

Школьников Владислав Ігорович

доктор філософії у галузі права, доцент, завідувач кафедри кримінології та інформаційних технологій
Національна академія внутрішніх справ, Київ, Україна
ORCID: 0000-0003-2041-9450
E-mail: shkolnikov.v.i@navs.edu.ua

Гуськова Віра Геннадіївна

доктор філософії з комп'ютерних наук, доцент кафедри штучного інтелекту
Національний технічний університет України «Київський політехнічний інститут
імені Ігоря Сікорського», Київ, Україна
ORCID: 0000-0001-7637-201X
E-mail: huskova.vira@lil.kpi.ua

Халигов Артем Азимович

аспірант відділення прикладної інформатики
Інститут телекомунікацій і глобального інформаційного простору НАН України, Київ, Україна
ORCID: 0009-0006-5465-4650
E-mail: khalygovartem@gmail.com

Лисов Богдан Сергійович

аспірант кафедри штучного інтелекту
Національний технічний університет України «Київський політехнічний інститут
імені Ігоря Сікорського», Київ, Україна
ORCID: 0009-0007-7963-6958
E-mail: bogukraine@gmail.com

ІНТЕЛЕКТУАЛЬНА СИСТЕМА ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ ДЛЯ ОЦІНЮВАННЯ КІБЕРРИЗИКІВ ОБ'ЄКТІВ КРИТИЧНОЇ ІНФРАСТРУКТУРИ НА ОСНОВІ АНАЛІЗУ ПОВЕДІНКОВОЇ АКТИВНОСТІ КОРИСТУВАЧІВ

У статті запропоновано інтелектуальну систему підтримки прийняття рішень для оцінювання кіберризиків об'єктів критичної інфраструктури на основі аналізу поведінки користувачів. Запропонований підхід поєднує двоетапний поведінковий аналіз користувачів до та після авторизації, методи машинного навчання для виявлення аномалій та байєсівське оцінювання кіберризиків у межах єдиної аналітичної архітектури. Розроблено механізм інтеграції технічних подій безпеки, поведінкових характеристик користувачів і контекстних факторів для формування інтегрального ризикового профілю. Для виявлення поведінкових аномалій використано моделі Isolation Forest, XGBoost та LSTM, а для оцінювання ризиків – байєсівську мережу, що забезпечує врахування невизначеності та причинно-наслідкових зв'язків між факторами ризику. Результати експериментальних досліджень на відкритих наборах даних CERT Insider Threat, LANL User Authentication та CICIDS2017 підтвердили ефективність запропонованого підходу. Найкращі результати виявлення аномалій продемонструвала модель LSTM, яка досягла значень ROC-AUC 0,96 для передавтентифікаційного аналізу та 0,97 для післяавтентифікаційного аналізу. Запропонована система забезпечила точність класифікації кіберризиків 0,91 та якість рекомендацій 0,89 при середньому часі реагування 87 мс. Практична цінність роботи полягає у можливості переходу від виявлення аномалій до комплексного оцінювання кіберризиків та формування рекомендацій щодо реагування на потенційні загрози в середовищах критичної інфраструктури.

Ключові слова: кібербезпека, критична інфраструктура, система підтримки прийняття рішень, оцінювання кіберризиків, аналіз поведінки користувачів, UBA, UEBA, машинне навчання; виявлення аномалій, байєсівські мережі, LSTM, XGBoost.

Вступ

Цифровізація об'єктів критичної інфраструктури супроводжується зростанням кількості та складності кіберзагроз, здатних впливати на функціонування енергетичних, транспортних, фінансових та державних інформаційних систем. У таких умовах кіберінциденти можуть призводити не лише до фінансових втрат, а й до порушення роботи критично важливих сервісів та виникнення ризиків для національної безпеки. Традиційні засоби кіберзахисту переважно орієнтовані на аналіз мережевого трафіку та моніторинг подій безпеки. Водночас

суттєва частина кіберінцидентів пов'язана з людським фактором, зокрема компрометацією облікових записів, порушенням політик безпеки, несанкціонованим використанням привілеїв та внутрішніми загрозами. У зв'язку з цим активно розвиваються підходи аналізу поведінки користувачів (User Behavior Analytics, UBA) та аналізу поведінки користувачів і сутностей (User and Entity Behavior Analytics, UEBA). Вони орієнтовані на виявлення відхилень від типових шаблонів активності на основі таких характеристик, як частота авторизацій, час доступу, використання привілейованих функцій, звернення до критичних ресурсів та інші параметри взаємодії з інформаційною системою. Використання таких ознак дозволяє підвищити ефективність раннього виявлення загроз, пов'язаних із людським фактором, компрометацією облікових записів або внутрішніми порушеннями.

Незважаючи на розвиток методів поведінкового аналізу, більшість існуючих рішень зосереджена на виявленні аномалій і не забезпечує комплексного оцінювання кіберризиків з урахуванням технічних, поведінкових та контекстних факторів. Це обумовлює необхідність розроблення інтелектуальних систем підтримки прийняття рішень, здатних інтегрувати результати поведінкового аналізу, методи машинного навчання та механізми оцінювання ризиків у єдиному аналітичному середовищі.

Постановка проблеми

Нехай об'єкт критичної інфраструктури розглядається як складна інформаційно-комунікаційна система, у якій взаємодіють користувачі, мережеві сервіси, прикладні компоненти та засоби контролю доступу. Для формалізації процесу оцінювання кіберризиків введемо множину користувачів, множину мережевих подій та множину системних подій. Множина користувачів визначається як:

$$U = \{u_1, u_2, \dots, u_n\}, \quad (1)$$

де кожен елемент u_i відповідає окремому користувачу системи. Множина мережевих подій задається як:

$$N = \{n_1, n_2, \dots, n_p\}, \quad (2)$$

де n_j описує окрему подію мережевої активності, наприклад з'єднання, передачу даних, підозрілий запит або спрацювання засобу мережевого моніторингу. Множина системних подій визначається як:

$$S = \{s_1, s_2, \dots, s_q\}, \quad (3)$$

де s_k відповідає окремій події системного рівня, зокрема авторизації, зміні прав доступу, запуску сервісу, зверненню до файлу або запису в журналі безпеки. Для кожного користувача u_i формується вектор поведінкових характеристик:

$$B = \{b_{i1}, b_{i2}, \dots, b_{im}\}. \quad (4)$$

Паралельно для інтервалу часу формується вектор технічних характеристик:

$$T = \{t_1, t_2, \dots, t_j\}, \quad (5)$$

де t_j описує технічні індикатори безпеки, отримані на основі мережевих і системних подій.

Після виконання процедур очищення, нормалізації, синхронізації часових міток та інженерії ознак формується інтегрований вектор стану системи:

$$X = [T, B], \quad (6)$$

де T відображає технічний стан інформаційної інфраструктури, а B – поведінковий профіль користувача або групи користувачів. Отриманий вектор X використовується як вхідна

інформація для модулів виявлення аномалій, байєсівського оцінювання кіберризиків та формування рекомендацій у системі підтримки прийняття рішень.

Необхідно розробити функцію оцінювання кіберризиків, яка на основі технічних подій безпеки, поведінкових характеристик користувачів та контекстних факторів забезпечує визначення рівня ризику для об'єкта критичної інфраструктури:

$$F: (T, B, C) \rightarrow R, \quad (7)$$

де T – множина технічних характеристик та подій безпеки, отриманих із мережевих і системних журналів; B – множина поведінкових характеристик користувачів; C – множина контекстних факторів, що характеризують критичність ресурсів, роль користувача, тип операції та інші особливості середовища; R – інтегральна оцінка кіберризиків.

При цьому функція F повинна забезпечувати інтеграцію різнорідних джерел даних, виявлення поведінкових аномалій, оцінювання ризику в умовах невизначеності та формування інформаційної основи для підтримки прийняття рішень щодо реагування на потенційні кіберзагрози.

Аналіз останніх досліджень та публікацій

Аналіз поведінки користувачів (UBA/UEBA) є одним із перспективних напрямів розвитку кібербезпеки, спрямованим на виявлення відхилень від типових шаблонів активності користувачів. Дослідження Eberle та Holder [1], Salem і Stolfo [2], Legg та ін. [3] показали ефективність поведінкового аналізу для виявлення внутрішніх загроз, компрометації облікових записів та несанкціонованих дій. Проте більшість існуючих рішень орієнтована переважно на виявлення аномалій без подальшого оцінювання кіберризиків. Для аналізу поведінкових характеристик користувачів широко застосовуються методи машинного навчання, зокрема Isolation Forest [5], Random Forest [6], XGBoost [7], Autoencoder [9] та LSTM [8]. Ці підходи дозволяють виявляти аномальну активність як за наявності розмічених даних, так і в умовах навчання без учителя. Водночас результати їх роботи зазвичай представлені у вигляді оцінок аномальності та не забезпечують комплексного оцінювання рівня кіберризиків. Для оцінювання кіберризиків в умовах невизначеності широко використовуються байєсівські мережі [9-10], які дозволяють моделювати причинно-наслідкові залежності між подіями безпеки та інтегрувати різнорідні джерела інформації. Більшість існуючих моделей ризику зосереджена переважно на технічних факторах та недостатньо враховує поведінкові характеристики користувачів. Системи підтримки прийняття рішень (DSS) застосовуються для оцінювання ризиків, пріоритизації інцидентів та вибору заходів реагування [10]. Сучасні DSS дедалі частіше інтегрують методи штучного інтелекту та машинного навчання, однак поведінковий контекст користувачів у більшості випадків залишається недостатньо врахованим.

Аналіз літератури показав, що задачі поведінкового аналізу, виявлення аномалій, оцінювання кіберризиків та підтримки прийняття рішень переважно розглядаються окремо. Недостатньо дослідженим залишається питання побудови єдиної DSS-архітектури для об'єктів критичної інфраструктури, яка б інтегрувала поведінкові характеристики користувачів, результати машинного навчання та байєсівське оцінювання кіберризиків. У роботі запропоновано інтелектуальну систему підтримки прийняття рішень, що забезпечує інтеграцію поведінкового аналізу користувачів, методів машинного навчання для виявлення аномалій та байєсівського оцінювання кіберризиків у межах єдиної аналітичної архітектури. Для формального подання запропонованого підходу доцільно визначити взаємозв'язки між поведінковими характеристиками користувачів, технічними подіями безпеки та параметрами стану інформаційної інфраструктури, що становить основу математичної постановки задачі оцінювання кіберризиків об'єктів критичної інфраструктури.

Метою статті є розроблення інтелектуальної системи підтримки прийняття рішень для оцінювання кіберризиків об'єктів критичної інфраструктури. Система базується на аналізі

поведінки користувачів до та після авторизації, використанні методів машинного навчання для виявлення аномалій і байєсівському підході для комплексного врахування технічних, поведінкових та контекстних факторів ризику.

Виклад основного матеріалу

Запропонована архітектура інтелектуальної системи підтримки прийняття рішень орієнтована на комплексне оцінювання кіберризиків об'єктів критичної інфраструктури з урахуванням технічних, системних, поведінкових та контекстних характеристик. Основна ідея підходу полягає в тому, що рівень кіберризиків формується не лише внаслідок наявності технічних індикаторів, а й під впливом аномальної або нетипової поведінки користувачів, особливо тих, які мають доступ до критично важливих ресурсів.

Представлений ілюструє архітектуру інтелектуальної системи підтримки прийняття рішень для оцінювання кіберризиків об'єктів критичної інфраструктури з урахуванням поведінкових характеристик користувачів (рис. 1).

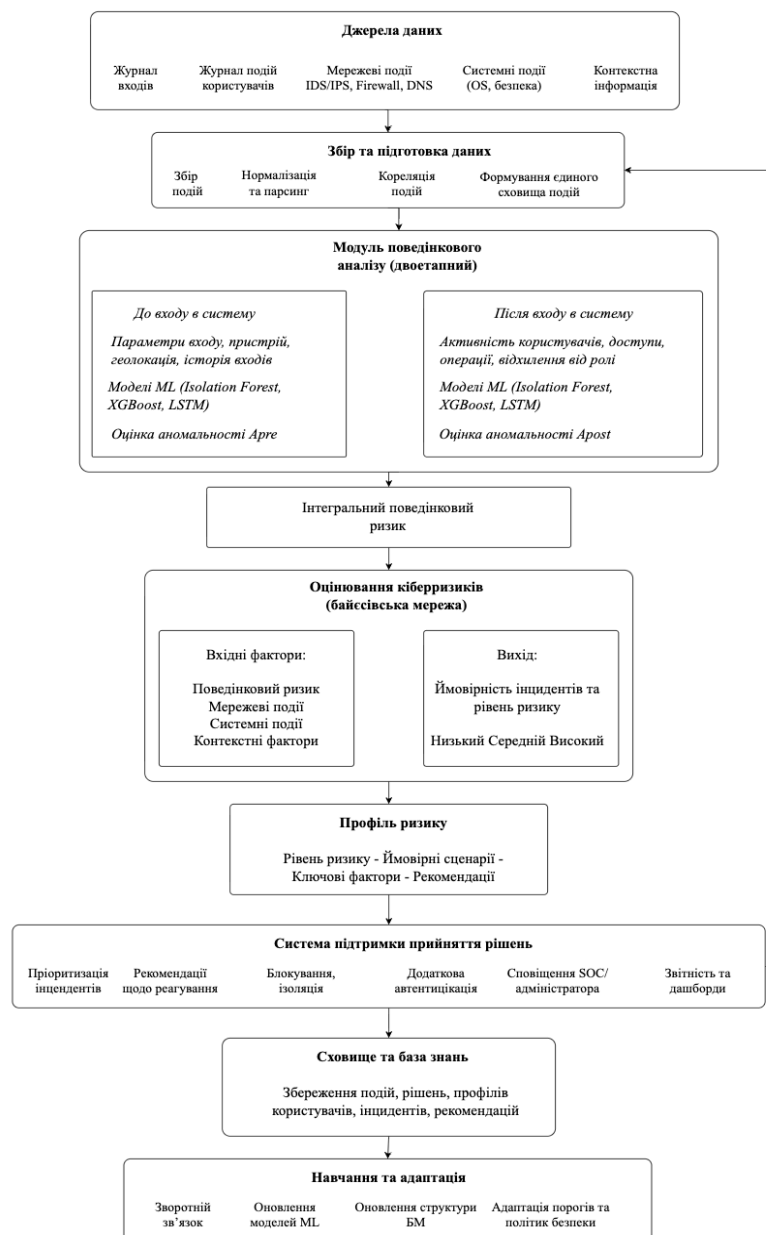


Рис. 1. Архітектура інтелектуальної системи підтримки прийняття рішень для оцінювання кіберризиків об'єктів критичної інфраструктури

Підхід, побудований як послідовний аналітичний конвеєр, у якому дані з різних джерел перетворюються на інтегральну оцінку кіберризиків та рекомендації щодо реагування. На першому етапі здійснюється збір даних із журналів авторизації, мережних подій, системних журналів, SIEM-систем та джерел контекстної інформації. Далі дані проходять нормалізацію, кореляцію подій, синхронізацію часових міток і формування єдиного сховища подій.

Ключовим елементом архітектури є двоетапний модуль поведінкового аналізу користувачів, який забезпечує оцінювання ризику як до авторизації, так і після входу до системи. На першому етапі аналізуються параметри авторизації, геолокація, пристрій та історія попередніх входів, а на другому – подальша активність користувача, зокрема доступ до критичних ресурсів, використання привілейованих функцій та відхилення від рольового профілю. Детальніше цей механізм розглянуто в наступному розділі.

Результати поведінкового аналізу використовуються для формування інтегрального поведінкового ризику користувача, який разом із технічними, системними та контекстними факторами передається до модуля байєсівського оцінювання ризиків. На цьому етапі визначається ймовірність інциденту та рівень ризику: низький, середній або високий. Отримана оцінка використовується системою підтримки прийняття рішень для пріоритизації інцидентів, оцінювання впливу, аналізу можливих сценаріїв розвитку подій та формування рекомендацій щодо реагування. До таких рекомендацій можуть належати додаткова автентифікація, обмеження доступу, блокування або ізоляція користувача, сповіщення SOC/адміністраторів та інші контрзаходи.

Додатково архітектура передбачає модуль навчання та адаптації, який забезпечує оновлення моделей машинного навчання, коригування порогів аномальності, уточнення структури байєсівської мережі та врахування експертного зворотного зв'язку. Це дозволяє системі адаптуватися до змін поведінкових профілів користувачів, нових типів загроз і специфіки конкретного об'єкта критичної інфраструктури.

Моніторинг поведінки користувачів і профілювання користувачів

Особливістю запропонованого підходу є використання двоетапного механізму поведінкового контролю користувачів, який охоплює як етап автентифікації, так і подальшу діяльність користувача після отримання доступу до інформаційної системи.

У процесі функціонування системи для кожного користувача формується та постійно оновлюється поведінковий профіль, що накопичує статистичні характеристики його взаємодії з інфраструктурою. Такий профіль є основою для виявлення відхилень від типової поведінки та оцінювання потенційного кіберризиків.

Запропонований механізм складається з двох взаємопов'язаних етапів, представлених на рисунку нижче (рис. 2).

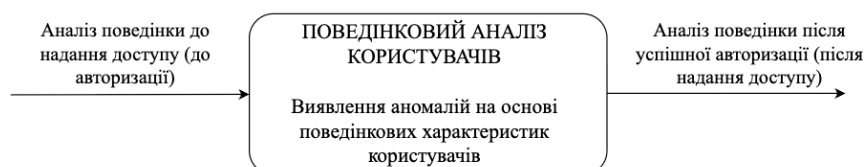


Рис. 2. Двоетапний механізм поведінкового аналізу користувачів у системі оцінювання кіберризиків

Поведінковий аналіз до авторизації

На першому етапі оцінюється ризик користувача до моменту отримання доступу до системи. Основною метою є виявлення ознак компрометації облікового запису або нетипового сценарію входу ще до початку роботи користувача з інформаційними ресурсами. Для виконання передавтентифікаційного аналізу використовуються характеристики, що описують процес входу користувача до системи, зокрема частота авторизацій, кількість невдалих спроб

входу, часові параметри доступу, географічна локація, використовуваний пристрій, IP-адреса, тип автентифікації та історія попередніх сеансів.

Водночас перелік ознак не є фіксованим і може розширюватися або коригуватися залежно від особливостей інформаційної системи, доступних джерел даних та результатів інженерії ознак. Важливу роль на цьому етапі відіграють експерти з кібербезпеки, які беруть участь у відборі найбільш інформативних характеристик, налаштуванні порогових значень та інтерпретації виявлених аномалій. На основі сформованого набору параметрів будується вектор передавтентифікаційних ознак:

$$B_i^{pre} = [LF_i, FLR_i, ATD_i, GAI_i], \quad (8)$$

де LF_i – частота входів; FLR_i – частка невдалих спроб входу; ATD_i – відхилення часу доступу; GAI_i – невідповідність географічної локації.

Для даного вектора обчислюється показник аномальності:

$$A_i^{pre} = M_{pre}(B_i^{pre}), \quad (9)$$

де M_{pre} – модель машинного навчання для аналізу авторизаційної активності.

Якщо:

$$A_i^{pre} > T_{pre}, \quad (10)$$

де T_{pre} – встановлений поріг ризику, система може вимагати багатофакторну автентифікацію; тимчасово обмежити доступ; заблокувати сесію та передати подію оператору SOC. В такий підхід реалізується превентивний захист ще до отримання доступу до ресурсів системи.

Поведінковий аналіз після авторизації

Після успішної авторизації здійснюється аналіз поведінки користувача безпосередньо під час роботи в інформаційній системі. На цьому етапі враховуються характеристики, що відображають фактичні дії користувача, зокрема тривалість сесій, використання привілейованих функцій, звернення до критичних ресурсів, відповідність дій рольовій моделі доступу, інтенсивність операцій із даними, спроби зміни прав доступу та інші нетипові сценарії роботи

Аналогічно до передавтентифікаційного аналізу, склад ознак може змінюватися залежно від специфіки об'єкта критичної інфраструктури, доступних джерел даних та результатів інженерії ознак. Відповідний вектор післяавтентифікаційних ознак визначається як:

$$B_i^{post} = [SD_i, RAD_i, CRA_i, PUF_i], \quad (11)$$

де SD_i – середня тривалість сесії; RAD_i – відхилення від рольового профілю; CRA_i – кількість звернень до критичних ресурсів; PUF_i – частота використання привілейованих операцій. Показник аномальності після входу визначається як

$$A_i^{post} = M_{post}(B_i^{post}), \quad (12)$$

де M_{post} – модель поведінкового аналізу активності користувача.

Якщо:

$$A_i^{post} > T_{post}. \quad (13)$$

За результатами оцінювання ризику система може ініціювати додаткові заходи захисту, зокрема обмеження доступу до критичних ресурсів, коригування привілеїв користувача, посилений моніторинг активності, пріоритизацію інциденту або передачу події до SOC для подальшого реагування.

Інтегральний поведінковий ризик

Для комплексного врахування обох рівнів контролю формується інтегральна оцінка поведінкового ризику:

$$BRisk_i = \alpha A_i^{pre} + (1 - \alpha) A_i^{post}, \quad (14)$$

де $BRisk_i$ – інтегральний поведінковий ризик користувача; A_i^{pre} – ризик до авторизації; A_i^{post} – ризик після авторизації; α – ваговий коефіцієнт.

Отримана оцінка надалі використовується модулем байєсівського оцінювання кіберризиків разом із технічними та контекстними факторами.

Байєсівське оцінювання кіберризиків

Після виявлення поведінкових аномалій виконується байєсівське оцінювання кіберризиків. Запропонований підхід передбачає комплексне врахування технічних факторів безпеки, інтегральної оцінки поведінкової аномальності та контекстних характеристик середовища. Використання байєсівських мереж є доцільним, оскільки оцінювання кіберризиків відбувається в умовах невизначеності, неповноти даних та складних причинно-наслідкових зв'язків між подіями безпеки.

Оцінка ризику визначається як умовна ймовірність:

$$P(R|T, A, C), \quad (15)$$

де R – рівень кіберризиків, T – множина технічних факторів, A – інтегральна оцінка поведінкової аномальності, C – множина контекстних факторів.

Відповідно до теореми Байєса:

$$P(R|T, A, C) = \frac{P(T, A, C|R) \cdot P(R)}{P(T, A, C)}. \quad (16)$$

Для використання в системі підтримки прийняття рішень формується інтегральний показник ризику:

$$R_{score} = F(T, A, C), \quad (17)$$

або у ваговій формі:

$$R_{score} = \beta T_{score} + \gamma A + \delta C_{score}, \quad (18)$$

де T_{score} – технічний ризик, A – поведінковий ризик, C_{score} – контекстний ризик, а $\beta + \gamma + \delta = 1$.

Отримане значення R_{score} використовується для класифікації рівня ризику, пріоритезації інцидентів і формування рекомендацій щодо реагування на потенційні кіберзагрози.

Алгоритм поведінково-орієнтованого оцінювання кіберризиків

Для узагальнення процесу оцінювання кіберризиків та підтримки прийняття рішень розроблено алгоритм Behavioral-Aware Risk Assessment. Алгоритм об'єднує результати двоетапного поведінкового аналізу, оцінювання технічних і контекстних факторів, байєсівське оцінювання ризику та формування рекомендацій щодо реагування на потенційні кіберзагрози.

Algorithm 1 Behavioral-Aware Risk Assessment

Input:

T – technical factors

B_{pre} – pre-authentication behavioral features

B_{post} – post-authentication behavioral features

C – contextual factors

*Output:**Risk Level**Recommended Actions*

1: Calculate pre-authentication anomaly score

$$A_{pre} = M_{pre}(B_{pre})$$

2: Calculate post-authentication anomaly score

$$A_{post} = M_{post}(B_{post})$$

3: Compute integrated behavioral anomaly score

$$A = \alpha A_{pre} + (1-\alpha)A_{post}$$

4: Estimate cyber risk using Bayesian Network

$$R_{score} = F(T, A, C)$$

5: Generate risk profile

$$RP = \{A, T, C, R_{score}\}$$

6: Classify risk level

Low / Medium / High / Critical

7: Generate recommendations

$$Rec = G(RP)$$

8: Return Risk Level and Recommended Actions

Запропонований алгоритм реалізує повний цикл оцінювання кіберризиків – від аналізу поведінки користувача до формування рекомендацій щодо реагування. Особливістю підходу є використання двоетапного поведінкового контролю, що дозволяє виявляти потенційні загрози як на етапі авторизації користувача, так і під час його подальшої роботи в інформаційній системі. Результати алгоритму використовуються для побудови профілю, пріоритетизації інцидентів та підтримки прийняття рішень у середовищах критичної інфраструктури.

Підтримка прийняття рішень та формування рекомендацій

Після оцінювання поведінкових аномалій та розрахунку інтегрального показника кіберризиків система переходить до формування рекомендацій щодо реагування на потенційні загрози. На відміну від традиційних підходів, запропонована архітектура підтримує прийняття рішень як на етапі авторизації користувача, так і під час його подальшої роботи в інформаційній системі.

На етапі авторизації рекомендації формуються на основі результатів передавтентифікаційного аналізу та можуть включати багатофакторну автентифікацію, додаткову перевірку користувача, підтвердження входу з довіреного пристрою, аналіз географічної локації або тимчасове блокування спроби входу.

Після успішної авторизації система використовує результати поведінкового аналізу та оцінювання ризику для формування коригувальних дій. Залежно від рівня ризику можуть застосовуватися обмеження доступу до критичних ресурсів, зниження рівня привілеїв, додатковий моніторинг активності, аудит дій користувача або передача інциденту до центру моніторингу безпеки (SOC).

Загалом процес формування рекомендацій базується на інтегральному ризиковому профілі користувача, який враховує результати передавтентифікаційного та післяавтентифікаційного аналізу, технічні індикатори безпеки та контекстні фактори. Це дозволяє перейти від простого виявлення аномалій до інтелектуальної підтримки прийняття рішень щодо реагування на кіберзагрози в середовищах критичної інфраструктури.

На рисунку нижче представлено процес формування рекомендацій на основі результатів двоетапного поведінкового аналізу користувачів, технічних та контекстних факторів (рис. 3). Оцінки аномальності до авторизації A_i^{pre} та після авторизації A_i^{post}

поєднуються з технічними подіями безпеки та контекстною інформацією в межах байєсівської моделі оцінювання ризику. На основі інтегрального показника ризику R_{score} формується ризиковий профіль користувача, який використовується модулем підтримки прийняття рішень для генерації рекомендацій щодо реагування на потенційні кіберзагрози, зокрема додаткової автентифікації, обмеження доступу, зниження привілеїв, посиленого моніторингу або передачі інциденту до центру моніторингу безпеки (SOC).

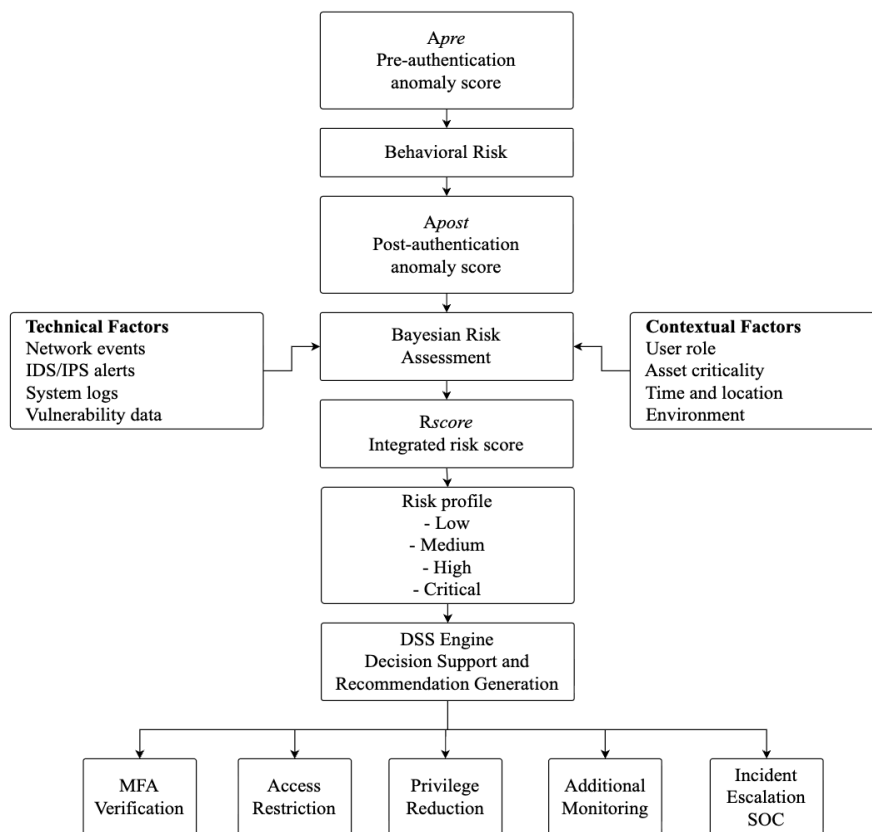


Рис. 3. Механізм формування рекомендацій у системі підтримки прийняття рішень для оцінювання кіберризиків

Експериментальна оцінка запропонованого підходу

Налаштування експерименту

Для експериментальної перевірки запропонованої архітектури використовувалися відкриті набори даних CERT Insider Threat Dataset, LANL User Authentication Dataset та CICIDS2017. Для моделювання післяавтентифікаційної активності користувачів було сформовано додатковий синтетичний набір поведінкових ознак на основі статистичних характеристик реальних журналів активності.

Після інтеграції даних було сформовано 48 732 записи користувацької активності, які містили технічні, поведінкові та контекстні характеристики. Вибірка була розподілена за схемою: Training set – 70%; Validation set – 15%; Test set – 15%.

Для забезпечення відтворюваності результатів використовувалась фіксована ініціалізація генератора випадкових чисел (random seed = 42).

Процедура навчання моделей

Для оцінювання передавтентифікаційних та післяавтентифікаційних аномалій були використані моделі Isolation Forest, XGBoost та LSTM. Наведено основні гіперпараметри моделей машинного навчання, використаних під час експериментального дослідження (табл. 1).

Таблиця 1

Параметри моделей для виявлення поведінкових аномалій

Модель	Параметр	Значення
Isolation Forest	n_estimators	200
	contamination	0.05
	max_samples	auto
XGBoost	max_depth	6
	learning_rate	0.1
	n_estimators	300
	subsample	0.8
LSTM	Кількість прихованих шарів	2
	Кількість нейронів у шарі	64
	batch size	64
	epochs	50
	optimizer	Adam

Значення гіперпараметрів були обрані на основі попередніх експериментів та рекомендацій, наведених у відповідних наукових дослідженнях. Для всіх моделей використовувалася однакова процедура попередньої обробки даних, що включала стандартизацію числових ознак та кодування категоріальних атрибутів методом One-Hot Encoding.

Побудова байєсівської мережі

Для оцінювання кіберризиків у запропонованій системі використовується байєсівська мережа, яка забезпечує інтеграцію технічних, поведінкових та контекстних факторів у межах єдиної ймовірнісної моделі. Використання байєсівського підходу дозволяє враховувати невизначеність, неповноту даних та складні причинно-наслідкові зв'язки між подіями безпеки.

Структура мережі включає чотири основні вузли:

1. **Technical Risk** – ризик, сформований на основі мережевих аномалій, спрацювань IDS/IPS та інших технічних індикаторів безпеки;
2. **Behavioral Risk** – ризик, сформований на основі результатів передавентифікаційного та післяавтентифікаційного аналізу поведінки користувача;
3. **Contextual Risk** – ризик, пов'язаний із критичністю ресурсу, роллю користувача та характеристиками середовища функціонування;
4. **Cyber Risk** – інтегральний показник кіберризиків.

Структура причинно-наслідкових залежностей між вузлами має вигляд:

Technical Risk → *Cyber Risk* ,
Behavioral Risk → *Cyber Risk* ,
Contextual Risk → *Cyber Risk* .

Для параметризації мережі використовувалися таблиці умовних ймовірностей (Conditional Probability Tables, CPT), сформовані на основі експертних оцінок та статистичного аналізу навчальної вибірки. Фрагмент CPT для вузла Cyber Risk наведено в таблиці 2.

Таблиця 2

Умовні ймовірності оцінювання кіберризiku залежно від технічних, поведінкових та контекстних факторів

Technical Risk	Behavioral Risk	Context Risk	P(Cyber Risk = High)
Low	Low	Low	0.05
High	Low	Low	0.35
Low	High	Medium	0.48
Medium	High	High	0.76
High	High	High	0.95

Навчання параметрів байєсівської мережі виконувалося методом Maximum Likelihood Estimation (MLE). Для оцінювання апостеріорних ймовірностей та обчислення інтегрального показника ризику використовувався алгоритм Variable Elimination. Отримані значення ймовірностей використовувалися для формування інтегральної оцінки кіберризiku та подальшого побудови ризикового профілю користувача, який надалі передавався до модуля підтримки прийняття рішень.

Для наочного подання причинно-наслідкових зв'язків між факторами ризику на рисунку наведено структуру байєсівської мережі, що використовується для оцінювання кіберризиків. Мережа об'єднує технічні, поведінкові та контекстні фактори, які впливають на формування інтегрального рівня кіберризiku (рис. 4).

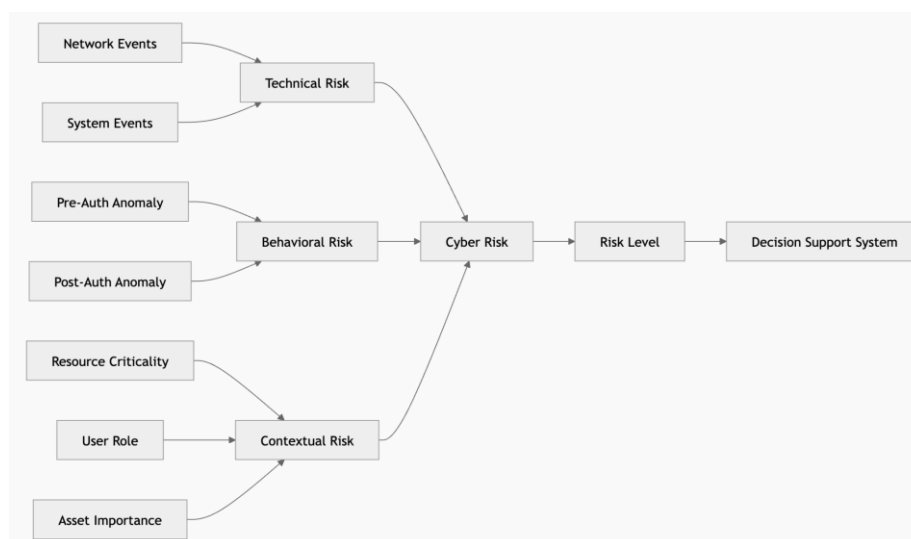


Рис. 4. Структура байєсівської мережі для оцінювання кіберризиків

Як показано на рисунку 4, вузол Technical Risk формується на основі мережевих аномалій, спрацювань IDS/IPS та системних подій. Вузол Behavioral Risk залежить від результатів передавентифікаційного та післяавтентифікаційного аналізу поведінки користувача. Вузол Contextual Risk враховує критичність ресурсу, роль користувача та важливість інформаційного активу. Сукупний вплив цих трьох груп факторів визначає вузол

Cyber Risk, який надалі використовується для класифікації рівня ризику та формування рекомендацій у DSS.

Порівняльний аналіз

Для оцінювання особливостей запропонованого підходу було проведено порівняльний аналіз існуючих рішень у сфері поведінкового аналізу користувачів, оцінювання кіберризиків та моніторингу безпеки. До порівняння було включено типові UEBA-системи, моделі виявлення аномалій на основі Isolation Forest, байєсівські моделі оцінювання ризиків та традиційні SIEM-рішення. Основними критеріями порівняння виступали наявність механізмів поведінкового аналізу, підтримка оцінювання кіберризиків, функції підтримки прийняття рішень та можливість виконання двоетапного аналізу поведінки користувача до та після авторизації (табл. 3).

Таблиця 3

Порівняння запропонованого підходу з існуючими рішеннями

Approach	Behavioral Analysis	Risk Assessment	DSS	Pre/Post Authentication
UEBA Systems	+	-	-	-
Isolation Forest	+	-	-	-
Bayesian Risk Models	-	+	-	-
Traditional SIEM	Partial	Partial	-	-
Proposed DSS	+	+	+	+

Результати порівняння свідчать, що більшість існуючих рішень орієнтована на вирішення окремих задач виявлення аномалій або оцінювання ризиків. На відміну від них запропонована система забезпечує комплексну інтеграцію поведінкового аналізу, оцінювання кіберризиків та механізмів підтримки прийняття рішень. Додатковою перевагою є реалізація двоетапного контролю поведінки користувача, який дозволяє виконувати аналіз як на етапі авторизації, так і під час подальшої роботи в інформаційній системі.

Експеримент 1. Виявлення аномалій до авторизації

Для оцінювання ефективності виявлення аномальної поведінки користувачів до авторизації використовувалися набори даних LANL User Authentication Dataset та Synthetic Authentication Dataset, які містять інформацію про події автентифікації користувачів. Аналіз виконувався на основі поведінкових ознак, що характеризують частоту входів до системи, кількість невдалих спроб авторизації, відхилення часу доступу від типового профілю та невідповідність географічної локації користувача. Для виявлення аномалій були використані моделі Isolation Forest, XGBoost та LSTM. Ефективність моделей оцінювалася за метриками Accuracy, Precision, Recall, F1-score та ROC-AUC. Наведено результати експерименту (табл. 4).

Таблиця 4

Ефективність моделей машинного навчання під час виявлення передавтентифікаційних аномалій

Model	Accuracy	Precision	Recall	F1	ROC-AUC
Isolation Forest	0.88	0.86	0.84	0.85	0.89
XGBoost	0.93	0.92	0.91	0.92	0.95
LSTM	0.94	0.93	0.92	0.93	0.96

Отримані результати свідчать, що найкращі показники продемонструвала модель LSTM, яка досягла значення ROC-AUC 0.96 та F1-міри 0.93. Це підтверджує доцільність використання рекурентних нейронних мереж для аналізу часових характеристик процесу авторизації та виявлення нетипових сценаріїв входу користувачів до системи. Водночас модель XGBoost також показала високий рівень точності, тоді як Isolation Forest забезпечив прийнятні результати без використання розмічених даних.

Експеримент 2. Виявлення аномалій після авторизації

Для оцінювання ефективності виявлення аномальної активності користувачів після авторизації використовувалися набори даних CERT Insider Threat Dataset та Synthetic Behavioral Dataset. Аналіз здійснювався на основі післяавтентифікаційних поведінкових ознак, що характеризують тривалість сесій, відхилення дій користувача від типової рольової моделі доступу, кількість звернень до критичних ресурсів та частоту використання привілейованих функцій. Для порівняння якості виявлення аномалій були застосовані моделі Isolation Forest, XGBoost та LSTM (табл. 5).

Таблиця 5

Ефективність моделей машинного навчання під час виявлення післяавтентифікаційних аномалій

Model	Accuracy	Precision	Recall	F1	ROC-AUC
Isolation Forest	0.86	0.84	0.83	0.83	0.87
XGBoost	0.92	0.91	0.90	0.91	0.94
LSTM	0.95	0.94	0.93	0.94	0.97

Результати показали, що найвищу якість виявлення післяавтентифікаційних аномалій забезпечила модель LSTM, яка досягла значення ROC-AUC 0.97 та F1-міри 0.94. Це підтверджує ефективність використання рекурентних нейронних мереж для аналізу послідовностей дій користувачів та виявлення складних поведінкових відхилень. Модель XGBoost також продемонструвала високі результати, тоді як Isolation Forest забезпечив прийнятну якість аналізу без необхідності використання розмічених даних.

Експеримент 3. Інтегральна оцінка ризику

Після виконання передавтентифікаційного та післяавтентифікаційного аналізу отримані оцінки аномальності об'єднуються в єдиний інтегральний показник поведінкового ризику.

Таблиця 6

Результати байєсівського оцінювання кіберризиків для типових сценаріїв активності користувачів

Scenario	Pre-Anomaly	Post-Anomaly	Risk
Normal Activity	0.08	0.12	Low
Unusual Login Time	0.42	0.18	Medium
Unusual Geolocation	0.57	0.21	Medium
Privilege Abuse	0.15	0.81	High
Account Compromise	0.88	0.90	Critical
Insider Threat	0.35	0.96	Critical

Такий підхід дозволяє враховувати як особливості процесу авторизації користувача, так і його подальшу активність у системі. Сформований інтегральний показник надалі використовується як один із вхідних параметрів байєсівської мережі для оцінювання загального рівня кіберризiku (табл. 6).

Експеримент 4. Оцінювання DSS

Окрім оцінювання якості моделей машинного навчання було проведено оцінювання ефективності системи підтримки прийняття рішень у цілому. Для цього використовувалися метрики точності класифікації ризику, якості сформованих рекомендацій, часу реагування та частоти хибнопозитивних спрацювань (табл. 7).

Таблиця 7

Результати оцінювання ефективності системи підтримки прийняття рішень

Metric	Value
Risk Classification Accuracy	0.91
Recommendation Precision	0.89
Average Response Time	87 ms
False Positive Rate	0.06

Результати свідчать про високу ефективність запропонованої DSS-архітектури. Зокрема, точність класифікації ризиків перевищує 90 %, а середній час обробки події не перевищує 100 мс, що дозволяє використовувати систему в режимі, близькому до реального часу.

Обговорення результатів

Результати експериментального дослідження підтвердили доцільність використання двоетапного поведінкового аналізу користувачів для оцінювання кіберризиків об'єктів критичної інфраструктури. Отримані значення метрик Accuracy, Precision, Recall, F1-score та ROC-AUC свідчать про здатність моделей машинного навчання ефективно виявляти аномальні шаблони активності як на етапі авторизації користувача, так і після його входу до інформаційної системи.

Порівняння моделей показало, що XGBoost та LSTM забезпечують найвищу якість виявлення поведінкових аномалій. При цьому модель LSTM продемонструвала найкращі результати завдяки здатності враховувати часову послідовність подій та динаміку поведінки користувачів. Це є важливим для задач кібербезпеки, де загрози часто проявляються у вигляді послідовності взаємопов'язаних дій, а не окремих подій.

Отримані результати також свідчать, що післяавтентифікаційний аналіз забезпечує дещо вищу точність порівняно з передавтентифікаційним. Це пояснюється використанням ширшого набору поведінкових характеристик, зокрема даних про доступ до критичних ресурсів, використання привілейованих функцій та відхилення від типової ролі моделі. Водночас передавтентифікаційний аналіз виконує важливу превентивну функцію та дозволяє виявляти потенційно небезпечні спроби доступу ще до отримання користувачем доступу до інформаційних ресурсів.

Перевагою запропонованого підходу є інтеграція результатів поведінкового аналізу з технічними та контекстними факторами в межах байєсівської моделі оцінювання кіберризiku. Це забезпечує більш повне врахування особливостей функціонування інформаційної системи та дозволяє формувати обґрунтовані оцінки ризику навіть за умов неповноти або невизначеності вхідних даних.

Оцінювання системи підтримки прийняття рішень показало, що запропонована архітектура дозволяє не лише виявляти аномалії, а й формувати ризикові профілі користувачів

та рекомендації щодо реагування на потенційні загрози. Це створює передумови для зменшення навантаження на операторів безпеки, підвищення ефективності пріоритезації інцидентів та скорочення часу реагування.

Разом із тим результати дослідження мають певні обмеження. Частина поведінкових сценаріїв була сформована на основі синтетичних даних, що може не повністю відображати складність реальних середовищ критичної інфраструктури. Крім того, запропонований підхід потребує подальшої апробації на реальних потокових даних SOC, SIEM та SOAR-систем.

Висновки

У статті запропоновано інтелектуальну систему підтримки прийняття рішень для оцінювання кіберризиків об'єктів критичної інфраструктури на основі аналізу поведінки користувачів. Розроблено двоетапний механізм поведінкового аналізу, який охоплює передавентифікаційний та післяавтентифікаційний контроль активності користувачів і дозволяє своєчасно виявляти потенційні загрози. Запропоновано підхід до інтеграції технічних подій безпеки, поведінкових характеристик користувачів та контекстних факторів у межах єдиного процесу оцінювання кіберризиків. Для цього використано байєсівську мережу, яка забезпечує врахування невизначеності та причинно-наслідкових зв'язків між різними групами факторів ризику.

Проведене експериментальне дослідження підтвердило ефективність запропонованого підходу. Найкращі результати виявлення поведінкових аномалій продемонструвала модель LSTM, яка досягла значень ROC-AUC 0,96 для передавентифікаційного аналізу та 0,97 для післяавтентифікаційного аналізу. Точність класифікації кіберризиків у системі підтримки прийняття рішень склала 0,91, а точність сформованих рекомендацій – 0,89.

Практична цінність роботи полягає у можливості використання запропонованої системи для моніторингу поведінки користувачів, оцінювання кіберризиків, пріоритезації інцидентів та підтримки прийняття рішень у середовищах критичної інфраструктури. Подальші дослідження доцільно спрямувати на розширення набору поведінкових ознак, адаптивне налаштування моделей машинного навчання, автоматизоване навчання структури байєсівських мереж та інтеграцію запропонованого підходу з реальними SOC, SIEM і SOAR-платформами.

Перелік посилань:

1. Eberle W., Holder L. Insider Threat Detection Using Graph-Based Approaches // *Cybersecurity Applications & Technology Conference for Homeland Security*. 2009. P. 237–241. DOI: 10.1109/CATCH.2009.7.
2. Salem M. B., Stolfo S. J. Modeling User Search Behavior for Masquerade Detection // *Recent Advances in Intrusion Detection*. Berlin : Springer, 2011. P. 181–200. DOI: 10.1007/978-3-642-23644-0_10.
3. Legg P. A., Buckley O., Goldsmith M., Creese S. Automated Insider Threat Detection System Using User and Role-Based Profile Assessment // *IEEE Systems Journal*. 2017. Vol. 11, No. 2. P. 503–512. DOI: 10.1109/JSYST.2015.2438442.
4. Liu F. T., Ting K. M., Zhou Z.-H. Isolation Forest // *Proceedings of the Eighth IEEE International Conference on Data Mining*. 2008. P. 413–422. DOI: 10.1109/ICDM.2008.17.
5. Breiman L. Random Forests // *Machine Learning*. 2001. Vol. 45, No. 1. P. 5–32. DOI: 10.1023/A:1010933404324.
6. Chen T., Guestrin C. XGBoost: A Scalable Tree Boosting System // *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 2016. P. 785–794. DOI: 10.1145/2939672.2939785.
7. Malhotra P., Ramakrishnan A., Anand G., Vig L., Agarwal P., Shroff G. LSTM-Based Encoder-Decoder for Multi-Sensor Anomaly Detection // *ICML Workshop on Anomaly Detection*. 2016.
8. Sakurada M., Yairi T. Anomaly Detection Using Autoencoders with Nonlinear Dimensionality Reduction // *Proceedings of the 2nd Workshop on Machine Learning for Sensory Data Analysis*. 2014. P. 4–11. DOI: 10.1145/2689746.2689747.
9. Poolsappasit N., Dewri R., Ray I. Dynamic Security Risk Management Using Bayesian Attack Graphs // *IEEE Transactions on Dependable and Secure Computing*. 2012. Vol. 9, No. 1. P. 61–74. DOI: 10.1109/TDSC.2011.34.
10. Frigault M., Wang L. Measuring Network Security Using Bayesian Network-Based Attack Graphs // *Proceedings of the 32nd Annual IEEE International Computer Software and Applications Conference*. 2008. P. 698–703. DOI: 10.1109/COMPSAC.2008.89.

Vladyslav Shkolnikov

Doctor of Philosophy in Law, Associate Professor, Head of the Department of Criminology and Information Technologies
National Academy of Internal Affairs, Kyiv, Ukraine
ORCID: 0000-0003-2041-9450
E-mail: shkolnikov.v.i@navs.edu.ua

Vira Huskova

Doctor of Philosophy in Computer Science, Associate Professor of the Department of Artificial Intelligence
National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Kyiv, Ukraine
ORCID: 0000-0001-7637-201X
E-mail: huskova.vira@lil.kpi.ua

Artem Khalygov

PhD Student Department of Applied Informatics, Institute of Telecommunications and Global Information Space
National Academy of Sciences of Ukraine, Kyiv, Ukraine
ORCID: 0009-0006-5465-4650
E-mail: khalygovartem@gmail.com

Bohdan Lysov

PhD Student of the Department of Artificial Intelligence
National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute", Kyiv, Ukraine
ORCID: 0009-0007-7963-6958
E-mail: bogukraine@gmail.com

BEHAVIORAL-AWARE DECISION SUPPORT SYSTEM FOR CYBER RISK ASSESSMENT IN CRITICAL INFRASTRUCTURE BASED ON USER ACTIVITY ANALYSIS

This paper proposes an intelligent decision support system for cyber risk assessment in critical infrastructure based on user behavior analysis. The proposed approach combines two-stage behavioral analysis of users before and after authentication, machine learning methods for anomaly detection, and Bayesian cyber risk assessment within a unified analytical architecture. A mechanism for integrating technical security events, user behavioral characteristics, and contextual factors to form an integrated risk profile is developed. Isolation Forest, XGBoost, and LSTM models are employed for behavioral anomaly detection, while a Bayesian network is used for risk assessment, enabling the consideration of uncertainty and causal relationships among risk factors. Experimental evaluation conducted on the CERT Insider Threat, LANL User Authentication, and CICIDS2017 datasets confirmed the effectiveness of the proposed approach. The LSTM model achieved the best anomaly detection performance, reaching ROC-AUC values of 0.96 for pre-authentication analysis and 0.97 for post-authentication analysis. The proposed system achieved a cyber risk classification accuracy of 0.91 and a recommendation precision of 0.89 with an average response time of 87 ms. The practical significance of the study lies in enabling the transition from anomaly detection to comprehensive cyber risk assessment and the generation of response recommendations for potential threats in critical infrastructure environments.

Keywords: cybersecurity, critical infrastructure, decision support system, cyber risk assessment, user behavior analysis, UBA, UEBA, machine learning, anomaly detection, Bayesian networks, LSTM, XGBoost.

Надійшла до редакції (Received): 23.08.2026

Прийнята до друку (Accepted): 08.09.2026

Опубліковано онлайн (Available online): 15.09.2026

<http://creativecommons.org/licenses/by/4.0/>

This work is licensed under Creative Commons Attribution-noncommercial-sharealike 4.0 International License