

*Віталій ЗАЯЦЬ**Західноукраїнський національний університет***ІНТЕГРАЦІЯ ШТУЧНОГО ІНТЕЛЕКТУ В ПРОТОКОЛИ БЕЗПЕКИ
ІНТЕРНЕТУ РЕЧЕЙ**

Вступ. В епоху стрімкого розвитку цифрових технологій Інтернет речей (IoT) став невід'ємною частиною нашого повсякденного життя, трансформуючи способи взаємодії пристроїв та обробки даних. Однак із зростанням кількості підключених пристроїв та обсягів даних, що передаються між ними, традиційні методи забезпечення безпеки стають все менш ефективними. Кібератаки стають більш витонченими, а вразливості в IoT-системах можуть призвести до серйозних наслідків, від витоку конфіденційних даних до порушення роботи критичної інфраструктури.

Штучний інтелект (ШІ) пропонує революційний підхід до вирішення проблем безпеки в IoT-середовищі. Завдяки здатності обробляти величезні масиви даних у реальному часі, виявляти аномалії та адаптуватися до нових загроз, ШІ може значно підвищити рівень захисту IoT-систем. Інтеграція ШІ в протоколи безпеки IoT відкриває нові можливості для створення більш надійних та стійких до атак систем.

Мета: дослідження ефективності застосування машинного навчання для виявлення та реагування на кіберзагрози в системах IoT, а також аналіз етичних і правових аспектів інтеграції штучного інтелекту в безпеку IoT.

**1. Аналіз використання машинного навчання для виявлення та
реагування на кіберзагрози в системах IoT**

В сучасному світі IoT став невід'ємною частиною нашого життя, охоплюючи різноманітні сфери від розумних будинків до промислових систем управління. З кожним роком кількість підключених пристроїв стрімко зростає, створюючи складну екосистему взаємопов'язаних об'єктів. Однак разом із розширенням можливостей IoT зростають і ризики кібербезпеки. Традиційні методи захисту часто виявляються неефективними через обмежені обчислювальні ресурси IoT-пристроїв та специфіку їх роботи. В цьому контексті машинне навчання стає потужним інструментом для виявлення та протидії кіберзагрозам в системах IoT [1].

Основною проблемою безпеки IoT-систем є їх вразливість до різноманітних типів атак, включаючи DDoS-атаки, спроби несанкціонованого доступу, маніпуляції з даними та шкідливе програмне забезпечення [2]. Традиційні системи безпеки, засновані на сигнатурному аналізі та статичних правилах, не здатні ефективно виявляти нові види загроз та адаптуватися до змін у поведінці зловмисників. Саме тому застосування методів машинного навчання для аналізу поведінки мережі та виявлення аномалій стає все більш актуальним. Алгоритми машинного навчання здатні обробляти величезні масиви даних, виявляти приховані закономірності та автоматично адаптуватися до нових типів загроз. На рисунку 1 зображено типову схему DDoS-атаки.

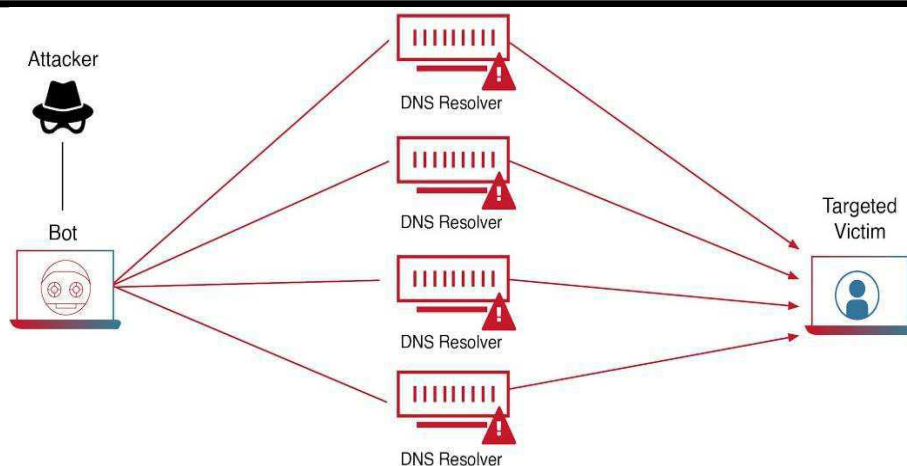


Рисунок 1 - Схема DDoS-атаки

В контексті IoT-безпеки особливо ефективними виявляються такі методи машинного навчання як глибокі нейронні мережі, алгоритми кластеризації та класифікації. Наприклад, згорткові нейронні мережі успішно застосовуються для аналізу мережевого трафіку та виявлення аномальної поведінки пристроїв. Алгоритми кластеризації допомагають групувати схожі патерни поведінки та виявляти відхилення від нормальної роботи системи. Методи класифікації, такі як випадкові ліси та методи опорних векторів, дозволяють з високою точністю категоризувати різні типи атак та визначати рівень загрози [3].

Важливим аспектом впровадження машинного навчання в системи безпеки IoT є збір та попередня обробка даних. IoT-пристрої генерують величезні обсяги різномірної інформації, включаючи показники сенсорів, логи подій, мережевий трафік та дані про стан системи. Для ефективного навчання моделей необхідно правильно відбирати та попередньо обробляти ці дані, видаляючи шуми та виділяючи найбільш інформативні ознаки. Крім того, важливо забезпечити баланс між точністю виявлення загроз та обчислювальними ресурсами, доступними на IoT-пристроях [4].

Одним з перспективних напрямків розвитку систем безпеки на основі машинного навчання є створення розподілених систем виявлення вторгнень. Такі системи дозволяють розподілити навантаження між різними вузлами мережі та забезпечити більш ефективний захист від розподілених атак. При цьому кожен вузол може мати свою локальну модель машинного навчання, яка обмінюється інформацією з іншими вузлами для покращення загальної ефективності системи. Це особливо важливо в контексті промислового Інтернету речей, де необхідно забезпечити безперервну роботу критично важливих систем [5].

Реагування на виявлені загрози також може бути автоматизовано за допомогою методів машинного навчання. Алгоритми можуть автоматично класифікувати рівень небезпеки та приймати рішення про необхідні заходи протидії. Наприклад, при виявленні підозрілої активності система може автоматично ізолювати скомпрометований пристрій, блокувати підозрілий трафік або змінювати параметри мережевої конфігурації. При цьому важливо забезпечити можливість людського контролю над автоматизованими рішеннями системи безпеки.

Однією з проблем використання машинного навчання в IoT-безпеці є

необхідність постійного оновлення моделей для адаптації до нових видів загроз. Зловмисники постійно розробляють нові методи атак, і системи безпеки повинні вміти виявляти раніше невідомі типи загроз. Для цього використовуються методи онлайн-навчання та трансферного навчання, які дозволяють моделям адаптуватися до змін у характері атак без необхідності повного перенавчання [6].

Важливим аспектом впровадження систем безпеки на основі машинного навчання є забезпечення конфіденційності даних. IoT-пристрої часто обробляють чутливу інформацію, і необхідно гарантувати, що системи безпеки не створюють додаткових ризиків витоку даних. Для цього використовуються методи федеративного навчання, які дозволяють навчати моделі без необхідності централізованого збору даних, а також методи диференційної приватності, що забезпечують захист конфіденційної інформації при навчанні моделей.

Також варто відзначити важливість правильної валідації та тестування систем безпеки на основі машинного навчання. Необхідно проводити регулярне тестування на стійкість до різних типів атак, оцінювати точність виявлення загроз та аналізувати можливі помилкові спрацьовування. Особливу увагу слід приділяти тестуванню на стійкість до атак на самі моделі машинного навчання, таких як отруєння даних навчання або обхід системи виявлення.

В майбутньому можна очікувати подальшого розвитку методів машинного навчання для забезпечення безпеки IoT-систем. Перспективними напрямками є розробка більш ефективних алгоритмів для роботи з обмеженими ресурсами, створення гібридних систем, що поєднують різні методи машинного навчання, та вдосконалення механізмів автоматичного реагування на загрози. Також важливим напрямком є стандартизація підходів до використання машинного навчання в IoT-безпеці та розробка відповідних рекомендацій та best practices.

2. Етичні та правові аспекти впровадження штучного інтелекту в системи безпеки IoT

Впровадження ШІ в системи безпеки IoT створює не лише технологічні виклики, але й піднімає важливі етичні та правові питання, які потребують ретельного розгляду та аналізу. В епоху, коли розумні пристрої стають невід'ємною частиною нашого повсякденного життя, забезпечення їхньої безпеки за допомогою ШІ має відбуватися з урахуванням фундаментальних прав людини, етичних норм та законодавчих вимог. Системи безпеки на основі ШІ мають унікальні можливості для захисту IoT-пристроїв та мереж, але їх впровадження повинно відбуватися відповідально та прозоро, з належною увагою до потенційних ризиків та наслідків для суспільства [7].

Одним з ключових етичних питань при використанні ШІ в системах безпеки IoT є забезпечення приватності користувачів. Розумні пристрої збирають величезну кількість персональних даних, включаючи інформацію про місцезнаходження, звички, стан здоров'я та інші аспекти приватного життя. Системи безпеки на основі ШІ, аналізуючи ці дані для виявлення потенційних загроз, повинні забезпечувати належний рівень захисту конфіденційної інформації. Це включає не лише технічні аспекти захисту даних, але й етичні принципи щодо збору, обробки та зберігання інформації. Важливо забезпечити прозорість щодо того, які дані збираються, як вони використовуються та хто має

до них доступ. Користувачі повинні мати можливість контролювати свої персональні дані та розуміти, як системи ШІ використовують цю інформацію для забезпечення безпеки [8].

Справедливість та недискримінація є іншим важливим етичним аспектом впровадження ШІ в системи безпеки IoT. Алгоритми машинного навчання, які використовуються для виявлення загроз та аномальної поведінки, можуть ненавмисно відтворювати або посилювати існуючі упередження. Наприклад, система безпеки може помилково класифікувати певні групи користувачів як потенційно небезпечні на основі історичних даних, які містять упереджену інформацію. Тому при розробці та впровадженні таких систем необхідно приділяти особливу увагу забезпеченню справедливості та рівного ставлення до всіх користувачів. Це включає регулярний аудит алгоритмів на наявність упереджень, тестування на різних групах користувачів та впровадження механізмів корекції виявлених проблем.

Правове регулювання використання ШІ в системах безпеки IoT також створює значні виклики. Законодавство в багатьох країнах не встигає за швидким розвитком технологій, створюючи правові прогалини та невизначеності. Особливо складними є питання відповідальності за помилки або збої в роботі систем ШІ. Хто несе відповідальність, якщо система безпеки на основі ШІ помилково блокує легітимний пристрій або, навпаки, не виявляє реальну загрозу? Ці питання потребують чіткого правового регулювання та встановлення механізмів відповідальності. Крім того, важливо забезпечити відповідність систем безпеки існуючим законам про захист персональних даних, таким як GDPR в Європейському Союзі, та галузевим стандартам безпеки. На рисунку 2 зображено блок-схему застосування GDPR.

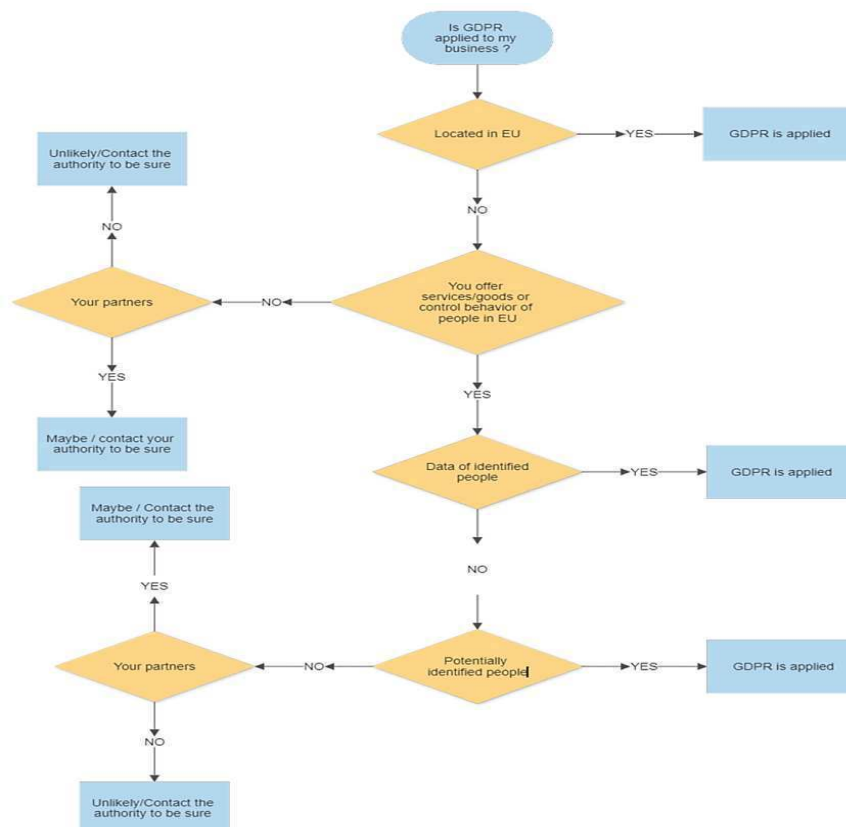


Рисунок 2 - Блок-схема застосування GDPR

Прозорість та підзвітність систем ШІ є критично важливими аспектами їх впровадження в системи безпеки IoT. Користувачі та зацікавлені сторони повинні розуміти, як приймаються рішення щодо безпеки, які критерії використовуються для виявлення загроз та які механізми існують для оскарження автоматизованих рішень. Це особливо важливо в контексті промислових IoT-систем, де помилкові спрацьовування систем безпеки можуть призвести до значних економічних втрат. Необхідно розробити механізми аудиту та верифікації рішень ШІ, а також забезпечити можливість людського нагляду за критично важливими рішеннями щодо безпеки.

Важливим аспектом є також етичне використання даних при навчанні моделей ШІ для систем безпеки. Збір та використання даних для навчання алгоритмів повинні відбуватися з дотриманням принципів інформованої згоди та поваги до приватності. Необхідно забезпечити належне знеособлення даних та захист від можливості реідентифікації особи. Крім того, важливо враховувати культурні та соціальні особливості різних регіонів при розробці та впровадженні систем безпеки на основі ШІ. Те, що вважається прийнятним в одному культурному контексті, може бути неприйнятним в іншому.

Окремої уваги заслуговує питання взаємодії людини та ШІ в системах безпеки IoT. Важливо знайти правильний баланс між автоматизацією та людським контролем. Повна автоматизація процесів безпеки може призвести до втрати критичного мислення та здатності реагувати на нестандартні ситуації. З іншого боку, надмірне покладання на людський контроль може знизити ефективність систем безпеки. Необхідно розробити ефективні моделі взаємодії, де ШІ доповнює та підсилює людські можливості, а не замінює їх повністю.

Стійкість та надійність систем безпеки на основі ШІ також мають важливе етичне значення. Системи повинні бути стійкими до спроб маніпуляції та атак на самі алгоритми ШІ. Це включає захист від спроб отруєння навчальних даних, протидію спробам обману системи та забезпечення стабільної роботи в умовах невизначеності. Також важливо забезпечити можливість оновлення та адаптації систем безпеки до нових загроз без порушення етичних принципів та правових норм.

Глобальний характер IoT-систем створює додаткові правові та етичні виклики. Різні країни мають різні підходи до регулювання ШІ та захисту даних, що може створювати складнощі при впровадженні глобальних систем безпеки. Необхідна міжнародна співпраця та гармонізація підходів до регулювання використання ШІ в системах безпеки IoT. Це включає розробку міжнародних стандартів, обмін найкращими практиками та створення механізмів транскордонного співробітництва у сфері кібербезпеки.

Дивлячись у майбутнє, можна прогнозувати, що етичні та правові аспекти використання ШІ в системах безпеки IoT будуть набувати все більшого значення. З розвитком технологій та збільшенням кількості підключених пристроїв зростатиме потреба в більш досконалих механізмах регулювання та етичних framework'ах. Важливо, щоб розвиток технологій відбувався паралельно з розвитком етичних норм та правового регулювання, забезпечуючи баланс між інноваціями та захистом прав та інтересів усіх зацікавлених сторін.

Висновок. Впровадження штучного інтелекту в системи безпеки Інтернету речей знаменує собою новий етап у розвитку цифрової безпеки, що характеризується переходом від реактивних до проактивних методів захисту. Використання методів машинного навчання для виявлення та реагування на кіберзагрози демонструє значний потенціал у підвищенні ефективності захисту IoT-систем, дозволяючи автоматично виявляти аномалії, передбачати потенційні атаки та адаптуватися до нових видів загроз. Проте технологічний прогрес у цій сфері нерозривно пов'язаний з необхідністю вирішення складних етичних дилем та правових питань.

Успішна інтеграція ШІ в системи безпеки IoT вимагає комплексного підходу, що враховує не лише технічні аспекти, але й забезпечує баланс між ефективністю захисту та повагою до приватності користувачів, дотриманням етичних норм та відповідністю правовим вимогам. Особливо важливим є забезпечення прозорості роботи алгоритмів, справедливості прийняття рішень та захисту персональних даних. При цьому необхідно враховувати глобальний характер IoT-систем та різноманітність правових та культурних контекстів їх використання. Подальший розвиток цього напрямку потребує тісної співпраці між технічними спеціалістами, етиками, юристами та представниками регуляторних органів для створення збалансованих рішень, які забезпечують високий рівень безпеки без компромісу щодо фундаментальних прав та свобод користувачів.

Перелік використаних джерел.

1. Домарев В.С. Проблеми та перспективи забезпечення кібербезпеки в умовах розвитку Інтернету речей // Вісник Київського національного університету імені Тараса Шевченка. - 2021. - № 3. - С. 14-22.
2. Божок Н.М., Коваленко О.В. Використання технологій машинного навчання для виявлення кіберзагроз у системах IoT // Кібербезпека та інформаційні технології. - 2022. - № 1. - С. 34-41.
3. Сергієнко І.В., Попов С.А. Методи машинного навчання для виявлення аномалій у трафіку Інтернету речей // Наукові записки Національного університету Києво-Могилянська академія. - 2020. - № 5. - С. 18-25.
4. Коваленко В.С., Іванова М. П. Використання методів штучного інтелекту для забезпечення інформаційної безпеки в IoT // Системи управління, навігації та зв'язку. - 2022. - № 2. - С. 8-16.
5. Лавренюк О.В., Романчук Ю.П. Розподілені системи виявлення вторгнень у промисловому IoT: новітні підходи та виклики // Проблеми інформаційних технологій. - 2023. - № 1. - С. 29-36.
6. Ткачук В.М., Міщенко П.І. Трансферне навчання для адаптації моделей машинного навчання в системах кібербезпеки IoT // Журнал наукових досліджень. - 2022. - № 6. - С. 52-59.
8. Григорчук О.С. Етичні та правові аспекти впровадження штучного інтелекту в системи кібербезпеки // Інформаційна безпека України. - 2023. - № 4. - С. 12-19.
9. Романюк Т.М., Остапенко О.М. Приватність і конфіденційність у системах безпеки IoT на основі ШІ // Наукові праці Національного авіаційного університету. - 2023. - № 7. - С. 45-53.