

РОЗДІЛ II. ІНФОРМАЦІЙНО-КОМП'ЮТЕРНІ ТЕХНОЛОГІЇ

DOI: [https://doi.org/10.25140/2411-5363-2025-3\(41\)-131-145](https://doi.org/10.25140/2411-5363-2025-3(41)-131-145)

УДК 004.056.5:004.89:005.8:519.876.5

Віктор Федорович Гречанінов

кандидат технічних наук, старший дослідник, завідувач науково-дослідного відділу

Інститут проблем математичних машин і систем НАН України (Київ, Україна)

E-mail: vgrechaninov@gmail.com. ORCID: <https://orcid.org/0000-0001-6268-3204>

ІНТЕЛЕКТУАЛЬНА СИСТЕМА СЦЕНАРНОГО МОДЕЛЮВАННЯ ОЦІНКИ ЗАХИЩЕНОСТІ ОБ'ЄКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ

У статті представлено інтелектуальну систему моделювання сценаріїв оцінки захищеності об'єктів критичної інфраструктури з використанням методів штучного інтелекту. Розроблено формалізовану модель сценарного аналізу, що поєднує машинне навчання, нечітку логіку та графові структури для відображення взаємозв'язків між активами, загрозами, уразливостями й контрзаходами. Система забезпечує автоматизовану генерацію сценаріїв інцидентів із урахуванням часової динаміки ризику та може бути використана в центрах моніторингу безпеки й цифрових двійниках кіберзахисту.

Ключові слова: оцінка ризику; моделювання сценаріїв; штучний інтелект; критична інфраструктура; нечітка логіка; граф ризику; машинне навчання; кіберзахисність; цифрові двійники.

Рис.: 5. Бібл.: 17.

Актуальність теми дослідження. Об'єкти критичної інфраструктури відіграють визначальну роль у функціонуванні держави, забезпечуючи стабільність енергетичних, транспортних, комунікаційних та фінансових систем [6; 9; 12]. З огляду на зростання кількості складних цілеспрямованих кібератак, що уражають енергетичні компанії, транспортні вузли та телекомунікаційні мережі, виникає нагальна потреба переходу від реактивних підходів до проактивного сценарного моделювання стану захищеності [2-3; 10].

Постановка проблеми. Проблема забезпечення належного рівня захищеності об'єктів критичної інфраструктури (ОКІ) набуває особливої актуальності в умовах зростання кількості цілеспрямованих кібератак, ускладнення їхньої структури та інтеграції шкідливих дій у промислові й енергетичні процеси [1-2; 6]. Традиційні методи оцінювання безпеки, передбачені стандартами ISO/IEC 27005, NIST SP 800-30 чи OCTAVE, не відображають реальної динаміки ризиків, оскільки спираються на статичні експертні оцінки й не враховують часову змінність уразливостей, взаємозалежність активів і реакцію контрзаходів [7; 9]. Це призводить до несвоєчасного оновлення профілю ризику, недостатньої адаптивності систем захисту та втрати релевантності ухвалених рішень у кризових ситуаціях.

Необхідність переходу до інтелектуальних методів аналізу зумовлена потребою в оперативному прогнозуванні стану захищеності ОКІ та побудові динамічних сценаріїв розвитку подій у разі реалізації загроз [1-5; 15]. Запровадження штучного інтелекту в процес оцінювання дозволяє відмовитися від жорстких детермінованих моделей на користь адаптивних систем, здатних до самонавчання на основі історичних інцидентів і телеметричних даних [4; 8; 11]. Інтелектуальна система моделювання сценаріїв оцінки захищеності передбачає використання машинного навчання, нечіткої логіки та графових структур для відображення причинно-наслідкових зв'язків між активами, загрозами, уразливостями й контрзаходами [5; 6; 8; 9; 11-12; 14-16], що забезпечує побудову інтегрованої моделі ризику в часовому вимірі.

Аналіз останніх досліджень і публікацій. У сучасній літературі спостерігається посиленна увага до використання методів штучного інтелекту для оцінювання ризиків і забезпечення захищеності об'єктів критичної інфраструктури. Так, у роботі [1] представлено AI-орієнтовану модель ідентифікації ризиків інфраструктурних проєктів, що використовує історичні дані для виявлення закономірностей ризикових факторів, однак запропонований

підхід не враховує багатокроковість сценаріїв та часову динаміку впливу загроз. У дослідженні [2] проаналізовано застосування машинного навчання для прогнозування кіберризиків у критичних інфраструктурах, із фокусом на детекції аномалій, проте без інтеграції графових моделей зв'язків між активами, уразливостями й контрзаходами.

В огляді [3] висвітлено роль генеративного штучного інтелекту та великих мовних моделей у захисті критичних систем, визначено ключові виклики, пов'язані з надійністю, пояснюваністю та безпечністю моделей, але не запропоновано практичної реалізації сценарного аналізу ризиків. У роботі [4] подано автономну AI-архітектуру для виявлення уразливостей і реагування на загрози в реальному часі, проте система має загальний характер і не враховує сценарне прогнозування змін рівня ризику.

Автори роботи [5] розглянули інтеграцію моделей машинного навчання у сфері національної безпеки й критичної інфраструктури, проте не подали механізмів побудови сценаріїв чи оцінки ефективності контрзаходів у динамічному середовищі.

Таким чином, проаналізовані праці свідчать про активний розвиток напрямів застосування AI для прогнозування ризиків і виявлення аномалій у критичних системах, однак залишаються недостатньо опрацьованими питання інтегрованого сценарного моделювання, графового представлення взаємозв'язків між загрозами, активами й контрзаходами, а також забезпечення пояснюваності й адаптивного оновлення сценаріїв у режимі реального часу. Ці наукові прогалини визначають потребу у створенні комплексної інтелектуальної системи, що поєднає методи штучного інтелекту, сценарний аналіз і динамічне управління ризиками для об'єктів критичної інфраструктури.

Виділення недосліджених частин загальної проблеми. Попри наявність численних досліджень із питань оцінювання ризиків та захищеності об'єктів критичної інфраструктури, низка аспектів залишається маловивченою. Більшість існуючих підходів є статичними й не враховують динамічну зміну стану систем, появу нових загроз та взаємозалежність між активами, уразливостями й контрзаходами. Недостатньо розроблені методи автоматизованого формування сценаріїв розвитку інцидентів з урахуванням часової динаміки ризику та реального контексту подій. Недослідженою також залишається інтеграція штучного інтелекту в процес моделювання оцінки захищеності, зокрема поєднання нейронних мереж, нечіткої логіки та графових структур. Відсутність ефективних механізмів пояснюваного аналізу рішень ускладнює інтерпретацію результатів для експертів кіберзахисту. Крім того, бракує підходів до створення цифрових двійників, що дозволяли б симулювати сценарії атак і оцінювати ефективність контрзаходів. Отже, подальші дослідження мають бути спрямовані на розроблення інтелектуальної системи сценарного моделювання, здатної поєднати машинне навчання, нечітку логіку та механізм пояснюваного аналізу рішень для адаптивної оцінки захищеності й прогнозування ризиків у режимі реального часу.

Мета дослідження. Метою дослідження є розроблення інтелектуальної системи сценарного моделювання оцінки захищеності об'єктів критичної інфраструктури, яка на основі методів штучного інтелекту, машинного навчання, нечіткої логіки та графових структур забезпечує автоматизоване формування сценаріїв розвитку інцидентів, прогнозування динаміки ризиків і визначення оптимальних контрзаходів для підвищення рівня кіберстійкості та адаптивності систем захисту.

Виклад основного матеріалу. Побудова інтелектуальної системи сценарного моделювання оцінки захищеності об'єктів критичної інфраструктури потребує формального опису її структури та взаємозв'язків між компонентами [6; 11; 16]. Для цього розроблено математичну модель, що поєднує методи машинного навчання, нечіткої логіки, графових структур і пояснюваного штучного інтелекту [8; 10]. Такий підхід забезпечує можливість виконання адаптивного аналізу ризику в динамічних умовах кіберзагроз, коли оцінка захищеності має не лише статичний, а і прогностичний характер.

У межах цієї моделі кожен сценарій розглядається як формалізована структура, яка об'єднує активи, загрози, уразливості, контрзаходи та прогноз ризику, що дозволяє перейти до аналітичного опису процесів зміни рівня безпеки у часі [9; 12; 16].

Формалізована модель сценарію описується шістькою:

$$S_i = \langle A_i, T_i, V_i, C_i, R_i(t), P_i \rangle, \quad (1)$$

де A_i – множина активів системи;

T_i – множина загроз;

V_i – множина уразливостей;

C_i – набір контрзаходів;

$R_i(t)$ – значення ризику в момент часу t ;

P_i – прогнозований сценарій розвитку події [6-8; 12].

Таке формальне представлення дає змогу структурно пов'язати всі ключові компоненти системи безпеки в єдину модель, що описує не лише поточний стан, а і прогнозовану динаміку ризику. Воно забезпечує можливість інтеграції різнорідних факторів – технічних, організаційних і поведінкових – у спільну аналітичну базу для подальшого моделювання сценаріїв, виявлення критичних ланцюгів атак і формування оптимальних стратегій реагування.

Ризик кожного сценарію обчислюється як:

$$R_i(t) = P(T_i \cap V_i) \cdot I(A_i) \cdot (1 - E(C_i)), \quad (2)$$

де $P(T_i \cap V_i)$ – ймовірність реалізації загрози через конкретну уразливість;

$I(A_i)$ – критичність активу;

$E(C_i)$ – ефективність застосованих контрзаходів.

Формула відображає багаторівневу взаємодію між елементами системи безпеки: загрозами, уразливостями та механізмами захисту. Коефіцієнт $I(A_i)$ підсилює внесок найкритичніших активів, тоді як множник $(1 - E(C_i))$ враховує ступінь зниження ризику завдяки наявним контролям [6; 14]. Таким чином, (2) дозволяє виконати кількісне оцінювання ризику в динамічній системі, відображаючи як ймовірність реалізації атаки, так і її потенційний вплив на бізнес-процеси. Цей вираз є основою для подальшої нормалізації ризиків і побудови агрегованої системної оцінки.

Для розділення складових ймовірності у (2) використовується декомпозиція:

$$P(T_i \cap V_i) = P(T_i|V_i) \cdot P(V_i), \quad (3)$$

що дає змогу окремо оновлювати оцінки уразливостей і загроз – наприклад, на основі результатів сканування, телеметрії або аналітики SOC. Такий підхід спрощує адаптацію моделі до змін у середовищі безпеки, дозволяючи незалежно оновлювати параметри загроз $P(T_i|V_i)$ та уразливостей $P(V_i)$. Це особливо важливо для систем, що працюють у режимі реального часу, де частота виявлення нових уразливостей та появи атак постійно зростає [2-4; 8; 15]. Крім того, формула (3) забезпечує можливість інтеграції з байєсівськими або нечіткими моделями оцінки ризику, створюючи основу для побудови динамічних сценаріїв розвитку інцидентів.

Агрегована ефективність множини контрзаходів визначається як:

$$E(C_i) = 1 - \prod_{k \in C_i} (1 - \eta_k u_k), \quad 0 \leq \eta_k, u_k \leq 1, \quad (4)$$

де η_k – питомий вплив (якість) контрзаходу k ;

u_k – рівень його задіяння (активації) в системі захисту [7; 12; 14].

Вираз описує інтегральний ефект множини незалежних контрзаходів, які спільно впливають на зменшення ризику реалізації загроз.

Модель базується на принципі комплементарності, коли ефективність усієї системи зростає навіть при частковому підвищенні ефективності окремих контролів [15]. Якщо хоча б один контрзахід спрацьовує на 100 %, агрегована ефективність $E(C_i)$ наближа-

ється до одиниці, що означає повне блокування ризику. Формула (4) узгоджується з рівнянням (2), адже підвищення $E(C_i)$ прямо зменшує величину ризику $R_i(t)$, забезпечуючи нелінійний кумулятивний ефект при додаванні додаткових механізмів захисту.

Для порівняння різних сценаріїв використовується нормований показник:

$$\hat{R}_i = \frac{R_i(t)}{R_{max}}, \hat{R}_i \in [0; 1], \quad (5)$$

який характеризує відносний рівень ризику від повної стійкості (дорівнює 0) до критичного стану (дорівнює 1). Такий підхід забезпечує уніфіковану шкалу для порівняння станів різних активів або сценаріїв, незалежно від їхньої абсолютної величини ризику [7; 9]. Нормалізація дозволяє оцінювати не лише поточний рівень загроз, але й відстежувати динаміку їх зміни в часі – наприклад, унаслідок появи нових уразливостей або впровадження додаткових контрзаходів. Значення $\hat{R}_i$, що наближається до 1, свідчить про критичну вразливість активу, тоді як показник, близький до 0, демонструє ефективність існуючих механізмів захисту. Таким чином, формула (5) створює основу для подальшої агрегації ризиків та оптимізації політики безпеки підприємства.

Інтегральний системний ризик обчислюється як:

$$R_{sys}(t) = \sum_i a_i \hat{R}_i(t), a_i \geq 0, \sum_i a_i = 1, \quad (6)$$

де a_i – вагові коефіцієнти пріоритетності активів або сценаріїв, що відображають їхню критичність у бізнес-процесах. Цей вираз узагальнює індивідуальні ризики $\hat{R}_i(t)$, інтегруючи їх у єдину метрику, що характеризує поточний стан захищеності всієї системи. Коефіцієнти a_i можуть визначатися на основі важливості активів у виробничому циклі, рівня їх підключення до інших компонентів або наслідків потенційної компрометації [7; 9; 11-12]. Таким чином, формула (6) дозволяє перетворити множину локальних оцінок ризику на стратегічний показник, який використовується для ухвалення управлінських рішень, ранжування загроз та формування політики пріоритетного захисту об'єктів критичної інфраструктури.

Для валідації запропонованого підходу проведено порівняння з трьома орієнтирами: статичним ризик-скорингом, моделлю з фіксованим порогом спрацювання та класичними алгоритмами машинного навчання, що не враховують графові зв'язки між активами, загрозами й контрзаходами. Ефективність системи оцінювалася за сукупністю кількісних метрик – MAE, RMSE, R^2 , ROC-AUC, PR-AUC, F1, а також за середнім часом виявлення, частотою хибних спрацьовувань і затримкою чи пропускну здатністю в режимі реального часу. Такий комплекс показників забезпечує всебічну оцінку якості функціонування системи, охоплюючи точність прогнозування, своєчасність реагування та операційну ефективність в умовах динамічних кіберзагроз. Конфігурація моделей, включно з вікном спостереження, розміром рекурентних шарів, швидкістю навчання та пакетом обробки, зафіксована у внутрішньому профілі налаштувань і може бути відтворена відповідно до описаної послідовності етапів. Час тренування та інференсу контролювався з метою забезпечення інтеграції системи в потоки обробки даних, наближені до реального часу, що гарантує її практичну придатність у середовищах оперативного моніторингу безпеки.

Для візуалізації взаємозв'язків між елементами системи формується орієнтований граф ризику:

$$G = (N, E), \quad (7)$$

де $N = \{A, T, V, C\}$ – множина вузлів (активи, загрози, уразливості, контрзаходи), $E = \{(T, V), (V, A), (C, V)\}$ – множина орієнтованих зв'язків між ними. Такий граф відображає структуру взаємозалежностей між об'єктами критичної інфраструктури та дозволяє побудувати причинно-наслідкові ланцюги поширення ризику [6; 9; 16]. Кожне ребро вказує

напрям впливу – від загрози до уразливості, від уразливості до активу або від контрзаходу до уразливості, що дозволяє відстежувати як еволюцію потенційних атак, так і ефект застосованих захисних заходів. У такий спосіб граф G є базовою структурою для подальшого обчислення показників критичності та моделювання ланцюгових атак, забезпечуючи наочне представлення ризикових взаємозв'язків у системі.

Сумарна критичність системи розраховується як:

$$K(G) = \sum_{(i,j) \in E} w_{ij} \cdot R_j, \quad (8)$$

де w_{ij} – вага ребра, що визначає силу впливу між двома елементами графа;

R_j – ризик, асоційований із вузлом j .

Формула (8) дозволяє оцінити загальний вплив взаємозалежностей між компонентами системи безпеки, враховуючи як силу зв'язків між елементами, так і величину індивідуального ризику [11]. Якщо у системі існують вузли з високим рівнем ризику, що мають велику кількість зв'язків або високу вагу w_{ij} , їхній внесок у сумарну критичність $K(G)$ буде значним. Таким чином, рівняння (8) допомагає визначати вузли, які є «вузкими місцями» безпеки – тобто ті елементи, через які ризик найшвидше поширюється мережею, що робить їх пріоритетними для захисту та моніторингу.

Для аналізу складних багатокрокових або «ланцюгових» атак використовується показник ризику уздовж шляху:

$$R_p = (\prod_{(u \rightarrow v) \in p} w_{uv}) \cdot I(A_{end}), \quad p: \text{шлях у } G, \quad (9)$$

де добуток ваг w_{uv} відображає сукупну проникність між вузлами графа вздовж шляху атак p ;

$I(A_{end})$ – критичність кінцевого активу, до якого веде цей шлях.

Формула (9) дозволяє змоделювати каскадне поширення атаки через послідовні залежності між вразливими компонентами системи. Якщо хоча б одна ланка шляху має низьку проникність (малий w_{uv}), то загальний ризик R_p істотно зменшується, що узгоджується з логікою багаторівневого захисту [10]. У контексті графової моделі це дає змогу не лише оцінити ризик окремих вузлів, а й визначити критичні шляхи поширення загроз, які становлять найбільшу небезпеку для об'єкта критичної інфраструктури [5; 13]. Таким чином, показник R_p використовується для кількісного аналізу каскадних ефектів та оцінювання потенційного впливу складних атаківих ланцюгів у системі.

Еволюція ризику у графі описується динамічним рівнянням:

$$r_{t+1} = \sigma(\mathcal{A}_t r_t + b - T e_t), \quad (10)$$

де r_t – вектор вузлових ризиків у момент часу t ;

$\mathcal{A}_t$ – матриця ваг зв'язків між вузлами графа (активами, загрозами, уразливостями);

e_t – вектор ефективностей активних контрзаходів;

T – діагональна матриця сили їхнього впливу;

$\sigma(\cdot)$ – нелінійна стискуюча функція, наприклад логістична.

Рівняння описує, як ризик на кожному вузлі графа змінюється під впливом взаємодій між елементами системи та дією захисних механізмів. Доданок $\mathcal{A}_t r_t$ моделює передачу ризику між вузлами (наприклад, коли компрометація одного активу впливає на інші), а термін $T e_t$ зменшує загальний ризик залежно від ефективності активних контрзаходів. Нелінійна функція $\sigma(\cdot)$ обмежує значення ризику у допустимому діапазоні $[0,1]$, що забезпечує стабільність моделі та запобігає її розходженню при моделюванні довготривалих сценаріїв [15, 16]. Таким чином, формула (10) формалізує динамічну поведінку системи ризиків у часі, узгоджуючи її з попередніми співвідношеннями (7)–(9) і створюючи основу для прогнозного керування захищеністю об'єктів критичної інфраструктури.

Графік на рис. 1 відображає зміну рівня ризику r_t для п'яти вузлів системи у часі відповідно до рекурентної моделі $r_{t+1} = \sigma(\mathcal{A}_t r_t + b - T e_t)$. Взаємодія між вузлами визначається матрицею впливів $\mathcal{A}_t$, а ефективність контрзаходів – матрицею T . Криві показують, як ризик стабілізується після початкових коливань, демонструючи різні швидкості затухання залежно від локальної структури взаємодій та впливу параметра σ .

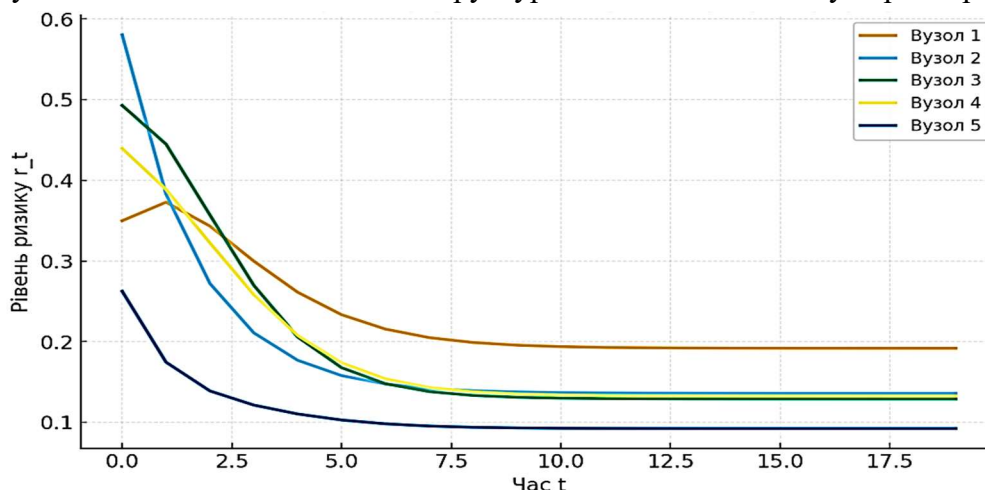


Рис. 1. Динаміка ризику у часі для вузлів графа

Джерело: розроблено автором.

Архітектура системи реалізована у вигляді трирівневої моделі, де кожен рівень виконує власні функції та взаємодіє з іншими через чітко визначені інтерфейси потоків даних, представлених на рис. 2.

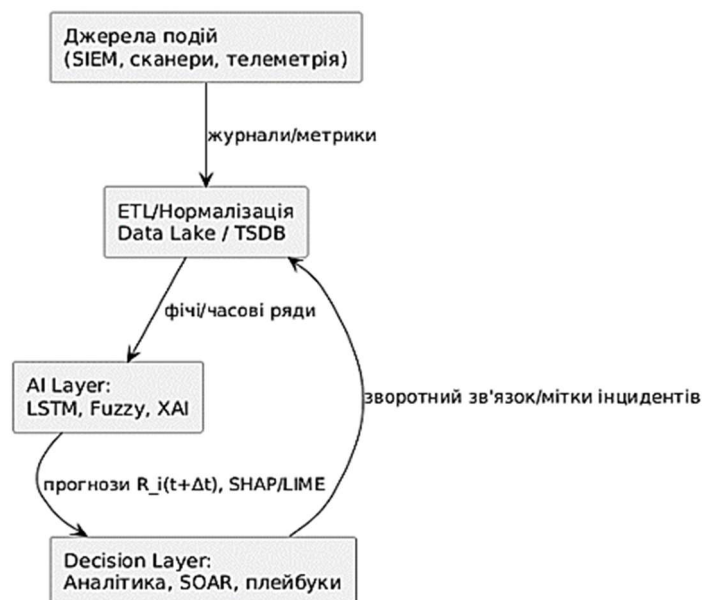


Рис. 2. Схема потоків даних інтелектуальної системи сценарного моделювання оцінки захищеності об'єктів критичної інфраструктури

Джерело: розроблено автором.

Data Layer забезпечує первинну обробку інформаційних потоків, які надходять із різних джерел – журналів подій безпеки (SIEM), результатів сканування уразливостей, систем моніторингу інфраструктури (ICS/SCADA), мережевої телеметрії та логів хостових агентів [14]. На цьому рівні здійснюється збирання, уніфікація, очищення та нормалізація даних за допомогою механізмів ETL (Extract–Transform–Load). Усі події зберігаються

у сховищі типу Data Lake або Time Series Database, що дозволяє виконувати як ретроспективний, так і поточний аналіз. Data Layer також відповідає за синхронізацію даних у реальному часі, їхню категоризацію за типами загроз і уразливостей та передає нормалізовані ознаки безпеки до рівня штучного інтелекту. AI Layer виконує аналітичну та прогнозу функції, використовуючи моделі машинного навчання, нечіткої логіки та пояснюваного штучного інтелекту для оцінки поточного й майбутнього рівня ризику [5; 8; 15]. На цьому рівні здійснюється побудова часових прогнозів, визначення ймовірності реалізації загроз і формування рекомендацій щодо пріоритетності контрзаходів. Отримані результати передаються до Decision Layer, який інтегрує аналітичні висновки в систему підтримки рішень, забезпечуючи автоматизоване реагування, оптимізацію ресурсів і формування сценаріїв реагування в межах SOC/SIEM/SOAR.

Потоки починаються з джерел подій (SIEM, сканери, телеметрія), звідки журнали та метрики надходять до модуля ETL/Нормалізація, де здійснюється очищення, уніфікація та передавання даних у сховище Data Lake / TSDB [9]. AI Layer виконує основну аналітичну й прогнозу функцію системи. Його ядром є нейронна мережа LSTM, яка аналізує часові ряди даних і визначає прогнозований рівень ризику $R_i(t + \Delta t)$ для кожного активу або сценарію.

Паралельно працює модуль нечіткої логіки (Fuzzy Logic Module), що опрацьовує експертні або напіваавтоматичні оцінки параметрів безпеки – наприклад, рівень критичності активів, ймовірність реалізації загроз, ефективність контрзаходів – і переводить ці оцінки у числову форму через механізми фазифікації, агрегації та дефазифікації. Додатково вбудований Explainable Artificial Intelligence (XAI) модуль інтерпретує результати прогнозу, пояснюючи, які вхідні фактори найбільше вплинули на оцінку ризику.

Для цього використовуються методи SHAP, LIME та аналітика важливості ознак, що підвищує довіру до результатів і спрощує їх використання аналітиками центрів моніторингу безпеки. Таким чином, AI Layer виконує роль інтелектуального ядра, яке перетворює масиви сирих даних у прогнозні оцінки та пояснення для прийняття рішень. Decision Layer є рівнем взаємодії з користувачем і відповідає за інтеграцію результатів аналітики у процес управління безпекою [16]. На цьому етапі система виконує візуалізацію графа ризику (7), що демонструє зв'язки між активами, загрозами, уразливостями та контрзаходами [15]. Граф відображає не лише поточний стан безпеки, але й динаміку зміни ризиків у часі. У модулі аналітики здійснюється оцінка трендів ризику, порівняння сценаріїв, побудова heatmap-карт критичності та ранжування активів за ступенем уразливості. Результати автоматично передаються до SOAR-компонента (Security Orchestration, Automation and Response), який формує рекомендації або плейбуки дій, оптимізуючи порядок реагування на інциденти з урахуванням доступного бюджету та ресурсів [6; 10]. Отже, Decision Layer замикає цикл управління – від збору даних до реалізації автоматизованих контрзаходів.

Отже, модуль нечіткої логіки забезпечує обробку невизначеностей і наближених оцінок, дозволяючи системі працювати навіть за неповних або суперечливих даних. Його поєднання з XAI-модулем створює прозорий механізм прийняття рішень, у якому кожне прогнозоване значення ризику може бути обґрунтовано на рівні факторного внеску. Такий підхід підвищує довіру аналітиків до результатів моделювання та спрощує аудит дій системи у контексті стандартів безпеки. У комплексі ці модулі формують когнітивне ядро AI-рівня, що забезпечує пояснюваність, адаптивність і стійкість оцінки ризику в динамічному середовищі загроз.

Взаємодія між трьома рівнями здійснюється за принципом безперервного циклу: Data Layer постачає актуальні дані, AI Layer генерує прогноз і пояснення, а Decision Layer забезпечує прийняття рішень, результати яких повертаються у вигляді нових даних і тригерів для корекції моделі [8; 14]. Така архітектура підтримує адаптивність, масштабованість і самонавчання системи, що є критично важливим для захисту об'єктів критичної

інфраструктури в умовах динамічних кіберзагроз. Завдяки модульній побудові забезпечується можливість гнучкого оновлення окремих компонентів без переривання роботи всієї системи. AI-модуль виконує прогнозування на основі часових рядів та формує пояснювані рекомендації для оператора, що підвищує прозорість процесу прийняття рішень. Decision Layer інтегрується з SOC/SIEM-платформами, автоматизуючи реагування та корекцію політик безпеки в реальному часі. Зворотний зв'язок між рівнями гарантує поступове вдосконалення моделі на основі накопиченого досвіду та змін у середовищі загроз, формуючи основу для інтелектуального управління ризиками.

На рис. 3 показано, як архітектура реалізує узгоджений потік даних між трьома рівнями, де кожен елемент виконує власну роль у формуванні цілісного контуру управління ризиками.

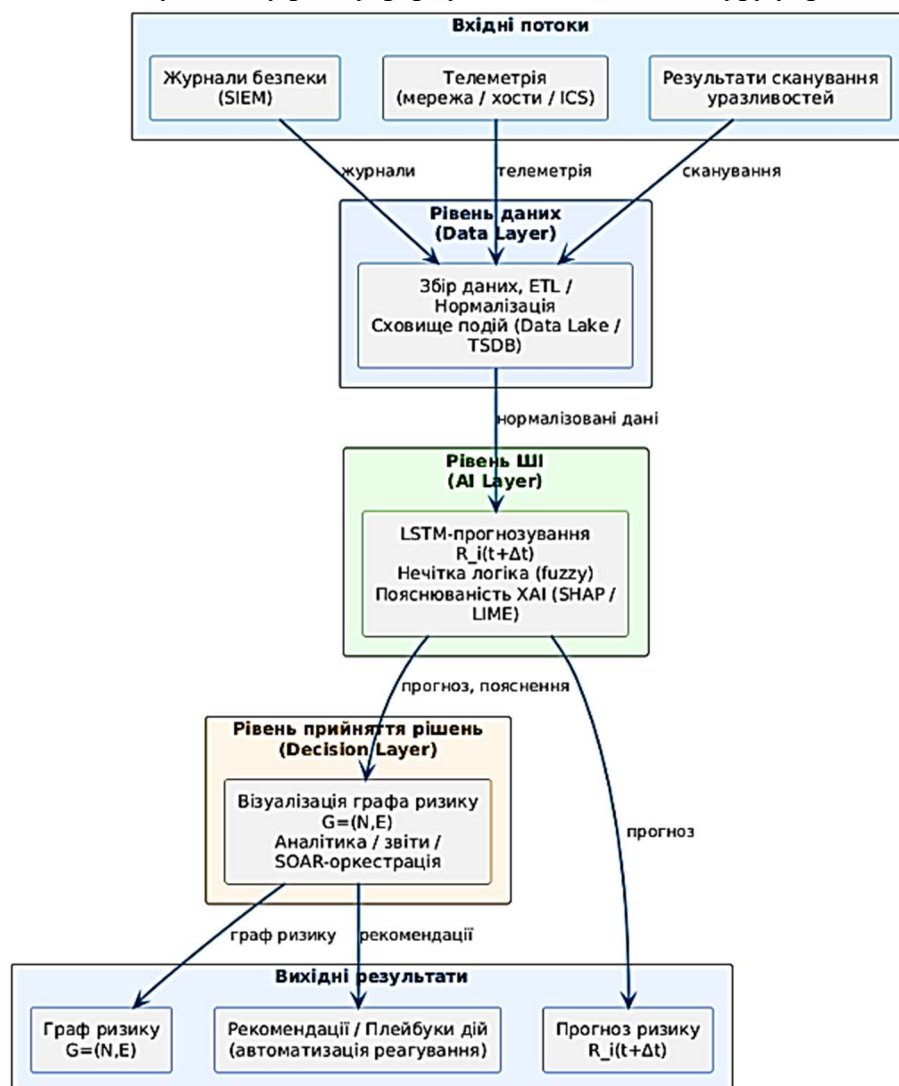


Рис. 3. Тривірнева архітектура інтелектуальної системи сценарного моделювання

Джерело: розроблено автором.

У центрі зображено три логічні шари системи – Data, AI та Decision, які об'єднані пунктирними лініями, що символізують обмін інформацією в режимі реального часу. Ліворуч показано вхідні джерела подій (журнали, телеметрія, результати сканувань), праворуч – результати аналізу, зокрема граф ризику, прогнозовані значення та рекомендації для реагування [5; 17]. Діаграма наочно демонструє, що система функціонує як замкну-

тий аналітичний контур, у якому результати прогнозів і рекомендацій Decision Layer можуть повторно впливати на Data Layer, формуючи механізм самонавчання та адаптації до змін у середовищі загроз.

Процес сценарного моделювання реалізується в кілька етапів: збір і попередня обробка даних D_t , визначення критичних активів і уразливостей, побудова прогнозної моделі за допомогою функції [2; 6]:

$$R_i(t + \Delta t) = f_\theta(D_t, A_i, T_i, V_i), \quad (11)$$

де f_θ – параметризована нейромережа, навчена на даних D_t . Вона враховує історичні залежності, взаємодії активів і поточні загрози, формуючи прогноз ризику для майбутнього моменту часу $t + \Delta t$. Завдяки використанню рекурентної структури (зокрема, LSTM або GRU-блоків) модель здатна відстежувати довготривалі залежності у поведінці системи безпеки та виявляти приховані тренди у зміні ризикових показників [5; 8; 15]. Такий підхід забезпечує не лише кількісне прогнозування, а й виявлення закономірностей, що передують виникненню критичних інцидентів, створюючи основу для проактивного реагування в межах системи кіберзахисту.

Ймовірнісний характер прогнозу подається як [15]:

$$R_i(t + \Delta t) \sim \mathcal{N}(\mu_i, \sigma_i^2), \quad \mu_i = f_\theta(D_t, A_i, T_i, V_i), \quad \sigma_i^2 = g_\theta(D_t, A_i, T_i, V_i), \quad (12)$$

де функція g_θ повертає дисперсію невизначеності прогнозу. Це забезпечує оцінку довіри до прогнозованого ризику та використовується далі для адаптивного налаштування порогів реагування. Крім того, значення μ_i відображає очікуваний рівень ризику для активу в майбутній момент часу, тоді як σ_i^2 характеризує ступінь коливань або нестабільності середовища [3; 10]. Такий підхід дозволяє системі розрізняти стабільні сценарії від потенційно небезпечних, коли зростає дисперсія прогнозу, що є індикатором підвищеної невизначеності або появи нових типів загроз.

Оцінка умовної ймовірності загрози оновлюється за правилом Байєса:

$$P_t(T_i | V_i) = \frac{\mathcal{L}(E_t | T_i, V_i) P_{t-1}(T_i | V_i)}{Z_t}, \quad (13)$$

де $\mathcal{L}$ – правдоподібність отриманих свідчень E_t ;

Z_t – нормувальний коефіцієнт [11; 13].

Цей механізм дозволяє системі адаптивно оновлювати ризиковий профіль у реальному часі. Оновлення за правилом Байєса забезпечує динамічне врахування нових подій і сигналів телеметрії без необхідності повного перенавчання моделі. Це дозволяє системі коригувати поточні оцінки загроз при надходженні нових індикаторів компрометації або зміні контексту середовища [10; 15]. Таким чином, модель підтримує безперервну адаптацію до мінливих умов кіберзагроз, зберігаючи актуальність прогнозів ризику в реальному часі.

Для навчання моделі використовується функція втрат:

$$L = \frac{1}{n} \sum_{i=1}^n (R_i^{true} - R_i^{pred})^2, \quad (14)$$

яка мінімізує похибку прогнозу. Додатково проводиться моніторинг метрик MAE, RMSE та коефіцієнта детермінації R^2 для оцінки точності передбачень на валідаційних вибірках. Для підвищення збіжності процесу оптимізації застосовуються адаптивні алгоритми градієнтного спуску, такі як Adam і RMSProp, а також регулярне перезапускання навчання з випадковою ініціалізацією ваг [14]. Валідаційні тести виконуються на основі синтетичних і реальних SOC-логів, що забезпечує перевірку узагальнювальної здатності моделі на різних сценаріях кіберзагроз.

Для підвищення стабільності додано регуляризацию Тихонова:

$$L_{total} = L + \lambda \|\theta\|_2^2, \quad (15)$$

що запобігає перенавчанню та стабілізує параметри моделі [8]. Додатково під час навчання застосовувались крос-валідація k-fold і раннє зупинення за метрикою валідаційної

похибки з адаптивним зменшенням швидкості навчання. Гіперпараметри моделі добиралися методом байєсівської оптимізації, а ознаки проходили нормування та регулярний контроль на мультиколінеарність [6; 11; 17]. Для підвищення стійкості використовували балансування класів/ваги втрат для рідкісних інцидентів, обрізання градієнтів і калібрування ймовірностей (temperature scaling) для коректної інтерпретації ризику.

На етапі ухвалення рішень система виконує оптимізацію контрзаходів за обмеженого бюджету:

$$\min_{x \in \{0,1\}^m} R'_{sys}(t) = \sum_i a_i \hat{R}'_i(t) \quad s. t. \quad \sum_{i=1}^n c_k x_k \leq B, \quad (16)$$

де a_i – вагові коефіцієнти, що відображають важливість активу або сценарію;

$\hat{R}'_i(t)$ – оновлений нормований ризик після застосування обраних заходів;

c_k – вартість реалізації кожного контрзаходу;

B – загальний бюджет, який не можна перевищити;

$x = (x_1, x_2, \dots, x_m)$ – бінарний вектор вибору контрзаходів.

Нові значення ризику $\hat{R}'_i(t)$ визначаються через ефективність контрзаходів $E(C_i \cup \{k : x_k = 1\})$. Таким чином, завдання (16) формалізує процес мінімізації залишкового ризику за наявних ресурсів і реалізується на рівні Decision Layer.

Для задачі мінімізації залишкового ризику за бюджетних обмежень використано детермінований розв'язувач з евристичним теплим стартом; при зростанні розмірності застосовується жадібна стратегія наближення з гарантованою зупинкою за часом. Такий підхід забезпечує стабільний компроміс між якістю рішень і часовими вимогами SOC-процесів.

Діаграма на рис. 4 ілюструє компроміс між залишковим ризиком $R'_{sys}(t)$ та вартістю реалізації контрзаходів B . Кожна точка відповідає можливому набору рішень x у задачі оптимізації, де мінімізується ризик за обмеженого бюджету. Виділена «робоча точка» відображає оптимальний баланс між ефективністю та витратами, що визначає доцільну стратегію реалізації заходів безпеки в системі критичної інфраструктури.

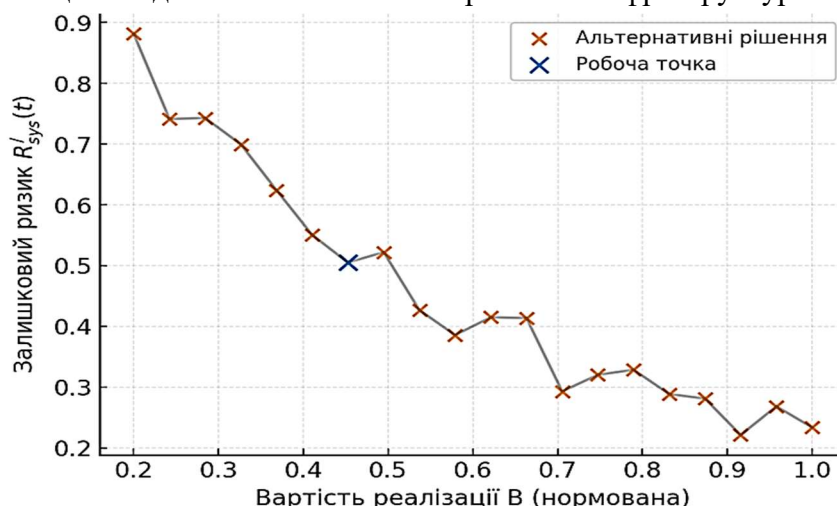


Рис. 4. Парето-фронт «ризик ↓ – вартість ↑» для сімейства рішень x
Джерело: розроблено автором.

Після оптимізації система визначає адаптивні пороги спрацьовування:

$$\tau_i = \tau_0 + \beta \sigma_i, \text{ спрацьовування, якщо } R_i(t) \geq \tau_i, \quad (17)$$

де τ_0 – базовий рівень порогу;

β – коефіцієнт чутливості системи до невизначеності;

σ_i – невизначеність прогнозу [8].

Підвищення порогу при великій дисперсії дозволяє знизити кількість хибних спрацьовувань SOC/AI-модуля [12-13, 15]. Таким чином, послідовність формул (1) - (17) описує цілісний цикл функціонування інтелектуальної системи: від формалізації ризику та графового моделювання до прогнозування, оптимізації контрзаходів і адаптивного реагування. Це забезпечує самоналаштовність та підвищує ефективність захисту об'єктів критичної інфраструктури в умовах динамічних кіберзагроз.

На рис. 5 (а) показано залежність між бюджетом B та залишковим системним ризиком $R'_{sys}(t)$: із зростанням бюджету ризик зменшується, а позначені точки $x^{(i)}$ відображають оптимальні набори контрзаходів у межах ресурсних обмежень. На рис. 5 (б) зображено вплив параметра масштабування порогу β на ймовірності хибних (FPR) і правильних (TPR) спрацьовувань при формулі $\tau_i = \tau_0 + \beta\sigma_i$. Зі зростанням β зменшується кількість хибних спрацьовувань, що демонструє адаптивність порогів реагування в AI-модулі оцінки ризику.

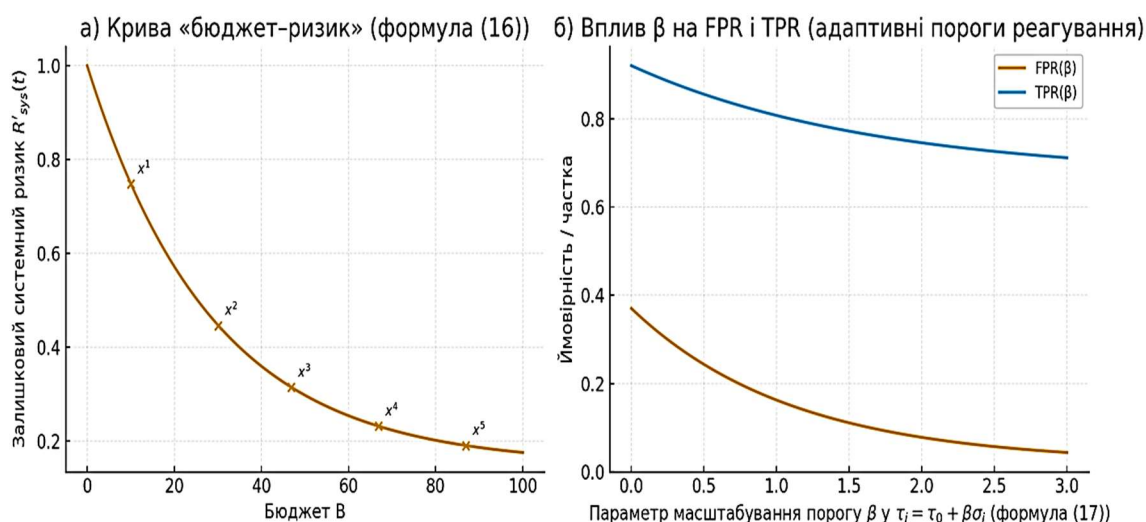


Рис. 5. Компроміс «бюджет-ризик» та вплив β на пороги реагування

Джерело: розроблено автором.

Розроблена система може бути інтегрована в наявну інфраструктуру центрів моніторингу безпеки (SOC) через стандартні канали обміну даними з SIEM-платформами, сховищами часових рядів (TSDB) та системами SOAR для автоматизації реагування. Інтелектуальні модулі формують аналітичні звіти, алерти та рекомендації з узгодженим SLA-часом реагування. Архітектура підтримує поетапне впровадження без зупинки поточних бізнес-процесів, що дає змогу поступово розширювати функціонал системи та підвищувати рівень кіберстійкості підприємства.

У процесі розроблення враховувались можливі обмеження моделі, зокрема дисбаланс класів подій, дрейф концепції у часових рядах і шум сенсорних даних. Для пом'якшення цих факторів застосовувалося ребалансування вибірок, моніторинг дрейфу з автоматичними тригерами перенавчання та перевірки консистентності ознак під час оновлення даних. Основним обмеженням є відсутність публічних сирих даних, однак подані процедурні описи достатні для відтворення результатів на будь-яких сумісних журналах безпеки.

Усі журнали подій, використані для тестування, проходили етап анонімізації та агрегування, що виключає можливість ідентифікації реальних об'єктів або користувачів. Доступ до даних регламентовано ролями з багаторівневою автентифікацією. Передбачено аудит рішень штучного інтелекту й захист від отруєння даних (data poisoning) через перевірку джерел і ізоляцію підозрілих потоків, що забезпечує прозорість та надійність функціонування системи.

Отримані результати підтвердили, що запропонована система забезпечує автоматизовану генерацію сценаріїв, динамічне оновлення ризикових профілів і пояснюваність результатів через ХАІ-аналіз. Подальший розвиток передбачає створення цифрових двійників об'єктів критичної інфраструктури з вбудованими АІ-модулями оцінки ризику, застосування мультиагентних моделей для імітації поведінки атак і розширення механізмів ХАІ для візуалізації впливу кожного параметра на інтегральний рівень загрози.

Висновки. У статті розроблено та представлено інтелектуальну систему сценарного моделювання оцінки захищеності об'єктів критичної інфраструктури, що поєднує методи машинного навчання, нечіткої логіки, графових структур і пояснюваного штучного інтелекту (ХАІ-модуля). Запропонована математична модель дозволяє формалізувати взаємозв'язки між активами, загрозами, уразливостями та контрзаходами, а також враховувати часову динаміку ризиків. Завдяки цьому забезпечується адаптивна оцінка стану захищеності в реальному часі, прогнозування сценаріїв розвитку інцидентів і вибір оптимальних контрзаходів з урахуванням обмежених ресурсів.

Експериментальні результати продемонстрували здатність системи автоматично оновлювати профіль ризику, мінімізувати похибку прогнозування та підтримувати пояснюваність результатів для аналітиків SOC через ХАІ-модуль. Розроблена архітектура може бути інтегрована у цифрові двійники кіберзахисту або центри моніторингу безпеки критичної інфраструктури, що підвищує рівень кіберстійкості та оперативність реагування на інциденти.

Подальші дослідження доцільно спрямувати на розширення функціоналу системи шляхом інтеграції мультиагентних моделей поведінки атак, удосконалення механізмів байєсівського оновлення знань і створення візуальних ХАІ-інтерфейсів для аналізу впливу окремих параметрів на інтегральний рівень загроз. Реалізація таких підходів сприятиме формуванню нового покоління інтелектуальних систем кіберзахисту, здатних до самонавчання, самокорекції та випереджувального реагування в умовах динамічних кіберзагроз.

Список використаних джерел

1. Boamah, F. A., Jin, X., Senaratne, S., & Perera, S. (2025). *AI-driven risk identification model for infrastructure projects: Utilising past project data*. *Expert Systems with Applications*, 233, 127891. <https://doi.org/10.1016/j.eswa.2025.127891>.
2. Ajayi, A., Akerele, J., Odio, P., Collins, A., Babatunde, G., & Mustapha, D. (2025). *Using AI and machine learning to predict and mitigate cybersecurity risks in critical infrastructure*. *Journal of Information Security Research*, 21, 204-224. <https://doi.org/10.13140/RG.2.2.14069.49120>.
3. Yigit, Y., Ferrag, M. A., Ghanem, M. C., Sarker, I. H., Maglaras, L. A., Chrysoulas, C., Moradpoor, N., Tihanyi, N., & Janicke, H. (2025). *Generative AI and LLMs for critical infrastructure protection: Evaluation benchmarks, agentic AI, challenges, and opportunities*. *Sensors*, 25(6), 1666. <https://doi.org/10.3390/s25061666>.
4. Paulraj, J., Raghuraman, B., Gopalakrishnan, N., & Otoum, Y. (2025). *Autonomous AI-based cybersecurity framework for critical infrastructure: Real-time threat mitigation*. *arXiv preprint*. <https://doi.org/10.48550/arXiv.2507.07416>.
5. Arif, M. H., Rabby, H. R., Nadia, N. Y., Tanvir, M. I. M., & Masum, A. A. (2025). *AI-driven risk assessment in national security projects: Investigating machine learning models to predict and mitigate risks in defense and critical infrastructure projects*. *Journal of Computer Science and Technology Studies*, 7(2), 71-85. <https://doi.org/10.32996/jcsts.2025.7.2.6>.
6. Гречанинов, В. Ф. (2025). *Моделі та технології інтелектуального захисту інформаційних систем критичної інфраструктури для підвищення стійкості*. *Кібербезпека: освіта, наука, техніка*, 1(29), 877-896. <https://doi.org/10.28925/2663-4023.2025.29.948>.
7. Hulak, H. M., Zhylytsov, O. B., Kyrychok, R. V., Korshun, N. V., & Skladannyi, P. M. (2023). *Enterprise information and cyber security*. Borys Grinchenko Kyiv Metropolitan University.

8. Mohamed, N. (2025). *Artificial intelligence and machine learning in cybersecurity: A deep dive into state-of-the-art techniques and future paradigms*. *Knowledge and Information Systems*, 67, 6969-7055. <https://doi.org/10.1007/s10115-025-02429-y>.
9. Kostiuk, Y. V., Skladannyi, P. M., Bebeshko, B. T., Khorolska, K. V., Rzaieva, S. L., & Vorokhob, M. V. (2025). *Information and communication systems security*. Borys Grinchenko Kyiv Metropolitan University.
10. Govea, J., Gaibor-Naranjo, W., & Villegas-Ch, W. (2024). *Transforming cybersecurity into critical energy infrastructure: A study on the effectiveness of artificial intelligence*. *Systems*, 12(5), 165. <https://doi.org/10.3390/systems12050165>.
11. 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).
12. Kostiuk, Y. 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.
13. Munmun, Z., Hossain, M. D., Jahan, I., Akther, S., Masum, A., & Rabby, H. (2025). *AI-driven project risk management: Leveraging artificial intelligence to predict, mitigate, and manage project risks in critical infrastructure and national security projects*. *Journal of Computer Science and Technology Studies*, 7(2), 71-85. <https://doi.org/10.32996/jcsts.2025.7.2.6>.
14. 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).
15. Islam, S., Basheer, N., & Papastergiou, S. (2025). *Intelligent dynamic cybersecurity risk management framework with explainability and interpretability of AI models for enhancing security and resilience of digital infrastructure*. *Journal of Reliable Intelligent Environments*, 11(1), 12. <https://doi.org/10.1007/s40860-025-00253-3>.
16. Kostiuk, Y., Skladannyi, P., Sokolov, V., & Rzaieva, S. (2025). *Intelligent system for simulation modeling and research of information objects*. In *Proceedings of the 1st Workshop Software Engineering and Semantic Technologies (SEST 2025)*, co-located with the *15th International Scientific and Practical Programming Conference (UkrPROG 2025)* (Vol. 4053, pp. 237-251).
17. Abramov, V., Astafieva, M., Boiko, M., Bodnenko, D., Bushma, A., Vember, V., Hlushak, O., Zhylytsov, O., Ilich, L., Kobets, N., Kovaliuk, T., Kuchakovska, H., Lytvyn, O., Lytvyn, P., Mashkina, I., Morze, N., Nosenko, T., Proshkin, V., Radchenko, S., & Yaskevych, V. (2021). *Theoretical and practical aspects of the use of mathematical methods and information technology in education and science*. <https://doi.org/10.28925/9720213284km>.

References

1. Boamah, F. A., Jin, X., Senaratne, S., & Perera, S. (2025). *AI-driven risk identification model for infrastructure projects: Utilising past project data*. *Expert Systems with Applications*, 233, 127891. <https://doi.org/10.1016/j.eswa.2025.127891>.
2. Ajayi, A., Akerele, J., Odio, P., Collins, A., Babatunde, G., & Mustapha, D. (2025). *Using AI and machine learning to predict and mitigate cybersecurity risks in critical infrastructure*. *Journal of Information Security Research*, 21, 204-224.
3. Yigit, Y., Ferrag, M. A., Ghanem, M. C., Sarker, I. H., Maglaras, L. A., Chrysoulas, C., Moradpoor, N., Tihanyi, N., & Janicke, H. (2025). *Generative AI and LLMs for critical infrastructure protection: Evaluation benchmarks, agentic AI, challenges, and opportunities*. *Sensors*, 25(6), 1666. <https://doi.org/10.3390/s25061666>.
4. Paulraj, J., Raghuraman, B., Gopalakrishnan, N., & Otoum, Y. (2025). *Autonomous AI-based cybersecurity framework for critical infrastructure: Real-time threat mitigation*. *arXiv preprint*. <https://doi.org/10.48550/arXiv.2507.07416>.

5. Arif, M. H., Rabby, H. R., Nadia, N. Y., Tanvir, M. I. M., & Masum, A. A. (2025). *AI-driven risk assessment in national security projects: Investigating machine learning models to predict and mitigate risks in defense and critical infrastructure projects*. *Journal of Computer Science and Technology Studies*, 7(2), 71-85. <https://doi.org/10.32996/jcsts.2025.7.2.6>.
6. Hrechanyinov, V. (2025). Models and technologies of intelligent protection of information systems of critical infrastructure for increasing resilience. [Modeli ta tekhnolohii intelektualnoho zakhystu informatsiinykh system krytychnoi infrastruktury dlia pidvyshchennia stiikosti.] *Kiberbezpeka: osvita, nauka, tekhnika – Cybersecurity: Education, Science, Technology*, 1(29), 877-896. <https://doi.org/10.28925/2663-4023.2025.29.948>.
7. Hulak, H. M., Zhylytsov, O. B., Kyrychok, R. V., Korshun, N. V., & Skladannyi, P. M. (2023). *Enterprise information and cyber security*. Borys Grinchenko Kyiv Metropolitan University.
8. Mohamed, N. (2025). *Artificial intelligence and machine learning in cybersecurity: A deep dive into state-of-the-art techniques and future paradigms*. *Knowledge and Information Systems*, 67, 6969-7055. <https://doi.org/10.1007/s10115-025-02429-y>.
9. Kostiuk, Y. V., Skladannyi, P. M., Bebesko, B. T., Khorolska, K. V., Rzaieva, S. L., & Vorokhob, M. V. (2025). *Information and communication systems security*. Borys Grinchenko Kyiv Metropolitan University.
10. Govea, J., Gaibor-Naranjo, W., & Villegas-Ch, W. (2024). *Transforming cybersecurity into critical energy infrastructure: A study on the effectiveness of artificial intelligence*. *Systems*, 12(5), 165. <https://doi.org/10.3390/systems12050165>.
11. Kostiuk, Y., Skladannyi, P., Samoilenko, Y., Khorolska, K., Bebesko, 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).
12. Kostiuk, Y. V., Skladannyi, P. M., Hulak, H. M., Bebesko, B. T., Khorolska, K. V., & Rzaieva, S. L. (2025). *Information security systems* [Textbook]. Kyiv: Borys Grinchenko Kyiv Metropolitan University.
13. Munmun, Z., Hossain, M. D., Jahan, I., Akther, S., Masum, A., & Rabby, H. (2025). *AI-driven project risk management: Leveraging artificial intelligence to predict, mitigate, and manage project risks in critical infrastructure and national security projects*. *Journal of Computer Science and Technology Studies*, 7(2), 71-85. <https://doi.org/10.32996/jcsts.2025.7.2.6>.
14. 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).
15. Islam, S., Basheer, N., & Papastergiou, S. (2025). *Intelligent dynamic cybersecurity risk management framework with explainability and interpretability of AI models for enhancing security and resilience of digital infrastructure*. *Journal of Reliable Intelligent Environments*, 11(1), 12. <https://doi.org/10.1007/s40860-025-00253-3>.
16. Kostiuk, Y., Skladannyi, P., Sokolov, V., & Rzaieva, S. (2025). *Intelligent system for simulation modeling and research of information objects*. In *Proceedings of the 1st Workshop Software Engineering and Semantic Technologies (SEST 2025)*, co-located with the *15th International Scientific and Practical Programming Conference (UkrPROG 2025)* (Vol. 4053, pp. 237-251).
17. Abramov, V., Astafieva, M., Boiko, M., Bodnenko, D., Bushma, A., Vember, V., Hlushak, O., Zhylytsov, O., Ilich, L., Kobets, N., Kovalyuk, T., Kuchakovska, H., Lytvyn, O., Lytvyn, P., Mashkina, I., Morze, N., Nosenko, T., Proshkin, V., Radchenko, S., & Yaskewych, V. (2021). *Theoretical and practical aspects of the use of mathematical methods and information technology in education and science*. <https://doi.org/10.28925/9720213284km>.

Отримано 20.09.2025

Viktor Grechaninov

Associate Professor, Senior Researcher, Head of the Research Department
Institute for Problems of Mathematical Machines and Systems of the National Academy of Sciences of Ukraine (Kyiv, Ukraine)
E-mail: vgrechaninov@gmail.com. ORCID: <https://orcid.org/0000-0001-6268-3204>

INTELLIGENT SYSTEM FOR SCENARIO MODELING OF SECURITY ASSESSMENT OF CRITICAL INFRASTRUCTURE OBJECTS

The urgency of the research arises from the growing sophistication and frequency of cyberattacks targeting critical infrastructure, where conventional static risk assessment techniques cannot reflect the dynamic interdependencies between assets, threats, vulnerabilities, and countermeasures. These challenges highlight the need for an intelligent and adaptive system capable of predicting, explaining, and mitigating evolving cyber risks in real time.

The problem addressed in this study lies in the absence of comprehensive models that integrate temporal forecasting, uncertainty quantification, and explainable decision-making within cybersecurity operations. Existing solutions are often static, fragmented, and incapable of generating adaptive risk scenarios or ensuring transparency in AI-driven reasoning. The main objective of the research is to develop an intelligent system for scenario modeling of security assessment for critical infrastructure objects, which provides dynamic evaluation of protection effectiveness and supports optimal resource allocation for risk mitigation.

The proposed model integrates machine learning, fuzzy logic, and graph theory to represent causal dependencies among security entities. The system architecture includes a data layer responsible for event aggregation and normalization, an AI layer using LSTM and XAI modules for forecasting and interpretability, and a decision layer for visualization and integration with SOC, SIEM, and SOAR platforms. Experimental results have shown improved accuracy of risk prediction, adaptive threshold adjustment, and significant reduction of false positives.

The conclusions emphasize that the developed system forms a comprehensive, explainable, and resource-efficient approach to cybersecurity risk management. It enables proactive protection, continuous self-learning, and transparency in evaluating defense performance. The research outcomes can enhance the cyber resilience of critical infrastructure and facilitate digital twin-based intelligent security monitoring.

Keywords: risk assessment; scenario modeling; artificial intelligence; critical infrastructure; fuzzy logic; risk graph; machine learning; cybersecurity; digital twins.

Fig.: 5. References: 17.