

МЕТОДИ ТА МОДЕЛІ ЗАХИСТУ ДАНИХ ІОТ В УМОВАХ ІНФОРМАЦІЙНОЇ НЕВИЗНАЧЕНОСТІ НА ОСНОВІ ФУНКЦІЙ ШТУЧНОГО ІНТЕЛЕКТУ

Анотація: Швидкий розвиток Інтернету речей (IoT) створює нові виклики в галузі безпеки даних. У роботі досліджуються методи та моделі захисту даних IoT на базі штучного інтелекту (ШІ) для виявлення загроз у реальному часі, адаптивної оцінки ризиків та прийняття рішень. Запропоновано ШІ-структуру, що підвищує стійкість та автономність захисту у IoT-системах.

Ключові слова: IoT, екосистема, інформаційна невизначеність, мережева екосистема, безпека даних, штучний інтелект, нейронні мережі.

Вступ

У сучасних умовах стрімкого розвитку цифрових технологій та розширення масштабів функціонування інформаційних систем (ІС) зростає потреба в їх здатності ефективно опрацьовувати великі обсяги даних, генерованих у режимі реального часу, з різних джерел, часто неоднорідних за структурою, якістю та рівнем достовірності. Одним з ключових викликів, що постає перед інформаційними системами, є інформаційна невизначеність, яка все частіше виявляється як критичний фактор, що обмежує ефективність, надійність і безпеку функціонування таких систем.

Під інформаційною невизначеністю у цьому контексті розуміється стан, за якого доступна інформація є неповною, неточною, суперечливою або неоднозначною, що ускладнює прийняття обґрунтованих рішень як людиною, так і автоматизованими агентами. Джерелами такої невизначеності можуть бути:

- 1) Технічна несправність обладнання;
- 2) Відсутність або недостатність контексту;
- 3) Часові обмеження на обробку;
- 4) Відсутність довіри до джерел інформації;
- 5) Пошкоджені або застарілі дані та ін.

Ця проблема є особливо гострою в умовах інформаційного перевантаження, високої динаміки середовища, а також функціонування систем в умовах кіберзагроз чи соціотехнічних факторів.

Наукова й практична її значущість визначається наступними умовами:

- 1) підвищення точності автоматизованих рішень у критичних сферах (охорона здоров'я, енергетика, фінанси, оборона);
- 2) забезпечення прозорості та пояснюваності ІС;
- 3) формування принципів розробки гнучких, адаптивних, самонавчальних систем, здатних функціонувати в умовах нестачі або надлишку інформації;
- 4) інтеграції різномірних джерел даних із різним ступенем надійності.

Аналіз останніх досліджень та публікацій

Під час роботи над написанням статті, було оглянуто ряд джерел, які підіймали тему безпеки та цілісності даних в динамічних екосистемах. Зокрема аналіз публікацій показав наступне, хоч питання цілісності та повноти даних активно підіймається, при цьому дуже рідко враховується інформаційна невизначеність як загроза мережі [1-2], а також, хоч використання адаптивних моделей та методів захисту на основі функцій ШІ періодично пропонується [1-5], пропозиції та способу їх використання для усунення інформаційної невизначеності або зменшення її впливу не було. Окрім цього, не було знайдено вирішення наступних проблем:

1. Недостатній рівень адаптивності захисту мережі;
2. Низький рівень інтерпретованості моделей ШІ, що обмежує можливість використання їх у критичних інфраструктурах із вимогами до аудиту прийнятих рішень;
3. Недостатній рівень уваги інформаційній невизначеності як джерела загроз мережі. Питання когнітивної перевантаженості мережі, помилкової інтерпретації сигналів або навмисних дій залишаються актуальними;
4. Низький рівень готовності моделей ШІ до протидії сучасним – динамічним загрозам.

Таким чином мета даної статті полягає в наступному:

1. Провести аналіз загроз інформаційної невизначеності інформаційним системам;
2. Визначити причини її виникнення та яку небезпеку вона становить інформаційній системі;
3. Провести аналіз існуючих моделей та методів використання ШІ що можуть бути використані або взяті за основу для усунення ІН.

Вплив інформаційної невизначеності

Інформаційна невизначеність становить значну і зростаючу загрозу стабільності, ефективності та безпеці сучасних мережевих систем на основі ІоТ. Її можна визначити як виникнення відхилень, спотворень, втрат або пошкоджень даних на етапах їх створення, передачі або обробки. Це явище, яке в цій статті називають «інформаційною невизначеністю», є серйозною перешкодою для точного виявлення, пом'якшення та нейтралізації нових загроз безпеці в складних цифрових екосистемах. До таких факторів можна віднести:

- 1) Фрагментація даних: часткова доступність телеметрії, журналів або інформації про події;
- 2) Шуми: наявність великої кількості подій, які не є інцидентами, але ускладнюють аналіз;
- 3) Динаміка оточення: Часта зміна конфігурацій, політик доступу, топологій мережі;

4) Атаки нульового дня: Невідомі (раніше не зареєстровані) вразливості, які не мають сигнатур;

5) Невидимий супротивник: використання методів ухилення від виявлення, таких як шифрування, обфускація, стелс-атаки;

Архітектура середовищ IoT характеризується величезною кількістю взаємопов'язаних, різномірних пристроїв, що робить їх особливо вразливими до впливу інформаційної невизначеності. Навіть незначні невідповідності даних можуть швидко поширюватися, що призводить до помилкової поведінки системи, зниження продуктивності або повного збою в роботі. Крім того, природа мереж IoT, що постійно розвивається, посилює цей ризик оскільки такі фактори як збої в апаратному забезпеченні, нестабільність мережі, аномалії програмного забезпечення та зловмисна діяльність сприяють підвищеній невизначеності.

Екосистема IoT

Екосистеми IoT, за задумом, створюються з можливістю розширення, що включає як горизонтальне (додавання нових пристроїв), так і вертикальне (додавання нових функцій і можливостей) масштабування (рис. 1) [1]. Крім того, однією з його важливих переваг є здатність виконувати крос – функціональну інтеграцію. Хоча це робить її універсальною, ця здатність також створює потенційно слабкі місця, проблеми для реалізації та серйозні загрози безпеці.

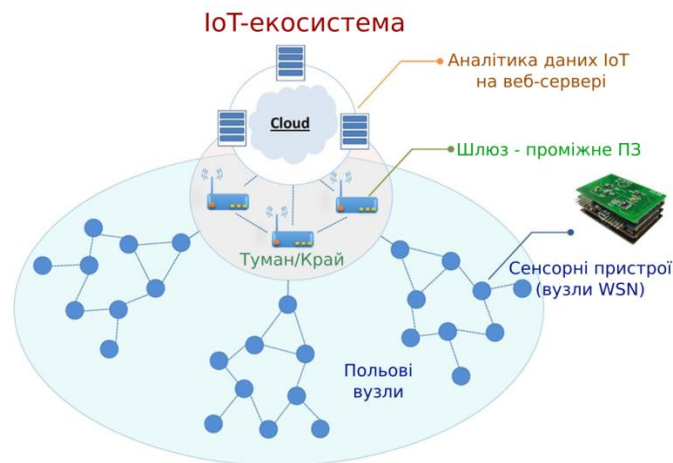


Рисунок 1. Візуальний приклад екосистеми IoT

Під час процесу розширення має бути суворий контроль, щоб забезпечити безперебійну інтеграцію нових пристроїв на всіх рівнях екосистеми. До них відносяться (рис. 2) [5]:

- 1) Шар застосунку;
- 2) Шар обробки даних;
- 3) Мережевий шар;
- 4) Сенсорний шар.

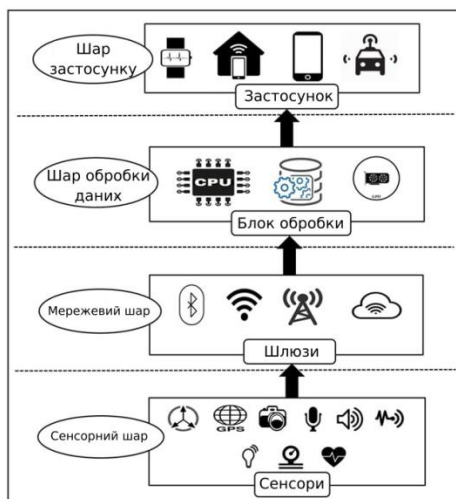


Рисунок 2. Шари та компоненти архітектури IoT

Динамічне розширення створює динамічне середовище, в якому дані часто можуть:

- 1) Неповні;
- 2) Несумісні;
- 3) Застарілі;
- 4) Потенційно піддавались маніпуляціям.

Причина його виникнення може бути різною, від некоректного маніпулювання або несанкціоноване втручання у програмне забезпечення, до втручання в роботу її обладнання та інтеграція старого або нового, несумісності з системою в цілому. Існує безліч потенційних причин для того, щоб це сталося, але наразі найважливішим є проблема, яку воно викликає, і той факт, що його можна вирішити.

Інформаційна невизначеність

Вплив інформаційної невизначеності має подвійний характер. По-перше, це ускладнює точну ідентифікацію інцидентів безпеки, маскуючи справжні аномалії на тлі непередбачуваної поведінки системи. По-друге, це створює нові слабкі місця, якими зловмисники можуть скористатися, використовуючи невизначеність, щоб обійти встановлені засоби захисту, ввести неправдиві дані або спровокувати збої на рівні системи.

Традиційні системи безпеки, які переважно покладаються на статичні набори правил і заздалегідь визначені сигнатури загроз, стають все більш неефективними у вирішенні динамічних проблем, пов'язаних з інформаційною невизначеністю. Ці методи, як правило, є реакційними, а не прогностичними, і їм важко пристосуватися до швидко мінливого ландшафту загроз, характерного для сучасних систем IoT, які також частіше не вимагають здатності адаптуватися[2][3].

Наслідки некерованої інформаційної невизначеності є глибокими, особливо в таких критично важливих секторах, як кібербезпека, охорона здоров'я, транспорт,

енергетична інфраструктура та національна оборона. У цих сферах цілісність даних безпосередньо пов'язана з громадською безпекою, економічною безпекою та операційною ефективністю.

Щоб краще зрозуміти його небезпеку та потенційні ризики, наведемо приклади надзвичайних ситуацій, спричинених ним у кожному зі згаданих секторів:

1) Транспортний сектор, збій у роботі системи FAA NOTAM (США, січень 2023 р.):

Система FAA Notice to Air Missions (NOTAM), яка надає пілотам інформацію про небезпеку польоту в режимі реального часу, вийшла з ладу через пошкоджений файл у критично важливій базі даних. Через це пілоти та авіадиспетчери не могли довіряти наявним навігаційним сповіщенням, що призводило до масових затримок. Понад 11000 рейсів було затримано та 1300 скасовано. Непослідовна та потенційно неточна передача даних призвела до *невизначеного середовища безпеки польотів*, що призвело до приземлення літаків[6].

2) Сектор кібербезпеки, операція «Тріангуляція: iOS Zero – Click Exploit» (2023):

Складна кібератака, що отримала назву «*Операція Тріангуляція*», була націлена на пристрої iOS за допомогою ланцюжка з чотирьох вразливостей нульового дня, доставлених через невидимі повідомлення iMessage. Атака використовувала незадокументовані апаратні функції та вразливості програмного забезпечення, щоб отримати підвищені привілеї, дозволяючи зловмисникам отримувати повідомлення, паролі та дані геолокації без втручання користувача. Прихований характер атаки створив невизначеність у виявленні та приписуванні порушення. В результаті тисячі пристроїв, у тому числі комерційних, урядових та дипломатичних організацій, були скомпрометовані, що призвело до значних проблем зі шпигунством[7][8].

3) Сектор охорони здоров'я, атака програм – вимагачів Change Healthcare (лютий 2024 року):

У деяких випадках інформаційна невизначеність є побічним продуктом кібератак, проте її наслідки можуть бути не менш руйнівними та небезпечними. Наприклад, у лютому 2024 року, після того, як група програм-вимагачів «The ALPHV/BlackCat» націлилася на Change Healthcare, велику американську компанію з виставлення рахунків у сфері медичних послуг. Атака порушила доступ до записів пацієнтів і платіжних систем, що призвело до плутанини щодо права пацієнтів на отримання допомоги та обробки платежів. Програма – вимагач шифрувала критично важливу інформацію та порушувала обмін даними між постачальниками, страховиками та пацієнтами, що призвело до порушення довіри та доступності даних у мережі. Як наслідок, дані понад 100 мільйонів пацієнтів були скомпрометовані, що спричинило значні затримки в наданні медичної допомоги та фінансове навантаження на постачальників медичних послуг[9][10].

4) Енергетичний сектор, атака на данські енергомережі (травень 2023 року):

Атака на данську енергетичну мережу у травні 2023 року є наймасштабнішою кібератакою на критично важливу інфраструктуру Данії на сьогоднішній день. Протягом кількох тижнів 22 організації енергетичного сектору були скомпрометовані

в ході добре скоординованої кампанії, яка використовувала вразливості в широко використовуваних мережевих пристроях. Зловмисники скористалися вразливістю брандмауера Zyxel, після чого оператори не змогли відрізнити законні системні команди від шкідливих мережевих інструкцій. Системи моніторингу відображали суперечливі або оманливі дані. Як наслідок, статус мережі став неоднозначним, і групи реагування не були впевнені, чи системи працюють безпечно, чи були скомпрометовані. Скомпрометувавши шляхи зв'язку (наприклад, маршрутизатори/брандмауери), зловмисники ефективно порушили надійні потоки даних, ще більше заплутавши інструменти моніторингу та приховавши ведення журналів[11].

5) Національна оборона, *Zombie Zero*: шкідливе програмне забезпечення, вбудоване в пристрої IoT:

Під час цієї складної кібератаки сканери штрих-кодів, виготовлені за кордоном, були вбудовані зі шкідливим програмним забезпеченням під час виробництва. Як тільки ці скомпрометовані пристрої були інтегровані в корпоративні мережі, вони ініціювали несанкціонований зв'язок із зовнішніми серверами, що призвело до викрадення даних[12].

– Вразливості ланцюгів поставок: Присутність шкідливого програмного забезпечення залишилася непоміченою під час стандартних перевірок безпеки, що підкреслює невизначеність у цілісності ланцюжків поставок і складність перевірки безпеки сторонніх апаратних компонентів.

– Експлуатація довіри до мережі: скомпрометовані пристрої були довіреними компонентами в безпечних мережах. Їхня зловмисна поведінка внесла невпевненість у надійність внутрішніх мережевих пристроїв, ускладнивши зусилля з виявлення загроз і реагування на них.

– Цілісність даних та ризики конфіденційності: Несанкціонована передача даних з цих пристроїв викликала занепокоєння щодо конфіденційності та цілісності конфіденційної інформації, оскільки було незрозуміло, до яких даних було отримано доступ або передано дані.

Тому розробка та впровадження динамічних, інтелектуальних та адаптивних механізмів безпеки є обов'язковими. Штучний інтелект з його можливостями для аналізу в реальному часі, виявлення аномалій і прогнозного моделювання є ключовою технологією в боротьбі з негативними наслідками інформаційної невизначеності.

Питання безпеки даних та обладнання

Безпека як даних, так і компонентів обладнання є фундаментальним викликом при розгортанні екосистем IoT, особливо в умовах, що характеризуються високим рівнем інформаційної невизначеності. Ці середовища, які покладаються на безперебійну взаємодію в реальному часі між взаємопов'язаними пристроями та

центральними пристроями обробки даних, за своєю суттю вразливі до збоїв, які можуть мати каскадний вплив на цілісність системи та фізичну продуктивність.

Значний ризик для його цілісності виникає через потенційну можливість зловмисних кібератак порушити роботу або маніпулювати потоками даних. Такі вторгнення не тільки ставлять під загрозу конфіденційність і точність переданої інформації, але й здатні примусово або частково деактивувати систему, як це було показано в прикладах надзвичайних ситуацій раніше.

Такі сценарії є особливо критичними в важливих системах, таких як автономні транспортні засоби, інфраструктура моніторингу охорони здоров'я, промислова автоматизація та військові мережі, де переривання послуг може призвести до катастрофічного збою та непоправної шкоди обладнання.

Пристрої IoT зазвичай обмежені з точки зору обчислювальної потужності та вбудованої пам'яті, що обмежує їхню здатність розміщувати складні протоколи шифрування або надійні системи брандмауерів. Як наслідок, вони часто служать точками входу для зловмисників, збільшуючи вразливість ширшої мережі. Крім того, різноманітність і масштаб розгортання IoT часто призводять до фрагментованого управління безпекою та непослідовних практик оновлень.

Окрім кіберзагроз, умови навколишнього середовища та апаратні обмеження створюють незначні ризики. Пристрої IoT часто розгортаються в складних або неконтрольованих умовах, де вплив коливань температури, вологості або електромагнітних перешкод може вплинути на продуктивність датчиків і точність даних. Ці стресові фактори можуть призвести до деградації обладнання та можливого виходу з ладу, що ставить під загрозу як цілісність даних, так і безперервність роботи.

Щоб пом'якшити ці проблеми, можна використовувати рішення на основі штучного інтелекту для моніторингу стану системи, виявлення ранніх ознак несправності обладнання та управління пріоритезацією ризиків між вузлами мережі. Моделі прогнозованого технічного обслуговування, засновані на алгоритмах машинного навчання, можуть прогнозувати збої в апаратному забезпеченні до того, як вони виникнуть, дозволяючи вжити превентивних заходів, які зменшують час простою системи та зберігають точність даних.

Зрештою, захист інфраструктури IoT вимагає інтегрованого підходу, який поєднує аналітику кібербезпеки, стійкість навколишнього середовища та діагностичні інструменти на основі штучного інтелекту. Забезпечення спільного захисту даних і фізичних активів має важливе значення для підтримки довгострокової надійності та безпеки інтелектуальних систем, що працюють у невизначених і динамічних умовах.

Моделі та методи використання ШІ у захисті даних

Оскільки вимоги до безпеки середовищ IoT продовжують ускладнюватися, традиційні механізми захисту, засновані на правилах, виявилися недостатніми для ефективної протидії загрозам, що виникають. Штучний інтелект (ШІ) з його здатністю

до автономного навчання та адаптивного прийняття рішень пропонує трансформаційний підхід до захисту даних у цих системах. У цьому розділі розглядаються основні моделі та методології на основі штучного інтелекту, які зараз використовуються для підвищення кібербезпеки, особливо в середовищах, що страждають від інформаційної невизначеності.

Системи виявлення вторгнень

Одне з найвідоміших застосувань штучного інтелекту в кібербезпеці – у розробці інтелектуальних систем виявлення вторгнень (IDS) (рис. 3) [13]. На відміну від статичних, заснованих на сигнатурах систем, IDS з посиленням штучним інтелектом може динамічно відстежувати мережевий трафік, виявляти аномальні закономірності та виявляти потенційні вторгнення в режимі реального часу. Для полегшення цього процесу зазвичай використовуються методи машинного навчання (ML) і глибокого навчання (DL).

Нейронні мережі

Серед архітектур глибокого навчання особливо виділяються згорткові нейронні мережі (CNNs) (рис. 4) [14] і рекурентні нейронні мережі (RNNs) (рис. 5) [15]. CNN дуже ефективні в обробці структурованих вхідних даних, таких як журнали мережевого трафіку, і можуть виявляти просторові особливості, що свідчать про зловмисну активність. На противагу цьому, RNN, включаючи варіанти довготривалої короткострокової пам'яті (LSTM), чудово аналізують послідовні дані, що робить їх придатними для виявлення залежних від часу моделей загроз і поведінки, що розвиваються з часом.

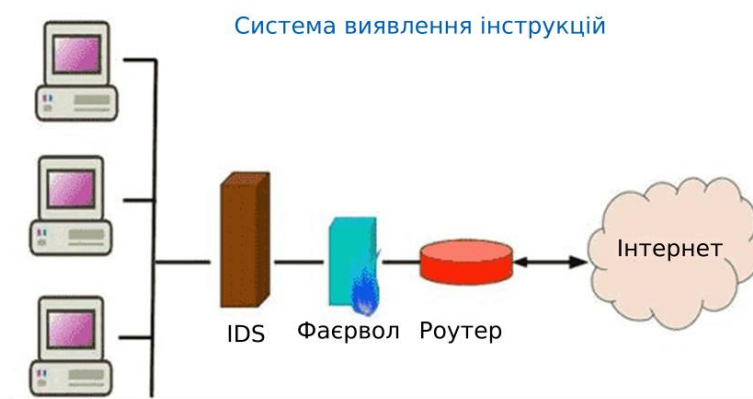


Рисунок 3. Візуальне представлення IDS

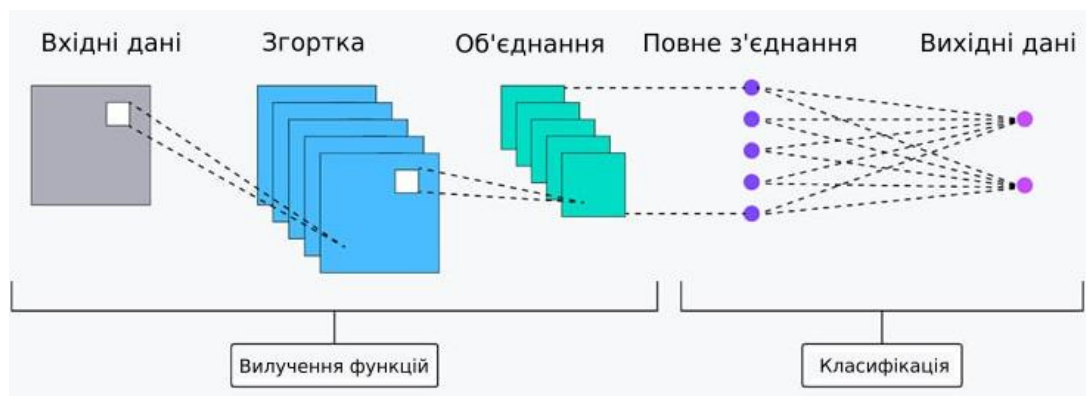


Рисунок 4. Архітектура CNN

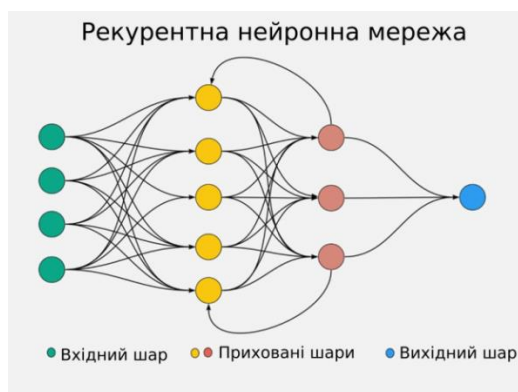


Рисунок 5. Архітектура RNN

Виявлення аномалій

На додаток до розпізнавання образів, виявлення аномалій (AD) залишається критично важливою областю застосування (рис. 6 [16]). Алгоритми штучного інтелекту можна навчити розпізнавати нормальні робочі базові лінії та позначати відхилення, які свідчать про несанкціонований доступ, маніпуляції даними або збої в обслуговуванні. Такі моделі особливо вигідні для виявлення загроз нульового дня та просунутих постійних загроз (APT), які часто обходять традиційні інструменти виявлення.

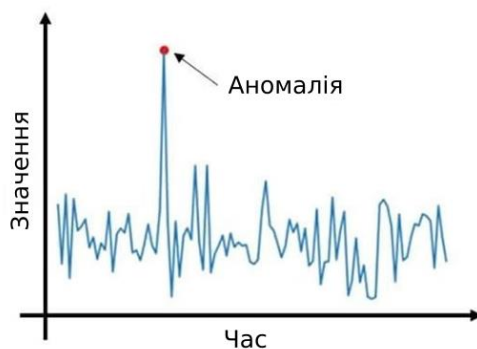


Рисунок 6. Візуальне представлення AD

Інші моделі на основі штучного інтелекту

Прогнозна аналітика на основі штучного інтелекту ще більше підвищує стійкість системи, прогножуючи потенційні порушення безпеки на основі історичних тенденцій даних. Це уможливорює проактивні стратегії пом'якшення наслідків і зміцнює надійність системи в довгостроковій перспективі.

Крім того, просунуті моделі, такі як навчання з підкріпленням, дозволяють системам удосконалювати свої стратегії захисту з часом через взаємодію з навколишнім середовищем, тоді як системи з нечіткою логікою (рис. 7 [17]) забезпечують гнучку основу для прийняття рішень у контексті часткових або неоднозначних даних - загальна характеристика невизначених екосистем IoT.

Методи на основі штучного інтелекту

Context – Aware Access Control (CAAC) – це динамічний метод контролю доступу, який коригує дозволи користувача/пристрою на основі контексту реального часу, такого як місцезнаходження, поведінка, стан пристрою, час або рівень ризику. CAAC не вимагає повної впевненості в даних, замість цього він приймає рішення про доступ, використовуючи наявну інформацію, часто призначаючи оцінки ризику або рівні впевненості, що, хоча і не без ризиків, робить його досить ефективним у випадках інформаційної невизначеності.

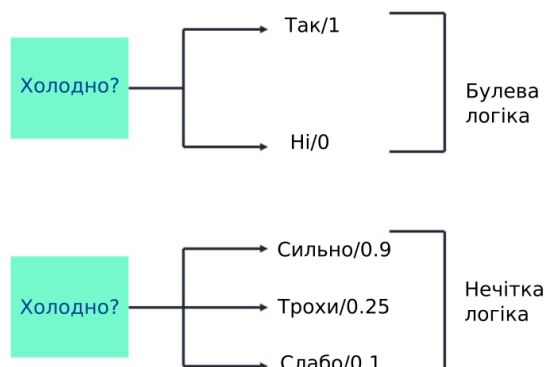


Рисунок 7. Візуальне представлення нечіткої логіки

Штучний інтелект дозволяє CAAC перейти від системи, заснованої на правилах, до інтелектуального та адаптивного процесу прийняття рішень. Більш умироутворено, це дозволяє CAAC обробляти неповні або зашумлені дані, вивчати поведінкові патерни, приймати рішення з управління ризиками та здійснювати безперервний моніторинг та адаптацію[18][19].

AI – driven deception and honeynet deployment (DDHD) – це ще один метод безпеки даних на основі штучного інтелекту, який може бути ефективним проти IU.

Обман, керований штучним інтелектом, – це використання штучного інтелекту для створення динамічних, реалістичних та адаптивних фальшивих систем, даних та

послуг, які призначені для залучення, введення в оману та захоплення зловмисників та їхнього програмного забезпечення. Вони часто покладаються на передбачувані структури та поведінку, щоб непомітно рухатися мережею. Обман обертає ці очікування проти них, вводячи фальшиві елементи, які виглядають і поведуться як законні цілі, але служать лише для виявлення, затримки та дезінформації[20].

“Honeynet” мережа “ханіпотів (медові горщики)” — системи - приманки, створені для імітації реальних пристроїв і сервісів — з метою залучення кіберзловмисників і контрольованого вивчення їхніх методів.

У той час як “ханіпоти” є окремими пастками (як фальшива сторінка входу), “honeynet” моделює ціле цифрове середовище: веб – сервери, бази даних, пристрої IoT, навіть фальшивих користувачів і робочі процеси.

DDHD приймають цю невизначеність. Замість того, щоб намагатися повністю усунути двозначність (що часто неможливо), вони створюють контрольоване середовище, де зловмисники добровільно викривають себе.

Втягуючи порушників у фальшивий, контрольований простір, захисники отримують перевагу - вивчають прийоми нападника, спостерігають за поведінкою та розробляють реакції - і все це без ризику для реальних систем.

Ще одним цікавим методом є злиття розвідувальних даних про загрози (TIF). TIF - це процес збору та об'єднання інформації з кількох джерел для створення цілісного уявлення про ландшафт загроз у реальному часі. Коли цей процес працює на основі AI, він стає швидшим, адаптивнішим і здатним справлятися з невизначеністю та складністю даних [21].

Злиття, кероване штучним інтелектом, не потребує ідеальних даних. Він процвітає в умовах невизначеності за рахунок:

- Оцінка найбільш ймовірних сценаріїв розвитку подій
- Постійне навчання та адаптація на основі зворотного зв'язку
- Виявлення ледь помітних закономірностей, які люди не помічають

Це перетворює невизначеність із тягара на перевагу, дозволяючи захисникам діяти раніше, впевненіше та ефективніше.

У сукупності розглянуті методології та моделі штучного інтелекту забезпечують основу для інтелектуальних, масштабованих та адаптивних фреймворків безпеки. Їхня здатність самовдосконалюватися та працювати в режимі реального часу позиціонує їх як важливі компоненти майбутніх – готових стратегій захисту IoT, особливо в середовищах, скомпрометованих ІУ.

Робота над рішенням

Оцінка ефективності моделей та методів захисту даних IoT в умовах інформаційної невизначеності на основі штучного інтелекту потребує багатогранного підходу. Це передбачає порівняння продуктивності з традиційними системами безпеки,

аналіз адаптивності в режимі реального часу та стійкість до мінливих загроз. Моделі, розглянуті в попередніх розділах, включаючи виявлення аномалій на основі машинного навчання та системи виявлення вторгнень на основі глибокого навчання (IDS), демонструють значні переваги у швидкості, точності та автономності системи.

Алгоритми машинного навчання здатні підвищити безпеку даних, вико ристовуючи свою здатність для постійного навчання на поведінці системи та як реальних, так і змодельованих атаках, адаптуючись до раніше небачених шаблонів загроз. Рекурентні нейронні мережі (RNN) та згорткові нейронні мережі (CNN) сприяють класифікації аномалій у реальному часі та розпізнаванню образів, що має важливе значення для захисту динамічних екосистем IoT. На момент написання даної статті це основні моделі, які використовуються в роботі над створенням першої версії адаптивної моделі безпеки даних на основі штучного інтелекту.

На даний момент двома основними цілями є:

1. Завершити аналіз методів та моделей безпеки даних на основі штучного інтелекту для виявлення найбільш ефективних та здатних виконувати поставлені завдання у разі їх використання;
2. Створити архітектуру методу захисту даних на основі штучного інтелекту.

У порівнянні зі статичними системами, заснованими на правилах, методи, керовані штучним інтелектом, виявляються кращими масштабованими та можливостями прогнозування, особливо в середовищах з високою інформаційною невизначеністю. Однак необхідно визнавати такі обмеження, як обчислювальні накладні витрати, залежність від великих наборів даних для навчання та потенційну вразливість зловмисників. Отже, для збалансованої продуктивності рекомендуються гібридні фреймворки, що інтегрують штучний інтелект із традиційними засобами керування та моделями прийняття рішень на основі ризику.

Емпіричні оцінки, отримані на основі симуляційних середовищ і контрольованих розгортань, показують до 40–60% покращення ранніх показників виявлення загроз і значне зниження кількості помилкових спрацьовувань при впровадженні систем, вдосконалених штучним інтелектом. Це підтверджує їх потенціал як важливих компонентів у сучасних архітектурах безпеки IoT.

Майбутні перспективи

Майбутнє захисту даних IoT в середовищах, що характеризуються високою інформаційною невизначеністю, все більше залежатиме від інтеграції технологій штучного інтелекту. У міру того, як екосистеми IoT стають все більш складними та взаємопов'язаними, традиційні статичні рамки безпеки будуть недостатніми. Очікується, що майбутні досягнення будуть зосереджені на підвищенні адаптивності, автономності та контекстуальної обізнаності систем, керованих штучним інтелектом.

Нові сфери розвитку включають федеративне навчання для децентралізованого навчання, пояснюваний штучний інтелект (XAI) для прозорого прийняття рішень

і периферійний штучний інтелект для локалізованого реагування на загрози з мінімальною затримкою. Ці технології сприятимуть виявленню та реагуванню в режимі реального часу, зменшать залежність від центральної хмарної обробки та підвищать довіру між користувачами та регулюючими органами.

Крім того, штучний інтелект відіграватиме ключову роль у прогностичній безпеці - прогнозуванні можливих векторів атак до того, як вони відбудуться, а також у самовідновлювальних мережах, які можуть автономно ізолюватися та відновлюватися після кібератак. Міждисциплінарна співпраця між кібербезпекою, етикою штучного інтелекту та системною інженерією матиме вирішальне значення для побудови безпечних, стійких та етично регульованих розумних середовищ.

У випадку з адаптивними системами безпеки, поки тривають роботи, є великі перспективи. Незважаючи на те, що вони знаходяться на ранніх стадіях і потребують відповідей на питання про необхідне обладнання, щоб зробити його портативним і надійним, моделі та методи на основі штучного інтелекту фактично показують, що вони набагато ефективніші для боротьби з загрозами кібербезпеки, проте вони не позбавлені слабких місць. Ідея взяти з них найкраще аж ніяк не нова, проте вона довела свою ефективність, і використання її для усунення небезпеки, пов'язаної з інформаційною невизначеністю, дає великі перспективи для підвищення загальної безпеки даних в екосистемі IoT, і, потенційно, може бути використане в інших мережевих і дата-орієнтованих середовищах.

Інвестиції в дослідження штучного інтелекту та створення всесвітньо визнаних стандартів кібербезпеки сформують траєкторію безпечної інфраструктури IoT у наступному десятилітті. Конвергенція штучного інтелекту та IoT не лише переосмислить стратегії безпеки, але й вплине на те, як захищені критично важливі системи в охороні здоров'я, енергетиці, транспорті та обороні.

Висновки

Зростаюча складність і масштаб екосистем IoT, особливо в умовах інформаційної невизначеності, вимагають передових і інтелектуальних рішень безпеки. Традиційних методів захисту вже недостатньо для протидії динамічним кіберзагрозам, які використовують непередбачуваність потоку даних і поведінки системи. У цьому документі продемонстровано, як штучний інтелект — за допомогою машинного навчання, глибокого навчання, виявлення аномалій і прогностичної аналітики — може бути використаний для значного підвищення стійкості та адаптивності систем IoT.

Інтегруючи штучний інтелект у системи безпеки IoT, організації можуть отримати вигоду від виявлення загроз у реальному часі, можливостей автономного реагування та адаптивних стратегій управління ризиками. Крім того, у дослідженні наголошується на важливості гібридних підходів, які поєднують штучний інтелект із звичайними моделями для усунення як поточних, так і виникаючих вразливостей.

У міру того, як технології штучного інтелекту продовжують розвиватися, їхня роль у захисті цифрової інфраструктури розширюватиметься, забезпечуючи більш інтелектуальний, масштабований і проактивний захист. Впровадження систем безпеки на основі штучного інтелекту не лише підвищує операційну ефективність, але й сприяє надійності критично важливої інфраструктури в різних секторах, від громадської безпеки до промислової автоматизації та національної оборони.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. *Abbas, A.* AI and IoT: Strengthening Information Security in Smart Ecosystems [Electronic resource] / Abbas, A. H. D. Roest, // ResearchGate, December 2024 <http://dx.doi.org/10.13140/RG.2.2.24583.92328>
2. *Arnott B.* 7 Examples of How AI is Improving Data Security Forcepoint: website URL: <https://www.forcepoint.com/blog/insights/ai-data-security-examples> (application date: 04.06.2025).
3. *Bolannavar J.* Smart Security : A Comprehensive Exploration of AI and IoT in Cybersecurity [Electronic resource] / J. Bolannavar, S. R. Gowda, // International Journal of Scientific Research in Computer Science, Engineering and Information Technology. No 10(6), 2024. P. 635-640. <http://dx.doi.org/10.32628/CSEIT2410447>
4. *Mccall A.* Cybersecurity in the Age of AI and IoT: Emerging Threats and Defense Strategies [Electronic resource] / A. Mccall // ResearchGate, November 21, 2024 – Available from: https://www.researchgate.net/publication/386050391_Cybersecurity_in_the_Age_of_AI_and_IoT_Emerging_Threats_and_Defense_Strategies
5. *Sikder A. K.* A Survey on Sensor-based Threats to Internet-of-Things (IoT) Devices and Applications [Electronic resource] / A. K. Sikder, G. Petracca, H. Aksu, T. Jaeger, and A. S. Uluagac, // 18 February 2018, <http://dx.doi.org/10.48550/arXiv.1802.02041>
6. *Muntean P.* A corrupt file led to the FAA ground stoppage. It was also found in the backup system CNN travel: website URL: <https://edition.cnn.com/travel/article/faa-ground-stop-causes/index.html> (application date: 04.06.2025).
7. *Eddy N.* Operation Triangulation' Spyware Attackers Bypass iPhone Memory Protections *DarkReading*: website URL: <https://www.darkreading.com/application-security/operation-triangulation-spyware-attackers-bypass-iphone-memory-protections> (application date: 04.06.2025)
8. *Lakshmanan R.* iOS Zero-Day Attacks: Experts Uncover Deeper Insights into Operation Triangulation *The Hacker News*: website URL: https://thehackernews.com/2023/10/operation-triangulation-experts-uncover.html?utm_source=chatgpt.com (application date: 04.06.2025)
9. Understanding the Change Healthcare Breach and Its Impact on Security Compliance *Hyperproof*: website URL: <https://hyperproof.io/resource/understanding-the-change-healthcare-breach/> (application date: 04.06.2025)

10. Office for Civil Rights (OCR), Change Healthcare Cybersecurity Incident Frequently Asked Questions *U.S. Department of Health and Human Service*: website URL: <https://www.hhs.gov/hipaa/for-professionals/special-topics/change-healthcare-cybersecurity-incident-frequently-asked-questions/index.html> (application date: 04.06.2025)
11. *Kasabji D.* A Sophisticated Cyber Offensive: Unraveling the Layers *Conscia*: website URL: https://conscia.com/blog/deep-dive-into-the-may-2023-cyber-attack-on-danish-energy-infrastructure/?utm_source=chatgpt.com (application date: 04.06.2025)
12. What is an Intrusion Detection System (IDS)? *Comodo*: website URL: <https://www.comodo.com/ids-in-security.php> (application date: 04.06.2025)
13. *Suryawanshi P. K.* Federated Learning and Hybrid IDS: A Novel Approach to IoT Security [Electronic resource] / P. K. Suryawanshi, S. K. Jagtap // SSRN Electronic Journal. No5, 2025, P.1-14 <http://dx.doi.org/10.2139/ssrn.5241753>
14. What is a Convolutional Neural Network? An Engineer's Guide *Zilliz*: website URL: <https://zilliz.com/glossary/convolutional-neural-network> (application date: 04.06.2025)
15. *Tang Ch.* DeepSCNN: a simplicial convolutional neural network for deep learning [Electronic resource] / Ch. Tang, Zh. Ye, Libing Bai, H. Zhao, // Applied Intelligence, No 55(4), 2025, <http://dx.doi.org/10.1007/s10489-024-06121-6>
16. Countants, How Machine Learning Can Enable Anomaly Detection *Medium*: website URL: <https://medium.datadriveninvestor.com/how-machine-learning-can-enable-anomaly-detection-eed9286c5306> (application date: 04.06.2025)
17. Deb Sayantini, What is Fuzzy Logic in AI and What are its Applications?, *Edureka!*: website URL: <https://www.edureka.co/blog/fuzzy-logic-ai/> (application date: 04.06.2025)
18. *Li B.* Context-Aware Risk Attribute Access Control [Electronic resource] / B. Li, Fan Yang, S. Zhang, // Mathematics, No 12(16):2541, 2024, 20 p. <http://dx.doi.org/10.3390/math12162541>
19. *Sondes B.* BIG-ABAC: Leveraging Big Data for Adaptive, Scalable, and Context-Aware Access Control [Electronic resource] / B. Sondes, T. Abdellatif, //Computer Modeling in Engineering & Sciences. 143(1), 2025, pp.1071-1093. <http://dx.doi.org/10.32604/cmes.2025.062902>
20. *Ebunoluwa A.* AI-Powered Honeypots: Enhancing Deception Technologies for Cyber Defense [Electronic resource]/ A. Ebunoluwa, A. James // ResearchGate, February 2025, 5 p Available from: URL: https://www.researchgate.net/publication/390113901_AI-Powered_Honeypots_Enhancing_Deception_Technologies_for_Cyber_Defense
21. *Benjamin M.* Real-Time Threat Intelligence Feeds and Machine Learning Fusion: Proactive Security in Dynamic Network Environments [Electronic resource] / M. Benjamin // ResearchGate, April 2025, 19 p Available from: URL: https://www.researchgate.net/publication/390887238_Real-Time_Threat_Intelligence_Feeds_and_Machine_Learning_Fusion_Proactive_Security_in_Dynamic_Network_Environments