

Федорченко В.М.

Харківський національний університет радіоелектроніки;

Харківський національний економічний університет імені Семена Кузнеця

Єрошенко О.А.

Харківський національний університет радіоелектроніки

ЗАСТОСУВАННЯ АЛГОРИТМІВ ШТУЧНОГО ІНТЕЛЕКТУ ДЛЯ МОДЕЛЮВАННЯ ЗАГРОЗ ІНФОРМАЦІЙНИХ СИСТЕМ

У статті розглянуто сучасні підходи до побудови моделей загроз інформаційних систем із використанням технологій штучного інтелекту. Проаналізовано основні напрями розвитку інтелектуальних методів у сфері кіберзахисту, зокрема застосування машинного та глибокого навчання, а також алгоритмів обробки природної мови для ідентифікації та прогнозування потенційних атак. Описано архітектуру та принципи роботи сучасних систем виявлення та запобігання вторгненням (IDS/IPS), зокрема Snort і Suricata, що забезпечують моніторинг мережевої активності в реальному часі. Визначено ключові особливості та переваги системи Snort, яка обрана як базова для реалізації запропонованої моделі через простоту налаштування, відкритий вихідний код і сумісність із нейронними мережами.

У роботі розроблено інтегровану модель захисту, яка поєднує можливості IDS Snort із нейронною мережею типу багатошарового перцептрона. Використано навчальний датасет NSL-KDD, що містить дані про мережеві з'єднання різних типів, включно з відомими видами атак. Проведено попередню обробку даних, нормалізацію та кодування ознак для ефективного навчання моделі. Нейронна мережа навчається розпізнавати шкідливу активність за характеристиками мережевого трафіку та класифікує події як нормальні або небезпечні. Для забезпечення стійкості моделі використано Dropout-шари, функцію втрат binary_crossentropy та оптимізатор Adam.

У статті наведено результати тестування моделі на реальних даних, що підтвердили її ефективність для виявлення як відомих, так і нових кіберзагроз. Запропоноване рішення може бути використане для створення адаптивних систем безпеки, здатних до самонавчання та інтеграції з існуючими платформами моніторингу. Отримані результати демонструють перспективність застосування штучного інтелекту у сфері кібербезпеки та підтверджують доцільність поєднання IDS Snort з нейронними мережами для побудови моделей загроз інформаційних систем.

Ключові слова: штучний інтелект, нейронна мережа, інформаційна системи, IDS/IPS, модель загрози, трафік, безпека.

Постановка проблеми. Сучасні інформаційні системи (ІС) дедалі глибше проникають у різні сфери діяльності, забезпечуючи автоматизацію процесів, обробку значних масивів даних та підтримку оперативного прийняття рішень. Водночас із розширенням функціональності та збільшенням інтенсивності інформаційних потоків посилюються загрози, що стосуються захисту даних та гарантування їхньої конфіденційності, цілісності й доступності. У таких умовах застосування технологій штучного інтелекту для формування моделей загроз в ІС набуває особливої актуальності, оскільки дає змогу своєчасно виявляти, прогнозувати та мінімізувати потенційні ризики.

Аналіз останніх досліджень і публікацій. Інформаційні системи сьогодні виступають ключовим елементом бізнес-процесів, управлінських механізмів та соціальних структур, забезпечуючи оброблення, накопичення й передачу значних масивів даних. Однак стрімкий технологічний прогрес і зростання складності ІС роблять їх усе більш вразливими та привабливими для кіберправопорушників. Постійне збільшення кількості кібератак, серед яких DDoS-атаки, фішинг, шкідливе програмне забезпечення та експлуатація вразливостей у ПЗ, формує додаткові ризики й ускладнює завдання забезпечення інформаційної безпеки. Попри те, що класичні інструменти кіберзахисту – такі як брандмауери, антивірусні системи та засоби виявлення

загроз – залишаються важливою складовою, їхніх можливостей вже недостатньо для протидії сучасним загрозам [1].

У таких умовах усе більшу роль відіграють інтелектуальні підходи до захисту, засновані на технологіях штучного інтелекту. Використання алгоритмів машинного та глибинного навчання, методів обробки природної мови та інших інструментів ШІ дає змогу створювати адаптивні системи, здатні ідентифікувати, попереджувати та прогнозувати широкий спектр кіберзагроз. Такі моделі загроз не лише виявляють відомі типи атак, а й дозволяють передбачати появу нових, забезпечуючи комплексний та гнучкий підхід до кібербезпеки.

З метою захисту інформаційних систем формується цілий комплекс заходів, спрямованих на забезпечення цілісності даних та стабільного функціонування інфраструктури. До нього належать організаційні компоненти, що охоплюють розроблення політик безпеки, навчання персоналу, управління доступом і регламентацію взаємодії користувачів з інформаційними ресурсами [2].

Технічний складник системи безпеки включає застосування міжмережевих екранів, інструментів для виявлення загроз, технологій шифрування, антивірусних рішень та інших механізмів, спрямованих на протидію спробам несанкціонованого доступу. Програмний аспект базується на регулярному оновленні ПЗ, усуненні вразливостей та постійному контролі мережевого трафіку з метою фіксації нетипової активності. Фізичний захист зосереджується на безпеці обладнання – зокрема, контролі доступу до серверних приміщень і запобіганні ризикам, пов'язаним із природними чинниками.

Одним із ключових елементів є створення резервних копій, що забезпечує можливість відновлення даних у разі технічних збоїв або дії кібератак. Безперервний моніторинг і аналіз роботи системи сприяють оперативному виявленню потенційних загроз, використовуючи сучасні інструменти автоматизації та штучного інтелекту. Важливою складовою залишається також своєчасне реагування на інциденти, ізоляція компрометованих компонентів і впровадження заходів для недопущення повторних атак [3–4].

Ефективний захист інформаційних систем можливий лише тоді, коли всі складові безпеки функціонують узгоджено та взаємодоповнюють одна одну, формуючи єдиний комплексний механізм. Саме для цього застосовуються спеціальні моделі загроз, які дають змогу системно аналізувати ризики.

Існує низка традиційних підходів до виявлення загроз в ІС, що широко застосовуються для формування базового рівня кіберзахисту та визначення потенційно небезпечних сценаріїв. Процес моделювання загроз рекомендовано виконувати вже на початкових етапах розроблення системи, що дозволяє своєчасно виявити слабкі місця та усунути їх до появи критичних наслідків, мінімізуючи витрати на ліквідацію результатів можливих атак.

Для таких завдань найчастіше використовують два відомі інструменти: OWASP Threat Dragon та Microsoft Threat Modeling Tool.

OWASP Threat Dragon – це засіб для побудови моделей загроз і створення відповідних діаграм у межах концепції безпечного життєвого циклу розробки програмного забезпечення. Інструмент ґрунтується на принципах та ідеях маніфесту з моделювання загроз (рисунок 1). Його функціонал дозволяє документувати потенційні загрози, визначати способи їх пом'якшення та відображати компоненти системи та поверхні атак у зручній візуальній формі. Threat Dragon може працювати як у вигляді вебзастосунку, так і як локальна настільна програма [5].

Microsoft Threat Modeling Tool – це інструмент моделювання загроз, який є основним елементом життєвого циклу розробки безпеки Microsoft (SDL). Це дозволяє архітекторам програмного забезпечення виявляти та пом'якшувати потенційні проблеми безпеки на ранній стадії, коли їх відносно легко та економічно ефективно вирішити. В результаті це значно знижує загальну вартість розробки. Крім того, було розроблено інструмент з урахуванням експертів, не пов'язаних із безпекою, що полегшує моделювання загроз для всіх розробників, надаючи чіткі вказівки щодо створення та аналізу моделей загроз (рисунок 2) [6].

Аналіз роботи програми дає змогу виокремити загальний механізм формування системи захисту. Він передбачає моделювання потенційних загроз, їхній детальний аналіз та подальшу розробку відповідних заходів протидії. Такий підхід забезпечує можливість заздалегідь сформулювати ефективні засоби захисту інформаційної системи та оперативно реагувати у разі виникнення реальної небезпеки.

Головною перевагою застосування технологій штучного інтелекту в сфері кібербезпеки є високий рівень автоматизації. На відміну від традиційних захисних механізмів, які потребують ручного налаштування параметрів і постійного втручання спеціаліста, системи на основі ШІ здатні само-

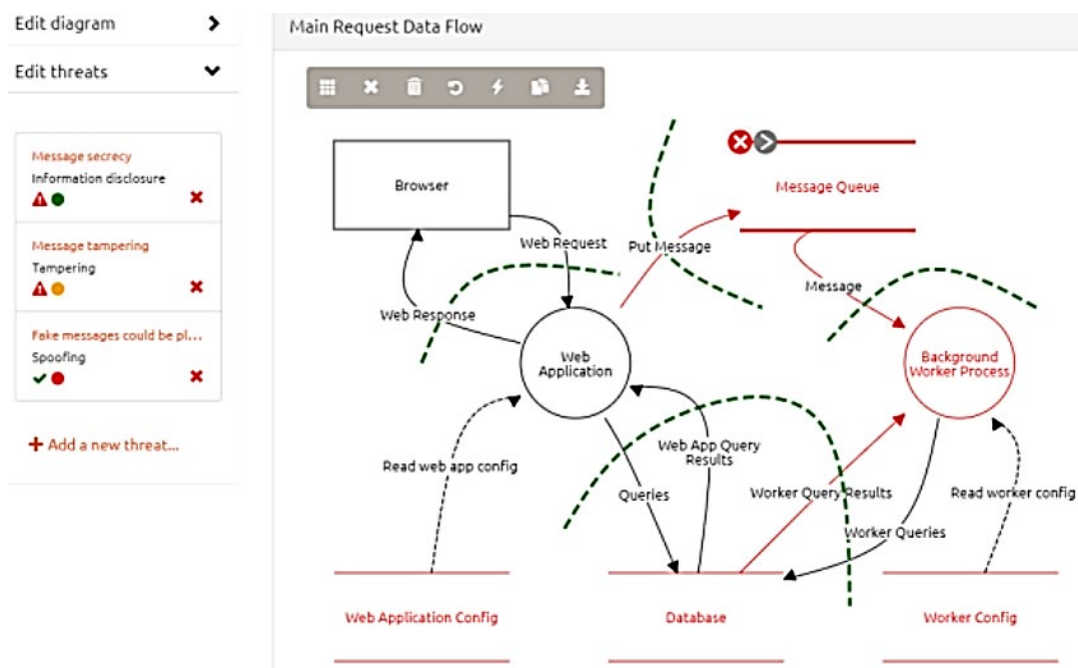


Рис. 1. Приклад побудованої діаграми загроз за допомогою програми OWASP Threat Dragon

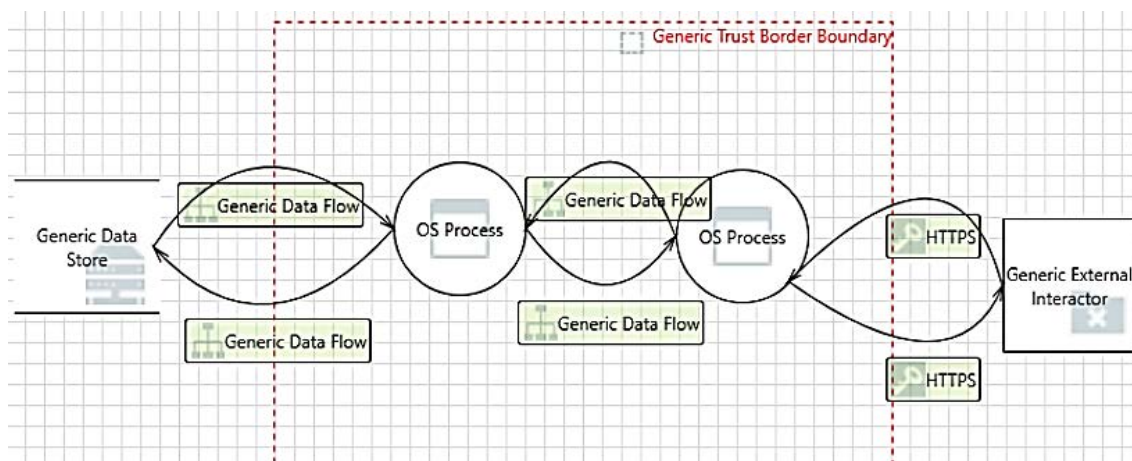


Рис. 2. Приклад побудованої діаграми загроз за допомогою утиліти Microsoft Threat Modeling Tool

стійно виконувати виявлення загроз. Це істотно зменшує залежність від людського фактора, мінімізує ризик помилок та підвищує загальну точність і ефективність захисту [7].

Крім того, нейронні мережі мають властивість самонавчання, що забезпечує їхню адаптивність до нових умов. Завдяки здатності аналізувати дані в режимі реального часу та навчатися на нових прикладах, інтелектуальні системи можуть ідентифікувати раніше невідомі типи атак та оперативно реагувати на них, забезпечуючи динамічний та гнучкий підхід до кібербезпеки.

На відміну від людини, ШІ, завдяки використанню алгоритмів машинного навчання, здатні

прогнозувати майбутні загрози, дозволяючи вжити необхідних заходів ще до настання потенційного ураження.

Також, якщо використовувати глибинні нейронні мережі, можна досягти високого рівня точності у визначенні аномальної поведінки та загроз, що забезпечує надійний захист ІС [8].

Однак використання штучного інтелекту для пошуку загроз має певні недоліки. Моделі сильно залежать від якості даних, на яких вони навчаються, тому можуть бути менш ефективними у виявленні рідкісних атак. Через складність нейронних мереж їх важко інтерпретувати, що ускладнює пояснення рішень системи. Також

можливі помилкові спрацьовування, які перевантажують операторів. ШІ вимагає значних обчислювальних ресурсів і регулярного оновлення моделей, щоб залишатися актуальним у змінних умовах. Крім того, його можуть обійти зловмисники, використовуючи спеціальні методи модифікації атак. Інтеграція таких рішень у наявні системи безпеки може бути складною, а аналіз мережевого трафіку іноді викликає етичні чи правові питання [9].

Таким чином, технології ШІ мають як суттєві недоліки, так і значні переваги над традиційними методами кібербезпеки та відіграють важливу роль у створенні комплексних систем захисту ІС, здатних забезпечувати високий рівень безпеки в умовах постійно зростаючих кіберзагроз.

Постановка завдання. Метою статті є аналіз існуючих підходів до побудови моделей загроз на основі ШІ та розробка власної моделі для підвищення рівня безпеки ІС.

Виклад основного матеріалу. Сучасні ІС вимагають застосування більш досконалих методів захисту, що зумовлено стрімким зростанням обсягів даних і зростаючою кількістю кібератак. Використання технологій штучного інтелекту стає одним із ключових підходів у формуванні моделей загроз та забезпеченні високого рівня інформаційної безпеки. Більшість сучасних методів виявлення та нейтралізації загроз ґрунтуються на машинному навчанні (ML), глибинному навчанні (DL) та технологіях обробки природної мови (NLP).

У практичній реалізації цих підходів важливу роль відіграють системи виявлення та запобігання вторгненням, призначені для безперервного контролю мережевого трафіку та виявлення потенційно шкідливої активності в режимі реального часу.

Одним із найпоширеніших інструментів цього класу є Snort – система з відкритим програмним кодом, що здійснює аналіз мережевих пакетів і виявлення ознак вторгнення. Snort працює на основі зіставлення трафіку з бібліотекою сигнатур атак і набором правил, що дає змогу фіксувати підозрілу поведінку в мережі. Вона підтримує декілька режимів роботи: простий перегляд трафіку, ведення журналів та сигнатурний аналіз. У поєднанні з іншими компонентами мережевої інфраструктури Snort може виконувати і функції системи запобігання вторгненням (IPS).

Ключовою складовою Snort є система правил, що визначає шаблони атак на основі IP-адрес, портів, протоколів та вмісту трафіку. Завдяки

широкій підтримці протоколів Snort здатний фіксувати інциденти у вигляді журналів або сповіщень, які можуть бути інтегровані з платформами керування подіями безпеки (SIEM), забезпечуючи комплексний контроль за станом кіберзахисту.

Suricata – це сучасна IDS/IPS-платформа з відкритим кодом, розроблена з урахуванням високої продуктивності та можливості масштабування. Завдяки багатопотоковій архітектурі система здатна ефективно використовувати ресурси багатоядерних процесорів, що забезпечує обробку значних обсягів мережевого трафіку без зниження швидкодії. Suricata підтримує глибокий аналіз пакетів, працює з широким спектром поширених протоколів, зокрема HTTP, DNS, TLS, FTP та SMB, а також повністю сумісна з правилковою базою Snort. Додатково система забезпечує розширені можливості аналізу завдяки використанню скриптів Lua та потоковій обробці даних, що дозволяє виявляти багатокрокові та складні типи атак.

Розвиток та оновлення Suricata здійснюються спільнотою Open Information Security Foundation, яка регулярно вдосконалює сигнатури й функціональні можливості системи.

На основі порівняльного аналізу для подальшого дослідження було обрано систему Snort, оскільки вона характеризується простотою налаштування, швидкою інтеграцією та меншими витратами часу на підготовку експериментального середовища. Крім того, Snort добре поєднується з нейронними мережами під час створення моделей безпеки інформаційних систем, забезпечуючи детальне журналювання мережевої активності для подальшого аналітичного опрацювання.

З огляду на проведений аналіз, для побудови моделі було вирішено застосувати багатоваріантну перцептронну нейронну мережу. Цей тип штучних нейронних мереж належить до технологій глибокого навчання й широко використовується для вирішення задач класифікації, регресії та інших прикладних завдань машинного навчання.

Для побудови даної моделі був обраний ноутбук із GPU компанії NVIDIA, а саме GeForce RTX 2050 із графічною пам'яттю 4 гігабайти та CPU компанії Intel, а саме Core i5-12450H із частотою 2 ГГц. Щодо операційної пам'яті було зроблено висновок, що 16 гігабайт для створення прототипу буде достатньо, але для створення повноцінної мережі для захисту буде потрібно більший обсяг. У якості носія інформації був використаний SSD накопичувач компанії KINGSTON із форм-фактором M.2 ємністю в 1 терабайт. У якості операційної системи була обрана Windows 10 pro.

Програмне забезпечення, що використовувалось під час роботи над проектом, містила в собі середовище Visual Studio Code із встановленим заздалегідь пакетом Python та бібліотеками TensorFlow та Keras. Саме цей набір був обраний тому що дозволив створити ШІ найшвидшими та відносно легшими методами на відміну від інших.

Для забезпечення можливості роботи нейронної мережі з реальними даними було розгорнуто систему виявлення та запобігання вторгненням Snort. Оскільки дане програмне забезпечення спочатку розроблене для Linux-середовища, виникла потреба встановити Npcap – сучасну версію WinPcap, що забезпечує підтримку необхідних бібліотек для роботи з мережевими пакетами у Windows. Snort було обрано завдяки його гнучкості в налаштуванні правил для аналізу трафіку та можливості формувати структуровані журнали подій, придатні для подальшого зчитування й опрацювання нейронною мережею.

У межах дослідження застосовано комбінований підхід до захисту інформаційної системи, який поєднує класичні способи виявлення вторгнень (IDS) із сучасними технологіями штучного інтелекту. Основне завдання полягає у створенні моделі, здатної аналізувати мережеву активність та ідентифікувати потенційні загрози у режимі реального часу. Методологія передбачає навчання нейронної мережі на історичних даних атак, інтеграцію отриманої моделі з IDS Snort і забезпечення автоматичного моніторингу безпекових подій.

Для навчання моделі використано датасет NSL-KDD, який є одним із найбільш відомих стандартних наборів даних у сфері кібербезпеки. Він містить приклади як нормальної мережевої активності, так і різних типів атак, і включає широкий спектр параметрів мережевих сесій: тип протоколу, кількість переданих байтів, TCP-прапори та інші характеристики. Структурованість даних робить їх зручними для використання у задачах машинного навчання.

Перед моделюванням було виконано попередню обробку даних: числові показники нормалізовано для уніфікації масштабів, а категоріальні характеристики (наприклад, тип протоколу) перетворено у числові представлення за допомогою методу one-hot encoding. Після цього дані поділено на тренувальний та тестовий набори для коректного оцінювання якості побудованої моделі.

Для здійснення класифікації мережевого трафіку було використано багатошаровий перцептрон. Цей різновид штучних нейронних мереж

добре підходить для розв'язання задач, де необхідно обробляти велику кількість вхідних параметрів. Архітектура моделі містила кілька прихованих шарів із активаційною функцією ReLU, що дало змогу враховувати нелінійні взаємозв'язки між ознаками. У вихідному шарі застосовано сигмоїдальну активацію, яка формує зручний формат для бінарного вирішення – чи є трафік атакою, чи він належить до нормального. Щоб зменшити ризик перенавчання, були використані Dropout-шари, які під час тренування випадково деактивують частину нейронів.

Навчання проводилося із застосуванням функції втрат `binary_crossentropy`, що є стандартною для задач бінарної класифікації. Для оптимізації параметрів моделі було обрано алгоритм Adam, який забезпечує швидку та стабільну збіжність. Після завершення тренування ефективність нейронної мережі перевіряли на тестовій вибірці, де розраховували точність, повноту, AUC (ROC-криву) та інші показники, що характеризують здатність системи коректно виявляти атаки.

Підготовлену модель інтегрували в систему Snort, яка виконувала функцію основного засобу контролю та аналізу мережевого трафіку. Snort зберігав інформацію про виявлені підозрілі дії у файлі журналу `alerts.txt`. Для автоматизованої обробки було створено Python-скрипт, який у режимі реального часу відстежував нові записи в цьому журналі, виділяв із них ключові ознаки та передавав їх до нейронної мережі. У випадку, коли модель ідентифікувала підозрілу активність як атаку, система формувала сповіщення, яке за потреби може супроводжуватися виконанням додаткових дій, наприклад надсиланням повідомлення адміністратору.

Такий підхід поєднує переваги Snort у роботі з реальним трафіком та можливості штучного інтелекту для підвищення точності виявлення загроз. Система здатна функціонувати автономно, аналізуючи нові події в режимі реального часу, та залишатися гнучкою для подальшого вдосконалення під нові типи атак.

Проектування моделі загроз на основі поєднання Snort та нейронної мережі здійснювалося поетапно, при цьому кожен етап відіграв важливу роль у забезпеченні коректної взаємодії між системою аналізу трафіку та алгоритмами штучного інтелекту.

Першим кроком стало налаштування середовища для роботи Snort. Було виконано завантаження та інсталяцію Snort на операційну систему Windows 10. Оскільки Snort спочатку орієнтована

```

18 #-----
19 # LOCAL RULES
20 #-----
21
22 alert tcp any any -> any any (msg:"High traffic detected"; dsize:>1500; sid:1000002; rev:1;)
23
24 alert icmp any any -> any any (msg:"ICMP Ping detected"; itype:8; sid:1000003; rev:1;)
25
26 alert tcp any any -> 192.168.1.10 80 (msg:"Access to forbidden site"; sid:1000004; rev:1;)
27
28 alert tcp any any -> any 23 (msg:"Telnet access detected"; sid:1000005; rev:1;)
29

```

Рис. 3. Індивідуальні правила фільтрації трафіку

ний на Linux-платформи, додатково встановлено бібліотеку NPcap, яка забезпечує можливість перехоплення та обробки мережесих пакетів у середовищі Windows.

Також до конфігурації були включені власні правила (рисунк 3), за допомогою яких можна відсіяти конкретні типи трафіку або зафіксувати нетипову активність. До прикладів таких дій належать сканування портів, спроби звернення до ресурсів із забороненим доступом та інші підозрілі запити.

Після завершення конфігурації Snort наступним кроком стало формування нейронної мережі, призначеної для аналізу потенційних загроз. Для створення моделі застосовували фреймворки TensorFlow та Keras. Неймережа отримувала на вхід набір мережесих параметрів, серед яких – IP-адреси, номери портів, типи протоколів та інші характеристики пакетів.

Структура моделі складалася з трьох основних шарів: двох прихованих шарів із активацією ReLU та механізмом Dropout для мінімізації ризику перенавчання, а також вихідного шару із сигмоїдальною активацією, що забезпечує визначення того, належить трафік до категорії загроз чи є нормальним (рисунк 4).

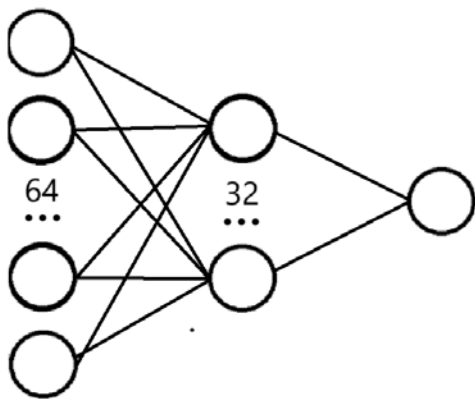


Рис. 4. Схема шарів розробленого ШІ

Процес навчання моделі здійснювався на датасеті NSL-KDD, який містить 123 записи, що описують як коректну активність у мережі, так і різні типи атак (рисунк 5). Перед подачею даних у мережу було виконано низку підготовчих кроків: нормалізацію числових параметрів, перетворення категоріальних полів у числовий формат та поділ набору даних на тренувальну й тестову частини.

Після завершення навчання отриману модель було збережено у форматі `attack_detection_model.h5`, що дозволяє надалі використовувати її для автоматичного визначення загроз у мережевому трафіку.

Інтеграція Snort із нейронною мережею була реалізована шляхом обробки журналів Snort у режимі реального часу. Завантажена модель перевіряє кожен новий запис журналу та визначає, чи містить він ознаки потенційної атаки. У випадку виявлення підозрілої активності система формує відповідне повідомлення, що забезпечує оперативне реагування на загрози (рисунк 6).

На фінальному етапі система проходила перевірку на реальних мережесих даних, що дало змогу оцінити її ефективність у виявленні як уже відомих, так і раніше невідомих типів атак. Приклади шкідливих подій, а також датасет NSL-KDD, були отримані з ресурсу Kaggle. На основі результатів тестування проводилося подальше вдосконалення як архітектури нейронної мережі, так і набору правил Snort, що дало змогу зменшити кількість хибних спрацьовувань.

У підсумку сформовано інтегровану модель загроз, яка поєднує класичні механізми аналізу трафіку, реалізовані у Snort, із можливостями сучасних методів штучного інтелекту. Такий підхід дозволив підвищити точність і швидкість виявлення загроз, забезпечуючи більш ефективний захист мережевої інфраструктури.

Завдяки використанню розширеного набору характеристик під час навчання модель загроз здатна одночасно аналізувати події як на транспортному,

Серед можливих напрямів удосконалення можна виокремити створення більш зручного інтерфейсу користувача, оскільки у поточній реалізації вся інформація подається переважно через консольні вікна та текстові файли.

Список літератури:

1. Яровий К.О., Гончар Л.В., Бабаян Д.П. Інформаційні системи і технології як невід’ємна частина в управлінні підприємством. *Інноваційна економіка*. 2021. № 7-8. С. 119–23. <https://doi.org/10.37332/2309-1533.2021.7-8.16>
2. Барковська О.Ю., Ні Я.С., Янковський О.А., Романенко А.О., Перетяка Є.О. Модель системи автоматизованого навантажувального тестування програмних застосунків із використанням методів штучного інтелекту. *Інформаційно-керуючі системи на залізничному транспорті*. 2025. № 1(30). С. 47–58. <https://doi.org/10.18664/iksz.v30i1.326699>
3. Fedorchenko V., Yeroshenko O., Shmatko O., Kolomiitsev O., Omarov M. Password hashing methods and algorithms on the .Net platform. *Advanced Information Systems*. 2024. № 8(4). Pp. 82–92. <https://doi.org/10.20998/2522-9052.2024.4.11>
4. Barkovska O., Ruban I., Tymoshenko D., Holovchenko O., Yankovskyi O. Research on mobile machine learning platforms for human gesture recognition in human-machine interaction systems. *Technology Audit and Production Reserves*. 2025. №2(2(82)). Pp. 6–14. <https://doi.org/10.15587/2706-5448.2025.325423>
5. Idris M., Syarif I., Winarno I. Web Application Security Education Platform Based on OWASP API Security Project. *EMITTER International Journal of Engineering Technology*. 2022. №10(2). Pp. 246–261. <https://doi.org/10.24003/emitter.v10i2.705>
6. Barkovska O. Двофакторна автентифікація на основі методу KWS та голосової верифікації. *Сучасний стан наукових досліджень та технологій в промисловості*. 2025. № 3(33). С. 5–18. <https://doi.org/10.30837/2522-9818.2025.3.005>
7. Choi W., Pandey S., Kim J. Detecting Cybersecurity Threats for Industrial Control Systems Using Machine Learning. *IEEE Access*. 2024. № 12. Pp. 153550–153563. <https://doi.org/10.1109/ACCESS.2024.3478830>
8. Єрошенко О.А. Особливості розпізнавання облич на зображеннях з використанням нейронних мереж. *Сучасні напрями розвитку інформаційно-комунікаційних технологій та засобів управління*. 2024. С. 96.
9. Михайліченко І.В., Ляшенко О.С. Модель прогнозування використання ресурсів у хмарних обчисленнях з використанням архітектури Informer. *Автоматизовані системи управління та прилади автоматизації*. 2025. № 186. С. 17–28. <https://doi.org/10.30837/0135-1710.2025.186.017>

Fedorchenko V.M., Yeroshenko O.A. APPLICATION OF ARTIFICIAL INTELLIGENCE ALGORITHMS FOR MODELING INFORMATION SYSTEM THREATS

The article examines modern approaches to building threat models for information systems using artificial intelligence technologies. The main directions of the development of intelligent methods in the field of cybersecurity are analyzed, including the application of machine learning, deep learning, and natural language processing algorithms for the identification and prediction of potential attacks. The architecture and operating principles of modern Intrusion Detection and Prevention Systems (IDS/IPS), such as Snort and Suricata, which provide real-time network activity monitoring, are described. The key features and advantages of the Snort system are identified, as it was selected as the basis for implementing the proposed model due to its ease of configuration, open-source nature, and compatibility with neural networks.

An integrated protection model combining the capabilities of the Snort IDS with a multilayer perceptron neural network has been developed. The NSL-KDD training dataset, containing data on various types of network connections, including known attack types, was used. Data preprocessing, normalization, and feature encoding were performed to ensure efficient model training. The neural network is trained to recognize malicious activity based on network traffic characteristics and classifies events as normal or dangerous. To improve model stability, Dropout layers, the binary crossentropy loss function, and the Adam optimizer were applied.

The paper presents testing results on real data, confirming the model's effectiveness in detecting both known and novel cyber threats. The proposed solution can be used to create adaptive, self-learning security systems capable of integrating with existing monitoring platforms. The obtained results demonstrate the potential of artificial intelligence in cybersecurity and confirm the feasibility of combining the Snort IDS with neural networks to build threat models for information systems.

Key words: artificial intelligence, neural network, information system, IDS/IPS, threat model, traffic, security.

Дата надходження статті: 14.11.2025

Дата прийняття статті: 02.12.2025

Опубліковано: 30.12.2025