

УДК 004.85

DOI: 10.32626/2308-5916.2025-27.82-102

В. В. Палагін, д-р техн. наук, професор,**О. В. Івченко**, канд. техн. наук,**О. А. Палагіна**, канд. техн. наук,**В. В. Філіпов**, канд. техн. наук,**О. С. Гавриш**, канд. фіз.-мат. наук,**А. В. Байрак**,**Р. Л. Пташкін**

Черкаський державний технологічний університет, м. Черкаси

РОЗРОБКА МЕТОДУ АНТИСПУФІНГУ ЗОБРАЖЕНЬ В СИСТЕМАХ БІОМЕТРИЧНОЇ БЕЗПЕКИ З ВИКОРИСТАННЯМ ML

Розглянуто проблему виявлення презентаційних атак (спуфінгу) в системах біометричної автентифікації на основі розпізнавання облич. Зростання спуфінг-атак із використанням статичних зображень, відеозаписів та 3D-масок створює нагальну потребу у розробці стійких до підробок систем їх виявлення. Метою дослідження є розробка багаторівневого комбінованого методу антиспуфінгу, який інтегрує фізичні та поведінкові ознаки обличчя з використанням моделей машинного навчання (ML) для адаптивного прийняття рішень.

Запропонована система включає чотири функціональні модулі аналізу зображень: виявлення країв (*Edge Detection*), аналіз мікрорухів (*Motion Analysis*), виявлення моргання (*Blink Detection*) та ідентифікація посмішки (*Smile Detection*). Кожен з модулів формує бінарне часткове рішення, яке далі передається до інтеграційного блоку. На відміну від класичних підходів зі сталими ваговими коефіцієнтами, у запропонованій системі ці ваги обчислюються адаптивно за допомогою навченої ML-моделі. Це забезпечує динамічну зміну впливу кожного модуля залежно від умов середовища, якості відеопотоку та характеристик обличчя користувача.

Наукова новизна роботи полягає в розробці гнучкої моделі комбінованої біометричної системи безпеки при впровадженні механізму оптимізації зважених коефіцієнтів модулів антиспуфінгу на основі реалізації методів машинного навчання. Система демонструє здатність до самоналаштування, що дозволяє підвищити точність і знизити частоту помилкових спрацьовувань у складних умовах експлуатації. Запропонована модель дозволяє балансувати між фізичними та поведінковими ознаками, адаптуючи їхній внесок у фінальне рішення у режимі реального часу.

© В. В. Палагін, О. В. Івченко, О. А. Палагіна,
В. В. Філіпов, О. С. Гавриш, А. В. Байрак, Р. Л. Пташкін, 2025

Система реалізована у вигляді модульної архітектури та протестована на контрольованому наборі даних, який містить кілька типів атак. Результати експериментів демонструють високу точність виявлення підробок (до 100% у тестових умовах) та стійкість до варіативності вхідних даних. Описано алгоритм функціонування, математичні моделі модулів, принцип інтеграції та можливості масштабування. Запропоноване рішення має перспективи практичного впровадження в мобільні пристрої, системи контролю доступу та онлайн-ідентифікації.

Ключові слова: *технічний захист інформації, біометрична автентифікація, машинне навчання, виявлення живості, комп'ютерний зір, багаторівневий захист, антиспуфінг.*

Вступ. Системи біометричної автентифікації відіграють дедалі важливішу роль у захисті інформаційних ресурсів, критичної інфраструктури та персональних даних. Розпізнавання обличчя стало одним із найбільш поширених біометричних методів через свою зручність, швидкодію та інтеграцію в масові пристрої. Однак зростання поширеності таких систем супроводжується збільшенням кількості атак на них, зокрема спуфінг-атак, у яких зловмисник імітує обличчя користувача за допомогою фотографій, відео чи 3D-масок. Ці загрози створюють серйозний виклик для надійності біометричних систем, знижуючи довіру до них (рис. 1) [1].

*Оригінальне
зображення*

Різновиди спуфінгу зображення



Рис. 1. Приклади різновидів спуфінгу зображення: атака 2D-друку (Фото), атака на 2D-відтворення (Відео), атака за допомогою 3D маски (3D-маска)

Згідно з останніми дослідженнями, презентаційні атаки є найбільш поширеним типом загроз у біометричних системах, і на них припадає до 80% усіх виявлених спроб несанкціонованого доступу [2-5]. Зважаючи на це, дедалі більше уваги приділяється розробці систем антиспуфінгу (*Presentation Attack Detection, PAD*), здатних виявляти підроблені спроби автентифікації. Проте сучасні підходи мають суттєві обмеження: вони або занадто спеціалізовані, або погано масштабуються до нових типів атак.

Традиційні методи, як-от LBP (*Local Binary Patterns*), HOG (*Histogram of Oriented Gradients*) та SIFT (*Scale-Invariant Feature Transform*), активно застосовувалися у задачах антиспуфінгу. Вони

дозволяють виділяти текстурні чи геометричні особливості зображень, виявляючи аномалії, характерні для підроблених матеріалів. Наприклад, LBP може виявити регулярні візерунки друку, а HOG – неприродні градієнти освітлення на 3D-масці. Водночас SIFT дозволяє знаходити стійкі ознаки на обличчі незалежно від масштабу. Проте всі ці методи залежать від ручного вибору ознак і не мають адаптивності до складних, раніше не спостережуваних атак [6-10].

Основною проблемою традиційних підходів є їх слабка узагальненість та нездатність до динамічної адаптації. Крім того, вони не враховують часові та поведінкові характеристики обличчя, обмежуючись лише текстурним аналізом. Це створює прогалину у здатності систем виявляти складні презентаційні атаки, зокрема з використанням високоякісних відеозаписів або глибоко детальних 3D-масок.

У відповідь на ці виклики, в науковому середовищі активно розвиваються підходи, засновані на *машинному та глибинному навчанні*. Вони дозволяють автоматично вивчати дискримінативні ознаки на основі великих обсягів даних, забезпечуючи більшу стійкість до нових атак. Проте навіть найсучасніші глибинні моделі, зокрема CNN, мають обмеження: складність навчання, висока залежність від обсягу якісних тренувальних даних, а також зниження точності в реальних умовах через зміни освітлення, носіння аксесуарів тощо [11, 12].

Враховуючи викладене, усе більше дослідників звертаються до *комбінованих підходів*, які поєднують фізичні ознаки (краї, текстури) та поведінкові (моргання, посмішка, рух голови) для побудови стійкої системи виявлення підробок. Такі підходи дозволяють компенсувати слабкі сторони кожного окремого методу, досягаючи високої точності навіть у складних середовищах, що особливо актуально для реального використання.

У цьому дослідженні пропонується багаторівнева система антиспуфінгу, яка інтегрує чотири незалежні модулі перевірки зображення: *виявлення країв, аналіз руху, виявлення моргання та розпізнавання посмішки*. Така система використовує як пасивні, так і активні індикатори живості, що дозволяє не лише виявляти широке коло атак, але й зберігати зручність взаємодії з користувачем.

Запропонований підхід демонструє високу ефективність у виявленні спроб підміни в різних умовах – від яскраво освітлених сцен до темних приміщень, з користувачами в окулярах чи без. Модульна архітектура дозволяє адаптувати систему під специфіку пристрою або сценарію, а результати експериментального тестування свідчать про 100% точність виявлення атак в контрольованих тестах. Робота має не лише прикладне, але й методологічне значення – вона формалізує комбінований підхід до виявлення спуфінгу та окреслює перспективи його подальшого розвитку.

Метою роботи є підвищення точності виявлення презентаційних атак у системах біометричної автентифікації шляхом розробки багаторівневої комбінованої моделі антиспуфінгу, яка інтегрує фізичні та поведінкові ознаки обличчя, забезпечуючи стійку і точну ідентифікацію справжніх користувачів навіть у складних умовах експлуатації.

1. Побудова структури та методів комбінованої системи виявлення спуфінгу. Запропонована система виявлення біометричного спуфінгу зображень реалізована на основі багаторівневої архітектури, яка поєднує кілька незалежних модулів верифікації для забезпечення стійкості до різних типів презентаційних атак. Кожен модуль виконує аналіз певного типу ознак – фізичних або поведінкових – і формує власне рішення щодо ймовірності спуфінгу. Інтегрована модель об'єднує ці локальні оцінки в узагальнене рішення, що базується на адаптивному пороговому підході (рис. 2).

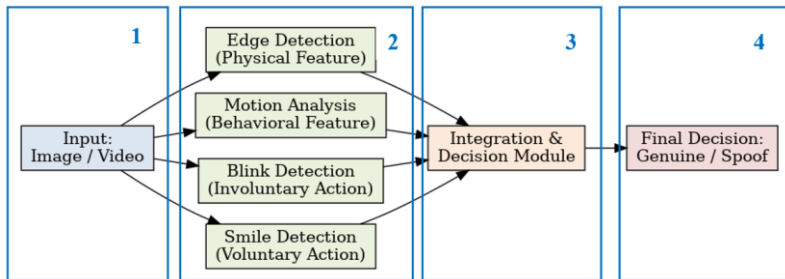


Рис. 2. Структура комбінованої системи виявлення біометричного спуфінгу

Загальна структура системи включає наступні компоненти:

1. **Вхідні дані** (*Input Image/Video*): відеопотік або послідовність зображень з фронтальної камери.
2. **Модулі виявлення:**
 - виявлення країв (*Edge Detection*);
 - аналіз руху (*Motion Analysis*);
 - виявлення моргання (*Blink Detection*);
 - виявлення посмішки (*Smile Detection*).
3. **Інтеграційний блок прийняття рішень** (*Integration & Decision Module*).
4. **Фінальне рішення** (*Final Decision*): *Справжнє (Genuine) / Підрване (Spoof)* зображення.

Запропонований комбінований підхід характеризується наступними особливостями:

- багаторівневий аналіз – методи виявлення країв, моргання, посмішки та аналізу руху забезпечують комплексний підхід, який

враховує як фізичні, так і поведінкові характеристики, що дозволяє ефективно виявляти як статичні атаки (наприклад, друківані фотографії), так і динамічні (відео, 3D-маски);

- адаптивність до різних сценаріїв – методи, такі як аналіз руху та виявлення посмішки, забезпечують активну участь користувача, що додає додатковий рівень захисту.

Розглянемо особливості побудови та застосування запропонованих модулів.

Алгоритм Канні (Canny Edge Detection) є одним із найпопулярніших та найефективніших методів виявлення країв у зображеннях, мінімізує шум та водночас забезпечує чітке і точне виявлення контурів. Алгоритм включає кілька послідовних етапів: згладжування зображення за допомогою гаусового фільтра, обчислення градієнтів яскравості, придушення немаксимумів та двопорогову гістерезис-фільтрацію. Такий підхід дозволяє точно визначити області з різким перепадом інтенсивності, які часто вказують на межі об'єктів або артефакти.

У контексті задачі антиспуфінгу, алгоритм Канні використовується для виявлення неприродних або надмірно різких країв, які характерні для друківаних зображень, екранів мобільних пристроїв або інших підобрених медіа. Наприклад, роздруковане фото (рис. 3-а) може мати надлишкову кількість високочастотних компонентів або нерівномірну текстуру країв, що виявляються через аналіз щільності виявлених країв. Такі візуальні особливості суттєво відрізняються від природного розмиття або плавного переходу градієнтів у справжніх обличчях (рис. 3-б).



Рис. 3. Візуалізація зображення градієнтної амплітудної характеристики алгоритмом Канні

Наступні вирази описують основні обчислення, які лежать в основі виявлення країв за допомогою алгоритму Канні. Для цього обчислюється *градієнт інтенсивності* $G(x, y)$, який визначається як евклідова норма похідних інтенсивності пікселів зображення по горизонталі $\partial I / \partial x$ та вертикалі $\partial I / \partial y$. Це дозволяє кількісно оцінити, наскільки різко зміню-

ється яскравість у кожній точці зображення, що і вказує на наявність краю. Високі значення градієнта вказує на різкий перехід в інтенсивності, характерний для меж об'єктів або текстурних аномалій:

$$G(x, y) = \sqrt{(\partial I / \partial x)^2 + (\partial I / \partial y)^2}. \quad (1)$$

Середня щільність країв D в області S – це середнє значення градієнтів по всій області аналізу і визначається як:

$$D = \sum_{(x, y) \in S} G(x, y) / |S|. \quad (2)$$

Це узагальнення дозволяє оцінити, наскільки загалом «різким» є зображення. Зіставлення цієї щільності з пороговим значенням T_{edges} дозволяє системі зробити висновок про ймовірність підробки (спуфінгу):

$$Edges_{flag} = \begin{cases} 1, D > T_{edges}, \\ 0, \text{інакше}. \end{cases} \quad (3)$$

Якщо D перевищує поріг T_{edges} , активується прапорець $Edges_{flag}$, що свідчить про потенційну аномалію. Такий підхід є ефективним для виявлення штучних артефактів, характерних для друкованих чи екранних зображень.

Однією з ключових переваг алгоритму Канні є його стійкість до шуму та здатність виділяти саме ті краї, які мають інформаційне значення. Водночас він дозволяє параметризувати рівень чутливості за допомогою низького та високого порогів, що робить його гнучким до умов освітлення, якості відео та характеристик камери. У межах багаторівневої системи виявлення спуфінгу модуль на основі Санпу виступає як швидкий та ефективний інструмент первинного аналізу фізичних характеристик, забезпечуючи базовий фільтр для двовимірних атак.

Метод **аналізу руху (Motion Analysis)** в задачах антиспуфінгу спрямований на виявлення мікродинаміки обличчя користувача в часовому вимірі та полягає у вимірюванні змін положення ключових орієнтирів обличчя між послідовними кадрами відеопотоку (рис. 4).

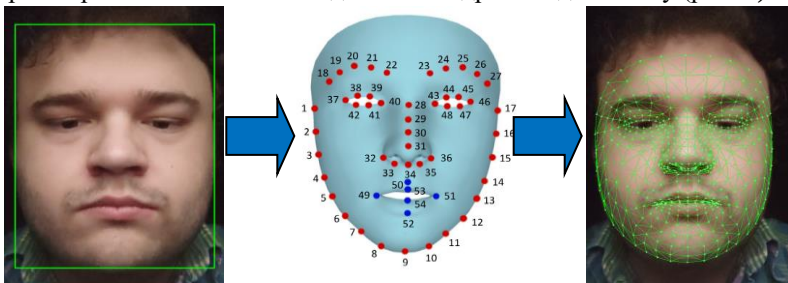


Рис. 4. Виявлення обличчя та виділення його орієнтирів

Нехай $P_i(t)$ – це координати i -ї контрольної точки обличчя в момент часу t . Тоді рух кожної точки між кадрами t та $t + 1$ можна оцінити за допомогою евклідової метрики:

$$\Delta_i(t) = \|P_i(t+1) - P_i(t)\|_2. \quad (4)$$

Ця величина показує, наскільки сильно зсунулась точка між кадрами, що є ознакою руху – мимовільного або навмисного. Щоб оцінити загальний рух обличчя в кадрі, обчислюється *середнє зміщення* всіх ключових точок у поточному часовому інтервалі:

$$M(t) = \frac{1}{k} \sum_{i=1}^k \Delta_i(t), \quad (5)$$

де k – кількість обраних орієнтирів обличчя. Це дає інтегральну оцінку рухливості об'єкта в кадрі. У реальних умовах обличчя живої людини демонструє природну, хоч і мінімальну, рухливість через мікропомітку, тремтіння, дихання або мікропереміщення голови. Натомість у випадках спуфінгу (наприклад, статичного зображення чи екранного відтворення) значення $M(t)$ буде близьке до нуля або мати неприродну регулярність.

Для прийняття рішення використовується порогова модель. Якщо середнє значення зміщення $M(t)$ перевищує заздалегідь визначене або адаптивне порогове значення θ_M , то це свідчить про присутність природного руху. Цей підхід є достатньо ефективним для виявлення 2D-атак, де зображення чи відео не мають глибини або руху. У складніших випадках, наприклад, з 3D-масками, метод може виявити аномальну синхронність рухів або їхню механічну регулярність. Метод є обчислювально простим, швидким у реалізації та добре інтегрується з іншими рівнями захисту в комбінованій системі.

Метод **виявлення моргання очей (Blink Detection)** базується на аналізі змін у геометрії ока протягом часу. Ключовим параметром є *Eye Aspect Ratio (EAR)* – безрозмірна метрика, яка оцінює ступінь відкритості ока. *EAR* обчислюється за координатами шести контрольних точок на контурі ока, які визначаються з використанням алгоритмів детекції орієнтирів обличчя:

$$EAR = \frac{\|p_2 - p_6\| + \|p_3 - p_5\|}{2\|p_1 - p_4\|}, \quad (6)$$

де p_1 і p_4 – координати кутів ока (горизонталь), а p_2, p_3, p_5, p_6 – координати точок на верхній та нижній повіках (вертикаль) (рис. 5) [13].

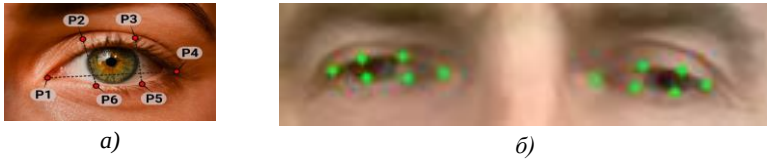


Рис. 5. Розмітка ока (а) та вилучення орієнтирів для розрахунку метрики Eye Aspect Ratio (EAR) (б)

Коли око відкрите – EAR має стабільне значення; під час моргання – вертикальні відстані різко зменшуються, і EAR падає.

Для виявлення моргання система відстежує часову послідовність значень EAR . Якщо в певному інтервалі часу значення EAR падає нижче за поріг θ_{Blink} , це інтерпретується як моргання:

$$BlinkDecision = \begin{cases} 1, & \exists t : EAR < \theta_{Blink}, \\ 0, & \text{інакше.} \end{cases} \quad (7)$$

Поріг θ_{Blink} може бути як фіксованим (наприклад, 0.2-0.3), так і адаптивним – залежно від анатомії користувача чи умов зйомки. Оскільки моргання є мимовільною поведінковою ознакою, його важко відтворити на фотографіях або попередньо записаному відео, що робить цей метод ефективним інструментом для виявлення спуфінгу. Він пасивний, не вимагає від користувача спеціальної дії, і добре поєднується з іншими методами в комбінованій системі.

Метрика EAR – є важливим компонентом виявлення живості, оскільки вони представляють мимовільну поведінку людини, яка слідує за конкретними зразками. В середньому людина моргає 15-20 разів на хвилину, причому кожне моргання триває від 100 до 400 мілісекунд. Для виявлення моргання співвідношення сторін розраховується динамічно шляхом вимірювання вертикальної відстані між ключовими точками на верхній і нижній повіках відносно горизонтальної відстані між кутами очей. Коли повіки змикаються під час моргання, вертикальний вимір наближається до нуля, що спричиняє значну зміну співвідношення сторін. Це падіння є характерною ознакою моргання.

Для аналізу EAR окремих осіб система використовує динамічний поріг зміни співвідношення сторін. Замість того, щоб покладатися на фіксоване значення, поріг налаштовується на основі пропорцій базової лінії користувача, які встановлюються під час сеансу аутентифікації. Цей підхід гарантує, що система може точно виявляти моргання незалежно від відмінностей у формі очей, розмірі або позиціонуванні користувача. Це також запобігає помилковим спрацьовуванням у випадках, коли око частково закривається через освітлення, втому або інші умови.

Виявлення моргання найкраще працює в режимі реального часу, при цьому система постійно відстежує рух очей користувача під час

процесу аутентифікації. Система відстежує тимчасову прогресію змін співвідношення сторін, щоб визначити чіткий малюнок моргання – раптове і коротке падіння співвідношення з подальшим поверненням до базової лінії. Ця тимчасова консистенція гарантує, що система не плутає моргання з іншими швидкими рухами очей або частковими закриттями.

Хоча виявлення моргання не захищене від крайових випадків – таких як підміни атак з використанням передових анімацій або високоякісних відео – воно залишається цінним компонентом багатоспіввідношення системи. Зосереджуючись на природній, мимовільній дії, яку важко повторити штучно, вона додає унікальний поведінковий вимір до загальної системи безпеки. Це робить його особливо ефективним в парі з іншими методами, забезпечуючи ефективний захист від спроб спуфінгу.

Метод **виявлення посмішки (Smile Detection)** використовує активну взаємодію з користувачем, вимагаючи виконати певну дію – посміхнутись. Така дія змінює геометричну структуру нижньої частини обличчя, зокрема – ротової області. Основна ідея методу полягає у фіксації пропорційних змін на обличчі через *Smile Ratio (SR)* – метрику, що оцінює співвідношення ширини рота до його висоти (рис. 6):

$$SR = \frac{\|p_{left} - p_{right}\|}{\|p_{top} - p_{bottom}\|}, \quad (8)$$

де p_{left} і p_{right} – крайні позиції рота відповідно, а p_{top} і p_{bottom} – найвищі та найнижчі точки ротового отвору відповідно. У нормальному (нейтральному) стані *SR* має помірне значення; при посмішці кути рота піднімаються і розсуваються в сторони, а висота відкритого рота збільшується, що призводить до значного зростання *SR*.

Система реєструє початкове значення *SR*, потім аналізує зміну цього значення у часовій послідовності. Якщо приріст *SR* перевищує порогове значення θ_{Smile} , фіксується позитивне рішення:

$$SmileDecision = \begin{cases} 1, & \text{якщо } \Delta SR > \theta_{Smile}, \\ 0, & \text{інакше.} \end{cases} \quad (9)$$

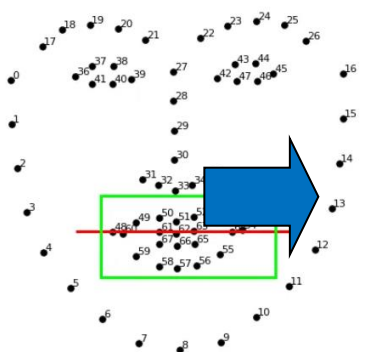
Така активна перевірка є ефективною для запобігання атакам з використанням статичних зображень чи відео, оскільки виконання посмішки вимагає *свідомої участі* користувача, чого неможливо досягти в підроблених презентаціях. Крім того, використання пропорційних вимірів (а не абсолютних піксельних значень) робить метод стійким до варіацій відстані до камери та масштабу обличчя.

Система відстежує співвідношення між певними точками обличчя на роті і навколо нього, розраховуючи співвідношення між горизонтальною шириною посмішки і вертикальною відстанню висоти рота. Посмішка виробляє передбачуване і виразне збільшення цього співвідношення, коли кути рота піднімаються вгору. Оскільки дві

особи не ідентичні, система використовує динамічний поріг, калібрований кожному користувачеві під час сеансу, адаптуючись до окремих структур обличчя та забезпечуючи надійне виявлення незалежно від варіацій інтенсивності посмішки.



а)



б)

Рис. 6. Розмітка посмішки (а) та вилучення орієнтирів (б) для розрахунку метрики Smile Ratio (SR)

Для посилення своєї надійності система вимагає від користувача виконання декількох посмішок під час процесу аутентифікації, при цьому кожен екземпляр відстежується і перевіряється. Це послідовне відстеження служить захистом від спроб підміни, які можуть покладатися на статичні фотографії, попередньо записані відео або анімації, що імітують одну посмішку. Застосовуючи повторну участь користувачів, виявлення посмішки мінімізує ризик шахрайського доступу, зберігаючи процес швидким та інтуїтивно зрозумілим для законних користувачів.

Виявлення посмішки особливо ефективно в реальних умовах, які можуть кинути виклик іншим методам. Він залишається міцним, коли користувачі носять окуляри або маски, які затінюють частини обличчя, оскільки він фокусується виключно на пропорційних рухах в області рота. Його здатність виявляти динамічні зміни обличчя гарантує, що він може адаптуватися до різних середовищ і умов освітлення без значної втрати точності.

Хоча тільки виявлення посмішки не може вирішити всі можливі сценарії спуфінгу, такі як розширені 3D-маски з імітацією міміки, воно служить цінним компонентом багатопланової системи. Вимагаючи навмисної, керованої користувачем дії, вона доповнює пасивні методи, такі як виявлення моргання та аналіз руху, зміцнюючи загальну структуру безпеки, зберігаючи простоту використання для справжніх користувачів.

Інтеграційний блок прийняття рішень. Третій інтеграційний блок виконує злиття результатів чотирьох окремих модулів верифікації з метою формування єдиної узагальненої оцінки автентичності зображення. Кожен модуль – *виявлення країв*, *аналіз руху*, *виявлення моргання та посмішки* – формує бінарний результат $d \in [0,1]$, який відображає наявність або відсутність ознак живої людини. Ці часткові рішення об'єднуються за допомогою їх зваженої лінійної комбінації, де кожному модулю надається вага w_i , що відображає його значимість у контексті конкретного сценарію:

$$S = w_E d_E + w_M d_M + w_B d_B + w_S d_S, \quad (10)$$

де S – функція агрегованої втрати, а відповідні складові виразу відображають ваги (w_i) та результат обробки чотирьох попередніх модулів (d_i) виявлення живості відповідно.

Значення ваг w_i можуть бути встановлені емпірично або оптимізовані автоматично методами машинного навчання (за допомогою логістичної регресії). Такий гнучкий підхід дозволяє адаптувати систему до різних умов використання, враховуючи, наприклад, слабке освітлення або наявність окулярів, які можуть знижувати ефективність окремих методів.

Фінальне рішення: Genuine/Spoof. Після обчислення інтегрального показника S , система порівнює його з заздалегідь визначеним або адаптивним порогом θ_{final} . Якщо обчислений показник перевищує поріг, користувача ідентифікують як справжнього (*Genuine*), інакше – як підозрілого (*Spoof*):

$$FinalDecision = \begin{cases} \text{Справжнє зображення (Genuine)}, & \text{якщо } S > \theta_{Final}, \\ \text{Підrobлене зображення (Spoof)}, & \text{інакше.} \end{cases} \quad (11)$$

Оптимізація ваг w_i у функції агрегованої втрати спрямована на мінімізацію False Acceptance Rate (FAR) та False Rejection Rate (FRR) при фіксованому порозі θ_{final} .

Таким чином, фінальне рішення відображає агреговану оцінку системи щодо автентичності взаємодії, забезпечуючи надійність виявлення спуфінгу шляхом використання як фізичних, так і поведінкових ознак. Такий підхід забезпечує високий рівень точності при мінімальному навантаженні на користувача.

2. Алгоритм функціонування комбінованої системи антиспуфінгу з адаптивним ваговим об'єднанням ознак. Запропонований комбінований підхід до виявлення спуфінгу реалізується через послідовний алгоритм, який включає як *пасивні*, так і *активні* перевірки живості. На першому етапі користувач проходить стандартну

процедуру входу в систему, після чого ініціюється процес верифікації. Система формує випадкове завдання для користувача (*моргнути, посміхнутись, повернути голову*), що дозволяє зменшити ризики підробки, особливо у випадках заздалегідь записаного відео.

Чотири окремі модулі – *виявлення країв, аналіз руху, виявлення моргання та посмішки* – працюють паралельно, кожен з яких виконує оцінку відповідної групи ознак. *Edge Detection* аналізує текстурні властивості, характерні для фізичних артефактів, тоді як *Motion Analysis* фіксує природні рухи обличчя. *Blink Detection* виявляє мимовільну активність очей, а *Smile Detection* вимагає короткої активної взаємодії з користувачем. Усі модулі формують проміжне рішення – значення 0 або 1, що вказує на наявність (*Genuine*) або відсутність (*Spoof*) живої поведінки (рис. 7).

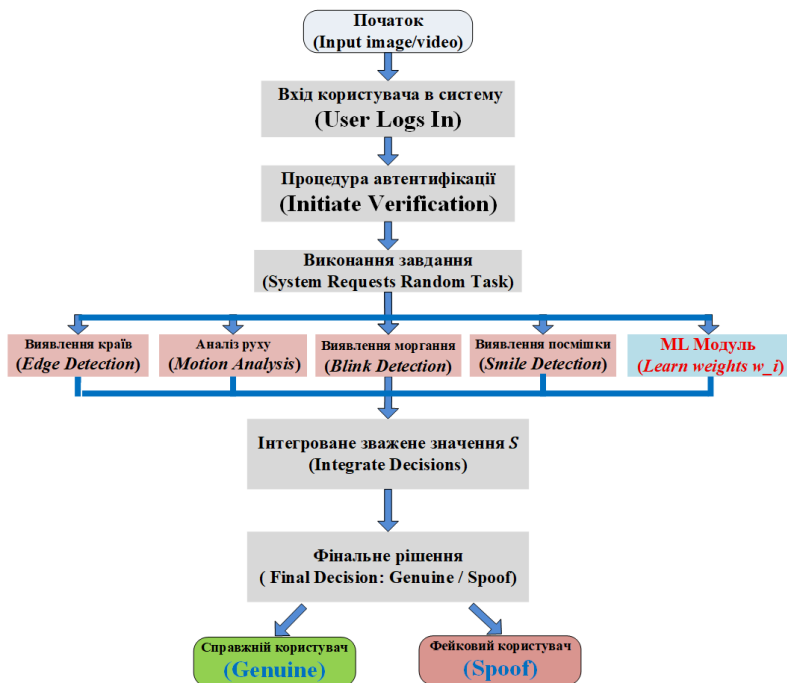


Рис. 7. Алгоритм функціонування комбінованої системи антиспуфінгу

Ці локальні рішення надходять до **інтеграційного блоку (Integrate Decisions)**, де реалізовано механізм зваженого об'єднання. Кожен модуль має свою вагу w_i , яка визначає степінь його впливу на фінальне рішення (**Final Decisions**). На відміну від фіксованих значень, що застосовувалися в базових системах, у цьому підході ваги

формується динамічно на основі попереднього навчання. Такий принцип дозволяє компенсувати похибки окремих модулів в умовах змінного середовища або якості даних.

Ключовим компонентом є **модуль машинного навчання (ML module)**, який формує оптимальні значення ваг w_i на основі попередньо зібраного набору даних. Наприклад, при слабкому освітленні зменшується вплив *Blink-модуля*, а більшу вагу отримують фізичні ознаки. Для цього був використаний алгоритм логістичної регресії, який навчався на мітках «справжній/підробка» з реальних тестових сценаріїв.

Після зваженого об'єднання модуль обчислює інтегрований показник S , який порівнюється з пороговим значенням θ_{final} . Якщо сумарна оцінка перевищує поріг – система ідентифікує користувача як справжнього. Якщо ж значення є нижчим – користувач отримує відмову у доступі. Такий механізм є стійким до маніпуляцій, зокрема, до часткових атак, коли підробка намагається відтворити лише окремі ознаки живості, або модифікувати зображення при застосуванні сучасних методів фільтрації зображень [14].

Таким чином, побудований алгоритм об'єднує надійність класичних детекторів із гнучкістю адаптивного машинного навчання. Комбінована система забезпечує не тільки високий рівень точності, але й адаптивність до нових типів атак, зміни середовища та індивідуальних особливостей користувачів. Впровадження такого підходу дозволяє перейти від жорстких статичних методів до інтелектуальної системи виявлення спуфінгу нового покоління.

Для оцінки якості роботи запропонованої комбінованої системи антиспуфінгу використовуються такі метрики біометричних систем, що рекомендовані стандартами ISO/IEC 30107 для *Presentation Attack Detection (PAD)* – систем [15]:

- коефіцієнт помилок класифікації презентації атаки (APCER – *Attack Presentation Classification Error Rate*)

$$APCER = FN / (TP + FN); \quad (12)$$

- коефіцієнт помилок класифікації нормального (добросовісного) представлення BPCER (*Bona Fide Presentation Classification Error Rate*)

$$BPCER = FP / (FP + TN); \quad (13)$$

- середній коефіцієнт помилок класифікації (ACER)

$$ACER = (APCER + NPCER) / 2; \quad (14)$$

- коефіцієнт помилкових відхилень (FRR) – частка справжніх користувачів, відхилених системою

$$FRR = FN / (TP + FN); \quad (15)$$

- коефіцієнт хибного прийняття (FAR) – частка підробок, помилково прийнятих як справжні

$$FAR = FP / (FP + TN), \quad (16)$$

де TP – атаки розпізнаються як атаки, TN – справжні зразки розпізнаються як справжні, FP – реальні зразки розпізнаються як атаки, FN – атаки визнані справжніми зразками.

3. Результати експериментального дослідження. На основі запропонованих методів комбінованої системи виявлення спуфінгу зображень реалізована послідовна PAD система, де спочатку приймається рішення про справжність зображення, і лише якщо визначено, що зразки походять від живої людини (*Accept*), тоді вони обробляються системою розпізнавання обличчя – FRS (рис. 8). На відміну від паралельного підходу, у послідовній схемі середній час спроби доступу буде довшим через послідовні затримки PAD та модулів розпізнавання облич. Однак цей підхід дозволяє уникнути додаткової роботи системи розпізнавання облич (FRS) у разі атаки PAD, оскільки її слід виявити на ранній стадії.

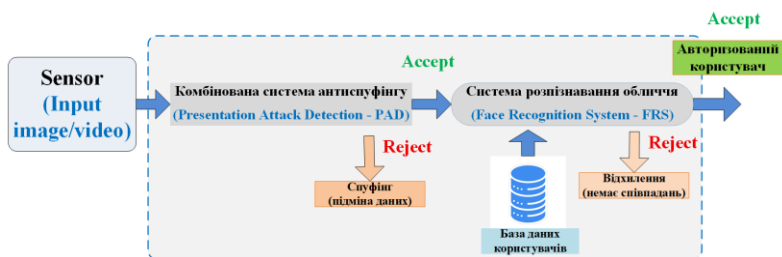


Рис. 8. Схема послідовного об'єднання PAD та системи розпізнавання обличчя (FRS)

Розроблений експериментальний підхід до аналізу текстур для вивчення додаткових можливостей проти спуфінгу. Цей метод розрізняє справжні обличчя та цифрові/друковані репродукції (рис. 9), аналізуючи три ключові характеристики текстури:

- *локальна дисперсія* – допомагає покращити здатність системи виявляти справжні біометричні ознаки, фокусуючись на текстурних і крайових деталях, які вказують на живого суб'єкта;
- *градієнтний аналіз* – виявлення тонких відмінностей в текстурних переходах, які могли б розрізнити реальну шкіру і друковані/відображувані зображення;
- *аналіз частотної області* – швидкий аналіз перетворення Фур'є був реалізований для виявлення періодичних моделей, які можуть бути присутніми в цифрових дисплеях або друкованих матеріалах.

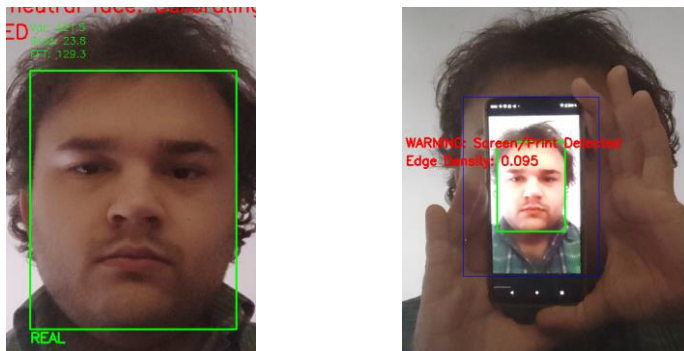


Рис. 9. Приклад роботи текстурного аналізу з оптимізованими порогоми

Необхідно відмітити, що даний підхід виявився складним для ефективного впровадження через кілька факторів:

- висока чутливість до умов освітлення;
- несумісні результати з різною якістю відео;
- труднощі у встановленні надійних порогів у різних середовищах тестування;
- відсутність послідовної фонові ізоляції в тестовому наборі даних.

В роботі проведені експерименти із впровадженням наступних методів виявлення живості, зокрема «*Motion Analysis*», «*Blink Detection*». Система виявлення моргання реалізує надійний метод ідентифікації природнього моргання очей за допомогою метрики співвідношення сторін ока в поєднанні з часовим аналізом. Цей підхід фокусується на вимірюванні зв'язку між вертикальними та горизонтальними дистанціями між мітками отворів очей для виявлення справжніх візерунків моргання (рис. 10-а).

Для врахування індивідуальних варіацій форми очей та кутів камери система реалізує адаптивний базовий підхід:

- підтримує буфер значень EAR під час станів з відкритими очима;
- обчислює поточну базову лінію з попередніх вимірювань;
- виявляє моргання на основі відносних змін з цієї базової лінії.

Впроваджено метод «*Smile Detection*», який адаптується до індивідуальних рис обличчя через фазу калібрування, встановлюючи персоналізовану базову лінію для більш точного виявлення. Запропонована функція вимірювання зв'язку між шириною рота і висотою в декількох точках, фіксуючи характерні зміни форми рота під час посмішки. Використання трьох вимірювань висоти (ліворуч, праворуч і посередині) забезпечується більш надійне виявлення асиметричних виразів (рис. 10-б).

Ключовою особливістю системи є її калібрувальна фаза, яка встановлює персоналізовану базову лінію для кожного користувача та вдається значно зменшити помилкові спрацювання, які можуть виникнути внаслідок незначних рухів рота. Під час калібрування система:

- збирає значення коефіцієнта посмішки, поки користувач зберігає нейтральний вираз;
- обчислює середню базову лінію;
- встановлює поріг.

Ця реалізація (рис. 10-б) дає значні переваги завдяки адаптивному підходу. Система успішно адаптує індивідуальні риси обличчя та вирази спокою. Такий підхід забезпечує стійкість до варіацій освітлення та зберігає ефективність під різними кутами обличчя.

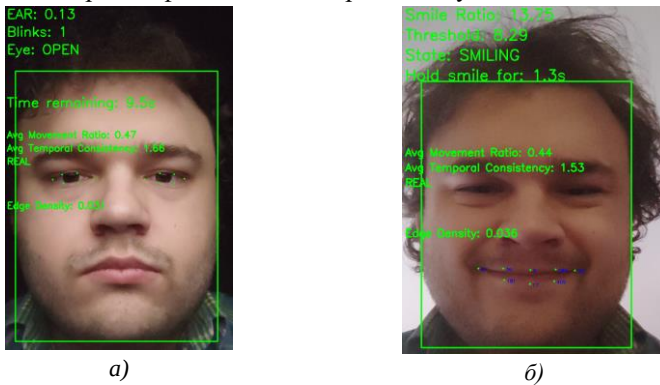


Рис. 10. Приклад роботи методів Blink Detection (а) та Smile Detection (б) для виявлення живості

Таким чином, система виявлення посмішки демонструє надійну продуктивність у диференціації автентичних посмішок від інших рухів рота або статичних зображень, тим самим підвищуючи загальну можливість виявлення живості системи безпеки. Ця надійність робить його ефективним компонентом у більш широкій системі боротьби зі спуфінгом.

4. Обговорення результатів. Розробка та впровадження комбінованої системи антиспуфінгу зображень успішно інтегрували кілька рівнів захисту від презентаційних атак. Система поєднує в собі розпізнавання обличчя за допомогою YOLOv8-Face з розпізнаванням обличчя MediaPipe, щоб створити основу для перевірки живості.

Процес тестування використовував повний набір даних з 32 відеозаписів, призначених для оцінки продуктивності системи в різних реальних сценаріях. Відео були записані за допомогою камери мобільного телефону і включали як успішні спроби аутентифікації, так і потенційні сценарії спуфінгу.

Для справжніх спроб автентифікації записи робилися в різних умовах навколишнього середовища та включали варіації освітлення (*чітке, нормальне, тьмяне та темне*), відстань від камери (*нормальне та віддалене положення*) та наявність або відсутність окулярів. Кожен запис фіксував користувача за конкретними підказками автентифікації: для системи виявлення моргання користувачам було доручено двічі моргати протягом десяти секунд, тоді як для системи виявлення посмішки користувачам було запропоновано підтримувати посмішку протягом двох секунд.

Тестовий набір даних також включав ряд спроб підміни. Вони включали презентаційні атаки з використанням як статичних зображень, так і відеозаписів, що відображаються на різних відстанях від камери. Додаткові сценарії відмов були записані для перевірки здатності системи відхиляти недійсні спроби автентифікації, такі як одиночні моргання, постійне закриття очей, підморгування та надмірний рух голови.

Параметри тестування були стандартизовані на всіх відео, з ключовими порогоми 1,9 для виявлення послідовності руху та 0,08 для аналізу щільності краю. Під час тестування система збирала бінарні результати *pass/fail* для кожної спроби автентифікації.

Тестовий відео ряд, який досліджувався в роботі, характеризується широким спектром умов проведення та впливом зовнішніх факторів (рис. 11).



Рис. 11. Класифікація тестових сценаріїв, проведених при дослідженні комбінованого методу

Процес тестування оцінював ці реалізації за різними сценаріями, включаючи різні умови освітлення, відстані та наявність окулярів, а також кілька типів спроб спуфінгу. Це комплексне тестування виявило як сильні сторони, так і обмеження кожного підходу, забезпечуючи цінне розуміння їх практичної ефективності в реальних застосуваннях.

Результати тестування показали різні характеристики продуктивності для кожного методу виявлення живості. Система виявлення

посмішки продемонструвала чудову продуктивність у складних умовах, особливо коли користувачі носили окуляри, досягаючи 80% успіху для справжніх спроб аутентифікації. Система виявлення моргання, хоча і трохи більш чутлива до факторів навколишнього середовища, все ще підтримувала 70% успіху для справжніх спроб.

Примітно, що комбінація всіх наведених методів показали ідеальну 100% точність у виявленні та відхиленні недійсних спроб аутентифікації та підміни атак. Це включає успішну відмову від атак презентації, використовуючи як статичні зображення, так і відеоповтори, а також неприпустимі спроби взаємодії, такі як поодинокі моргання або підморгування.

Наявність окулярів зробило помітний вплив на методи. Виявлення посмішки залишалося ефективним, коли користувачі носили окуляри, зберігаючи свою здатність виявляти рухи рота і пропорційні зміни, пов'язані з посмішкою. І навпаки, точність виявлення моргання значно знизилася за допомогою окулярів, оскільки рамки частково затуляли повіки користувача та зменшили точність відстеження моргання.

Варіації відстані до датчика не вплинули негативно на продуктивність. Обидва методи продемонстрували послідовні результати в відео, де користувач був розташований ближче або далі від датчика. Це вказує на те, що система може розміщувати користувачів на різних відстанях без значної втрати точності, що є важливим для реального використання.

Основним обмеженням, визначеним за допомогою тестування, була деградація продуктивності в умовах низького освітлення, що впливає на методи виявлення.

Висновки. У роботі представлено комбінований метод виявлення спуфінгу в системах біометричної автентифікації, який базується на інтеграції фізичних і поведінкових ознак обличчя. Запропонована система включає чотири окремі модулі перевірки живості – виявлення країв, аналіз мікрорухів, детекцію моргання та посмішки – що забезпечують високий рівень достовірності верифікації користувача. Кожен з модулів виконує локальну оцінку, яка у подальшому враховується в інтегрованому рішенні про справжність або підробку зображення.

Однією з ключових відмінностей запропонованої системи є використання механізму зваженого об'єднання рішень, де вагові коефіцієнти w_i визначаються динамічно за допомогою навченої моделі машинного навчання. Навчання моделі здійснюється на анотованих даних з використанням функції втрат, орієнтованої на мінімізацію помилок класифікації в задачі детекції презентаційних атак. Такий підхід забезпечує адаптивність системи до змінних умов зйомки, індивідуальних особливостей обличчя, типів атак і підвищує здатність до узагальнення без необхідності повного перенавчання. Запропонований комбінований метод антиспуфінгу істотно підвищує точність і стійкість системи до різних типів атак – від простих 2D-фотографій до високоякісних відео чи 3D-масок.

Проведене експериментальне тестування продемонструвало високу ефективність розробленої системи, зокрема точність виявлення спуфінгу на рівні до 100% у контрольованих умовах. Модульна архітектура дає змогу легко адаптувати систему до різних платформ і типів застосування, у тому числі мобільних пристроїв і систем безпечного доступу. Гнучкість у виборі модулів і стратегій об'єднання дозволяє налаштовувати систему під конкретні вимоги й обмеження.

У подальших дослідженнях планується розширити експериментальну базу, інтегрувати нові типи біометричних ознак та удосконалити модель машинного навчання з урахуванням гетерогенності користувачів і умов експлуатації, використання технологій Deep Learning для оптимізації архітектури системи та її навчання в умовах обмежених даних.

Список використаних джерел:

1. Hernandez-Ortega J., Fierrez J., Morales A., Galbally J. Introduction to Face Presentation Attack Detection. *Handbook of Biometric Anti-Spoofing. Advances in Computer Vision and Pattern Recognition*. Springer, Cham. 2019. URL: https://doi.org/10.1007/978-3-319-92627-8_9.
2. Jain A. Ross, Prabhakar S. An introduction to biometric recognition. *IEEE Transactions on Circuits and Systems for Video Technology*. 2004. Vol. 14. № 1. P. 4-20.
3. Akhtar Z., Micheloni C. Foresti G. L. Biometric Liveness Detection: Challenges and Research Opportunities. *IEEE Security & Privacy*. 2015. Vol. 13. № 5. P. 63-72. DOI: 10.1109/MSP.2015.116.
4. Suchetha V., Anusha S., Deekshitha P., Deepika S., Naik P. Survey on face and fingerprint based person identification system. *Journal of Computer Science Engineering and Software Testing*. 2022. Vol. 8. P. 51-56.
5. Kumar S. Jain, Kumar M. Face and gait biometrics authentication system based on simplified deep neural networks. *International Journal of Information Technology*. 2022. Vol. 15. № 1. P. 1005-1014.
6. Sharma D., Selwal A. A survey on face presentation attack detection mechanisms: hitherto and future perspectives. *Multimedia Systems*. 2023. Vol. 29. № 3. P. 1527-1577.
7. Günay Yılmaz A., Turhal U., Nabiye V. Face presentation attack detection performances of facial regions with multi-block LBP features. *Multimed Tools Appl*. 2023. Vol. 82. P. 40039-40063. URL: <https://doi.org/10.1007/s11042-023-14453-7>.
8. Erdogmus N., Marcel S. Spoofing Face Recognition With 3D Masks. *IEEE Transactions on Information Forensics and Security*. 2014. Vol. 9. № 7. P. 1084-1097. DOI: 10.1109/TIFS.2014.2322255.
9. Faseela Abdullakutty, Eyad Elyan, Pamela Johnston. A review of state-of-the-art in Face Presentation Attack Detection: From early development to advanced deep learning and multi-modal fusion methods. *Information Fusion*. 2021. Vol. 75. P. 55-69. URL: <https://doi.org/10.1016/j.inffus.2021.04.015>.
10. Vijayan R. S. Image-Driven Face Presentation Attack Detection in Biometric Systems. *2025 Emerging Technologies for Intelligent Systems (ETIS)*, Trivandrum, India, 2025. P. 1-6. DOI: 10.1109/ETIS64005.2025.10961132.

11. Hoffman S., Sharma R., Ross A. Convolutional Neural Networks for Iris Presentation Attack Detection: Toward Cross-Dataset and Cross-Sensor Generalization. *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*, Salt Lake City, UT, 2018. P. 1701-17018. DOI: 10.1109/CVPRW.2018.00213.
12. Zhang Y., Zhao M., Yan L., Gao T., Chen J. CNN-Based Anomaly Detection for Face Presentation Attack Detection with Multi-Channel Images. *2020 IEEE International Conference on Visual Communications and Image Processing, VCIP*. 2020. P. 189-192.
13. Mahmood M. et al. A Novel Approach to Liveness Detection: Real-time Eye Blink, Head Movement, and Lip Movement Analysis. *2025 International Conference on Electrical, Computer and Communication Engineering (ECCE)*, Chittagong, Bangladesh, 2025. P. 1-6. DOI: 10.1109/ECCE64574.2025.11014007.
14. Palahin D., Palahina E., Palahin V. Development of Methods for Image Filtering in Noise and their Implementation for a Web Service. *2024 IEEE 42nd International Conference on Electronics and Nanotechnology (ELNANO)*, Kyiv Ukraine, 2024. Pp. 391-395. DOI: 10.1109/ELNANO63394.2024.10756816.
15. Presentation Attack Detection: Measuring Performance with APCER and BPCER. URL: <https://www.idrnd.ai/presentation-attack-detection>.

DEVELOPMENT OF AN IMAGE ANTISPOOFING METHOD IN BIOMETRIC SECURITY SYSTEMS USING MACHINE LEARNING

This article addresses the problem of detecting presentation attacks (spoofing) in facial biometric authentication systems. Given the growing prevalence of spoofing attacks using printed photos, video replays, and 3D masks, there is an urgent need to develop robust anti-spoofing mechanisms. The objective of this study is to develop a multi-level combined method for spoofing detection that integrates both physical and behavioral facial features using machine learning models for adaptive decision-making.

The proposed system includes four functional modules: edge detection, motion analysis, blink detection, and smile detection. Each module generates a binary decision regarding the presence of liveness indicators, which are subsequently passed to an integration block. Unlike traditional approaches with fixed weighting schemes, this system computes weights adaptively based on a trained machine learning model. This enables the dynamic adjustment of each module's influence depending on environmental conditions, video quality, and individual facial characteristics.

The scientific novelty of this work lies in the development of a flexible mechanism for optimizing weighted coefficients of anti-spoofing modules based on the outcome of the learning phase. The system demonstrates the ability to self-adjust, enhancing overall detection accuracy while reducing false acceptance/rejection rates in challenging scenarios. The proposed model dynamically balances the contributions of physical and behavioral features to the final decision in real time.

The system is implemented as a modular architecture and tested on a controlled dataset containing various spoofing scenarios. Experimental results demonstrate high detection accuracy (up to 100% in test settings) and

resilience to variability in input data. The article presents the operational algorithm, mathematical formulation of the modules, the principle of decision integration, and scalability potential. The proposed solution shows promise for practical implementation in mobile devices, access control systems, and online identity verification platforms.

Key words: *technical information protection; biometric authentication; machine learning; liveness detection; computer vision; multi-layered protection; anti-spoofing.*

Отримано: 9.06.2025

УДК 004.85

DOI: 10.32626/2308-5916.2025-27.102-121

В. В. Палагін, д-р техн. наук, професор,

О. В. Івченко, канд. техн. наук,

О. А. Палагіна, канд. техн. наук,

В. В. Філіпов, канд. техн. наук,

А. В. Байрак,

О. А. Проценко

Черкаський державний технологічний університет, м. Черкаси

МЕТОД МАШИННОГО НАВЧАННЯ ДЛЯ АНАЛІЗУ ШКІДЛИВОГО МЕРЕЖЕВОГО ТРАФІКУ НА ПРИКЛАДНОМУ РІВНІ (DHCP SPOOF)

DHCP (Dynamic Host Configuration Protocol) є критично важливим елементом мережевої інфраструктури, що забезпечує автоматичну видачу IP-адрес і параметрів конфігурації клієнтам. Проте через відсутність механізмів аутентифікації в базовому протоколі він є вразливим до атак типу DHCP spoofing, коли зловмисник імітує роботу легітимного сервера та видає клієнтам шкідливі налаштування. У цій роботі запропоновано підхід до виявлення таких атак із використанням методів машинного навчання, що дозволяє ідентифікувати як типові варіанти атак із фальшивими IP-адресами, так і складніші – із підміною MAC-адрес при справжньому IP сервера.

У межах дослідження розроблено інструмент для генерації DHCP-трафіку з різноманітними сценаріями поведінки, включно з нормальними сесіями, атаками з різними IP-адресами зловмисників та маскуванням під легітимного сервера. Запропоновано застосування методів машинного навчання (*Machine learning – ML*) для аналізу даних про трафіку, які отримані в режимі реального часу. На основі отриманого PCAP-файлу було автоматизовано процес побудови датасету зі стисненим, але інформативним набором ознак (кількість унікальних IP/MAC-адрес, перший MAC-відповідач, відповідність IP до MAC тощо). Побудована модель

© В. В. Палагін, О. В. Івченко, О. А. Палагіна,
В. В. Філіпов, А. В. Байрак, О. А. Проценко, 2025