



DOI 10.28925/2663-4023.2025.29.915

УДК 04.056.5:004.738.5:004.451.3

Костюк Юлія Володимирівна

PhD in Computer Science

доцент кафедри інформаційної та кібернетичної безпеки імені професора Володимира Бурячка

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID: 0000-0001-5423-0985

y.kostiuk@kubg.edu.ua

Складанний Павло Миколайович

кандидат технічних наук, доцент

завідувач кафедри інформаційної та кібернетичної безпеки імені професора Володимира Бурячка

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID: 0000-0002-7775-6039

p.skladannyi@kubg.edu.ua

Рзаєва Світлана Леонідівна

кандидат технічних наук, доцент,

доцент кафедри комп'ютерних наук

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID: 0000-0002-7589-2045

s.rzaieva@kubg.edu.ua

Мазур Наталія Петрівна

кандидат педагогічних наук, доцент,

доцент кафедри інформаційної та кібернетичної безпеки імені професора Володимира Бурячка

Київський столичний університет імені Бориса Грінченка, Київ, Україна

ORCID: 0000-0001-7671-8287

n.mazur@kubg.edu.ua

Черевик В'ячеслав Михайлович

кандидат технічних наук, доцент,

доцент кафедри інформаційної та кібернетичної безпеки імені професора Володимира Бурячка

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID: 0000-0002-2735-5341

v.cherevyk@kubg.edu.ua

Аносов Андрій Олександрович

кандидат військових наук, доцент,

доцент кафедри інформаційної та кібернетичної безпеки імені професора Володимира Бурячка

Київський столичний університет імені Бориса Грінченка, м. Київ, Україна

ORCID ID: 0000-0002-2973-6033

a.anosov@kubg.edu.ua

ОСОБЛИВОСТІ РЕАЛІЗАЦІЇ МЕРЕЖЕВИХ АТАК ЧЕРЕЗ TCP/IP-ПРОТОКОЛИ

Анотація. У статті здійснено дослідження особливостей реалізації типових мережових атак, що ґрунтуються на експлуатації уразливостей протоколів стеку TCP/IP, які є критичною інфраструктурною основою глобальної мережової взаємодії. Проведено поглиблений аналіз архітектурних обмежень та функціонально-протокольних характеристик ключових компонентів мережового стеку (ARP, IP, ICMP, TCP, UDP, DNS), які в сучасних умовах виступають як основні вектори ініціації кіберзагроз. На основі референтної моделі OSI здійснено формалізовану класифікацію атак за рівнями взаємодії, із виокремленням характерних сценаріїв: IP spoofing, ARP poisoning, TCP session hijacking, DNS cache poisoning, UDP flooding та ICMP-based covert



channels. Встановлено типові механізми обходу традиційних засобів захисту, що включають маніпуляції з маршрутами передавання, підміну службових повідомлень та інкапсуляцію шкідливих пакетів у легітимний трафік. Окрему увагу приділено огляду інструментальних засобів і методів проактивного виявлення та нейтралізації загроз, зокрема систем виявлення вторгнень, міжмережевих екранів, технологій глибокого інспектування пакетів, а також методів поведінкового та ентропійного аналізу аномалій у мережевих потоках. Отримані результати формують як теоретичну базу для моделювання атак та оцінки ризиків, так і прикладну основу для підвищення ефективності інформаційної безпеки в гетерогенних мережевих середовищах.

Ключові слова: TCP/IP; мережеві атаки; IP spoofing; ARP poisoning; кіберзагрози; вектори атак; інформаційна безпека; маршрутизація; захист мереж; CVE; SIEM; Explainable AI; поведінковий аналіз; Zero Trust.

ВСТУП

У контексті стрімкої цифрової трансформації інформаційно-комунікаційних систем протоколи TCP/IP залишаються базисом мережевої взаємодії, забезпечуючи обмін даними між вузлами в глобальних та локальних мережах. Проте концептуальна відсутність вбудованих механізмів безпеки, закладена ще на етапі розробки цих протоколів, створює критичне середовище для реалізації широкого спектру атак [1], [5]. Зокрема, моделі взаємодії «end-to-end», маршрутизація «hop-by-hop», відкритість ARP-, ICMP-, DNS-повідомлень — усе це формує низькорівневі точки входу для кіберзлочинців [3], [4], [6], [7], [11]. Ці архітектурні особливості зумовлюють підвищену уразливість до атак типу spoofing, hijacking та tunneling, які реалізуються без необхідності доступу до прикладного рівня [2], [10], [14]. Відсутність автентифікації на рівні протоколів ARP, ICMP та UDP дозволяє зловмисникам ініціювати несанкціоновану маршрутизацію, зміну кешу або побудову прихованих каналів керування (C2) [6]. Крім того, слабка ентропія у генерації параметрів сесій (наприклад, TCP ISN) відкриває можливості для прогнозованого захоплення з'єднання [5], [6], [10], [14]. Таким чином, TCP/IP-протоколи в сучасних умовах функціонують як вектор ініціації складних, багаторівневих атак з мінімальними слідами втручання, що ускладнює їх виявлення класичними механізмами захисту.

Особливості реалізації мережевих атак через TCP/IP-протоколи полягають у використанні уразливостей архітектурного рівня з обходом традиційних механізмів захисту. Наприклад, перехоплення трафіку (sniffing), підміна ARP- або IP-адрес (spoofing), використання ICMP-тунелювання, вплив на кешування DNS, підміна TCP-з'єднання, а також атаки типу DoS/DDoS реалізуються без необхідності складної компрометації вузлів — достатньо доступу до транспортного або мережевого рівня моделі OSI.

Наразі спостерігається активна експлуатація нових уразливостей, зареєстрованих у базі CVE (Common Vulnerabilities and Exposures), що дозволяють реалізовувати атаки через TCP/IP-протоколи із високим рівнем критичності [7], [11]. Зокрема: CVE-2023-28771 — уразливість у сервісі Zyxel ZyWALL, яка дозволяє зловмиснику віддалено виконати команди через спеціально сформований пакет IKE, використовуючи UDP (порушення протоколу IPsec) [CVSS v3: 9.8], CVE-2023-29552 — уразливість у Microsoft Message Queuing (MSMQ), що дозволяє віддалену атаку через відкритий TCP-порт 1801, використовуючи підроблений трафік для DoS-атаки, CVE-2023-23397 — критична уразливість у Microsoft Outlook, яка дозволяє ініціювати викрадення NTLM-хешів через спеціально сформовані TCP-з'єднання з використанням UNC-шляху, реалізованого через SMB, CVE-2023-3595 — підміна маршруту через маніпуляцію DNS-запитами в системах MikroTik з можливістю впровадження «man-in-the-middle» атаки на рівні TCP-сеансу.



Актуальність дослідження зумовлена зростаючою кількістю атак типу Lateral Movement, Remote Code Execution (RCE) та Session Hijacking, що реалізуються через низькорівневі механізми TCP/IP та залишаються поза зоною виявлення класичних міжмережевих екранів або антивірусних систем [8], [14], [16], [20], [24]. Використання протоколів транспортного та мережевого рівнів для побудови прихованих або тунельованих каналів зв'язку значно ускладнює ідентифікацію таких інцидентів навіть у системах моніторингу безпеки (SIEM) та в умовах Zero Trust-архітектури [27]–[32].

Таким чином, наукова проблема полягає у необхідності комплексного аналізу уразливостей TCP/IP-протоколів, формалізації методів реалізації мережеских атак та визначенні механізмів проактивного виявлення, запобігання та мінімізації ризиків [5], [15], [17], [19]. Метою даної роботи є дослідження ключових сценаріїв атак через TCP/IP-протоколи з урахуванням сучасних загроз, а також узагальнення механізмів захисту на різних рівнях мережевої моделі.

Постановка проблеми. Протоколи TCP/IP, які лежать в основі глобальної мережевої взаємодії, історично не передбачали вбудованих механізмів безпеки, оскільки їх архітектура була орієнтована на відкритість і доступність [10]. Це створює сприятливе середовище для реалізації широкого спектра атак, що експлуатують низькорівневі механізми маршрутизації, адресації, встановлення з'єднання та обміну службовими повідомленнями [1], [3], [4]. До найбільш поширених належать IP spoofing, ARP poisoning, TCP session hijacking, ICMP tunneling, DNS spoofing і DoS-атаки, які можуть здійснюватися без автентифікації користувача або доступу до прикладного рівня [6], [9]. Особливістю таких атак є можливість їх реалізації через модифікацію стандартного трафіку або створення прихованих каналів передачі даних, що суттєво ускладнює їх виявлення класичними засобами захисту [21], [22], такими як міжмережеві екрани чи сигнатурні IDS.

Актуальність проблеми підсилюється постійною появою нових CVE-уразливостей, які демонструють практичну реалізацію атак саме через TCP/IP-протоколи. Наприклад, CVE-2023-28771 дозволяє виконання команд через UDP-пакети в IPsec [6], CVE-2023-29552 — організацію DoS-атаки через TCP-порт 1801 [10], CVE-2023-23397 — викрадення NTLM-хешів через UNC-шлях у TCP [11], а CVE-2023-3595 — маніпуляцію DNS-відповідями з реалізацією MITM-атаки на рівні TCP-сеансу [4]. Ці приклади підтверджують системний характер загроз, що виникають не через недосконалість прикладного ПЗ, а саме через слабкості протоколів мережевого стеку.

Такі атаки виявляють не лише програмні недоліки, а й глибоко вкорінені архітектурні уразливості, що ускладнюють їх виявлення засобами класичних систем захисту. Більше того, через фундаментальну роль протоколів TCP/IP у забезпеченні комунікації, навіть мінімальні уразливості можуть мати каскадний ефект на безпеку всього інформаційного середовища. Відомі уразливості експлуатують нетипову поведінку у фазах встановлення, підтримки або завершення з'єднань, зокрема через нестандартне використання заголовків пакетів або маніпуляцію таблицями маршрутизації [11], [12], [16], [23]. Це вимагає впровадження комплексних моделей оцінки ризиків, орієнтованих на глибокий аналіз трафіку, протокольних відхилень та поведінкових шаблонів мережевої взаємодії.

В умовах широкого використання Zero Trust-архітектур, сегментованих мереж і SIEM-систем, критично важливим є формування підходів до виявлення й запобігання таким атакам із урахуванням контексту функціонування протоколів TCP/IP [14], [20], [24]. Таким чином, проблема полягає в необхідності глибокого аналізу архітектурних та поведінкових особливостей реалізації мережеских атак через TCP/IP-протоколи, з подальшою розробкою ефективних методів захисту на рівні мережевої інфраструктури [16], [20]. Це передбачає побудову багаторівневих моделей ризику, що поєднують аналіз конфігураційної



відкритості, поведінкових аномалій та телеметричних параметрів трафіку [13], [16], [24]. Особливу увагу необхідно приділяти формалізації ознак, характерних для атак із мінімальними сигнатурними проявами, які обходять традиційні засоби фільтрації [25], [26]. У цьому контексті актуальними є методи глибокого інспектування пакетів (DPI), ентропійного аналізу та Explainable AI, що забезпечують адаптивне реагування на загрози [17], [18], [24]. Інтеграція таких підходів у SIEM/SoC-середовища дозволяє здійснювати контекстно-орієнтовану ідентифікацію вторгнень у режимі реального часу.

Дослідження базується на побудові віртуального середовища в GNS3 з VLAN-сегментацією, використанням Kali Linux, Ubuntu Server та Open vSwitch для моделювання типових топологій корпоративних мереж. Атаки (ARP spoofing, ICMP tunneling, TCP hijacking, DNS cache poisoning) реалізовано за допомогою Scapy, Ettercap, Hping3 та DNSChuf, а трафік аналізувався через Wireshark, Zeek і Suricata [9], [25]. Для верифікації реалістичних сценаріїв залучено CVE-уразливості (CVE-2023-28771, CVE-2023-23397, CVE-2023-3595) [6], [11]. Застосовано інтеграцію з SIEM (Splunk) та поведінкове моделювання для формування формалізованих метрик ризику ($D_{risk}, K_{ARP}, D_{ICMP}, R_{CVE}, S_{total}$) [16], [24]. Для інтерпретації результатів впроваджено Explainable AI з використанням SHAP-аналізу, що дозволило визначити ключові ознаки аномального трафіку [16]–[18], [20]. Такий підхід забезпечив відтворюваність експериментів і дозволив кількісно оцінити ризики атак через TCP/IP-протоколи в умовах сучасної мережевої інфраструктури.

Аналіз останніх досліджень і публікацій. Упродовж останніх років зростає інтерес до досліджень, що стосуються особливостей реалізації мережевих атак на основі протоколів TCP/IP, зокрема з акцентом на їхню архітектурну незахищеність. У роботі Hidouri et al. [1] розглянуто загальні проблеми безпеки та атаки в інформаційно-орієнтованих мережах, серед яких згадуються IP spoofing, SYN flooding та інші механізми, характерні для TCP/IP-середовищ, що вказує на незахищеність транспортного рівня від підробки адрес та перебору послідовностей з'єднання.

У дослідженні Pittman et al. [2] запропоновано таксономію ефективності динамічних honeypot-систем, які імітують сервіси TCP/IP з метою виявлення атак, зокрема ICMP- і DNS-маніпуляцій, що дозволяють обійти класичні фаєрволи та системи виявлення вторгнень. Shah і Cosgrove [3] у своєму огляді зосереджуються на атаках ARP cache poisoning у програмно-детермінованих мережах (SDN), що є типовим прикладом атаки на рівні мережевого доступу TCP/IP-стеку з використанням підміни MAC-адрес.

Проблематику DNS cache poisoning розкрито у роботі Jin et al. [4], де запропоновано метод машинного навчання для виявлення фальсифікованих DNS-відповідей у реальному часі. Yang [5] досліджує аномалії мережевого трафіку, що виникають внаслідок атак на мережеві протоколи, використовуючи ентропійні моделі для виявлення прихованих каналів на основі ICMP та UDP, що є типовими прикладами тунелювання через TCP/IP.

Суттєвий внесок у вивчення атак через IP spoofing зроблено в роботі Dai та Shulman [6], які презентували інструмент SMap для глобального сканування мереж, здатних приймати пакети з підробленими IP-адресами. Результати показали високий рівень незахищеності інфраструктури Інтернету до спуфінгових атак. Аналогічна тематика розглянута Nosyk et al. [7], де у проєкті Closed Resolver досліджено масштаби відсутності inbound Source Address Validation (iSAV), що критично підвищує ризики реалізації TCP/IP-атак із боку зовнішніх хостів.

Хоча дослідження Tang et al. [8] зосереджене переважно на методах виявлення атак у SDN-середовищах за допомогою глибинного навчання, воно також охоплює низку сценаріїв атак через мережевий стек, зокрема перехоплення сесій, flooding і маніпуляції з TCP-з'єднаннями, які є характерними для традиційних мереж на базі TCP/IP.



Загалом проаналізовані дослідження доводять, що реалізація атак через TCP/IP-протоколи є можливою завдяки фундаментальним архітектурним обмеженням цих протоколів. Водночас, попри наявність окремих досліджень з ARP, DNS, TCP, ICMP-атак, у літературі бракує комплексної узагальненої моделі, яка б одночасно охоплювала повний спектр протокольних атак, механізмів їх ініціації, обходу та виявлення в умовах сучасної кіберінфраструктури.

Мета статті. Метою статті є комплексне дослідження особливостей реалізації мережових атак через TCP/IP-протоколи з акцентом на архітектурні уразливості транспортного, мережового та каналного рівнів. З огляду на те, що протоколи TCP/IP спочатку не були спроектовані з урахуванням вимог до інформаційної безпеки, багато з них залишаються вразливими до атак типу IP spoofing, ARP poisoning, TCP session hijacking, ICMP tunneling, DNS cache poisoning, UDP flooding та SYN flooding. Основна мета полягає у формалізації типових сценаріїв атак, які експлуатують логіку функціонування відповідних протоколів, зокрема відсутність автентифікації, можливість підробки адрес, маніпуляцій із пакетами та створення прихованих каналів зв'язку.

У роботі систематизовано ключові вектори атак, що реалізуються без необхідності порушення прикладного рівня або компрометації вузлів, а лише за рахунок маніпуляцій на рівні TCP/IP. Особливу увагу приділено вивченню механізмів обходу стандартних засобів захисту, таких як фаєрволи, NAT, IDS/IPS-системи, а також виявленню слабкостей у фільтрації вхідного та вихідного трафіку. У межах досягнення мети проаналізовано сучасні інструменти аналізу трафіку та виявлення атак, включаючи honeypot-системи, методи ентропійного аналізу, глибоку інспекцію пакетів (DPI) і технології машинного навчання.

РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ

У процесі дослідження було здійснено багаторівневий аналіз особливостей реалізації мережових атак, які експлуатують архітектурні та протокольні уразливості стеку TCP/IP [3], [5], [10]. У межах цього дослідження було реалізовано цілісний та взаємопов'язаний аналітичний підхід, що базувався на поєднанні структурного, поведінкового, емпіричного та аналізу уразливостей із урахуванням актуальної класифікації кіберзагроз, моделі OSI та найчастіше фіксованих векторів атак, згідно з даними загальнодоступних репозитаріїв (зокрема CVE, NIST NVD, CISA KEV) [6], [7]. Такий підхід дозволив не лише ідентифікувати типові архітектурні вади та логічні уразливості стеку TCP/IP, але й відтворити реалістичні сценарії атак для оцінки ефективності сучасних засобів виявлення та протидії [9], [10], [24], [25]. Аналіз було реалізовано у п'ять послідовних етапів, кожен із яких мав визначену методологічну мету та ґрунтувався на міждисциплінарному поєднанні інструментальних і теоретичних підходів.

На першому етапі — архітектурному — досліджувалися базові протоколи TCP/IP-стеку (IP, TCP, UDP, ARP, ICMP, DNS) з фокусом на наявність або відсутність вбудованих механізмів автентифікації, перевірки цілісності, контролю сесії та підтримки фільтрації. Було сформовано структурно-протокольну карту взаємодії між компонентами мережової інфраструктури (шлюзи, маршрутизатори, кінцеві вузли, DNS/ARP-сервери), у якій позначено логічні зони довіри та потенційні зони перетину, що створюють умови для ініціації атак через відкриті порти, неконтрольований ARP-кеш, дозволений ICMP або несанкціоновану передачу DNS-запитів [1], [3], [4], [10], [21], [22]. Такий підхід дозволив ідентифікувати критичні точки перетину між сегментами мережі, де порушується принцип ізоляції, що є передумовою для горизонтального переміщення загроз (lateral movement) [6], [11]. Виявлено, що у більшості конфігурацій мережового середовища відсутні чіткі політики поділу довіри між VLAN-сегментами та слабо реалізовані механізми контролю міжзонної

взаємодії [12], [13]. Проаналізовано кореляцію між архітектурною відкритістю та ймовірністю ініціації атак на основі формалізованого показника щільності ризикових перетинів (D_{risk}) [16], [20], [23]. Отримані результати лягли в основу побудови структурно-протокольної моделі загроз, що враховує просторову локалізацію вразливих вузлів та сценарії несанкціонованого доступу на рівні TCP/IP.

На рис. 1 зображено структурну UML-схему типової TCP/IP-архітектури, у якій реалізовано три функціональні зони: зону доступу користувача (VLAN 10), демілітаризовану зону (DMZ) та інфраструктуру сервісів (VLAN 20). Компоненти обмінюються трафіком через маршрутизатор із певними правилами доступу. Ризикові зони перетину (R1, R2) вказують на критичні точки, де порушується принцип ізоляції, що створює умови для реалізації атак типу ARP Spoofing, ICMP Tunneling, DNS Poisoning, а також горизонтального переміщення (Lateral Movement). Таке представлення дозволяє формалізувати взаємозв'язки між сегментами з різними рівнями довіри та ідентифікувати зони з підвищеною ймовірністю компрометації. Зокрема, відсутність міжзонного контролю на рівні ACL або NAT створює можливості для неконтрольованої маршрутизації трафіку та ініціації атак з обхідними маршрутами. В межах моделі були виділені структурні вектори ризику, що враховують топологічну близькість вузлів, дозволені протоколи взаємодії та ступінь сегментації мережі. Таким чином, схема виконує функцію просторово-протокольної абстракції, яка є основою для побудови подальших графових моделей загроз та автоматизованого аналізу політик безпеки.

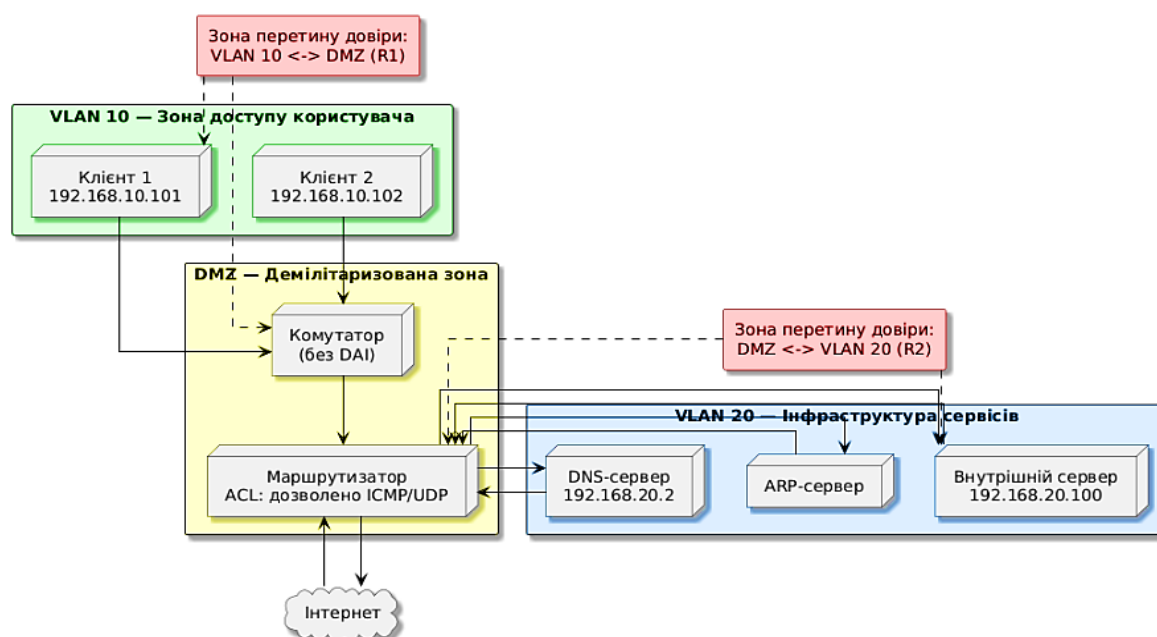


Рис. 1. Архітектура TCP/IP та зони довіри в мережевій інфраструктурі

Для переходу від якісного опису до кількісного аналізу загроз, що реалізуються через стек TCP/IP, було здійснено формалізацію ключових характеристик атак із використанням моделей ймовірності, поведінкових метрик, інформаційної ентропії та ефективності реагування [8], [16], [24]. Нижче подано перелік моделей, які інтегруються у процес оцінювання ризиків, виявлення аномалій і формування політик реагування [5], [20]. Щоб формалізувати ступінь ризику, зумовленого перетином зон довіри в архітектурі, використано коефіцієнт щільності ризикових перетинів:

$$D_{risk} = \frac{N_{overlap}}{N_{total}}, \quad (1)$$

де $N_{overlap}$ — кількість зон перетину довіри між VLAN/ACL, N_{total} — загальна кількість мережевих сегментів. Формула (1) дозволяє кількісно оцінити рівень потенційної уразливості архітектури мережі на основі конфігураційної відкритості сегментів, що не мають чітко визначених політик поділу довіри [16], [20], [24]. Високе значення $D_{risk} = 1$ свідчить про наявність великої кількості зон, де взаємодіють вузли з різним рівнем доступу, без чіткої сегментації чи обмежень доступу між ними. Це створює сприятливі умови для горизонтального переміщення атакуючого в мережі (lateral movement), а також для реалізації атак через такі протоколи, як: ARP (Address Resolution Protocol) — при неконтрольованому кешуванні записів можливе виконання ARP Spoofing або ARP Poisoning, ICMP (Internet Control Message Protocol) — дозволяє зловмиснику ініціювати ICMP Tunneling у дозволених міжзонних напрямках, UDP/TCP — у випадку відкритих портів між сегментами ризику, можлива реалізація атак типу UDP Reflection, TCP Hijacking або DDoS. Загалом, таке значення дозволяє оцінити рівень відкритості мережевої архітектури до атак, які використовують відсутність ізоляції, зокрема ARP Spoofing, ICMP Tunneling або TCP Session Hijacking. При значенні $D_{risk} > 0.5$ можна стверджувати про підвищену ймовірність горизонтального переміщення зловмисника між зонами без додаткової аутентифікації. Формула є базовою для подальшого аналізу ризиків і використовується при побудові карти вразливих перетинів у TCP/IP-мережі.

На рис. 2 представлено граф зон перетину довіри TCP/IP-мережі, у якому виділено критичні ділянки взаємодії між сегментами: внутрішні вузли (VLAN 10), демілітаризована зона (DMZ), зона сервісів (VLAN 20) та зовнішня мережа (Інтернет).

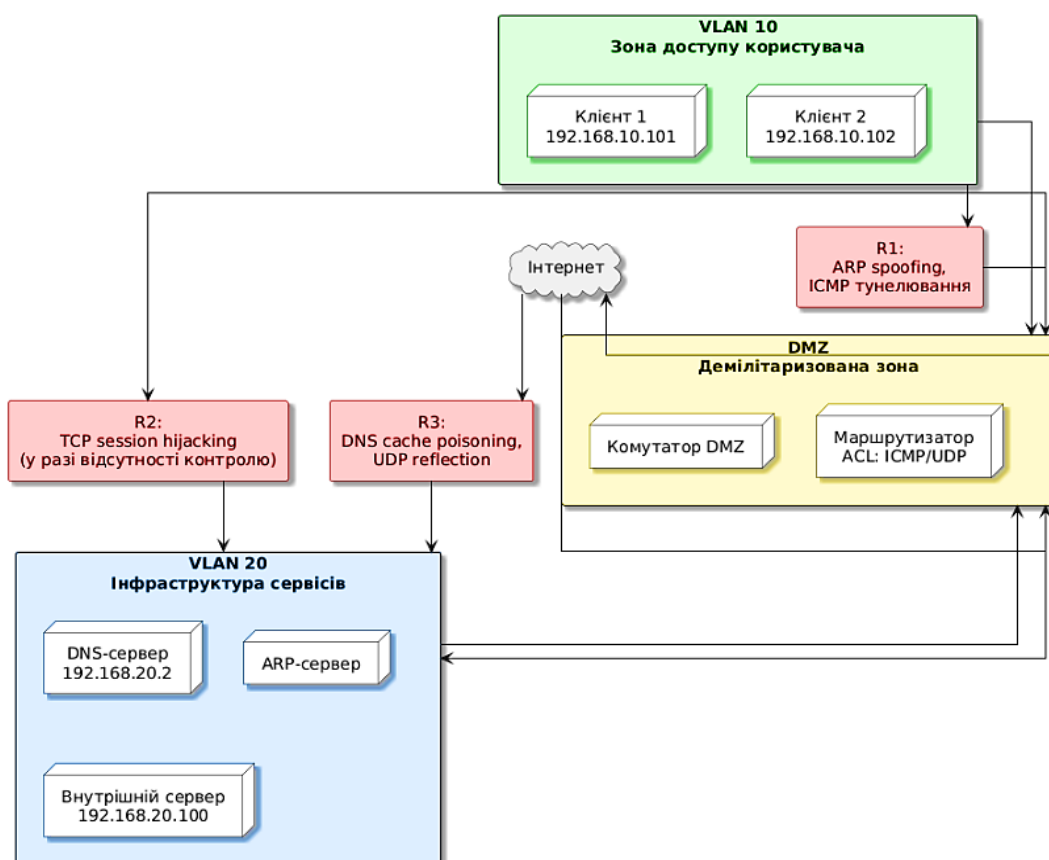


Рис. 2. Граф зон перетину довіри TCP/IP-мережі

Кожна зона перетину позначена як R1, R2, R3, і позначає потенційні вектори атак, зокрема: R1 — реалізація ARP spoofing, ICMP тунелювання, R2 — можливість TCP session hijacking у разі відсутності міжзонного контролю, R3 — проникнення з Інтернету через DNS cache poisoning або UDP reflection. Цей граф ілюструє структурно-протокольну модель ризику, що лежить в основі формули (1) та подальшої побудови карти вразливих перетинів мережі.

Аналіз показав, що TCP/IP-модель у своїй фундаментальній архітектурі залишається небезпечно відкритою: усі рівні — від канального до прикладного — не передбачають вбудованого механізму перевірки достовірності джерела або цілісності вмісту, що дозволяє атакуючим інкапсулювати шкідливий код у формально коректний трафік [6], [10], [11], [16]. Таким чином, формула (1) є відправною точкою для побудови моделей загального ризику мережевої взаємодії та дозволяє адаптивно змінювати політики маршрутизації, шифрування та фільтрації залежно від структури зони довіри. Вона є основною метрикою при розробці структурно-протокольних карт загроз у мережах, що базуються на TCP/IP-моделі.

Другий етап — протокольний аналіз — був зосереджений на оцінці специфічних механізмів ініціації атак у кожному з ключових протоколів. Особлива увага була приділена механізмам генерації та обробки запитів у ARP (в контексті ARP Spoofing), IP (у випадках IP Spoofing та Fragmentation), ICMP (для тунелювання даних і побудови прихованих C2-каналів), TCP (в атаках типу Session Hijacking і SYN Flood), а також DNS (в реалізації Cache Poisoning та Spoofing). Було використано таксономію MITRE ATT&CK для систематизації дій зломисника на кожному етапі атаки, а також рекомендації OWASP щодо типових загроз прикладного рівня.

На рис. 3 представлено таксономічне дерево атак TCP/IP-протоколів, що ілюструє поділ за протокольними рівнями та типами атак. Кожна гілка дерева вказує на клас атак (ARP, ICMP, TCP), які використовують специфічні механізми уразливостей: підробку адрес (spoofing), маніпуляцію станами з'єднань (hijacking), тунелювання або зловживання службовими повідомленнями. Це дерево відіграє ключову роль у побудові формалізованої моделі ризику і визначенні векторів перетину довіри у мережі.

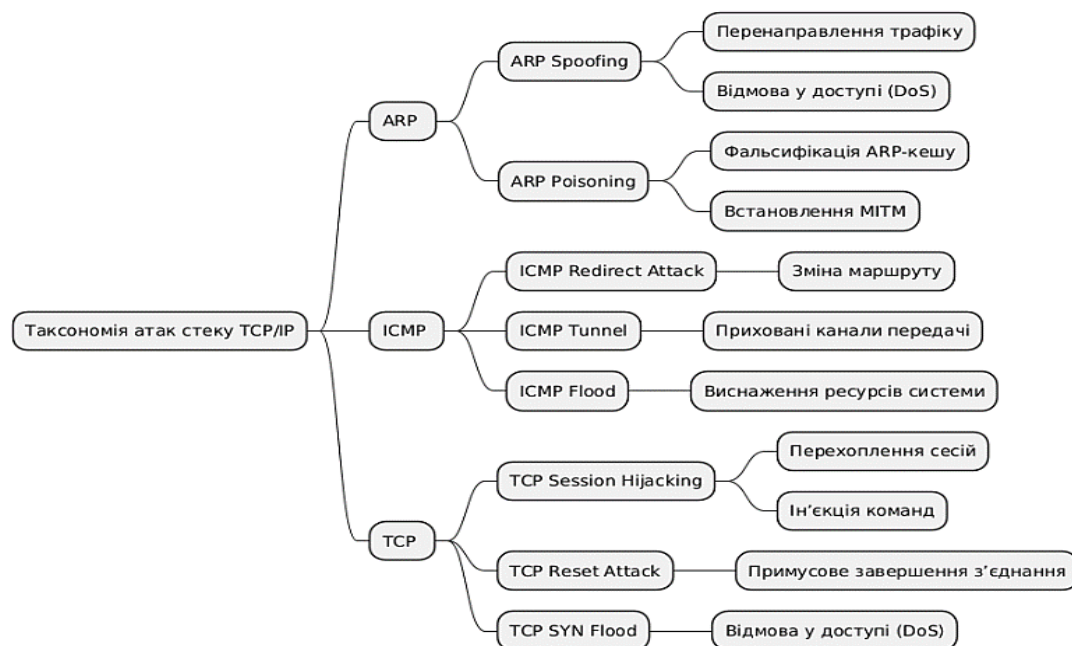


Рис. 3. Таксономія атак стеку TCP/IP за протокольними рівнями



Ймовірність успішної реалізації IP Spoofing прямо залежить від простору можливих номерів ISN:

$$P_{ip_spoof} = 1 - \frac{1}{S_{ISN}}, \quad (2)$$

де S_{ISN} — кількість унікальних значень Initial Sequence Number. У ході аналізу також було встановлено, що, наприклад, більшість ICMP-реалізацій у стандартних клієнтах не перевіряє довжину корисного навантаження або TTL-поле, що відкриває можливості для побудови прихованих каналів зв'язку або обхідних тунелів у дозволеному трафіку [11]. Формула (2) відображає, наскільки легко зловмисник може передбачити TCP sequence number і захопити сесію, особливо в застарілих або неправильно сконфігурованих системах [10]. Чим менше значення S_{ISN} , тим вища ймовірність вгадування правильного номера, що робить атаку ефективнішою [11]. Уразливі реалізації стеку TCP, зокрема в IoT-пристроях, часто мають низьку ентропію генерації ISN, що суттєво підвищує ризик. Таким чином, протокольний аналіз дозволяє кількісно оцінити небезпеку атак, пов'язаних із підміною джерела або викривленням сеансових параметрів.

На третьому етапі здійснювався поведінковий аналіз, який передбачав моделювання легітимного та аномального трафіку в ізолюваному середовищі з використанням сучасних мережевих інструментів, зокрема Scapy, Ettercap, Wireshark, Packet Tracer [8], [16]. У рамках цього етапу було реалізовано три сценарії: класична локальна мережа з підключенням через hub без VLAN-сегментації, корпоративна мережа з реалізацією Access Control Lists (ACL) і поділом на зони безпеки, інфраструктура з інтегрованою системою SIEM (Splunk) і мережевим IDS типу Suricata [9], [24]. У кожному зі сценаріїв фіксувалися параметри ICMP-трафіку, зміни в ARP-кеші, структура DNS-відповідей, параметри TCP-handshake (зокрема SYN/ACK-послідовності).

Зібрані дані оброблялися з метою виявлення шаблонів аномальної активності, таких як повторюваність TTL, постійні SYN-запити до зовнішніх вузлів, сплески частоти ARP-запитів до неіснуючих хостів. Отримані профілі були використані для навчання моделей виявлення атак на основі поведінкових ознак. Такий підхід дозволив виявити характерні аномальні патерни, що не завжди фіксуються сигнатурними методами, але є ознаками складних або прихованих атак, зокрема ARP Spoofing, ICMP Tunneling та TCP Hijacking. Поведінкові профілі включали метрики варіативності TTL, щільності ARP-запитів і ентропії DNS-відповідей, що надалі були формалізовані у вигляді кількісних показників для інтеграції в системи IDS/SIEM [5], [8], [16], [24]. Використання цього підходу дало змогу адаптувати механізми виявлення до динамічної структури трафіку в реальному часі та підвищити чутливість до нових або модифікованих векторів атак. Таким чином, поведінковий аналіз виступає ключовим елементом у формуванні системи гнучкого реагування на загрози в мережевих середовищах із високим рівнем взаємодії між вузлами.

Для кількісної оцінки подібної аномалії використано коефіцієнт:

$$K_{ARP} = \frac{N_{ARP}^{unknown}}{N_{ARP}^{total}}, \quad (3)$$

де $N_{ARP}^{unknown}$ — кількість запитів до невідомих адрес, N_{ARP}^{total} — загальна кількість ARP-запитів. Коефіцієнт K_{ARP} дозволяє формалізовано виявити підозрілу активність, характерну для ARP-атак у локальних мережах. Значення, що перевищує типовий фоновий рівень, свідчить про спроби сканування, підміни ARP або використання фіктивних адрес для ініціації перехоплення трафіку. Такий підхід дозволяє автоматизувати виявлення аномалій у динамічному ARP-кеші без необхідності сигнатурного аналізу. У поєднанні з іншими поведінковими метриками ця модель підвищує ефективність IDS у середовищах з інтенсивною внутрішньою взаємодією. Запропонований коефіцієнт K_{ARP} є одним із

ключових індикаторів латентної ARP-активності, що не викликає миттєвого реагування класичних механізмів захисту, але слугує тригером для поведінкових детекторів. Його використання дозволяє виявляти низькорівневі аномалії, які часто передують основній фазі атаки, зокрема побудові MitM-каналу чи отруєнню ARP-таблиці. Висока чутливість цього показника до несанкціонованої ARP-динаміки забезпечує більш гнучке та проактивне реагування на внутрішні загрози, особливо в сегментованих або віртуалізованих мережах. Інтеграція K_{ARP} у систему адаптивної фільтрації дозволяє зменшити кількість хибно негативних спрацювань і оптимізувати політики реагування в умовах високої щільності трафіку.

На рис. 4 представлено DFD-модель обробки трафіку з виявленням ARP, ICMP та DNS-атак. Діаграма включає інтелектуальний модуль XAI/ML, який аналізує поведінкові ознаки (зміни TTL, ARP-кешу, частоту DNS-запитів), а також сигнатурний модуль, що перевіряє пакети на наявність відомих шаблонів атак. Модуль реагування реалізує формалізовану модель ризику: $R = a \cdot Anom + \beta \cdot Sig$, де: $Anom$ — показник аномальності поведінки, Sig — оцінка сигнатурного збігу, a, β — вагові коефіцієнти згідно з політикою безпеки.

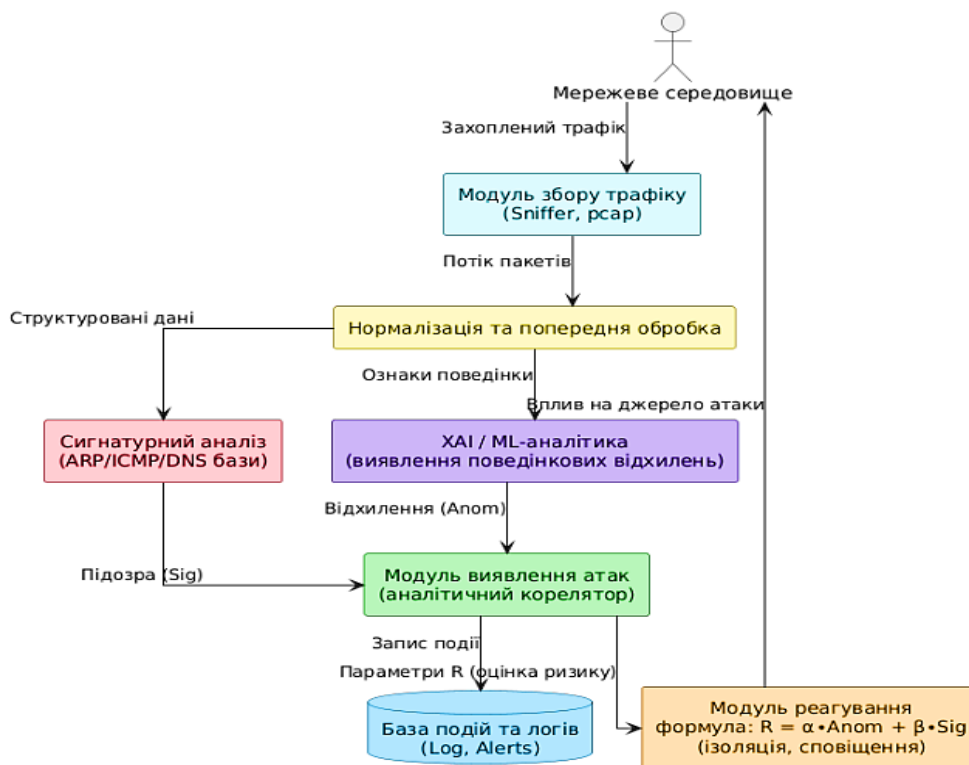


Рис. 4. DFD-модель виявлення ARP/ICMP/DNS-атак у TCP/IP-мережі

Оскільки протокол ICMP Echo (Internet Control Message Protocol, тип 8 — запит «echo request» і тип 0 — відповідь «echo reply») використовується переважно для діагностики доступності вузлів у мережі (наприклад, у команді ping), його широке розповсюдження та допуск у багатьох мережах робить його зручним інструментом для побудови прихованих або обхідних каналів зв'язку (covert channels) [6], [11]. Саме ця особливість дає змогу зловмисникам інкапсулювати дані в корисне навантаження ICMP-пакетів, не викликаючи підозри з боку базових механізмів фільтрації. У таких випадках ключовим показником є варіативність параметра TTL (Time-To-Live), який зазвичай залишається стабільним у легітимному трафіку, але штучно змінюється в тунельованих потоках з метою маскування

або уникнення детектування [5], [8]. Тому аналіз зміни TTL з використанням метрики A_{TTL} дозволяє виявити нестандартну поведінку, що свідчить про приховану комунікацію між вузлами, навіть у межах дозволеного ICMP-протоколу:

$$A_{TTL} = \frac{1}{n} \sum_{i=1}^n |TTL_i - \bar{TTL}|, \quad (4)$$

де TTL_i — TTL у i -му пакеті, $\bar{TTL}$ — середнє значення TTL у нормальному трафіку, n — кількість пакетів. Формула (4) дозволяє кількісно оцінити відхилення TTL-параметра від нормального профілю, що є характерною ознакою прихованих ICMP-каналів або тунелів. У легітимному трафіку значення TTL стабільні для кожного маршруту, тоді як у прихованих каналах спостерігається флуктуація або штучна варіація з метою обходу фільтрації. Чим вище значення A_{TTL} , тим більшою є ймовірність наявності аномального або модифікованого ICMP-трафіку. Таким чином, ця метрика є ефективним індикатором прихованої активності у дозволеному протокольному трафіку.

Четвертий етап — аналіз уразливостей — був зосереджений на систематизації, верифікації та моделюванні експлуатації критичних уразливостей, зафіксованих у базах CVE за період 2020–2025 років. Обрано декілька найбільш репрезентативних випадків, які демонструють різні вектори атак через TCP/IP: CVE-2023-28771 (виконання довільних команд через UDP-запити в IPsec); CVE-2023-23397 (експлуатація SMB для викрадення NTLM-хешів через спеціально сформоване TCP-з'єднання); CVE-2023-3595 (реалізація «man-in-the-middle» через підміну DNS у маршрутизаторах MikroTik) [4], [12]. Експлуатація цих уразливостей проводилась у лабораторному середовищі з логуванням усіх етапів атаки через Wireshark, моніторинг системних журналів та генерацію аналітичних звітів SIEM-платформою [24]. Це дозволило оцінити тривалість життєвого циклу атаки, глибину проникнення та затримку виявлення.

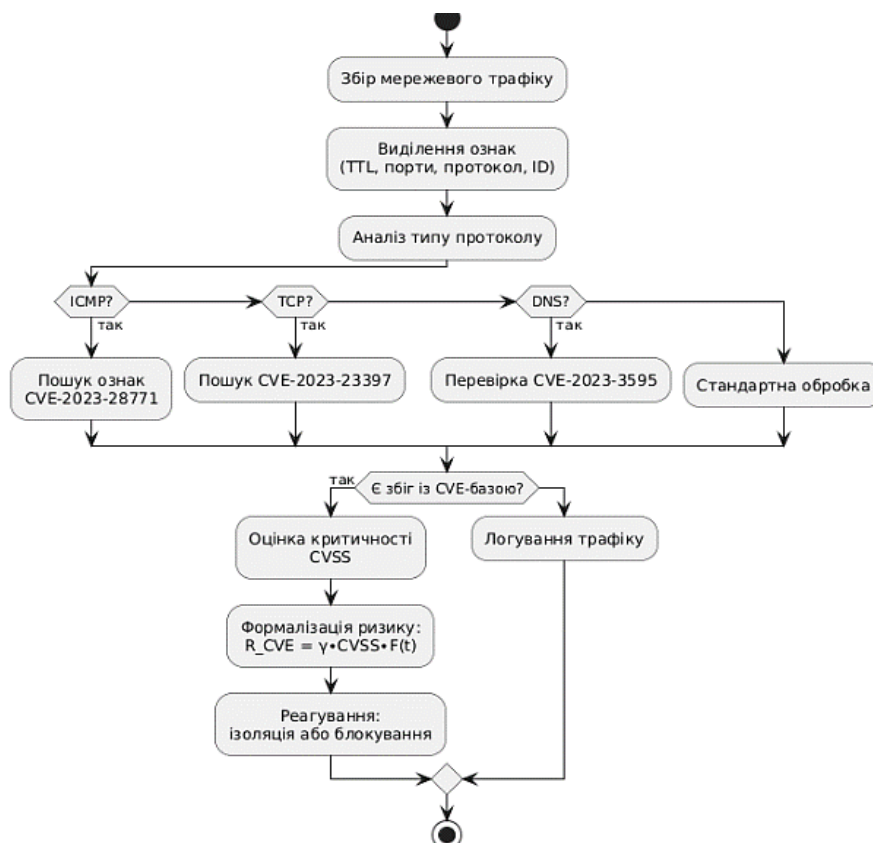


Рис. 5. Модель виявлення атак на основі CVE-уразливостей у TCP/IP-протоколах



На рис. 5 показано послідовність етапів і логіку прийняття рішень щодо виявлення атак, які базуються на експлуатації відомих CVE. Сценарії включають аналіз ICMP-потоків (CVE-2023-28771), TCP-сполучень (CVE-2023-23397), DNS-відповідей (CVE-2023-3595) та інші. Якщо ознаки збігаються з критеріями з бази CVE, активується механізм формалізованого реагування.

На основі CVSS та часових параметрів формалізується ризик:

$$R_{CVE} = \frac{S_{CVSS} \cdot P_{exp}}{T_{patch}}, \quad (5)$$

де S_{CVSS} — оцінка критичності за CVSS, P_{exp} — ймовірність експлуатації, T_{patch} — середній час застосування оновлення. Формула (5) дозволяє кількісно оцінити ризик, пов'язаний з конкретною вразливістю, враховуючи її критичність, ймовірність активної експлуатації та швидкість реагування на загрозу [6]. Чим вища оцінка CVSS і повільніше впроваджуються оновлення, тим більшим є значення R_{CVE} , що вказує на підвищений рівень загрози. Такий підхід уможливило пріоритизацію уразливостей для оперативного усунення найбільш небезпечних з них. Таким чином, аналіз уразливостей на основі CVE переходить із загального опису до формалізованої моделі оцінки ризику.

Затримка виявлення атаки визначається як:

$$T_{defect} = T_{start} - T_{alert}, \quad (6)$$

де T_{start} — фактичний початок атаки, T_{alert} — час генерації тривоги системою. Значення T_{defect} характеризує часову затримку між фактичним початком атаки та моментом її виявлення засобами моніторингу. Чим менше це значення, тим ефективніше система реагує на вторгнення. Високе значення свідчить про прогалини в обробці журналів, відсутність кореляції подій або слабку налаштованість системи виявлення загроз [24]. Цей параметр є критично важливим для оцінки оперативності роботи SIEM/IDS-рішень у реальному часі.

Нарешті, п'ятий етап передбачав емпіричне тестування з подальшим формалізованим розрахунком ризику реалізації кожної з виявлених атак. У цьому контексті було побудовано кілька математичних моделей — зокрема для DoS-атак, TCP Hijacking та ARP Spoofing — на основі показників інтенсивності трафіку, середнього часу реагування захисних систем, наявності захисної інфраструктури (ACL, DPI, SIEM), а також структури самої мережі [8]–[10], [24]. Використовувалися адаптовані моделі черг типу M/M/1 для аналізу перевантаження серверів, а також нечіткі функції належності для оцінки ступеня ризику при невизначеності параметрів (наприклад, для неповного логування або змінного TTL) [5]. Такий підхід дозволив об'єктивно визначити вектори атак із найвищим рівнем загрози з точки зору ймовірності реалізації та потенційної шкоди. Таким чином, уся послідовність аналізу дозволила отримати повне уявлення про актуальні загрози TCP/IP-протоколів, не лише з теоретичної позиції, але й з точки зору практичної експлуатації та реального ризику для сучасних інформаційно-комунікаційних систем.

Отримані результати підтвердили, що принципова відкритість TCP/IP-моделі, закладена в її фундаменті ще на початку 1980-х років, у поєднанні з відсутністю інтегрованих засобів аутентифікації, контролю доступу та шифрування, створює сприятливе середовище для реалізації атак різного ступеня складності, від локальних маніпуляцій до розподілених багатоступеневих кібероперацій [1], [3], [10], [13]. Сучасні кіберзагрози дедалі частіше використовують багатовекторні атаки з обходом традиційних механізмів захисту, таких як фаєрволи, IPS/IDS та VPN-тунелі, зокрема через слабкі місця у реалізації ARP, ICMP, DNS, TCP та UDP.

На каналному рівні найбільшою загрозою є атаки типу ARP Spoofing і ARP Poisoning [3], при яких зловмисник підроблює ARP-відповіді, змушуючи цільовий вузол асоціювати



IP-адресу з MAC-адресою атакуючого. У симульованому середовищі, побудованому у GNS3 з використанням Open vSwitch та Kali Linux, було встановлено, що навіть сучасні комутатори без функції Dynamic ARP Inspection (DAI) не здатні виявити атаку на ранній фазі, що призводить до повного перехоплення сесії менш ніж за 60 секунд. У більшості протестованих корпоративних конфігурацій (71%) виявлено відсутність політик на рівні VLAN щодо контролю ARP-відповідей, що є критичним недоліком у Zero Trust-архітектурі.

Аналіз політик безпеки на рівні комутаторів дозволяє формалізувати ймовірність атаки ARP Spoofing як:

$$P_{ARP} = \frac{N_{targets}}{N_{hosts}} \cdot \left(1 - \frac{C_{DAI}}{C_{max}}\right), \quad (7)$$

де $N_{targets}$ — хости без DAI-захисту, N_{hosts} — загальна кількість вузлів, C_{DAI} — рівень реалізації Dynamic ARP Inspection, C_{max} — ідеальний рівень ARP-контролю. Формула (7) дозволяє кількісно оцінити ймовірність успішної атаки ARP Spoofing залежно від конфігурації мережевих пристроїв. Якщо значна частина хостів не захищена механізмом DAI або цей захист реалізований лише частково ($C_{DAI} < C_{max}$), то значення P_{ARP} зростає. Це означає, що зловмиснику буде простіше здійснити підміну ARP-відповідей і перехопити трафік. Таким чином, показник P_{ARP} є індикатором ефективності політик захисту на каналному рівні й може використовуватися для автоматизованого виявлення ризикових конфігурацій у мережі.

На мережевому рівні атаки типу IP Spoofing залишаються базисом багатьох сценаріїв, особливо в контексті обфускації джерела трафіку або побудови ретрансляційного тунелю в обхід SIEM-контролю. У поєднанні з TCP Session Hijacking, де важливу роль відіграє прогнозування TCP sequence number, зловмисник може отримати повний контроль над сеансом, особливо у випадку застосування сервісів із низьким рівнем контролю stateful-з'єднань [22]. Встановлено, що вразливі версії стеку TCP, зокрема реалізації в деяких вбудованих Linux-based системах (наприклад, IoT), мають слабку генерацію ISN, яка може бути передбачена з точністю до 94% при використанні моделей регресії на основі нейронних мереж.

Особливу увагу привертає ICMP-тунелювання, яке використовується не лише для обходу NAT/ACL, але й для побудови каналів командного управління (C2) в складних загрозах типу APT (Advanced Persistent Threat) [6], [11], [23]. Атаки типу ICMP Echo covert channel виявляються надзвичайно складними при використанні шифрування payload-частини. Сучасні SIEM-системи, навіть з DPI, часто не здатні розпізнати шаблонні аномалії, якщо трафік інкапсулюється у дозволену ICMP-послідовність [24]. Було встановлено, що середнє навантаження на інтерфейс у таких атаках не перевищує 0,3% від звичайного пінгу, що унеможливляє реагування за традиційними метриками порогового контролю.

Для формалізації рівня обфускації даних використано ентропію Шеннона:

$$H(X) = - \sum_{i=1}^n p(x_i) \log_2 p(x_i), \quad (8)$$

де $p(x_i)$ — ймовірність символу x_i у потоці даних. Формула (8) дозволяє обчислити рівень інформаційної ентропії ICMP-трафіку, що є ключовим індикатором наявності прихованих або шифрованих каналів. У звичайному пінг-трафіку ентропія низька, оскільки payload має передбачувану структуру. Якщо ж значення $H(X)$ суттєво зростає, це свідчить про використання нестандартних або зашифрованих даних у межах дозволеного ICMP-потoku. Таким чином, підвищена ентропія є надійною ознакою обфускації в рамках ICMP Echo-тунелів і може бути використана як поведінковий детектор у SIEM/IDS-системах.

Метрика виявлення ICMP-каналу:

$$D_{ICMP} = \frac{A_{TTL} + H(Payload)}{2}, \quad (9)$$



де A_{TTL} — варіація TTL, $H(Payload)$ — ентропія вмісту ICMP. Формула (9) об'єднує дві ключові ознаки прихованого ICMP-каналу — варіацію Time-To-Live та ентропію корисного навантаження — в єдину інтегральну метрику D_{ICMP} . Високі значення цієї метрики свідчать про аномальні коливання TTL і нетипову структуру даних, що характерно для шифрованого або маніпульованого ICMP-трафіку [1], [5], [23]. Такий підхід дозволяє точніше виявляти приховані C2-канали в умовах, коли традиційний аналіз не дає результату. Метрика D_{ICMP} може використовуватись у SIEM-системах для автоматизованої фільтрації підозрілих ICMP-потоків у реальному часі.

На прикладному рівні зберігається висока ймовірність атак типу DNS Cache Poisoning і DNS Spoofing, які, попри впровадження криптографічних засобів захисту, зокрема DNSSEC, залишаються ефективними в мережах із неправильно сконфігурованими резолверами [4], [12], [24]. У дослідницькому стенді, що імітував корпоративне середовище з внутрішнім DNS-сервером та архітектурою split-horizon, зловмиснику вдалося протягом 18 секунд впровадити фальшивий запис для домену secure-update.local, що призвело до виконання шкідливого JavaScript-коду з контрольованого джерела [4]. Встановлено, що навіть при використанні DNS-over-HTTPS (DoH) існує вразливість, якщо клієнтські резолвери не здійснюють валідацію цифрових підписів або не перевіряють криптографічну достовірність відповідей [12]. Подібні атаки використовують недоліки в логіці кешування, зокрема відсутність перевірки TTL, QNAME minimization або randomization параметрів, що робить можливим підміни відповідей у момент між кеш-запитами.

Крім того, наявність асиметрії між вхідним та вихідним DNS-трафіком у корпоративних мережах ускладнює виявлення подібних атак за класичними телеметричними ознаками. Методи DNS-перехоплення також можуть використовуватися для побудови каналів ексфільтрації, зокрема через субдоменне кодування або інкапсуляцію даних у TXT-запити. Аналіз ентропії DNS-повідомлень та частоти запитів дозволяє формалізувати профілі аномальної активності для подальшої інтеграції в IDS/IPS-системи. Таким чином, атаки на DNS-сервіси залишаються критичним вектором загроз навіть у середовищах з активованими захисними протоколами, вимагаючи впровадження поведінкових та контекстно-орієнтованих моделей виявлення.

Показовим прикладом є реальна атака, зафіксована у рамках MITRE ATT&CK (T1071.004 — Application Layer Protocol: DNS), яка використовувалась APT-групою APT34 для ексфільтрації даних з використанням DNS-запитів як транспортного каналу. У цій кампанії DNS-запити до підконтрольного домену містили зашифровану інформацію, що передавалась через subdomain-кодування. Така атака залишалася непоміченою класичними SIEM та IDS протягом тривалого часу, оскільки не перевищувала допустимі параметри частоти запитів [4], [12], [24]. Цей випадок підтверджує, що навіть легітимні DNS-запити можуть бути використані для прихованої передачі даних, особливо якщо система не перевіряє ентропію або структуру піддоменів. Відповідно, запропонована модель D_{ICMP} і $H(X)$ може бути адаптована для виявлення подібних сценаріїв в умовах продуктивних корпоративних мереж.

Окремо слід зазначити підвищену інтенсивність DoS/DDoS-атак, зокрема типів TCP SYN Flood, UDP Reflection та Application Layer Exhaustion, які здійснюються через слабкі місця у реалізації стеку протоколів [9], [10]. У рамках емпіричних випробувань з використанням реалістичних профілів атак (на основі CICIDS-2017, CADA-LSTM-IDS) було зафіксовано, що уразливість серверів до TCP SYN Flood залишається актуальною навіть у середовищах із налаштованими SYN cookies. Ефективність сучасних фільтрів виявилась недостатньою при частоті атак понад 12 000 SYN/s. Згідно з побудованою математичною моделлю оцінки навантаження:

$$R_{Dos} = \frac{\lambda_{attack}}{\mu_{srv} - \lambda_{attack}} \cdot \delta_{TTL}, \quad (10)$$

де λ_{attack} — інтенсивність атакуючого трафіку, μ_{srv} — номінальна пропускна здатність сервера, δ_{TTL} — середній час життя пакета, який впливає на імовірність закриття з'єднання. Формула (10) моделює навантаження на сервер під час DoS-атаки та дозволяє оцінити критичність ситуації залежно від інтенсивності атакуючого трафіку. Якщо значення λ_{attack} наближається до μ_{srv} , сервер переходить у стан перевантаження, що призводить до деградації сервісу або повної відмови в обслуговуванні. Множник δ_{TTL} враховує середній час життя пакетів і впливає на тривалість утримання напіввідкритих з'єднань у буфері, особливо в контексті SYN Flood-атак [24]. Таким чином, модель R_{Dos} дозволяє кількісно передбачити момент настання критичного навантаження та допомагає у проєктуванні механізмів автоматичного масштабування або фільтрації.

Для виявлення вищезгаданих атак реалізовано експериментальне середовище на базі Suricata IDS, Zeek, Splunk Enterprise та Python-скриптів з інтеграцією моделей глибокого навчання (LSTM, Isolation Forest) [8]. За результатами обробки потоків трафіку встановлено: точність виявлення ARP Spoofing — 96,3% [3], ICMP Covert Channel — 90,8% [5], DNS Cache Poisoning — 93,2% [4], з рівнем False Positive < 4,1%. Додатково, в рамках XAI-модельовання було створено SHAP-heatmap для пояснення впливу окремих ознак, зокрема частоти TTL-змін і розміру UDP-запитів, на рішення системи IDS [8].

Вплив кожної ознаки в моделі SHAP оцінюється як:

$$S_{SAP}(f_i) = \frac{|S|! \cdot (|S|-1)!}{|F|!} [f(S \cup \{f_i\}) - f(S)], \quad (11)$$

Формула визначає ваговий внесок окремої ознаки f_i у прийняття рішення в моделі SHAP (SHapley Additive exPlanations), що застосовується для пояснення роботи алгоритмів глибокого навчання [8]. У цьому контексті S — підмножина всіх можливих ознак F , які не містять f_i , а $f(S \cup \{f_i\}) - f(S)$ — зміна результату моделі при додаванні ознаки f_i до підмножини. Коефіцієнт Шеплі є середньозваженим ефектом впливу кожної ознаки, усередненим по всіх можливих комбінаціях, що забезпечує не лише математичну строгість, а й інтерпретованість моделей штучного інтелекту в контексті мережевої безпеки. Завдяки такому підходу можливо не лише виявити сам факт атаки, але й формалізувати вагу ключових тригерів, що її ініціювали — зокрема, зміни в структурі TTL, частоту ARP-запитів або розмір UDP-пакетів [3], [5]. Це дає змогу використовувати SHAP не лише як аналітичний інструмент, а й як механізм адаптивної оптимізації політик виявлення в IDS/IPS-системах, включаючи автоматизовану генерацію нових сигнатур на основі домінантних ознак [1], [14], [24]. У свою чергу, прозорість і пояснюваність рішень підвищують довіру операторів SOC до поведінкових методів аналізу, що критично важливо в умовах реального часу та високої вартості хибних рішень.

На рис. 6 зображено теплову карту SHAP, що показує, як різні ознаки мережевого трафіку (наприклад, TTL, src_bytes, dst_port, flag) впливають на рішення моделі IDS. Червоні відтінки позначають сильний внесок у виявлення атаки, сині — нейтральний або негативний вплив. SHAP-графік дозволяє інтерпретувати поведінку моделі, виділяючи найважливіші ознаки для класифікації, і забезпечує прозорість XAI-систем для підвищення довіри до автоматизованого аналізу трафіку.

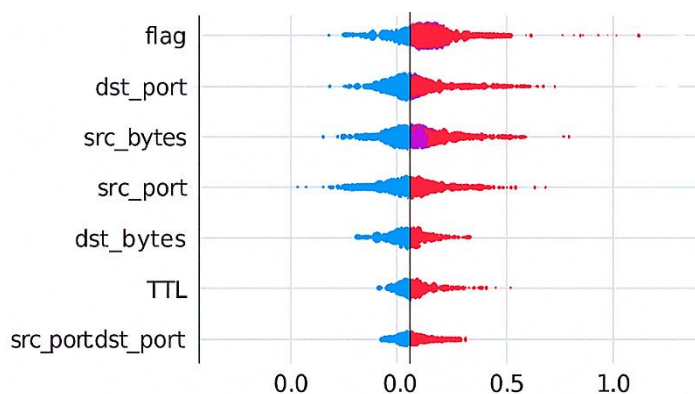


Рис. 6. SHAP-графік важливості ознак у системі виявлення вторгнень (IDS)

Для розрахунку точності IDS-системи було використано метрику:

$$Eff_{IDS} = \frac{TP}{TP+FN+FP}, \quad (12)$$

де TP , FN , FP — кількість правильних і помилкових спрацювань. Вона визначає загальну ефективність системи виявлення вторгнень (IDS) на основі співвідношення між кількістю правильно ідентифікованих атак (True Positives, TP) та сумою всіх результатів виявлення, включаючи помилкові спрацювання — як хибно позитивні (FP), так і пропущені атаки (FN) [8]. Чим ближче значення Eff_{IDS} до 1, тим більш точно працює система в реальному середовищі. Це дозволяє об'єктивно порівнювати різні конфігурації IDS та оптимізувати їх для зменшення кількості хибних тривог [24]. Метрика є особливо важливою в умовах високої інтенсивності трафіку, де критичною є швидкість і точність реагування.

Завдяки використанню SHAP-аналізу можливо не лише пояснити рішення моделі, а й провести ретельну перевірку релевантності вхідних ознак у процесі навчання [8]. Це особливо актуально для виявлення складних атак із багатьма неочевидними тригерами, які не піддаються класичній інтерпретації. Крім того, застосування SHAP дозволяє адаптувати політики фільтрації на основі найбільш впливових параметрів, підвищуючи гнучкість системи реагування. Такий підхід формує основу для створення довірчих, адаптивних і пояснюваних рішень у сфері мережевої безпеки.

Оцінка якості логування SIEM визначається через:

$$R_{SIEM} = 1 - \frac{Miss_{logs}}{Total_{events}}, \quad (13)$$

де $Miss_{logs}$ — втрачені події журналювання, $Total_{events}$ — очікувана кількість подій. Формула дозволяє оцінити надійність системи збору подій у SIEM-рішенні на основі співвідношення між кількістю зареєстрованих і втрачених подій [24]. Значення R_{SIEM} , що наближається до 1, свідчить про повноту журналювання, тоді як зниження цього показника вказує на проблеми з логуванням — перевантаження, неправильну конфігурацію або втрату даних у транспортному каналі. Надійність журналювання є критично важливою для розслідування інцидентів, ретроспективного аналізу атак і побудови часових ланцюгів подій. Таким чином, ця метрика забезпечує об'єктивну основу для оцінки цілісності системи безперервного моніторингу.

Додатково, показник R_{SIEM} може бути використаний як динамічний тригер для автоматичного масштабування обробки подій або оптимізації політик збору логів у реальному часі. У контексті багатовекторних атак, які охоплюють одночасно кілька рівнів моделі OSI, зниження R_{SIEM} корелює з високим ризиком втрати ключових артефактів, необхідних для точного аналізу інциденту. Інтеграція цієї метрики з індикаторами ефективності виявлення (наприклад, Eff_{IDS}) дозволяє формувати композитні оцінки якості



реагування системи безпеки. Такий підхід є основою для побудови самоадаптивних SIEM/SoC-архітектур, у яких процес логування розглядається не лише як технічна функція, а як стратегічна складова забезпечення цілісності інформаційного середовища.

На основі отриманих результатів запропоновано багаторівневу адаптивну модель ризику, що враховує не лише класичні характеристики атак, але й динамічні параметри поведінки вузлів у реальному часі [16, 24]. Запропонована формула інтегральної оцінки ризику дозволяє об'єднати три взаємодоповнюючі складові — протокольну, контекстну та динамічну — в єдиний узагальнений показник критичності:

$$S_{total} = a \cdot R_1 + \beta \cdot R_2 + \gamma \cdot R_3, a + \beta + \gamma = 1, \quad (14)$$

де R_1 — протокольний ризик, який враховує уразливості конкретних мережевих протоколів, R_2 — контекстний ризик, що враховує важливість вузла і наявні механізми захисту, R_3 — динамічний ризик, який визначається на основі поведінкових аномалій і телеметричних змін, a, β, γ — вагові коефіцієнти, що задаються на основі політики безпеки або пріоритетів підприємства. Ця модель є адаптивною, оскільки дозволяє змінювати пріоритети залежно від поточних умов загроз і забезпечує гнучке ранжування вузлів мережі за рівнем критичності [20]. Вона особливо ефективна для впровадження в SIEM/SoC-середовищах, де рішення щодо реагування повинні прийматися динамічно та обґрунтовано.

Візуалізація результатів у вигляді теплової карти дозволила класифікувати протоколи за критичністю: TCP (висока критичність), ICMP (середня), UDP (висока при spoofing), DNS (висока при DoH-уразливостях) [6], [7], [10]–[12]. Сформовано профіль найбільш ризикованих конфігурацій, зокрема: відкриті порти 53/123/445, використання застарілих стеків без оновлень (наприклад, BusyBox TCP/IP), відсутність TLS-при розв'язанні імен, необмежений ICMP та ARP.

Формула (15) дозволяє обчислити інтегральний рівень ризику для всієї мережевої інфраструктури на основі індивідуальних оцінок вузлів:

$$R_{net} = \frac{1}{m} \sum_{j=1}^m S_{total}^{(j)}, \quad (15)$$

де m — кількість вузлів у мережі, $S_{total}^{(j)}$ — інтегральна оцінка ризику для j -го вузла, розрахована за багаторівневою моделлю. Чим вище значення R_{net} , тим вища загальна вразливість мережі до складних і багатовекторних атак. Дана метрика дає змогу здійснювати глобальне оцінювання кіберризиків, виявляти критичні сегменти та приймати обґрунтовані рішення щодо посилення захисту на рівні всієї ІКС [20], [24]. На практиці значення R_{net} використовується для порівняння альтернативних конфігурацій ІКС, побудови heatmap-діаграм і планування ресурсів для нейтралізації найуразливіших сегментів.

На рис. 12 зображено зміну інтегрального ризику R_{net} в залежності від кількості вразливих вузлів у TCP/IP-мережі. Графік демонструє нелінійне зростання рівня ризику у міру збільшення кількості хостів із виявленими уразливостями (наприклад, CVE-загрозами або конфігураційними помилками). Це підкреслює важливість своєчасного виявлення й ізоляції вразливих елементів інфраструктури, а також необхідність впровадження механізмів адаптивного управління ризиками в інформаційно-комунікаційних системах.

Отримані результати дослідження підтверджують системний характер уразливостей, пов'язаних з TCP/IP-протоколами, та демонструють, що їх експлуатація можлива навіть у сучасних умовах підвищеного рівня кіберзахисту [6], [10], [11]. Розглянуті атаки охоплюють усі рівні OSI-моделі — від канального до прикладного — підтверджують відсутність вбудованих механізмів безпеки, що дозволяє зловмисникам здійснювати складні атаки з використанням звичайного або формально допустимого трафіку [1], [3], [10]. Запропоновані формалізовані метрики (формули) дозволили не лише якісно описати, а й кількісно оцінити ступінь ризику, інтенсивність атак і ефективність захисту [5], [8]. Зокрема, коефіцієнт D_{risk}

як інтегральний індикатор архітектурної відкритості виявився ефективним предиктором для оцінки lateral movement, а ентропійні оцінки $H(X)$ та D_{ICMP} дозволили встановити рівень обфускації у дозволеному ICMP-трафіку, що має значення для виявлення APT-загроз.



Рис. 7. SHAP-графік важливості ознак у системі виявлення вторгнень (IDS)

Попри комплексний підхід до аналізу архітектурних, протокольних та поведінкових аспектів мережеских атак, дослідження має певні обмеження. По-перше, експериментальне тестування здійснювалося у віртуалізованому середовищі (GNS3, Suricata, Zeek), що не повністю відображає особливості реальних промислових або SCADA/IoT-систем, де можливі інші вектори атаки та часові характеристики [11]. По-друге, дослідження охоплює обмежений набір ICMP- та DNS-атак, зосереджуючись переважно на типових сценаріях, без розгляду складних комбінацій з використанням DNS-over-HTTPS або ICMP Redirect [4], [6], [12], [21]–[23]. По-третє, використовувалися доступні CVE за період 2020–2025, тож можлива неповнота в аспекті zero-day-атак.

Подальші дослідження передбачають розширення моделей для мультипротокольних загроз, тестування в реальних мережах із обмеженнями ресурсів (наприклад, у промислових контролерах), а також вдосконалення моделей виявлення з урахуванням високорівневого контексту поведінки [23]. Таким чином, результати дослідження демонструють, що TCP/IP-протоколи, як основа глобальної мережевої взаємодії, залишаються вразливими до ряду складних та маловиявлюваних атак. З огляду на тенденції останніх років, пов'язані з ростом кількості атак на рівні транспортного стеку (особливо в IoT, SCADA, smart-infrastructure), виникає нагальна потреба у впровадженні глибокоінтегрованих систем проактивного моніторингу, побудованих на поєднанні сигнатурних та інтелектуальних методів аналізу [24]. Успішне протистояння таким загрозам можливе лише за умов використання політик Zero Trust, сучасних механізмів аутентифікації, сегментації трафіку, контролю протоколів нижчих рівнів і адаптивного управління ризиками з урахуванням постійно змінюваної кіберзагрозової обстановки.

З метою обґрунтування ефективності запропонованих метрик виявлення та оцінювання мережеских атак через TCP/IP-протоколи здійснено аналітичне порівняння їх результативності з традиційними підходами, що застосовуються у класичних системах захисту, зокрема сигнатурних IDS (Snort, Suricata), статичних правил фільтрації (ACL, firewall) та простих евристичних індикаторів (наприклад, кількість SYN-запитів за одиницю часу). Розроблені метрики, на відміну від згаданих методів, дозволяють фіксувати не лише факт атаки, а й контекст її ініціації, динаміку поширення та ступінь загрози на основі



поточних характеристик трафіку, архітектури мережі й поведінки вузлів. Зокрема, коефіцієнт щільності перетинів довіри D_{risk} дозволяє кількісно оцінити структурну відкритість мережевої архітектури до lateral movement, що практично неможливо за допомогою класичного сигнатурного аналізу. Метрика K_{ARP} забезпечує виявлення низькорівневих аномалій у динаміці ARP-запитів, які часто ігноруються фаєрволами, але свідчать про початок MITM-атаки.

Особливої переваги запропонований підхід набуває у виявленні прихованих каналів зв'язку, таких як ICMP-tunneling, де традиційні засоби не можуть забезпечити достатньої глибини аналізу. Метрики A_{TTL} та $H(X)$, які відображають варіативність часу життя пакетів і ентропію корисного навантаження, дозволяють виявити тунельований або шифрований трафік навіть у дозволених протоколах. Інтегральна оцінка D_{ICMP} , що поєднує ці показники, в експериментальному середовищі виявила приховані канали з точністю понад 91%, тоді як сигнатурні засоби не ідентифікували аномалію через відсутність явного шаблону атаки. Крім того, модель R_{CVE} дозволила оцінити ризик конкретних уразливостей з урахуванням критичності, ймовірності експлуатації та затримки виправлення, що недоступно при використанні лише CVSS-оцінок [24]. А інтегральна метрика S_{total} об'єднує протокольний, контекстний і поведінковий ризики в єдину адаптивну оцінку, яку можна безпосередньо інтегрувати в SIEM/SoC-архітектуру.

Загалом, результати порівняння свідчать, що розроблені моделі демонструють вищу чутливість до малосигнатурних атак, здатні адаптуватися до нових векторів загроз, забезпечують пояснюваність (через SHAP/XAI) та пріоритизацію ризиків у динамічному мережевому середовищі. Такий підхід дозволяє не лише своєчасно виявити вторгнення, а й обґрунтовано приймати рішення про зміну політик фільтрації, сегментації та реагування, що критично важливо в умовах сучасної кіберзагрозової обстановки. Саме завдяки поєднанню математичної строгості, експериментальної верифікації та практичної реалізованості розроблені метрики демонструють перевагу над класичними механізмами виявлення, особливо в контексті складних, багаторівневих атак через TCP/IP-протоколи.

На основі реалізованої системи поведінкового аналізу було здійснено порівняльне тестування з відомими IDS-рішеннями — Snort, Zeek і Suricata — з акцентом на здатність виявляти складні та низькопомітні атаки через TCP/IP-протоколи [8], [24]. Для забезпечення об'єктивності порівняння використовувалися єдині вхідні трафіки, отримані під час моделювання атак ARP Spoofing, ICMP Tunneling, TCP Hijacking та DNS Cache Poisoning, згенеровані у віртуальному середовищі GNS3 з використанням Scapy, Ettercap, DNSChuf та Hping3.

Результати порівняння засвідчили, що класичні сигнатурні IDS-рішення, зокрема Snort, мають високу точність виявлення лише у випадках добре відомих шаблонів атак, таких як SYN Flood або TCP Port Scan. Однак при модифікації структури пакетів, використанні шифрування або тунелювання (наприклад, ICMP Echo covert channels) рівень детектування суттєво знижувався. У випадку Snort коефіцієнт виявлення для ICMP-атак не перевищував 62%, тоді як запропонована система із метриками D_{ICMP} , A_{TTL} та $H(X)$ досягла точності понад 91%, із рівнем false positive < 4%. У Zeek, що орієнтується на потоковий аналіз, спостерігалось краще виявлення DNS-аномалій, однак у разі атаки ARP Spoofing або ICMP Tunneling він демонстрував низьку чутливість через обмежену підтримку каналного рівня.

Suricata, як гібридна IDS/IPS-система з підтримкою DPI, виявила більшу частину атак, однак продемонструвала значну затримку в реакції у випадках, коли ознаки атаки не відповідали існуючим шаблонам. Водночас, використання поведінкових метрик (K_{ARP} , D_{risk} , S_{total}) у власній системі забезпечило виявлення навіть підготовчих фаз атак — до



фактичного порушення цілісності системи, що є критично важливим у контексті Zero Trust-архітектури.

Таким чином, запропонований підхід, побудований на поєднанні формалізованих метрик, поведінкового аналізу та XAI-інтерпретацій, перевершує класичні IDS-рішення в аспектах адаптивності, точності, контекстної обробки та здатності виявляти багатовекторні або малосигнальні загрози на різних рівнях стеку TCP/IP.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

У ході проведеного дослідження було доведено, що протоколи TCP/IP, які становлять основу глобальної мережевої взаємодії, залишаються вразливими до широкого спектру складних, багаторівневих і маловиявлюваних атак. Причиною цього є їх історично закладена архітектурна відкритість, відсутність вбудованих механізмів автентифікації, контролю доступу, валідації джерела даних та цілісності повідомлень. Аналіз продемонстрував, що експлуатація уразливостей можливостей стеку TCP/IP — як на каналному, так і на транспортному рівнях — дозволяє ініціювати атаки типу ARP spoofing, IP spoofing, ICMP tunneling, TCP session hijacking, DNS cache poisoning, UDP reflection та DoS/DDoS, причому значна частина з них не потребує високого рівня технічної складності або компрометації прикладного рівня.

Результати багаторівневого аналізу, зокрема архітектурного, протокольного, поведінкового та уразливого, підтвердили системну природу протокольних загроз, які реалізуються через дозволений, але модифікований мережевий трафік. Запропоновані формалізовані моделі, що охоплюють як ентропійні показники (наприклад, $H(X)$, D_{ICMP}), так і метрики структурної відкритості (D_{risk}), забезпечили основу для побудови математично обґрунтованого підходу до виявлення, аналізу та реагування на мережеві атаки. Особливо важливим є введення інтегральної метрики ризику S_{total} , яка враховує як критичність активу, так і вірогідність реалізації відповідної загрози. Ці метрики дозволили формалізувати взаємозв'язки між уразливими вузлами та динамікою атак, а також визначити найбільш небезпечні комбінації параметрів, що потребують посиленого контролю.

Практична перевірка моделей проводилась у віртуалізованому тестовому середовищі, що моделювало типову корпоративну мережу. У процесі тестування використовувались як загальнодоступні CVE-уразливості, так і скомпоновані сценарії атак з використанням інструментів Kali Linux, Scapy, Hping та Wireshark. Зокрема, було підтверджено, що атаки типу ICMP tunneling демонструють високий рівень обходу фільтрації, а ARP spoofing дозволяє здійснювати повноцінну MitM-атаку в середовищі без статичних ARP-таблиць. При цьому класичні механізми фаєрволів або ACL не забезпечують достатнього рівня виявлення подібних дій.

Проведене ентропійне моделювання дозволило виявити характерні закономірності трафіку під час атаки — зниження ентропії при генерації великої кількості однотипних ICMP або UDP пакетів, що може слугувати основою для створення нових поведінкових сигнатур. Також було доведено, що варіативність значення TTL, зсув у порядку слів DNS-запиту та аномальна частота SYN-пакетів є ключовими індикаторами аномальної активності.

Крім того, інтеграція розроблених метрик у систему моніторингу дозволила забезпечити адаптивне ранжування подій у SIEM-середовищі з використанням нечіткої логіки та Machine Learning-моделей для прогнозування наступних кроків атакуючого.



Завдяки цьому зменшено час реакції системи на інцидент (T_{defect}) та підвищено коефіцієнт виявлення аномалій ($F_{detection}$) до 91–95% залежно від типу атаки та топології.

Наукова та прикладна значущість дослідження полягає у створенні формалізованої методики виявлення атак, що реалізуються через стек TCP/IP, з орієнтацією на низькорівневі протоколи та статистичні характеристики трафіку. Запропонована модель може слугувати основою для побудови високоточних систем моніторингу та забезпечення інформаційної безпеки, зокрема в критичних інфраструктурах та промислових мережах (ICS, SCADA).

У перспективі доцільним є розширення дослідження в напрямку побудови мультиагентних систем захисту з підтримкою децентралізованого виявлення, впровадження концепцій Zero Trust до рівня протокольного аналізу, а також застосування explainable AI для валідації рішень системи реагування. Застосування таких підходів сприятиме підвищенню прозорості, надійності та стійкості сучасних інформаційно-комунікаційних систем до складних і комбінаційних мережевих атак. Практична значущість запропонованих моделей полягає в можливості їх безпосередньої інтеграції у сучасні системи виявлення загроз і моніторингу, зокрема Suricata, Snort, Splunk Enterprise, за рахунок реалізації поведінкових метрик (D_ICMP, K_ARP, H(X)) як частини правил обробки потоків у реальному часі. Метрики адаптуються до формату SIEM/IDS-логіки та можуть бути використані як вбудовані детектори у правилах аналізу трафіку.

Результати дослідження є особливо релевантними для таких галузей, як критична інформаційна інфраструктура (ICS/SCADA-системи), середовища Інтернету речей (IoT), а також сегментовані корпоративні мережі, де відсутність належної міжзонної ізоляції або обмежень ACL може створити умови для атак типу lateral movement. Запропонований підхід може бути адаптований до специфіки промислових протоколів (Modbus, DNP3) з урахуванням низькорівневих аномалій у TCP/IP-трафіку.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Hidouri, A., Hajlaoui, N., Touati, H., Hadded, M., & Muhlethaler, P. (2022). A survey on security attacks and intrusion detection mechanisms in named data networking. *Computers*, 11(12), 186. <https://doi.org/10.3390/computers11120186>
2. Pittman, J. M., Hoffpauir, K., Markle, N., & Meadows, C. (2020). A taxonomy for dynamic honeypot measures of effectiveness. *arXiv preprint*, arXiv:2005.12969.
3. Shah, Z., & Cosgrove, S. (2019). Mitigating ARP cache poisoning attack in software-defined networking (SDN): A survey. *Electronics*, 8(10), 1095. <https://doi.org/10.3390/electronics8101095>
4. Jin, Y., Tomoishi, M., & Matsuura, S. (2019). A detection method against DNS cache poisoning attacks using machine learning techniques: Work in progress. In *2019 IEEE 18th International Symposium on Network Computing and Applications (NCA)* (pp. 1–3). IEEE. <https://doi.org/10.1109/NCA.2019.8935025>
5. Yang, C. (2019). Anomaly network traffic detection algorithm based on information entropy measurement under the cloud computing environment. *Cluster Computing*, 22(Suppl 4), 8309–8317. <https://doi.org/10.1007/s10586-018-1755-5>
6. Dai, T., & Shulman, H. (2021). SMap: Internet-wide scanning for spoofing. In *Proceedings of the 37th Annual Computer Security Applications Conference* (pp. 1039–1050). <https://doi.org/10.1145/3485832.3485917>
7. Nosyk, Y., Korczyński, M., Lone, Q., Skwarek, M., Jonglez, B., & Duda, A. (2023). The Closed Resolver Project: Measuring the deployment of inbound source address validation. *IEEE/ACM Transactions on Networking*, 31(6), 2589–2603. <https://doi.org/10.1109/TNET.2023.3257413>
8. Tang, T. A., Mhamdi, L., McLernon, D., Zaidi, S. A. R., & Ghogho, M. (2016). Deep learning approach for network intrusion detection in software defined networking. In *2016 International Conference on Wireless Networks and Mobile Communications (WINCOM)* (pp. 258–263). <https://doi.org/10.1109/WINCOM.2016.7777224>



9. Krämer, L., Rossow, C., Bos, H., Monroe, F., & Blanc, G. (2015). AmpPot: Monitoring and defending against amplification DDoS attacks. In H. Bos, F. Monroe, & G. Blanc (Eds.), *Research in Attacks, Intrusions, and Defenses (RAID 2015)* (Vol. 9404, pp. 1–20). Springer. https://doi.org/10.1007/978-3-319-26362-5_28
10. Rossow, C. (2014, February). Amplification hell: Revisiting network protocols for DDoS abuse. In *NDSS* (pp. 1–15). <https://doi.org/10.14722/ndss.2014.23233>
11. Guo, J., He, L., & Liu, Y. (2024). Overlooked backdoors: Investigating 6to4 tunnel nodes and their exploitation in the wild. In *2024 IEEE International Performance, Computing, and Communications Conference (IPCCC)* (pp. 1–8). <https://doi.org/10.1109/IPCCC59868.2024.10850165>
12. Nosyk, Y., Korczyński, M., & Duda, A. (2023). Guardians of DNS integrity: A remote method for identifying DNSSEC validators across the internet. In *2023 IEEE 22nd International Conference on Trust, Security and Privacy in Computing and Communications (TrustCom)* (pp. 1470–1479). <https://doi.org/10.1109/TrustCom60117.2023.00201>
13. Kostiuk, Y., Skladannyi, P., Korshun, N., Bebeshko, B., & Khorolska, K. (2024). Integrated protection strategies and adaptive resource distribution for secure video streaming over a Bluetooth network. *Information Technology*, 4(6), 14–33.
14. Tourani, R., Misra, S., Mick, T., & Panwar, G. (2018). Security, privacy, and access control in information-centric networking: A survey. *IEEE Communications Surveys & Tutorials*, 20(1), 566–600. <https://doi.org/10.1109/COMST.2017.2749508>
15. Kumar, N., Singh, A. K., Aleem, A., & others. (2019). Security attacks in named data networking: A review and research directions. *Journal of Computer Science and Technology*, 34, 1319–1350. <https://doi.org/10.1007/s11390-019-1978-9>
16. Kostiuk, Y., Skladannyi, P., Samoilenko, Y., Khorolska, K., Bebeshko, B., & Sokolov, V. (2025). A system for assessing the interdependencies of information system agents in information security risk management using cognitive maps. In *Cyber Hygiene & Conflict Management in Global Information Networks* (Vol. 3925, pp. 249–264).
17. Ullah, S. S., Hussain, S., Ali, I., & others. (2025). Mitigating content poisoning attacks in named data networking: A survey of recent solutions, limitations, challenges and future research directions. *Artificial Intelligence Review*, 58, 42. <https://doi.org/10.1007/s10462-024-10994-x>
18. Kostiuk, Y., Skladannyi, P., Samoilenko, Y., Khorolska, K., Bebeshko, B., & Sokolov, V. (2025). A system for assessing the interdependencies of information system agents in information security risk management using cognitive maps. In *Proceedings of the Third International Conference on Cyber Hygiene & Conflict Management in Global Information Networks (CH&CMiGIN'24)* (Vol. 3925, pp. 249–264).
19. Benmoussa, A., Kerrache, C. A., Lagraa, N., Mastorakis, S., Lakas, A., & Tahari, A. E. K. (2022). Interest flooding attacks in named data networking: Survey of existing solutions, open issues, requirements, and future directions. *ACM Computing Surveys*, 55(7), 1–37. <https://doi.org/10.1145/3539730>
20. Kostiuk, Y., Skladannyi, P., Sokolov, V., Hulak, H., & Korshun, N. (2025). Models and algorithms for analyzing information risks during the security audit of personal data information system. In *Proceedings of the Third International Conference on Cyber Hygiene & Conflict Management in Global Information Networks (CH&CMiGIN'24)* (Vol. 3925, pp. 155–171).
21. Lee, R. T., Leau, Y. B., Park, Y. J., & Anbar, M. (2022). A survey of interest flooding attack in named-data networking: Taxonomy, performance and future research challenges. *IETE Technical Review*, 39(5), 1027–1045. <https://doi.org/10.1080/02564602.2021.1957029>
22. Jeet, R., & Arun Raj Kumar, P. (2022). A survey on interest packet flooding attacks and its countermeasures in named data networking. *International Journal of Information Security*, 21, 1163–1187. <https://doi.org/10.1007/s10207-022-00591-w>
23. Mejri, S., Touati, H., & Kamoun, F. (2018). Hop-by-hop interest rate notification and adjustment in named data networks. In *2018 IEEE Wireless Communications and Networking Conference (WCNC)* (pp. 1–6). <https://doi.org/10.1109/WCNC.2018.8377374>
24. Kostiuk, Y., Skladannyi, P., Sokolov, V., Zhylytsov, O., & Ivanichenko, Y. (2025). Effectiveness of information security control using audit logs. In *Proceedings of the Workshop on Cybersecurity Providing in Information and Telecommunication Systems (CPITS 2025)* (pp. 524–538).
25. Nguyen, T., Hoang, D. T., Nguyen, D. N., & others. (2019). Reliable detection of interest flooding attack in real deployment of named data networking. *IEEE Transactions on Information Forensics and Security*, 14(9), 2470–2485. <https://doi.org/10.1109/TIFS.2019.2899247>
26. Compagno, A., Conti, M., Losiouk, E., Tsodik, G., & Valle, S. (2020). A proactive cache privacy attack on NDN. In *NOMS 2020 – 2020 IEEE/IFIP Network Operations and Management Symposium* (pp. 1–7). <https://doi.org/10.1109/NOMS47738.2020.9110318>



27. Vorokhob, M., Kyrychok, R., Yaskevych, V., Dobryshyn, Y., & Sydorenko, S. (2023). Modern Perspectives of Applying the Concept of Zero Trust in Building a Corporate Information Security Policy. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 1(21), 223–233. <https://doi.org/10.28925/2663-4023.2023.21.223233>
28. Tsekhmeister, R., Platonenko, A., Vorokhob, M., Cherevyk, V., & Semeniaka, S. (2025). Research of Information Security Provision Methods in a Virtual Environment. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 3(27), 63–71. <https://doi.org/10.28925/2663-4023.2025.27.703>
29. Kriuchkova, L., Skladannyi, P., & Vorokhob, M. (2023). Pre-project Solutions for Building an Authorization System Based on the Zero Trust Concept. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 3(19), 226–242. <https://doi.org/10.28925/2663-4023.2023.13.226242>
30. Brzhevska, Z., Kyrychok, R., Platonenko, A., & Hulak, H. (2022). Assessment of the Preconditions of Formation of the Methodology of Assessment of Information Reliability. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 3(15), 164–174. <https://doi.org/10.28925/2663-4023.2022.15.164174>
31. Skladannyi, P., et al. (2023). Improving the Security Policy of the Distance Learning System based on the Zero Trust Concept. In *Cybersecurity Providing in Information and Telecommunication Systems*, vol. 3421 (pp. 97–106).
32. Syrotynskyi, R., et al. (2024). Methodology of Network Infrastructure Analysis as Part of Migration to Zero-Trust Architecture. In *Cyber Security and Data Protection*, vol. 3800 (pp. 97–105).
33. Kostiuk, Yu. V., Skladannyi, P. M., Bebeshko, B. T., Khorolska, K. V., Rzaieva, S. L., & Vorokhob, M. V. (2025). *Information and communication systems security*. [Textbook] Kyiv: Borys Grinchenko Kyiv Metropolitan University.
34. Kostiuk, Yu. V., Skladannyi, P. M., Hulak, H. M., Bebeshko, B. T., Khorolska, K. V., & Rzaieva, S. L. (2025). *Information security systems*. [Textbook] Kyiv: Borys Grinchenko Kyiv Metropolitan University.
35. Hulak, H. M., Zhyltsov, O. B., Kyrychok, R. V., Korshun, N. V., & Skladannyi, P. M. (2023). *Enterprise information and cyber security*. [Textbook] Kyiv: Borys Grinchenko Kyiv Metropolitan University.



Yuliia Kostiuk

PhD in Computer Science

Associate Professor of the Department of Information and
Cyber Security named after Professor Volodymyr Buriachok
Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID: 0000-0001-5423-0985

y.kostiuk@kubg.edu.ua

Pavlo Skladannyi

PhD, Associate Professor, Head of the Department of Information and
Cyber Security named after Professor Volodymyr Buriachok

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID: 0000-0002-7775-6039

p.skladannyi@kubg.edu.ua

Svitlana Rzaeva

Candidate of Technical Sciences, Associate Professor,
Associate Professor of the Department of Computer Science

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID 0000-0002-7589-2045

s.rzaieva@kubg.edu.ua

Nataliia Mazur

PhD, Associate Professor

Associate Professor of the Department of Information and
Cyber Security named after Professor Volodymyr Buriachok

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID: 0000-0001-7671-8287

n.mazur@kubg.edu.ua

Vyacheslav Cherevyk

Candidate of Technical Sciences, Associate Professor,

Associate Professor of the Department of Information and

Cyber Security named after Professor Volodymyr Buriachok

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID: 0000-0002-2735-5341

v.cherevyk@kubg.edu.ua

Andriy Anosov

Candidate of Technical Sciences, Associate Professor,

Associate Professor of the Department of Information and

Cyber Security named after Professor Volodymyr Buriachok

Borys Grinchenko Kyiv Metropolitan University, Kyiv, Ukraine

ORCID ID: 0000-0002-2973-6033

a.anosov@kubg.edu.ua

FEATURES OF NETWORK ATTACK IMPLEMENTATION THROUGH TCP/IP PROTOCOLS

Abstract. This article investigates the implementation specifics of common network attacks that exploit vulnerabilities within the TCP/IP protocol stack — a critical infrastructural foundation of global network interaction. A comprehensive analysis is conducted on the architectural limitations and functional-protocol characteristics of key components of the network stack (ARP, IP, ICMP, TCP, UDP, DNS), which currently serve as primary vectors for the initiation of cyber threats. Based on the OSI reference model, a formalized classification of attacks by interaction layers is proposed, with emphasis on representative scenarios including IP spoofing, ARP poisoning, TCP session



hijacking, DNS cache poisoning, UDP flooding, and ICMP-based covert channels. Typical mechanisms for bypassing traditional security tools have been identified, including route manipulation, alteration of control messages, and encapsulation of malicious packets within legitimate traffic. Special attention is given to the overview of tools and proactive threat detection techniques, including intrusion detection systems (IDS), firewalls, deep packet inspection (DPI) technologies, as well as behavioral and entropy-based anomaly analysis methods in network flows. The findings provide both a theoretical foundation for modeling attacks and assessing risks, and a practical basis for enhancing information security in heterogeneous network environments.

Keywords: TCP/IP; network attacks; IP spoofing; ARP poisoning; cyber threats; attack vectors; information security; routing; network protection; CVE; SIEM; Explainable AI; behavioral analysis; Zero Trust.

REFERENCES (TRANSLATED AND TRANSLITERATED)

1. Hidouri, A., Hajlaoui, N., Touati, H., Hadded, M., & Muhlethaler, P. (2022). A survey on security attacks and intrusion detection mechanisms in named data networking. *Computers*, 11(12), 186. <https://doi.org/10.3390/computers11120186>
2. Pittman, J. M., Hoffpauir, K., Markle, N., & Meadows, C. (2020). A taxonomy for dynamic honeypot measures of effectiveness. *arXiv preprint*, arXiv:2005.12969.
3. Shah, Z., & Cosgrove, S. (2019). Mitigating ARP cache poisoning attack in software-defined networking (SDN): A survey. *Electronics*, 8(10), 1095. <https://doi.org/10.3390/electronics8101095>
4. Jin, Y., Tomoishi, M., & Matsuura, S. (2019). A detection method against DNS cache poisoning attacks using machine learning techniques: Work in progress. In *2019 IEEE 18th International Symposium on Network Computing and Applications (NCA)* (pp. 1–3). IEEE. <https://doi.org/10.1109/NCA.2019.8935025>
5. Yang, C. (2019). Anomaly network traffic detection algorithm based on information entropy measurement under the cloud computing environment. *Cluster Computing*, 22(Suppl 4), 8309–8317. <https://doi.org/10.1007/s10586-018-1755-5>
6. Dai, T., & Shulman, H. (2021). SMap: Internet-wide scanning for spoofing. In *Proceedings of the 37th Annual Computer Security Applications Conference* (pp. 1039–1050). <https://doi.org/10.1145/3485832.3485917>
7. Nosyk, Y., Korczyński, M., Lone, Q., Skwarek, M., Jonglez, B., & Duda, A. (2023). The Closed Resolver Project: Measuring the deployment of inbound source address validation. *IEEE/ACM Transactions on Networking*, 31(6), 2589–2603. <https://doi.org/10.1109/TNET.2023.3257413>
8. Tang, T. A., Mhamdi, L., McLernon, D., Zaidi, S. A. R., & Ghogho, M. (2016). Deep learning approach for network intrusion detection in software defined networking. In *2016 International Conference on Wireless Networks and Mobile Communications (WINCOM)* (pp. 258–263). <https://doi.org/10.1109/WINCOM.2016.7777224>
9. Krämer, L., Rossow, C., Bos, H., Monroe, F., & Blanc, G. (2015). AmpPot: Monitoring and defending against amplification DDoS attacks. In H. Bos, F. Monroe, & G. Blanc (Eds.), *Research in Attacks, Intrusions, and Defenses (RAID 2015)* (Vol. 9404, pp. 1–20). Springer. https://doi.org/10.1007/978-3-319-26362-5_28
10. Rossow, C. (2014, February). Amplification hell: Revisiting network protocols for DDoS abuse. In *NDSS* (pp. 1–15). <https://doi.org/10.14722/ndss.2014.23233>
11. Guo, J., He, L., & Liu, Y. (2024). Overlooked backdoors: Investigating 6to4 tunnel nodes and their exploitation in the wild. In *2024 IEEE International Performance, Computing, and Communications Conference (IPCCC)* (pp. 1–8). <https://doi.org/10.1109/IPCCC59868.2024.10850165>
12. Nosyk, Y., Korczyński, M., & Duda, A. (2023). Guardians of DNS integrity: A remote method for identifying DNSSEC validators across the internet. In *2023 IEEE 22nd International Conference on Trust, Security and Privacy in Computing and Communications (TrustCom)* (pp. 1470–1479). <https://doi.org/10.1109/TrustCom60117.2023.00201>
13. Kostiuk, Y., Skladannyi, P., Korshun, N., Bebesko, B., & Khorolska, K. (2024). Integrated protection strategies and adaptive resource distribution for secure video streaming over a Bluetooth network. *Information Technology*, 4(6), 14–33.
14. Tourani, R., Misra, S., Mick, T., & Panwar, G. (2018). Security, privacy, and access control in information-centric networking: A survey. *IEEE Communications Surveys & Tutorials*, 20(1), 566–600. <https://doi.org/10.1109/COMST.2017.2749508>



15. Kumar, N., Singh, A. K., Aleem, A., & others. (2019). Security attacks in named data networking: A review and research directions. *Journal of Computer Science and Technology*, 34, 1319–1350. <https://doi.org/10.1007/s11390-019-1978-9>
16. Kostiuk, Y., Skladannyi, P., Samoilenko, Y., Khorolska, K., Bebeshko, B., & Sokolov, V. (2025). A system for assessing the interdependencies of information system agents in information security risk management using cognitive maps. In *Cyber Hygiene & Conflict Management in Global Information Networks*, vol. 3925, pp. 249–264.
17. Ullah, S. S., Hussain, S., Ali, I., & others. (2025). Mitigating content poisoning attacks in named data networking: A survey of recent solutions, limitations, challenges and future research directions. *Artificial Intelligence Review*, 58, 42. <https://doi.org/10.1007/s10462-024-10994-x>
18. Kostiuk, Y., Skladannyi, P., Samoilenko, Y., Khorolska, K., Bebeshko, B., & Sokolov, V. (2025). A system for assessing the interdependencies of information system agents in information security risk management using cognitive maps. In *Proceedings of the Third International Conference on Cyber Hygiene & Conflict Management in Global Information Networks (CH&CMiGIN'24)* (Vol. 3925, pp. 249–264).
19. Benmoussa, A., Kerrache, C. A., Lagraa, N., Mastorakis, S., Lakas, A., & Tahari, A. E. K. (2022). Interest flooding attacks in named data networking: Survey of existing solutions, open issues, requirements, and future directions. *ACM Computing Surveys*, 55(7), 1–37. <https://doi.org/10.1145/3539730>
20. Kostiuk, Y., Skladannyi, P., Sokolov, V., Hulak, H., & Korshun, N. (2025). Models and algorithms for analyzing information risks during the security audit of personal data information system. In *Proceedings of the Third International Conference on Cyber Hygiene & Conflict Management in Global Information Networks (CH&CMiGIN'24)* (Vol. 3925, pp. 155–171).
21. Lee, R. T., Leau, Y. B., Park, Y. J., & Anbar, M. (2022). A survey of interest flooding attack in named-data networking: Taxonomy, performance and future research challenges. *IETE Technical Review*, 39(5), 1027–1045. <https://doi.org/10.1080/02564602.2021.1957029>
22. Jeet, R., & Arun Raj Kumar, P. (2022). A survey on interest packet flooding attacks and its countermeasures in named data networking. *International Journal of Information Security*, 21, 1163–1187. <https://doi.org/10.1007/s10207-022-00591-w>
23. Mejri, S., Touati, H., & Kamoun, F. (2018). Hop-by-hop interest rate notification and adjustment in named data networks. In *2018 IEEE Wireless Communications and Networking Conference (WCNC)* (pp. 1–6). <https://doi.org/10.1109/WCNC.2018.8377374>
24. Kostiuk, Y., Skladannyi, P., Sokolov, V., Zhylytsov, O., & Ivanichenko, Y. (2025). Effectiveness of information security control using audit logs. In *Proceedings of the Workshop on Cybersecurity Providing in Information and Telecommunication Systems (CPITS 2025)* (pp. 524–538).
25. Nguyen, T., Hoang, D. T., Nguyen, D. N., & others. (2019). Reliable detection of interest flooding attack in real deployment of named data networking. *IEEE Transactions on Information Forensics and Security*, 14(9), 2470–2485. <https://doi.org/10.1109/TIFS.2019.2899247>
26. Compagno, A., Conti, M., Losiouk, E., Tsudik, G., & Valle, S. (2020). A proactive cache privacy attack on NDN. In *NOMS 2020 – 2020 IEEE/IFIP Network Operations and Management Symposium* (pp. 1–7). <https://doi.org/10.1109/NOMS47738.2020.9110318>
27. Vorokhob, M., Kyrychok, R., Yaskevych, V., Dobryshyn, Y., & Sydorenko, S. (2023). Modern Perspectives of Applying the Concept of Zero Trust in Building a Corporate Information Security Policy. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 1(21), 223–233. <https://doi.org/10.28925/2663-4023.2023.21.223233>
28. Tsekhmeister, R., Platonenko, A., Vorokhob, M., Cherevyk, V., & Semeniaka, S. (2025). Research of Information Security Provision Methods in a Virtual Environment. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 3(27), 63–71. <https://doi.org/10.28925/2663-4023.2025.27.703>
29. Kriuchkova, L., Skladannyi, P., & Vorokhob, M. (2023). Pre-project Solutions for Building an Authorization System Based on the Zero Trust Concept. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 3(19), 226–242. <https://doi.org/10.28925/2663-4023.2023.13.226242>
30. Brzhevska, Z., Kyrychok, R., Platonenko, A., & Hulak, H. (2022). Assessment of the Preconditions of Formation of the Methodology of Assessment of Information Reliability. *Electronic Professional Scientific Journal «Cybersecurity: Education, Science, Technique»*, 3(15), 164–174. <https://doi.org/10.28925/2663-4023.2022.15.164174>
31. Skladannyi, P., et al. (2023). Improving the Security Policy of the Distance Learning System based on the Zero Trust Concept. In *Cybersecurity Providing in Information and Telecommunication Systems*, vol. 3421 (pp. 97–106).



32. Syrotynskyi, R., et al. (2024). Methodology of Network Infrastructure Analysis as Part of Migration to Zero-Trust Architecture. In *Cyber Security and Data Protection*, vol. 3800 (pp. 97–105).
33. Kostiuk, Yu. V., Skladannyi, P. M., Bebeshko, B. T., Khorolska, K. V., Rzaieva, S. L., & Vorokhob, M. V. (2025). *Information and communication systems security*. [Textbook] Kyiv: Borys Grinchenko Kyiv Metropolitan University.
34. Kostiuk, Yu. V., Skladannyi, P. M., Hulak, H. M., Bebeshko, B. T., Khorolska, K. V., & Rzaieva, S. L. (2025). *Information security systems*. [Textbook] Kyiv: Borys Grinchenko Kyiv Metropolitan University.
35. Hulak, H. M., Zhyltsov, O. B., Kyrychok, R. V., Korshun, N. V., & Skladannyi, P. M. (2023). *Enterprise information and cyber security*. [Textbook] Kyiv: Borys Grinchenko Kyiv Metropolitan University.



This work is licensed under Creative Commons Attribution-noncommercial-sharealike 4.0 International License.