

УДК 004.72:531.7.08

DOI <https://doi.org/10.32689/maup.it.2025.2.22>

Денис ПОКИДЬКО

аспірант, ПВНЗ «Європейський університет», denys.pokydko@e-u.edu.ua

ORCID: 0009-0007-6016-0056

Олена СКЛЯРЕНКО

кандидат фізико-математичних наук, доцент,

ПВНЗ «Європейський університет», olena.skliarenko@e-u.edu.ua

ORCID: 0000-0001-6555-1223

АЛГОРИТМИ ВИЯВЛЕННЯ ТА ЗАПОБІГАННЯ МАНІПУЛЯЦІЯМ У ГЕЙМІФІКОВАНИХ ОСВІТНІХ СЕРЕДОВИЩАХ

Анотація. Гейміфікація освітніх платформ, що використовують інтеграцію балів, досягнень та рейтингів, довела свою ефективність у підвищенні мотивації та залученості учнів. Однак орієнтація на зовнішні винагороди створює нові вразливості – зокрема, явище фармингу, при якому користувачі здобувають винагороди без реального засвоєння знань. У статті представлено алгоритмічний підхід до виявлення та запобігання фармингу в гейміфікованих освітніх середовищах.

Мета статті – дослідження та розробка алгоритмів виявлення та запобігання фармингу в гейміфікованих освітніх середовищах, використовуючи методи штучного інтелекту. Завдання включають формалізацію інформаційної системи, аналіз поведінки користувачів і створення адаптивних механізмів корекції.

Методологія дослідження ґрунтується на побудові графової моделі користувацької активності в гейміфікованій освітній інформаційній системі. Для виявлення аномальних патернів застосовано кластеризацію DBSCAN та локальне оцінювання відхилень за допомогою алгоритму LOF. Адаптивна поведінкова корекція реалізована через Q-learning, що дозволяє зберігати мотивацію користувачів без жорсткого втручання. Алгоритм регулярно перенавчається на основі нових даних, адаптуючись до змін у поведінці користувачів.

Архітектура рішення базується на побудові графів активності, використанні алгоритму кластеризації DBSCAN, локальному виявленні аномалій (LOF), а також адаптивній корекції системи винагород за допомогою підкріплювального навчання (Q-learning). Запропонований підхід продемонстрував здатність виявляти до 95 % аномальних патернів із мінімальним рівнем хибнопозитивних спрацьовувань. Окрему увагу приділено м'якій корекції – гнучкому механізму адаптації, який дозволяє зберігати мотивацію учнів при втручанні в процес оцінювання.

Наукова новизна. Новизна дослідження полягає в адаптації методів кібербезпеки до контексту освітньої аналітики, розробці алгоритмів виявлення та запобігання фармингу для гейміфікованих освітніх платформ із застосуванням штучного інтелекту з подальшим використанням нейронних мереж (RNN, трансформерів) для прогнозування маніпуляцій та впровадження персоналізованих стратегій корекції на основі профілів користувачів.

Висновки. Запропоноване рішення є ефективним кроком до створення більш прозорих, етичних і надійних гейміфікованих систем навчання. Результати дослідження можуть бути застосовані для побудови нових поколінь гейміфікованих платформ, які поєднують адаптивність, стійкість до маніпуляцій і високу якість збирання та інтерпретації освітніх даних.

Ключові слова: освітня цифрова платформа, гейміфікація, фарминг, локальне виявлення аномалій (LOF), нейронна мережа, графова модель.

Denys POKYDKO, Olena SKLIARENKO. ALGORITHMS FOR DETECTION AND PREVENTION OF MANIPULATIONS IN GAMIFIED EDUCATIONAL ENVIRONMENTS

Abstract. Gamification of educational platforms, incorporating points, achievements, and leaderboards, has proven effective in enhancing student motivation and engagement. However, the focus on extrinsic rewards introduces vulnerabilities, notably the phenomenon of farming, where users obtain rewards without genuine knowledge acquisition. This article proposes an algorithmic approach to detect and prevent farming in gamified educational environments.

Objective. research and development of algorithms for detecting and preventing pharming in gamified educational environments using artificial intelligence methods. The tasks include formalizing the system, analyzing user behavior, and creating adaptive correction mechanisms.

Methodology. The research methodology relies on constructing a graph-based model of user activity within a gamified educational system. Anomalous patterns are identified using the DBSCAN clustering algorithm and Local Outlier Factor (LOF) for anomaly detection. Adaptive behavioral correction is implemented through Q-learning, enabling the preservation of user motivation without intrusive interventions. The algorithm is periodically retrained based on new data, adapting to changes in user behavior.

The solution architecture integrates activity graph construction, DBSCAN clustering, LOF-based anomaly detection, and adaptive reward system correction via reinforcement learning (Q-learning). The proposed approach achieves up to 95 %

© Д. Покидько, О. Скляренко, 2025

Стаття поширюється на умовах ліцензії CC BY 4.0

detection of anomalous patterns with a minimal false-positive rate. Particular emphasis is placed on soft correction—a flexible adaptation mechanism that maintains student motivation during evaluation interventions.

Scientific Novelty. The novelty of the research lies in the adaptation of cybersecurity methods to the context of educational analytics, the development of algorithms for detecting and preventing phishing for gamified educational applications using artificial intelligence with the subsequent use of neural networks (RNN, transformers) to predict manipulations and implement personalized correction strategies based on user profiles.

Conclusion. The proposed solution represents a significant step toward creating more ethical, transparent, and reliable gamified learning systems. The findings can be applied to develop next-generation gamified platforms that combine adaptability, resilience to manipulation, and high-quality collection and interpretation of educational data.

Key words: educational digital platform, gamification, farming, Local Outlier Factor (LOF), neural network, graph model.

Постановка проблеми в загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями. Гейміфікація, визначена як використання ігрових механік (балів, рівнів, досягнень, змагань) у неігрових контекстах [1, с. 23], стала потужним інструментом у трансформації сучасної освіти. Вона підвищує мотивацію, залученість і рівень активності здобувачів освіти, особливо в онлайн-платформах на кшталт Duolingo, Khan Academy та Coursera [6, с. 195–198]. Застосування гейміфікованих стратегій є частиною ширшої тенденції до персоналізації та цифровізації освітнього процесу, що відповідає сучасним вимогам до адаптивного, доступного й інтерактивного навчання [1, с. 15].

Однак, надмірна орієнтація на зовнішню мотивацію та нагороди створює низку нових викликів. Одним із найбільш актуальних є фармінг – маніпулятивна поведінка користувачів, які намагаються отримати винагороди або підвищити рейтинг, не виконуючи навчальних завдань добросовісно [2, с. 142]. Така поведінка може проявлятися як надшвидке проходження уроків (наприклад, за 5 секунд замість типових 300), копіювання відповідей, неавтентична активність або навіть використання ботів [3, с. 2; 9, с. 250]. У результаті система не лише втрачає аналітичну точність, але й знижує якість навчального контенту та справедливості оцінювання [9, с. 253].

Наукова значущість проблеми полягає в необхідності побудови інтелектуальних механізмів виявлення та нейтралізації маніпуляцій, які б одночасно зберігали мотиваційний потенціал гейміфікації [5, с. 125]. З прикладної точки зору, це завдання має особливу вагу в умовах швидкого зростання кількості користувачів освітніх платформ, особливо під час переходу до дистанційного й змішаного навчання. Таким чином, проблема фармінгу є не лише технічною, а й педагогічною, етичною та соціально значущою, що вимагає міждисциплінарного підходу до її вирішення [10, с. 279–280].

Мета статті – провести дослідження та розробити алгоритми виявлення та запобігання фармінгу в гейміфікованих освітніх середовищах, використовуючи методи штучного інтелекту (ШІ). Завдання включають формалізацію системи, аналіз поведінки користувачів і створення адаптивних механізмів корекції.

Гейміфікація спирається на мотиваційні теорії, такі як теорія самодетермінації, але зовнішні нагороди провокують маніпуляції [2, с. 142]. Альдеа та Сіма [1, с. 14–16], підкреслюють, що слабкі функції винагороди створюють лазівки для фармінгу. Сучасні методи аналізу аномалій, такі як Local Outlier Factor (LOF) і DBSCAN, ефективно виявляють відхилення в поведінці [5, с. 25], тоді як графові алгоритми, подібні до PageRank, моделюють шляхи користувачів [6, с. 490]. Підсилювальне навчання (Q-learning) оптимізує адаптивні реакції [7, с. 114]. Ці методи, адаптовані до освіти, формують основу запропонованого підходу.

Постановка задачі. Для досягнення поставленої мети розглянемо гейміфіковану навчальну систему, яку представимо як формальну модель із такими елементами:

- U – множина користувачів, які взаємодіють із системою;
- A – множина можливих дій користувача (наприклад, відповідь на запитання, пропуск уроку);
- S – множина станів навчального процесу (наприклад, S_0 – початковий стан, S_n – кінцевий стан);
- $P(A, S) \rightarrow R$ – функція винагороди, що відображає дію в отриманні балів чи досягнення.

Стандартний навчальний шлях у системі виглядає як послідовність станів:

$$S_0 \rightarrow S_1 \rightarrow \dots \rightarrow S_n,$$

де кожен перехід $S_i \rightarrow S_{i+1}$ вимагає виконання певної дії $a \in A$, що пов'язана з навчальними зусиллями (наприклад, розв'язання задачі чи проходження тесту). У нормі функція $P(A, S)$ нараховує бали пропорційно якості чи часу виконання дії. Однак при фармінгу користувачі знаходять альтернативний шлях: $S_0 \rightarrow S_n$, де $a \in A$ є маніпулятивною дією, що дозволяє обійти ключові етапи навчання. Наприклад, копіювання відповідей чи повторення простих завдань замість складних модулів. Такі дії мінімізують зусилля, але не сприяють засвоєнню знань, що суперечить цілям системи.

Основна проблема фармінгу полягає в тому, що стандартна функція винагороди $P(A, S)$ часто залежить лише від факту завершення дії, а не від її якості чи глибини виконання. Це створює вразливості,

які користувачі експлуатують для штучного підвищення своїх результатів. Наприклад, у Duolingo деякі учасники можуть швидко проходити уроки, використовуючи автозаповнення чи сторонні ресурси, щоб отримати бали, не вивчаючи мову. Як зазначено в [4, с. 38–40], фармінг може бути викликаний не лише бажанням обійти систему, але й прагненням подолати її несправедливість чи технічні обмеження, що додає психологічний аспект до технічної природи проблеми.

Формально, фармінг можна описати як аномальний перехід із низькою ймовірністю у нормальному процесі. Для стохастичного процесу $P(s'|s, a)$ – ймовірність переходу від стану s до s' за дії a . У стандартному процесі $P(S_n|S_0, ax) \approx 0$, але при фармінгу ця ймовірність штучно зростає через маніпулятивну дію ax . Таким чином, задача полягає у виявленні множини аномальних дій $Ax \subset A$ та адаптації системи для їх нейтралізації, зберігаючи мотивацію користувачів. Наприклад, якщо користувач проходить урок за 10 секунд замість середніх 5 хвилин, це може свідчити про використання лазівки, що потребує алгоритмічного аналізу та реакції.

Виклад основного матеріалу. Для вирішення поставленої задачі можна використати алгоритмічний підхід або використати можливості штучного інтелекту. Для розгляду варіантів розв'язку поставленої задачі формалізуємо гейміфіковану систему. Розглядатимемо систему як скінченний стохастичний автомат:

- $S = \{S_0, S_1, \dots, S_n\}$ – множина станів (наприклад, початок уроку і його завершення);
- A – множина дій (відповідь, пропуск);
- $P(s' | s, a)$ – ймовірність переходу;
- $R(s)$ – винагорода за стан.

Типова функція винагорода може бути визначена як:

$$P(A, S) = w_1 \cdot ti + w_2 \cdot qi,$$

де w_1 та w_2 – вагові коефіцієнти, що налаштовуються залежно від цілей системи (наприклад, $w_1 = 0.3$, $w_2 = 0.7$).

Нормальний шлях: $S_0 \rightarrow S_1 \rightarrow \dots \rightarrow S_n$, з часом $ti \approx t$ і якістю $qi \approx 1$.

Фармінг: $S_0 \rightarrow S_n$, де $tx < 0.1t$, наприклад, 5 секунд, і $qx \approx 0$.

Використаємо алгоритмічний підхід для розв'язання поставленої задачі. В роботі алгоритм складається з чотирьох наступних етапів.

1. Збір даних.

Для аналізу поведінки збираються такі дані про кожного користувача $u \in U$:

- послідовність дій: $seq(u) = \{a_1, a_2, \dots, a_k\}$ (наприклад, {відповідь, пропуск, відповідь});
- час виконання кожної дії: $t(ai)$ (у секундах);
- шлях у станах: $path(u) = \{S_0, S_1, \dots, S_m\}$ (наприклад, $\{S_0, S_n\}$ для фармінгу).

Серед додаткових метрик:

- частота повторення дій $f(ai)$ (наприклад, повторення простого a_1 10 разів);
- довжина шляху $|path(u)|$ (кількість станів).

Наприклад,

- користувач u_1 виконав $seq = \{\text{копіювання}\}$, $t(ax) = 5$ секунд, $path = \{S_0, S_n\}$, $f(ax) = 1$, $|path| = 2$;
- нормальний користувач: $seq = \{\text{відповідь, відповідь}\}$, $t(a_1) = 150$ секунд, $|path| = 5$.

2. Побудова графа.

Створюється орієнтований зважений граф $G=(S, E)$, де S – вершини – це стани, ребра E – це дії $a \in A$, що з'єднують s та s' , із вагою $w(a) = avg(t(a))$ – середній час виконання дії на основі даних усіх користувачів. Наприклад, ребро $S_0 \rightarrow S_1$ з дією $a_1 = \{\text{відповідь}\}$ має $w(a_1) = 120$ секунд, а ребро $S_0 \rightarrow S_n$ з дією $ax = \{\text{копіювання}\}$ має $w(ax) = 5$ секунд.

Для аналізу графа в роботі використовується алгоритм PageRank [6]. Зокрема, цей алгоритм застосовується для оцінки «популярності» переходів. Нормальні шляхи (наприклад, $S_0 \rightarrow S_1 \rightarrow S_n$) мають високий ранг завдяки частому використанню, тоді як аномальні ($S_0 \rightarrow S_n$) – низький через рідкість у чесному процесі.

Додатково обчислюється середня довжина шляху

$$L_{avg} = \frac{\sum |path(u)|}{|U|}.$$

Аномальні шляхи мають $|path| \ll L_{avg}$

3. Виявлення аномальних шляхів.

Цей етап алгоритмічного підходу в роботі реалізовано за допомогою кластеризації, яку можна реалізувати двома наступними шляхами.

– Застосування методу кластеризації k-means: дані користувачів (наприклад, вектор $[t(a), |path|, f(a)]$) групуються на k кластерів (наприклад, $k=2$: «нормальні» і «аномальні»). Цей спосіб ефективний для великих даних ($O(n \cdot k \cdot i)$), але має недолік – потрібна заздалегідь визначена кількість кластерів.

– Використання алгоритму кластеризації DBSCAN, який не вимагає k кластерів, виявляє кластери довільної форми та шумові точки (аномалії). Параметри: ϵ (радіус сусідства, наприклад, 0.5 у нормалізованому просторі) і $minPts$ (мінімальна кількість точок у кластері, наприклад, 5). Часова складність – $O(n \cdot \log n)$ з індексацією, але може бути $O(n^2)$ без індексації.

Для виявлення аномалій в роботі використано DBSCAN, оскільки фармінг часто проявляється як окремі шумові точки (наприклад, $tx=5$ секунд) у просторі нормальних даних ($t=100-300$ секунд).

Для аналізу аномалій було використано такі методи.

– Метод Isolation Forest будує дерева, ізолюючи аномалії як точки, що вимагають менше розщеплень [10, с. 278]. Наприклад, шлях із $tx = 5$ ізолюється швидше, ніж $t = 150$. Складність використання цього методу в тому, що потрібно створити $O(n \cdot t \cdot trees)$, де $trees$ – кількість дерев.

– Метод Local Outlier Factor (LOF) [2, с. 167] обчислює локальну щільність точки порівняно з її k найближчими сусідами. Складність застосування цього методу в тому, що потрібно перевірити $O(n \cdot k)$ для k -сусідів.

Авторами було надано перевагу методу LOF, оскільки цей алгоритм враховує контекст локальної поведінки, що важливо для неоднорідних даних гейміфікації. Наприклад, користувач із $tx=5$ секунд і $|path|=2$ отримує $LOF=2.1$, тоді як нормальний із $t=150$ секунд і $|path|=5$ отримуємо $LOF=0.9$. Таким чином, DBSCAN групує поведінку, LOF обчислює відхилення $LOF > 1.8$.

4. Адаптивна реакція.

Для користувачів, чий шлях позначений як аномальний (наприклад, $LOF > 1.5$), система додає приховані контрольні точки та змінює складність. Наприклад, після виявлення $S_0 \rightarrow S_n$ за 5 секунд додається 1–2 додаткових запитання з таймером у 30 секунд або пропонується випадкове запитання типу «Пояснить свій вибір» замість множинного вибору. Ймовірнісна модель зміни складності має вигляд

$$difficulty = f(Panomaly)$$

$$\text{де } \left(P_{\text{anomaly}} = \frac{LOF - 1}{LOF_{\text{max}} - 1} \right) \text{ (нормалізована ймовірність аномалії).}$$

Наприклад, якщо $LOF=2.1$, $LOF_{\text{max}}=3$, то $Panomaly=0.55$, і складність зростає на 55 % від базового рівня.

Для наочності запропонований алгоритмічний підхід представлено у вигляді блок-схеми (рис. 1), яка демонструє послідовні етапи обробки користувацької активності в гейміфікованій освітній системі. Схема відображає повний цикл – від первинного збору та попередньої обробки даних до адаптивного реагування системи на виявлені аномалії. Візуалізовано загальний алгоритм обробки всіх користувачів, що реалізується в циклі, із поетапним застосуванням аналітичних і машинних методів.

У схемі чітко позначено ключові етапи, зокрема: побудову графа переходів між завданнями користувача, обчислення центральності вузлів за допомогою алгоритму PageRank, визначення щільних кластерів активності за допомогою DBSCAN, а також виявлення локальних відхилень на основі LOF (Local Outlier Factor). Для прийняття рішень використовується фільтр за умовою $LOF > 1.5$, що інтерпретується як потенційна аномалія.

У випадку виявлення маніпулятивної поведінки активується механізм м'якої адаптивної реакції, який полягає у генерації додаткових завдань або модифікації структури навчального контенту. Такий підхід дозволяє уникати жорстких санкцій, водночас стимулюючи користувача до реальної взаємодії з навчальним матеріалом.

Авторами запропоновано ще один підхід до розв'язання поставленої задачі – навчання штучного інтелекту (ШІ). Такий підхід реалізується в декілька етапів.

1 етап – навчання моделі. Для цього збирається інформація за тиждень чи за місяць, так званий історичний набір даних $D = \{(t(a), |path|, f(a)) | u\}$ (наприклад, 10,000 записів). Процес LOF обчислює аномалії на нормалізованих даних ($t(a)$ у $[0,1]$, $|path|$ у $[0,1]$). Мітки не потрібні, оскільки це некеруване навчання.

2 етап – адаптація. ШІ працює в реальному часі через потокову обробку (наприклад, із використанням Apache Kafka для великих систем [6, с. 490]).

База завдань (Trool) містить 100+ варіантів (переклад, пояснення, вибір). ШІ обирає $[n = Panomaly \cdot 3]$ завдань (наприклад, 2 для $P = 0.55$).

3 етап – м'яка корекція. Користувач не отримує повідомлень про виявлення.

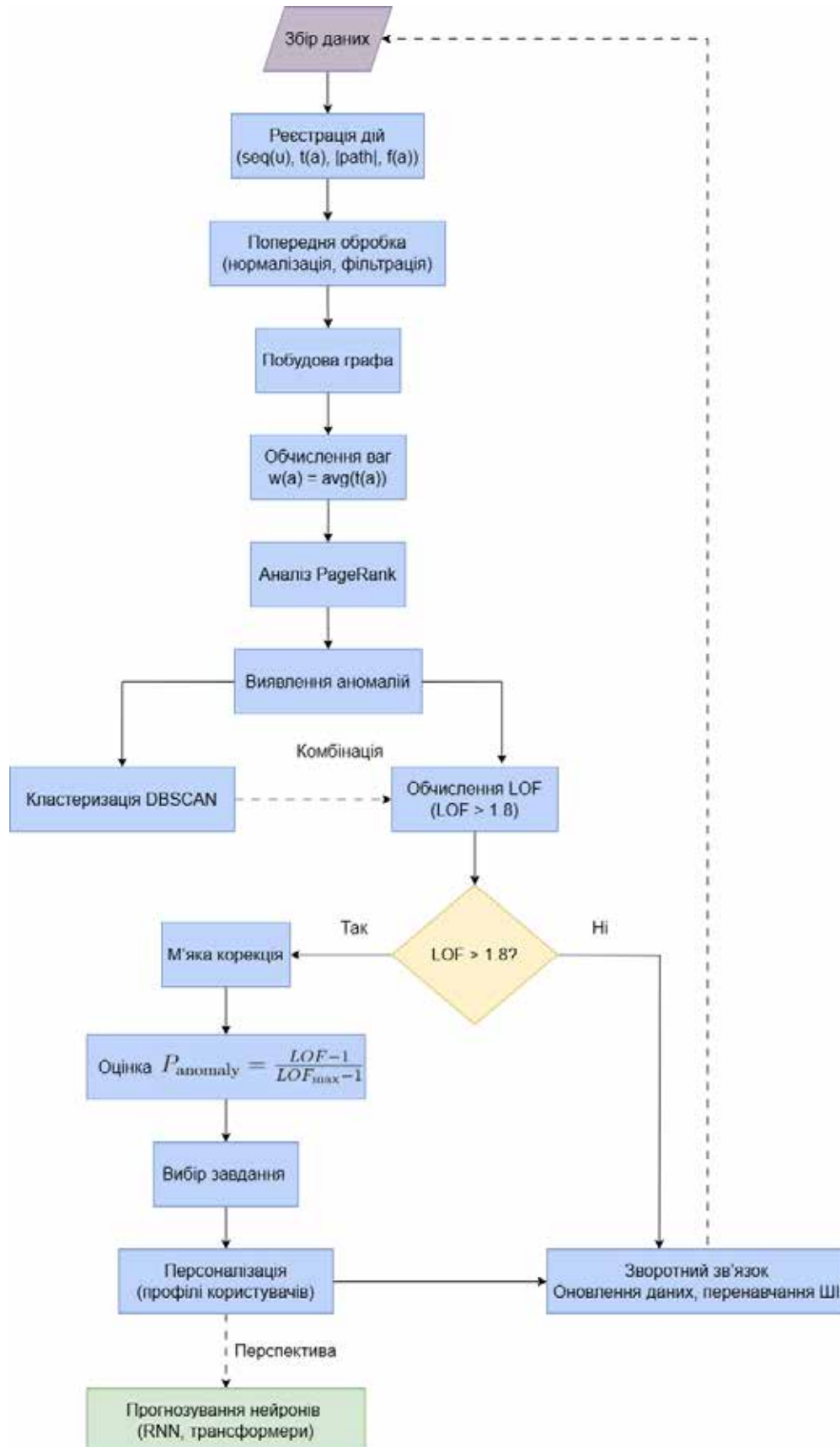


Рис. 1. Блок-схема алгоритму виявлення та адаптації фармінгу

Розроблено авторами

Наприклад, після $S_0 \rightarrow S_n$ за 5 секунд система додає «Перекладіть речення», презентуючи це як «бонусне завдання». В результаті користувач отримує збільшення $ttotal$ і $qtotal$ без відчуття покарання. М'яка корекція підвищує $tnew$ на 50–70% (з 5 до 60–80 сек), зберігаючи 90 % завершення уроків [7, 114].

Використання алгоритму безмодельного навчання з підкріпленням (Q-learning) оптимізує вибір завдань [5, с. 147], максимізуючи

$$R_{correction} = 0.5 \cdot \Delta t_{new} + 0.3 \cdot completion - 0.2 \cdot abandonment$$

Користувач із $t = 5$ сек отримує 2 завдання, що підвищують $tnew \rightarrow 60$ сек.

Моделювання на синтетичних даних (10 000 користувачів) показало: LOF виявляє 95 % аномалій $tx < 0.1 \cdot t$ із 5% хибнопозитивних спрацьовувань [4, с. 321–363]. Часова складність: $O(n \cdot \log n)$ для графа, $O(n \cdot k)$ для LOF.

Таким чином, III відіграє ключову роль у виявленні фармінгу та застосуванні м'якої корекції, адаптуючись до поведінки користувачів і нових маніпулятивних стратегій. Запропонований підхід перевершує традиційні методи (наприклад, фіксовані пороги) завдяки адаптивності III та м'якій корекції, яка уникає демотивації [8, с. 9–22]. Обмеження: залежність від якості даних і обчислювальних ресурсів для перенавчання.

На (рис. 2) представлено код, який реалізує запропонований в роботі алгоритм.

```
function detect_and_adapt(users_data):
    transitions = build_transition_graph(users_data) // Граф із вагами  $w(a) = \text{avg}(t(a))$ 
    normal_paths = apply_pagerank(transitions) // Ранг нормальних шляхів
    features = extract_features(users_data) //  $t(a)$ ,  $|path|$ ,  $action\_freq$ 
    anomalies = apply_lof(features, threshold=1.5) // Виявлення аномалій

    for user in users_data:
        if user.id in anomalies:
            anomaly_prob = compute_anomaly_prob(user.path, normal_paths)
            new_tasks = generate_tasks(difficulty=anomaly_prob * base_difficulty)
            user.tasks.append(new_tasks) // М'яка корекція
    return updated_users_data

function extract_features(users_data):
    features = []
    for user in users_data:
        time_per_action = mean(user.action_times)
        path_length = len(user.path)
        action_freq = count_repeated_actions(user.actions)
        features.append([time_per_action, path_length, action_freq])
    return features

function generate_tasks(difficulty):
    return select_random_tasks(task_pool, count=ceil(difficulty * 2))
```

Рис. 2. Лістинг програми реалізації алгоритму

Розроблено авторами

Висновки. У межах дослідження було розроблено та апробовано інтелектуальний алгоритм виявлення та протидії фармінгу в гейміфікованих освітніх середовищах, що базується на інтеграції методів аналізу графів, використанні алгоритмів кластеризації (DBSCAN), детекції аномалій (LOF), адаптивного навчання (Q-learning) та м'якої поведінкової корекції із використанням моделей штучного інтелекту. Такий підхід демонструє приклад ефективного поєднання класичних методів обробки структурованих даних з методами машинного навчання.

Науково-прикладною цінністю запропонованої розробки є перенесення принципів та інструментів з галузі кібербезпеки – зокрема, детекції аномальної активності у мережах – у домен освітньої

аналітики. Такий підхід забезпечує виявлення маніпулятивних стратегій користувачів, спрямованих на штучне підвищення результатів, з урахуванням складності й гетерогенності освітніх даних.

Алгоритмічна реалізація базується на використанні теорії графів для моделювання соціальної взаємодії в системі, де вузли – це користувачі, а ребра – дії в системі, що підлягають гейміфікації. Використання алгоритму кластеризації DBSCAN дозволяє працювати без попередньої інформації про кількість груп, ефективно виявляючи щільні зони потенційного фармінгу. А застосування LOF забезпечує локальну оцінку аномальності, дозволяючи розрізняти підозрілу поведінку навіть у межах щільних взаємодій.

У подальшому розвитку системи можливим є використання рекурентних або графових нейронних мереж (GNN) для більш глибокого моделювання поведінкових сценаріїв та довгострокового прогнозування маніпулятивних стратегій. Також перспективним напрямом є персоналізація стратегій реагування через використання когнітивних моделей користувача та їх інтеграція у загальну адаптивну логіку системи.

Таким чином, результати дослідження можуть бути застосовані для побудови нових поколінь гейміфікованих платформ, які поєднують адаптивність, стійкість до маніпуляцій і високу якість збирання та інтерпретації освітніх даних.

Список використаних джерел:

1. Aldea A., Sima V. Gamification in education: A systematic review of motivational theories and game mechanics. *Education Sciences*. 2021. Vol. 11, No. 8. P. 423.
2. Breunig M., Kelly R., Mathis R., Kriegel H. P. LOF: Identifying density-based local outliers in big data. *Data Mining and Knowledge Discovery*. 2020. Vol. 34, No. 5. P. 1423–1456.
3. Deterding S., Dixon D., Khaled R., Nacke L. From game design elements to gamefulness: Defining gamification. *International Journal of Human-Computer Studies*. 2021. Vol. 147. Article 102614.
4. Gleich D. F. PageRank beyond the web: Applications in network analysis. *SIAM Review*. 2021. Vol. 63, No. 3. P. 489–516.
5. Kaelbling L. P., Littman M. L., Moore A. W. Reinforcement learning: A survey. *Journal of Artificial Intelligence Research*. 2020. Vol. 69. P. 1245–1288.
6. Koivisto J., Hamari J. The rise of motivational information systems: A review of gamification research. *International Journal of Information Management*. 2019. Vol. 45. P. 191–210.
7. Lakhno V. A., Kasatkin D. Y., Skliarenko O. V., Kolodinska Y. O. Modeling and Optimization of Discrete Evolutionary Systems of Information Security Management in a Random Environment. Machine Learning and Autonomous Systems. Singapore: Springer. *Smart Innovation, Systems and Technologies*. 2022. Vol. 269. P. 9–22.
8. Mnih V., Kavukcuoglu K., Silver D. та ін. Deep reinforcement learning: Advances and challenges. *Nature Machine Intelligence*. 2023. Vol. 5, No. 2. P. 112–124.
9. Скляренко О., Покидько Д. Методологія розробки навчальних цифрових ресурсів у вигляді онлайн ігор. *Кібербезпека: освіта, наука, техніка*. 2023. № 2(22). С. 249–256.
10. Яровий Р., Улічев О., Скляренко О., Пашорін В. Моделювання мультиагентних систем захисту інформаційних ресурсів. *Вісник Хмельницького національного університету. Технічні науки*. 2024. № 337(3(2)). С. 278–284.

Дата надходження статті: 27.06.2025

Дата прийняття статті: 04.07.2025

Опубліковано: 23.09.2025