

<https://doi.org/10.31891/2307-5732-2025-359-80>
УДК 004.056

ЯМНИЧ АНДРІЙ

Національний університет «Львівська політехніка»

<https://orcid.org/0009-0005-7226-1896>

e-mail: andrii.b.yamnych@lpnu.ua

КОРОБЕЙНИКОВА ТЕТЯНА

Національний університет «Львівська політехніка»

<https://orcid.org/0000-0003-2487-8742>

e-mail: tetianakorobeinikova@gmail.com

ФОРМАЛІЗОВАНИЙ ОПИС ПРОЦЕСНОЇ МОДЕЛІ ДИНАМІЧНОГО АНАЛІЗУ ТА ПРОГНОЗУВАННЯ РИЗИКІВ ІНФОРМАЦІЙНОЇ БЕЗПЕКИ ДЛЯ ПЕРСОНАЛУ

У статті розглянуто проблематику динамічного оцінювання та прогнозування ризиків інформаційної безпеки (ІБ), спричинену зростанням ролі людського фактора, зокрема персоналу, в умовах цифровізації, гібридної роботи та мінливого контексту доступу. Запропоновано процесну модель, у якій послідовність функцій f_1 - f_8 безперервно перетворює технічні й поведінкові дані на ризик-адаптивні рішення управління доступом, навчання та реагування. Модель передбачає багатовимірну класифікацію ресурсів, формування вектора ознак Q , побудову цифрового двійника користувача з одержанням матриці ризиків R , генерацію контрзаходів та персоналізованих політик, збір зворотного зв'язку F_{back} і самонавчання. Інтеграція RBAC-блокчейну забезпечує прозорість, незмінність і відтворюваність аудиту доступу, що узгоджується з принципами Zero Trust.

Ключові слова: динамічне оцінювання ризику (DRA); цифровий двійник; RBAC-блокчейн; UEBA; Zero Trust; ризик-адаптивні політики доступу.

YAMNYCH ANDRII

KOROBEGINIKOVA TETIANA

Lviv Polytechnic National University

FORMALIZED DESCRIPTION OF THE PROCESS MODEL FOR DYNAMIC ANALYSIS AND PREDICTION OF INFORMATION SECURITY RISKS FOR PERSONNEL

This article addresses the challenge of dynamically assessing and forecasting information-security risks stemming from the human factor amid accelerated digitalization, hybrid work patterns, and evolving access contexts. We propose a process model in which a sequence of functions (f_1 - f_8) continuously transforms technical and behavioral evidence into risk-adaptive decisions for access governance, guidance, and response-closing the loop with verifiable auditability and self-learning. The model starts with a multidimensional resource-classification matrix (f_1), proceeds with the acquisition and unification of behavioral/technical signals (f_2), and produces a normalized feature vector Q (f_3). A user digital twin (f_4) runs "what-if" simulations to estimate a probability matrix R over threat classes while an RBAC-blockchain immutably records access transactions. Based on R , the system generates adaptive countermeasures and personalized policies and training (f_5 - f_6), collects feedback F_{back} on effectiveness and behavioral change (f_7), and updates weights, models, and RBAC rules (f_8).

The approach conforms to Zero Trust principles by replacing implicit trust with continuous validation of subjects, devices, and requests, incorporating context and risk. We introduce a thresholding scheme that activates preventive, detective, or corrective controls according to asset criticality and predicted risk; we also outline fine-tuning and transfer-learning procedures to keep models current without excessive computational cost. Personalized dashboards and multichannel delivery reduce the "risk window," whereas qualitative feedback (e.g., content clarity and user satisfaction) exposes elements of security culture.

The proposed model establishes an "analysis-forecast-action-feedback-self-correction" cycle that improves risk-assessment accuracy, enhances response timeliness, and advances transparency in access governance via blockchain-backed audit trails. The results are directly integrable with SIEM/UEBA ecosystems and enterprise access-management platforms and can support organization-wide cyber-literacy programs. By combining classical statistics, modern machine learning, digital-twin simulation, and distributed-ledger auditability within a single engineered workflow, the model delivers an interpretable and evolvable pathway to human-centric, risk-adaptive security in corporate environments.

Keywords: dynamic risk assessment (DRA); digital twin; RBAC-blockchain; UEBA; Zero Trust; risk-adaptive access policies;

Стаття надійшла до редакції / Received 25.09.2025

Прийнята до друку / Accepted 15.11.2025

Постановка проблеми

Швидка цифровізація бізнес-процесів, перехід до гібридних форматів роботи та фрагментація ІТ-ландшафтів призвели до різкого зростання ролі людського фактора в інцидентах ІБ [1-3]. Статичні моделі контролю доступу й періодичні аудити не відповідають динаміці середовища. Водночас регуляторні та галузеві стандарти наголошують на необхідності постійної верифікації суб'єктів і запитів, відмові від імпліцитної довіри та прозорості аудиту (Zero Trust) [4-7]. У практиці підприємств часто відсутній цілісний процес, який би поєднував класифікацію активів, багатоканальний збір поведінкових/технічних сигналів, їх нормалізацію до аналітичного вектора ознак, прогноз ризикових сценаріїв із використанням цифрового двійника та незмінний облік транзакцій доступу [8]. Наявні рішення фрагментарні: або фокусуються на технічних логах, ігноруючи поведінкові показники, або застосовують складні ML-моделі без пояснюваності та відтворюваності аудиту.

Таким чином, виникає потреба у процесній моделі, яка 1) формалізує шлях від даних до управлінських дій; 2) забезпечує безперервний перерахунок ризику та адаптацію політик; 3) фіксує всі зміни доступу у верифікованому реєстрі; 4) вбудовує механізми самонавчання на основі вимірюваного зворотного зв'язку. Окремим викликом є персоналізація: політики мають враховувати роль, контекст і індивідуальний профіль

ризик користувача, щоби мінімізувати як хибнопозитивні спрацювання (надмірні обмеження), так і хибнонегативні (пропущені загрози). Потрібно також скоротити "вікно ризику" за рахунок оперативної доставки рекомендацій і навчального контенту саме тим співробітникам, чії дії формують поточний ризиковий фон.

Запропонована у статті процесна модель адресує зазначені розриви, об'єднуючи в один керований ланцюг класифікацію активів, UEBA-сигнали, цифровий двійник користувача, блокчейн-аудит і замкнений контур самонавчання. Формалізоване подання даних (Q , R , F_{back}), порогові схеми активації контрзаходів, а також інтеграція з існуючими SIEM/UEBA/IAM-сервісами створюють основу для інженерно керованої, масштабованої та прозорої системи ризик-адаптивного доступу.

Аналіз останніх джерел

Сучасні огляди підтверджують, що статичні методики оцінювання ризику не відповідають швидкості змін у хмарних та мережевих середовищах. Систематичний огляд авторів Cheimonidis & Rantos (2023) проаналізував 50 моделей DRA та показав переваги безперервного перерахунку ризикових показників на основі потоків подій, контекстних атрибутів і ML/BN-методів у (майже) реальному часі [9]. Це корелює з потребою ризик-адаптивних політик та інтеграції з каналами розвідданих і Zero Trust-архітектурою.

Рамка Zero Trust, формалізована в NIST SP 800-207 (2020), визначає відмову від імпліцитної довіри та постійну верифікацію суб'єктів, активів і запитів на доступ з урахуванням контексту та ризику; додаткові настанови щодо реалізації політик на рівні застосунків наведено в SP 800-207A (2023) [10]. Такі документи засвідчують рух індустрії до гранульованого, ризик-орієнтованого контролю доступу, що безпосередньо підтримує функції f_4 - f_6 запропонованої моделі.

UEBA підтверджено як ефективний інструмент для профілювання користувачів та безперервної автентифікації. Зокрема, Martín та ін. (2022) показано, що комбінування динаміки клавіатури та миші на рівні ознак підсилює стабільність детекції [11], що узгоджується з етапами f_2 - f_3 (нормалізація й побудова Q).

Використання цифрових двійників (DT) для моделювання поведінки та "what-if"-аналізу у безпеці отримало систематизацію в огляді авторів Alcaraz & Lopez (2022). Автори описують загрози DT і одночасно підкреслюють прогностичну користь симуляцій [12], що співпадає з нашим застосуванням DT у функції f_4 для отримання матриці ймовірностей R . (NICS Lab)

На рівні інсайдерських загроз актуальний систематичний огляд Al-Mhiqani, M. N. та ін. [13] (2024) фіксує потребу в мультиджерельних технічних і поведінкових сигналах і вдосконалених алгоритмах аналізу, що підтримує рішення типу UEBA+DT у запропонованій моделі.

Незмінність журналів доступу та відтворюваність аудиту посилюються блокчейн-підходами. NIST IR 8403 окреслює застосування блокчейну для контролю доступу, підкреслюючи переваги децентралізації, простежуваності та тампер-резистентності. На рівні оглядових праць Namane та ін. (2022) [14] систематизують підходи до блокчейн-керованого доступу в IoT; Ullah та ін. (2023) [15] узагальнюють стан досліджень BE-ABAC, а Punia та ін. (2024) [16] виокремлюють 12 парадигм блокчейн-AC для хмарних середовищ. Ці результати обґрунтовують заміну "централізованого журналювання" на розподілений реєстр у нашому RBAC-блокчейні (f_4 , f_8) [17].

Запропонована нами у [18] багатовимірна матриця класифікації ресурсів для оцінки ризиків ІБ (2024), де обґрунтовано підхід до формування ваг і ознак, безпосередньо лягає в основу f_1 ; а робота про інтеграцію RBAC і блокчейну для підвищення прозорості керування доступом (2024) [17] підтверджує практичну релевантність блокчейн-аудиту у корпоративних системах. Додатково, огляд концепції нульової довіри в українському академічному виданні відображає коректність впровадження ZTA-підходів у корпоративні мережі [7].

Підсумовуючи, сучасний науковий ландшафт підтверджує необхідність DRA у ZTA-рамці, доцільність UEBA для безперервної автентифікації та інсайдерського моніторингу, прогностичний потенціал цифрових двійників, а також доречність блокчейн-механізмів для незмінного аудиту та автоматизації політик.

Тема дослідження є актуальною, оскільки інтенсифікація загроз, пов'язаних із людським фактором, постійно зростає, водночас вимога Zero Trust до постійної верифікації та прозорого аудиту також росте. Розвиток технічних засобів для реалізації ризик-адаптивних рішень (UEBA, цифрові двійники, блокчейн-аудит) може забезпечити зниження ризиків і відтворювану основу для вдосконалення

Мета роботи полягає у підвищенні точності та своєчасності оцінювання ризиків для персоналу, покращенні керованості доступу через персоналізовані ризик-адаптивні політики та вдосконаленні прозорості й відтворюваності аудиту на основі незмінного реєстру транзакцій.

Виклад основного матеріалу

Процесна модель динамічного аналізу та прогнозування ризиків ІБ для персоналу - це керований набір пов'язаних процесів (що зреалізовані набором функцій f_1 - f_8), що безперервно перетворюють технічні та поведінкові дані користувачів на ризик-адаптивні рішення з доступу, навчання й реагування, із замкненим контуром самонавчання та аудиту (рис. 1).

Основними задачами запропонованої процесної моделі динамічного аналізу та прогнозування ризиків інформаційної безпеки для персоналу є: безперервне оцінювання та прогнозування ризиків, зумовлених діями персоналу; персоналізація політик доступу та навчального контенту під актуальний профіль ризику; оперативне реагування (контрзаходи) й аудит усіх змін і транзакцій доступу; самонавчання системи (оновлення ваг, моделей, правил доступу).

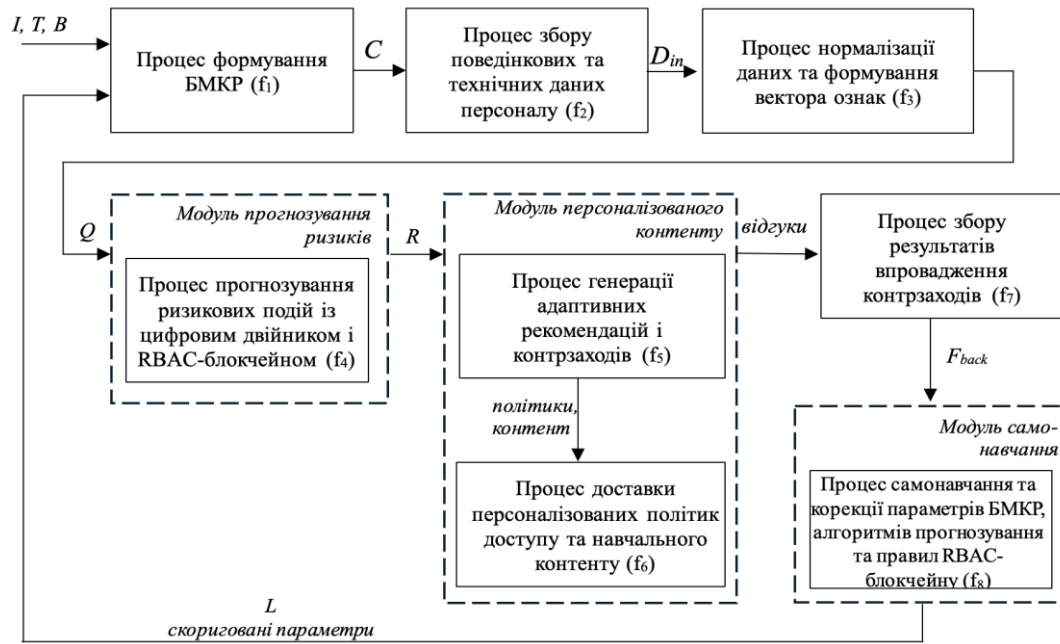


Рис. 1 - Процесна модель динамічного аналізу та прогнозування ризиків інформаційної безпеки для персоналу

Ключовими характеристиками процесної моделі є:

- модульність (забезпечений функціями f_1 - f_8), контекстність і ризик-адаптивність (Zero Trust-підхід) [6-7];
- реальний час і проактивність (прогноз, не лише фіксація);
- персоніфікація (вектор ознак Q , цифровий двійник);
- верифікованість і простежуваність (RBAC-блокчейн) [17];
- масштабованість та інтегрованість (стандартизовані формати даних);
- безперервне вдосконалення (матриця значень F_{back} , що активує модуль самонавчання).

Складові процесної моделі динамічного аналізу та прогнозування ризиків інформаційної безпеки для персоналу: функція, процес, що вона підтримує, результат її роботи та взаємозв'язок ключових функцій - наведено у таблиці 1.

Таблиця 1

Складові процесної моделі динамічного аналізу та прогнозування ризиків інформаційної безпеки для персоналу

Функція	Процес	Результат
f_1	формування багатовимірної матриці класифікації ресурсів [18]	Багатовимірна матриця класифікації ресурсів. Структурує активи за ознаками (критичність, конфіденційність, цінність тощо) і задає ваги/пороги для подальших рішень.
f_2	збору поведінкових та технічних даних персоналу	Збір поведінкових і технічних даних. Агрегує журнали, телеметрію, опитування; формує «сирий» масив подій на рівні користувача/сесії/пристрою.
f_3	нормалізації даних та формування вектора ознак	Нормалізація даних та формування вектора ознак Q . Очищає, уніфікує, масштабує дані; створює Q як стандартизоване подання профілю користувача (з відбором інформативних ознак).
f_4	прогнозування ризикових подій із цифровим двійником і RBAC-блокчейном [17]	Прогноз ризикових подій (цифровий двійник + RBAC-блокчейн). Буде цифровий двійник користувача, оцінює ймовірності загроз (матриця R); фіксує транзакції доступу в блокчейні.
f_5	генерації адаптивних рекомендацій і контрзаходів та	Генерація адаптивних контрзаходів. Перетворює R на конкретні дії (превентивні/ детективні/коригувальні) з урахуванням порогів і класу активу (з f_1).
f_6	доставки персоналізованих політик доступу та навчального контенту	Доставка персоналізованих політик і навчання. Надає користувачу оновлені правила доступу й освітній контент, скорочуючи «вікно ризику».
f_7	збору результатів впровадження контрзаходів	Збір результатів впровадження (F_{back}). Вимірює ефективність контрзаходів/навчання (кількісні й якісні метрики), фіксує поведінкові зміни.
f_8	самонавчання та корекції параметрів БМКР, алгоритмів прогнозування та правил RBAC-блокчейну	Самонавчання і корекція параметрів/правил. Оновлює ваги, моделі, RBAC-правила з урахуванням F_{back} ; результати повертаються до f_1 (векторизація ознак/ваги) та f_4 (прогнозу), закривши цикл.

Взаємозв'язок у цілісний ланцюг: f_1 забезпечує контекст і ваги для f_4 - f_6 та оновлюється через f_8 . $f_1 \rightarrow f_2 \rightarrow f_3 (Q) \rightarrow f_4 (R) \rightarrow f_5$ (контрзаходи) $\rightarrow f_6$ (доставка) $\rightarrow f_7 (F_{\text{back}}) \rightarrow f_8$ (оновлення).

З метою перетворити запропоновані підходи на керовану інженерну систему, яку можна вимірювати, порівнювати, сертифікувати й безпечно еволюціонувати, пропонується формалізований опис моделі динамічного аналізу та прогнозування ризиків інформаційної безпеки для персоналу.

Впровадження формалізованого опису дасть можливість забезпечувати таке:

- Відтворюваність і аудит. Чітко визначені входи/виходи, метрики та правила забезпечують прозорий трасувальний аудит (хто, коли, чому отримав/втратив доступ), що критично для нагляду та відповідності вимогам;
- Метричність і керованість. Формальні змінні (Q, R, F_{back}), пороги, функції втрат і KPI дозволяють вимірювати ефективність (FP/FN, час реакції, зниження ризику) і обґрунтовано їх оптимізувати;
- Інтєроперабельність і масштабування. Уніфіковані інтерфейси/формати дають змогу інтегрувати модель із SIEM/UEBA/IAM та переносити її між підрозділами/організаціями;
- Перевірюваність і валідація. Можна коректно проводити what-if-симуляції (через цифровий двійник), порівнювати алгоритми, виконувати А/В-тести й формальну оцінку ризику;
- Зменшення неоднозначностей. Формальна специфікація усуває різночитання між безпекою, ІТ та бізнесом, чітко розмежовує ролі, дані, політики та відповідальність;
- Підтримка Zero Trust. Формалізація контекстних правил і перевірок робить ризик-адаптивні рішення відтворюваними, а їх зміни - контрольованими та обґрунтованими.

Вхідними даними для процесу формування багатовимірної матриці класифікації ресурсів (БМКЛ) є сукупність множин: I - інформаційні ресурси підприємства, T - технічні параметри, отримані з лог-файлів та телеметрії, і B - поведінкові метрики, зібрані в процесі анкетування й опитування персоналу (1).

$$I = \{i_1, i_2, \dots, i_{n_I}\}, \quad T = \{t_1, t_2, \dots, t_{n_T}\}, \quad B = \{b_1, b_2, \dots, b_{n_B}\} \quad (1)$$

Кожен елемент $i_k \in I$ характеризується атрибутами конфіденційності, цінності, доступності та цілісності. Параметри $t_k \in T$ відображають низькорівневі технічні показники (IP-адреса, порт, обсяг переданих даних тощо), а $b_k \in B$ фіксують поведінкові ознаки (час заходів, частота повторних входів у систему, відхилення від добових шаблонів активності). Первинні вагові коефіцієнти ознак зберігаються в множині (2):

$$W_0 = \{w_{ij}^0 \mid i_j \in I, j = 1, \dots, m\}, \quad (2)$$

де w_{ij}^0 - початковий ваговий коефіцієнт для j -ої ознаки i_j .

Функція f_1 формує багатовимірну матрицю класифікації $C \in \mathbb{R}^{m \times k}$, де $m = |I|$ - кількість класифікованих ресурсів, а k - число обраних класифікаційних ознак. Формально (3) та (4):

$$C = f_1(I, W_0), \quad C = [c_{jk}], \quad c_{jk} = w_{jk} \cdot \phi(i_j), \quad (3)$$

$$\phi(i_j) = \text{norm}(\psi(i_j)), \quad \psi(i_j) = (\text{conf}_{i_j}, \text{val}_{i_j}, \text{avail}_{i_j}, \text{integ}_{i_j}), \quad (4)$$

де оператор norm приводить початкові оцінки $\psi(i_j)$ до єдиної шкали, а w_{jk} - підсумкові вагові коефіцієнти, отримані внаслідок поєднання експертних оцінок і кореляційного аналізу інцидентів безпеки ψ .

Функція f_2 відповідає за збори даних у єдину вхідну множину (5):

$$D_{\text{in}} = D_{\text{logs}} \cup D_{\text{tele}} \cup D_{\text{survey}}, \quad (5)$$

де (див. вираз (6) нижче):

$$D_{\text{logs}} = \{(\tau_\ell, p_\ell) \mid \ell = 1, \dots, N_\ell\}, \quad D_{\text{tele}} = \{(\tau_t, \theta_t)\}, \quad D_{\text{survey}} = \{(\tau_s, \beta_s)\} \quad (6)$$

Кожен запис відображається як пара часової мітки τ та значення параметра (порт, IP, обсяг, метрика). На цьому етапі застосовуються процедури видалення дублікатів і некоректних записів, а також вирівнювання часових міток щодо єдиного UTC, що гарантує узгодженість і цілісність даних. Формально фільтрація задається як оператор (7):

$$D_{\text{in}}^* = \{d \in D_{\text{in}} \mid \text{valid}(d)\}, \quad (7)$$

де $\text{valid}(d)$ - булева функція коректності, що перевіряє формат, цілісність і відсутність аномальних перепадів у переданих значеннях.

Функція f_3 нормалізує відфільтровану множину D_{in}^* та формує вектор ознак (8):

$$Q \in \mathbb{R}^r, \quad Q = f_3(D_{\text{in}}^*), \quad (8)$$

де r - кількість відібраних ознак після процедур відбору. Нормалізація містить перетворення Min-Max (9),

$$x_{\text{norm}} = \frac{x - \min x}{\max x - \min x}, \quad (9)$$

а також z-оцінювання (10),

$$x_{\text{std}} = \frac{x - \mu_X}{\sigma_X}, \quad (10)$$

а також one-hot кодування дискретних категорій. Після розрахунку ознак $\xi_j \in Q$ їх відбирають за допомогою Random Forest Importance або методів відбору головних компонент (PCA), що зменшують розмір простору ознак до r змінних із мінімальною втратою інформації. Таким чином, (11):

$$Q = [\xi_1, \xi_2, \dots, \xi_r]^T, \quad \xi_j = \text{select}(\text{transform}(D_{\text{in}}^*)) \quad (11)$$

Функція f_4 реалізує прогнозування ризикових подій на основі цифрового двійника D_{tw} та механізму RBAC-блокчейну. Спочатку будується цифровий двійник як простір (12):

$$D_{\text{tw}} = \{\delta_1, \dots, \delta_s\}, \quad \delta_j = g_{\text{scenario}}(Q, j), \quad (12)$$

де $\mathcal{G}_{\text{scenario}}$ генерує дозвільні й недозвільні сценарії. Паралельно визначається функція контролю доступу (13):

$$g_{\text{RBAC}}: \text{Roles} \times \text{Perms} \rightarrow \{0,1\}, \quad (13)$$

згідно з якою кожна транзакція tr записується в блокчейн-ланцюжок $\mathcal{B}$. Прогнозна матриця R обчислюється:

$$R = f_4(Q, \mathcal{D}_{\text{tw}}, g_{\text{RBAC}}), \quad R \in [0,1]^{s \times p}, \quad (14)$$

де p – число типів загроз. Конкретно (15):

$$R_{jk} = h(Q, \delta_j, \theta_k), \quad h \in \{\text{LSTM, XGBoost, SVM}\}, \quad (15)$$

де θ_k – параметри моделі, навченої на історичних інцидентах, а запис у блокчейн гарантує незмінність журналів доступу [17].

Функція f_5 трансформує прогнозну матрицю R у набір рекомендацій і контрзаходів. Формалізація відбувається через оператор:

$$P_{\text{access}} = \Phi(R), \quad L_{\text{content}} = \Psi(R), \quad (16)$$

де Φ – мапінг ймовірностей у правила адаптивного доступу (обмеження, багатофакторна автентифікація), а Ψ – генерація навчального контенту за принципом «learning path». Кожен контрзахід $p_i \in P_{\text{access}}$ описується як кортеж (17):

$$p_i = (r_i, \delta_i), \quad (17)$$

де r_i – правило доступу (наприклад, «обмежити доступ до ресурсу i », «вимагати багатофакторної автентифікації»), а $\delta_i \in [0,1]$ – ступінь його критичності, визначений через ймовірність загрози R_{jk} . Паралельно генерується навчальний контент як множина модулів:

$$L_{\text{content}} = \{\ell_1, \ell_2, \dots, \ell_q\}, \quad (18)$$

де кожний модуль ℓ_j формується за правилом (19):

$$\ell_j = (c_j, \rho_j), \quad (19)$$

Де c_j – власне навчальний матеріал (наприклад, тестове завдання або інтерактивний кейс), а ρ_j – оцінка його пріоритетності, що обчислюється як (20):

$$\rho_j = \max_k (R_k) \cdot \omega_j, \quad (20)$$

де ω_j – коефіцієнт ефективності відповідного формату контенту, отриманий унаслідок апробації в підготовчих експериментах.

Функція f_6 доставляє комплект $(P_{\text{access}}, L_{\text{content}})$ користувачу через набір протоколів та інтерфейсів. Формально задати канали доставлення можна як множину (21):

$$H = \{h_{\text{UI}}, h_{\text{MSG}}, h_{\text{EMAIL}}, h_{\text{MOBILE}}\}, \quad (21)$$

де кожен канал $h \in H$ реалізує функцію (22):

$$H: (P_{\text{access}}, L_{\text{content}}) \mapsto \text{DeliveryReceipt}, \quad (22)$$

а результатом h є підтвердження доставлення DR , що містить часову мітку τ , ідентифікатор користувача u та статус виконання $s \in \{\text{передано, прочитано, підтверджено}\}$. Система агрегує ці підтвердження в множину (23):

$$DRS = \{DR_1, \dots, DR_n\}, \quad (23)$$

що слугує основою для оцінки оперативності доставлення та рівня залученості персоналу.

Функція f_7 збирає зворотний зв'язок як множину (24):

$$F_{\text{back}} = \{f_1^{\text{back}}, f_2^{\text{back}}, \dots, f_m^{\text{back}}\}, \quad (24)$$

де кожний елемент f_j^{back} є кортежем (25):

$$f_j^{\text{back}} = (\tau_j, u_j, r_j, \Delta b_j), \quad (25)$$

де τ_j – час отримання зворотного зв'язку, u_j – ідентифікатор користувача, r_j – оцінка ефективності контрзаходу (наприклад, проходження навчального модуля зі швидкістю v_j та правильністю p_j), а Δb_j – вектор змін у поведінкових ознаках, обчислений як (26):

$$\Delta b_j = b_j^{\text{post}} - b_j^{\text{pre}}, \quad (26)$$

де b_j^{pre} та b_j^{post} – значення відповідних метрик до і після впровадження контрзаходів.

Нарешті, функція f_8 відповідає за самонавчання та корекцію системи. Формально цей процес описується як ітераційне оновлення параметрів: $\Theta = \{w_{jk}, \theta_k, \dots\}$, де w_{jk} – вагові коефіцієнти матриці C , а θ_k – параметри моделей прогнозування. Механізм reinforcement learning задає правило (27):

$$\Theta_{t+1} = \Theta_t + \alpha \nabla \Theta L(\Theta; F_{\text{back}}), \quad (27)$$

де α – швидкість навчання, а L – функція втрат, що враховує помилку прогнозування та відхилення поведінки:

$$L(\Theta; F_{\text{back}}) = \sum_{j=1}^m \|h(Q_j; \Theta) - \Delta b_j\|^2 + \lambda \sum_k w_{jk} - w_{jk}^0\|^2, \quad (28)$$

де перший доданок мінімізує різницю між прогнозованою й реальною зміною поведінкових ознак, а другий вводить регуляризацию щодо початкових ваг W_0 із коефіцієнтом λ . Одночасно коригуються правила RBAC-блокчейну як множина транзакцій (29):

$$B_{t+1} = B_t \cup \{\text{new_rules}\}, \quad (29)$$

де new_rules генеруються на основі аналізу аномальних сценаріїв з Δb .

Запропонований формалізований опис демонструє цілісний перехід від первинних даних до безперервного циклу вдосконалення, упроваджуючи науково нові поєднання статистичних, машинно-навчальних і блокчейн-технологій для підвищення точності, швидкості реагування та адаптивності системи захисту інформаційних ресурсів персоналу.

Узагальнення формалізованого опису доповнюється наведенням конкретних прикладів матричних

структур, які ілюструють перетворення даних на різних етапах технологічного ланцюга. Нехай багатовимірною матрицею класифікації C набуває вигляду:

$$C = \begin{pmatrix} c_{11} & c_{12} & \dots & c_{1k} \\ c_{21} & c_{22} & \dots & c_{2k} \\ \vdots & \vdots & \ddots & \vdots \\ c_{m1} & c_{m2} & \dots & c_{mk} \end{pmatrix}, \quad (30)$$

де кожен елемент c_{jk} обчислюється як добуток нормованої оцінки ознаки та вагового коефіцієнта (31):

$$c_{jk} = w_{jk} \frac{\psi_{jk} - \min(\psi_{j\bullet})}{\max(\psi_{j\bullet}) - \min(\psi_{j\bullet})}, \quad (31)$$

а вектор початкових нормованих ознак $\psi_{j\bullet}$ формується за допомогою експертної функції ψ та статистичного оператора $norm$. Після збору й очищення даних утворюється матриця технічних і поведінкових подій (32):

$$D_{in}^* = \begin{pmatrix} \tau_1 & p_1 & b_1 \\ \tau_2 & p_2 & b_2 \\ \vdots & \vdots & \vdots \\ \tau_N & p_N & b_N \end{pmatrix}, \quad (32)$$

де кожний рядок містить часову мітку τ_i , технічний параметр p_i та поведінкову метрику b_i . Відповідні стовпці нормалізуються за схемою Min–Max та z-оцінювання, після чого кожний разом сформований вектор ознак Q подається як стовпцевий (33):

$$Q = \begin{pmatrix} \xi_1 \\ \xi_2 \\ \vdots \\ \xi_r \end{pmatrix}, \quad (33)$$

де $\xi_j = select(transform(D_{in}^*))$ – відфільтрована й нормована ознака.

У модулі прогнозування ризикових подій формується ймовірнісна матриця R розмірності $s \times p$, що відображає ймовірність кожного з p типів загроз у s симульованих сценаріях цифрового двійника (34):

$$R = \begin{pmatrix} R_{11} & R_{12} & \dots & R_{1p} \\ R_{21} & R_{22} & \dots & R_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ R_{s1} & R_{s2} & \dots & R_{sp} \end{pmatrix}, R_{jk} = h(Q, \delta_j, \theta_k). \quad (34)$$

Запис кожної транзакції доступу фіксується в розподіленому реєстрі B , що гарантує незмінність і прозорість, а сама матриця R слугує основою для обчислення ступеня критичності контрзаходів. Після генерації рекомендацій та політик утворюється набір контрзаходів (35):

$$P_{access} = \begin{pmatrix} p_1 & \delta_1 \\ p_2 & \delta_2 \\ \vdots & \vdots \\ p_q & \delta_q \end{pmatrix}, \quad (35)$$

у якому кожний рядок містить правило p_i і відповідний коефіцієнт важливості δ_i . Додатково генерується матриця навчальних модулів (36):

$$L_{content} = \begin{pmatrix} c_1 & \rho_1 \\ c_2 & \rho_2 \\ \vdots & \vdots \\ c_q & \rho_q \end{pmatrix}, \quad (36)$$

де c_j – контент, а ρ_j – його пріоритетність. У результаті доставлення й виконання контрзаходів формується матриця зворотного зв'язку (37):

$$F_{back} = \begin{pmatrix} \tau_1 & u_1 & r_1 & \Delta b_1 \\ \tau_2 & u_2 & r_2 & \Delta b_2 \\ \vdots & \vdots & \vdots & \vdots \\ \tau_m & u_m & r_m & \Delta b_m \end{pmatrix}, \quad (37)$$

що об'єднує часові мітки, ідентифікатори користувачів, оцінки ефективності контрзаходів та вектори змін у поведінці. За допомогою алгоритму (38)

$$\theta_{t+1} = \theta_t + \alpha \nabla_{\theta} \sum_{j=1}^m \|h(Q_j; \theta_t) - \Delta b_j\|^2 \quad (38)$$

відбувається корекція параметрів θ , а нові правила безпеки new_rules додаються до блокчейну B_{t+1} , забезпечуючи безперервне вдосконалення системи.

Таким чином, наведені матриці ілюструють повний ланцюг перетворення даних у межах архітектури системи оцінювання ризиків та управління доступом: від первинної матриці класифікації C , що визначає базову таксономію об'єктів, ролей і типів активності, через очищену й нормалізовану техніко-поведінкову матрицю D_{in}^* , що акумулює перевірені значення параметрів, до вектора ознак Q , який є аналітичним ядром цифрового профілю працівника.

Подальше перетворення цього вектора в матрицю ймовірностей R дає змогу оцінити ризиковий потенціал поведінки користувача в розрізі різних сценаріїв взаємодії з інформаційними ресурсами. На основі R система формує набір рекомендацій P_{access} , що визначають допустимі або обмежені рівні доступу до ресурсів, а також інформаційний вектор $L_{content}$, що відображає персоналізовані освітні або когнітивні втручання,

наприклад, запуск тренінгу, відображення інструкцій або ініціацію м'якого контролю.

Завершальним етапом цього циклу є побудова матриці зворотного зв'язку F_{back} , яка акумулює реакцію користувача на прийняті рішення: поведінкові зміни після застосування обмежень, ефективність навчальних впливів, дотримання нових політик доступу. Ця матриця не лише уможливорює оцінку результативності застосованих заходів, а й слугує основою для адаптивного оновлення профілю, забезпечуючи циклічність і самонавчальну здатність усієї системи.

Така багаторівнева, формалізована послідовність гарантує не лише чітке відстеження інформаційних потоків і ступенів ризику, але й логічну прозорість прийнятих рішень, верифікованість контрзаходів та інтеграцію поведінкових чинників у стратегію управління довірою. Це забезпечує не лише відповідність стандартам інформаційної безпеки, а й відкриває можливості для створення гнучкої, інтелектуально керованої системи, здатної до адаптації в умовах мінливої загрозової динаміки та трансформації організаційного середовища.

Висновки

Загалом, наукова новизна полягає у створенні комплексної адаптивної процесної моделі оцінювання ризиків інформаційної безпеки для персоналу, що поєднує багатовимірний аналіз технічних і поведінкових ознак користувачів, механізми блокчейн-аудиту транзакцій доступу та алгоритми глибинного машинного навчання в єдиний замкнутий контур управління. Запропонований підхід забезпечує динамічне формування та автоматичне коригування профілів доступу в режимі реального часу з урахуванням змін контексту використання інформаційних ресурсів та виявлених аномалій у поведінці користувачів. Така інтеграція класичних статистичних методів, моделей прогнозування на основі цифрових двійників і розподілених реєстрових технологій дає змогу не лише підвищити точність і швидкість реагування на загрози, але й формувати персоналізовані стратегії зниження ризиків, орієнтовані на конкретні сценарії взаємодії працівника з інформаційною системою.

Запропонована процесна модель забезпечує цілісний, формалізований і відтворюваний цикл управління ризиками, у якому дані про поведінку та технічний стан систем переводяться у вектор ознак Q , далі у прогноз ризиків R за допомогою цифрового двійника, після чого автоматично генеруються контрзаходи та персоналізовані політики доступу з їхньою оперативною доставкою користувачам. Інтеграція RBAC-блокчейну гарантує незмінність журналів доступу, прозорість змін і перевірюваність рішень, що відповідає принципам Zero Trust. Зворотний зв'язок F_{back} і модуль самонавчання знижують кількість хибних спрацьовувань і адаптують модель до змін контексту без значних обчислювальних витрат завдяки процедурам fine-tuning і transfer learning.

Отримані результати підвищують точність оцінювання ризику для персоналу (завдяки багатовимірному профілюванню та симуляціям), покращують оперативність реагування (через порогові схеми активації превентивних, детективних і коригувальних дій) та вдосконалюють прозорість управління доступом (завдяки блокчейн-аудиту). Персоналізовані панелі та мультиканальна доставка рекомендацій скорочують "вікно ризику", а якісні метрики зворотного зв'язку висвітлюють культурні аспекти безпеки. Сформовано основу для інтеграції з наявними SIEM/UEBA/IAM-екосистемами та для масштабування між підрозділами й організаціями.

У майбутніх роботах доцільно детальніше верифікувати ефект від різних комбінацій ознак у Q , дослідити вплив різних архітектур цифрових двійників на стабільність R та порівняти моделі журналювання на блокчейні з огляду на продуктивність і приватність. Загалом, представлена модель демонструє практичну життєздатність ризик-адаптивного, людиноцентричного підходу до безпеки доступу, відображаючи сучасні тенденції динамічного ризик-менеджменту та вимоги до відтворюваності аудиту в корпоративних мережах.

Література

1. Papanikolaou, A., Varvarousi, E., & Gavala, E. (2024). Postal sector digitalisation: Security and vulnerabilities. *International Journal of Applied Systemic Studies*, 11(1), 42–51. <https://doi.org/10.1504/IJASS.2024.139211>
2. Sèdes, F., & Degrace, J. (2024). Social engineering and security: From human vulnerabilities to malicious threats. In *2024 20th International Conference on Wireless and Mobile Computing, Networking and Communications (WiMob)* (pp. 301–305). IEEE. <https://doi.org/10.1109/WiMob61911.2024.10770451>
3. Korobeinikova, T., Tachenko, I., Romanyuk, O., Romanyuk, S., Stakhov, O., & Reyda, O. (2024). Assessing network security risks: A technological chain perspective. In *2024 14th International Conference on Advanced Computer Information Technologies (ACIT)* (pp. 565–570). IEEE. <https://doi.org/10.1109/ACIT62333.2024.10712586>
4. Alnafea, F. S. M., Sengar, A. S., Hamid, N. K., Selladurai, K. M., Hassan, S. I., & Saravanakumar, R. (2025). Improving multi-factor authentication security with deep learning-based user behaviour analysis. In *2025 3rd International Conference on Integrated Circuits and Communication Systems (ICICACS)* (pp. 1–5). IEEE. <https://doi.org/10.1109/ICICACS65178.2025.10968961>
5. Kaur, M., & Garg, P. (2025). Exploring behavioral patterns for security in cloud computing: A case study. In *2025 3rd International Conference on Advancement in Computation & Computer Technologies (InCACCT)* (pp. 720–725). IEEE. <https://doi.org/10.1109/InCACCT65424.2025.11011337>

6. Korobeinikova, T. I. (2025). The zero trust model: Theory, practice, and prospect. In *Heritage of European Science 2025: Innovative technology, computer science, cybernetics and automation, security systems, transport development, architecture and construction, physics and mathematics* (Monographic series “European Science,” Book 37, Part 1, pp. 126–147). European Science.
7. Коробейнікова, Т., Журавель, І., Бодак, А., & Борошенко, Д. (2024). Концепція нульової довіри: сучасні методи забезпечення кібербезпеки в корпоративних мережах [The zero trust concept: Modern methods for ensuring cybersecurity in corporate networks]. *Вісник Львівського державного університету безпеки життєдіяльності*, 30, 67–77. <https://journal.ldubgd.edu.ua/index.php/Visnuk/article/view/2769>
8. Terumalasetti, S., & Reeja, S. R. (2025). Enhancing social media user's trust: A comprehensive framework for detecting malicious profiles using multi-dimensional analytics. *IEEE Access*, 13, 7071–7093. <https://doi.org/10.1109/ACCESS.2024.3521951>
9. Cheimonidis, P., & Rantos, K. (2023). Dynamic risk assessment in cybersecurity: A systematic literature review. *Future Internet*, 15(10), 324. <https://doi.org/10.3390/fi15100324>
10. Rose, S., Borchert, O., Mitchell, S., & Connelly, S. (2020). *Zero Trust Architecture (NIST Special Publication 800-207)*. National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-207>
11. Martín, A. G., Martín-de-Diego, I., Fernández-Isabel, A., et al. (2022). Combining user behavioural information at the feature level to enhance continuous authentication systems. *Knowledge-Based Systems*, 244, 108544. <https://doi.org/10.1016/j.knosys.2022.108544>
12. Alcaraz, C., & Lopez, J. (2022). Digital twin: A comprehensive survey of security threats. *IEEE Communications Surveys & Tutorials*, 24(3), 1475–1503. <https://doi.org/10.1109/COMST.2022.3171465>
13. Al-Mhiqani, M. N., Alsabou, T. A. A., Al-Shehari, T., Abdulkareem, K. H., Ahmad, R., & Mohammed, M. A. (2024). Insider threat detection in cyber-physical systems: A systematic literature review. *Computers & Electrical Engineering*, 119, 109489. <https://doi.org/10.1016/j.compeleceng.2024.109489>
14. Namane, S., Derhab, A., Guerroumi, M., & Challal, Y. (2022). Blockchain-based access control techniques for IoT: A systematic review. *Electronics*, 11(14), 2225. <https://doi.org/10.3390/electronics11142225>
15. Ullah, S. S., Oleshchuk, V., & Pussewalage, H. S. G. (2023). A survey on blockchain-envisioned attribute-based access control for Internet of Things: Overview, comparative analysis, and open research challenges. *Computer Networks*, 235, 109994. <https://doi.org/10.1016/j.comnet.2023.109994>
16. Punia, A., Hoda, M., Kaushik, K., & Tomar, D. (2024). A systematic review on blockchain-based access control systems in cloud environment. *Journal of Cloud Computing*, 13, 62. <https://doi.org/10.1186/s13677-024-00697-7>
17. Ямнич, А. Б., & Коробейнікова, Т. І. (2024). Модель контролю доступу персоналу до інформаційних ресурсів підприємств на основі RBAC та технології BLOCKCHAIN [A personnel access control model for enterprise information resources based on RBAC and blockchain]. *Вісник Хмельницького національного університету*, 343(6(1)), 380–386. <https://doi.org/10.31891/2307-5732-2024-343-6-56>
18. Коробейнікова, Т. І., & Ямнич, А. Б. (2024). Багатовимірна матриця класифікації інформації для оцінки ризиків інформаційної безпеки [A multidimensional information-classification matrix for information security risk assessment]. *Інформаційні технології та комп'ютерна інженерія*, (2), 91–106.

References

1. Papanikolaou, A., Varvarousi, E., & Gavala, E. (2024). Postal sector digitalisation: Security and vulnerabilities. *International Journal of Applied Systemic Studies*, 11(1), 42–51. <https://doi.org/10.1504/IJASS.2024.139211>
2. Sèdes, F., & Degrace, J. (2024). Social engineering and security: From human vulnerabilities to malicious threats. In *2024 20th International Conference on Wireless and Mobile Computing, Networking and Communications (WiMob)* (pp. 301–305). IEEE. <https://doi.org/10.1109/WiMob61911.2024.10770451>
3. Korobeinikova, T., Tachenko, I., Romanyuk, O., Romanyuk, S., Stakhov, O., & Reyda, O. (2024). Assessing network security risks: A technological chain perspective. In *2024 14th International Conference on Advanced Computer Information Technologies (ACIT)* (pp. 565–570). IEEE. <https://doi.org/10.1109/ACIT62333.2024.10712586>
4. Alnafea, F. S. M., Sengar, A. S., Hamid, N. K., Selladurai, K. M., Hassan, S. I., & Saravanakumar, R. (2025). Improving multi-factor authentication security with deep learning-based user behaviour analysis. In *2025 3rd International Conference on Integrated Circuits and Communication Systems (ICICACS)* (pp. 1–5). IEEE. <https://doi.org/10.1109/ICICACS65178.2025.10968961>
5. Kaur, M., & Garg, P. (2025). Exploring behavioral patterns for security in cloud computing: A case study. In *2025 3rd International Conference on Advancement in Computation & Computer Technologies (InCACCT)* (pp. 720–725). IEEE. <https://doi.org/10.1109/InCACCT65424.2025.11011337>
6. Korobeinikova, T. I. (2025). The zero trust model: Theory, practice, and prospect. In *Heritage of European Science 2025: Innovative technology, computer science, cybernetics and automation, security systems, transport development, architecture and construction, physics and mathematics* (Monographic series “European Science,” Book 37, Part 1, pp. 126–147). European Science.
7. Korobeinikova, T., Zhuravel, I., Bodak, A., & Borodenco, D. (2024). Kontsepsiia nulovoi doviry: suchasni metody zabezpechennia kiberbezpeky v korporatyvnykh merezhakh [The zero trust concept: Modern methods for ensuring cybersecurity in corporate networks]. *Visnyk Lvivskoho derzhavnoho universytetu bezpeky zhyttiedialnosti*, 30, 67–77. <https://journal.ldubgd.edu.ua/index.php/Visnuk/article/view/2769>
8. Terumalasetti, S., & Reeja, S. R. (2025). Enhancing social media user's trust: A comprehensive framework for detecting malicious profiles using multi-dimensional analytics. *IEEE Access*, 13, 7071–7093. <https://doi.org/10.1109/ACCESS.2024.3521951>
9. Cheimonidis, P., & Rantos, K. (2023). Dynamic risk assessment in cybersecurity: A systematic literature review. *Future Internet*, 15(10), 324. <https://doi.org/10.3390/fi15100324>
10. Rose, S., Borchert, O., Mitchell, S., & Connelly, S. (2020). *Zero Trust Architecture (NIST Special Publication 800-207)*. National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.800-207>
11. Martín, A. G., Martín-de-Diego, I., Fernández-Isabel, A., et al. (2022). Combining user behavioural information at the feature level to enhance continuous authentication systems. *Knowledge-Based Systems*, 244, 108544. <https://doi.org/10.1016/j.knosys.2022.108544>

12. Alcaraz, C., & Lopez, J. (2022). Digital twin: A comprehensive survey of security threats. *IEEE Communications Surveys & Tutorials*, 24(3), 1475–1503. <https://doi.org/10.1109/COMST.2022.3171465>
13. Al-Mhiqani, M. N., Alsoufi, T. A. A., Al-Shehari, T., Abdulkareem, K. H., Ahmad, R., & Mohammed, M. A. (2024). Insider threat detection in cyber-physical systems: A systematic literature review. *Computers & Electrical Engineering*, 119, 109489. <https://doi.org/10.1016/j.compeleceng.2024.109489>
14. Namane, S., Derhab, A., Guerroumi, M., & Challal, Y. (2022). Blockchain-based access control techniques for IoT: A systematic review. *Electronics*, 11(14), 2225. <https://doi.org/10.3390/electronics11142225>
15. Ullah, S. S., Oleshchuk, V., & Pussewalage, H. S. G. (2023). A survey on blockchain-envisioned attribute-based access control for Internet of Things: Overview, comparative analysis, and open research challenges. *Computer Networks*, 235, 109994. <https://doi.org/10.1016/j.comnet.2023.109994>
16. Punia, A., Hoda, M., Kaushik, K., & Tomar, D. (2024). A systematic review on blockchain-based access control systems in cloud environment. *Journal of Cloud Computing*, 13, 62. <https://doi.org/10.1186/s13677-024-00697-7>
17. Yamnych, A. B., & Korobeinikova, T. I. (2024). Model kontroliu dostupu personalu do informatsiinykh resursiv pidpriemstv na osnovi RBAC ta tekhnolohii BLOCKCHAIN [A personnel access control model for enterprise information resources based on RBAC and blockchain]. *Visnyk Khmelnytskoho natsionalnoho universytetu*, 343(6(1)), 380–386. <https://doi.org/10.31891/2307-5732-2024-343-6-56>
18. Korobeinikova, T. I., & Yamnych, A. B. (2024). Bahatovymirna matrytsia klasyfikatsii informatsii dlia otsinky ryzykiv informatsiinoi bezpeky [A multidimensional information-classification matrix for information security risk assessment]. *Informatsiini tekhnolohii ta komp'iuterna inzheneriia*, (2), 91–106.