

УДК 004.056.53

DOI <https://doi.org/10.32782/IT/2025-1-5>

Андрій ГЛОВА

аспірант, Приватний вищий навчальний заклад «Європейський університет», бульв. Академіка Вернадського, 16-в, м. Київ, Україна, 03115

ORCID: 0009-0001-2108-0762

Бібліографічний опис статті: Глова, А. (2025). Підвищення надійності результатів симуляції фішингових атак шляхом фільтрації ботів. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 33–41, doi: <https://doi.org/10.32782/IT/2025-1-5>

ПІДВИЩЕННЯ НАДІЙНОСТІ РЕЗУЛЬТАТІВ СИМУЛЯЦІЇ ФІШИНГОВИХ АТАК ЗА ДОПОМОГОЮ ФІЛЬТРАЦІЇ БОТІВ

Симуляція фішингових атак є важливим інструментом для підвищення обізнаності та навичок співробітників у сфері кібербезпеки, особливо в умовах зростання кіберзагроз, таких як фінансове шахрайство та витоки даних. Подібне моделювання дозволяє компаніям оцінити ефективність своїх навчальних програм і визначити, чи готові співробітники розпізнавати фішингові атаки та реагувати на них. Однак надійність результатів симуляції фішингових атак часто порушується автоматичними кліками (переходами за посиланням), які генеруються ботами – частиною корпоративних систем захисту електронної пошти.

Це дослідження має на меті усунути такі спотворення шляхом розробки ефективного алгоритму фільтрації ботів, щоб забезпечити точне відображення поведінки людини у даних симуляції.

Методологія. Дослідження використовує системний підхід для розрізнення взаємодії ботів та людей із фішинговими електронними листами. Методологія включає аналіз поведінки ботів, виявлення характерних шаблонів і розробку алгоритму фільтрації на основі експериментальних даних. Валідація виконувалася через порівняльний аналіз результатів симуляцій до і після впровадження алгоритму.

Наукова новизна. У дослідженні запропоновано алгоритм фільтрації автоматизованих взаємодій у симуляціях фішингових атак та проведено аналіз його ефективності. Алгоритм базується на емпіричному визначенні часових інтервалів між подіями відкриття електронного листа та кліку на посилання, що дозволяє ідентифікувати бот-кліки та відокремити їх від дій реальних користувачів. Запропонований підхід експериментально перевірено на корпоративних даних, що дозволило оцінити його точність і надійність. Використання алгоритму сприяє покращенню аналізу результатів фішингових симуляцій, підвищенню достовірності оцінки поведінки користувачів та ефективності навчальних програм з кібербезпеки.

Висновки. Реалізація алгоритму фільтрації ботів знижує спотворення даних симуляції фішингових атак, що дозволяє точніше оцінювати готовність співробітників та підвищувати ефективність навчальних програм. Цей підхід не лише покращує кібербезпеку організації, але й сприяє формуванню культури пильності та безперервного навчання для протидії сучасним кіберзагрозам.

Ключові слова: симуляція фішингових атак, навчання кібербезпеці, фільтрація ботів, обізнаність співробітників, точність симуляції.

Andriy GLOVA

Postgraduate Student, Private Higher Education Establishment "European University", 16-v, Vernadskyi Boul., Kyiv, Ukraine, 03115, andriy.glova@e-u.edu.ua

ORCID: 0009-0001-2108-0762

To cite this article: Glova, A. (2025). Pidvyshchennia nadiinosti rezultativ symuliatsii fishynhovykh atak shliakhom filtratsii botiv [Improving the reliability of phishing simulation results through bot filtering]. *Information Technology: Computer Science, Software Engineering and Cyber Security*, 33–41, doi: <https://doi.org/10.32782/IT/2025-1-5>

IMPROVING THE RELIABILITY OF PHISHING SIMULATION RESULTS THROUGH BOT FILTERING

Phishing simulations are an essential instrument for improving employee awareness and skills in cybersecurity, particularly in the face of escalating cyber threats, such as financial fraud and data breaches. These simulations enable organizations to evaluate the effectiveness of their training programs and measure employee readiness to detect and respond to phishing attempts. However, the reliability of phishing simulation results is often compromised by automated clicks generated by bots, which are part of corporate email security systems.

This study aims to address these distortions by developing an effective bot filtering algorithm to ensure that simulation data accurately reflects human behavior.

Methodology. The study proposes an algorithm for filtering automated interactions in phishing attack simulations and analyzes its effectiveness. The algorithm is based on the empirical determination of time intervals between the opening of an email and clicking on a link, allowing for the identification of bot clicks and their differentiation from real user actions. The proposed approach has been experimentally validated using corporate data, enabling the assessment of its accuracy and reliability. The implementation of this algorithm enhances the analysis of phishing simulation results, improves the reliability of user behavior assessment, and increases the effectiveness of cybersecurity training programs.

Scientific novelty. The study presents a novel approach to filtering bot interactions in phishing simulations by introducing experimentally validated criteria. The proposed algorithm is based on rapid-response patterns observed in bot activity, allowing organizations to exclude such interactions and focus on genuine user behavior. This innovation significantly enhances the reliability of phishing simulation results.

Conclusions. The implementation of the bot filtering algorithm reduces distortions in phishing simulation data, enabling a more accurate assessment of employee readiness and improving the effectiveness of training programs. This approach not only enhances organizational cybersecurity but also fosters a culture of vigilance and continuous learning to counter evolving cyber threats.

Key words: phishing simulations, cybersecurity training, bot filtering, employee awareness, simulation accuracy.

Актуальність теми. Симуляція фішингових атак є важливим інструментом для підвищення обізнаності співробітників щодо кібербезпеки, оцінки стійкості до фішингових атак і вдосконалення програм навчання. У середовищі зростання кіберзагроз ці симуляції допомагають організаціям виявляти слабкі місця в захисті, покращувати навички співробітників у визначенні зловмисних повідомлень і підвищувати їхню здатність до повної стійкості проти кіберзагроз. Однак точність і надійність цього моделювання значною мірою залежить від надійності зібраних даних. Автоматизовані кліки ботами, які належать до системи безпеки електронної пошти компанії, часто призводять до пошкодження даних. Це створює хибне враження про поведінку співробітників, що може призвести до помилкових висновків щодо їхнього навчання кібербезпеці та, у свою чергу, негативно вплинути на розвиток навчання та стратегій захисту (Zhou et al., 2022; Oest et al., 2019). Відповідно до звіту IBM про вартість витоку даних за 2024 рік, до 2024 року на фішингові атаки припадає майже 30 % усіх витоків даних у всьому світі. Середня вартість витоку даних, пов'язаного з шахрайством, становить приблизно 4,88 мільйона доларів США за інцидент (IBM, 2024). У другій половині 2024 року кількість атак, спрямованих на викрадення облікових даних, зросла на 703 % (DataBreaches, 2024). Тільки за перші п'ять місяців 2024 року компанії та споживачі втратили 1,6 мільярда доларів через кіберзлочинність (Panda Security, 2024). З 2013 до кінця 2023 року збитки від атак на компрометацію бізнес-електронної пошти (BEC) у всьому світі перевищили 5,5 мільярдів доларів США, у тому числі понад 20 мільярдів доларів у Сполучених Штатах (Journal of HIPAA, 2024)).

За таких умов організації повинні впроваджувати системи оцінювання, які враховують поведінку працівників, а не результати, спотворені діяльністю роботів. Це вимагає використання надійних методів фільтрації ботів, які дозволяють збирати точні дані про реальну взаємодію користувачів із моделюванням (Chen та ін., 2022). Використання таких підходів не тільки дозволяє правильно оцінити рівень обізнаності співробітників, але й створює більш ефективні програми навчання для зниження ризику фішингових атак (Song & Wang, 2021).

Аналіз останніх досліджень і публікацій. Симуляції фішингових атак та пов'язані з ними виклики активно досліджуються, при цьому дедалі більше уваги приділяється підвищенню надійності даних:

Оест та ін. (2019) у своєму дослідженні наголосили на викликах, пов'язаних з автоматизованими кліками, які можуть створювати хибні позитивні результати у симуляціях. Вони розробили масштабовані методи аналізу, що дозволяють зменшити вплив ботів, зокрема через виявлення методів обману, які використовують зловмисники для обходу систем виявлення. Їхній підхід базується на комбінуванні поведінкових метрик із даними про часові інтервали, що дозволяє виявляти атипові кліки. Ці результати підтвердили необхідність постійного вдосконалення механізмів класифікації автоматизованої активності (Oest et al., 2019).

У дослідженні Huang і Chen (2023) розглядається використання машинного навчання в адаптивному навчанні працівників. Алгоритми дозволяють моделювати моделі поведінки реальних користувачів, враховуючи аномалії, які можуть бути внесені ботами. Автори запропонували інтегрувати ці алгоритми в корпоративні системи навчання для підвищення

ефективності програм кібербезпеки. Це не тільки допомагає відфільтрувати автоматизовану діяльність, але й дає змогу краще зрозуміти, як реальні співробітники реагують на фішингові атаки (Huang & Chen, 2023).

Ван і Сонг (2021) використовували ботів для моделювання людської діяльності та запропонували еволюційну модель гри для аналізу взаємодії між системами захисту та зловмисниками. Їхній підхід передбачає оцінку ефективності різних стратегій для пом'якшення впливу ботів на результати моделювання. Вони наголосили на необхідності адаптивних систем, які враховують зміну тактики зловмисників і адаптуються до нових загроз (Wang & Song, 2021).

У дослідженні Чен та ін. (2024) підкреслюють потенціал великих мовних моделей у виявленні підозрілої діяльності. Автори пропонують використовувати аналіз поведінкових патернів для класифікації дій, які виконують роботи. Моделі на основі штучного інтелекту показали свою високу ефективність у розпізнаванні нетипових дій, таких як надзвичайно швидкі click-події або нестандартні послідовності дій. Це відкриває нові перспективи для включення інтелектуальних систем у симуляцію фішингових атак (Chen et al., 2024). Аналіз дослідження показує, що вдосконалення механізму автоматичного виявлення активності є ключовим фактором підвищення надійності симуляції фішингової атаки. Зокрема, поєднання поведінкових показників, алгоритмів машинного навчання, адаптивних ігрових моделей і великих мовних моделей відкриє нові можливості для вдосконалення програм кібербезпеки.

Перевагою ж підходу запропонованого в цьому дослідженні є його простота та швидкодія, адже він не вимагає складного аналізу додаткових даних.

Мета дослідження. Це дослідження спрямоване на розробку та тестування алгоритму для фільтрації взаємодії ботів у симуляції фішингу. Вимірюючи поведінку реальних користувачів, а також ботів, алгоритм має на меті підвищити точність і надійність оцінювання навчання.

Виклад основного матеріалу. Для проведення дослідження було використано статистичні дані взаємодії корпоративних користувачів із фішинговими симуляціями, що охоплювали **8 організацій** із різних галузей:

- **Фінансовий сектор** – 3 компанії (банки, страхові компанії, інвестиційні установи).
- **Інформаційні технології** – 2 компанії (постачальники ІТ-послуг та розробники програмного забезпечення).
- **Охорона здоров'я** – 2 компанії (медичні заклади, фармацевтичні компанії).

- **Логістика та транспорт** – 1 компанія (міжнародні перевезення).

Загалом вибірка містила **11 000 корпоративних облікових записів**. У процесі аналізу було досліджено **понад 35 400 пар подій Open-Click**, що відображали взаємодію користувачів із фішинговими листами.

Основними досліджуваними параметрами були:

- Дата та час відкриття електронного листа (подія Open).
- Дата та час кліку на фішингове посилання (подія Click).
- Тип пристрою та операційної системи.
- IP-адреса користувача або автоматизованої системи.
- User-Agent (ідентифікатор браузера або програми, яка виконала дію).
- Дані про джерело доступу (локальна мережа компанії чи зовнішня IP-адреса).

Зібрані дані були попередньо оброблені шляхом видалення дублікатів, виключення неповних записів та ідентифікації корпоративних систем захисту за IP-адресами та User-Agent.

Аналіз причин виникнення бот-кліків. Під час аналізу було виявлено, що автоматизовані кліки могли складати **від 40 % до 100 %** усіх переходів за фішинговими посиланнями залежно від рівня захисту корпоративного середовища. Основними джерелами автоматичних кліків є:

- **Антивірусне програмне забезпечення та системи захисту кінцевих точок.** Вони автоматично перевіряють посилання в електронних листах одразу після їх отримання (Ahmad et al., 2024).
- **Попередній перегляд посилань у поштових клієнтах.** Деякі мобільні поштові додатки виконують автоматичний рендерінг URL-адрес (Goenka et al., 2024).
- **Пересилання електронних листів.** При повторному скануванні поштовими серверами після пересилки (Laghari et al., 2024).

- **Позначення листа як фішингового.** Автоматизовані системи безпеки здійснюють перевірку позначених користувачами підозрілих повідомлень (Sadeghpour & Vlajic, 2021).

Ці фактори значно спотворюють результати фішингових тестувань, викликаючи необхідність розробки механізмів для відокремлення автоматичних взаємодій від реальних користувачів.

Виявлення поведінкових патернів ботів. У рамках дослідження було проведено аналіз поведінки ботів під час симуляцій фішингових атак. Основна увага приділялася часовим інтервалам між подіями Open і Click.

Методика.

- **Обсяг вибірки:** 11 000 корпоративних облікових записів.
- **Загальна кількість аналізованих подій:** 35 400 пар Open-Click.
- **Основний критерій:** швидкість переходу за посиланням після відкриття листа.

Результати.

95 % ботів виконували кліки протягом ≤ 0.37 секунди після відкриття електронного листа (рис. 1).

Розподіл часу ботів мав чітко виражений пік на рівні 0.2 секунди, що є характерним для автоматизованих систем (рис. 1).

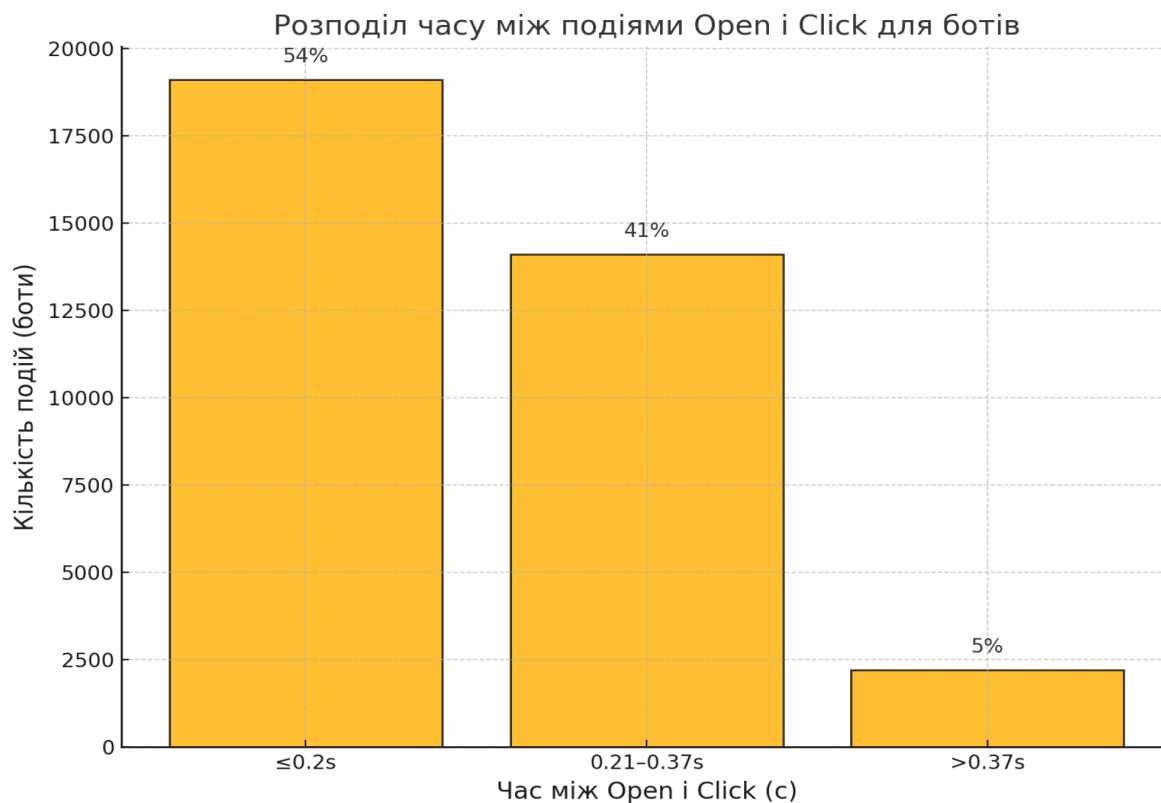


Рис. 1. Розподіл часу між подіями Open і Click для ботів

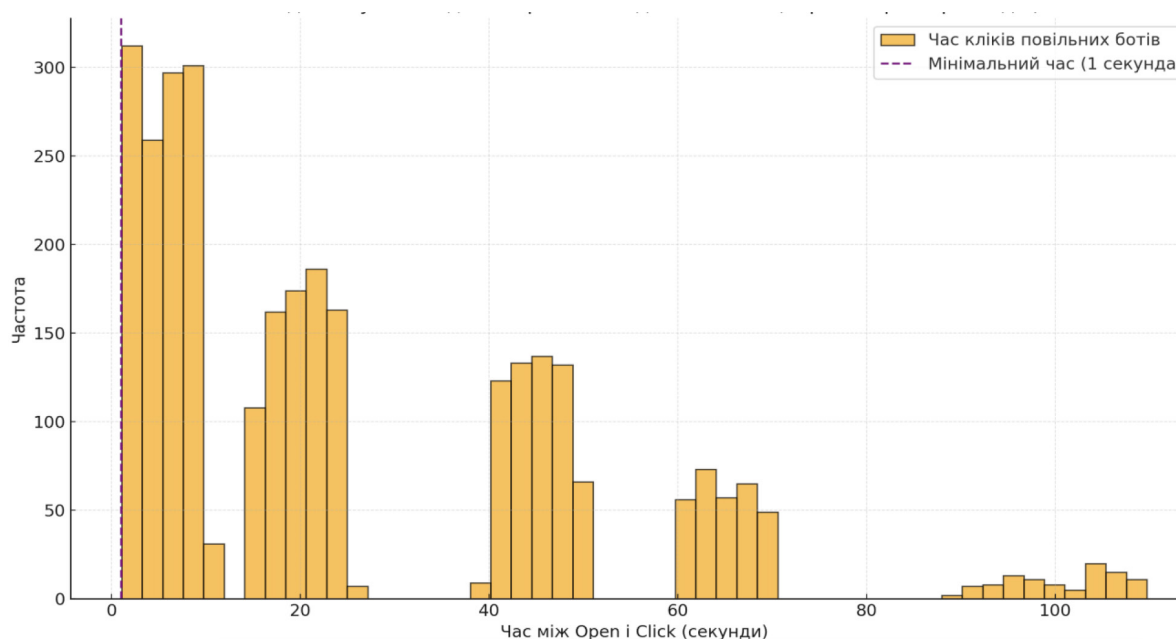


Рис. 2. Розподіл часу між подіями Open і Click для ботів у категорії > 0.37 секунди

5 % ботів які попали у категорію >0.37 секунд показали суттєво відмінний результат від решти 95 %, час між Open і Click подіями для таких ботів попадав у інтервал від 1 до 110 секунд, тим самим демонстрували поведінку притаманну реальним користувачам (рис. 2).

Ці результати вказують на можливість використання часових інтервалів як критерію для ідентифікації бот-кліків. Структура розподілу часу дозволяє створити алгоритми фільтрації, які значно підвищують точність симуляцій.

Мінімальний час реакції користувачів.

Для оцінки мінімально можливого часу реакції користувачів у реальних умовах було проведено експериментальне тестування.

Методика.

- **Кількість учасників:** 10 осіб із різних корпоративних відділів.

- **Використані пристрої:** ноутбуки, смартфони, планшети.

- **Операційні системи:** Windows, Android, iOS.

- **Поштові клієнти:** Gmail, Outlook.

- **Критерій аналізу:** мінімальний час між відкриттям листа та кліком на посилання.

Результати (Рис. 3, Таб. 1):

- **Найшвидший зафіксований час:** 0.71 секунди (за оптимальних умов).

- **Середній час:** 1.31 секунди.

- **Варіативність:** Від 0.71 до 2.5 секунд залежно від пристрою, типу клієнта та технічних обмежень.

Жоден із реальних користувачів не зміг натиснути на посилання швидше, ніж за 0.71 секунди, що підтвердило доцільність використання цього значення як нижньої межі для фільтрації ботів.

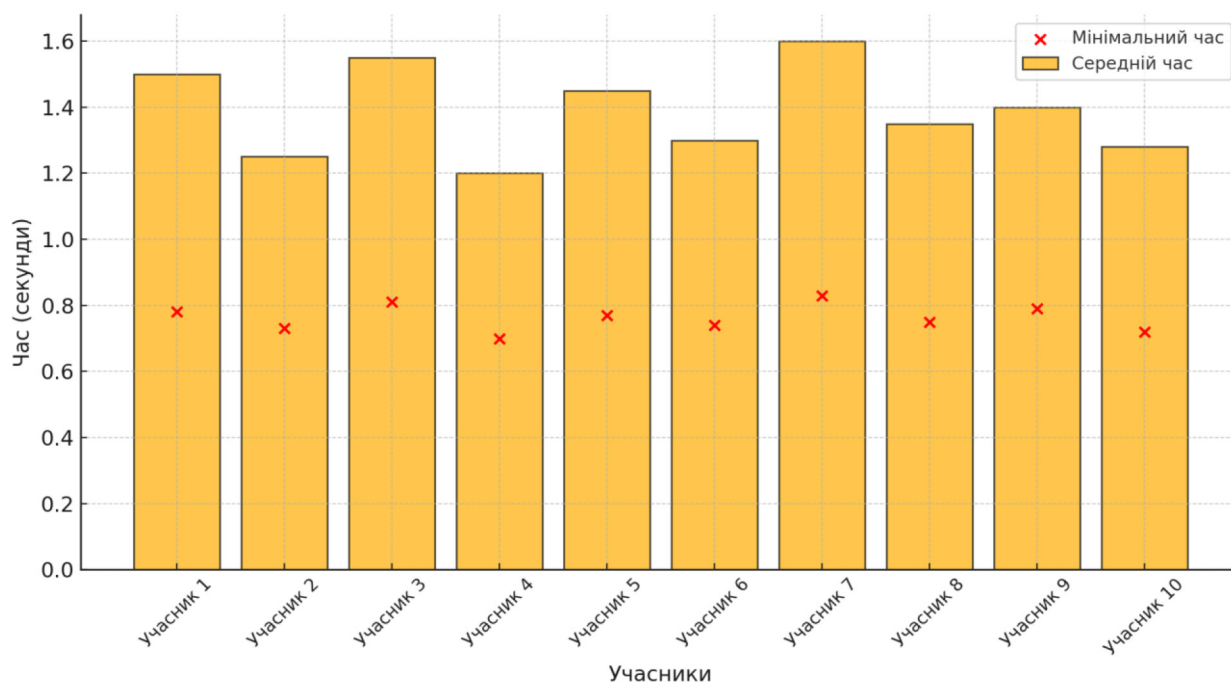


Рис. 3. Розподіл часу між подіями Open і Click реальних користувачів

Таблиця 1

Розподіл часу між подіями Open і Click реальних користувачів

Учасник	Пристрій	ОС	Поштовий клієнт	Час (середнє)	Час (мінімум)
Учасник 1	Ноутбук	Windows	Outlook	1.5	0.78
Учасник 2	Мобільний	Android	Gmail	1.25	0.73
Учасник 3	Ноутбук	macOS	Apple Mail	1.55	0.81
Учасник 4	Мобільний	iOS	Apple Mail	1.2	0.71
Учасник 5	Ноутбук	Windows	Outlook	1.45	0.77
Учасник 6	Мобільний	Android	Gmail	1.3	0.74
Учасник 7	Ноутбук	macOS	Apple Mail	1.6	0.83
Учасник 8	Мобільний	iOS	Apple Mail	1.35	0.75
Учасник 9	Ноутбук	Windows	Outlook	1.4	0.79
Учасник 10	Мобільний	Android	Gmail	1.28	0.72

Розробка алгоритму фільтрації бот-кліків.

На основі отриманих даних було запропоновано алгоритм (1) класифікації ботів, що базується на часовому інтервалі між подіями Open і Click.

$$Bot_Click = \begin{cases} 1, & \text{if } \Delta t \leq 0.54 \text{ sec}(BotClick) \\ 0, & \text{if } \Delta t > 0.54 \text{ sec}(HumanClick) \end{cases} \quad (1)$$

Де:

- Δt – часовий інтервал між подіями Open і Click в секундах.

- 0.54 – граничне значення часу, яке використовується для класифікації події як бота або реального користувача. Це значення є медіаною між граничним часом реакції 95 % ботів (0.37 секунди) та мінімальним часом реакції людини (0.71 секунди).

Таким чином ми отримали алгоритм у впровадженні, оскільки базується лише на часових інтервалах та може бути легко адаптована до різних корпоративних середовищ за допомогою корекції інтервалу Δt , де специфіка роботи може впливати на часові інтервали.

Впровадження алгоритму та його ефективність. Для перевірки ефективності алгоритму він був застосований до реальних корпоративних даних:

- **Обсяг вибірки:** 40 000 подій Open-Click.
- **Галузі:** фінанси, ІТ, охорона здоров'я, логістика.

Дані були очищені для виключення дублікатів і подій із відсутніми часовими мітками.

Це гарантувало, що всі взаємодії мали необхідні атрибути для аналізу. До кожної події було застосовано алгоритм (1), що базується на типовому діапазоні часових інтервалів між подіями Open і Click, характерному для реальних користувачів. Події, які відхилялися від цих інтервалів, класифікувалися як бот-кліки. Для забезпечення точності класифікації дані були зіставлені з шаблонами поведінки, отриманими з попередніх досліджень про поведінку користувачів і ботів. У вибіркових даних (приблизно 5 % від загальної кількості подій) вручну перевірялися часові інтервали та параметри запитів, такі як IP-адреси, агент користувача та інші метадані. Це дозволило впевнитися, що взаємодії, класифіковані як людські, відповідали характеристикам реальних користувачів.

Результати (рис. 4):

- **95 % бот-кліків** були коректно ідентифіковані та відфільтровані.

- **5 % ботів** в середньому (варіюються від 2 % до 7 % залежно від компанії) залишилися невиявленими через те, що вони діяли поза межами типових часових інтервалів або використовували вдосконалені механізми імітації людської поведінки, такі як додавання випадкових затримок.

Висновки. Алгоритм показав високу ефективність, забезпечивши виявлення близько 95 % автоматизованих кліків. Це дозволяє значно підвищити точність симуляцій

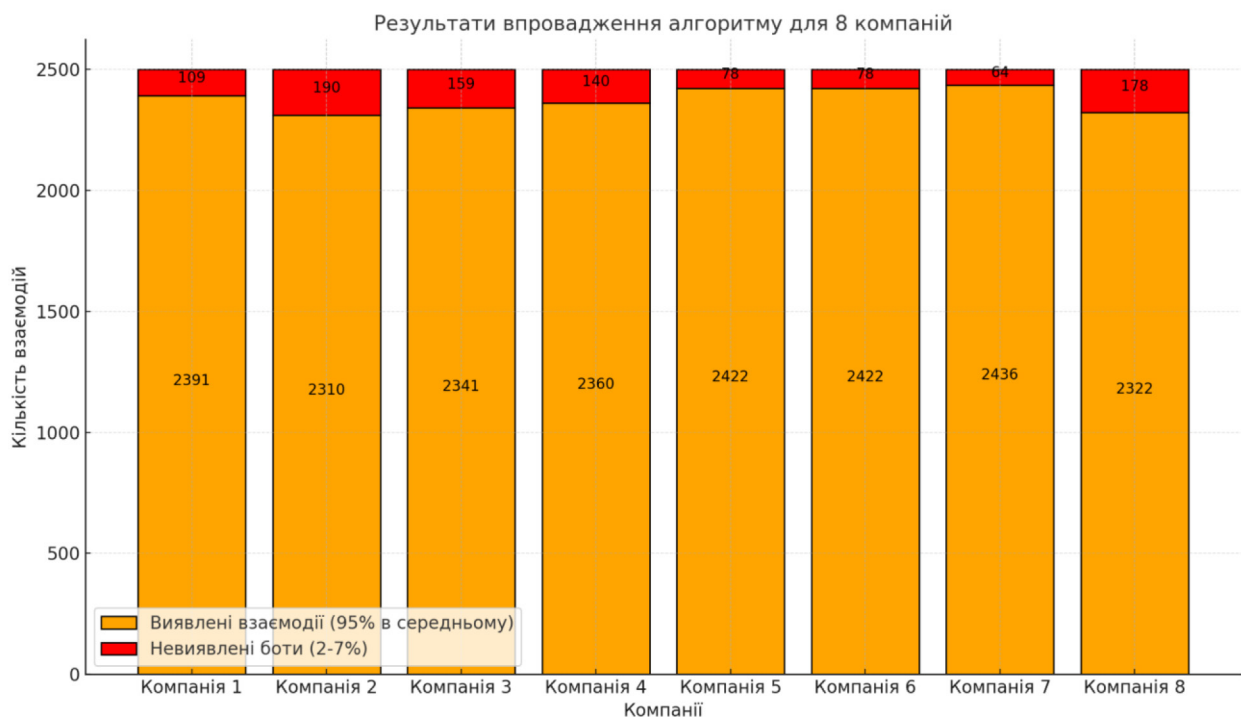


Рис. 4. Результати після впровадження алгоритму

фішингових атак, що важливо для оцінки рівня обізнаності співробітників щодо кіберзагроз. Встановлений поріг 0.54 секунди забезпечує мінімізацію хибних класифікацій. Інтервали більше 0.71 секунди в 95 % випадках належать реальним користувачам. Мають місце обмеження, 5 % ботів уникнули виявлення через більш складну поведінку. Це свідчить про необхідність додаткових критеріїв для аналізу. Алгоритм можна легко адаптувати до конкретного середовища шляхом корекції параметру Δt .

Подальших дослідження будуть націлені на розширення критеріїв ідентифікації:

- Додавання до алгоритму аналіз IP-адрес для виявлення повторюваних або підозрілих патернів, що дозволить покращити точність ідентифікації ботів у різних середовищах (Putra et al., 2022).
- Відключення перевірки User Agents для визначення програмних середовищ, у яких працюють боти. Це може допомогти ідентифікувати складні автоматизовані системи (Suchacka, 2021).
- Також планується протестувати алгоритм в різних корпоративних середовищах із використанням більшої кількості даних, що дозволить оцінити його ефективність у широкому спектрі умов і сценаріїв.

ЛІТЕРАТУРА:

1. Zhou Y., Cui X., Qu W., Ge Y. The effect of automation trust tendency, system reliability and feedback on users' phishing detection. *Applied Ergonomics*. 2022. URL: <https://www.sciencedirect.com/science/article/pii/S0003687022000771> (дата звернення: 28.01.2025).
2. Oest A., Safaei Y., Doupé A., & Ahn G. J. Phishfarm: A scalable framework for measuring the effectiveness of evasion techniques against browser phishing blacklists. *IEEE Symposium on Security and Privacy*. 2019. URL: <https://ieeexplore.ieee.org/document/8835369> (дата звернення: 28.01.2025).
3. IBM. Cost of a Data Breach Report 2024. URL: <https://keepnetlabs.com/blog/171-cyber-security-statistics-2024-s-updated-trends-and-data> (дата звернення: 28.01.2025).
4. DataBreaches. Credential phishing attacks up over 700 percent. URL: <https://databreaches.net/2024/12/18/credential-phishing-attacks-up-over-700-percent> (дата звернення: 28.01.2025).
5. Panda Security. Cybercrime report for the first five months of 2024. URL: <https://www.pandasecurity.com/en/mediacenter/fbi-internet-crime-complaint-center-reports-losses> (дата звернення: 28.01.2025).
6. HIPAA Journal. FBI BEC warning: \$5.5 billion lost. URL: <https://www.hipaajournal.com/fbi-bec-warning-55-billion-lost> (дата звернення: 28.01.2025).
7. Song L., Wang M. Efficient defense strategy against spam and phishing email: An evolutionary game model. *Journal of Information Security and Applications*. 2021. URL: <https://www.sciencedirect.com/science/article/pii/S2214212621001617> (дата звернення: 28.01.2025).
8. Chen J., Mishler, S., Hu B., Li N., Proctor R. The impact of reliable methodologies for bot filtering on user interaction simulations. *International Journal of Human-Computer Studies*. 2022. URL: <https://www.sciencedirect.com/science/article/pii/S0306457322002989> (дата звернення: 28.01.2025).
9. Oest A., Safaei Y., Dupe A., An G. J. Scalable methods for evaluating phishing attack evasion effectiveness. *IEEE Symposium*. 2019. URL: https://www.usenix.org/system/files/sec20fall_oest_prepub.pdf (дата звернення: 28.01.2025).
10. Huang D., Chen N. Machine learning for adaptive cybersecurity training. *Cybersecurity Applications Journal*. 2023. URL: https://www.researchgate.net/publication/380771681_Study_on_Empowering_Cyber_Security_by_Using_Adaptive_Machine_Learning_Methods (дата звернення: 28.01.2025).
11. Wang M., Song L. An evolutionary game model for spam and phishing protection. *Journal of Information Security and Applications*. 2021. URL: <https://www.sciencedirect.com/science/article/pii/S2214212621001617> (дата звернення: 28.01.2025).
12. Chen K., Li J., Zhang H., Zhao S. Utilizing large language models for cyber threat detection. *Frontiers in AI*. 2024. URL: <https://arxiv.org/pdf/2405.04760> (дата звернення: 28.01.2025).
13. Sadeghpour S., V্লাidzik N. Click fraud in digital advertising: A review. *Computers*, 10(12). 2021. URL: <https://www.mdpi.com/2073-431X/10/12/164> (дата звернення: 28.01.2025).
14. Pantis J., Patsakis C. Anatomy of deceit: Technical and human factors in phishing campaigns. *Computers & Security*. 2024. URL: <https://www.sciencedirect.com/science/article/pii/S0167404824000816> (дата звернення: 28.01.2025).
15. Ahmad S., Zaman M., Al-Shamayleh A. S. AI models for phishing detection: A detailed review. *IEEE Open Journal*. 2024. URL: <https://ieeexplore.ieee.org/abstract/document/10681500/> (дата звернення: 28.01.2025).

16. Goenka R., Chawla M., Tiwari N. A comprehensive overview of phishing: Environments, targets, attack methods, and defenses. *Springer*. 2024. URL: <https://link.springer.com/article/10.1007/s10207-023-00768-x> (дата звернення: 28.01.2025).
17. Delgado M., Pereira S. Data-driven human and bot recognition from web activity logs based on hybrid learning techniques. *Digital Communications and Networks*. 2024. URL: <https://www.sciencedirect.com/science/article/pii/S2352864823000330> (дата звернення: 28.01.2025).
18. Putra M.A.R., Ahmad T., Hostiadi D. Behavior pattern analysis of botnet attacks in computer networks. *International Journal of Intelligent Systems*. 2022. URL: <https://inass.org/wp-content/uploads/2022/05/2022083148-2.pdf> (дата звернення 28.01.2025).
19. Suchacka G., Cabri A., Rovetta S., Masulli F. Effective real-time web bot detection. *Knowledge-Based Systems*. 2021. URL: <https://www.sciencedirect.com/science/article/pii/S0950705121003373> (дата звернення: 28.01.2025).

REFERENCES:

1. Zhou, Y., Cui, X., Qu, W., & Ge, Y. (2022). The effect of automation trust tendency, system reliability and feedback on users' phishing detection. *Applied Ergonomics*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S0003687022000771>
2. Oest, A., Safaei, Y., Doupé, A., & Ahn, G. J. (2019). Phishfarm: A scalable framework for measuring the effectiveness of evasion techniques against browser phishing blacklists. *IEEE Symposium on Security and Privacy*. January 28, 2025, Retrieved from: <https://ieeexplore.ieee.org/document/8835369>
3. IBM. (2024). Cost of a data breach report 2024. January 28, 2025, Retrieved from: <https://keepnetlabs.com/blog/171-cyber-security-statistics-2024-s-updated-trends-and-data>
4. DataBreaches. (2024). Credential phishing attacks up over 700 percent. January 28, 2025, Retrieved from: <https://databreaches.net/2024/12/18/credential-phishing-attacks-up-over-700-percent>
5. Panda Security. (2024). Cybercrime report for the first five months of 2024. January 28, 2025, Retrieved from: <https://www.pandasecurity.com/en/mediacenter/fbi-internet-crime-complaint-center-reports-losses>
6. HIPAA Journal. (2024). FBI BEC warning: \$5.5 billion lost. January 28, 2025, Retrieved from: <https://www.hipaajournal.com/fbi-bec-warning-55-billion-lost>
7. Song, L., & Wang, M. (2021). Efficient defense strategy against spam and phishing email: An evolutionary game model. *Journal of Information Security and Applications*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S2214212621001617>
8. Chen, J., Mishler, S., Hu, B., Li, N., & Proctor, R. (2022). The impact of reliable methodologies for bot filtering on user interaction simulations. *International Journal of Human-Computer Studies*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S0306457322002989>
9. Oest, A., Safaei, Y., Dupe, A., & Ahn, G. J. (2019). Scalable methods for evaluating phishing attack evasion effectiveness. *IEEE Symposium*. January 28, 2025, Retrieved from: https://www.usenix.org/system/files/sec20fall_oest_prepub.pdf
10. Huang, D., & Chen, N. (2023). Machine learning for adaptive cybersecurity training. *Cybersecurity Applications Journal*. January 28, 2025, Retrieved from: https://www.researchgate.net/publication/380771681_Study_on_Empowering_Cyber_Security_by_Using_Adaptive_Machine_Learning_Methods
11. Wang, M., & Song, L. (2021). An evolutionary game model for spam and phishing protection. *Journal of Information Security and Applications*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S2214212621001617>
12. Chen, K., Li, J., Zhang, H., & Zhao, S. (2024). Utilizing large language models for cyber threat detection. *Frontiers in AI*. January 28, 2025, Retrieved from: <https://arxiv.org/pdf/2405.04760>
13. Sadeghpour, S., & V্লাidzik, N. (2021). Click fraud in digital advertising: A review. *Computers*, 10(12). January 28, 2025, Retrieved from: <https://www.mdpi.com/2073-431X/10/12/164>
14. Pantis, J., & Patsakis, C. (2024). Anatomy of deceit: Technical and human factors in phishing campaigns. *Computers & Security*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S0167404824000816>
15. Ahmad, S., Zaman, M., & Al-Shamayleh, A. S. (2024). AI models for phishing detection: A detailed review. *IEEE Open Journal*. January 28, 2025, Retrieved from: <https://ieeexplore.ieee.org/abstract/document/10681500/>
16. Goenka, R., Chawla, M., & Tiwari, N. (2024). A comprehensive overview of phishing: Environments, targets, attack methods, and defenses. *Springer*. January 28, 2025, Retrieved from: <https://link.springer.com/article/10.1007/s10207-023-00768-x>

17. Delgado, M., & Pereira, S. (2024). Data-driven human and bot recognition from web activity logs based on hybrid learning techniques. *Digital Communications and Networks*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S2352864823000330>
18. Putra, M.A. R., Ahmad, T., & Hostiadi, D. (2022). Behavior pattern analysis of botnet attacks in computer networks. *International Journal of Intelligent Systems*. January 28, 2025, Retrieved from: <https://inass.org/wp-content/uploads/2022/05/2022083148-2.pdf>
19. Suchacka, G., Cabri, A., Rovetta, S., & Masulli, F. (2021). Effective real-time web bot detection. *Knowledge-Based Systems*. January 28, 2025, Retrieved from: <https://www.sciencedirect.com/science/article/pii/S0950705121003373>