

НАУКОВА РОБОТА

**на тему: «Покращення та використання алгоритму машинного навчання
Isolation Forest для виявлення мережевих атак»**

Шифр: Forest Watchdog

ЗМІСТ

ВСТУП	3
РОЗДІЛ 1. ШИФРОВАНИЙ МЕРЕЖЕВИЙ ТРАФІК, АНОМАЛІЇ ТА ВИКЛИКИ АНАЛІЗУ.....	5
РОЗДІЛ 2. МАТЕМАТИЧНІ ЗАСАДИ АЛГОРИТМУ ISOLATION FOREST	9
РОЗДІЛ 3. МОДИФІКАЦІЯ ПРАВИЛА СПЛІТУ В ISOLATION FOREST: МІНІМІЗАЦІЯ ЧАСТКИ МЕНШОЇ ГІЛКИ ТА ЇЇ ВЛАСТИВОСТІ.....	16
ВИСНОВКИ.....	19
СПИСОК ВИКОРИСТАНОЇ ЛІТЕРАТУРИ.....	21
Додаток 1.....	24

ВСТУП

У сучасному цифровому суспільстві інформація стала одним із найцінніших ресурсів, а її захист – одним із ключових завдань кібербезпеки. Швидкий розвиток інформаційно-телекомунікаційних технологій, зростання кількості онлайн-сервісів та глобалізація мережевих інфраструктур зумовлюють постійне збільшення обсягів передавання даних у мережах загального користування [1]. Задля захисту конфіденційності користувачів і забезпечення цілісності даних активно впроваджуються криптографічні протоколи, які переводять більшість мережевого трафіку в зашифрований вигляд. За статистикою, станом на 2023 рік понад 90 % веб-трафіку у світі передавалося через протокол HTTPS, що ґрунтується на TLS [2].

З одного боку, широке застосування шифрування суттєво підвищує рівень безпеки користувачів, адже зломисники не мають прямого доступу до вмісту пакетів. З іншого - це створює серйозні виклики для систем моніторингу та виявлення загроз. Традиційні засоби захисту, що спиралися на аналіз сигнатур або портів, фактично втрачають ефективність, оскільки вміст зашифрованого трафіку є прихованим, а службові заголовки протоколів – обмеженими [3].

Таким чином, виникає так звана «сліпа зона» для класичних IDS/IPS-систем, у якій зломисники можуть маскувати шкідливу активність. Це стосується як приховування команд і шкідливого коду всередині TLS-сесій, так і організації розподілених атак типу DDoS, замаскованих під легітимні потоки даних [4].

В умовах постійного зростання кількості та складності кіберзагроз необхідно переходити від контент-орієнтованого аналізу до поведінкових підходів, що базуються на дослідженні статистико-часових характеристик трафіку [5]. Такі методи дозволяють виявляти відхилення від типової поведінки навіть тоді, коли зміст переданих даних недоступний [6].

У цьому контексті дедалі більшого значення набувають технології машинного навчання, здатні формувати моделі «нормального» трафіку й автоматично виокремлювати аномалії [7]. Особливу увагу дослідників

привертають алгоритми некерованого навчання, які не потребують попередньої розмітки даних і дають змогу знаходити раніше невідомі загрози (zero-day атаки) [8].

Отже, актуальність дослідження зумовлена поєднанням двох тенденцій: стрімким зростанням частки шифрованого трафіку в глобальних мережах та постійним ускладненням кіберзагроз, які використовують ці канали для маскування. Це створює потребу в нових теоретичних і практичних підходах до аналізу зашифрованих потоків, зокрема із застосуванням методів машинного навчання [9].

РОЗДІЛ 1

ШИФРОВАНІЙ МЕРЕЖЕВИЙ ТРАФІК, АНОМАЛІЇ ТА ВИКЛИКИ АНАЛІЗУ

Основним шляхом подолання обмежень класичних IDS/IPS у зашифрованих каналах є застосування методів машинного навчання. Вони аналізують виключно метадані: розміри пакетів, часові інтервали, напрямки потоків – і дозволяють будувати поведінкові моделі без доступу до аналізу корисного навантаження мережесих пакетів [1, 23]. Такий підхід компенсує «сліпі зони» сигнатурних методів і дає змогу виявляти загрози там, де аналіз корисного навантаження є недоступним [5].

У практиці розрізняють два основні напрями: кероване та некероване навчання. Керовані методи ґрунтуються на попередньо розмічених наборах даних, що дає змогу навчати моделі розрізняти нормальний і аномальний трафік. Прикладом є алгоритм k-Nearest Neighbors, який класифікує нові зразки за близькістю до відомих спостережень. Його сильними сторонами є простота та інтерпретованість, але він чутливий до вибору параметра k і до масштабування ознак [10]. Подібні характеристики мають дерева рішень та їх ансамблеві модифікації (Random Forest), а також нейронні мережі, які демонструють високу точність при розпізнаванні відомих атак, проте вимагають великих обсягів якісно розмічених даних [8].

Некеровані методи не потребують апріорної розмітки й орієнтовані на навчання на вибірках нормального трафіку. Відхилення від очікуваної поведінки вони визначають як потенційні аномалії. До них належать Local Outlier Factor, що оцінює локальну щільність об'єкта відносно сусідів, One-Class SVM, який будує граничну поверхню між нормальними і відхиленими точками, та Autoencoder, що виявляє аномалії за величиною похибки відновлення вхідних даних. LOF добре працює у середовищах з рівномірним розподілом даних, але втрачає ефективність при високій розмірності; One-Class SVM має високу здатність до виявлення складних залежностей, проте обчислювально затратний;

Autoencoder здатен відтворювати багатовимірні залежності, але потребує значних ресурсів [11, 15].

Окремо виділяють ансамблеві методи, серед яких найбільшу увагу отримав Isolation Forest, запропонований Liu, Ting та Zhou [13]. Його принцип полягає в тому, що рідкісні об'єкти ізолюються за допомогою випадкових розбиттів значно швидше, ніж ті, що належать до «нормальних». Це забезпечує алгоритму низку переваг: лінійно-логарифмічну масштабованість відносно розміру вибірки, здатність працювати у високовимірних просторах, відсутність потреби у попередньо розмічених даних і високу ефективність при дисбалансі класів [19, 14]. Порівняльні дослідження підтверджують, що Isolation Forest стабільно перевершує LOF та One-Class SVM за швидкодією та здатністю масштабуватися на великі вибірки [9].

Наукові публікації останніх років підкреслюють, що перспективи розвитку систем моніторингу трафіку полягають у поєднанні кількох підходів: використанні некерованого навчання для виявлення невідомих атак та застосуванні глибокого навчання для підвищення точності. Поєднання статистико-часових характеристик із алгоритмами машинного навчання дозволяє компенсувати відсутність доступу до вмісту пакетів і забезпечувати ефективний захист навіть у середовищі повсюдного шифрування [6].

У сфері інформаційної безпеки поняття аномалії трактується як відхилення від очікуваної або типової поведінки системи. У контексті мережевого трафіку це означає будь-яку подію чи послідовність пакетів, які не відповідають встановленій «нормі» для даної мережі. Виявлення таких відхилень дозволяє ідентифікувати як технічні проблеми, так і потенційні кіберзагрози [12].

Аномалії у мережевому трафіку прийнято класифікувати на три основні групи [2]. До точкових належать окремі спостереження, які суттєво відрізняються від інших (наприклад, поява пакета із нетипово великим розміром). Контекстуальні аномалії проявляються лише у певному часовому чи просторовому контексті, як-от різке зростання обсягу трафіку вночі. Нарешті,

колективні аномалії утворюються сукупністю подій, що в ізоляції можуть здаватися нормальними, але разом формують нетипову поведінку - типовий приклад становлять DDoS-атаки.

Джерела аномальної поведінки можуть бути різноманітними [2]. Вони виникають як через технічні збої (помилки в роботі маршрутизаторів, зациклення пакетів, перевантаження каналів зв'язку), так і внаслідок навмисних дій зловмисників (атаки типу DDoS, сканування портів, експлуатація вразливостей). Окремим фактором є нетипова поведінка користувачів, наприклад, передавання великих обсягів даних у неробочий час чи підключення до сумнівних вузлів.

Особливу складність становить аналіз зашифрованого трафіку, де зміст пакетів недоступний для перевірки. У таких випадках системи безпеки працюють переважно з метаданими – розмірами пакетів, часовими інтервалами та напрямками потоків. Це значно знижує ефективність класичних сигнатурних методів, проте створює умови для застосування поведінкових моделей та алгоритмів машинного навчання [8, 12].

Сучасні дослідження вказують, що найбільш інформативними характеристиками для виявлення аномалій у зашифрованому трафіку є розмір і частота пакетів, часові інтервали між ними, а також обсяг переданих даних [2, 12]. Додатково використовують аналіз напрямів з'єднань та особливостей протоколів, оскільки нетипові IP-адреси, географічно віддалені вузли або використання нестандартних портів можуть сигналізувати про спробу обходу засобів захисту. Сукупність цих параметрів дає змогу будувати профілі «нормальної» поведінки і виявляти відхилення навіть без розшифрування даних.

У практиці кіберзахисту найбільш поширеними прикладами аномалій є DDoS-атаки, масове сканування портів, функціонування ботнетів та нетипова користувацька активність, наприклад, завантаження великих обсягів даних у нетиповий час [2, 12]. Такі явища практично неможливо виявити лише сигнатурними методами, проте вони добре простежуються за статистичними характеристиками потоків.

Таким чином, аналіз аномалій у мережевому трафіку є ключовим напрямом розвитку сучасних систем кіберзахисту. В умовах повсюдного застосування шифрування класичні методи втрачають ефективність, і на перший план виходить вивчення статистико-часових характеристик потоків та використання алгоритмів машинного навчання, особливо некерованих методів, що дозволяють будувати моделі поведінки мережі та виявляти раніше невідомі загрози [8].

РОЗДІЛ 2

МАТЕМАТИЧНІ ЗАСАДИ АЛГОРИТМУ ISOLATION FOREST

Isolation Forest ґрунтується на припущенні, що аномальні об'єкти можна відокремити випадковими розбиттями значно швидше, ніж звичайні спостереження. Під «відокремити (ізолювати)» мається на увазі послідовно ділити простір випадковими осьовими розбиттями (обираємо випадкову ознаку і випадковий поріг) доти, доки цільовий об'єкт не опиниться один у листі дерева.

Аномалії рідкісні та «віддалені» від масиву нормальних точок, тож для їх ізоляції потрібно менше поділів (коротший шлях від кореня до листа), ніж для типових спостережень [13]. Для вибірки без міток класів норма/аномалія

$$X = \{x^i \in R^d | i = 1, \dots, n\} \quad (2.1)$$

де d - це кількість ознак (розмірність кожного вектора спостережень), а n - загальна кількість точок у наборі. Isolation Forest це ансамблевий метод основна ціль якого побудувати ансамбль ізоляційних дерев. Для кожного з t дерев ансамблю незалежно формується власна підвибірка $S \subset X$ розміру ψ простим випадковим відбором, для іншого дерева береться інша, незалежна підвибірка.

Розмір підмножини ψ визначає, наскільки глибоким може стати дерево під час ізоляції спостережень: максимальна глибина оцінюється як $\log_2 \psi$. Зі зменшенням ψ дерева стають неглибокими, отже, шлях для потенційної аномалії коротшає, що підвищує чутливість до дрібних відхилень у даних. Водночас менша підмножина містить менше інформації, тож збільшується випадкова мінливість оцінок - аномальний бал стає менш стабільним від запуску до запуску. Таким чином, вибір ψ – компроміс між здатністю «помічати» слабкі аномалії та надійністю оцінки.

Оскільки аномальні мережеві спостереження, зазвичай, відрізняються від більшості, вони виявляються «ізолюваними» (відокремленими) зазвичай мають статистики, що істотно відхиляються від типової норми, і вже після невеликої кількості випадкових поділів, відповідно, опиняються ближче до кореня дерева.

Натомість спостереження, характерні для «нормальних» потоків, вимагають більшої кількості розгалужень для відокремлення [2, 5].

Кожне з цих дерев будується на власній випадковій підвибірці даних. У кожному вузлі випадково обирається ознака f і порогове значення розбиття τ - скаляр, рівномірно вибраний у межах $[\min f(U), \max f(U)]$ для поточної підвибірки U у вузлі; далі виконується поділ: ліворуч $\{x : f(x) < \tau\}$, праворуч $\{x : f(x) \geq \tau\}$. Оскільки і вибір ознаки, і вибір τ робиться заново в кожному вузлі та для кожного дерева, їхні структури відрізняються [13].

Ключовими параметрами алгоритму є t (кількість дерев), ψ (розмір вибірки S) та максимальна глибина ($\log_2 \psi$), [5].

Часова складність. Побудова одного дерева на підвибірці розміру ψ має складність $O(\psi \log_2 \psi)$, оскільки кожний поділ вимагає $O(1)$, а висота дерева в середньому пропорційна $\log_2 \psi$. [13]

Якщо $\psi \approx n$, складність наближається до $O(t \cdot n \log n)$, де n – загальна кількість спостережень у початковому наборі даних.

На рис. 2.1 позначені основні етапи роботи алгоритму.

Після формування ансамблю з t ізоляційних дерев для кожного спостереження $x \in R^d$ обчислюється середня довжина шляху :

$$\bar{h}(x) = \frac{1}{t} \sum_{j=1}^t h_j(x) \quad (2.2) [5]$$

де $h_j(x)$ - кількість ребер від кореня j -го дерева до листа, що ізолює x .

Малі значення $\bar{h}(x)$ означають, що точка ізолюється швидко – отже, з великою ймовірністю є аномальною, тоді як великі значення $\bar{h}(x)$ вказують на глибшу інтеграцію точки в загальний розподіл даних.

Завдяки такому ансамблевому підходу, коли ми усереднюємо довжини шляхів по t незалежних деревах, випадкові коливання розподілу даних згладжуються, і одержуємо статистично стійку оцінку «ізолюваності» кожного спостереження.

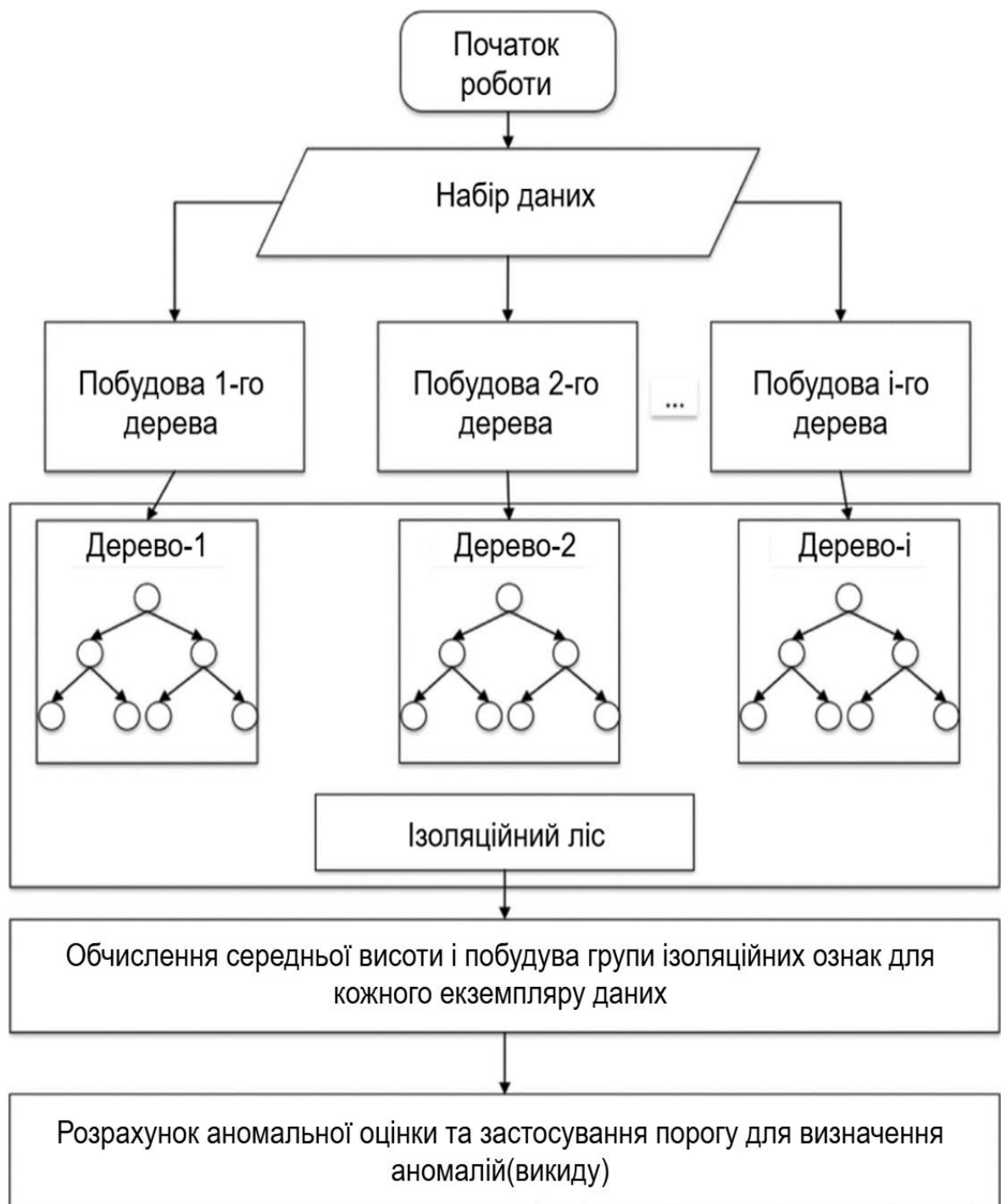


Рис .2.1 основні етапи роботи Isolation Forest

Гранична глибина дерева h_{max} встановлює максимальну кількість поділів. Звичайно беруть $h_{max} = \lceil \log_2 \psi \rceil$, що гарантує повне бінарне розбиття підвибірki. Зменшення h_{max} скорочує час навчання з $O(\psi \log \psi)$ до $O(2^{h_{max}})$, але може призвести до неповної ізоляції глибоко вкладених аномалій [5].

Щоб отримати аномальний бал, $\bar{h}(x)$ нормалізують за допомогою гармонічної поправки $C(\psi)$, що забезпечує порівнянність довжин шляхів на підвбірках різного розміру ψ . Ця поправка походить із теорії випадкових бінарних дерев і пов'язана з гармонічним числом [5]:

$$H_m = 1 + \frac{1}{2} + \dots + \frac{1}{m} \quad (2.3) [5]$$

а саме

$$C(\psi) = 2 * \left[H_{(\psi-1)} - \frac{\psi-1}{\psi} \right] \quad (2.4) [5]$$

Для великих ψ $H_{(\psi-1)}$ прийнято замінити як

$$H_{\psi-1} \approx \ln(\psi - 1) + \gamma \quad (2.5) [5]$$

із сталою Ейлера–Маскероні $\gamma \approx 0,5772156649$.

Формально $C(\psi)$ є математичним очікуванням довжини шляху до самого глибокого листа в незбалансованому бінарному дереві на ψ об'єктах і дозволяє перетворити шляхи в узагальнений аномальний показник.

Далі аномальний бал визначається як

$$s(x) = 2^{-\bar{h}(x)/c(\psi)} \quad (2.6) [5]$$

Завдяки такому перетворенню $s(x) \in (0,1]$: значення, наближені до 1, відповідають спостереженням із короткими шляхами (легко ізольованим, аномальними), тоді як $s(x) \rightarrow 0$ характеризує глибокі в дереві звичайні зразки [5].

Для прийняття рішення про аномалію вводять порогове значення ε та використовують бінарне правило

$$\delta(x; \varepsilon) = \begin{cases} 1, & s(x) \geq \varepsilon \\ 0, & s(x) \leq \varepsilon \end{cases} \quad (2.7) [5]$$

Значення ε зазвичай обирають одним із способів. Один з них - квантильне правило: надати ε на рівні 99-го процентиля розподілу $s(x)$ на навчальній вибірці, тоді саме верхні 1% балів вважають аномаліями.

Оскільки нормальний мережевий трафік зазвичай переважає аномальний на кілька порядків, традиційна точність (Assurasy) не відображає реальної ефективності детектора: модель може демонструвати понад 99 % правильно класифікованих випадків, просто завжди відповідаючи «норма». Саме тому для

оцінки виявлення рідкісних подій і роботи в умовах нерівномірного розподілу варто використовувати метрики, чутливі до балансу класів.

Спершу будується матриця сплутаності, у якій фіксуються чотири величини:

- True Positive (TP) – кількість реальних атак, правильно виявлених системою;
- False Positive (FP) – кількість нормальних спостережень, помилково класифікованих як атаки;
- True Negative (TN) – кількість нормальних спостережень, правильно визнаних безпечними;
- False Negative (FN) – кількість реальних атак, що залишилися непоміченими.

На їх основі обчислюють:

True Positive Rate (Recall), або чутливість, яка показує частку знайдених аномалій серед всіх реальних

$$TPR = \frac{TP}{TP+FN} \quad (2.9) [15]$$

False Positive Rate, або питома частка помилково позначених атак серед нормальних прикладів

$$FPR = \frac{FP}{FP+TN} \quad (2.10) [15]$$

Precision, що відображає частку справді аномальних спрацьовувань серед усіх позитивних прогнозів

$$Precision = \frac{TP}{TP+FP} \quad (2.11) [15]$$

Найбільш узагальненою метрикою в мережевій безпеці є F1-score, яке поєднує Precision і Recall у єдину оцінку якості та обчислюється як гармонійне середнє

$$F1 = \frac{2*Precision*Recall}{Precision+Recall} \quad (2.12) [15]$$

Таким чином, замість загальної Ассурасу ми орієнтуємося на комбінацію TPR, FPR і F1-score, що дає змогу адекватніше оцінити здатність системи виявляти атаки й мінімізувати хибні спрацьовування в умовах дисбалансу класів.

Оскільки аномалій дуже мало, досить важливим стає аналіз поведінки класифікатора при зміні порога між класами. З цією метою будують ROC-криву (залежність TPR від FPR) і обчислюють площу під нею, ROC-AUC.

ROC-крива (Receiver Operating Characteristic) будується як графік залежності показника чутливості (True Positive Rate, TPR) від частки хибних тривог (False Positive Rate, FPR) при зміні порога класифікації. Формально для кожного порога τ обчислюють

$$TPR(\tau) = \frac{TP(\tau)}{TP(\tau) + FN(\tau)} \quad (2.13) [15]$$

$$FPR(\tau) = \frac{FP(\tau)}{FP(\tau) + TN(\tau)} \quad (2.14) [15]$$

і відкладають точку $FPR(\tau), TPR(\tau)$. З'єднуючи ці точки в порядку спадання τ , отримуємо ROC-криву.

ROC-AUC (Area Under the ROC Curve) – це площа під цією кривою, що кількісно характеризує здатність моделі відрізняти позитивні та негативні приклади. У інтегральній формі

$$AUC = \int_0^1 TPR(FPR) d(FPR) \quad (2.15) [15]$$

На практиці для дискретних точок (FPR_i, TPR_i) AUC обчислюють як суму площ трапецій

$$AUC \approx \sum_{i=1}^{n-1} (FPR_{i+1} - FPR_i) * \frac{TPR_{i+1} + TPR_i}{2} \quad (2.16) [15]$$

Проте коли негативний клас явно переважає, навіть відмінна модель може демонструвати високу ROC–AUC за рахунок малого впливу FP на FPR. У таких умовах доцільніше аналізувати Precision–Recall діаграму та інтегральний показник PR-AUC (або Average Precision) , який фокусується лише на позитивному класі й відображає зміну точності (Precision) залежно від повноти (Recall) у всьому діапазоні порогів. Формально це записується як

$$PR - AUC = \int_0^1 Precision(r) dr \quad (2.17) [7]$$

де r пробігає значення Recall від 0 до 1.

Інколи дослідники вдаються до додаткових, більш зведених показників. Так, G-mean

$$G - \text{mean} = \sqrt{TPR * TNR} \quad (2.18) [16]$$

дає зрозумілу «середню» оцінку чутливості та специфічності, особливо популярну у статтях з детекції DDoS-трафіку. Матриця сплутаності дає змогу вирахувати ще один надзвичайно корисний показник для задач із сильним дисбалансом класів – коефіцієнт кореляції Метьюза (Matthews Correlation Coefficient, MCC). Цей індекс фактично є кореляційним коефіцієнтом між двома бінарними векторами – спостереженими мітками та прогнозами моделі - і обчислюється за формулою .

$$MCC = \frac{TP*TN-FP*FN}{\sqrt{(TP+FP)*(TP+FN)*(TN+FP)*(TN+FN)}} \quad (2.19) [16]$$

Практичні рекомендації виглядають так: для порівняння різних моделей доцільно подавати F1 та PR-AUC, де F1 фіксує ефективність у вибраній робочій точці, а PR-AUC характеризує поведінку в усьому діапазоні порогів. Для вибору самого робочого порога зручно спершу обмежити припустимий рівень хибних спрацювань (наприклад, $FPR \leq 0,01$), а потім шукати максимальний Recall за цієї умови. І, нарешті, оскільки дисперсія F1 зменшується лише як $1/\sqrt{t}$. (де t – кількість дерев у ансамблі) [5], у виробничих системах рекомендовано публікувати не лише середні значення, а й довірчі інтервали для основних метрик.

Таким чином, належний вибір і коректне трактування метрик дають змогу об’єктивно оцінити алгоритм виявлення аномалій у середовищі з гострим дисбалансом класів і великою вартістю помилок двох типів.

РОЗДІЛ 3

МОДИФІКАЦІЯ ПРАВИЛА СПЛІТУ В ISOLATION FOREST: МІНІМІЗАЦІЯ ЧАСТКИ МЕНШОЇ ГІЛКИ ТА ЇЇ ВЛАСТИВОСТІ

Базовий Isolation Forest (iForest) ізолює об'єкти послідовними випадковими осьовими розбиттями; аномальні (рідкісні та «віддалені») точки в середньому ізолюються за меншу кількість кроків, що відображається коротшими шляхами у дереві [13].

У запропонованій реалізації збережено ансамблеву конструкцію Isolation Forest, але змінено правило вибору розбиття у вузлі. Нехай у вузол надійшла підвибірка $U \subset X$ розміру N , де $X = \{x^i \in R^d | i = 1, \dots, n\}$ – повний набір даних, x^i – вектори ознак.

У класичному варіанті для вузла випадково обирають ознаку f і випадковий поріг τ , після чого будують ліву та праву гілки за правилом $U_L = \{x \in U : x_f < \tau\}$, $U_R = \{x \in U : x_f \geq \tau \text{ або } x = NaN\}$ [13].

У нашому варіанті в кожному вузлі перебираються всі ознаки $f \in \{1, \dots, d\}$ у випадковому порядку; для кожної ознаки генерується один випадковий поріг τ з інтервалу $[\min f(U), \max f(U)]$ і обчислюється частка меншої гілки

$$r(f, \tau) = \frac{\min(\{u_L, u_R\})}{N} \in (0, 0.5] \quad (3.1)$$

Далі обирається пара $(f, \tau) = \operatorname{argmin}(r(f, \tau))$; тобто фіксується той спліт, який дає найбільш однобоке (дисбалансне) розбиття у поточному вузлі. Такий локальний критерій прямо реалізує ідею Isolation Forest: швидко «відкусити» малу підмножину, якщо вона наявна, і тим самим скоротити очікуваний шлях ізоляції потенційних аномалій.

Обчислювальна складність базового Isolation Forest на підвибірці ψ для одного дерева оцінюється як $O(\psi \log \psi)$: середня висота дерева пропорційна $\log \psi$, а робота на вузол приймається сталою в моделі випадкових розбиттів [13]. У модифікованому варіанті кожен вузол оцінює один випадковий поріг для кожної з d ознак (і обирає найкращий за критерієм $\min(\{u_L, u_R\})$), що додає множник d до вартості вузла. Отже, навчання одного дерева оцінюється як $O(d \psi \log \psi)$, а

ансамблю з t дерев - як $O(t d \psi \log \psi)$. Складність оцінювання одного нового зразка не змінюється та лишається $O(t \log \psi)$, оскільки проходження шляху в кожному дереві має середню довжину порядку $\log \psi$.

Перевага описаного правила вибору спліта полягає в тому, що за наявності невеликої «периферійної» групи спостережень у вузлі існує осьовий поріг, який відокремлює її у відносно малу гілку; систематичний вибір найменшої гілки на кожному кроці зменшує очікувані довжини шляхів для таких спостережень та збільшує віддільність спектра балів $s(x)$ між нормою та аномаліями. На мережевих даних це особливо корисно, коли відхилення проявляються тонкими зсувами у декількох часово-статистичних ознаках, притаманних зашифрованим потокам. Ціна за підвищену роздільну здатність - зростання часу побудови дерев на множник, близький до d , однак ансамблева усередненість та ліміт глибини $\lceil \log_2 \psi \rceil$ зберігають придатність до великих вибірок.

Для оцінювання ефективності запропонованого покращення Isolation Forest (локально «жадібний» спліт) проведено порівняння зі стандартним Isolation Forest із sklearn, а також із K-Nearest Neighbors (K-NN) та Local Outlier Factor (LOF). Усі експерименти виконано на реальних даних зашифрованого мережевого трафіку (датасет CICDDoS2018 із платформи Kaggle)[19], метрики – F1, TPR, FPR. Результати експериментів показано в таблиці 3.1.

Таблиця 3.1 Результати експериментів

Модель / Метрика	F1	TPR	FPR
Isolation Forest (sklearn)	0.555	0.554	0.023
Покращений Isolation Forest	0.611	0.610	0.020

Порівняно зі стандартним Isolation Forest, запропонований варіант підвищує F1 та TPR приблизно на +0.056 абсолютних пунктів ($\approx +10\%$ відносно 0.555/0.554) і одночасно зменшує FPR з 0.023 до 0.020 ($\approx -13\%$ відносно). Це означає кращий компроміс «виявлення $\leftrightarrow$ хибні спрацьовування» при збереженні тієї ж інтерпретації бала $s(x)$ та ансамблевої схеми.

Локально «жадібний» вибір спліта мінімізує частку меншої гілки у вузлі, знаходить невеликі периферійні підмножини, якщо вони присутні. Це скорочує очікувані довжини шляхів для рідкісних спостережень, зміщуючи їхні $s(x)$ ближче до 1 і збільшуючи віддільність від норми. На зашифрованому трафіку, де відхилення виражені тонкими зсувами статистико-часових ознак, такий механізм ізоляції працює краще за чисто випадкові спліти базового iForest.

Зауважимо, що виграш у точності досягається ціною вищої навчальної складності: якщо базовий Isolation Forest навчається за $O(t \psi \log \psi)$, то у запропонованій модифікації – за $O(t d \psi \log \psi)$, оскільки в кожному вузлі перевіряється по одному кандидатному порогу для кожної з d ознак; водночас етап оцінювання зразка зберігає $O(t \log \psi)$, адже середня довжина шляху ізоляції лишається логарифмічною.

Отже Запропоноване покращення Isolation Forest полягає у локально-«жадібному» виборі спліта в кожному вузлі: серед випадково згенерованих порогів для всіх ознак фіксується той, що мінімізує розмір меншої гілки. Це зберігає ідею Isolation Forest (швидка ізоляція рідкісних точок) і, в середньому, скорочує шляхи для периферійних спостережень, підвищуючи віддільність аномальних балів без зміни ансамблевої схеми, нормалізації шляхів та способу обчислення $s(x)$.

ВИСНОВКИ

У роботі обґрунтовано, що повсюдне шифрування мережевого трафіку радикально обмежує ефективність класичних сигнатурних підходів і зміщує акцент на поведінковий аналіз метаданих. Показано, що статистико-часові характеристики потоків (розподіли розмірів пакетів, інтервали, інтенсивність, напрямки з'єднань) достатні для побудови профілів «норми» та виявлення відхилень без доступу до корисного навантаження. Теоретичний розбір Isolation Forest підтвердив його придатність до таких умов: алгоритм не потребує розмітки, стійкий до дисбалансу класів, масштабується майже лінійно та природно відокремлює рідкісні спостереження коротшими шляхами в деревах.

На основі цієї теорії запропоновано та реалізовано покращення Isolation Forest: у кожному вузлі з усіх ознак вибирається той випадковий поріг, що мінімізує частку меншої гілки. Такий локально «жадібний» вибір підсилює індуктивну ідею Isolation Forest – швидку ізоляцію малих підмножин – і практично підвищує віддільність аномальних балів у порівнянні з повністю випадковими сплітами. При цьому збережено ключові елементи: ансамблеву схему, обмеження глибини, гармонічну нормалізацію довжин шляхів і формулу аномального балу. Аналітично встановлено, що ціна за точнішу ізоляцію – збільшення часу навчання від $O(t \psi \log \psi)$ до $O(t d \psi \log \psi)$.

Практична цінність полягає в тому, що модифікований iForest краще відокремлює тонкі зсуви у метаданих зашифрованих потоків, що критично для виявлення невідомих або контекстуальних загроз. Модифікований iForest покращує роздільну здатність між нормою й рідкісними відхиленнями у зашифрованих потоках, зберігаючи простоту реалізації та майже лінійну масштабованість. Це робить підхід придатним для вбудування в існуючі NDR/IDS конвеєри як легкий поведінковий детектор поверх flow-метрик.

Наукова новизна роботи полягає в локальному підсиленні ізоляції через критерій «мінімальної меншої гілки» у вузлах і в процедурі керованого вибору порога рішення, що разом дає приріст точності без змін у статистичній інтерпретації алгоритму. Обмеження пов'язані з додатковою вартістю навчання

та потенційною чутливістю до вибору ψ і глибини; однак вони компенсуються масштабованістю ансамблю та стабільністю оцінювання. Отже, запропонований підхід є збалансованим компромісом «швидкість $\rightarrow$ точність» і відповідає практичним вимогам моніторингу зашифрованого трафіку: забезпечує вищу якість виявлення при збереженні придатності до великих мереже.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Могильний, С. Б. Інформаційна безпека: лабораторний практикум [Електронний ресурс]. Електронні текстові дані (1 файл: 5,1 Мбайт). Київ: КПІ ім. Ігоря Сікорського, 2023, 145с. Режим доступу: <https://ela.kpi.ua/server/api/core/bitstreams/69fcd5b9-92ce-4468-89de-c3b9016c11d2/content>
2. A Survey of Methods for Encrypted Traffic Classification and Analysis. [Електронний ресурс] / P. Velan, C. Milan, P. Celeda, M. Drasar // Wiley InterScience. 2014. Режим доступу: <https://is.muni.cz/publication/1305020/a-survey-of-methods-for-encrypted-traffic-classification-and-analysis.pdf>.
3. Подвисоцька, О. П. Застосування алгоритмів машинного навчання для виявлення аномалій мережного трафіку / О. П. Подвисоцька, С. О. Носок // Теоретичні і прикладні проблеми фізики, математики та інформатики. Київ, 2024. С. 288–290. Режим доступу: <https://ela.kpi.ua/server/api/core/bitstreams/53d7170e-df2d-423d-b815-805b9a9b7dfe/content>.
4. A Cluster-then-label Semi-supervised Learning Approach for Pathology Image Classification. [Електронний ресурс] / M. Peikari, S. Salama, S. Nofech-Moze, A. Martel // Scientific Reports. 2018. Режим доступу: <https://www.nature.com/articles/s41598-018-24876-0>.
5. Cheong Hee Park. Comparative study for outlier-detection methods in high-dimensional text data // Journal of Artificial Intelligence and Soft Computing Research. 2023. Vol. 13, No. 4. p. 235–247. DOI: 10.32674/jaiscr.v13i4.5678.
6. Isolation Forest. // Proceedings of the 2008 8th IEEE International Conference on Data Mining (ICDM), Pisa (Italy), 15-19 Dec. 2008. Piscataway, NJ: IEEE, 2008. p. 413–422. DOI: 10.1109/ICDM.2008.17.
7. Жилін, А. В. Технології захисту інформації в інформаційно-телекомунікаційних системах [Електронний ресурс]: навч. пос. Електронні текстові дані (1 файл: 5,23 Мбайт). Київ: КПІ ім. Ігоря Сікорського, 2021, 213 с. Режим доступу: <https://ela.kpi.ua/server/api/core/bitstreams/d99a0045-e907-4d17-afc1-5431f67d2444/content>.

8. Network Traffic Classification techniques and comparative analysis using Machine Learning algorithms. [Електронний ресурс] / [M. Shafiq, X. Yu, L. Asif Ali Laghari та ін.] // IEEE. 2016. Режим доступу до ресурсу: <https://ieeexplore.ieee.org/document/7925139>.
9. Practical evaluation of encrypted traffic classification based on a combined method of entropy estimation and neural networks. [Електронний ресурс] / K. Zhou, W. Wang, C. Wu, T. Hu // ETRI Journal. 2020. Режим доступу: <https://onlinelibrary.wiley.com/doi/full/10.4218/etrij.2019-0190>.
10. Остапов, С. Е. Кібербезпека: сучасні технології захисту: навч. пос. Львів: «Новий Світ-2000», 2024, 678 с. Режим доступу: https://opac.kpi.ua/F/?func=direct&doc_number=000645277&local_base=KPI01.
11. Lightweight, Payload-Based Traffic Classification: An Experimental Evaluation. [Електронний ресурс] / [F. Risso, M. Baldi, O. Morandi та ін.] // IEEE. 2008. Режим доступу до ресурсу: <https://ieeexplore.ieee.org/document/4534133>.
12. Khalife J. A multilevel taxonomy and requirements for an optimal traffic-classification model [Електронний ресурс] / J. Khalife, A. Hajjar, J. Diaz-Verdejo // Wiley Online Library. 2014. Режим доступу: <https://onlinelibrary.wiley.com/doi/abs/10.1002/nem.1855>
13. Rezaei S.; Liu X. Deep Learning for Encrypted Traffic Classification: An Overview [Електронний ресурс] // IEEE Communications Magazine. 2019, Vol. 57, No. 5, p. 76–81. Режим доступу: <https://arxiv.org/abs/1810.07906>.
14. Powers D. M. W. Evaluation: From Precision, Recall and F-Measure to ROC, Informedness and Correlation // Journal of Machine Learning Technologies. 2011, Vol. 2, No. 1, p. 37-63.
15. Airlangga G. Advanced machine learning techniques for seismic anomaly detection in Indonesia: a comparative study of LOF, Isolation Forest, and One-Class SVM [Електронний ресурс] / G. Airlangga // Proceedings of Geo-AI 2024, Jakarta (Indonesia), 5-7 Apr. 2024. Jakarta: Geo-AI Press, 2024, p. 112-118.
16. El Hairach M. L.; Bellamine I.; Tmiri A. Anomaly Detection in PV Modules: A Comparative Study of DBSCAN, k-means, Isolation Forest and LOF // Proceedings of

the 2023 7th IEEE Congress on Information Science and Technology (CiSt), Agadir–Essaouira (Morocco), 16-22 Dec. 2023. IEEE, 2023. DOI: 10.1109/CiSt56084.2023.10409931

17. Saito T.; Rehmsmeier M. The Precision–Recall Plot is More Informative than the ROC Plot in Imbalanced Datasets // PLOS ONE. 2015. Vol. 10, No. 3, e0118432.

18. Xu L.; Li Y. Hybrid iForest-DBSCAN for SCADA anomaly detection // Knowledge-Based Systems. 2025, Vol. 278, Article 109452. DOI: 10.1016/j.knosys.2025.109452

19. IDS 2018 Intrusion CSVs (CSE-CIC-IDS2018). *Kaggle*. URL: <https://www.kaggle.com/datasets/solarmainframe/ids-intrusion-csv?select=02-20-2018.csv> (дата звертання: 10.09.2025).

Додаток 1

МАТЕРІАЛИ

... науково-практичної конференції курсантів (студентів), аспірантів,
докторантів та молодих учених

“ ...

...”

... 2025 року

Київ – 2025

Секція № ...“ ...
...”

...

...

...

...

...

...

...

...

...

...

...

...

...

...

...

...

...

..

Forest Watchdog

...

...

...

...

...

...

...

...

...

...

...

...

...

Анотація. В роботі розглянуто проблему аналізу та виявлення мережових атак у зашифрованому мережевому трафіку з використанням методів керованого машинного навчання та методів глибокого навчання.

Summary. The thesis considers the problem of analyzing and detecting network attacks in encrypted network traffic using supervised machine learning and deep learning methods.

Ключові слова: кібербезпека, шифрування, мережевий трафік, машинне навчання, виявлення аномалій.

Сучасні протоколи TLS 1.3, IPsec і більшість VPN-рішень шифрують не лише корисне навантаження, а й значну частину службових заголовків, фактично усуваючи можливість сигнатурного або порт-орієнтованого аналізу без розшифрування трафіку.

Зловмисники активно експлуатують цю можливість, приховуючи шкідливі навантаження й DDoS-флуди всередині легітимних TLS-сесій, що значно ускладнює їх ідентифікацію.

В сучасних мережах фіксують зростання атак, спрямованих саме на зашифровані канали, через це є необхідність переходу від контент-орієнтованої перевірки до обробки статистико-часових ознак трафіку й поведінкових патернів.

Наразі популярними методами для аналізу шифрованого трафіку є методи машинного навчання а саме некеровані методи (unsupervised) та керовані методи (supervised) Основна ідеї використання вище названих методів це в пошуку аномалій трафіку.

Методи керованого машинного навчання потребують використання заздалегідь розмічені вибірки шифрованого трафіку та витягнуті з нього статистико-часові ознаки наприклад, розмір пакетів, інтервал між пакетами, кількість повідомлень у TLS-рукописанні, щоб навчити алгоритми відрізняти нормальну поведінку від аномальної та виявляти атаки в зашифрованих каналах.

Прикладом таких алгоритмів може слугувати алгоритм K найближчих сусідів (K-Nearest Neighbors). Алгоритм K найближчих сусідів класифікує новий пакет, порівнюючи його статистико-часові ознаки з K найближчими пакетами у розміченій вибірці, і присвоює клас за принципом більшості. Такий алгоритм не робить жорстких припущень про розподіл даних, підтримує довільні метрики, легко пояснюється й адаптується до задач на кшталт класифікації мережевого трафіку, але потребує коректного вибору K, масштабування ознак.

Методи некерованого машинного навчання не потребують заздалегідь розмічених даних, як в випадку з методами з наглядом. Узагальнена схема роботи методів некерованого виявлення аномалій включає послідовні етапи: збирання та попереднє очищення даних, нормалізацію ознак, навчання моделі на

нормальних спостереженнях для отримання outlier-score (рахунок викидів).

Прикладом методів некерованого машинного навчання є алгоритм ізоляційного лісу (Isolation Forest). Алгоритм ізолюваного лісу працює на основі простого принципу: цей алгоритм будує аномальні дерева прийняття рішень, розподіляючи дані на різні підгрупи доти. Оскільки аномалії зазвичай потребують меншої кількості розподілень для їх виділення, вони будуть ближче до кореня дерева порівняно з нормальними об'єктами даних.

Висновки. Методи машинного навчання нині стали провідним інструментом для виявлення атак у зашифрованому трафіку. Ці методи аналізують статистичні й часові ознаки пакетів без розшифрування даних, ефективно виявляючи загрози.