, ключових слів та словосполучень.
2. Моніторинг мережевих ресурсів (файлові сховища)

3. Визначення власників інформаційних ресурсів, виявлення користувачів, які отримують доступ до цих ресурсів.

4. Застосування правил для робочих станцій (моніторинг, копіювання, відправлення тощо)

5. Моніторинг підключень до Web-сервісів, надсилання вкладень поштовими сервісами, соціальні мережі.

6. Налаштування правил для створення тінювих копій.

7. Налаштування правил для роботи користувачів поза корпоративною мережею

Тестові Політики:

Підготовка політик відповідно до поставленого завдання:

1.«Test» - забезпечує моніторинг всіх вкладень, які можуть бути копіювати на зовнішні носії, вивантажені в веб-сервіси, відправлені поштою або вкладенням в Skype (далі - переміщення). Ця політика має інформативний характер з метою зрозуміти хто, де і як часто відправляє файли, якими каналами.

2. Test IDM - забезпечує моніторинг переміщення конфіденційної інформації, з якої були зняті цифрові відбитки (неструктуровані). Для цього необхідно задіяти робочий проект із загальним обсягом __ Мб інформації різних форматів (документи, креслення, малюнки тощо). Поріг спрацьовування політики __%.

3. «Test EDM» – забезпечує моніторинг переміщення конфіденційної інформації, з якої було знято цифрові відбитки (структуровані). Для цього необхідно провести вивантаження даних із наявних БД (визначити задіяні БД для тестування).

4. Налагодження рішення:

Політики працюють у режимі моніторингу (дії користувача не блокуються, надсилання файлів відбувається із записом про подію в Enforce Platform). Користувачі не отримують жодних попереджувальних повідомлень.

1.3.3 Рекомендації щодо планування розгортання

Планування установки та вимоги до системи залежать від:

- Типу та обсягу інформації, яку ви хочете захистити
- Обсягу мережного трафіку, який ви хочете відстежувати
- Розміру організації
- Типу серверів виявлення, який ви вибираєте для встановлення.

Ці фактори впливають на:

- Тип рівня установки, який вибирається для розгортання (трирівневий, дворівневий чи однорівневий).
- Системні вимоги для встановлення продукту

1.3.4 Приклад технічного завдання як головного документу для розгляду організацією системи DLP з метою придбання та впровадження:

Endpoint DLP (для 500 робочих місць).

Підтримка та супровід 12 місяців - 24/7.

Загальні основні вимоги (програмне забезпечення повинне забезпечувати):

- Інтеграцію MS Active Directory.
- Підтримку термінального середовища.
- Налаштування та управління за рахунок центральної консолі адміністратора.
- Захист системи від втручання.
- Блокування спроб видалення програми без явного дозволу.
- Блокування спроб змін налаштування системи з кінцевої точки.
- Можливість працювати з архівними записами.

- Можливість інтеграції з SIEM.
- Можливість інтеграції з UEBA.
- Наявність механізмів користувальницької класифікації та маркування даних, забезпечення їх видимості, керованості та захисту; відповідність стандартам та нормативним вимогам (GDPR, ISO 27001: 2013, PCI, SOX, HIPAA).
- Створення white / black списку для налаштування політик.
- Можливість додавання і налаштування шаблонів або ключових слів для ідентифікації даних.
- Централізоване управління документами для виявлених файлів.
- Шифрування дисків, шифрування USB флеш-накопичувачів.
- Проведення аудиту журналу зовнішніх пристроїв, підключених до системи.
- Запобігання друку файлів з конфіденційною інформацією на основі політик.
- Можливість реакції на інциденти: автоматичне повідомлення відповідальної особи або власника інформації.
- Блокування в додатках інструменту Rootkit.
- Пошук журналів по вмісту файлу або типу файлу.
- Можливість створення "Retention Trend" і "Long-term Retention" звітів.
- Оптичне розпізнавання символів за рахунок вбудованого OCR модулю.
- Блокування запуску зареєстрованих процесів в стані входу / виходу (функція BlockingControl).
- PrintScreen контроль.
- Маскування інформації при відображенні конфіденційності даних.
- Технічна сервісна підтримка Виробника або офіційного дистриб'ютора Виробника строком не менше ніж 36 місяців в централізованому розгортанні програмного забезпечення, проведенні навчання співробітників Замовника (не менше 3-х спеціалістів) по роботі з ним.

- Постійний доступ до центра технічної підтримки Виробника через сайт, електронною поштою або за телефоном 24x7
- Постійний (24x7) авторизований доступ до сайту Виробника
- Отриманням від Виробника всіх необхідних оновлень

До питання вибору продукту та партнера потрібно підходити із максимальною серйозністю. Зазвичай консалтингова компанія із інформаційної безпеки з якою ви, як керівник чи представник компанії замовника, співпрацюєте, має допомогти вам визначити ваші потреби та запропонувати ряд рішень та подальшу стратегію розвитку безпеки.

Після визначення ваших цілей, вам будуть представлені рішення які пропонує ринок. Вас, як замовника, вочевидь будуть цікавити лідери серед постачальників продуктів у цьому напрямку. Із цим вам допоможе Radicati Market Quadrant.

Radicati – дослідницька компанія, що спеціалізується на ринках інформаційних технологій. Аналог Гартнера.

Radicati Market Quadrant проводить оцінку постачальників та рішень, класифікуючи їх як: Top Players, Mature Players, Specialists, Trail Blazers. Стосовно постачальників рішень DLP маємо таку картину.

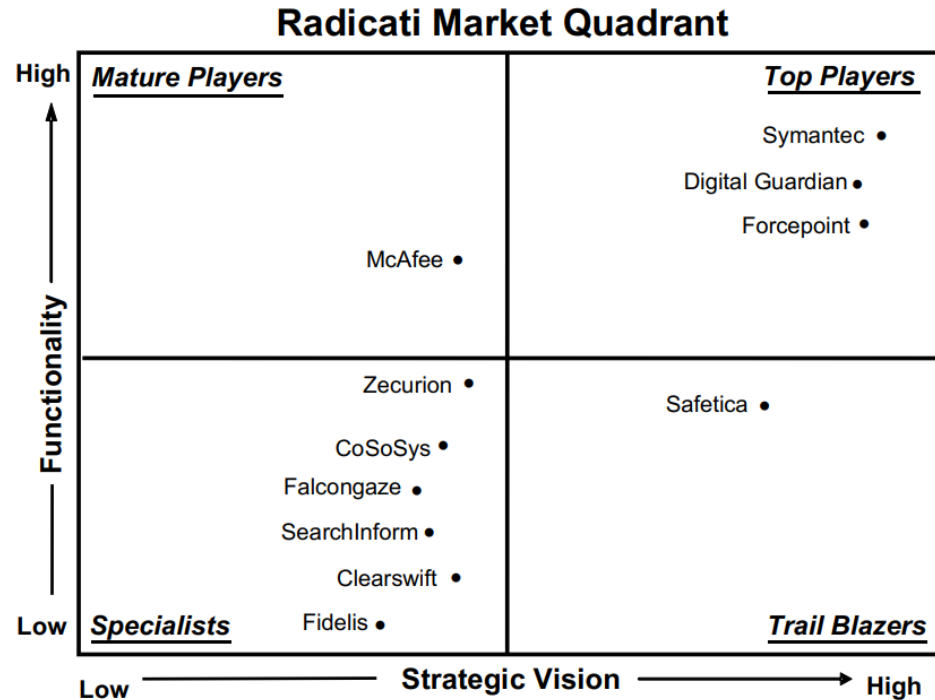


Рисунок 1.3.4.1 – DLP Marker Quadrant 2021

Після того як у вас з'являється приблизне уявлення про те, як на ринку йдуть справи, вам будуть продемонстровані порівняльні таблиці функціоналів продуктів (Battlecard's).

Параметр сравнения	Symantec	McAfee	WebSense	InfoWatch
Сертификаты	ФСТЭК. Сертификат соответствия № 2602 от 22.03.2012 ТУ+НДВ по 4 уровню контроля (ИСПДн 1)	Сертификат ФСТЭК с НДВ по 4 уровню контроля, но только на агентов	Сертификат соответствия ФСТЭК по ТУ N 2645 от 28.08.2012 г	ФСТЭК. Сертификат соответствия НДВ 4 и ИСПДн 1, Газпромсерт, Аккредитация ЦБ, Экспертное заключение ОАЦ при Президенте Республики Беларусь
Общие сведения				
Место в рейтингах Gartner	Первое	Третье	Пятое	Последнее
Сравниваемые версии	Symantec Data Loss Prevention 12.5	McAfee DLP 9.3	WebSense DSS 7.8.3	InfoWatch Traffic Monitor Enterprise 5.1
Единый веб-интерфейс управления	Да	Да, с разделением хостовой и сетевой части в разных меню и с разными политиками	Да	Нет (для управления агентами своя "толстая" консоль InfoWatch Device Monitor)
Системные требования				
Системные требования для клиентской части	Celeron 1000 МГц, 256 МБ ОЗУ	CPU 1000 МГц, 1024 МБ ОЗУ	для Windows: Pentium 4 (1.8 GHz or above) At least 850 MB free hard disk space (250 MB for installation, 600 MB for operation) At least 512 MB RAM on Windows XP At least 1GB RAM on Windows Vista, Windows 7, Windows 8, Windows Server 2003, Windows Server 2008, Windows Server 2012 Standard Edition	Аппаратные требования, предъявляемые Клиентом InfoWatch Device Monitor, соответствуют минимальным требованиям, предъявляемым используемой операционной системой. Кроме этого: Материнская плата должна обеспечивать аппаратную поддержку USB; Свободное место на жестком диске должно составлять не менее 300 МБ
Поддержка агентом ОС Windows	Windows XP\Vista\7\8\8.1	Windows XP\Vista\7\8\8.1	Windows XP\Vista\7\8\8.1	Windows XP\Vista\7\8\8.1
Поддержка агентом Mac OS X	Да (только обнаружение хранимой информации на диске)	Да	Да	Нет
Возможности интеграции				
Сетевые каталоги LDAP	Active Directory и любые другие LDAP каталоги и без каталогов	Active Directory, OpenLDAP, не работает без директории	Active Directory, Novell eDirectory, Lotus Notes, Oracle Directory Services, Generic LDAP directories	Active Directory, Novell eDirectory
Возможность интеграции со сторонними решениями для защиты конечных точек	С любыми решениями с помощью скриптов Python на уровне конечных точек (Endpoint FlexResponse)	Нет	С любыми решениями с помощью скриптов Python на уровне конечных точек (Remedation Scripts)	DeviceLock, Lumension Device Control

Рисунок 1.3.4.2 – Приклад актуального Battlecard (1)

33	Электронная почта	Да	Да	Да	Нет
34	Буфер обмена	Да	Да	Да	Нет
35	Доступ приложений к файлам	Да	Да	Да	Нет
36	Снимок экрана (Print screen)	Нет, но можно отключить возможность делать снимки	Да	Да	Нет
37	Съемные носители	Да	Да	Да	Да, блокировка только на уровне управления устройствами
38	HTTP	Да	Да	Да	Да, без блокировки
39	HTTPS	Да	Да	Да	Да, без блокировки
40	FTP	Да	Да	Да	Да, без блокировки
41	Общие каталоги Windows	Да	Да	Да	Да, без блокировки
42	Печать	Да	Да	Да	Да, блокировка только на уровне управления устройствами
43	Возможность получить теньную копию информации	Да	Да	Да	Да
44	Ограничения доступа в зависимости от типа съемного носителя (производитель, серия, модель, экземпляр)	Да - по регулярному выражению Device ID	Да, критерии: имя устройства, производитель, серийный номер, класс устройства, ID и другие - всего 15 критериев	Да, критерии: имя устройства	Да, критерии: имя устройства, производитель, серийный номер
45	Возможность разрешать копирование только на доверенные носители	Да	Да	Да	Да, блокировка только на уровне управления устройствами, а не на базе контентного анализа
46	Шифрование файлов на USB-устройствах	Да (реакция на инцидент в интеграции с любыми сторонними средствами, используя Endpoint FlexResponse)	Да, при интеграции с McAfee Encryption	Да (собственными средствами)	Нет
47	Контролируемые каналы на мобильных устройствах				
48	Перечень поддерживаемых типов устройств	Устройства на базе iOS (iPad, iPhone), поддержка Android ожидается в 2014 году	iOS, Android, Windows Phone при использовании McAfee MDM	Устройства на базе iOS (iPad, iPhone), Android	Нет
49	Перечень поддерживаемых каналов	HTTP, почта, мобильные приложения (FaceBook, DropBox, Gmail, ActiveSync и т.д.)	HTTP, почта через почтовый клиент, веб-почта, IM-клиенты, Skype	почта через почтовый клиент (EAS)	Нет
50	Поиск конфиденциальной информации в сети предприятия				
51	Сканирование рабочих станций	Да	Да	Да	Да
52	Сканирование общих сетевых хранилищ	Да	Да	Да	Да
53	Сканирование хранилищ Microsoft SharePoint	Да	Да	Да	Да
54	Сканирование баз данных	Да	Да	Да	Нет
55	Сканирование LiveLink	Да	Нет	Нет	Нет
56	Сканирование Documentum	Да	Да	Нет	Нет
57	Сканирование Exchange	Да	Нет	Да	Нет
58	Сканирование Lotus Notes	Да	Нет	Нет	Нет

Рисунок 1.3.4.3 – Приклад актуального Battlecard (2)

Висновок до розділу 1

У цьому розділі були розглянуті основні поняття, пов'язані із напрямом DLP. Були розглянуті: стани в яких знаходяться дані (DAR, DIM, DIU); типи (кінцевих точок, мережевий та хмарний) та режими роботи рішень (активний та пасивний); класифікація даних (public, sensitive, confidential)

Були розглянуті особливості впровадження DLP систем.

2. Сучасні положення, проблематики та рішення

2.1. Огляд статистики збитків у результаті витоку даних

Проблема витоку інформації є дуже актуальною. Це викликано стабільно зростаючою кількістю зафіксованих випадків витоку конфіденційної інформації у всіх країнах світу. У першу чергу було проаналізовано розподіл навмисних та ненавмисних витоків інформації за період 2017 – 2021 рр. (рисунок 2.1.1)

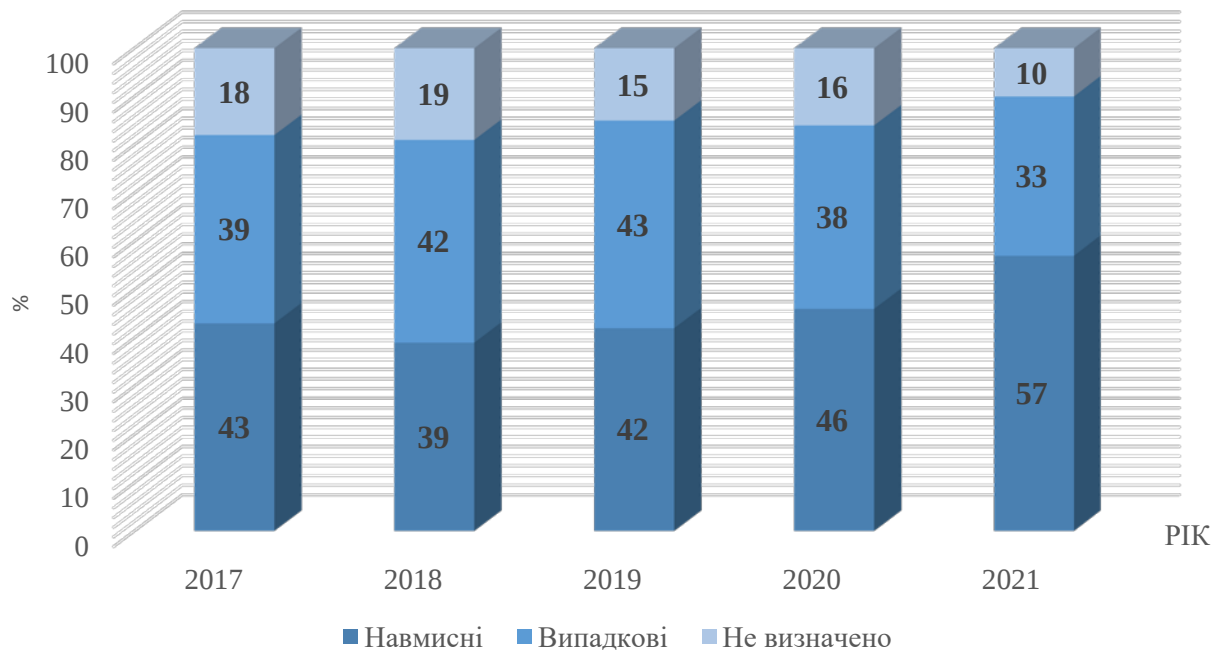


Рисунок 2.1.1 Розподіл навмисних та ненавмисних витоків інформації за період 2017 – 2021 рр, %*

**Складено автором на основі*

На рисунку 2.1.1 проведений динамічний аналіз розподілу навмисних та випадкових витоків інформації організацій за 2017 – 2021 рр. Згідно з наведеною статистикою, частка випадкових витоків інформації знижується протягом аналізованого періоду – у 2017 р. частка випадкових витоків складала 39% від загального об’єму, а у 2021 р. – 33%. З іншого боку, доля навмисних витоків інформації має тенденцію до зростання і склала 57% у 2021 р., що на 14% більше

ніж у 2017 р. Варто відзначити 2020 р., починаючи з якого почало стрімке зростання навмисних витоків інформації. Це, перш за все, пов'язано з поширенням Covid-19, що змусило організації оперативно перелаштуватися на віддалений режим роботи. Отже, можна зробити висновок, що основні проблеми інформаційної безпеки пов'язані, перш за все, з навмисними загрозами, так як вони є головною причиною злочинів і правопорушень. Впровадження системи DLP вплине на співвідношення випадкових і навмисних витоків.

Для того, щоб краще зрозуміти масштаби збитків та ризиків для організацій, було проаналізовано середню вартість фінансових збитків організацій через витік даних за період 2015 – 2021 рр (рисунок 2.1.2) Дані складені на основі показників 537 організацій у 17 галузях по всьому світу.

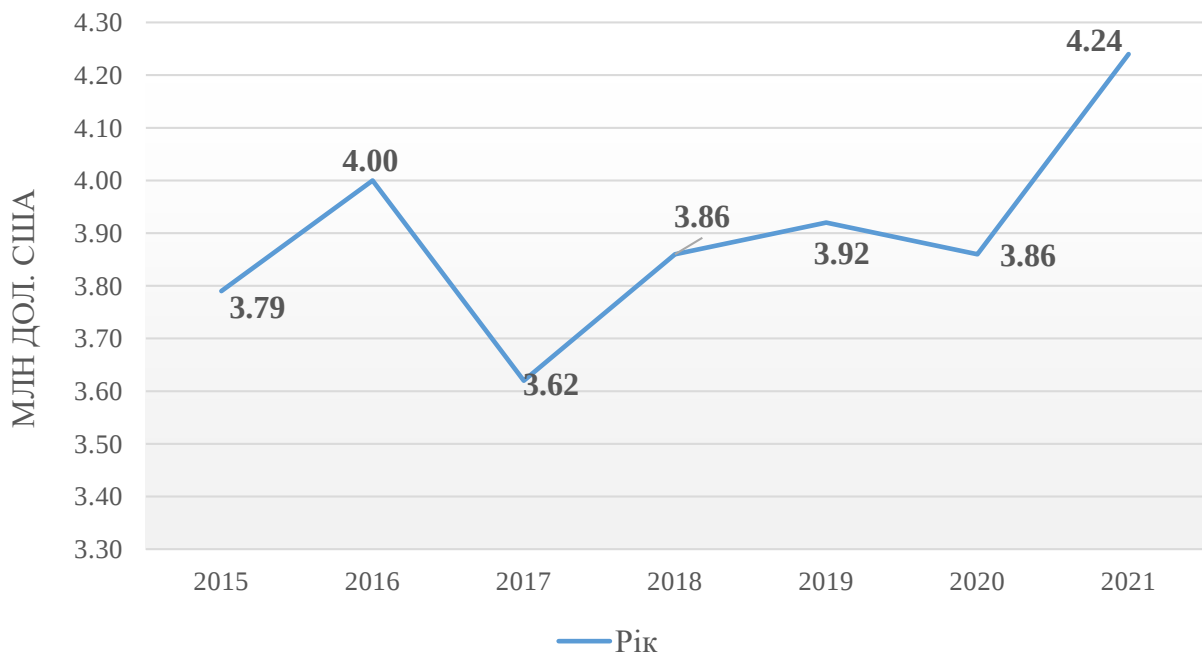


Рисунок 2.1.2 – Середня вартість збитків через витік даних по всьому світу за період 2015 – 2021 рр., млн дол. США*

**Складено автором на основі аналітичних досліджень*

Вартість втрати даних за період 2015 – 2021 рр. має тенденцію до зростання, і зросла за аналізований період з 3,79 млн дол. США до 4,24 млн дол.

США, тобто на 12%. При чому максимальні значення спостерігались у 2016 р. – 4,00 млн дол. США та у 2021 р. – 4,24 млн дол. США.

Для більш детальної аналітики витоку інформації було проведено дослідження з співвідношення кількості інцидентів та кількості витоків у галузевому розподілі за 2021 р. На рисунку 2.1.3 представлені галузі, на які прийшла найбільша частка інцидентів та витоків інформації у світі за результатами 2021 р.

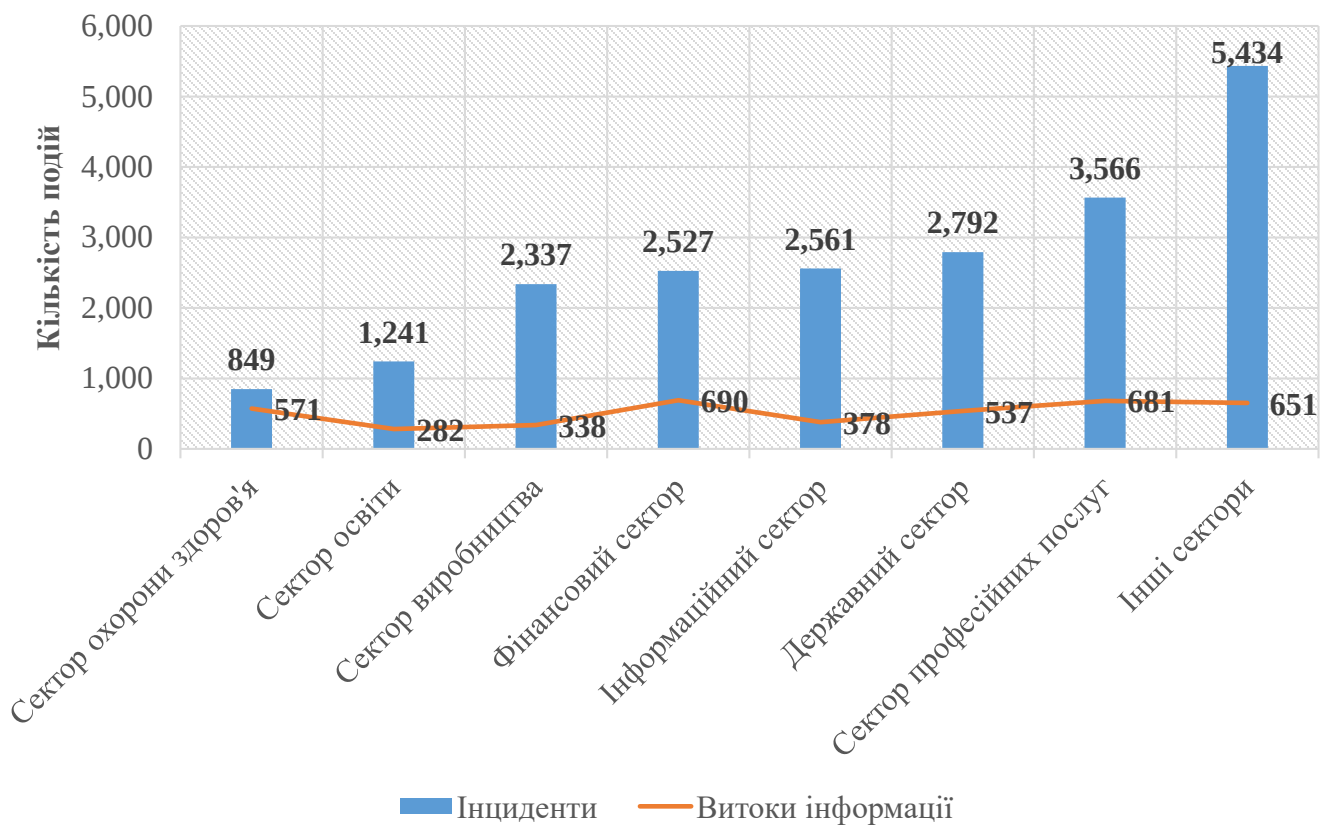


Рисунок 2.1.3 Кількість інцидентів і витоків за 2021 рік в залежності від індустрії*

**Складено автором на основі аналітичних досліджень*

За результатами проведеного дослідження можна зробити висновок, що найбільш незахищеною галуззю є Охорона здоров'я. Всього на кількість 849 інцидентів, кількість підтверджених витоків інформації склала 571, що становить

67% від загальної кількості інцидентів. У той час найбільш захищеними галузями виявились галузі Виробництва та Інформаційних послуг. При кількості 2337 та 2561 інцидентів відповідно, кількість підтверджених витоків інформації 338 та 378, що становить 14% та 15% від загальної кількості інцидентів. Найбільша кількість інцидентів у Секторі професійних послуг (3566) та Державному секторі (2792), але вони мають досить невеликий об'єм витоку інформації – по 19% кожний, що свідчить про достатній рівень захищеності.

Основними загрозами проаналізованих галузей є атаки на веб-сервери та додатки, втручання у роботу системи та соціальна інженерія, а також досить високий рівень помилок адміністрування, що свідчить про недостатній рівень обізнаності персоналу.

Аналізуючи збитки у результаті витоку даних, доцільним є звернути увагу на розподіл середньої вартості збитків організацій по країнах світу у 2021 р (рисунок 2.1.4)

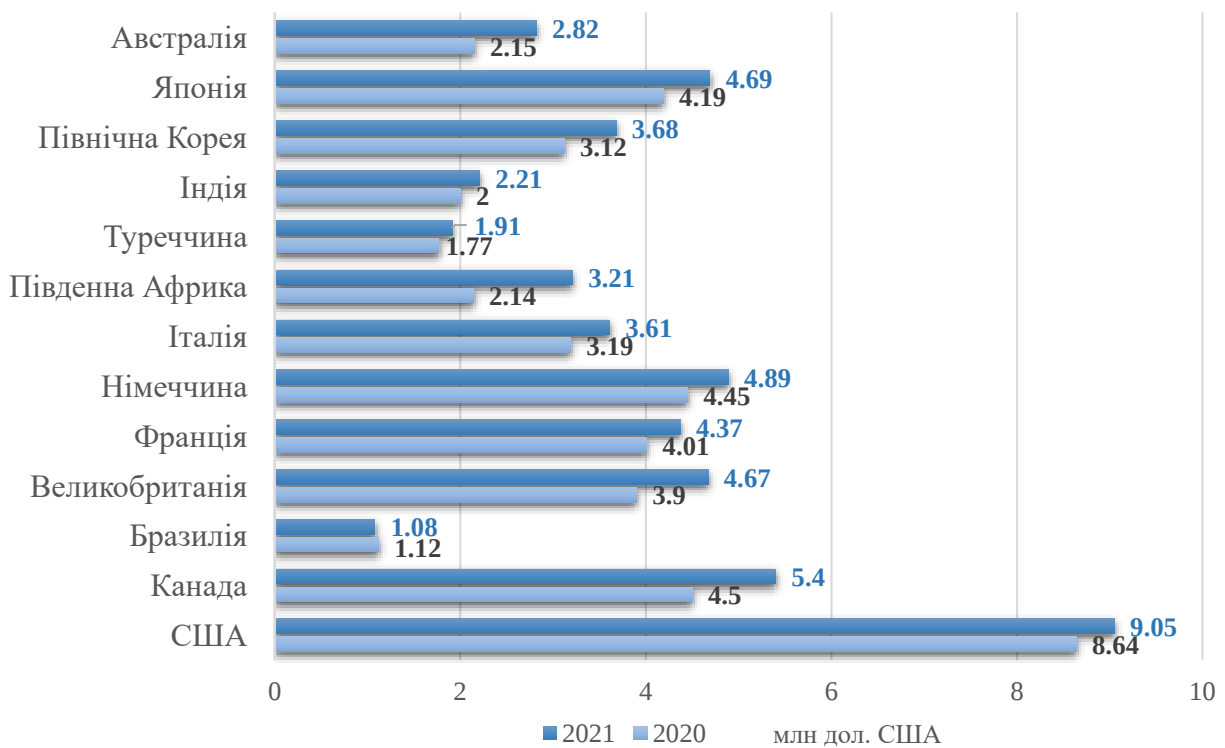


Рисунок 2.1.4 Середня вартість збитків по країнах за 2021*

**Складено автором на основі аналітичних досліджень*

У п'ятірку країн та регіонів за найвищим середнім показником витоку даних увійшли: США, Канада, Німеччина, Японія, Великобританія. Варто звернути увагу на загальну світову тенденцію зростання середньої вартості збитків, що продемонстрували дані всіх аналізованих країн, окрім Бразилії. За результатами дослідження, Бразилія зменшила середню вартість збитків за останні два роки, що склали 1,12 млн дол. США у 2020 р. та 1,08 млн. дол. США у 2021 р. Це може бути пов'язано з тим, що з початком Covid-19, основна увага на світовому рівні була сконцентрована на виробниках медичних препаратів, які розташовані у країнах, які увійшли у п'ятірку країн з найвищим середнім показником витоку даних.

2.2 Проблематика класифікації даних та False-positive спрацювання

Показник успішності DLP систем буде залежати від якості трьох факторів: DPI, Labeling та Rules (правил/політик).

Від якості підходу до маркування даних залежить подальший успіх роботи систем DLP. Щоб дані були правильно класифіковані спочатку потрібно їх визначити (інвентаризація). Після того, коли є загальне бачення активів, присутнє, виконується безпосередньо класифікація. В продуктах DLP часто є свої вбудовані класифікатори, а деколи продукт DLP інтегрується із продуктом класифікатором, таким як Boldon James.

Класифікація може бути ручна – із залученням людини, автоматизована – без залучення людини. Усі методи не досконалі і мають недоліки. Людина може примусово або випадково неправильно промаркувати документ, або ж, якщо має на це права, змінювати маркери документу. Для зниження ризиків пов'язаних із

людським фактором при маркуванні, потрібно впроваджувати додаткові перевірки та наголошувати на важливості цього маркування. Додатковою перевіркою є контекстний аналіз агентом DLP.

Автоматизований контекстний аналіз також не досконалий, оскільки проводиться по строго прописаним шаблонам. Строгість часом призводить до помилкового маркування через відсутність розуміння цими шаблонами сенсу повідомлень. Це в перспективі призведе до False-positive спрацювання агентом DLP. Тому цей автоматизований метод маркування також потребує додаткової перевірки. Використання міток для автоматичного маркування конфіденційних документів не вирішує елементарну задачу ідентифікації вмісту: файл без мітки, навіть якщо містить важливу інформацію, буде пропущений аналізатором. Наприклад, в результаті роботи компанії створюється конфіденційний документ, що забезпечується міткою. Але в процесі його підготовки на робочих місцях співробітників осідає безліч копій, чернеток і фрагментів текстів, які згодом переміщуються по мережі без жодного контролю – технологія міток «сліпа» по відношенню до вмісту.

Тому необхідно виділити можливість маркування документів за допомогою інших методів.

У якості додаткової перевірки вмісту документів пропонується використання машинного навчання.

2.3 Порівняння методів класифікації даних машинним навчанням

Двома найпопулярнішими методами текстової класифікації даних є Support Vector Machine (SVM) та Naïve Bayes (NB) Classifier. Їх часто порівнюють між собою та намагаються визначити який із методів кращий. Під час аналізу висновків досліджень стає зрозумілим, що дійсно обидва методи дуже добре (із

високою точністю) виконують поставлені перед ними задачі. Вибір на користь одного чи іншого методу буде залежати виключно від потреб.

Метод опорних векторів (SVM) – одна з найпопулярніших моделей класифікації. Вона застосовує геометричну інтерпретацію даних. За замовчуванням, це бінарний класифікатор. Він відображає точки даних у просторі, щоб максимізувати відстань між двома категоріями. Для SVM точки даних є N -мірні вектори, і метод шукає $N-1$ -мірну гіперплощину для їх розподілу. Це називається лінійним класифікатором. Цій умові можуть задовольняти багато гіперплощин. Таким чином, найкраща гіперплощина та, яка дає найбільший запас чи відстань між двома категоріями. Таким чином, це називається гіперплощиною з максимальним запасом.

Наївному Байєсу не потрібно вивчати всі можливі кореляції між функціями через припущення про незалежність. Якщо N – це кількість ознак, то загальний алгоритм вимагає аналізу 2^N можливих взаємодій функцій, тоді як NB потребує лише порядку N точок даних. Таким чином, класифікатори NB можуть легше навчатися з невеликих наборів навчальних даних завдяки припущенню про незалежність класу. У той же час на NB не впливає прокляття розмірності.

Наївний Байєс (NB) є дуже швидким методом. Це залежить від умовних ймовірностей, які легко реалізувати та оцінити. Тому для цього не потрібен ітераційний процес. NB підтримує двійкову класифікацію, а також мультиноміальну. NB передбачає, що ознаки між ними незалежні, але це припущення не завжди виконується.

SVM є більш потужним для вирішення завдань нелінійної класифікації. SVM добре узагальнюється у просторах великої розмірності, як тексти. Він ефективний з більшими розмірами, ніж зразки. Добре працює, коли класи добре розділені. SVM за своєю концепцією є бінарною моделлю, хоча її можна застосувати для класифікації кількох класів з хорошими результатами. Вартість

навчання SVM для великих наборів даних є недоліком. SVM займає багато часу під час навчання великих наборів даних. Це вимагає налаштування гіперпараметрів, що не є тривіальним і вимагає часу. SVM більш привабливий теоретично. І NB, і SVM дозволяють вибрати функцію ядра для кожного з них і чутливі до оптимізації параметрів. Порівняння точності SVM і NB у класифікації спау показало, що базовий алгоритм NB дав найкращі результати прогнозування (97,8%). У той же час алгоритми SVM і NB отримали точність значно вище 90%, використовуючи налаштування параметрів, коли потрібно. В різних дослідках на виході маємо різні результати.

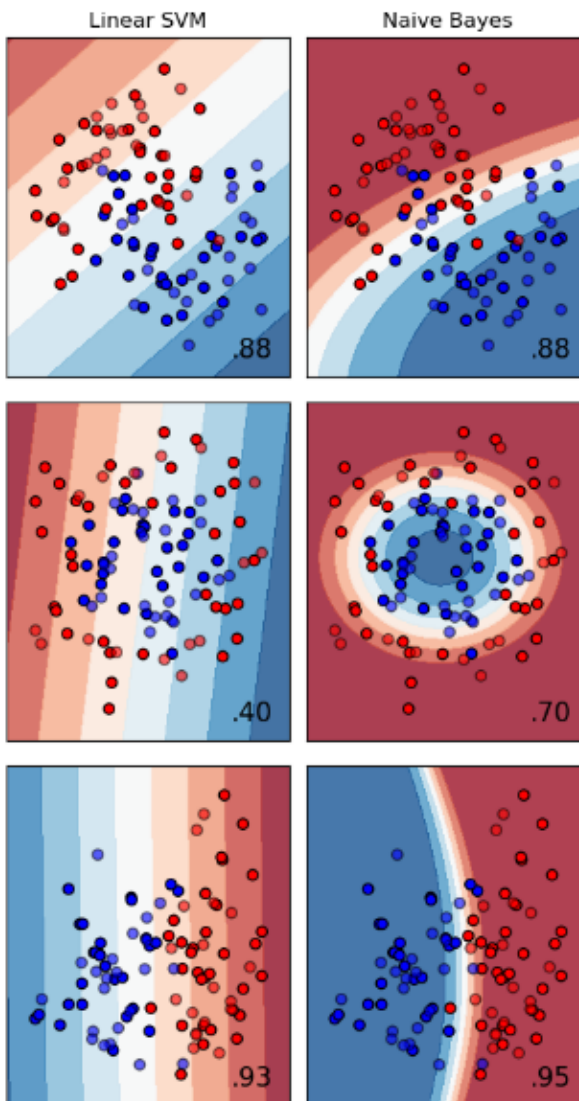


Рисунок 2.3.1 - Порівняння класифікаторів SVM та NB

Підсумовуючи хочеться зазначити, що обидва методи є ефективними. Вибір на користь одного чи іншого буде залежати від конкретно сформульованої потреби. В процесі дослідження визначилося, що метод SVM є кращим рішенням, коли потрібно класифікувати дані на багато категорій. Але його вартість та час, витрачений на навчання, є його головними недоліками. Модель Наївного Байєса навпроти, краще виконує функцію класифікації, коли дані необхідно розбити на невелику кількість категорій, як у нашому випадку – на дві. Вона працює швидше та не потребує багато вхідних даних для навчання завдяки припущенню про незалежність. Якщо чітко та правильно підібрати дані для навчання моделі Байєса, це рішення постає більш релевантним для використання у нашому випадку.

Висновок до розділу 2

У цьому розділі продемонстрована статистика збитків для організацій через витоки даних по всьому світу. З кожним роком збитки збільшуються та їх кількість зростає.

Це говорить перш за все про те, що розрив між технологіями атаки та протидії ним лише збільшується. Окрім старіння технологій захисту та розробок нових методів для атак, на кількісний показник збитків також впливає політична ситуація в світі між державами, та всередині країн. Через збільшення кількості конфліктів між державами та всередині країн, збільшується інтерес та з'являються мотиви для проведення атак з метою отримання переваги. Керівники підприємств та державних органів мають впроваджувати та підтримувати рішення з інформаційної та кібербезпеки задля зменшення ризиків пов'язаних з кібератаками, щоб зменшити вірогідність понести фінансових та репутаційних збитків, а у воєнний час ще й людських втрат.

Також була розглянута проблематика класифікації та False-positive спрацювання. Автоматичне маркування документів є проблемою, яку можна вирішити за допомогою додаткової перевірки вмісту машинним навчанням. Були розглянуті та порівняні два основних класифікатори машинного навчання. Загалом результативність їх дуже схожа. Різниця результатів буде залежати від конкретних потреб. В залежності від потреби був обраний метод Наївного Байєса для подальшого поєднання із DLP системою.

3. Експериментальний метод класифікації даних

3.1 Інтегрованість машинного навчання з рішенням DLP. Опис ідеї

Трапляється, що серед ІТ-активів компанії знаходиться багато НЕ промаркованих документів, особливо на початковому етапі впровадження DLP рішення.

Першочергово правила/політики, визначені та створені адміністратором, спрацьовують згідно з мітками на документах, які виставлені класифікатором даних. Однак якщо цих міток поки що немає, DLP автоматично проводить морфологічний аналіз вмісту цих документів з метою присвоєння йому мітки (публічний, внутрішній, конфіденційний).

Часто навіть у документі, який не є конфіденційним, знаходяться слова або фрази, через які система DLP визнає його конфіденційним та не дозволить маніпуляції над ним. Відомо, що зазвичай із конфіденційними даними не буде особливо багато дозволів, як: копіювання, редагування, видалення, пересилання, знімок екрану тощо. Це автоматичне маркування веде до False-positive спрацювання, що зупиняє бізнес діяльність організації.

Як потенційно можливий до розробки та реалізації варіант рішення цієї проблематики виступає інтегрування машинного навчання (ML) із системою DLP. Machine Learning – це піднапрямок штучного інтелекту в галузі інформаційних технологій, який здатен до самонавчання та вдосконалення без втручання у процес людини. Він створює алгоритми та виявляє закономірності під час аналізу великого масиву даних із метою самонавчання. В процесі його роботи використовуються статистичний та математичний методи навчання. Напрямок розвитку ML поступово виходить на новий рівень, тому такого роду задачі для нього вирішувані. Необхідно інтегрувати машинне навчання із DLP рішенням як додаткову перевірку вмісту документів, тоді і тільки тоді, коли DLP проаналізувало вміст документу, визначило там ключові слова/фрази, та має

намір промаркувати його як конфіденційний. Після того як DLP визначає певний шаблон, який відповідає прописаним вимогам конфіденційності, цей документ направляється на додаткову перевірку до ML.

ML само навчений на тисячах прикладів, та постійно удосконалюється. Його задачею у розрізі аналізу вмісту криється не лише у виявленні та підтвердженні наявності слів/фраз які підпадають під категорію конфіденційних, але і на розуміння загального сенсу контексту.

Коли машинне навчання розпізнає, що те, що виявив DLP, це просто слово/фраза, яке не має нічого спільно із конфіденційними даними, він виносить рішення НЕ маркувати цей документ як конфіденційний. Але оскільки перевага саме у DLP документ все ж таки буде промаркований як конфіденційний, але у зв'язку із розбіжностями рішень, додається наступний рівень перевірки – вручну адміністратором, який виносить кінцеве рішення. Задачею машинного навчання є не “сперечатися” із DLP, а бути його “радником”. ML має аналізувати вірогідності та виносити рішення. Якщо кінцевим результатом цього аналізу буде вірогідність конфіденційності цього документу вище (наприклад) 80%, то ML підтверджує намір системи DLP промаркувати документ відповідною міткою.

Як результат симбіозу між машинним навчанням та DLP, частота False-positive спрацювань значно зменшиться та маркування документів буде точнішим завдяки додатковому рівню перевірки.

Алгоритм їх взаємодії можна описати наступним чином:

1. Першим кроком є перевірка документу системою DLP;
2. Якщо вміст документу не співпадає із шаблонами конфіденційності, він буде промаркований як публічний, або додатково перевірений іншими методами;
3. Якщо вміст документу співпадає із шаблоном конфіденційності, цей документ відправляється на додатковий рівень перевірки до ML;

4. ML прораховує вірогідність того, чи документ і справді є конфіденційним;
5. ML виносить рішення на основі якого початкове рішення системи DLP або підтверджується, або спростовується;
6. Якщо підтверджується – документ промарковано конфіденційним;
7. Якщо спростовується – документ все одно маркується конфіденційним, але додається додатковий рівень перевірки людиною, яка винесе фінальне рішення;

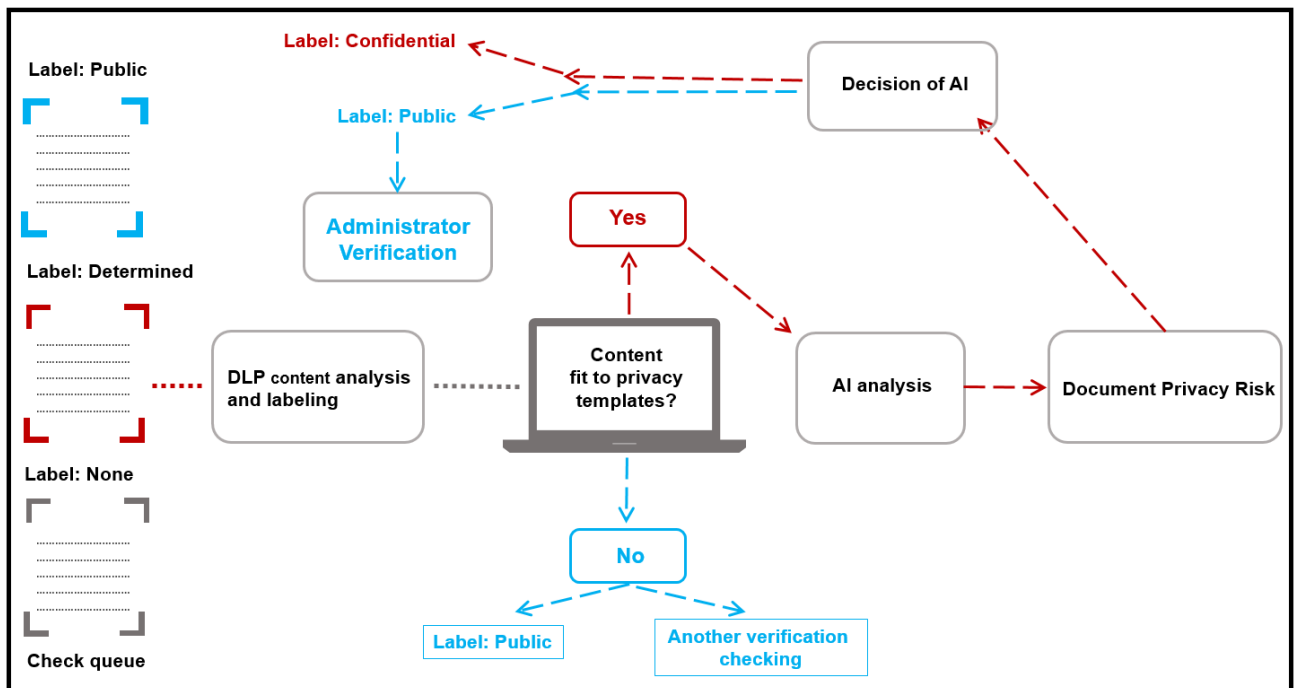


Рисунок 3.1.1 - Графічне представлення алгоритму взаємодії DLP та ML

Ідея полягає в тому, що нехай все ж таки публічний документ промаркується конфіденційним, але цей факт НЕ залишиться осторонь та буде перевірено вручну адміністратором. Сповіщення адміністратору іде тільки у разі розбіжності результатів перевірок.

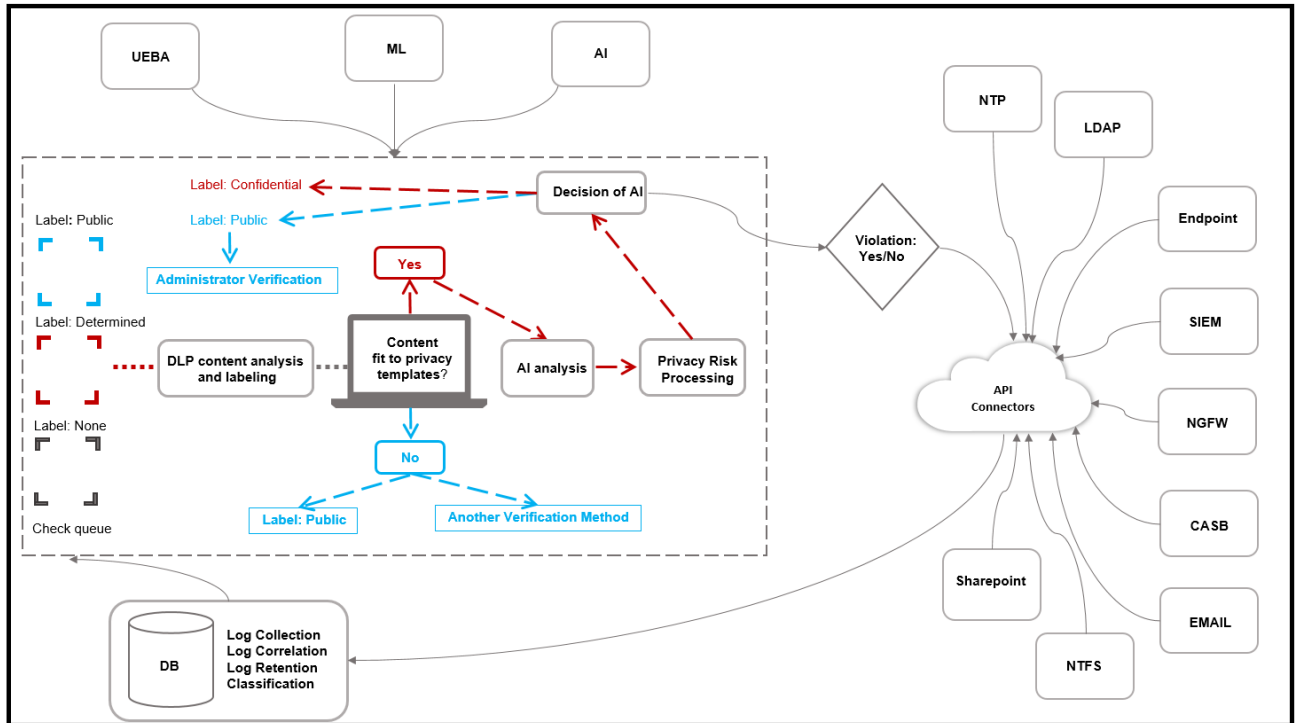


Рисунок 3.1.2 Схема інтеграції DLP+ML із архітектурою організації

3.2 Реалізація запропонованого методу маркування документів

Методи навчання ML

Існує два основні типи машинного навчання:

- Навчання із вчителем (Supervised Learning);
- Навчання без вчителя (Unsupervised Learning);

Навчання із вчителем. У цьому випадку машина має “наставника”, який вказує їй, як правильно вчинити. Наставник заздалегідь відмічає усі потрібні дані, щоб машина самонавчалася на конкретних прикладах. Так він показує, що на цій світлині автомобіль, на іншій – пішохід. У термінах машинного навчання вчитель – це саме втручання людини в процес обробки інформації. З учителем машина вчиться швидше і якісніше, тому для вирішення практичних завдань такі алгоритми використовують частіше. До алгоритмів навчання з учителем

належать такі типи задач, як регресія і **класифікація**, саме те, що нас і цікавить у розрізі DLP.

Навчання без вчителя. У навчанні без учителя машині просто вивалюють купу папок на стіл і кажуть: «розберись, що тут до чого і як». Дані в цьому випадку не розмічені, і машина сама намагається самотійно визначити закономірності. На практиці такі алгоритми використовують рідко, зазвичай як методи аналізу та підготовки даних, а не як основний алгоритм.

Задачі, які виконує машинне навчання: регресія, класифікація, кластеризація, прогнозування, зменшення розмірності, виявлення аномалій та пошук правил.

Не вдаючись у роз'яснення кожного із термінів, рухаємося у потрібному нам напрямі. **Навчання із вчителем -> Класифікація.**

Класифікація. Алгоритми класифікації дозволяють розділити об'єкти відповідно до заздалегідь зазначених класів, наприклад розділити кішок і собак, музику за жанром, **конфіденційні документи і публічні**. Сьогодні алгоритми класифікації використовують для великої кількості завдань: визначення мови, спам-фільтрів, визначення шахрайства (коли зловмисник сплачує за послуги вкраденими коштами), пошуку схожих документів тощо. Щоб класифікація спрацювала, потрібно мати розмічені дані з категоріями і ознаками, які машина буде вчитися визначати. Залежно від певних ознак алгоритм визначає, до якого з класів можна віднести об'єкт. І потрібно розуміти, що чим менша кількість класів, до яких об'єкти мають бути розподілені, тим більша вірогідність саме Правильного розподілення. Оскільки деякі об'єкти часто мають доволі багато ознак, які можуть одночасно відноситися до різних класів-розподілу. Тому, якщо є така можливість, краще зменшити кількість класів, на які потрібно розподілити об'єкти. В ідеалі на два.



Рисунок 3.2.1 Типи машинного навчання

Для ефективного машинного навчання необхідно ретельно підібрати приклади документів на яких система буде навчатися. Як публічних так і конфіденційних. Важливо визначити декілька основних критеріїв та правил вибору цих прикладів:

- Це мають бути текстові документи які чітко пов'язані із темою, що досліджується, або ж має інші спільні риси;
- У разі відсутності спільних рис між документами, алгоритм не зможе визначити закономірності та, як наслідок, ефективно класифікувати дані;
- Необхідна кількість прикладів залежить від рівня їх спільності між собою (конфіденційних і публічних). Якщо ці приклади мають багато поширених

термінів, які зустрічаються дуже рідко, то загалом достатньо невеликої кількості прикладів;

- З іншого боку, якщо відмінності між конфіденційним і публічним набором більш тонкі, буде потрібно більше прикладів.

Після визначення критеріїв підбору прикладів файлів для машинного навчання, можна розпочинати саме навчання. Після надсилання прикладів до ML, сканер починає перевіряти файли та надавати їх до алгоритмів навчання.

У якості методу класифікації був обраний (у минулому розділі) “Наївний Байєсівський класифікатор”. Це ймовірнісний класифікатор який побудований на теоремі Байєса для визначення ймовірності належності об’єкта до визначених класів. Його перевагою над багатьма іншими методами є мала необхідна кількість вхідних даних для навчання, оцінки даних та класифікації та висока швидкість навчання.

$$\text{classify}(f_1, \dots, f_n) = \arg \max_c p(C = c) \prod_{i=1}^n p(F_i = f_i \mid C = c)$$

Рисунок 3.2.2 Будування класифікатору по ймовірнісній моделі

Кероване машинне навчання для захисту даних вимагає, загалом, двох типів приклади: Контент, який потрібно захистити (“**позитивні**” приклади) Контрприклад (“**негативні**” приклади) Контрприклад – це документи, які тематично пов’язані з додатним набором, але не призначені для захисту.

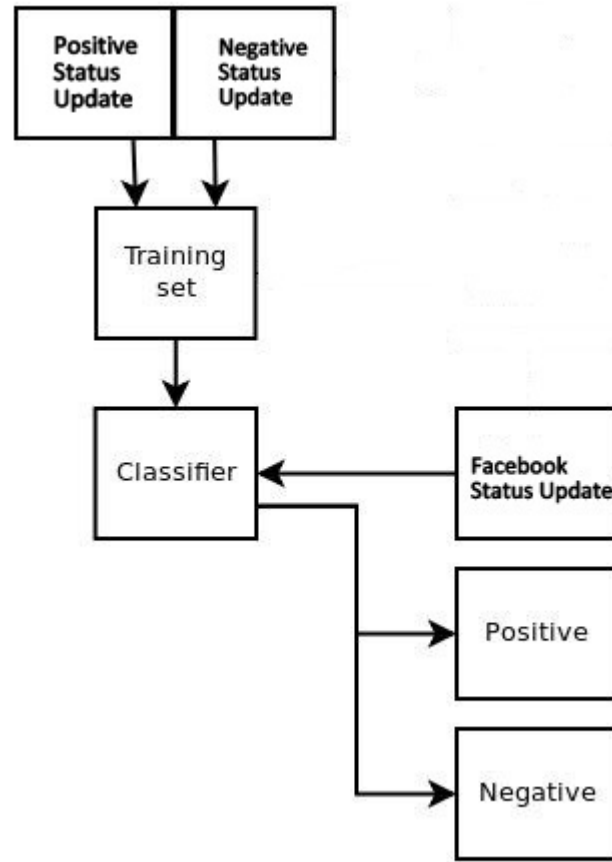


Рисунок 3.2.3 Методологія-алгоритм Наївного Байєсівського класифікатора

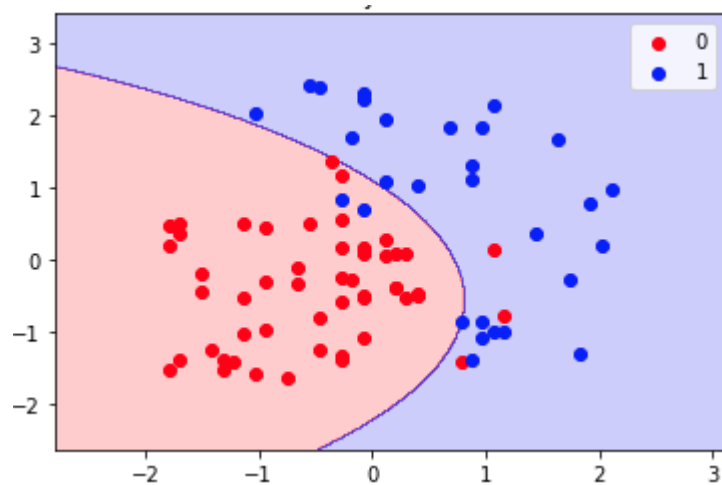


Рисунок 3.2.4 Тестування Баєсівської моделі

Висновок до розділу 3

Був представлений новий підхід до зниження False-positive спрацювань. Це досягається за допомогою інтегрованості машинного навчання із DLP рішенням. Задачею машинного навчання є додаткова перевірка документів на наявність у них конфіденційної інформації. У результаті виносяться рішення про маркування документу належним чином. Якщо DLP помилково виявляє конфіденційну інформацію в документах, він промаркує документ міткою “конфіденційно”, після чого права працівника на документ будуть обмежені під час його взаємодії із документом – це і є False-positive спрацювання якому це гібридне рішення буде протидіяти шляхом додаткових рівнів перевірки.

Машинне навчання в даній роботі було проведено за допомогою Байєсівського класифікатора. Оскільки потрібно розрізняти лише 2 типи документів, публічні та конфіденційні, цей метод має перевагу над іншими.

Висновок

До питань захисту в інформаційному полі завжди потрібно відноситись із достатньою серйозністю на будь-якому рівні: державному, корпоративному та особистому.

Організації мають розуміти свої сильні та слабкі сторони з точки зору безпеки. Впровадивши рішення DLP в архітектуру безпеки організації, зменшуються ризики втрати чутливої до витоку інформації, тим самим забезпечуючи відносну, але не стовідсоткову, стійкість до фінансових та репутаційних втрат.

У першому розділі були розкриті основні поняття DLP як напряду. Розглянутий загальний підхід до захисту організації. Розкриття теми класифікації, вибору продукту. Вивчено законодавство інформаційної безпеки.

У другому розділі наведена статистика збитків для організацій через витоки даних. Наголошено на проблематиці класифікації даних та False-positive спрацювань.

Третій розділ. Представлена під час дослідження концепція створена як додаткова перевірка документів, які класифікатор має намір промаркувати як конфіденційні. Якщо це перевірка підтверджує намір DLP – документ маркується як конфіденційний, якщо заперечує – документ все одно маркується конфіденційним, але додається наступний рівень перевірки вручну адміністратором, який винесе остаточне рішення.

Перелік джерел посилань

1. Міністерство внутрішніх справ України \ Національна академія внутрішніх справ \ РИБАЛЬСКИЙ О.В., СМАГЛЮК В.М., ХАХАНОВСЬКИЙ В.Г. / ОСНОВИ ІНФОРМАЦІЙНОЇ БЕЗПЕКИ. ПІДРУЧНИК ДЛЯ КУРСАНТІВ ВНЗ МВС УКРАЇНИ
2. Factum. Колегія детективів і фахівців безпеки бізнесу. Витік інформації [Електронний ресурс]. – Режим доступу: <http://ukr.detectiveua.com/vitik-inform> – Заголовок з екрану
3. Гержан С. Г., Масальская Е. А. Предотвращение утечки данных с помощью DLP-систем. Режим доступу: http://ir.nmu.org.ua/bitstream/handle/123456789/148756/masalska_gerjan.pdf
4. П.П. Вовчановський, В.В. Демчинський. Архітектура DLP-СИСТЕМ В УМОВАХ ПОЛІТИКИ BYOD
https://www.researchgate.net/publication/347525260_Arhitektura_DLP-sistem_v_umovah_politiki_BYOD
5. Впровадження DLP-систем [Електронний ресурс]. – Режим доступу: <https://techexpert.ua/our-services/implementation-of-dlp-systems/> – Заголовок з екрана.
6. Chapple M., Stewart J. M., Gibson D. (ISC) 2 CISSP Certified Information Systems Security Professional Official Study Guide. – John Wiley & Sons, 2018.
7. ITSM [електронний ресурс] – режим доступу до ресурсу: <https://www.isms.online/iso-27002/control-5-9-inventory-of-information-and-other-associated-assets/>
8. IBM: Cost of Data Breach Report 2021 [електронний ресурс] – режим доступу до ресурсу: <https://www.ibm.com/downloads/cas/OJDVQGRY>

9. Features of modern DLP systems [электронный ресурс] – режим доступа до ресурсу:

<https://www.forbes.com/sites/davidbalaban/2021/11/12/all-features-of-modern-dlp-systems/?sh=27666e2e7618>

10. Classification and DLP [электронный ресурс] – режим доступа до ресурсу: <https://www.endpointprotector.com/blog/how-data-classification-and-data-loss-prevention-go-hand-in-hand/>

11. Automated Classification [электронный ресурс] – режим доступа до ресурсу:

<https://www.dataclassification.com/solutions/automated-classification/#:~:text=What%20is%20automated%20classification%3F,optimal%20success%20in%20data%20classification>

12. Naive Bayes Classification [электронный ресурс] - режим доступа до ресурсу:

<https://towardsdatascience.com/naive-bayes-document-classification-in-python-e33ff50f937e>

Додаток Програмна реалізація

Data Classification **classification.py**

```

from nltk.corpus import brown, movie_reviews, reuters

import random

import math

import operator

import collections


#Function for separating the docs
# def separate(c,C,D):
#   docs_in_each_class=[]
#   for x in C:
#       if(x == c):
#           for d in D:
#               if(c == movie_reviews.categories(d)):
#                   docs_in_each_class.append(d)
#   return docs_in_each_class


def TrainBernoulli(C,D):
    #print D
    #print len(D)
    #voc=[]          #Contains vocabularies
    #voc=movie_reviews.words()
    #docs=[]         #Contains all docs irrespective of category
    #docs=movie_reviews.fileids()

```

```

#C=[]
#C=movie_review.categories() #Contains all the categories

V=[]
t=[]          #V is vocabulary
for x in D:
    t = movie_reviews.words(fileids=x)
    for i in t:
        V.append(i)
#print V
V=list(set(V))          #Removing duplicate elements
N=len(D)                #Count of docs
prior={}
condprob=collections.defaultdict(dict)
# Nc=0
countd={}
shah=[]
docs_in_each_class=[]
for c in C:
    shah=movie_reviews.fileids(categories=c)
    t=0
    for k in shah:
        if k in D:
            t=t+1
            docs_in_each_class.append(k)
    countd[c]=t

```

```

#print countd
for c in C:
    # Nc=Nc+1
    # print Nc
    prior[c] = countd[c]/N
    #print prior[c]

#Nc=len(movie_reviews.fileids(x))
# docs_in_each_class=[]          #contains docs of category c
# docs_in_each_class=movie_reviews.fileids(categories=c)

#    #print Nc

#Nctl=[]
for t in V:
    #condprob[t]={ }
    Nct=0
    for d in docs_in_each_class:
        if t in movie_reviews.words(fileids=d):
            Nct=Nct+1
    print t , Nct
    #print Nct
    #Nctl.append(Nct)

```

```

        Nc=countd[c]
        condprob[t][c]=(Nct+1)/(Nc+2)
    return V,prior,condprob

```

```

def ApplyBernoulliNB(C,V,prior,condprob,d):
    score_fin = []
    for x in d:
        Vd = movie_reviews.words(fileids=x)
        score = {}
        for c in C:
            score[c] = math.log(prior[c])
            for t in V:
                if t in Vd:
                    score[c]+=math.log(condprob[t][c])
                else:
                    score[c]+=math.log(1-condprob[t][c])
            k=max(score.iteritems(),key=operator.itemgetter(1))[0]
            score_fin.append(k)
    return score_fin

```

#Function for splitting the data set

```

def splitDataset(docs,splitRatio_train,splitRatio_test,splitRatio_cross):

```

```

trainSize = int(len(docs) * splitRatio_train)
testSize = int(len(docs) * splitRatio_test)
trainSet = []
testSet = []
copy = list(docs)

while len(trainSet) < trainSize:
    index = random.randrange(len(copy))
    trainSet.append(copy.pop(index))

while len(testSet) < testSize:
    #print len(copy)
    index = random.randrange(len(copy))
    testSet.append(copy.pop(index))

return [trainSet, testSet, copy]

#Function for calculating the accuracy
def getAccuracy(test, predictions):
    correct = 0
    p=0
    for x in test:
        #print "Hey1 ",movie_reviews.categories(fileids=x),
type(movie_reviews.categories(fileids=x))
        #print "Hey2 ",predictions,type(predictions)

```

```

s=movie_reviews.categories(fileids=x)[0]
s=s.encode('utf8')
if s == predictions[p]:
    correct += 1
p=p+1
return (correct/float(len(test))) * 100.0

```

#The calling code

```

voc=[]          #Contains vocabularies
voc=movie_reviews.words()

```

```

docs=[]        #Contains all docs irrespective of category
docs=movie_reviews.fileids()

```

```

C=[]
C=movie_reviews.categories()      ##Contains all the categories

```

#splitting the data set

```
splitRatio_train = 0.60
```

```
splitRatio_test = 0.20
```

```
splitRatio_cross = 0.20
```

```

train,          test          ,cross          =
splitDataset(docs,splitRatio_train,splitRatio_test,splitRatio_cross)

```

#Calling the TrainBernoulli function with the training data

```
V,prior,condprob = TrainBernoulli(C,train)
```

```
#Calling the ApplyBernoulliNB function
```

```
score=[] #List that stores the predicted class labels of the test data
```

```
score=ApplyBernoulliNB(C,V,prior,condprob,test)
```

```
#Calculating the accuracy
```

```
acc=getAccuracy(test, score)
```

```
print('Accuracy: {0}').format(acc)
```