

Міністерство освіти і науки України
Тернопільський національний технічний університет імені Івана Пулюя
(повне найменування вищого навчального закладу)

Факультет комп'ютерно-інформаційних систем і програмної інженерії
(назва факультету)

Кафедра кібербезпеки
(повна назва кафедри)

КВАЛІФІКАЦІЙНА РОБОТА

на здобуття освітнього ступеня

бакалавр
(освітній рівень)

на тему: *"Ідентифікація аномалій мережевого трафіку з використанням нейронних мереж"*

Виконав: студент

Спеціальності:

125 «Кібербезпека»

(шифр і назва напрямку підготовки, спеціальності)

Василишин В. В.

підпис

(прізвище та ініціали)

Керівник

Стадник М. А.

підпис

(прізвище та ініціали)

Нормоконтроль

підпис

(прізвище та ініціали)

Завідувач кафедри

Загородна Н.В.

підпис

(прізвище та ініціали)

Рецензент

підпис

(прізвище та ініціали)

Міністерство освіти і науки України
Тернопільський національний технічний університет імені Івана Пулюя

Факультет комп'ютерно-інформаційних систем і програмної інженерії
(повна назва факультету)
Кафедра кібербезпеки
(повна назва кафедри)

ЗАТВЕРДЖУЮ
Завідувач кафедри
Загородна Н.В.
(підпис) (прізвище та ініціали)
«__» _____ 2022 р.

**ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ**

на здобуття освітнього ступеня Бакалавр
(назва освітнього ступеня)
за спеціальністю 125 Кібербезпека
(шифр і назва спеціальності)
Студенту Василишину Вадиму Віталійовичу
(прізвище, ім'я, по батькові)
1. Тема роботи Ідентифікація аномалій мережевого трафіку з використанням нейронних мереж

Керівник роботи Стадник М. А., к.т.н.
(прізвище, ім'я, по батькові, науковий ступінь, вчене звання)

Затверджені наказом ректора від « 23 » березня 2022 року № 4/7-178
2. Термін подання студентом завершеної роботи 17.06.2022

3. Вихідні дані до роботи Промаркований набір даних із позначеннями аномалій

4. Зміст роботи (перелік питань, які потрібно розробити)

Вступ. 1. Огляд літературних джерел. 1.1 Open System Interconnection (OSI) Model

1.2 Мережеві загрози 1.2.1 Атаки на фізичному рівні 1.2.2 Пасивні атаки 1.3. Botnet

1.4 Основи захисту мережі 1.5 Постановка задачі 2 Теоретичні основи виявлення

аномалій у мережевому трафіку. 2.1 Типи аномалій. 2.2 Виявлення вторгнення в мережу

2.2.1 Ознаки для виявлення вторгнення в мережу. 2.2.2 Виявлення аномалій за допомогою

даних і алгоритмів. 2.3 Проблеми використання МН та штучного інтелекту для виявлення

аномалій. 3 Виявлення аномалій мережевого трафіку з використанням C-LSTM мережі.

3.1 Етапи розробки алгоритму ідентифікації аномалій мережевого трафіку. 3.2 Дослідження

вхідного набору даних та його попередня обробка 3.3 Обґрунтування обрання моделі

3.4 Навчання моделі класифікатора 3.5 Тестування C-LSTM для виявлення аномалій

мережевого трафіку. 4 Безпека життєдіяльності, основи хорони праці 4.1 Заходи захисту

обладнання від статичної електрики 4.2 Значення адаптації в трудовому процесі. Висновки.

Список літературних джерел.

5. Перелік графічного матеріалу. 1. Титулка. 2. Мета та задачі дослідження. 3. OSI модель.

4. Мережеві загрози. 5. Типи аномалій мережевого трафіку. 6. Виявлення мережевих

аномалій. 7. Етапи розробки алгоритму. 8. Візуальне дослідження характеристик. 9. 1D

послідовність мережевого трафіку. 10. Приклади аномалій трафіку. 11. Структура

комбінованої нейронної мережі C-LSTM. 12. Результати навчання C-LSTM. 13. Оцінка якості

класифікатора C-LSTM. 14. Порівняння точності класифікаторів після 10-кратної

перехресної перевірки. 15. Висновки

6. Консультанти розділів роботи

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		завдання видав	завдання прийняв
Безпека життєдіяльності, основи хорони праці			

7. Дата видачі завдання 23.03.2022 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів роботи	Термін виконання етапів роботи	Примітка
1.	Ознайомлення з завданням до кваліфікаційної роботи	23.03 – 28.03	Виконано
2.	Підбір джерел про аномалії мережевого трафіку	29.03 – 05.04	Виконано
3.	Опрацювання джерел в галузі дослідження	06.04 – 11.04	Виконано
4.	Виконання дослідження щодо ідентифікацій аномалій мережевого трафіку	12.04 – 18.04	Виконано
5.	Розроблення програмного коду	19.04 – 25.04	Виконано
6.	Оформлення розділу “Огляд літературних джерел”	26.04 – 29.04	Виконано
7.	Оформлення розділу “Теоретичні основи”	02.05 – 05.05	Виконано
8.	Оформлення розділу “Виявлення аномалій мережевого трафіку з використанням C-LSTM мереж”	06.05 – 11.05	Виконано
9.	Виконання завдання до підрозділу “Безпека життєдіяльності, основи хорони праці”	12.05 – 16.05	Виконано
10.	Оформлення кваліфікаційної роботи	17.05 – 08.06	Виконано
11.	Нормоконтроль	09.06 – 13.06	Виконано
12.	Перевірка на плагіат	14.06 – 15.06	Виконано
13.	Попередній захист кваліфікаційної роботи	16.06 – 17.06	Виконано
14.	Захист кваліфікаційної роботи	23.06	

Студент

(підпис)

Василишин В. В.

(прізвище та ініціали)

Керівник роботи

(підпис)

Стадник М. А.

(прізвище та ініціали)

АНОТАЦІЯ

Ідентифікація аномалій мережевого трафіку з використанням нейронних мереж
// Кваліфікаційна робота ОР «Бакалавр» // Васишин Вадим Віталійович // Тернопільський національний технічний університет імені Івана Пулюя, факультет комп'ютерно-інформаційних систем і програмної інженерії, кафедра кібербезпеки, група СБс-42 // Тернопіль, 2022 // С. – 56, рис. – 16 , табл. – 0, кресл. – 15, додат. – 0.

Ключові слова: ВИЯВЛЕННЯ АНОМАЛІЙ, МЕРЕЖЕВИЙ ТРАФІК, БОТНЕТ, ВИЯВЛЕННЯ ВТОРГНЕННЯ, OSI-МОДЕЛЬ, НЕЙРОННА МЕРЕЖА.

В кваліфікаційній роботі вирішується проблема виявлення аномалій мережевого трафіку з використанням C-LSTM нейронної мережі, яка є комбінацією згорткової та рекурентної нейронних мереж. В роботі наведено основні загрози, що можуть виникати на кожному з рівнів моделі OSI. Детально розглянуто класифікацію botnet мереж, що є одною із причин аномального трафіку. Наведено типи аномалій та здійснено систематичний огляд існуючих методів виявлення, таких як: класифікація, кластеризація, статистичний та на основі теорії інформації.

Імплементовано алгоритм виявлення аномалій мережевого трафіку з використанням нейронної мережі C-LSTM, представлено оцінку якості її класифікації за показниками точності, повноти та F1-ознаки.

ANNOTATION

Identification of network traffic anomalies using neuron networks // Qualification thesis of educational level “Bachelor” // Vasilyshyn Vadym Vitaliyovych // Ternopil Ivan Pulyuy National Technical University, Faculty of Computer Information Systems and Software Engineering, Department of Cybersecurity, CБс-42 group // Ternopil, 2022 // P. – 56 , fig. – 16, table – 0 , drawing – 15 , apendix – 0.

Key words: ANOMALY DETECTION, NETWORK TRAFFIC, BOTNET, INTRUSION DETECTION, OSI-MODEL, NEURAL NETWORK.

The qualification work solves the problem of detecting network traffic anomalies using C-LSTM neural network, which is a combination of convolutional and recurrent neural networks. The paper presents the main threats that may arise at each level of the OSI model. The classification of botnet networks, which is one of the causes of abnormal traffic, is considered in detail. The types of anomalies are given and a systematic review of existing detection methods is performed. They are classification, clustering, statistical, and information theory.

An algorithm for the network traffic anomalies identification has been implemented by using the C-LSTM neural network. An assessment of the quality of its classification in terms of accuracy, recall, and F1-score has been presented.

ЗМІСТ

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ, СКОРОЧЕНЬ І ТЕРМІНІВ	7
ВСТУП.....	8
1 ОГЛЯД ЛІТЕРАТУРНИХ ДЖЕРЕЛ.....	10
1.1 Open System Interconnection (OSI) Model.....	10
1.2 Мережеві загрози	11
1.2.1 Атаки на фізичному рівні	12
1.2.2 Пасивні атаки	12
1.2.1 Активні атаки	13
1.3 Botnet.....	17
1.4 Основи захисту мережі.....	21
1.5 Постановка задачі	24
2 ТЕОРЕТИЧНІ ОСНОВИ ВИЯВЛЕННЯ АНОМАЛІЙ У МЕРЕЖЕВОМУ ТРАФІКУ	26
2.1 Типи аномалій	26
2.2 Виявлення вторгнення в мережу	28
2.2.1 Ознаки для виявлення вторгнення в мережу	30
2.2.2 Виявлення аномалій за допомогою даних і алгоритмів	31
2.3 Проблеми використання МН та штучного інтелекту для виявлення аномалій.....	34
3 Виявлення аномалій мережевого трафіку з використанням C-LSTM мережі	36
3.1 Етапи розробки алгоритму ідентифікації аномалій мережевого трафіку.....	36
3.2 Дослідження вхідного набору даних та його попередня обробка	37
3.3 Обґрунтування обрання моделі.....	40
3.4 Навчання моделі класифікатора.....	42
3.5 Тестування C-LSTM для виявлення аномалій мережевого трафіку	45
4 БЕЗПЕКА ЖИТТЄДІЯЛЬНОСТІ, ОСНОВИ ХОРОНИ ПРАЦІ	48
4.1 Заходи захисту обладнання від статичної електрики.....	48
4.2 Значення адаптації в трудовому процесі	50
ВИСНОВКИ	54
СПИСОК ЛІТЕРАТУРНИХ ДЖЕРЕЛ	56

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ, СКОРОЧЕНЬ І ТЕРМІНІВ

ARP	Address Resolution Protocol
C&C	Command and Control
CNN	Convolution Neural Network
DPI	Deep Packet Inspection
IRC	Internet Relay Chat
LSTM	Long Short Term Memory
SVM	Support vector machine
TCP	Transmission Control Protocol
VLAN	Virtual Local Area Network
МН	Машинне навчання
ОС	Операційна система
ПЗ	Програмне забезпечення

ВСТУП

Інформаційні технології значно поширюються у повсякденному житті. Кожен з нас використовує смартфон для прочитання актуальних новин щодо війни в Україні, практично кожен українець користується додатком “Тривога” для реагування на ракетні напади ворога, обмінюється інформацією через електронні пошти, спілкується в месенджерах зі своїми рідними, робить оголошення в соціальних мережах щодо допомоги збройним силам України, деякі навіть будучи закордоном управляють дистанційно пристроями в розумному будинку. Все це було б неможливо без мережі та отримання мережевого трафіку. Саме тому надзвичайно важливо було отримати систему Starlink та використовувати її для військових та на територіях із обмеженим доступом до інтернету. Таке масове використання комп’ютеризованих пристроїв та систем не тільки в Україні, але і у всьому світі, призводить до критичних загроз, такі як вразливості нульового дня, деякі мобільні атаки, отримання доступу до персональних мереж.

Необхідність захисту комп’ютерної мережі завжди було актуальною задачею. Оскільки еволюція комп’ютерних мереж не тільки дозволила користувачам збільшити швидкість трафіку, отримати нові та покращені вбудовані системи захисту в активні компоненти мережі, але і спровокувала ще більшу кількість мережевих атак. Зазвичай всі нововведення та новітні технології загострюють проблеми з комп’ютерною безпекою, оскільки потребують додаткового захисту.

Найбільш ймовірно, що зловмисники перш за все отримують доступ до інфраструктури, тобто мережі, а вже потім до певних елементів на її вузлах. Безпека мережі – це широка практика захисту комп’ютерних мереж і доступних до мережі кінцевих точок від зловмисного, неправильного використання та відмови. Автоматизоване виявлення аномалій у мережевому трафіку є надзвичайно важливою задачею, оскільки вони вказують на значущі та рідкісні події і є підставою до виконання критичних дій, які необхідно виконати в кількох областях застосування для забезпечення необхідним рівнем безпеки.

Метою кваліфікаційної роботи є виявлення та ідентифікація аномалій мережевого трафіку з використанням машинного навчання та штучного інтелекту. Для досягнення цієї мети необхідно вирішити ряд завдань, які представлено нижче у вигляді списку:

- Дослідити предметну область та причини виникнення аномалій у мережевому трафіку.
- Проаналізувати основні методи виявлення аномалій у мережі
- Дослідити широкоживані алгоритми класифікації.
- Обґрунтувати обрання моделі класифікатора.
- Дослідити якість роботи класифікатора для виявлення аномалій.
- Протестувати роботу моделі класифікатора на тестувальному наборі даних.

1 ОГЛЯД ЛІТЕРАТУРНИХ ДЖЕРЕЛ

1.1 Open System Interconnection (OSI) Model

Зловмисники націлені на мережі здебільшого як засіб для досягнення ширших цілей, таких як шпигунство, вилучення даних, шкода. Через складність і повсюдність підключення між комп'ютерними системами зловмисники часто можуть знайти способи отримати доступ навіть до найбільш ретельно розроблених систем. Тому ми зосереджуємо наше дослідження безпеки мережі на сценаріях, які припускають, що зловмисник уже зламав периметр і активно атакує мережу зсередини, спричинюючи аномалії у трафіку.

Для аналізу мережевого трафіку та виявлення аномалій надзвичайно важливо розуміти взаємодію мережевого обладнання і програмного забезпечення (ПЗ), процес передачі даних через мережеві протоколи, основи роботи комп'ютерної мережі. Для опису цих питань найкраще підходить модель взаємозв'язку відкритих систем OSI [*]. Модель OSI містить сім шарів, що представлено на рис. 1.1.

Рівень 1: фізичний. Перетворює цифрові дані в електричні або механічні сигнали, які можна передавати по мережі, і перетворює сигнали, отримані по мережі, в цифрові дані.

Рівень 2: каналний. Здійснює передачу даних між сусідніми вузлами у фізичній мережі.

Рівень 3: мережевий. Маршрутизує пакети та здійснює контроль потоку між двома точками мереж.

Рівень 4: транспортний. Забезпечує зв'язок між хостом, диктує якість і надійність передачі даних.

Рівень 5: рівень сесії. Ініціює, підтримує та закриває сесії між процесами програми.

Рівень 6: рівень презентації. Перекладає двійкові дані у формати, зрозумілі додаткам (в ергономічному представленні).

Рівень 7: прикладний шар. Відображає користувачам дані, отримані з мережі.



Рис. 1.1 – Мережевий стек моделі OSI

При передачі на основі пакетів потік даних сегментується на менші одиниці, кожен з яких містить деякі метадані про джерело, призначення та вміст передачі. Кожен пакет передається через мережевий рівень і форматується у відповідному протоколі транспортним рівнем з реконструкцією інформаційного потоку з окремих пакетів, що відбуваються на рівні сесії або вище.

Безпека мережевого, транспортного та рівня сесій (рівні 3, 4 та 5 моделі OSI відповідно) є предметом дослідження кваліфікаційної роботи, оскільки саме на цих рівнях атаки характеризуються аномаліями.

1.2 Мережеві загрози

Перш за все необхідно дослідити модель загроз перед аналізом мережевих даних, оскільки виявлення причини та джерела аномальних даних спрощують дослідження в декілька разів. Для початку проаналізуємо атаки, які стосуються лише мережевого, транспортного та рівня сесії. Незважаючи на те, що атаки фізичного рівня (рівень 1 [L1]) і канального рівня (рівень 2 [L2]) важливі для розгляду в загальних обговореннях безпеки мережі, проте їх реалізація та

безпеку виходять за рамки розробників додатків та інженерів безпеки. Таким чином, реалізувати захист на цих рівнях, як правило, неможливо для аналітиків кібербезпеки. За захист рівнів L1 та L2 відповідають спеціалісти при проектуванні та розробці мережевих пристроїв або програмного забезпечення до них.

Представимо загальну класифікацію загроз на пасивні та активні атаки, а також більш детальноше оглянемо чотири класи активних атак: порушення (breaches), спуфінг (spoofing), pivoting і відмова в обслуговуванні (DoS).

1.2.1 Атаки на фізичному рівні

Атаки фізичного рівня, наприклад атаки на 802.3 Ethernet, 802.11 WiFi і Bluetooth, часто обговорюються в публікаціях з безпеки мережі, оскільки вони дуже поширені. Крім фізичного знищення мережевих пристроїв або підключення електронних сніферів до сегментів мережі, атаки на бездротові механізми передачі, такі як WiFi і Bluetooth, особливо актуальні.

Aircrack-ng є прикладом програмного забезпечення для злому WiFi, яке автоматизує обхід певних слабких протоколів безпеки, що захищають мережі WiFi. MAC flooding є прикладом атаки канального рівня, яка переповнює таблиці MAC в мережевих комутаторах, щоб замінити законні записи нелегітимними, спричиняючи надсилання мережевих пакетів до ненавмисних частин мережі.

Інтернет прослуховування є формою пасивної мережевої атаки на фізичний рівень. Це передбачає перехоплення мережевого трафіку в певний момент передачі, передачу його зловмисникам без відома користувачів мережі.

Атаки “man-in-the-middle” часто використовуються для прослуховування Інтернету та дозволяють зловмисникам переглядати конфіденційний трафік між хостами та користувачами.

1.2.2 Пасивні атаки

Пасивні мережеві атаки не ініціюють зв'язок з вузлами в мережі і не взаємодіють з даними мережі та не змінюють їх. Зловмисники зазвичай використовують пасивні прийоми для збору інформації та розвідувальної діяльності. Сканування портів – це пасивна мережева атака, яку зловмисні особи використовують для пошуку відкритих портів, щоб ідентифікувати служби, що працюють на серверах. Враховуючи знання про порти за замовчуванням певних служб або програм, супротивник може дізнатися, що працює сервер, лише на основі його відкритих портів; наприклад, відкритий порт 27017 пропонує запущений екземпляр MongoDB. Зауважимо, що прослуховування Інтернету є формою пасивної атаки, яка проявляється на фізичному рівні.

1.2.1 Активні атаки

Активні атаки набагато різноманітніші і їх можна розділити на наступні категорії:

Порушення (breaches). Мережні атаки є, мабуть, найпоширенішими мережевими атаками. Термін “breach” може означати діру у периметрі внутрішньої мережі або дію зловмисника, який використовує таку діру для отримання несанкціонованого доступу до приватних систем. Для багатьох серверних інфраструктур вузли мережі лежать по периметру. Для прикладу, це можуть бути маршрутизатори, проксі, брандмауери, комутатори та балансувальники навантаження. Системи виявлення вторгнень є формою захисту периметра, яка намагається виявити, коли зловмисник активно використовує вразливості периметра, щоб отримати доступ до мережі.

Зловмисники можуть проникнути в мережі після виконання пасивного збору та розвідки інформації, виявляючи вразливості в загальнодоступних кінцевих точках, які дають їм доступ до системної оболонки або root. Команди, що видаються через мережу, які намагаються використати вразливості програм, можна виявити шляхом перевірки зв'язку між серверами, тому атаки на віддалені програми іноді класифікуються як атаки злому мережі. Наприклад, атаки переповнення буфера або *heap*, SQL-ін'єкції та атаки міжсайтових

сценаріїв (XSS) принципово не викликані проблемами з мережею, однак, перевіряючи трафік між хостами, системи безпеки мережі можуть потенційно виявити такі спроби злому серверів у мережі.

Основні спроби віддаленого переповнення буфера можна виявити шляхом перевірки вмісту мережових пакетів на наявність певних сигнатур атак і блокування або поміщення підозрілих пакетів на карантин. Тим не менш, поліморфізм таких дистанційно запусканих атак зробив метод перевірки сигнатур атак по суті марним.

Значну загрозу для мережі також становлять порушення даних з боку інсайдерів. Виявлення внутрішніх загроз має на меті виявити, коли законні люди в системі скомпрометують цю систему (наприклад, корумпований співробітник, який намагається продати бізнес-секрети конкурентам). Внутрішні загрози є особливо складною проблемою, оскільки системні адміністратори зазвичай мають безперешкодний доступ до інфраструктури, а також можуть контролювати системи безпеки, які запобігають незаконним спробам отримати доступ до цінних даних. Для зменшення рівня атак такого типу необхідно налаштувати належну політику контролю доступу на основі ролей і систему стримувань та протидії для контролю та аудиту внутрішніх систем безпеки, а також шифрування збережених даних.

Spoofing. Зловмисники використовують spoofing (тобто надсилання фальсифікованих даних) як механізм встановлення своєї присутності в середині довіреного каналу зв'язку між двома об'єктами. Spoofing DNS і spoofing ARP (також відомий як забруднення кешу) зловживають механізмами мережевого кешування, щоб змусити клієнта спілкуватися з підробленим об'єктом замість призначеного об'єкта. Видаючи себе передбачуваною стороною-отримувачем, зловмисники можуть брати участь у пасивних атаках прослуховування телефонних розмов, щоб вилучити інформацію з конфіденційних комунікацій. DNS-сервери транслюють читані людиною домени (наприклад, www.example.com) на IP-адреси серверів за допомогою протоколу дозволу DNS. Мережевий зловмисник може “забруднити” кеш DNS клієнта, тимчасово

скерувавши клієнта на нелегітимний DNS-сервер, який відповідає на запит на дозвіл DNS зі зловмисною IP-адресою. DNSSEC було введено, щоб гарантувати аутентифікацію та цілісність у DNS, що запобігає більшості атак підробки доменних імен.

User to Root (U2R): зловмисник прагне отримати незаконний доступ до адміністративного облікового запису для маніпулювання або зловживання важливими ресурсами. Використовуючи підхід соціальної інженерії або перебираючи пароль, зловмисник отримує доступ до звичайного облікового запису користувача, а потім використати чи деяку вразливість, щоб отримати привілей суперкористувача.

Віддалений доступ до користувача (Remote to User, R2U): зловмисник хоче отримати локальний доступ як користувач цільової машини, щоб мати привілей надсилання пакетів через її мережу. Найчастіше зловмисник використовує метод проб і помилок, щоб вгадати пароль за допомогою автоматизованих сценаріїв, методу грубої сили тощо. Існують також деякі складні атаки, за допомогою яких зловмисник встановлює інструмент для перехоплення пароля, перш ніж проникнути в систему.

Pivoting – це техніка переміщення між серверами в мережі після того, як зловмисник отримав доступ до точки входу. Інфраструктури, в яких межі між службами спроектовані або налаштовані неправильно, особливо вразливі до таких атак. У багатьох серйозних порушеннях даних зловмисники пройшли через мережу після отримання доступу до хоста з низькою цінністю. Хости низької цінності – це хости в мережі, які не надають багато інформації зловмисникам, навіть якщо їх зламати. Зазвичай це системи, спрямовані назовні, такі як веб-сервери або термінали в точках продажу. За конструкцією ці системи не зберігають інформацію, яка була б цінною для зловмисників. Отже, бар'єри безпеки навколо них зазвичай є більш послабленими в порівнянні з високоцінними системами, такими як бізнес-логіка або бази даних клієнтів. У добре спроектованій безпечній мережі зв'язок між вузлами низької вартості та високоцінними хостами має бути дозволено лише через невеликий і

контрольований набір каналів. Однак багато мереж не є ідеально сегментованими і містять сліпі зони, які дозволяють зловмисникам переходити з однієї віртуальної локальної мережі (VLAN) до іншої, зрештою знаходячи шлях до високоцінного хосту. Такі прийоми, як підробка комутаторів і подвійне тегування, дозволяють зловмисникам виконувати перехід VLAN на неадекватно налаштованих інтерфейсах VLAN. Зловмисники також можуть використовувати скомпрометований низькоцінний хост для виконання атак грубої сили, фаззінгу або сканування портів на іншій частині мережі, щоб знайти свій наступний перехід у мережі.

Meterpreter – це інструмент, призначений для автоматизації процесу pivot мережі; ви можете використовувати його в тестуванні на проникнення, щоб знайти головні вразливості у вашій мережі.

DoS. Атаки відмови в обслуговуванні (DoS, Denial of Service) спрямовані на загальну доступність системи, перешкоджаючи доступу до неї потенційним користувачам. Існує багато різновидів DoS-атак, включаючи розподілену атаку DoS (DDoS, Destibuted Denial of Service), яка відноситься до використання кількох IP-адрес, які можуть охоплювати широкий діапазон геолокацій для атаки служби.

SYN flooding — це метод DoS, який зловживає механізмом рукостискання TCP і використовує необережно реалізовані кінцеві точки. У процесі тристороннього рукостискання TCP кожне нове TCP-з'єднання починається з того, що клієнт надсилає на сервер SYN-пакет. Сервер відповідає клієнту пакетом SYN-ACK і чекає певний період часу на відповідь ACK від клієнта. Сервер повинен підтримувати напіввідкрите з'єднання під час очікування пакету ACK від клієнта, враховуючи, що затримка в отриманні пакетів може бути спричинена різними причинами, такими як проблеми з підключенням до мережі або проблеми з перевантаженням. Підтримка цих напіввідкритих з'єднань споживає системні ресурси протягом певного періоду часу. Шкідливі клієнти можуть постійно ініціювати TCP-з'єднання, наповнюючи сервер пакетами SYN і не відповідаючи ACK, в кінцевому підсумку вичерпуючи сервер його кінцевих

ресурсів і не даючи легітимним клієнтам ініціювати з'єднання. Існує багато інших різновидів DoS-атак, які подібним чином вичерпують доступні ресурси серверів, щоб перервати доступність послуги.

Botnet (ботнет) – це мережа скомпрометованих приватних комп'ютерів, які заражені шкідливим ПЗ і використовуються для віддаленого доступу. Ботнети можуть використовуватися в кількох варіантах, включаючи розсилку спаму, шахрайство з кліками та вилучення веб-вмісту, але звичайне використання – це запуск DDoS-атак. В наступному підпункті буде детальніше описано архітектуру ботнет мереж.

1.3 Botnet

Приблизно половину всього веб-трафіку генерують боти. Боти набувають багатьох форм, іноді таких простих, як сценарій Bash, що складається з команд curl, іноді у вигляді сценарію безголовного браузера, такого як PhantomJS або іноді навіть у вигляді широкомасштабного поширеного веб-сканера, що працює на базі, як-от Apache Nutch.

Важливо знати про ботнети через велику загрозу безпеці, яку вони становлять для корпоративних мереж і організацій. Зомбі, які є частиною мережі, можуть не вважатися активними мережевими зловмисниками, але можуть мати такі ж шкідливі наслідки, коли активується ботнет. Подібно до зловмисників АРТ та мережових зловмисників, бот-зомбі можуть запускати інсайдерські атаки, які отримують важливу інформацію, погіршують цілісність системи та спричиняють хаос у вашому середовищі.

Аналіз запитів DNS або аналіз шаблонів для пошуку характеристик мережевого трафіку можуть вказати на поведінку автономного бота. За допомогою МН можна також виявити ботів, які використовують мережі швидкого потоку для маскуванню серверів командно-управління (C&C). Крім того, методи кластеризації були застосовані для захоплення мережевого трафіку для виявлення зв'язку зомбі-контроль на основі знання топологій бот-мереж.

Ботнети – повністю розподілені системи, які зазвичай є відмовостійкі та високодоступні архітектури, які можуть протистояти будь-кому, хто намагається їх закрити. Зомбі-машини в бот-мережах (тобто боти) зазвичай контролюються делегаторами завдань (серверами C&C). Коли комп'ютер заражений шкідливим програмним забезпеченням бот-мережі, відбувається інсталяція прихованого процесу-демона, який очікує вказівок від сервера C&C. Перші ботнети використовували протоколи Internet Relay Chat (IRC) для зв'язку з C&C. IRC був легким вибором на початку бот-мереж через його повсюдність і простоту розгортання. Крім того, це була технологія, з якою зловмисники ботнетів були знайомі, оскільки на каналах IRC відбувається велика кількість підпільної активності. Загалом, існує чотири категорії архітектур C&C:

Зірка/централізована мережа: єдиний централізований C&C-сервер, який видає команди всім зомбі. Ця топологія, яка називається топологією зірки, забезпечує найбільш прямий зв'язок із зомбі, але, очевидно, є вразливою, оскільки в системі є одна точка збою. Якби сервер C&C було видалено, адміністратори більше не матимуть доступу до зомбі. Зомбі в цій конфігурації зазвичай використовують DNS як механізм для пошуку свого C&C-сервера, оскільки адреси C&C мають бути жорстко закодовані в програмі ботнету, а використання імені DNS замість IP-адреси забезпечує інший рівень опосередкованості в цій системі, яка часто змінюється. Дослідники безпеки виявляють зомбі у мережі, шукаючи підозрілі запити DNS.

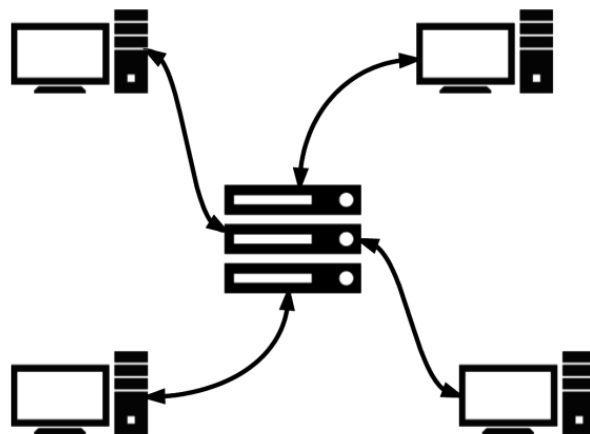


Рисунок 1.2 – Централізована схема ботнет C&C мережі

Топологія розширена зірка. Топологія бот-мережі з багатьма “зірками” дуже схожа на централізовану топологію С&С, але спеціально розроблена для усунення єдиної точки збою. Ця топологія значно підвищує рівень складності мережі (рис. 1.3): С&С сервери повинні часто спілкуватися між собою, синхронізація стає проблемою, а координація загалом вимагає більше зусиль. З іншого боку, розширена зірка зменшує проблеми, що виникають через фізичну відстань. Для бот-мереж, які охоплюють великі географічні регіони, наявність зомбі, які постійно спілкуються з С&С-сервером по всьому світу, є джерелом неефективності та збільшує шанси виявлення, оскільки кожне спілкування споживає більше системних ресурсів. Розповсюдження серверів С&С по всьому світу, як це роблять мережі доставки вмісту з веб-ресурсами, є хорошим способом вирішення цієї проблеми. Але все одно має бути якийсь DNS або балансування навантаження сервера, щоб зомбі могли спілкуватися з кластером С&С. Крім того, ця топологія не вирішує проблему, коли кожен зомбі повинен спілкуватися з центральним командуванням.

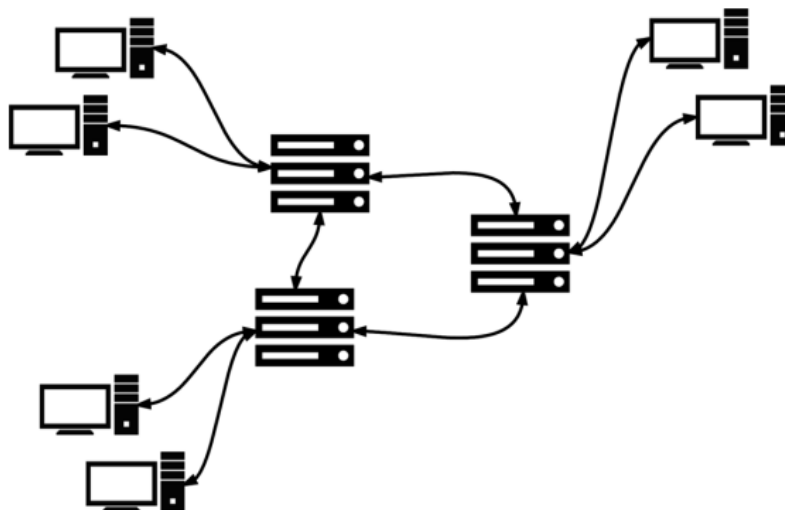


Рисунок 1.3 – Топологія розширена зірка для ботнет С&С мережі

Ієрархічні топології С&С були спеціально розроблені для вирішення проблеми централізованого командування. Як представлено на рис. 1.4, зомбі більше не є просто слухачами, але тепер також можуть діяти як ретранслятори

для отриманих команд. Адміністратори ботнетів можуть адмініструвати команди невеликій групі безпосередньо підключених зомбі, які потім, у свою чергу, поширюватимуть їх своїм дочірнім зомбі, поширюючи команди по всій мережі. Відстеження значно ускладнюється. Тим не менш, оскільки для розповсюдження команд по мережі потрібен час, топологія може бути непридатною для дій, які вимагають реакцій або відповідей у реальному часі. Крім того, видалення одного зомбі, який опиняється на вищому місці в ієрархії поширення, може зробити значну частину мережі недоступною для бот-мережі адміністраторів.

Рандомізовані мережі P2P. Наступним еволюційним кроком у топології бот-мереж є повністю децентралізована система зомбі, що нагадує однорангову (P2P) мережу (рис. 1.5). Ця топологія усуває будь-який центральний механізм управління та керування. Адміністратори ботнетів можуть видавати команди будь-якому боту або ботам у мережі, і ці команди потім поширюються по всій мережі в стилі багатоадресної передачі. Це налаштування призводить до високостійкої системи, оскільки видалення окремого бота не впливає на інших ботів у мережі. Тим не менш, ця топологія має свою частку складнощів. Проблема затримки розповсюдження команд не вирішена, і може бути важко забезпечити, щоб команди були видані та підтверджені всіма ботами в мережі, оскільки немає централізованої команди або офіційного протоколу потоку інформації.

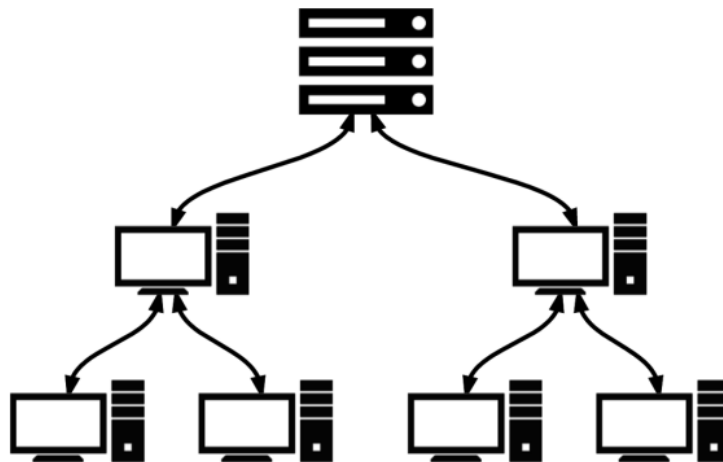


Рисунок 1.4 – Ієрархічна топологія для ботнет C&C мережі

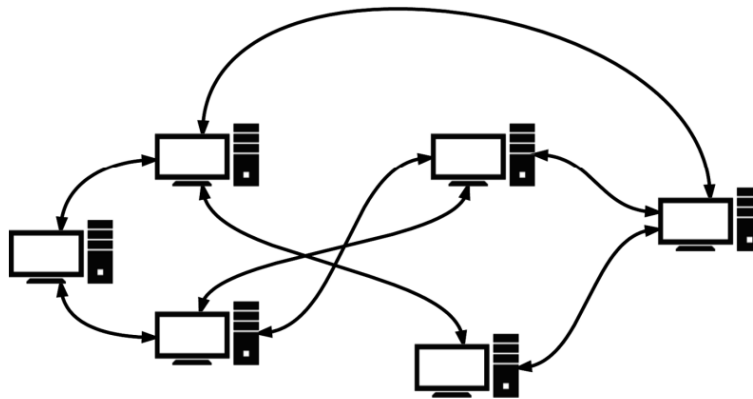


Рисунок 1.5 – Рандомізовані мережі P2P для ботнет C&C мережі

Крім того, автори ботнетів повинні розробити свій механізм видачі команд таким чином, щоб запобігти захопленню їхніх бот-мереж третіми сторонами. Оскільки команди для керування ботнетом можна надавати будь-якому окремому боту, у тому числі тим, які можуть знаходитися в несприятливих або ненадійних середовищах виконання, важливий надійний механізм для забезпечення автентичності виданих команд. Поширеним способом досягнення цього в сучасних P2P бот-мережах є підписання команд адміністраторами за допомогою асиметричної криптографії, тому аутентифікація команд може виконуватися децентралізованим і безпечним способом.

1.4 Основи захисту мережі

Мережі мають складну модель захисту через широкий діапазон можливих атак та векторів загроз. Як і у випадку з захистом будь-якої складної системи, адміністратори повинні взаємодіяти зі зловмисниками на кількох напрямках і не повинні робити припущень щодо надійності будь-якого окремого компонента рішення.

Контроль доступу та аутентифікація. Взаємодія клієнта з мережею починається з рівня контролю доступу. Контроль доступу – це форма

авторизації, за допомогою якої контролюються користувачі, ролі або хости в організації та їх доступ до кожного сегмента мережі. Брандмауери є засобом контролю доступу, який дотримується попередньо визначених політик щодо того, як хостам у мережі дозволено спілкуватися один з одним. Ядро Linux включає в себе вбудований брандмауер iptables, який забезпечує набір правил на рівні IP, який визначає вхідні та вихідні можливості хоста, які можна налаштувати в командному рядку. Наприклад, системний адміністратор, який хоче налаштувати iptables так, щоб дозволити лише вхідні з'єднання Secure Shell (SSH) (через порт 22) із певної підмережі, 92.168.100.0/24.2.

Активні методи аутентифікації збирають більше інформації про клієнта, що підключається, часто використовуючи криптографічні методи, приватні знання або розподілені механізми для досягнення надійнішої атестації клієнта. Наприклад, на додаток до використання джерела з'єднання як сигналу, системному адміністратору може знадобитися ключ SSH та облікові дані автентифікації, щоб дозволити запит на з'єднання. У деяких випадках багатофакторна аутентифікація (MFA) може бути найкращим методом для підвищення рівня безпеки. MFA розбиває автентифікацію на кілька частин, через що зломисникам доведеться використовувати кілька схем або пристроїв, щоб отримати обидві частини ключа, щоб отримати потрібний доступ.

Системи виявлення вторгнень виходять за межі початкових бар'єрів контролю доступу, щоб виявити спроби або успішні порушення мережі шляхом пасивного спостереження. Системи запобігання вторгненню продаються як удосконалення пасивних систем виявлення вторгнень, які мають здатність перехоплювати пряму лінію зв'язку між джерелом і призначенням та автоматично діяти на виявлені аномалії. Перегляд пакетів у режимі реального часу є вимогою для будь-якої системи виявлення або запобігання вторгненню, яка забезпечує як рівень видимості для вмісту, що протікає через периметр мережі, так і даних, на основі яких може здійснюватися виявлення загроз.

Виявлення зломисників всередині мережі. Якщо припустити, що зломисники зможуть пройти через механізми контролю доступу та уникнути

виявлення системою виявлення вторгнень, вони проникнуть через периметр мережі та працюватимуть у межах надійної інфраструктури. Добре спроектовані системи повинні припускати, що це завжди так, і працювати над виявленням зловмисників у мережі. Правильна сегментація мережі може допомогти обмежити шкоду, завдану зловмисниками всередині мережі. Відокремлюючи загальнодоступні системи від високоцінних центральних інформаційних серверів і дозволяючи лише ретельно охоронюваний і контрольований доступ через API між окремими сегментами мережі, адміністратори звужують канали, доступні для зловмисників, забезпечуючи більшу видимість і можливість аудиту трафіку, що протікає між вузлами. Мікросегментація – це практика сегментації мережі на різні ділянки на основі логічної функції кожного елемента. Правильна мікросегментація може спростити мережу та керування політикою безпеки, але її постійна ефективність залежить від суворих процесів зміни інфраструктури. Зміни в мережі повинні бути точно відображені в схемах сегментації, керувати якими може бути складно, оскільки складність систем зростає. Тим не менш, сегментація мережі дає адміністраторам можливість запровадити суворий контроль за кількістю можливих шляхів між вузлом А і вузлом В у мережі, а також забезпечує додаткову видимість, що дозволяє використовувати науку про дані для виявлення зловмисників.

Безпека на основі даних. Коли периметр скомпрометований, будь-які дані, що зберігаються в мережі, можуть бути вкрадені. Зловмисники, які отримують доступ до сегмента мережі, що містить облікові дані користувачів, мають безперешкодний доступ до інформації, що зберігається в ній, і єдиний спосіб обмежити збиток – це використовувати метод орієнтований на дані. Шифрування збережених даних – це спосіб досягнення безпеки, орієнтованої на дані, і стандартною практикою хешування.

Тим не менш, шифрування збережених даних підходить не для всіх контекстів і середовищ. Безперервне шифрування та дешифрування даних може мати нереально високі вимоги до ресурсів. Виконання аналізу даних із

зашифрованими даними – це мета, над якою працюють дослідницькі спільноти з питань безпеки та інтелект [8].

Honeypots – це приманки, призначені для збору інформації про зловмисників. Загальними цілями компонентів honeypot є вивчення методології та цілей атаки, а також збір криміналістичної інформації для аналізу дій зловмисника. Honeypots представляють інтерфейси, які імітують реальні системи, і іноді можуть бути дуже успішними в обмані зловмисників, щоб вони виявили характеристики їхньої атаки, що може допомогти в автономному аналізі даних або активних контрзаходах. Honeypots, стратегічно розміщені в середовищах із значним обсягом атак, можуть бути корисними для збору позначених навчальних даних для дослідження та вдосконалення моделей виявлення атак.

1.5 Постановка задачі

Проаналізувавши типові активні мережеві атаки, такі як breaches, spoofing, pivoting, DoS, атаки з використанням ботнет мережі, розуміємо, що захист мережі повинен відбуватись в комплексі. Звичайно адміністратори мережі повинні виконати ряд превентивних заходів: сегментація мережі, закриття типових портів, встановлення списку дозволених IP адрес, двохфакторної автентифікації та авторизації, встановлення прав доступу, проінформувати типових користувачі щодо зловмисного ПЗ та встановити рівні відповідальності за витік інформації. Проте у випадку масової атаки з використанням ботнет мережі надзвичайно важко зрозуміти причину та виявити місце в периметрі мережі, де відбувається атака. Автоматичне виявлення аномалій дозволяє спрогнозувати атаку на мережу та запобігти шкідливим наслідкам. Звичайно виявлення аномалій мережевого трафіку є також одним з напрямів захисту, що сприятиме підвищенню рівня безпеки і спростить деяку роботу її системним адміністраторам.

Метою даною роботи є розробка алгоритму виявлення аномалій мережевого трафіку з використанням нейронних мереж. Для досягнення цілі необхідно виконати наступні логічно-послідовні пункти:

- Дослідження алгоритмів виявлення аномалій у трафіку мережі.
- Обрання ознак для подальшої класифікації.
- Попередня підготовка набору даних.
- Розробка моделі нейронної мережі для виявлення аномалії.
- Тестування моделі на тестових даних та оцінка її якості.

2 ТЕОРЕТИЧНІ ОСНОВИ ВИЯВЛЕННЯ АНОМАЛІЙ У МЕРЕЖЕВОМУ ТРАФІКУ

2.1 Типи аномалій

Аномалії називаються шаблонами в даних, які не відповідають чітко визначеній характеристиці нормальних шаблонів. Вони породжуються різними зловмисними діями, наприклад, шахрайством з кредитними картками, мобільним телефоном, кібератаками, які є важливими для аналітиків даних. Важливим аспектом виявлення аномалії є природа аномалії. Аномалії можна класифікувати за такими способами [8]:

- Точкова аномалія (рис 2.1(а)): коли конкретний екземпляр даних відхиляється від звичайного шаблону набору даних, його можна вважати точковою аномалією. Для реалістичного прикладу, якщо нормальне споживання водди людиною становить 2 літри на день, але якщо воно стає п'ятдесяти літрами за будь-який випадковий день, то це точкова аномалія.

- Контекстна аномалія (рис 2.1(б)): коли екземпляр даних веде себе аномально в певному контексті, його називають контекстною або умовною аномалією; наприклад, витрати на кредитну картку в святковий період, наприклад, Різдво або Новий рік, зазвичай вищі, ніж протягом решти року. Хоча він може бути високим, він не може бути аномальним, оскільки високі витрати є контекстуально нормальними за своєю природою. З іншого боку, настільки ж високі витрати протягом несвяткового місяця можна вважати контекстною аномалією.

- Колективна аномалія (рис 2.1(в)): коли набір подібних екземплярів даних веде себе аномально по відношенню до всього набору даних, група екземплярів даних називається колективною аномалією. Наприклад, у результатах електрокардіограми (ЕКГ) людини наявність низьких значень протягом тривалого періоду часу вказує на зовнішній феномен, що відповідає

аномальному передчасному скороченню, тоді як одне низьке значення само по собі не є вважається аномальним.

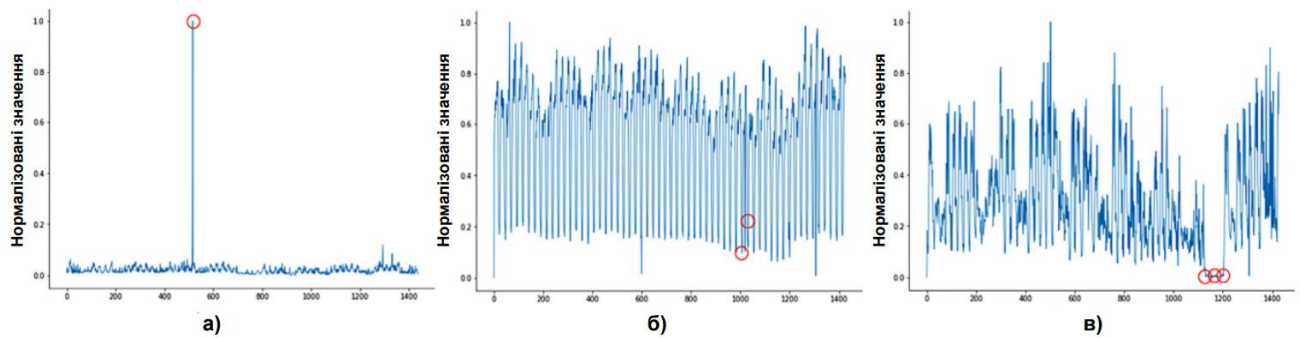


Рисунок 2.1 – Типи аномалій мережевого трафіку

На основі розглянутих типів аномалій і атак у п. 1.2 попереднього розділу кваліфікаційної роботи встановимо взаємозв'язок між атаками та аномаліями і представимо на рис. 2.2.

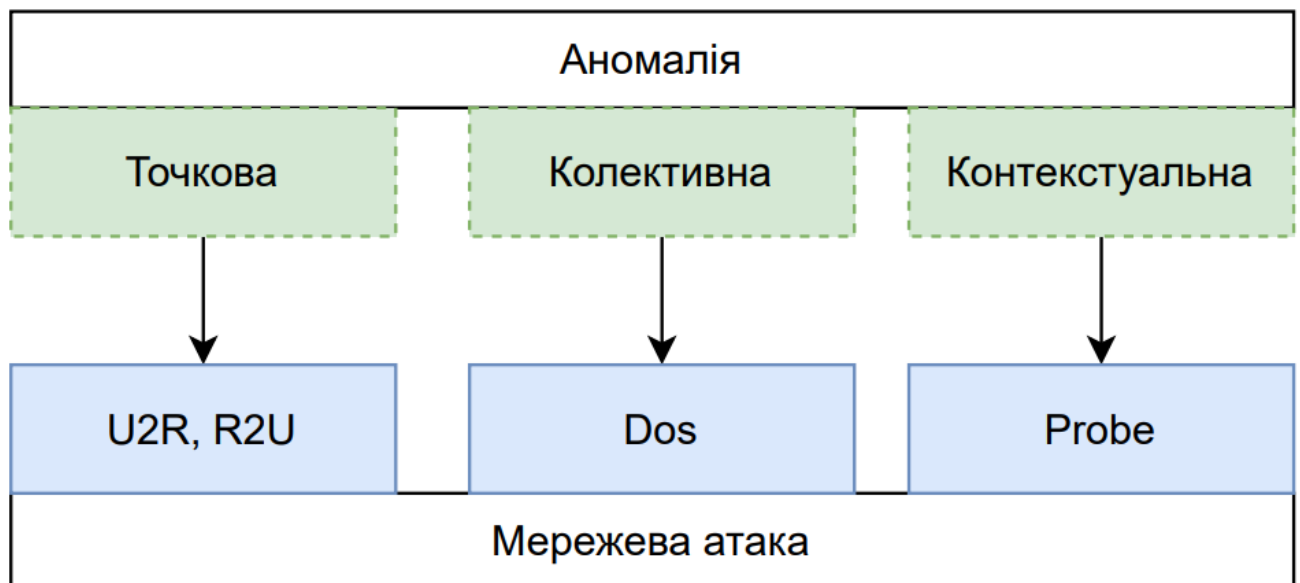


Рисунок 2.2 – Взаємозв'язок між атаками та аномаліями [8]

Характеристики DoS-атаки відповідають колективній аномалії. У разі DoS-атаки численні запити на підключення до веб-сервера є колективною аномалією, але один запит є законним. Тому будемо розглядати DoS-атаку як

колективну аномалію. Пасивні атаки (probe) засновані на конкретному намірі отримати інформацію та розвідку, здебільшого підпадає під контекстну аномалію. Атаки U2R і R2L є специфічними та складними. Запускати такі атаки нелегко в порівнянні з іншими. Тому ці атаки розглядаються як точкова аномалія.

2.2 Виявлення вторгнення в мережу

Інструменти виявлення мережових вторгнень працюють на основі базової концепції перевірки трафіку, який проходить між хостами. Від простих утиліт для прослуховування протоколів, таких як `tcpdump`, до деяких складніших інструментів перевірки, як-от `Bro`, що є екосистемою програмного забезпечення виявлення вторгнення в мережу. Вона має багато утиліт і пакетів програм, які можуть допомогти отримати сигнали від мережевого трафіку всіх видів.

Так само, як і виявлення вторгнень на хост, атаки можна ідентифікувати шляхом зіставлення трафіку з відомою сигнатурою шкідливого трафіку або шляхом виявлення аномалій, порівнюючи трафік із раніше встановленими базовими показниками.

Snort — це популярний IDS з відкритим кодом, який перевіряє пакети та мережевий трафік для виявлення аномалій у реальному часі. Це фактичний вибір для моніторингу виявлення вторгнень, що забезпечує гарний баланс зручності та функціональності. Крім того, він підтримується активною спільнотою користувачів і співавторів із відкритим кодом, які створили для нього надбудови та графічні інтерфейси. *Snort* має відносно просту архітектуру, що дозволяє користувачам виконувати аналіз трафіку в IP-мережах у режимі реального часу, писати правила, які можуть запускатися виявленими умовами, і порівнювати трафік із встановленим базовим рівнем звичайного профілю мережевого зв'язку.

При вилученні функцій для виявлення вторгнень в мережу існує суттєва різниця між вилученням метаданих мережевого трафіку та перевіркою вмісту мережевого трафіку. Перший використовується при перевірці пакетів із заповненням стану (SPI), працюючи на мережевому та транспортному рівнях,

досліджуючи заголовок і нижній колонтитул кожного мережевого пакета, не торкаючись контексту пакета. Цей підхід підтримує стан попередніх отриманих пакетів і, отже, може асоціювати нові отримані пакети з раніше побаченими пакетами. Системи SPI можуть знати, чи є пакет частиною рукостискання для встановлення нового з'єднання, ділянкою існуючого мережевого з'єднання або несподіваним шахрайським пакетом. Ці системи корисні для забезпечення контролю доступу – традиційної функції мережевих брандмауерів – оскільки вони мають чітке уявлення про IP-адреси та порти, які беруть участь у листуванні. Вони також можуть бути корисними для виявлення дещо складніших атак рівня 3/4, таких як спуфінг IP, атаки TCP/IP (наприклад, “отруєння” кешу ARP або переповнення SYN) і DDoS. Однак існують очевидні обмеження щодо обмеження аналізу лише заголовками та нижчими колонтитулами пакетів. Наприклад, SPI не може виявити ознаки порушень або вторгнень на рівні програми, оскільки для цього буде потрібно більш глибокий рівень перевірки.

Ще одним інструментом перевірки даних є *глибока перевірка пакетів (DPI)*. DPI є процесом перевірки даних, інкапсульованих у мережеві пакети, на додаток до заголовків і нижніх колонтитулів. Це дозволяє збирати сигнали та статистичні дані про мережеву кореспонденцію, що надходять із прикладного рівня. Завдяки цьому DPI здатний збирати сигнали, які можуть допомогти виявити спам, зловмисне програмне забезпечення, вторгнення та тонкі аномалії. DPI в режимі реального часу є складною проблемою інформатики через обчислювальні вимоги, необхідні для розшифровки, розбирання та аналізу пакетів, що проходять через перетин мережі.

Bro – одна з найперших систем, яка реалізувала систему пасивного моніторингу мережі для виявлення вторгнень в мережу. Bro складається з двох компонентів: ефективного механізму подій, який витягує сигнали з живого мережевого трафіку, і механізму політики, який використовує події та сценарії політики та виконує відповідні дії у відповідь на різні спостережувані сигнали.

Вро застосовується для виявлення підозрілої активності у веб-програмах, перевіряючи рядки, присутні в тілі POST запитів HTTP. Наприклад, виявлення ін'єкції SQL та атаки відображення міжсайтових сценаріїв (XSS), створивши профіль вмісту основного тексту POST для певної точки входу веб-програми [7].

Набір функцій, які генеруються через DPI для виявлення аномалій, сильно залежить від характеру програм, які працюють у мережі, а також від векторів загроз, що стосуються інфраструктури. Якщо мережа не містить жодних зовнішніх веб-серверів, використання DPI для виявлення атак XSS не має значення. Якщо ваша мережа містить лише системи точки продажу, підключені до баз даних PostgreSQL, у яких зберігаються дані клієнтів, то слід зосередитися на несподіваних мережевих з'єднаннях, які можуть свідчити про те, що злоумисник обертається у вашій мережі.

Якщо DPI використовується в середовищі з функцією безпеки транспортного рівня/рівнем безпечних сокетів (TLS/SSL), де пакети, що перевіряються, зашифровані, програма, яка виконує DPI, повинна зупинити дію SSL. DPI по суті вимагає, щоб система виявлення аномалій працювала як людина посередині, а це означає, що зв'язок, що проходить через точку перевірки, більше не є наскрізною захищеною. Ця архітектура може становити серйозну загрозу безпеці та/або продуктивності вашого середовища, особливо у випадках, коли припинення SSL та повторне шифрування пакетів реалізовано неправильно. Вам потрібно дуже уважно переглянути та перевірити методи генерації функцій, які перехоплюють трафік TLS/SSL, перш ніж розгортати їх у середовищі.

2.2.1 Ознаки для виявлення вторгнення в мережу

Група спеціальних інтересів із виявлення знань та інтелекту даних (SIGKDD) від Асоціації обчислювальної техніки (ACM) щороку проводить Кубок KDD, що ставить перед учасниками різні завдання. Кілька років тому темою було “виявлення вторгнень у комп'ютерну мережу”, завданням якої було “вивчити прогнозу модель, здатну розрізняти законні та нелегітимні з'єднання

в комп'ютерній мережі”. Цей штучний набір даних дуже старий проте з кожним роком додаються нові дані та характеристики: агрегування транзакцій за IP-адресами, геолокацією, мережевими блоками (наприклад, /16, /14), префіксами BFP, інформацією про номер автономної системи (ASN), кількість створених файлів, кількість оболонок, чи є входження як хост, кількість помилок, кількість файлів для відкриття, кількість оболонок.

2.2.2 Виявлення аномалій за допомогою даних і алгоритмів

Перш ніж вибрати алгоритм, важливо подумати про характер і якість джерела даних. Рівень “забрудненості” аномаліями прямо корельовано впливає на методологію виявлення. Якщо дані не містять аномалій (або мають аномалії, позначені, щоб ми могли їх видалити), то таке завдання є виявленням новизни. В іншому випадку це буде задача виявленням аномалій. При виявленні аномалій обраний алгоритм повинен бути нечутливим до малих відхилень, які погіршують якість навченої моделі. Часто визначити, який підхід слід прийняти, є нетривіальним рішенням. Очищення набору даних для видалення аномалій є трудомістким, а іноді й зовсім неможливим. Якщо не має уявлення про те, чи містять дані аномалії, то краще почати з припущення, що це так, і ітераційно рухатися до кращого рішення.

Класифікація технік виявлення мережових аномалій представлено на рис. 2.3.

Техніки виявлення мережових аномалій			
Статистичний метод	Класифікація	Кластеризація	Теорія інформації
Модель суміші	Байсова мережа		
Обробка сигналів	Support Vector Machine		
Principal Component Analysis	Нейронні мережі		
	На основі критерію		

Рисунок 2.3 – Класифікація технік виявлення мережових аномалій

Методи, засновані на класифікації, ґрунтуються на глибоких знаннях експертів про характеристики мережевих атак. Коли експерт надає ознак системі виявлення, атака з відомим шаблоном може бути виявлена відразу після її запуску. Це залежить виключно від сигнатури атаки як системи, яка здатна виявити атаку лише в тому випадку, якщо її підпис раніше надав експерт мережі. Це демонструє систему, яка може виявляти лише те, що їй відомо, при цьому вона буде уразливою до нових атак, які постійно з'являються в різних. Навіть якщо новий підпис атаки буде створено та включено в систему, початкові втрати є незамінними, а процедура відновлення надзвичайно дорога.

Підходи, засновані на класифікації, спираються на звичайний шаблон мережевого трафіку, який створює базу знань, і розглядає діяльність, яка відхиляється від базового шаблону, як аномальну. Перевага полягає в їх здатності виявляти атаки, які є абсолютно новими, за умови, що вони демонструють значні відхилення від нормального профілю. Крім того, оскільки звичайний трафік, не включений до бази знань, вважається атакою, виникнуть випадкові помилки. Відповідно, методам виявлення аномалій необхідне навчання, щоб побудувати звичайний шаблон активності, який займає багато часу, а також залежить від наявності повністю нормальних наборів даних про трафік. На практиці отримати екземпляри трафіку без атак дуже дорого. Більше того, у сучасних динамічних мережевих середовищах, що розвиваються, надзвичайно важко підтримувати звичайний шаблон в актуальному стані.

Основний принцип методу опорних векторів SVM полягає у виведенні гіперплощини, яка максимізує межу розділення між позитивними та негативними класами [9]. Цікавою властивістю SVM є те, що він є приблизною реалізацією принципу мінімізації структурного ризику, заснованого на статистичній теорії навчання. Стандартний алгоритм SVM – це методика навчання з вчителем, яка вимагає промаркованих даних для створення правила класифікації. Однак його також можна адаптувати як алгоритм навчання без вчителя, за допомогою якого він намагається відокремити весь набір навчальних

даних від його походження, тоді як звичайний контрольований SVM намагається розділити два класи даних у просторі ознак гіперплощиною. У статті [10] концепція неконтрольованого SVM використовується для виявлення аномальних подій. Алгоритм знаходить гіперплощини, які відокремлюють екземпляри даних від їх джерела з максимальним запасом, а потім вирішується задача оптимізації, щоб визначити найкращу гіперплощину [11].

Байєсова мережа є ефективним підходом до моделювання області, що містить невизначеність. Дискретна випадкова величина представлена за допомогою орієнтованого ациклічного графа (DAG), де кожен вузол відображає стан випадкової величини і містить таблицю умовних ймовірностей (Conditional Probability Table, CPT). Завдання CPT – забезпечити ймовірність перебування вузла в певному стані. У байєсівській мережі між вузлами існує зв'язок, що вказує на те, що змінна, представлена дочірнім вузлом, залежить від змінної, представленої батьківськими вузлами. Оскільки цю мережу можна використовувати для схеми класифікації подій, вона також застосовна для виявлення аномалій мережі. У статті [12], визначено дві основні проблеми, викликані високими хибними результатами в техніках виявлення аномалій. Передбачається, що системи виявлення аномалій містять ряд моделей для аналізу різних ознак події. Перша проблема полягає в тому, що моделі, які забезпечують оцінку або ймовірність нормальності/ненормальності події, вимагають, щоб система виявлення аномалій об'єднувала свої різні результати, що призводить до високої кількості помилкових спрацьовувань. По-друге, системи виявлення аномалій не можуть впоратися з незвичайними, але законними поведінками, наприклад, раптовим збільшенням використання ЦП, використання пам'яті тощо. Якщо ця проблема виникає, додаткова інформація може пояснити незвичайну поведінку, яка не є аномальним, ігнорується.

Властивості нейронної мережі для класифікації даних також використовувались для виявлення аномалій мережі. Нейронні мережі використовуються в різних областях застосування, таких як обробка зображень і мовлення, але вони мають високі обчислювальні вимоги. Для виявлення

аномалій мережі нейронна мережа була об'єднана з іншими методами, такими як статистичний підхід та його варіанти. У [13], нейронна мережа RNN використовується для надання коефіцієнта віддаленості для аномального мережевого трафіку. Це багат шарове сприйняття з прямим зв'язком із трьома прихованими шарами, розміщеними між вхідним і вихідним шарами. Його мета - відтворити шаблон вхідних даних на вихідному рівні з мінімізованою помилкою за допомогою навчання.

Хоча обробка сигналів є цікавою областю дослідження, використання такої техніки для виявлення аномалій майже не досліджувалося. У Thottan and Ji (2003) представлена методика статистичної обробки сигналів, заснована на виявленні різких змін. Автори описують мережеві аномалії двома способами: аномалії відповідають збоєм і проблемам продуктивності; проблеми, пов'язані з безпекою, наприклад атаки DoS.

Інформаційно-теоретичні заходи можуть бути використані для створення відповідної моделі виявлення аномалій. У [14] для пояснення характеристик набору даних використовуються кілька показників, таких як ентропія, умовна ентропія, відносна ентропія, отримання інформації та вартість інформації для виявлення аномалій.

2.3 Проблеми використання МН та штучного інтелекту для виявлення аномалій

Виявлення аномалій ґрунтується на принципово іншій парадигмі ніж просто машинне навчання для прикладу в рекомендаційних системах. Ціна помилок у виявленні вторгнень чи аномалій величезна. Неправильна класифікація однієї аномалії може призвести до серйозного порушення в системі. Створення помилкових сповіщень має менш різкий вплив, але помилкові хибні спрацьовування можуть швидко погіршити довіру до системи, навіть призводячи до того, що сповіщення повністю ігноруються. Через високу вартість помилок класифікації повністю автоматизовані наскрізні системи

виявлення аномалій, які працюють виключно за допомогою машинного навчання чи штучного інтелекту, зустрічаються не часто – майже завжди є людина, яка перевіряє, чи є оповіщення актуальними, перш ніж виконувати будь-які дії.

Семантичний розрив є справжньою проблемою МН та штучного інтелекту в багатьох середовищах. Порівняно зі статичними наборами правил або евристичними методами, іноді може бути важко пояснити, чому подія була позначена як аномалія, що призводить до більш тривалих циклів розслідування інцидентів. У практичних випадках інтерпретація або пояснення результатів часто настільки ж важлива, як і точність результатів. Особливо для систем виявлення аномалій, які з часом постійно розвивають свої моделі прийняття рішень, варто інвестувати інженерні ресурси в компоненти системи, які можуть генерувати зрозумілі людині пояснення для попереджень, створених системою машинного навчання. Наприклад, якщо попередження подається системою виявлення викидів, що працює на основі SVM одного класу, використовуючи приховану комбінацію функцій, вибраних за допомогою методів визначення розмірності, людям може бути важко зрозуміти, які комбінації явних сигналів шукає система. Наскільки це можливо, враховуючи непрозорість багатьох процесів штучного інтелекту буде корисно створити пояснення, чому модель прийняла рішення.

Розробити надійну схему оцінки для систем виявлення аномалій може бути навіть складніше, ніж побудувати саму систему. Оскільки виявлення аномалій на даних часових рядів передбачає, що існує можливість введення даних, яких ніколи не було в минулому, не існує комплексного способу оцінки системи, враховуючи величезні можливості різних аномалій, з якими система може зіткнутися в реальності.

З ВИЯВЛЕННЯ АНОМАЛІЙ МЕРЕЖЕВОГО ТРАФІКУ З ВИКОРИСТАННЯМ C-LSTM МЕРЕЖІ

3.1 Етапи розробки алгоритму ідентифікації аномалій мережевого трафіку

В пункті 1.5 було чітко встановлено поетапні завдання, вирішення яких в комплексі дозволить побудувати адекватну модель для виявлення аномалій на основі промаркованого набору даних. Розглянемо їх детальніше та сформуємо перелік необхідних інструментів для їх досягнення (рис 3.1).



Рисунок 3.1 – Етапи ідентифікації аномалій мережевого трафіку

Спершу потрібно візуально дослідити характеристики даних та дослідити їх природу виникнення, щоб в подальшому обрати правильну модель класифікатора. Наступним кроком є попередньої обробки даних: видалення зразків із відсутніми даними, їх нормалізація, об'єднання в одну часову серію. Для виконання таких дій було використано бібліотека Pandas та NumPy.

Другий етап характеризується застосуванням моделі, її тренуванням на тренувальному наборі даних, порівняння результатів із іншими відомими моделями МН та штучного інтелекту. Для виконання задачі другого етапу було використано бібліотеку Keras та її методи Sequential, Dense, Reshape, Conv2D, Flatten, MaxPooling1D, Conv1D, LSTM та набір оптимізаторів.

Результати порівняння якості класифікації обраної моделі та відомих класифікаторів буде виконано на третьому етапі дослідження з метою підтвердження обґрунтування щодо моделі.

Для ергономічного представлення даних та візуалізації були використані бібліотеки matplotlib і seaborn на кожному з етапів.

3.2 Дослідження вхідного набору даних та його попередня обробка

Для виявлення аномалій у роботі було використано набір даних від Yahoo Webscope S5. (США). Він складається із 367 часових рядів, кожен з яких з 1500 точок даних, що становить 5 050 000 точок даних (<https://research.yahoo.com/>) відповідно. Набір даних містить чотири класи: A1, A2, A3 і A4. Класи містять 67, 100, 100 і 100 файлів відповідно. Клас A1 містить дані про трафік із реальних серверів, але класи A2, A3 і A4 містять синтетичні дані про аномалії зі зростаючим рівнем складності [15].

Для перевірки запропонованої архітектури виявлення аномалій ми використали клас A1 тестового набору даних Yahoo Webscope S5, який складається з 67 файлів.

Файли цього класу складаються з реальних даних трафіку. Кожен файл зображував значення з різними часовими мітками, на основі яких вони були класифіковані на нормальні та аномальні. Зібрані дані представлені часовими рядами значень вимірювання трафіку з фактичних серверів у одиницях однієї години. Аномальні значення були позначені вручну.

Для початку роботи з набором даних та класифікаторами було імпортовано у робочу область необхідні бібліотеки (Лістингу 3.1)

Лістинг 3.1 – Імпорт необхідних бібліотек

```
import pandas as pd
import numpy as np
from sklearn import preprocessing
import glob
import matplotlib.pyplot as plt
from keras.models import Sequential
from keras.layers import Dense, Reshape, Conv2D, Flatten, MaxPooling1D, Conv1D, LSTM
from keras import optimizers
from keras.utils import to_categorical
from sklearn.metrics import confusion_matrix, accuracy_score, precision_score, recall_score, f1_score
```

Проглянемо структуру даних, записану в одному файлі та конкатинуємо один за одним всі 67 файлів в один набір даних (лістинг 3.2).

Лістинг 3.2 – Дослідження даних та конкатинація 67 файлів даних

```
#view of a sample of data
df=pd.read_csv("/content/drive/My Drive/dataset/A1Benchmark/real_1.csv")
df.head()
from sklearn import preprocessing
import glob
path = r'/content/drive/My Drive/dataset/A1Benchmark' #set the path accordingly
all_files=glob.glob(path+"/*.csv")

dataset_conc=[]
for filename in all_files:
    df=pd.read_csv(filename,index_col=None,header=0)
    df=df.replace(0,np.nan)
    df=df.dropna(axis=0, how='any',subset=['value'])
    df.value = preprocessing.normalize([df.value]).T
    dataset_conc.append(convert_2d(df))
frame=pd.concat(dataset_conc,axis=0,ignore_index=True)
```

	timestamp	value	is_anomaly
0	1	0.000000	0
1	2	0.091758	0
2	3	0.172297	0
3	4	0.226219	0
4	5	0.176358	0

На рис. 3.2 візуально представлено нормалізовані дані одного з файлів (по осі абсцис – час, по осі ординат – значення трафіку).

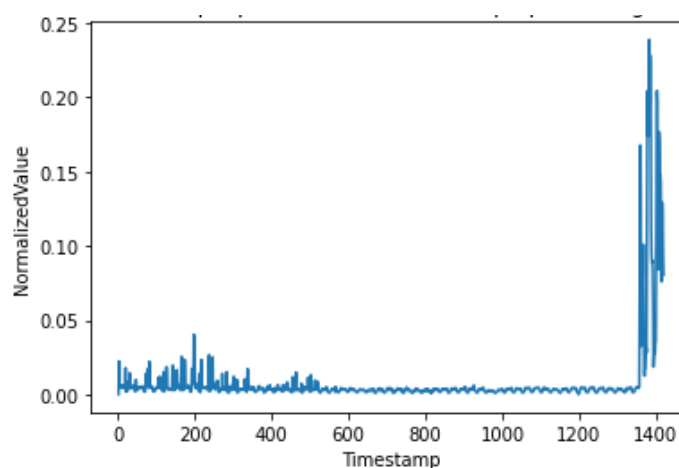


Рисунок 3.2 – Графік залежності трафіку від часу на основі даних файлу 3.csv із набору Yahoo Webscope S5

Необхідно зазначити, що перед процедурою конкатинації дані з кожного файлу були нормалізовані (`preprocessing.normalize` з лістингу 3.2), а зразки з відсутніми значеннями були видаленими (функції `replace` та `dropna` з лістингу 3.2).

В результаті отримали набір даних представлений у наступному форматі:

Лістинг 3.3 – Формат даних для подальшого навчання класифікатора

frame.head()																			
	42	43	44	45	46	47	48	49	50	51	52	53	54	55	56	57	58	59	60
23410	0.024310	0.022960	0.025661	0.025661	0.026111	0.021834	0.028137	0.024761	0.026786	0.025211	0.032414	0.028812	0.032864	0.029488	0.032414	0.029938	0.029713	0.029938	0.027687
24310	0.022960	0.025661	0.025661	0.026111	0.021834	0.028137	0.024761	0.026786	0.025211	0.032414	0.028812	0.032864	0.029488	0.032414	0.029938	0.029713	0.029938	0.029488	0.027687
22960	0.025661	0.025661	0.026111	0.021834	0.028137	0.024761	0.026786	0.025211	0.032414	0.028812	0.032864	0.029488	0.032414	0.029938	0.029713	0.029938	0.029488	0.027687	0.029938
25661	0.025661	0.026111	0.021834	0.028137	0.024761	0.026786	0.025211	0.032414	0.028812	0.032864	0.029488	0.032414	0.029938	0.029713	0.029938	0.029488	0.027687	0.029938	0.027687
25661	0.026111	0.021834	0.028137	0.024761	0.026786	0.025211	0.032414	0.028812	0.032864	0.029488	0.032414	0.029938	0.029713	0.029938	0.029488	0.027687	0.029938	0.027687	0.029938

Попередньо опрацьований та отриманий набір даних `frame` будемо використовувати як вхід на класифікатор. Візуальне представлення 1D послідовності трафіку зображено на рис. 3.3.

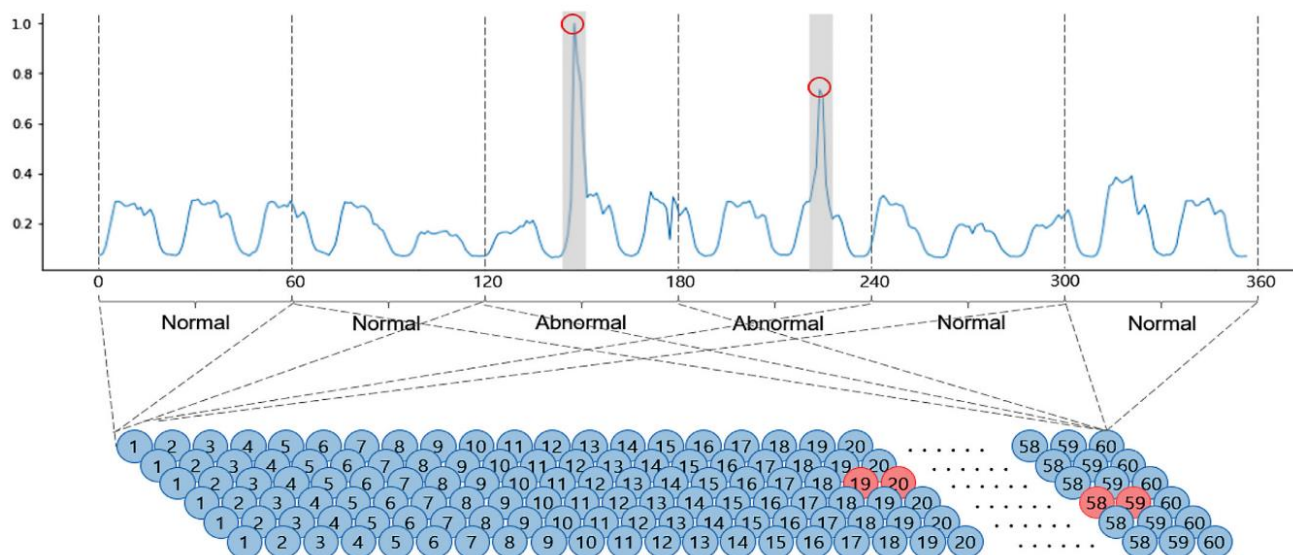


Рисунок 3.3 – 1D послідовності трафіку

Результати дослідження розподілу аномалій та нормальних значень у наборі даних свідчать про незбалансованість: всього записів є 94886, лише 1669

є аномальними, що становить 0,02% від загальної кількості. Для вирішення дисбалансу даних було використано алгоритм ковзного вікна, завдяки якому було створено 90910 вікон, з яких 8470 були промаркованими, як аномальні.

3.3 Обґрунтування обрання моделі

Аномальні шаблони чи “сплекси” (екстремуми) можуть містити як локальні, так і глобальні відхилення. Хоча глобальні аномалії можна легко ідентифікувати шляхом звичайного огляду, локальні аномалії мають шаблони, подібні до нормального сигналу, а це означає, що виявити на основі екстремумів складно.

На рис. 3.4 представлені два графіки аномалій трафіку. Рис. 3.4 (а) показує локальну аномалію, де аномалія існує всередині трафіку, а рис. 3.4 (б) показує глобальну аномалію, в якій аномалія існує поза трафіком.

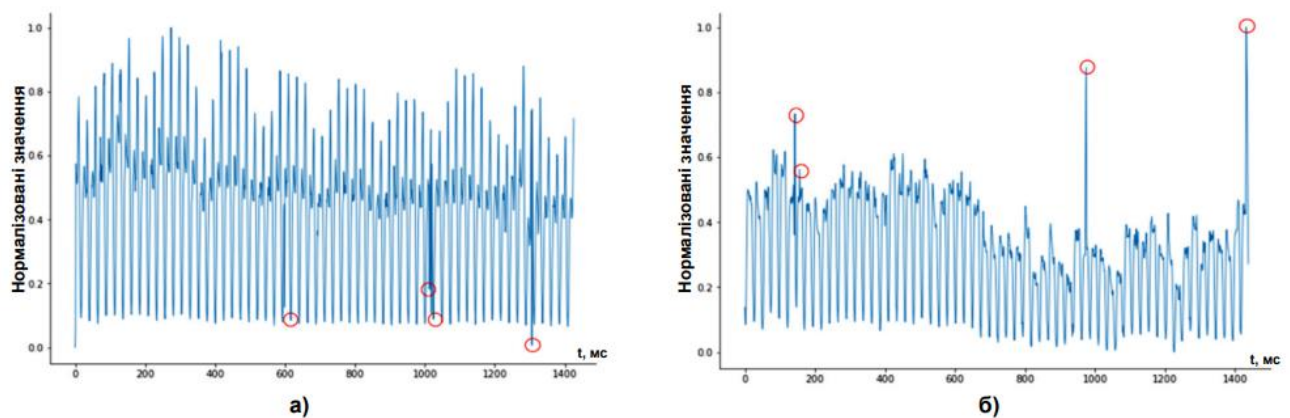


Рисунок 3.4 – Приклади графіків аномалій трафіку

Проаналізувавши рис. 3.4 було зроблено висновок, що для виявлення аномалій потрібно працювати одночасно і в просторовій і в часовій області. Оскільки, щоб виявити глобальну аномалію достатньо встановити певний рівень і слідкувати за його перевищенням, що і буде ідентифікатором аномалії. Проте з локальною аномалією не все так просто, і для її виявлення потрібен певний комплекс моделей.

Останнім часом глибоке навчання активно вивчається для обробки зображень і сигналів. Це техніка, яка знаходить ключові функції у великій кількості даних або дуже складних даних за допомогою комбінації кількох методів нелінійного перетворення. Глибоке навчання найкраще представлено двома алгоритмами: згортковими нейронними мережами (CNN) для розпізнавання зображень і рекурентними нейронними мережами (RNN), які в основному використовуються для обробки природної мови та розпізнавання мовлення [15]. CNN мають локальне сприйнятливое поле та фільтри, які можуть відображати просторові характеристики шляхом вилучення основних візуальних ознак, таких як орієнтовані краї, кінцеві точки та кути [16]. RNN має дуже глибоку структуру, яка з'єднує основні нейронні одиниці в хронологічному порядку і, як правило, ефективна для моделювання послідовних даних шляхом навчання за допомогою блоків шлюзів, таких як одиниці довгострокової пам'яті (LSTM).

Крім того, вивчається комбінація шарів CNN і LSTM, щоб виділити часові та просторові особливості, що і є необхідно для досліджуваного набору даних. Оскільки розпізнавання мовлення та обробка природної мови мають тимчасову та просторову інформацію, комбінація CNN і LSTM може ефективно витягувати ознаки [17]. В даний час ці переваги застосовуються для класифікації та прогнозування сигналів з датчиків, які використовуються в індустріальній сфері [18].

Розглянемо коротко специфіку роботи LSTM, вперше яка була представлена Зеппом Хохрайтер і Юргеном Шмідхубер в 1997 році [19]. Дана нейронна мережа є підвидом рекурентної нейронної мережі (рис. 3.5).

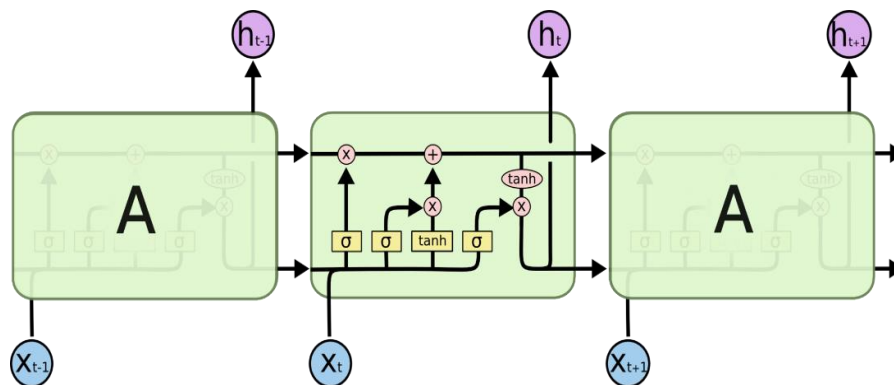


Рисунок 3.5 – Повторювальний модуль LSTM

Ключовим компонентом є стан комірки – горизонтальна лінія, що знаходиться у верхній частині модуля. Інформація може проходити через стан комірки і не змінюватись в ній. Проте модуль може і змінювати інформацію в ній за допомогою фільтрів (gates), що складаються із шару сигмоїдальної нейронної мережі та операції множення. Сигмоїдальний шар на виході завжди матиме або “0” або “1”, що позначатиме видалення чи додавання частки інформації з модуля до стану комірки. Як бачимо з рисунку у повторювального модуля є три таких фільтри. Стан комірки відіграє роль довгої пам’яті, а результати роботи фільтрів відображає коротку пам’ять, звідси і походить назва мережі LSTM.

LSTM надзвичайно добре виконує задачі розпізнавання відео контексту разом із CNN, також застосовується для прогнозування даних на основі їх попередніх значень, для розпізнавання природньої мови. За аналогією до відео контексту використаємо досліджуваний набір даних, за ідентифікацію аномалій в часовій області відповідатиме LSTM, а за просторову область CNN. Таким чином можливо буде визначити локальні аномалії.

3.4 Навчання моделі класифікатора

Представимо ідею комбінованої нейронної мережі описану у п. 3.2 у вигляді рис. 3.6.

Запропонована комбінована модель C-LSTM складається з шарів CNN і LSTM, що з'єднані у лінійну структуру. Просторові ознаки у виділяються шарами згортки та об'єднання. Часові ознаки потім витягуються шарами LSTM. За шаром згортки слідує функція активації. Це дозволяє CNN фіксувати складні характеристики у вхідному набору даних.

Вхідними даними для C-LSTM є вектор даних довжиною 60, який проходить через LSTM після проходження шарів згортки та об'єднання. Було використано tanh як функцію активації C-LSTM. Tanh – це функція, яка масштабує та зміщує сигмовидну функцію, тому tanh навчається швидше, ніж сигмовидна, якщо вона використовується як функція активації. Потім навчена модель виконує виявлення аномалії на даних тесту за допомогою функції softmax.

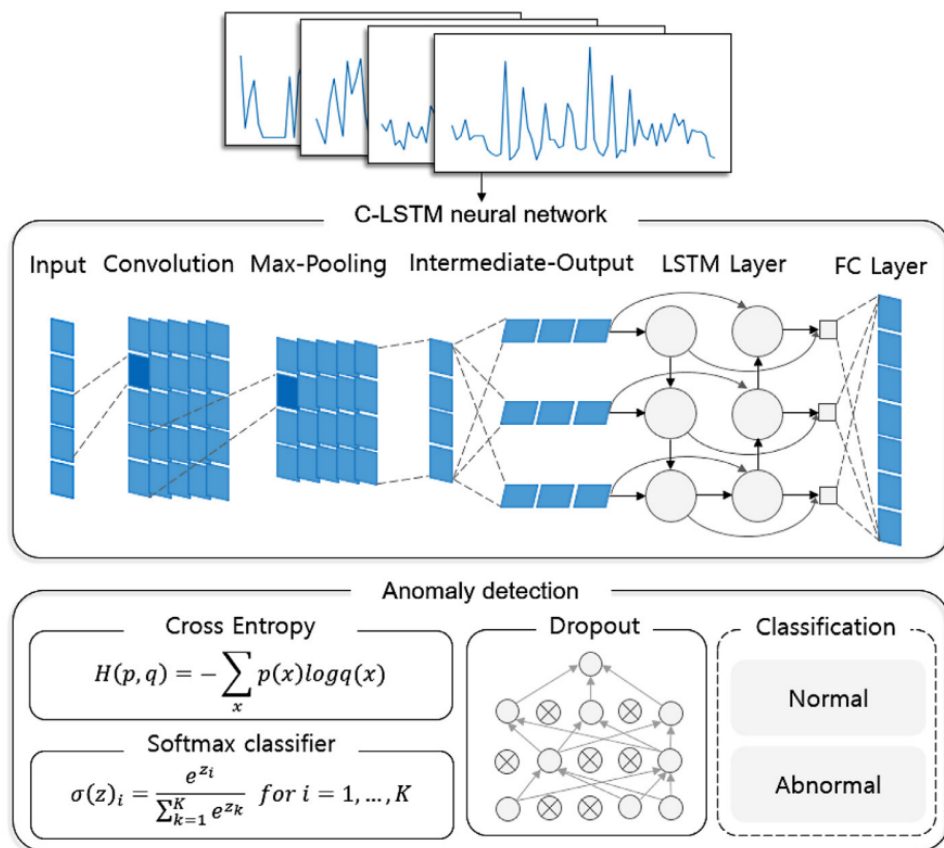


Рисунок 3.6 – Структура комбінованої нейронної мережі C-LSTM

Для формування тестового та тренувального набору було виконано операції:

Лістинг 3.4 – Формування наборів для навчання мережі

```
#train test split
y=frame.iloc[:, 60]
X=frame.iloc[:, 0:60]
X_train = X[:int(X.shape[0]*0.7)]
X_test = X[int(X.shape[0]*0.7):]
y_train = y[:int(X.shape[0]*0.7)]
y_test = y[int(X.shape[0]*0.7):]

#reshaping the data
X_train=X_train.to_numpy()
nrows, ncols = X_train.shape
X_train = X_train.reshape(nrows, ncols, 1)

X_test=X_test.to_numpy()
nrows, ncols = X_test.shape
X_test = X_test.reshape(nrows, ncols, 1)

y_test = y_test.to_numpy()
```

Комбінована нейронна мережа у вигляді програмного коду представлена лістингом 3.5.

Лістинг 3.5 – Створення моделі класифікатора C-LSTM

```
model=Sequential()
model.add(Conv1D(64, kernel_size=5, strides=1, padding='same', activation='tanh',input_shape=(60, 1)))
model.add(MaxPooling1D(pool_size=2, strides=2))
model.add(Conv1D(64, kernel_size=5, strides=1, padding='same', activation='tanh'))
model.add(MaxPooling1D(pool_size=2, strides=2))
model.add(Reshape((1,15*64)))
model.add(LSTM(64, activation='tanh',return_sequences='False'))
model.add(Flatten())
model.add(Dense(32, activation='tanh'))
model.add(Dense(2, activation='softmax'))
model.compile(loss='binary_crossentropy', optimizer='adam',metrics=['accuracy'])
model.summary()
```

Type	#Filter	Kernel size	Stride	# Param
Convolution	64	5	1	384
Activation (tanh)	–	–	–	0
Pooling	–	2	2	0
Convolution	64	5	1	20,544
Activation (tanh)	–	–	–	0
Pooling	–	2	2	0
LSTM (64)	–	–	–	262,400
Dense (32)	–	–	–	2080
Activation (tanh)	–	–	–	0
Dense (2)	–	–	–	66
Softmax	–	–	–	0
Total number of parameters				285,474

Для навчання нейронної мережі C-LSTM було виконано 500 епох навчання, метрикою оцінки якості роботи є точність, а для оцінки втрат використано crossentropyloss. Результати навчання наведено на рис. 3.7.

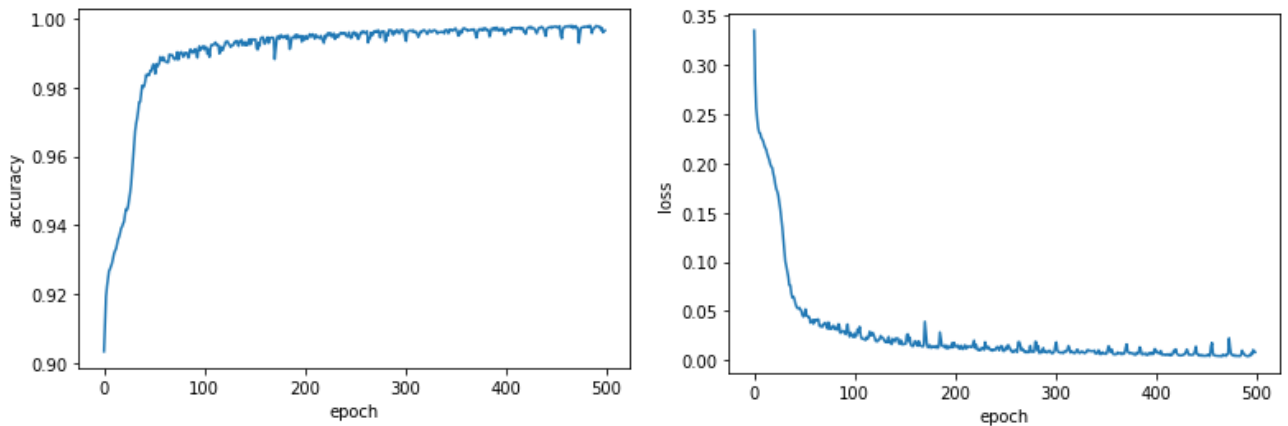


Рисунок 3.7 – Результати навчання нейронної мережі C-LSTM

3.5 Тестування C-LSTM для виявлення аномалій мережевого трафіку

В попередньому пункті модель навчалась на основі тренувальних даних, для оцінки якості класифікатора було використано 4 оцінки якості: точність, достовірність, повнота та F1-score (Лістинг 3.6) та сформовано матрицю помилок.

Лістинг 3.5 – Оцінка якості класифікатора C-LSTM

```
from sklearn.metrics import confusion_matrix, accuracy_score, precision_score, recall_score, f1_score

Confusion_Matrix = confusion_matrix(y_test, y_pred)
Accuracy = accuracy_score(y_test, y_pred)
precision = precision_score(y_test, y_pred, average='binary')
recall = recall_score(y_test, y_pred, average='binary')
F1_Score = f1_score(y_test, y_pred, average='binary')
```

```
print("Confusion_Matrix")
print(Confusion_Matrix)
print("Accuracy ", Accuracy)
print("Precision ", precision)
print("recall ", recall)
print("f1_score ", F1_Score)
```

```
Confusion_Matrix
[[22955 1230]
 [ 704 1729]]
Accuracy 0.9273423998797806
Precision 0.5843190266982089
recall 0.7106452938758734
f1_score 0.641320474774481
```

Також в роботі було проведено порівняльний аналіз із алгоритмами KNN, SVM, RF (Random Forest), DT(Decision Tree), MLP (Multilayered perceptron), результати представлені на рис. 3.8.

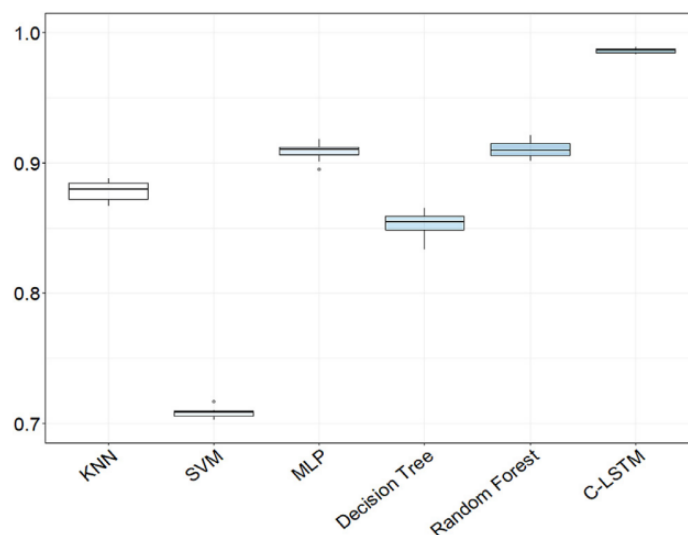


Рисунок 3.8 – Порівняння точності класифікаторів після 10-кратної перехресної перевірки

Ядро методу опорних векторів (SVM) було встановлено на лінійне значення, а L2-норма використовувалася як критерій. Критерій зупинки був встановлений на $1e-4$. Максимальна кількість ітерацій для виконання встановлена на 1000. Для KNN встановлено кількість сусідів рівним 7. Усі точки в кожній околиці зважуються однаково. Для Random Forest було встановлено кількість дерев у лісі до 20 та використано gini для вимірювання якості розколу. Запропонований C-LSTM має найвищу продуктивність порівняно з іншими

методами машинного навчання, за ним слідують алгоритми Random Forest, багатошарового персептрона (MLP) та KNN.

З використанням перехресної перевірки було досягнуто кращих результатів для класифікатора C-LSTM, а саме:

- Точність після навчання з 700 епохами становить ~94%.
- Зміна розміру ядра до 10 покращила продуктивність моделі, підвищивши точність до 95% і показник f1 до 70%.
- Після зміни функції активації на relu в кожному шарі, точність різко зросла до ~96%, а показник f1 збільшився з 64,2% до 73,5%.

Виявлення аномалій дозволяє нам обробляти дані швидше, ефективно та безпечно, що і підтверджує застосування класифікатора у вигляді нейронної мережі C-LSTM. Отже, завдяки штучному інтелекту підприємства можуть підвищити ефективність і безпеку своїх цифрових технологій операції.

4 БЕЗПЕКА ЖИТТЄДІЯЛЬНОСТІ, ОСНОВИ ХОРОНИ ПРАЦІ

4.1 Заходи захисту обладнання від статичної електрики

Накопичення зарядів статичної електрики відбувається під час користування одягом із штучного волокна, вовни, шовку, взуттям з підошвами, що не проводять електричного струму, виконання робіт з речовинами діелектриками та шліфувальною шкіркою.

Дія статичної електрики для людини безпечна, бо сила струму дуже мала, але: розряд енергії відбувається у вигляді помірного і сильного уколу або поштовху; вплив зарядів може призвести до тяжких нещасних випадків внаслідок рефлекторного руху поблизу незахищених та рухомих частин, перебування на висоті; іскрові розряди можуть спричинити спалах або вибух горючих речовин; вибухи при перевезенні рідин у незаземлених цистернах тощо.

Заходи щодо захисту від статичної електрики:

- заземлення технологічного устаткування, трубопроводів, апаратів;
- застосування загального і місцевого зволоження повітря в небезпечних приміщеннях робочої зони, якщо це допустимо за умовами виробництва;
- використання струмопровідної підлоги, а також спецвзуття зі струмопровідною підошвою, антистатичних рукавичок;
- іонізація повітря, застосування індукційних або тканинних нейтралізаторів.

Великою небезпекою є дії грозових розрядів. Всі можливі заходи для усунення небезпеки розряду атмосферної електрики, забезпечення безпеки людей, збереження будинків, устаткування і матеріалів від руйнування, вибухів і пожеж, називається захистом від блискавки.

Атмосферна електрика, електричні заряди виникають в процесі руху крапель води в атмосфері. Процес утворення, поділу і накопичення електричних зарядів у хмарах відбувається через виникнення в них могутніх висхідних

повітряних потоків з інтенсивною конденсацією водяного пару і розбризкуванням водяних крапель. Дрібнодисперсний водяний пил, що розбризкується, заряджається негативно, а важкі краплі води – позитивно.

Прямий удар блискавки зумовлює миттєве нагрівання струмо-провідних конструкцій до температури плавлення або навіть випару, вибух чи розщеплення непровідних конструкцій, вибух будинків і будівель. При такому прямому ударі блискавки в край металевого резервуару, тепла, що виділяється, достатньо для оплавлення в місці контакту сталевих листів товщиною 4 мм.

Вторинні прояви блискавки виникають через різницю потенціалів на металевих частинах устаткування, трубопроводах і струмопроводах в результаті електромагнітної й електростатичної індукції від прямого удару блискавки.

Електромагнітна індукція. При розряді блискавки в просторі виникає магнітне поле, що змінюється з часом. Це поле індукуює у металевих конструкціях, трубопроводах, електричних провідниках електрорушійну силу. У контурах із замкненою конфігурацією електричний струм викликає нагрівання конструкцій. У незамкнених металевих контурах, наприклад, у трубопровідних комунікаціях, прокладених на поверхні землі, електромагнітна індукція може викликати іскріння чи нагрівання в місцях зближення трубопроводів різних контурів.

Досвідом встановлено, що значна кількість пожеж цистерн і резервуарів з нафтопродуктами зумовлена головним чином вторинними проявами блискавки, а не прямими її ударами.

Електростатична індукція. Під грозовою хмарою у наземних об'єктах індукуються електричні заряди, рівні за величиною і протилежні за знаком зарядам хмари. Електростатичні заряди індукуються навіть на об'єктах, добре ізольованих від землі:

- на металевих дахах будинків;
- на проводах повітряних ліній зв'язку;
- на водопровідних і каналізаційних трубах;

- на електропровідниках усередині будинків та інших заземлених конструкціях.

Між об'єктом і землею потенціал тим вищий, чим вищий об'єкт. Всі металеві провідні елементи будинків і споруд заземлюють для боротьби з виникненням на спорудах і усередині будинків різниці потенціалів між устаткуванням внаслідок електростатичної індукції.

Занесення високих потенціалів у будинки і споруди відбувається через повітряні, надземні і підземні комунікації (труби, повітряні лінії зв'язку й електроенергії).

Пожежі і руйнування від розрядів атмосферної електрики відбуваються переважно від прямих ударів блискавки. Це найбільш небезпечний її прояв.

Небезпечна блискавка і для людей. Ураження можуть бути через виникнення високих напруг на окремих частинах устаткування усередині будинків й поза будинками, а також під час виникнення крокової напруги.

4.2 Значення адаптації в трудовому процесі

Сучасний розвиток технічного та технологічного стану виробництва передбачає постійну автоматизацію та оптимізацію виробничих процесів. Сьогодні, напевно, важко уявити компанію, господарська діяльність в якій здійснювалась би без використання комп'ютерної техніки. Через масовий характер робіт, що виконуються працівниками за допомогою комп'ютера, законодавством України чітко врегульовано норми та вимоги до використання комп'ютерної техніки на підприємстві, безпосередньо й охорона праці при роботі з комп'ютером.

Важливими вимоги до приміщення є наступними: приміщення, в яких планується установка та подальша робота з комп'ютером, повинні відповідати проектній документації будинку, погодженій з уповноваженими державними органами. Крім того, роботодавець повинен враховувати санітарні нормативи освітлення, вимоги до параметрів мікроклімату (температура, відносна

вологість), ступеня і сили вібрації, звукового шуму і вогнестійкості приміщення, а також характеристики електромагнітного, ультрафіолетового та інфрачервоного полів. Конкретні показники зазначених санітарних норм див. в Державних санітарних правилах і нормах роботи з візуальними дисплейними терміналами електронно-обчислювальних машин ДСанПН 3.3.2.007-98, затверджених Постановою Головного державного санітарного лікаря України №7 від 10 грудня 1998 року. Правила поширюються на умови й організацію праці при роботі з візуальними дисплейними терміналами (ВДТ) усіх типів вітчизняного та зарубіжного виробництва на основі електронно-променевих трубок (ЕПТ), що використовуються в електронно-обчислювальних машинах (ЕОМ) колективного використання та персональних ЕОМ (ПЕОМ). Так, наприклад, роботодавцю заборонено установлювати комп'ютери в приміщеннях, розташованих у підвалах будинків. Для уникнення можливих аварій та замикань, поряд з приміщеннями, де вестиметься робота з комп'ютером (над чи під ними), також не дозволяється проведення робіт, що потребують здійснення надмірно вологих технологічних процесів. Відповідне приміщення повинно бути укомплектоване системами центрального або індивідуального опалення, кондиціонування чи вентиляції повітря. Але при установці зазначених систем, необхідно переконатись, що батареї опалення, водопровідні труби, вентиляційні кабелі тощо, надійно сховані під захисними щитками, які перешкоджатимуть можливому потраплянню робітника під напругу.

У кожній кімнаті, де обладнуватимуться робочі місця співробітників, що працюватимуть на комп'ютері, повинні бути наявні елементи природного та штучного освітлення. При цьому, на вікнах слід встановити легко регульовані жалюзі чи штори, які дозволять працівникам коригувати рівень освітлення в приміщенні. Бажано розмістити комп'ютери в кімнаті таким чином, щоб світло потрапляло на екрани моніторів з півдня чи північного сходу. З метою досягнення максимального рівня безпечності і охорони праці при роботі з комп'ютером, виробничі приміщення необхідно обладнати аптечками першої медичної допомоги, системами автоматичної пожежної сигналізації і

вогнегасниками. В приміщенні, в якому разом працюють 5 або більше комп'ютерів, на видимому місці установлюється службовий вимикач, який у разі потреби дозволить повністю відключити електричне живлення кімнати.

Роботодавець, який використовує найману працю робітників, повинен забезпечити відповідність їхніх робочих місць комфортним та безпечним умовам. Розмір одного робочого місця має становити не менше 6 квадратних метрів. При необхідності, суміжні робочі місця співробітників, що працюють з комп'ютером, слід розділити перегородками висотою до 2 метрів. При визначенні достатнього розміру приміщення і робочого місця на одну особу необхідно додатково враховувати шафи, сейфи, тумби або інші предмети меблів чи обладнання, які знаходяться в кімнаті. На столі працівника можливо розмістити допоміжні для роботи пристрої (принтери, колонки, сканери), а також місця для зберігання документів, за умови, що це не обмежуватиме видимість екрану і не заважатиме працівнику. У разі надмірного шуму чи вібрації технічного обладнання, роботодавець повинен забезпечити працівників антивібраційними килимками. Робочий стілець співробітника має бути підйомно-поворотним, легко регульованим за висотою та забезпечувати належну підтримку та зручне положення спини і хребта особи. Щодня необхідно проводити вологе прибирання приміщення, та очищати робоче місце та безпосередньо монітор комп'ютера від запиленості.

На підприємстві забороняється: проводити ремонт та технічне обслуговування комп'ютера за робочим місцем працівника; самочинно ремонтувати або намагатись здійснити технічне налагодження комп'ютера без залучення компетентних спеціалістів; складувати на робочому місці зайві документи, деталі та предмети, що не потрібні для роботи; використовувати монітори з нечітким зображенням та монітори, у яких наявні поламки екрану; працювати з матричним принтером без антивібраційного покриття та зі знятою кришкою. Допускати до роботи осіб, які не пройшли затвердження на підприємстві курс охорони праці для роботи з комп'ютером, не дозволяється.

При прийнятті на роботу кожна особа має пройти лікарський огляд. Окрім того, при подальшій трудовій діяльності в компанії, така особа підлягає регулярному лікарському огляду не рідше ніж раз на 2 роки. Обов'язковим є проходження таких лікарів як терапевта, невропатолога та офтальмолога. В компанії мають бути чітко встановлені перерви для відпочинку працівників (окрім обідньої), як правило, тривалістю 10-15 хвилин раз на годину або дві, в залежності від складності роботи. В будь-якому випадку, роботодавець повинен передбачити такий розпорядок роботи на підприємстві, щоб час неперервної роботи з комп'ютером був не більше ніж 4 години. Додатково, для збереження належного рівня здоров'я та професійної придатності робітників, рекомендується виділити на підприємстві окреме побутове приміщення для перепочинку працівників і зняття ними нервово-емоційного напруження, що виникає при роботі з комп'ютером.

ВИСНОВКИ

Основу модель передачі даних в мережі OSI, типи мережевих загроз, починаючи від атак на фізичному рівні та завершуючи атаками ботнет мереж розглянуто у теоретичних розділах. Детально представлено основи захисту мережі, такі як: багатофакторна автентифікація, встановлення та розмежування рівнів доступів, сегментація мережі, захист портів від пасивного сканування, аналіз даних мережевого трафіку. Надзвичайно важливо спеціалісту кібербезпеки розуміти можливі причини появи аномалій у мережі та проводити комплексний захист превентивно.

Для автоматичного виявлення аномалій у мережевому трафіку надзвичайно важливо розуміти предметну область. В другому розділі висвітлено основні типи аномалій, представлено взаємозв'язок типів із реальними атаками. Також представлено основні проблеми використання алгоритмів на основі даних. Найчастіше розробник зустрічається із незбалансованими даними, оскільки відсоток аномальних зразків до загальної кількості даних є надзвичайно малим. Це призводить до зменшення показників якості роботи будь якого автоматизованого алгоритму виявлення аномалій. Аналогічна проблема була і у досліджуваному наборі даних, яку було вирішено з використанням техніки ковзного вікна.

В практичній частині представлено покроковий процес виявлення аномалій мережевого трафіку на основі набору даних від Yahoo з використанням комбінованої нейронної мережі C-LSTM. Її застосування обґрунтовано тим, що для виявлення локальних аномалій потрібно досліджувати як просторову так і часову область, що в комплексі є можливим для CNN та LSTM. Результати тренування нейронної мережі свідчать про можливість її впровадження у ідентифікаційні системи відстежування трафіку. Порівняльний аналіз з відомими класифікаторами МН за показниками точності свідчать про доцільність застосування комбінованої нейронної мережі C-LSTM.

В подальших дослідженнях необхідно дослідити параметри мережі та визначити оптимальні, частково ця робота була виконаною. Активаційна

функція Relu дозволяє збільшити точність класифікатора, аналогічно як і збільшення епох навчання. В результаті розроблено алгоритм для виявлення мережових аномалій з використанням C-LSTM. Імплементация цього методу виявлення значно спростить роботу системним адміністраторам та аналітикам.

СПИСОК ЛІТЕРАТУРНИХ ДЖЕРЕЛ

1. Dumas, M. and Schwartz, L., 2009. Principles Of Computer Networks And Communications. Upper Saddle River, NJ: Pearson Prentice Hall.
2. Chio, C.; Freeman, D. Machine Learning and Security, O'Reilly Media, 2018, 125-180.
3. Bhardwaj, Akashdeep & Kaushik, Keshav & Sapra, Varun. (2021). Majestic Meterpreter. 16. 106.
4. N.-Y. Lee and H.-J. Chiang, "The research of botnet detection and prevention," presented at the Computer Symposium (ICS), 2010, International, 2010, pp. 119– 124.
5. Rivest, R., Shamir, A. and Adleman, L. A method for obtaining digital signatures and public-key cryptosystems. Communications of the ACM 21, 2 (1978), 120--126.
6. Ahmed M, Mahmood A. Network traffic analysis based on collective anomaly detection. In: 2014 IEEE 9th conference on industrial electronics and applications (ICIEA), 2014. p. 1141–46
7. Anderson, N. (2007). Deep packet inspection meets 'Net Neutrality, CALEA'. Arstechnica. Retrieved from <http://arstechnica.com/articles/culture/Deep-packet-inspection-meets-netneutrality.ars>
8. Ahmed M, et al. A survey of network anomaly detection techniques. Journal of Network and Computer Applications (2015), <http://dx.doi.org/10.1016/j.jnca.2015.11.016>
9. Eskin E. Anomaly detection over noisy data using learned probability distributions. In: Proceedings of the seventeenth international conference on machine learning, ICML '00, San Francisco, CA, USA, Morgan Kaufmann Publishers Inc., 2000. p. 255–62.
10. Estevez-Tapiador JM, Garcia-Teodoro P, Diaz-Verdejo JE. Anomaly detection methods in wired networks: a survey and taxonomy. Comput Commun 2004; 27(16):1569–84.

11. Cristianini N, Shawe-Taylor J. An introduction to support vector machines: and other kernel-based learning methods. New York, NY, USA: Cambridge University Press; 2000
12. Krügel C, Toth T, Kirda E. Service specific anomaly detection for network intrusion detection. In: Proceedings of the 2002 ACM symposium on applied computing, SAC '02, ACM, New York, NY, USA; 2002. p. 201–8.
13. Hawkins S, He H, Williams G, Baxter R. Outlier detection using replicator neural networks. In: Kambayashi Y, Winiwarter W, Arikawa M, editors. Data warehousing and knowledge discovery, lecture notes in computer science, vol. 2454. Berlin, Heidelberg: Springer; 2002. p. 170–80
14. Lee W, Xiang D. Information-theoretic measures for anomaly detection. In: Proceedings of 2001 IEEE symposium on security and privacy, 2001 S P 2001; 2001. p. 130–43.
15. Laptev, N., & Amizadeh, S.. A labeled anomaly detection dataset S5 Yahoo Research, v1.
<https://webscope.sandbox.yahoo.com/catalog.php?datatype=s&did=70>.
16. Jeffrey Donahue, Lisa Anne Hendricks, Sergio Guadarrama, Marcus Rohrbach, Subhashini Venugopalan, Kate Saenko, Trevor Darrell; Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2015, pp. 2625-2634.
17. Y. Lecun, L. Bottou, Y. Bengio and P. Haffner, "Gradient-based learning applied to document recognition," in Proceedings of the IEEE, vol. 86, no. 11, pp. 2278-2324, Nov. 1998, doi: 10.1109/5.726791
18. Noble CC, Cook DJ. Graph-based anomaly detection. In: Proceedings of the ninth ACM SIGKDD international conference on knowledge discovery and data mining, KDD '03, ACM, New York, NY, USA; 2003. p. 631–6.
19. Zhou, C., Sun, C., Liu, Z., & Lau, F. (2015). A C-LSTM neural network for text classification. arXiv preprint arXiv:1511.08630