

Міністерство освіти і науки України
Харківський національний університет радіоелектроніки

Кваліфікаційна наукова
праця на правах рукопису

ШУПИЛЮК МИКИТА ВОЛОДИМИРОВИЧ

УДК 004.93.1:004.032.26

ДИСЕРТАЦІЯ

Модель та методи інформаційної технології автоматизованого
графологічного аналізу рукописного тексту

Спеціальність: 126 – Інформаційні системи та технології

Галузь знань: 12 – Інформаційні технології

Подається на здобуття ступеня доктора філософії

Дисертація містить результати власних досліджень.

Використання ідей, результатів і текстів інших авторів мають посилання на
відповідне джерело _____ М. В. Шупилюк

Науковий керівник
Мартовицький Віталій Олександрович,
кандидат технічних наук, доцент
Харків – 2026

АНОТАЦІЯ

Шупилюк М.В. Модель та методи інформаційної технології автоматизованого графологічного аналізу рукописного тексту. – Кваліфікаційна наукова праця на правах рукопису.

Дисертація на здобуття ступеня доктора філософії за спеціальністю 126 Інформаційні системи та технології. – Харківський національний університет радіоелектроніки, Міністерство освіти і науки України, Харків, 2026.

Дисертаційну роботу присвячено розробці моделі та удосконаленню методів інформаційної технології графологічного аналізу зображень рукописного тексту. У роботі представлено підходи до автоматизованого графологічного аналізу рукописного тексту, що забезпечують комплексний і ефективний аналіз, який окрім ідентифікації автора ще дозволяє здійснювати психологічне профілювання. Запропоновані зміни включають як метод визначення особистості, що набув подальшого розвитку так і удосконалений метод ідентифікації та пошуку автора.

Актуальність обраної проблематики зумовлена важливою роллю графологічного аналізу рукописного тексту в задачах криміналістичної експертизи, забезпечення безпеки та дослідження історичних документів. Однією з головних передумов ефективного графологічного аналізу є здатність системи враховувати високу внутрішньоавторську мінливість почерку, яка виникає під впливом фізичних, психологічних та ситуативних факторів. Обмеження існуючих методів має суттєвий вплив не тільки на точність ідентифікації автора та визначення особистості, а й на можливість комплексного аналізу рукописного тексту. У зв'язку з цим постає питання підвищення ефективності автоматизованого графологічного аналізу шляхом

розробки нових моделей та методів. У традиційних підходах до аналізу почерку (структурних, текстурних, графемних та заснованих на глибокому навчанні) поточні обмеження, такі як висока залежність від попередньої обробки зображень, потреба у великих обсягах навчальних даних та значні обчислювальні витрати, не враховуються на достатньому рівні. Це, звичайно, призводить до зниження надійності та точності аналізу. Також потрібно зазначити, що більшість досліджень зосереджені на окремих задачах (ідентифікація автора або верифікація підпису), тоді як комплексний підхід, що поєднує ідентифікацію автора з визначенням психологічного профілю особистості, залишається недостатньо дослідженим. Це можна вирішити за рахунок використання інструментів сучасних архітектур глибокого навчання, які здатні ефективно виявляти глобальні та локальні закономірності у зображеннях рукописного тексту. Самі ж підходи, які поєднують зорові трансформери з методами виявлення ключових точок, відкривають нові можливості для створення більш точних та комплексних систем графологічного аналізу. Таким чином, дослідження ефективних методів автоматизованого графологічного аналізу рукописного тексту за допомогою сучасних архітектур глибокого навчання робить роботу актуальною.

Метою дисертаційної роботи є підвищення точності та автоматизація аналізу зображень рукописного тексту шляхом розробки моделей та методів автоматизованого графологічного аналізу рукописного тексту.

Об'єкт дослідження – процес обробки та графологічного аналізу рукописного тексту.

Предмет дослідження – моделі та методи графологічного аналізу зображень рукописного тексту, які направлені на ефективність та автоматизацію аналізу.

У першому розділі дисертації представлено аналіз сучасного стану досліджень у сфері графологічного аналізу рукописного тексту. Розглянуто основні процеси аналізу почерку та особливості графологічного аналізу. Здійснено огляд та аналіз сучасних методів ідентифікації автора та методів визначення особистості. Проаналізовано психометричні моделі особистості, методи нанесення цифрового водяного знака, а також технології безпечного зберігання даних. На основі проведеного аналізу сформульовано мету та задачі дослідження.

У другому розділі представлено метод визначення психологічних характеристик особистості за рукописним текстом. Обґрунтовано вибір психометричної моделі «Велика п'ятірка» та встановлено кореляцію її рис із відповідними графологічними особливостями почерку. Запропоновано функціональну модель визначення особистості, яка включає етапи попередньої обробки вхідних даних із застосуванням масштабно-інваріантного ознакового перетворення. Для безпосереднього визначення особистості використано архітектуру глибокого навчання Зоровий трансформер. Обґрунтовано вибір програмних засобів та обрано набір даних CVL для навчання. Представлено результати експериментального дослідження, які підтверджують ефективність удосконаленого методу з максимальною точністю класифікації до 99,48%.

У третьому розділі представлено удосконалений метод ідентифікації та пошуку автора за рукописним текстом. Обґрунтовано вибір масштабно-інваріантних центрально-окружних детекторів як методу виявлення ознак зображень. Запропоновано підхід до ідентифікації автора, який інтегрує центрально-окружні детектори та архітектуру глибокого навчання Зоровий трансформер. Побудовано функціональну модель ідентифікації автора та проведено її формалізацію шляхом побудови математичної моделі.

Представлено результати експериментального дослідження, які підтверджують збільшення загальної кількості вилучених ознак на 71%, скорочення середнього часу попередньої обробки на 40% та підвищення показників ідентифікації.

У четвертому розділі представлено модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту. Побудовано функціональну та структурну моделі інформаційної технології, здійснено математичну формалізацію задачі. Розроблено архітектуру інформаційної системи на базі технологій хмарних обчислень, яка охоплює повний цикл обробки даних від завантаження зображення до формування комплексного звіту. Інтегровано цифрові водяні знаки як метод захисту зображень та реалізовано механізми забезпечення цілісності даних на основі криптографічного хешування.

У висновках дисертаційної роботи узагальнено результати дослідження. Підсумовано значення отриманих наукових та практичних результатів. Показано, що запропоновані модель інформаційної технології та удосконалені методи визначення особистості та ідентифікації автора з використанням Зорового трансформера та масштабно-інваріантних детекторів забезпечують підвищення точності та автоматизацію графологічного аналізу рукописного тексту. Впровадження розроблених методів суттєво покращує загальну ефективність автоматизованого графологічного аналізу.

Методи дослідження – В дослідженні було використано широкий спектр методів. Так, при проведенні аналізу сучасного стану досліджень використовувались методи системного аналізу, процесний підхід та порівняльний аналіз існуючих рішень, що дозволило систематизувати підходи до ідентифікації автора, визначення особистості та забезпечення цілісності цифрових зображень, а також сформулювати мету та задачі дослідження. У

межах розробки методу визначення особистості за рукописним текстом, було використано методи глибокого навчання, зокрема архітектура Зорового трансформера, психометрична модель «Велика п'ятірка» та масштабно-інваріантне ознакове перетворення, методи визначення порогу для бінаризації зображень та градації сірого для попередньої обробки вхідних даних. При розробці удосконаленого методу ідентифікації та пошуку автора за рукописним текстом використовувались масштабно-інваріантні центрально-окружні детектори в інтеграції із архітектурою глибокого навчання Зоровий трансформер, а також методи математичного моделювання та функціонального моделювання, що дозволило збільшити кількість вилучених ознак та скоротити час попередньої обробки. При розробці моделі інформаційної технології автоматизованого графологічного аналізу рукописного тексту використовувались методи математичної формалізації, функціонального та структурного моделювання, а також технології хмарних обчислень, цифрові водяні знаки та криптографічне хешування для забезпечення цілісності та захисту даних. Оцінка експериментальних даних, отриманих у ході роботи, проводилася на основі порівняльного аналізу показників ефективності (Hard Top K, Soft Top K, mAP) та метрик точності класифікації. Вибір методів досліджень забезпечив достовірність отриманих результатів і висновків, що підтверджується збіжністю результатів експериментальних досліджень, отриманих при програмній реалізації інформаційної технології автоматизованого графологічного аналізу рукописного тексту, з теоретичними розрахунками, а також узгодженням із прогностичними оцінками ефективності запропонованих підходів.

Завдання дослідження:

- провести аналіз сучасного стану досліджень у сфері аналізу рукописного тексту;

- розробити модель інформаційної технології автоматизованого аналізу рукописного тексту;
- удосконалити метод ідентифікації та пошуку автора на основі зображень рукописного тексту;
- розробити метод визначення психологічних характеристик особистості за рукописним текстом;
- провести експериментальну оцінку результатів дослідження.

Наукова новизна отриманих результатів:

1. Отримала подальший розвиток модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту, яка, на відміну від існуючих рішень, інтегрує в єдиному функціональному циклі метод визначення особистості на основі зорового трансформера та метод ідентифікації автора на основі масштабно-інваріантних центрально-окружних детекторів і зорового трансформера, що забезпечує комплексне виконання задач ідентифікації автора та його психологічного профілювання в межах однієї інформаційної системи.

2. Набув подальшого розвитку метод визначення психологічних характеристик особистості за рукописним текстом, який, на відміну від існуючих підходів на основі згорткових нейронних мереж, використовує архітектуру зорового трансформера та психометричну модель «Велика п'ятірка», що забезпечує одночасне врахування локальних і глобальних графологічних ознак почерку та дозволяє підвищити точність класифікації психологічних характеристик.

3. Удосконалено метод ідентифікації та пошуку автора, в якому, за рахунок інтеграції масштабно-інваріантних центрально-окружних детекторів отримується більша кількість ознак з зображення рукописного тексту. Запропонований метод забезпечив збільшення значень точності ідентифікації,

а також кількісне скорочення як середнього часу на елемент, так і загальної тривалості попередньої обробки, а також збільшення загальної кількості вилучених ознак.

Проведено аналіз предметної області за темою дослідження, розглянуто основні процеси аналізу почерку та особливості графологічного аналізу рукописного тексту. Здійснено огляд існуючих методів ідентифікації автора, визначення особистості, психометричних моделей особистості та методів забезпечення цілісності цифрових зображень. Сформульовано мету та завдання дисертаційної роботи.

Запропоновано метод визначення психологічних характеристик особистості за рукописним текстом, що набув подальшого розвитку на базі архітектури зорового трансформера, та обґрунтовано використання психометричної моделі «Велика п'ятірка» та встановлено кореляцію її рис із графологічними особливостями почерку.

Запропоновано удосконалений метод ідентифікації та пошуку автора, в якому за рахунок інтеграції масштабно-інваріантних центрально-окружних детекторів із зоровим трансформером отримується більша кількість ознак із зображення рукописного тексту. Запропонований метод забезпечив збільшення кількості вилучених ознак на 71% та скорочення середнього часу попередньої обробки на 40%.

Запропоновано модель інформаційної технології автоматизованого графологічного аналізу, яка об'єднує методи визначення особистості та ідентифікації автора для забезпечення комплексного аналізу, що включає як ідентифікацію автора, так і його психологічне профілювання. Побудовано функціональну та структурну моделі інформаційної технології, здійснено математичну формалізацію задачі та розроблено архітектуру інформаційної системи на базі хмарних технологій.

Практичне значення результатів роботи полягає в тому, що розроблені у дисертації модель інформаційної технології автоматизованого графологічного аналізу, удосконалені методи визначення особистості та ідентифікації автора є науково-практичною основою для подальшого удосконалення систем поведінкової біометрії та криміналістичної експертизи.

Запропоновані моделі та методи дозволяють: збільшити точність визначення психологічних характеристик особистості, збільшити точність ідентифікації та пошуку автора за рукописним текстом, збільшити кількість виділених ознак рукописного тексту, зменшити середній час на елемент та загальну тривалість попередньої обробки зображень рукописного тексту.

Отриманий метод визначення особистості на основі зорового трансформера підтверджено випробуваннями, які довели підвищення точності психологічного профілювання автора (ПРАТ «Бізнес Автоматика»). Впроваджена модель інформаційної технології автоматизованого аналізу рукописного тексту забезпечила повний цикл обробки даних, від автентифікації до формування вихідних звітів у межах єдиної архітектури, про що свідчить акт від 20.02.2026.

Матеріали дисертації повною мірою викладені у 10 публікаціях, зокрема 4 статті, з них – 2 статті у наукових фахових виданнях України категорії Б та 2 статті, що індексуються в Scopus, з них - одна - 3 квартиль, 6 тез доповідей у матеріалах міжнародних наукових конференцій.

Ключові слова: інформаційна технологія, інформаційна система, рекомендаційна система, машинне навчання, помилки першого та другого роду, розпізнавання, визначення особистості, ідентифікація автора, системи підтримки рішень, зоровий трансформер, згорткова нейронна мережа, глибоке навчання, програмне забезпечення, обробка зображення, обробка візуальних даних.

Список публікацій здобувача

1. Shupyliuk M., Martovytskyi V. Model of information technology for automated handwriting analysis. Scientific notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences. 2026. P. 442–447. DOI: <https://doi.org/10.32782/2663-5941/2026.1.2/55> (Фахове видання категорії Б).
2. Shupyliuk M., Martovytskyi V., Romanenkov Y. Enhancing writer identification and writer retrieval with CenSurE and Vision Transformers. Technology Audit and Production Reserves. 2025. Vol. 6, No. 2(86). P. 6–14. DOI: <https://doi.org/10.15587/2706-5448.2025.343943> (Входить до міжнародної наукометричної бази Scopus).
3. Shupyliuk M., Martovytskyi V. Analysis of personality detection and writer identification methods. Control, Navigation and Communication Systems. 2025. Vol. 79, No. 1. P. 138–142. DOI: <https://doi.org/10.26906/SUNZ.2025.1.138-142> (Фахове видання категорії Б).
4. Shupyliuk M., Martovytskyi V., Bolohova N., Romanenkov Y., Osiievskyi S., Liashenko S., Nesmiiian O., Nikiforov I., Sukhoteplyi V., Lapchenkov Y. Devising an approach to personality identification based on handwritten text using a vision transformer. Eastern-European Journal of Enterprise Technologies. 2025. Vol. 1, No. 2(133). P. 53–65. DOI: <https://doi.org/10.15587/1729-4061.2025.322726> (Входить до міжнародної наукометричної бази Scopus, 3 квартиль).
5. Smelyakov K., Shupyliuk M., Martovytskyi V., Tovchyrechko D., Ponomarenko O. Efficiency of image convolution. 2019 IEEE 8th International Conference on Advanced Optoelectronics and Lasers (CAOL). 2019. P. 578–583. DOI: <https://doi.org/10.1109/CAOL46282.2019.9019450> (Входить до міжнародної наукометричної бази Scopus).

6. Shupyliuk M., Martovytskyi V. Novel Vision Transformer based method for Personality detection from handwriting. Seventh International Scientific and Technical Conference COMPUTER AND INFORMATION SYSTEMS AND TECHNOLOGIES. 2024. P. 5–6.
7. Шупилюк М. В., Мартовицький В. О. Дослідження методів попередньої обробки зображень з рукописним текстом. Матеріали дванадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2(4), с. 111, 2024. DOI: <https://doi.org/10.32620/PI.24.t2>
8. Шупилюк М. В., Мартовицький В. О. Analysis of keypoint detection methods for images with handwritten text. Матеріали п'ятнадцятої міжнародної науково–технічної конференції Сучасні напрями розвитку інформаційно–комунікаційних технологій та засобів управління, том 2(2), с. 61, 2025. DOI: <https://doi.org/10.32620/ICT.25.t2>
9. Шупилюк М. В., Мартовицький В. О. Підхід до ідентифікації автора за допомогою sentence та зорового трансформера. Матеріали тринадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 3(4), с. 39, 2025. DOI: <https://doi.org/10.32620/PI.25.t3>
10. Шупилюк М. В., Мартовицький В. О. Model of information technology for automated handwriting analysis. Матеріали шістнадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2, с. 132, 2026. DOI: <https://doi.org/10.32620/ICT.26.t2>

ABSTRACT

Shupyliuk M. Model and methods of information technology for automated graphological analysis of handwritten text. - Qualification scientific work on the rights of manuscript.

Dissertation for the degree of Doctor of Philosophy in the specialty 126 “Information systems and technologies” (12 - Information Technologies) - Kharkiv National University of Radio Electronics, Ministry of Education and Science of Ukraine, Kharkiv, 2026.

The dissertation work is devoted to the development of a model and improvement of information technology methods for graphological analysis of handwritten text images. The work presents approaches to automated graphological analysis of handwritten text, which provide a comprehensive and effective analysis, which, in addition to writer identification, also allows for psychological profiling. The proposed changes include both an improved method for personality detection and an improved method for identification and retrieval of the writer.

The relevance of the chosen topic is determined by important role of graphological analysis of handwritten text in the tasks of forensic examination, security and research of historical documents. One of the main prerequisites for effective graphological analysis is the ability of the system to take into account the high intra-author variability of handwriting, which arises under the influence of physical, psychological and situational factors. The limitations of existing methods have a significant impact not only on the accuracy of writer identification and personality detection, but also on the possibility of comprehensive analysis of handwritten text. In this regard, the question arises of increasing the efficiency of automated graphological analysis by developing new models and methods. In

traditional approaches to handwriting analysis (structural, textural, grapheme and based on deep learning), current limitations, such as high dependence on pre-processing of images, the need for large amounts of training data and significant computational costs, are not taken into account at a sufficient level. This, of course, leads to a decrease in the reliability and accuracy of the analysis. It should also be noted that most studies focus on individual tasks (writer identification or signature verification), while a comprehensive approach that combines writer identification with the determination of a psychological profile of a person remains understudied. This can be solved by using tools from modern deep learning architectures, in particular visual transformers, which are able to effectively detect global and local patterns in handwritten text images. The very approaches that combine visual transformers with keypoint detection methods open up new opportunities for creating more accurate and complex graphological analysis systems. Thus, the study of effective methods for automated graphological analysis of handwritten text using modern deep learning architectures makes the work relevant.

The aim of the dissertation is to improve the accuracy and automation of the analysis of handwritten text images by developing models and methods of automated graphological analysis of handwritten text.

The object of the study is the process of processing and graphological analysis of handwritten text.

The subject of the study is models and methods of graphological analysis of handwritten text images, which are aimed at the efficiency and automation of the analysis.

The first section of the dissertation presents an analysis of the current state of research in the field of graphological analysis of handwritten text. The main processes of handwriting analysis and the features of graphological analysis are considered. A review and analysis of modern methods of writer identification and

methods of personality detection are carried out. Psychometric models of personality, methods of applying a digital watermark, as well as technologies for secure data storage are analyzed. Based on the analysis, the goal and objectives of the study are formulated, which are related to the development of a model and improvement of information technology methods for automated graphological analysis of handwritten text.

The second section presents an improved method for personality detection based on handwritten text. The choice of the Big Five psychometric model is justified and the correlation of its features with the corresponding graphological features of handwriting is established. A functional model for personality detection is proposed, which includes stages of preprocessing of input data using scale-invariant feature transformation. For direct determination of personality, the deep learning architecture Visual Transformer is used. The choice of software tools is justified and the CVL dataset is selected for training. The results of an experimental study are presented, which confirm the effectiveness of the improved method with a maximum classification accuracy of up to 99.48%.

The third section presents an improved method for writer identification based on handwritten text. The choice of scale-invariant center-surround detectors as a method for detecting image features is justified. An approach to writer identification is proposed that integrates center-surround detectors and the deep learning architecture Visual Transformer. A functional model of writer identification is constructed and its formalization is carried out by building a mathematical model. The results of an experimental study are presented, which confirm an increase in the total number of extracted features by 71%, a reduction in the average pre-processing time by 40%, and an increase in identification indicators.

The fourth section presents an information technology model of automated graphological analysis of handwritten text. A functional and structural model of

information technology has been built, and a mathematical formalization of the problem has been carried out. An information system architecture based on cloud technologies has been developed, which covers the full data processing cycle from image loading to the formation of a comprehensive report. Digital watermarks have been integrated as a method of image protection and mechanisms for ensuring data integrity based on cryptographic hashing have been implemented.

The conclusions of the dissertation work summarize the results of the research. The significance of the obtained scientific and practical results has been summarized. It is shown that the proposed information technology model and improved methods for personality detection and writer identification using the Visual Transformer and scale-invariant detectors provide increased accuracy and automation of graphological analysis of handwritten text. The implementation of the developed methods significantly improves the overall efficiency of automated graphological analysis.

Research methods – The study used a wide range of methods. Thus, when analyzing the current state of research the methods of system analysis, process approach and comparative analysis of existing solutions were used, which allowed to systematize approaches to writer identification, personality detection and ensuring the integrity of digital images, as well as to formulate the goal and objectives of the study. When developing an method for personality detection from handwritten text, deep learning methods were used, in particular the architecture of the Visual Transformer, the psychometric model “Big Five” and scale-invariant feature transformation, methods for determining the threshold for image binarization and grayscale for pre-processing of input data, which made it possible to effectively capture both local and global features of handwriting and increase the accuracy of personality determination. When developing an improved method for writer identification and retrieval from handwritten text, scale-invariant center-surround

detectors were used in integration with the deep learning architecture Visual Transformer, as well as methods of mathematical modeling and functional modeling, which allowed to increase the number of extracted features and reduce the pre-processing time. When developing an information technology model for automated graphological analysis of handwritten text, methods of mathematical formalization, functional and structural modeling were used, as well as cloud computing technologies, digital watermarks and cryptographic hashing to ensure data integrity and protection. The evaluation of the experimental data obtained in the course of the work was carried out on the basis of a comparative analysis of performance indicators (Hard Top K, Soft Top K, mAP) and classification accuracy metrics. The choice of research methods ensured the reliability of the obtained results and conclusions, which is confirmed by the convergence of the results of experimental studies obtained during the software implementation of the information technology of automated graphological analysis of handwritten text with theoretical calculations, as well as agreement with predictive estimates of the effectiveness of the proposed approaches.

Research objectives:

- to analyze the current state of research in the field of handwriting analysis;
- to develop an information technology model for automated handwriting analysis;
- to improve the method of author identification;
- to improve the method of personality determination;
- to conduct an experimental evaluation of the research results.

Scientific novelty of the results obtained:

1. The information technology model of automated graphological analysis of handwritten text has been further developed, which, unlike existing solutions,

integrates in a single functional cycle a method of personality determination based on a visual transformer and a method of writer identification based on scale-invariant central-peripheral detectors and a visual transformer, which ensures the comprehensive implementation of the tasks of author identification and his psychological profiling within the framework of a single information system.

2. The method for personality psychological characteristics detection from handwritten text has been further developed, which, unlike existing approaches based on convolutional neural networks, uses the architecture of a visual transformer and the “Big Five” psychometric model, which ensures simultaneous consideration of local and global graphological features of handwriting and allows for increased accuracy in the classification of psychological characteristics.

3. The method of writer identification and retrieval has been improved, in which, due to the integration of scale-invariant center surround extremas, a larger number of features is obtained from the handwritten text image. The proposed method provided an increase in the identification accuracy values, as well as a quantitative reduction in both the average time per element and the total duration of pre-processing, as well as an increase in the total number of extracted features.

The subject area of the research topic was analyzed, the main processes of handwriting analysis and the features of graphological analysis of handwritten text were considered. The existing methods of writer identification, personality detection, psychometric models of personality and methods of ensuring the integrity of digital images were reviewed. The goal and objectives of the dissertation were formulated.

An improved method of personality detection from handwritten text based on the architecture of the visual transformer was proposed, the use of the psychometric model “Big Five” was justified and the correlation of its features with graphological features of handwriting was established. This method is suitable for tasks of

automated detection of psychological characteristics due to the ability of the attention mechanism to record local and global features of handwriting.

An improved method of writer identification and retrieval is proposed, in which a larger number of features from the handwritten text image are obtained by integrating scale-invariant central-surround detectors with a visual transformer. The proposed method provided an increase in the number of extracted features by 71% and a reduction in the average pre-processing time by 40%.

An information technology model of automated graphological analysis is proposed, which combines advanced methods of personality detection and writer identification to provide a comprehensive analysis that includes both writer identification and psychological profiling. A functional and structural model of information technology is built, a mathematical formalization of the problem is carried out, and an information system architecture based on cloud technologies is developed.

The practical significance of the results of the work lies in the fact that the information technology model of automated graphological analysis, improved methods of personality detection and writer identification developed in the dissertation are a scientific and practical basis for further improvement of behavioral biometrics and forensic examination systems.

The proposed models and methods allow: to increase the accuracy of psychological characteristics detection when detecting personality, to increase the accuracy of identification and retrieval based on handwritten text, to increase the number of highlighted handwritten text features, to reduce the average time per element and the total duration of pre-processing of handwritten text images.

The obtained method of determining personality based on a visual transformer was confirmed by tests that proved an increase in the accuracy of the author's psychological profiling (PJSC "Business Automation"). The implemented

information technology model of automated handwritten text analysis provided a full cycle of data processing, from authentication to the formation of output reports within a single architecture, as evidenced by the act dated 02/20/2026.

The materials of the dissertation are fully presented in 10 publications, including 4 articles, of which 2 articles in Ukrainian scientific journals of category B and 2 articles indexed in Scopus, of which one is 3rd quartile, 6 abstracts of reports in the materials of international scientific conferences.

Keywords: information technology, information system, recommender system, machine learning, errors of the first and second kinds, recognition, writer identification, personality detection, decision support systems, visual transformer, convolution neural network, deep learning, software, image processing, visual data processing.

List of publications of the applicant

1. Shupyliuk M., Martovytskyi V. Model of information technology for automated handwriting analysis. Scientific notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences. 2026. P. 442–447. DOI: <https://doi.org/10.32782/2663-5941/2026.1.2/55> (Фахове видання категорії Б).
2. Shupyliuk M., Martovytskyi V., Romanenkov Y. Enhancing writer identification and writer retrieval with CenSurE and Vision Transformers. Technology Audit and Production Reserves. 2025. Vol. 6, No. 2(86). P. 6–14. DOI: <https://doi.org/10.15587/2706-5448.2025.343943> (Входить до міжнародної наукометричної бази Scopus).
3. Shupyliuk M., Martovytskyi V. Analysis of personality detection and writer identification methods. Control, Navigation and Communication Systems. 2025. Vol. 79, No. 1. P. 138–142. DOI: <https://doi.org/10.26906/SUNZ.2025.1.138-142> (Фахове видання категорії Б).
4. Shupyliuk M., Martovytskyi V., Bolohova N., Romanenkov Y., Osiievskyi S., Liashenko S., Nesmiian O., Nikiforov I., Sukhoteplyi V., Lapchenkov Y. Devising an approach to personality identification based on handwritten text using a vision transformer. Eastern-European Journal of Enterprise Technologies. 2025. Vol. 1, No. 2(133). P. 53–65. DOI: <https://doi.org/10.15587/1729-4061.2025.322726> (Входить до міжнародної наукометричної бази Scopus, 3 квартиль).
5. Smelyakov K., Shupyliuk M., Martovytskyi V., Tovchyrechko D., Ponomarenko O. Efficiency of image convolution. 2019 IEEE 8th International Conference on Advanced Optoelectronics and Lasers (CAOL). 2019. P. 578–583. DOI: <https://doi.org/10.1109/CAOL46282.2019.9019450> (Входить до міжнародної наукометричної бази Scopus).

6. Shupyliuk M., Martovytskyi V. Novel Vision Transformer based method for Personality detection from handwriting. Seventh International Scientific and Technical Conference COMPUTER AND INFORMATION SYSTEMS AND TECHNOLOGIES. 2024. P. 5–6.
7. Шупилюк М. В., Мартовицький В. О. Дослідження методів попередньої обробки зображень з рукописним текстом. Матеріали дванадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2(4), с. 111, 2024. DOI: <https://doi.org/10.32620/PI.24.t2>
8. Шупилюк М. В., Мартовицький В. О. Analysis of keypoint detection methods for images with handwritten text. Матеріали п'ятнадцятої міжнародної науково–технічної конференції Сучасні напрями розвитку інформаційно–комунікаційних технологій та засобів управління, том 2(2), с. 61, 2025. DOI: <https://doi.org/10.32620/ICT.25.t2>
9. Шупилюк М. В., Мартовицький В. О. Підхід до ідентифікації автора за допомогою sentence та зорового трансформера. Матеріали тринадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 3(4), с. 39, 2025. DOI: <https://doi.org/10.32620/PI.25.t3>
10. Шупилюк М. В., Мартовицький В. О. Model of information technology for automated handwriting analysis. Матеріали шістнадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2, с. 132, 2026. DOI: <https://doi.org/10.32620/ICT.26.t2>

ЗМІСТ

Перелік умовних скорочень	24
Вступ	29
1 Огляд стану проблеми та постановка завдання дослідження	38
1.1 Опис процесів розпізнавання рукописного тексту	38
1.2 Огляд методів ідентифікації автора	43
1.3 Аналіз існуючих моделей визначення особистості	46
1.4 Аналіз психометричних моделей особистості	49
1.5 Огляд методів нанесення цифрового водяного знака	52
1.6 Аналіз технологій безпечного зберігання даних	57
1.7 Постановка задач дослідження	61
1.8 Висновки за першим розділом	62
2 Набув подальшого розвитку метод визначення психологічних характеристик особистості	65
2.1 Постановка задачі визначення психологічних характеристик особистості	65
2.2 Підхід визначення психологічних характеристик	72
2.3 Метод визначення психологічних характеристик особистості	81
2.4 Експериментальне дослідження методу визначення особистості	90
2.5 Висновки за другим розділом	102
3 Удосконалений метод ідентифікації та пошуку автора	104
3.1 Постановка задачі ідентифікації та пошуку автора	104
3.2 Метод виявлення ознак зображень	112
3.3 Метод ідентифікації автора	118

3.4 Експериментальна перевірка метода ідентифікації автора	126
3.5 Висновки за третім розділом	135
4 Модель інформаційної технології автоматизованого аналізу рукописного тексту.....	137
4.1 Постановка задачі автоматизованого аналізу рукописного тексту	137
4.2 Інформаційна технологія автоматизованого графологічного аналізу рукописного тексту.....	141
4.3 Інформаційна система автоматизованого графологічного аналізу рукописного тексту.....	148
4.4 Висновки за четвертим розділом	155
Висновки	157
Перелік джерел посилання	160
Додаток А Список публікацій здобувача	174
Додаток Б Акт впровадження	176

ПЕРЕЛІК УМОВНИХ СКОРОЧЕНЬ

- ЗНМ – згорткова нейронна мережа
- КНДІСЕ – Київський науково-дослідний інститут судових експертиз
- ПМП – пряма нейронна мережа (перцептрон прямого поширення)
- ТІ – точки інтересу
- ANN – штучна нейронна мережа (англ. Artificial Neural Network)
- API – програмний інтерфейс застосунку (англ. Application Programming Interface)
- ASIFT – афінно-масштабно-інваріантне перетворення ознак (англ. Affine Scale-Invariant Feature Transform)
- AWS – хмарний сервіс Amazon (англ. Amazon Web Services)
- BER – коефіцієнт бітових помилок (англ. Bit Error Rate)
- BOF – пакет ознак (англ. Bag of Features)
- CDN – глобальна мережа доставки контенту (англ. Content Delivery Network)
- CEDAR – Центр передового досвіду в аналізі документів (англ. Center of Excellence for Document Analysis and Recognition)
- CEDAR-FOX – система судово-медичного аналізу документів CEDAR
- CENSURE – центрально-окружний детектор (англ. Center Surround Extremas)
- CJIS – Інформаційні служби кримінального правосуддя (англ. Criminal Justice Information Services)
- CNN – згорткова нейронна мережа (англ. Convolutional Neural Network)
- CSAFE – Центр статистики та застосувань у судових доказах (англ. Center for Statistics and Applications in Forensic Evidence)

CTC – конекціоністська часова класифікація (англ. Connectionist Temporal Classification)

CVL – набір даних лабораторії комп'ютерного зору (англ. Computer Vision Lab)

DCT – дискретне косинусне перетворення (англ. Discrete Cosine Transform)

DFT – дискретне перетворення Фур'є (англ. Discrete Fourier Transform)

DHT – розподілена хеш-таблиця (англ. Distributed Hash Table)

DL – глибинне навчання (англ. Deep Learning)

DWT – дискретне вейвлет-перетворення (англ. Discrete Wavelet Transform)

F1 – міра F1 (англ. F1-score)

FAST – детектор ключових точок (англ. Features from Accelerated Segment Test)

FFNN – нейронна мережа прямого зв'язку (англ. Feedforward Neural Network)

FFM – п'ятифакторна модель (англ. Five-Factor Model)

FN – хибний негативний результат (англ. False Negative)

FP – хибний позитивний результат (англ. False Positive)

GMM – модель гауссових сумішей (англ. Gaussian Mixture Model)

GR-RNN – глобально-контекстна залишкова рекурентна нейронна мережа (англ. Globally-contextual Residual Recurrent Neural Network)

HC – детектор кутів Гарріса (англ. Harris Corner)

HEXACO – шестифакторна модель особистості

HIPAA – Закон про перенесення та підзвітність медичного страхування (англ. Health Insurance Portability and Accountability Act)

HM – карта рукописного письма (англ. Handwriting Map)

HOG – гістограма орієнтованого градієнта (англ. Histogram of Oriented Gradients)

IAM – набір даних рукописного тексту (англ. IAM Handwriting Database)

INR – неявні нейронні представлення (англ. Implicit Neural Representations)

IPFS – Міжпланетна файлова система (англ. InterPlanetary File System)

JSON – формат обміну даними (англ. JavaScript Object Notation)

JWT – веб-токен JSON (англ. JSON Web Token)

KEK – ключ шифрування ключа (англ. Key Encryption Key)

KES – Служба шифрування ключів (англ. Key Encryption Service)

KMS – система керування ключами (англ. Key Management System)

KNN – класифікатор К-найближчих сусідів (англ. K-Nearest Neighbors)

LBP – локальний бінарний шаблон (англ. Local Binary Pattern)

LLR – коефіцієнт правдоподібності (англ. Log-Likelihood Ratio)

LSTM – довга короткочасна пам'ять (англ. Long Short-Term Memory)

mAP – середня середня точність (англ. mean Average Precision)

MBTI – Індикатор типу Майєрса-Бріггса (англ. Myers-Briggs Type Indicator)

MFA – багатофакторна автентифікація (англ. Multi-Factor Authentication)

MHA – механізм багаторівневої уваги (англ. Multi-Head Attention)

MLP – багатошаровий перцептрон (англ. Multi-Layer Perceptron)

MSA – механізм самоуваги (англ. Multi-Head Self-Attention)

NC – нормалізована кореляція (англ. Normalized Correlation)

OCEAN – модель особистості: Відкритість, Сумлінність, Екстраверсія, Співробітництво, Нейротизм (англ. Openness, Conscientiousness, Extraversion, Agreeableness, Neuroticism)

OEK – ключ шифрування об'єкта (англ. Object Encryption Key)

ORB – детектор та дескриптор ключових точок (англ. Oriented FAST and Rotated BRIEF)

PAN – мережа аналізу особистості (англ. Personality Analysis Network)

PDF – формат переносних документів (англ. Portable Document Format)

PII – персональна ідентифікаційна інформація (англ. Personally Identifiable Information)

PTLDM – модель виявлення рівня рис особистості (англ. Personality Trait Level Detection Model)

RBAC – контроль доступу на основі ролей (англ. Role-Based Access Control)

RNN – рекурентна нейронна мережа (англ. Recurrent Neural Network)

RST – масштабування, обертання, переміщення (англ. Rotation, Scaling, Translation)

S3 – хмарне сховище об'єктів Amazon (англ. Amazon Simple Storage Service)

SGAN – напівконтрольована генеративно-змагальна мережа (англ. Semi-supervised Generative Adversarial Network)

SGD – стохастичний градієнтний спуск (англ. Stochastic Gradient Descent)

SGX – розширення захисту програмного забезпечення Intel (англ. Software Guard Extensions)

SHA-256 – алгоритм криптографічного хешування (англ. Secure Hash Algorithm 256-bit)

SIFT – масштабно-інваріантне перетворення ознак (англ. Scale-Invariant Feature Transform)

SSE – шифрування на стороні сервера (англ. Server-Side Encryption)

SSL – напівконтрольоване навчання (англ. Semi-Supervised Learning)

SURF – прискорені стійкі ознаки (англ. Speeded Up Robust Features)

SVD – розкладання за сингулярними значеннями (англ. Singular Value Decomposition)

SVM – метод опорних векторів (англ. Support Vector Machine)

TCB – база довірених обчислень (англ. Trusted Computing Base)

TDX – розширення довіри Intel (англ. Trust Domain Extensions)

TEE – надійне середовище виконання (англ. Trusted Execution Environment)

TN – істинний негативний результат (англ. True Negative)

TP – істинний позитивний результат (англ. True Positive)

ViT – Зоровий трансформер (англ. Vision Transformer)

VLAD – вектор локально агрегованих дескрипторів (англ. Vector of Locally Aggregated Descriptors)

VPC – віртуальна приватна хмара (англ. Virtual Private Cloud)

WANDA – система аналізу рукописного тексту (англ. Writing Analysis and Description Automatized)

WLSR – зважена регуляризація згладжування міток (англ. Weighted Label Smoothing Regularization)

WORM – запис один раз, читання багато разів (англ. Write-Once-Read-Many)

XML – розширювана мова розмітки (англ. Extensible Markup Language)

ВСТУП

Актуальність теми дослідження.

В умовах стрімкого розвитку цифрових технологій та зростання диджиталізації історичних і юридичних архівів формуються нові вимоги до методів аналізу та ідентифікації рукописного тексту. Типовими прикладами застосувань таких методів є криміналістична експертиза, забезпечення безпеки, дослідження історичних документів та психологічне профілювання. Сучасний автоматизований графологічний аналіз передбачає обробку складних зображень рукописного тексту, що характеризуються високою внутрішньоавторською мінливістю почерку, яка виникає під впливом фізичних, психологічних та ситуативних факторів. Обмеження існуючих методів здатне не тільки знизити точність ідентифікації автора та визначення особистості, а й призвести до неможливості комплексного аналізу рукописного тексту. Саме тому питання підвищення точності розпізнавання, автоматизації аналізу зображень, а також забезпечення комплексності обробки є ключовими для побудови сучасних систем графологічного аналізу. Традиційні підходи до аналізу почерку (структурні, текстурні, графемні та засновані на класичних нейронних мережах) ґрунтуються на фіксованих стратегіях вилучення ознак і не завжди враховують значну варіативність характеристик почерку, зумовлену індивідуальними особливостями автора, умовами написання та якістю зображень. Це призводить до обмеження здатності забезпечувати високу надійність та точність автоматизованого графологічного аналізу.

Використання сучасних архітектур глибокого навчання, зокрема зорових трансформерів, відкриває нові можливості для створення адаптивних

методів, що здатні ефективно виявляти глобальні та локальні закономірності у зображеннях рукописного тексту і, в свою чергу, формувати рішення з урахуванням їх реальних характеристик. Одним із таких підходів є застосування механізму самоуваги зорового трансформера, який дозволяє відображати та аналізувати багатовимірні ознаки зображень почерку, виявляти приховані залежності й підвищувати точність класифікації. Проте класичні методи вилучення ознак із зображень не пристосовані до умов, коли критичним фактором є баланс між кількістю виявлених ознак, часом попередньої обробки та точністю ідентифікації. Це потребує їх удосконалення з урахуванням комплексних показників: масштабно-інваріантних характеристик ключових точок, просторових закономірностей розміщення ознак та взаємозв'язків між графологічними елементами. Удосконалення методів ідентифікації автора та визначення особистості з урахуванням цих параметрів дозволяє суттєво підвищити ефективність аналізу, покращити надійність розпізнавання та розширити функціональні можливості системи.

Також важливими є питання впровадження таких сучасних підходів у загальну модель інформаційної технології графологічного аналізу. З огляду на це, актуальною є розробка моделей, які поєднують ідентифікацію і пошук автора та визначення особистості на базі психометричних моделей особистості, забезпечуючи тим самим комплексний аналіз рукописного тексту включаючи психологічне профілювання автора. Це відповідає актуальним науковим і практичним потребам.

Таким чином, актуальність дослідження зумовлена потребою підвищення точності та автоматизації графологічного аналізу рукописного тексту, що потребує використання сучасних архітектур глибокого навчання та адаптивних підходів до вилучення та аналізу ознак почерку. Запропоновані у дослідженні модель інформаційної технології автоматизованого

графологічного аналізу, удосконалений метод визначення особистості на базі зорового трансформера та удосконалений метод ідентифікації автора з інтеграцією масштабно-інваріантних центрально-окружних детекторів забезпечують підвищення точності класифікації, скорочення часу обробки та комплексність аналізу рукописного тексту.

Метою дисертаційної роботи є підвищення точності та автоматизація аналізу зображень рукописного тексту шляхом розробки моделей та методів автоматизованого графологічного аналізу рукописного тексту.

Для досягнення поставленої мети необхідно вирішити **завдання дослідження**:

- провести аналіз сучасного стану досліджень у сфері аналізу рукописного тексту;
- розробити модель інформаційної технології автоматизованого аналізу рукописного тексту;
- удосконалити метод ідентифікації автора;
- удосконалити метод визначення особистості;
- провести експериментальну оцінку результатів дослідження.

Об'єкт дослідження – процес обробки та графологічного аналізу рукописного тексту.

Предмет дослідження – моделі та методи графологічного аналізу зображень рукописного тексту, які направлені на ефективність та автоматизацію аналізу.

Методи дослідження – В дослідженні було використано широкий спектр методів. Так, при проведенні аналізу сучасного стану досліджень у сфері графологічного аналізу рукописного тексту використовувались методи системного аналізу, процесний підхід та порівняльний аналіз існуючих рішень, що дозволило систематизувати підходи до ідентифікації автора,

визначення особистості та забезпечення цілісності цифрових зображень, а також сформулювати мету та задачі дослідження. При розробці удосконаленого методу визначення особистості за рукописним текстом використовувались методи глибокого навчання, зокрема архітектура Зорового трансформера з механізмом уваги, психометрична модель «Велика п'ятірка» та масштабно-інваріантне ознакове перетворення, методи визначення порогу для бінаризації зображень та градації сірого для попередньої обробки вхідних даних, що дало змогу ефективно фіксувати як локальні, так і глобальні ознаки почерку та підвищити точність визначення особистості. При розробці удосконаленого методу ідентифікації та пошуку автора за рукописним текстом використовувались масштабно-інваріантні центрально-окружні детектори в інтеграції із архітектурою глибокого навчання Зоровий трансформер, а також методи математичного моделювання та функціонального моделювання, що дозволило збільшити кількість вилучених ознак та скоротити час попередньої обробки. При розробці моделі інформаційної технології автоматизованого графологічного аналізу рукописного тексту використовувались методи математичної формалізації, функціонального та структурного моделювання, а також технології хмарних обчислень, цифрові водяні знаки та криптографічне хешування для забезпечення цілісності та захисту даних. Оцінка експериментальних даних, отриманих у ході роботи, проводилася на основі порівняльного аналізу показників ефективності (Hard Top K, Soft Top K, mAP) та метрик точності класифікації. Вибір методів досліджень забезпечив достовірність отриманих результатів і висновків, що підтверджується збіжністю результатів експериментальних досліджень, отриманих при програмній реалізації інформаційної технології автоматизованого графологічного аналізу рукописного тексту, з теоретичними розрахунками, а

також узгодженням із прогностичними оцінками ефективності запропонованих підходів.

Наукова новизна отриманих результатів:

1. Отримала подальший розвиток модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту, яка, на відміну від існуючих рішень, інтегрує в єдиному функціональному циклі метод визначення особистості на основі зорового трансформера та метод ідентифікації автора на основі масштабно-інваріантних центрально-окружних детекторів і зорового трансформера, що забезпечує комплексне виконання задач ідентифікації автора та його психологічного профілювання в межах однієї інформаційної системи.

2. Набув подальшого розвитку метод визначення психологічних характеристик особистості за рукописним текстом, який, на відміну від існуючих підходів на основі згорткових нейронних мереж, використовує архітектуру зорового трансформера та психометричну модель «Велика п'ятірка», що забезпечує одночасне врахування локальних і глобальних графологічних ознак почерку та дозволяє підвищити точність класифікації психологічних характеристик.

3. Удосконалено метод ідентифікації та пошуку автора, в якому, за рахунок інтеграції масштабно-інваріантних центрально-окружних детекторів отримується більша кількість ознак з зображення рукописного тексту. Запропонований метод забезпечив збільшення значень точності ідентифікації, а також кількісне скорочення як середнього часу на елемент, так і загальної тривалості попередньої обробки, а також збільшення загальної кількості вилучених ознак.

Проведено аналіз предметної області за темою дослідження, розглянуто основні процеси аналізу почерку та особливості графологічного аналізу

рукописного тексту. Здійснено огляд існуючих методів ідентифікації автора, визначення особистості, психометричних моделей особистості та методів забезпечення цілісності цифрових зображень. Сформульовано мету та завдання дисертаційної роботи.

Запропоновано удосконалений метод визначення особистості за рукописним текстом на базі архітектури зорового трансформера, обґрунтовано використання психометричної моделі «Велика п'ятірка» та встановлено кореляцію її рис із графологічними особливостями почерку. Даний метод придатний для задач автоматизованого визначення психологічних характеристик за рахунок здатності механізму уваги фіксувати локальні та глобальні ознаки почерку.

Запропоновано удосконалений метод ідентифікації та пошуку автора, в якому за рахунок інтеграції масштабно-інваріантних центрально-окружних детекторів із зоровим трансформером отримується більша кількість ознак із зображення рукописного тексту. Запропонований метод забезпечив збільшення кількості вилучених ознак на 71% та скорочення середнього часу попередньої обробки на 40%.

Запропоновано модель інформаційної технології автоматизованого графологічного аналізу, яка об'єднує удосконалені методи визначення особистості та ідентифікації автора для забезпечення комплексного аналізу, що включає як ідентифікацію автора, так і його психологічне профілювання. Побудовано функціональну та структурну моделі інформаційної технології, здійснено математичну формалізацію задачі та розроблено архітектуру інформаційної системи на базі технологій хмарних обчислень.

Запропоновано використання механізму захисту зображень на основі цифрових водяних знаків та механізмів забезпечення цілісності та незмінності

даних на основі криптографічного хешування та незмінного сховища з безперервним ланцюгом зберігання доказів

Практичне значення результатів роботи полягає в тому, що розроблені у дисертації модель інформаційної технології автоматизованого графологічного аналізу, удосконалені методи визначення особистості та ідентифікації автора є науково-практичною основою для подальшого удосконалення систем поведінкової біометрії та криміналістичної експертизи.

Запропоновані моделі та методи дозволяють:

- збільшити точність визначення психологічних характеристик при визначенні особистості;
- збільшити точність ідентифікації та пошуку за рукописним текстом;
- збільшити кількість виділених ознак рукописного тексту;
- зменшити середній час на елемент та загальну тривалість попередньої обробки зображень рукописного тексту.

Отриманий метод визначення особистості на основі зорового трансформера підтверджено випробуваннями, які довели підвищення точності психологічного профілювання автора (ПРАТ «Бізнес Автоматика»). Впроваджена модель інформаційної технології автоматизованого аналізу рукописного тексту забезпечила повний цикл обробки даних, від автентифікації до формування вихідних звітів у межах єдиної архітектури, про що свідчить акт від 20.02.2026.

Особистий внесок здобувача. Всі основні результати дисертаційної роботи, які виносяться на захист, отримано автором особисто. В [1] здобувач розробив модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту, побудував функціональну та структурну моделі, здійснив математичну формалізацію задачі та описав архітектуру

інформаційної системи на базі хмарних технологій. В [2] здобувач запропонував удосконалений метод ідентифікації та пошуку автора за рукописним текстом із застосуванням масштабно-інваріантних центрально-окружних детекторів CenSurE в інтеграції із зоровим трансформером, провів експериментальну перевірку методу та аналіз отриманих результатів. В [3] здобувач провів порівняльний аналіз існуючих методів визначення особистості та ідентифікації автора за рукописним текстом, систематизував підходи та обґрунтував доцільність застосування зорових трансформерів для задач графологічного аналізу. В [4] здобувач розробив підхід до визначення особистості за рукописним текстом на базі зорового трансформера з використанням психометричної моделі «Велика п'ятірка», провів експериментальне дослідження та аналіз результатів класифікації. В [5], що виконана у співавторстві, здобувач провів дослідження ефективності згорткових операцій над зображеннями, що стало основою для подальшої розробки методів попередньої обробки зображень рукописного тексту. В тезах доповіді [6] здобувач запропонував новий метод визначення особистості за почерком на базі зорового трансформера та провів попереднє експериментальне дослідження його ефективності. В тезах доповіді [7] здобувач дослідив методи попередньої обробки зображень із рукописним текстом, провів аналіз етапів підготовки даних для застосування методів глибокого навчання. В тезах доповіді [8] здобувач провів аналіз методів виявлення ключових точок на зображеннях із рукописним текстом та обґрунтував вибір детектора CenSurE для задач ідентифікації автора. В тезах доповіді [9] здобувач описав підхід до ідентифікації автора з інтеграцією детектора CenSurE та зорового трансформера, провів порівняльний аналіз із класичними методами вилучення ознак. В тезах доповіді [10] здобувач

представив загальну модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту та описав її основні компоненти.

Апробація результатів дисертації.

Основні положення дисертаційної роботи представлено на таких міжнародних конференціях:

- на Сімнадцятій міжнародній науково-технічній конференції «Комп'ютерні та інформаційні системи і технології», 2024 [6];
- на Дванадцятій міжнародній науково-технічній конференції «Проблеми інформатизації», 2024 [7];
- на П'ятнадцятій міжнародній науково-технічній конференції «Сучасні напрями розвитку інформаційно-комунікаційних технологій та засобів управління», 2025 [8];
- на Тринадцятій міжнародній науково-технічній конференції «Проблеми інформатизації», 2025 [9];
- на Шістнадцятій міжнародній науково-технічній конференції «Сучасні напрями розвитку інформаційно-комунікаційних технологій та засобів управління», 2026 [10];

Публікації. Результати дисертаційної роботи викладені у 10 публікаціях, зокрема 4 статті, з них – 2 статті у наукових фахових виданнях України категорії Б та 2 статті, що індексуються в Scopus, з них 1 виходить в 3 квартиль, 6 тез доповідей у матеріалах міжнародних наукових конференцій.

Структура та обсяг роботи. Дисертація складається зі вступу, чотирьох розділів, висновків, переліку джерел посилання, додатків. Загальний обсяг роботи складає 176 сторінок тексту, що містять 131 сторінку основного тексту, анотація на 20 сторінках, 34 рисунки, 15 таблиць, список використаних джерел із 94 найменувань на 14 сторінках, 2 додатки на 3 сторінках.

1 ОГЛЯД СТАНУ ПРОБЛЕМИ ТА ПОСТАНОВКА ЗАВДАННЯ ДОСЛІДЖЕННЯ

1.1 Опис процесів розпізнавання рукописного тексту

Аналіз почерку функціонує як складне поведінкове дослідження, яке інтерпретує нервово-м'язові шаблони особистості, зафіксовані на фізичному або цифровому носії. По суті, ця дисципліна поділена на окремі операційні галузі, призначені для виконання конкретних аналітичних цілей, зокрема ідентифікації автора, перевірки підпису та визначення особистості. У різноманітних застосуваннях процес аналізу значною мірою залежить від вилучення та перевірки унікальних характеристик почерку. Експерти оцінюють структурні компоненти, такі як поля сторінки, внутрішній міжрядковий інтервал, геометричний нахил окремих літер та узгодженість базової лінії, яка є невидимою лінійною межею, на якій лежать символи [11]. Крім того, при аналізі за допомогою відповідних середовищ захоплення, аналіз почерку може враховувати динамічні характеристики, такі як тиск пера, прискорення та плавність штрихів [12]. Незалежно від того, чи оцінюється класичний офлайн-рукопис, написаний ручкою на папері, чи цифровий динамічний підпис, фундаментальним процесом залишається вичерпне дослідження просторових, структурних та кінематичних параметрів для розуміння походження та природи написаного слова.

Аналіз почерку охоплює численні життєво важливі сектори суспільства, виступаючи незамінним інструментом у правовій, безпековій та академічній сферах. У сфері експертизи сумнівних документів аналіз почерку слугує фундаментальним стовпом криміналістики [13]. Він відіграє важливу роль у

кримінальних розслідуваннях, допомагаючи виявляти докази та визначати вину у справах, пов'язаних з фінансовим шахрайством, листами з вимаганням, крадіжкою особистих даних і навіть серйозними правопорушеннями, такими як вбивства або самогубства, де автентифікація рукописних нотаток є питанням життя і смерті [13]. Суди вже давно покладаються на систематичну оцінку цих фізичних артефактів для підтвердження або спростування свідчень, що робить цілісність цих експертиз життєво важливим компонентом ширшої системи кримінального правосуддя [14].

Окрім своєї безпосередньої юридичної корисності, аналіз почерку має глибоке культурне та академічне значення в дослідженні історичних документів. Палеографи та історики регулярно покладаються на ці процеси для проведення ідентифікації писарів та визначення автентичності рукописів вікової давнини [15]. Пов'язування різних аркушів історичних документів з однією особою дозволяє дослідникам відстежувати міграцію текстів, реконструювати стародавні бібліотечні колекції та розгадувати давні історичні таємниці [15]. Крім того, аналіз еволюції стилів почерку допомагає в точному датуванні недокументованих рукописів, співвідносячи фізичне виконання писарем із усталеними стилістичними прийомами певних історичних періодів [15]. Незалежно від того, чи використовується він для забезпечення безпеки сучасної банківської операції, чи для розкриття соціально-економічного контексту середньовічного письма, аналіз почерку слугує надійним зв'язком між людською поведінкою та документованими доказами.

При цьому аналіз почерку супроводжується серйозними проблемами, які ускладнюють як експертну оцінку людиною, так і механічну оцінку. Основна особливість впливає з плинної природи самого людського почерку, який страждає від величезної мінливості всередині автора [16]. Почерк людини ніколи не буває ідеально статичним; натомість він постійно

коливається через природні внутрішні та зовнішні фактори, такі як фізична втома, психологічний стрес, зміна емоційних станів, вибір письмового інструменту та навіть фізична текстура поверхні для письма [16]. Ця природна мінливість надзвичайно ускладнює встановлення єдиної, абсолютно надійної базової лінії того, що становить справжній стиль письма людини.

Додатковою проблемою є подібність між авторами та навмисні маніпуляції з почерком. Висококваліфіковані самозванці часто здатні точно імітувати мікроскопічні штрихи та геометричні символи іншої людини, розмиваючи межі між автентичним документом та витонченою підробкою [16]. Ця проблема особливо критична під час перевірки підпису, де система повинна відрізнити справжній підпис від імітації [12]. Крім того, експертам доводиться регулярно стикатися із замаскованими підписами – оманливим сценарієм, коли автор навмисно змінює власні природні графічні звички, щоб згодом заперечити право власності на юридичний контракт або фінансовий документ [17].

Технічний формат доступних доказів також вводить суворі обмеження. Традиційна офлайн-документація, яка спирається виключно на статичні скановані зображення готового тексту, не містить критичних динамічних даних, таких як швидкість кінчика пера, прискорення та миттєві зміни тиску, що значно ускладнює виявлення підробки, ніж у динамічних середовищах [16]. Крім того, текстовий зміст загального письма вносить набагато більшу мінливість, ніж структуровані підписи, оскільки текстовий почерк містить нескінченну комбінацію різних літер, розділових знаків та структурних макетів [18]. Під час аналізу історичних рукописів або реальних судово-медичних доказів дані часто сильно пошкоджені деградацією навколишнього середовища, такою як плями від води, вицвіле чорнило та фоновий шум, що робить чітку ізоляцію ознак надзвичайно важкою [15, 19].

Аналіз почерку є складною проблемою, тому було проведено багато досліджень. Для категоризації цих досліджень та визначення переваг і недоліків можна використовувати різні аспекти. Розглядаючи тип ознак як функціональний контекст, можна використовувати методи, описані в літературі які можуть бути класифіковані на чотири основні групи: структурні, текстурні, графемні та методи, що базуються на автоматичному навчанні [20].

Структурні методи враховують форми алографів, а потім застосовують розподіл ймовірності викидів графем. Для структурних систем вартість виконання висока через складність етапів попередньої обробки, які включають сегментацію, бінаризацію та виявлення країв. Сильна залежність від попередньої обробки є обмеженням цих методів, оскільки збій етапу попередньої обробки означає, що етапи вилучення ознак та класифікації схильні до невдач. Концентрація таких методів переважно на формах алографів означає, що додаткова інформація, намальована в одному слові між алографами, пропускається. Додатковим недоліком таких методів є вразливість до будь-яких варіацій характеристик символів або алографів (тобто нахилу, співвідношення сторін).

Текстурні (трансформаційні) методи працюють з оцифрованим зображенням рукописного тексту, які розглядаються як спеціальні текстури та витягують ознаки (після застосування деяких методів трансформації) з усього документа, областей інтересу (або ROI, таких як блоки, комірки сітки, зв'язані компоненти, слова тощо) або фрагментів письма (WF). Текстурні методи є ефективнішими з точки зору виконання, оскільки вони не потребують додаткових етапів попередньої обробки, таких як сегментація та бінаризація. Ще однією перевагою є навчання без параметрів. З іншого боку, текстурні методи зазвичай потребують більше даних для вилучення надійних та високодискримінаційних ознак, при цьому таких даних може не бути

(наприклад судово-медичні експерти, щоб передбачити автентичність, можуть мати справу з невеликими зразками запитуваного почерку). Додатковим недоліком таких методів є потреба у висококонтрастних зображеннях.

Методи на основі графем (на основі пакетів ознак (BOF)) спираються на генерацію кодової книги для всіх типів символів. У таких методах рукописний текст, сегментований на сегменти символів, та графеми генеруються для кожного символу, після чого генерується кодова книга за допомогою будь-якого алгоритму кластеризації, для кожної кодової книги обчислюються показники відстані гістограми, а отримана гістограма використовується для класифікації. Такі методи вважаються ефективними для ідентифікації автора, оскільки вони не потребують надмірної кількості оригінального зразка почерку автора для прогнозування автентичності запитуваного зразка почерку. Однак, цей тип методів має деякі недоліки, такі як необхідність витратити багато часу на вилучення та порівняння деталей графем. Крім того, такі методи вимагають значного обсягу пам'яті через велику кількість ознак графем, особливо для методів, які застосовують один алгоритм кластеризації.

Методи автоматичного навчання (на основі моделі) спираються на використання методів глибокого навчання для вилучення автоматично вивчених ознак глибокими моделями. Такі методи є надійними та показують високі показники ідентифікації, особливо для ідентифікації на рівні слів. Ще однією перевагою є те, що моделі автоматично навчаються представленням з даних (тобто автоматичне вилучення ознак), тому згорткові нейронні мережі широко використовуються для вилучення ознак із зображень (тобто слів, сторінок, ділянок). Основним недоліком методів на основі моделі є навчання, оскільки потрібна велика кількість зразків, навчання може бути дуже трудомістким, а моделі схильні до перенавчання.

Ще одним важливим недоліком методів на основі моделей є складність вибору найкращих значень для великої кількості параметрів. Крім того, цей тип методів вимагає дуже значних обчислювальних ресурсів та часу.

1.2 Огляд методів ідентифікації автора

Під час аналізу почерку для ідентифікації автора використовується метод індивідуальної ідентифікації особи, суть якого полягає у повному дослідженні індивідуальних характеристик документа, що досліджується, та порівнянні їх з відомими проявами цих характеристик в інших документах (порівняльний матеріал) [21].

Методи ідентифікації автора мають спільні/стандартні кроки попередньої обробки, вилучення ознак та ідентифікації автора (часто званої класифікацією), як показано на рис. 1.1.



Рисунок 1.1 – Спільні кроки методів ідентифікації автора

Сучасні методи ідентифікації авторів з високим рівнем розпізнавання на публічному наборі даних CVL з точністю Top-1 наведено в таблиці 1.1.

Таблиця 1.1 – Точність найсучасніших методів ідентифікації авторів з використанням набору даних CVL.

Автор	Метод	Точ (%)
Christlein et al. [22]	GMM super vector + E-SVM	99.4
Christlein et al. [23]	CNN + VLAD	99.5
Chen et al. [24]	ResNet-50	99.2
Helal et al. [25]	CNN	99.80
Sulaiman et al. [26]	CNN + LBP	97.55
Kumar and Sharma [27]	SEG WI model	99.35
Koepf et al. [28]	ViT	99.0
Semma et al. [29]	Resnet-34	99.5
He and Schomaker [30]	Adaptive CNN	79.1
He and Schomaker [31]	FragNet-64	99.1
He and Schomaker [32]	GR-RNN	99.3
Wirmanto et al. [33]	Xception+siamese	99.88

Дослідження в галузі ідентифікації автора також поділяються на чотири категорії методологій аналізу почерку.

У текстурній категорії Крістлайн та ін. [22] запропонували використовувати локальний дескриптор RootSIFT для фіксації локальних властивостей рукописних документів, використовувати супервектори GMM як метод кодування та використовувати методи опорних векторів Exemplar (SVM) для вимірювання подібності.

У категорії автоматичного навчання Крістлайн та ін. [23] запропонували використовувати згорткову нейронну мережу (CNN) для навчання потужного представлення ділянок, використовуючи членство в кластерах як цілі, а глобальний дескриптор зображення створюється за допомогою кодування Vector of Locally Aggregated Descriptors (VLAD). Чен та ін. [24] запропонували глибоку CNN (ResNet-50) та зважену регуляризацию згладжування міток

(WLSR) для доповнення даних, що дозволить вивчити більше розпізнавальних ознак для представлення властивостей різних стилів письма. Хелал та ін. [25] експериментували з CNN DL з додатковим підходом до відмінностей за допомогою класифікатора SVM. Сулейман та ін. [26] запропонували локальний бінарний шаблон (LBP) як ручний дескриптор ознак та структуру AlexNet для вилучення глибокого опису як другого дескриптора, обидва дескриптори потім зібралися в матрицю даних та закодувалися за допомогою VLAD. Модель без сегментації, представлена Кумаром та Шармою [27], де вони застосовують CNN для характеристики автора заданого зразка. Кепф та ін. [28] представили новий метод, заснований на зоровому трансформері (ViT) з класифікатором K-найближчих сусідів (KNN). Семма та ін. У роботі [29] представлена глибока нейронна мережа (CNN), навчена для захоплення глибоких ознак у невеликих ділянках, які були вилучені з рукописних зображень за допомогою ключових точок FAST та детектора Harris Corner (HC), а глибоко вивчені ознаки кодуються за допомогою VLAD та триангуляційного вбудовування з класифікатором KNN. Спочатку Хе та Шомакер [30] запропонували глибоку адаптивну CNN, засновану на багатозадачній адаптації структури AlexNet, після чого Хе та Шомакер [31] запропонували глибоку нейронну мережу (FragNet, натхненну Fraglets та BagNet) на основі фрагментів, сегментованих на вхідному зображенні документа, та пірамідальних карт ознак у CNN. У новішій роботі Хе та Шомакер [32] запропонували глобально-контекстну залишкову рекурентну нейронну мережу (GR-RNN), де просторовий зв'язок між послідовністю фрагментів моделюється рекурентною нейронною мережею (RNN) для посилення дискримінаційної здатності локальних ознак фрагментів, а також використовується додаткова інформація між глобально-контекстними та локальними фрагментами. Вірманто та ін. [33] запропонували сегментувати

дані на два входи та використовувати Xception як екстрактор ознак у сіамській мережі, останній генерує закодоване представлення та витягує ідентифікаційну мітку.

Сучасні методи ідентифікації авторів зазвичай вимірюють продуктивність на основі загальнодоступних наборів даних (наприклад, CVL, IAM тощо), а метрики, що використовуються для порівняння, - це точність Top-1 та точність Top-5. Де точність Top-1 означає, що очікуваний результат точно відповідає очікуваній відповіді (тобто відповіді з найвищою ймовірністю), а точність Top-5 означає, що очікуваний результат знаходиться серед перших п'яти відповідей з найвищою ймовірністю.

1.3 Аналіз існуючих моделей визначення особистості

Під час аналізу почерку з метою визначення психологічних характеристик особистості досліджуються відмінності письмових зразків від стандартного письма, і на основі цього визначається особистість особи, яка надає зразок письма [21].

Методології визначення особистості мають спільні/стандартні кроки попередньої обробки, вилучення ознак та визначення особистості (часто званого класифікацією), як показано на рис. 1.2.

Для проведення визначення особистості потрібно вибрати модель вимірювання психології особистості. Моделі особистості встановлюють теоретичну основу, визначаючи точні виміри, риси та метрики людської поведінки. З іншого боку, визначення особистості - це алгоритмічне виконання, воно використовує машинне навчання, обробку природної мови або поведінковий аналіз для аналізу необроблених даних (таких як текст, мовлення або динаміка натискання клавіш або зображення) та відображає їх

безпосередньо у встановлених психологічних категоріях. Тому, алгоритм визначення особистості без моделі для визначення цільових змінних не має сенсу.



Рисунок 1.2 – Спільні кроки методів визначення особистості

Дослідження в галузі визначення особистості показують широку різноманітність підходів, і вони, здебільшого, належать до категорії автоматичного навчання.

Отже, в категорії автоматичного навчання Гаврилеску [34] запропонував 3-рівневу архітектуру нейронної мережі з класифікацією ANN, SVM та KNN на 3-вимірному рівні. Полап та Возняк [35] запропонували гнучку нейронну мережу, де деякі нейрони адаптуються до цілей класифікації. Топалоглу та Екмекчі [36] запропонували використовувати методи дерева рішень та інтелектуального аналізу даних у поєднанні з алгоритмами ID3 та J48. Гаврилеску та Візірян [37] продовжили підхід 3-рівневої архітектури, і цей часовий базовий рівень включає нормалізацію та сегментацію, проміжний рівень генерує карту рукописного письма (НМ) за допомогою двійкового коду, верхній рівень виконує класифікацію за допомогою нейронної мережі прямого

зв'язку (FFNN). Джоші та ін. [38] розробили систему класифікації на основі SVM та використовували зіставлення шаблонів для вилучення літер. Віджая та ін. [39] порівнювали сегментацію лівого, правого, верхнього та нижнього полів та використовували SVM для класифікації. Фатіма та ін. [40] запропонували класифікацію на основі CNN з методами оптимізації SGD та AdaDelta. Чітлангія та Малаті [41] запропонували виконувати вилучення ознак за допомогою техніки гістограми орієнтованого градієнта (HOG) та класифікацію за допомогою SVM. Томас та ін. [42] зібрали тексти, написані десятима суб'єктами, використовуючи модель CNN та логістичну регресію, та вибрали сім ознак для завдання класифікації. Патхак та ін. [43] запропонували архітектуру глибокої нейронної мережі на основі CNN з довгою короткочасною пам'яттю (LSTM) та конекціоністською часовою класифікацією (CTC). Елнгар та ін. [44] запропонували нову модель виявлення рівня рис особистості (PTLDM), де для класифікації використовуються Мережа аналізу особистості (PAN) та PersonaNet. Бернардо та ін. [45] запропонували як екстрактор ознак використовувати модель легкої архітектури SqueezeNet, оптимізовану методом оптимізації гіперпараметрів, та SVM для класифікації. Рахман та ін. [46] запропонували вилучати ознаки за допомогою підходу представлення письма на основі графів та виконувати класифікацію за допомогою напівконтрольованої генеративно-змагальної мережі (SGAN). Самсуряді та ін. [47] провели складне порівняння підходів SVM (три варіації ядра: лінійний, RBF та поліноміальний), KNN та дерева рішень.

Методи визначення психологічних характеристик особистості з літератури з високим рівнем розпізнавання на наборах даних з метрикою точності наведено в таблиці 1.2.

Таблиця 1.2 – Методи визначення особистості в літературі.

Автор	Метод	Точ (%)
Gavrilescu [34]	ANN, SVM, and KNN	88.6
Polap and Wozniak [35]	Flexible neural network	93
Topaloglu and Ekmekci [36]	Decision tree	93.75
Gavrilescu and Vizireanu [37]	FFNN and template-matching	84.4
Joshi et al. [38]	SVM and template-matching	97
Wijaya et al. [39]	SVM	82.73
Fatimah et al. [40]	CNN	97.31
Chitlangia and Malathi [41]	HOG+SVM	80
Thomas et al. [42]	CNN	65
Pathak et al. [43]	Deep NN	97.7
Elngar et al. [44]	ANN + PersonaNet	65
Bernardo et al. [45]	Hybrid two-stage SqueezeNet and SVM	91.26
Rahman et al. [46]	SSL, SGAN	91.30
Samsuryadi et al. [47]	SVM, KNN, and decision tree	99

Дослідники визначення особистості, на відміну від дослідників ідентифікації автора, зазвичай замість загальнодоступних наборів даних вимірюють продуктивність на основі приватних наборів даних про почерк, а метрики, що використовуються для порівняння - це точність, влучність, повнота, бал F1, істинно позитивний (TP), істинно негативний (TN), та помилки першого та другого роду хибнопозитивний (FP) та хибнонегативний (FN) результат.

1.4 Аналіз психометричних моделей особистості

Як було згадано раніше вибір психометричної моделі особистості є критично важливою складовою для визначення особистості, особливо в галузі графології. В рамках аналізу почерку, точність передбачення обмежується надійністю позначок базової істини, що надається обраною психологічною моделлю [48]. Це відбувається незалежно від того які методи використовуються традиційні методи вилучення ознак чи архітектури глибокого навчання. Якщо моделі бракує внутрішньої узгодженості або міжкультурної валідності, будь-які проведені паралелі між механічними особливостями, такими як швидкість штриха, тиск пера або нахил, та рисами особистості будуть невідтворюваними або незначними з боку статистики. Тому вибір моделі в якій показник психометричної стійкості високий є важливим для того щоб почерк міг вважатися біометричним маркером для психологічної оцінки [49].

П'ятифакторна модель (FFM або OCEAN), або Велика п'ятірка, залишається найширше прийнятою структурою в сучасній психології. Вона класифікує особистість на п'ять широких областей: Відкритість, Сумлінність, Екстраверсія, Співробітництво та Нейротизм. Дослідження постійно демонструють, що ці фактори стабільні протягом усього дорослого життя та мають високу спадковість [48]. У комп'ютерній лінгвістиці та дослідженнях комп'ютерного зору Велика п'ятірка часто використовується, оскільки її безперервні шкали ознак (на відміну від бінарних типів) дозволяють проводити більш нюансований регресійний аналіз та вищу деталізацію в позначеннях машинного навчання [50].

Спираючись на FFM, модель HEXACO вводить шостий вимір: Чесність-Скромність. Цей фактор враховує риси, пов'язані зі щирістю, справедливістю та скромністю, які часто змішували з приємністю в попередніх моделях [51].

Дослідження, показують, що HEXACO має чудову прогностичну силу щодо просоціальної та антисоціальної поведінки. Виділяючи фактор Чесність-Скромність, дослідники можуть краще розрізняти справжню кооперативну поведінку та стратегічні соціальні маніпуляції, що є особливо актуальним в організаційній психології [51].

Індикатор типу Майєрса-Бріггса (MBTI), заснований на теорії Юнга, класифікує людей на 16 різних типів за допомогою чотирьох дихотомій (наприклад, інтроверсія проти екстраверсії). Цей індикатор популярний у корпоративному середовищі, при цьому він стикається зі значною академічною критикою. Більшість людей отримують бали посередині спектра, а не на крайніх значеннях, що ставить під сумнів надійність бінарної класифікації [50]. Модель 16 особистостей намагається подолати цей розрив, використовуючи знайому номенклатуру MBTI, одночасно зіставляючи базову оцінку з методологією «Великої п'ятірки», що базується на рисах, фактично створюючи гібрид, який зберігає популярність, одночасно покращуючи психометричну відповідність [52].

Хоча «Велика п'ятірка» описує, ким є людина, такі моделі, як DiSC та CliftonStrengths, зосереджуються на тому, як людина діє або робить свій внесок у професійне середовище. DiSC (Домінування, Вплив, Стійкість та Сумлінність). Вони вважаються поведінковою моделлю, яка використовується для покращення комунікації в команді та вирішення конфліктів [53]. CliftonStrengths (раніше StrengthsFinder) базується на позитивній психології та має на меті визначити найкращі таланти людини серед 34 можливих тем. Ці моделі менш поширені в клінічних дослідженнях, проте вони вважаються корисними при співставленні між талантом та завданням, тому ці моделі згадуються в літературі про управління [53].

Еннеаграма - це типологічна система, яка описує дев'ять взаємопов'язаних типів особистості, мотивованих глибинними страхами та бажаннями. Через такі особливості як психодинамічне походження та відсутність стандартизованого емпіричного тестування Еннеаграма в академічних колах піддавалася сумнівам. Проте нещодавні дослідження в яких вивчалась надійність свідчать про те, що, вона може запропонувати унікальне розуміння основних «фіксацій на ego» та мотивацій, які можуть бути пропущені моделями, заснованими на рисах. Це можливо за рахунок модернізації за допомогою ретельного психометричного тестування [54].

1.5 Огляд методів нанесення цифрового водяного знака

Зі зростанням популярності контенту, що генерується штучним інтелектом, та цифрового обміну, водяні знаки стали важливими для забезпечення цілісності даних, перевірки автентичності та захисту інтелектуальної власності. [55]. Вирішення цих проблем є важливими в контексті аналізу зображень та для сфер криміналістики і медицини. Існують про активні та пасивні підходи. Пасивний підхід прагне виявити сліди маніпуляцій після факту, а активний намагається захистити зображення заздалегідь, наприклад під час створення. Невидиме цифрове водяне маркування стало основною про активною технологією для вирішення цих проблем. Воно вбудовує прихований, надійний сигнал у дані зображення в момент створення або під час безпечної передачі [56]. Чіткий ланцюг зберігання (CoC) це одна з основних безпекових вимог серйозних систем, особливо правових [55]. Вбудовування цифрових підписів у зображення та інші документи може служити для підтвердження їхньої автентичності та запобігання несанкціонованим змінам [57]. Завдяки такому вбудовуванню

унікальних ідентифікаторів або даних автентифікації в медіафайли, можна виявляти підробки, відстежувати джерело несанкціонованого розповсюдження та встановлювати походження цифрового контенту [55].

Встановлення пріоритетів для певних параметрів залежить від області застосування. Сучасні методи балансують між ємністю, якістю та стійкістю. Для практичних систем важливо забезпечити достатню ємність без втрати якості зображення та зберегти стійкість до спотворень під час передачі або обробки [58]. Наприклад методи які використовуються в медицині надають пріоритет непомітності та зворотності, щоб забезпечити клінічну точність і зберегти діагностичну цінність зображень [59].

Геометричні атаки, такі як обрізання та розділення, перешкоджають стандартному водяному знаку, спричиняючи десинхронізацію, якій протидіють передові методи, що використовують стабільність домену перетворення, надлишкове вбудовування та синхронізацію точок ознак [60].

Сучасні системи працюють у частотній або трансформаційній областях, таких як дискретне вейвлет-перетворення (DWT), дискретне косинусне перетворення (DCT) або дискретне перетворення Фур'є (DFT) [61].

Розкладання за сингулярними значеннями (SVD) є взірцем для надійного водяного знаку. Воно розкладає зображення на матриці, сингулярні значення яких, залишаються надзвичайно стабільними до шуму або геометричних зрушень. Вбудовуючи водяний знак у ці стабільні сингулярні значення, досягається високий ступінь стійкості. У таких схемах як DWT-SVD він часто спеціально спрямований на низькочастотний піддіапазон, оскільки він містить більшу частину енергії зображення та стійкий до видалення [60].

У захисті від геометричних атак, таких як обрізання значний прогрес було досягнуто за рахунок використання локальних детекторів ознак для синхронізації водяного знака. Зазвичай вбудовування сигналу робиться

відносно меж зображення, але ці методи прив'язують водяний знак до стабільних ключових точок. Такі точки визначаються за допомогою таких алгоритмів, як SIFT (Scale-Invariant Feature Transform) або SURF (Speeded Up Robust Features) [60]. Різні підходи до синхронізації водяних знаків на основі ознак, узагальнено разом із їхніми ключовими властивостями стійкості та сферами застосування і представлено в таблиці 1.3.

Таблиця 1.3 – Порівняння методів водяних знаків на основі ознак

Метод	Технічна реалізація	Первинна стійкість	Корисність
SIFT-Based	Прив'язка сигналів до точок, інваріантних щодо масштабування/обертання	Масштабування, обертання, переміщення (RST).	Перевірка авторських прав.
SURF-SVD	Виявлення ключових точок у поєднанні з частотним вбудовуванням.	Гаусівський шум, 30% обрізання.	Цілісність медичних зображень.
ASIFT-DWT	Афінний SIFT для багатокутної стійкості.	Складні афінні перетворення, нахили.	Юридичні докази з різних ракурсів.
LSFR-DFT	Локальні квадратні області ознак з модуляцією Фур'є.	Атаки з використанням Screen-cam, 50% обрізання.	Сліди витоків за фотографіями.

Афінно-масштабно-інваріантне перетворення ознак (ASIFT) є дуже цінним для судово-медичних застосувань. Це можливо завдяки його інваріантності до афінних перетворень, наприклад, нахилених знімків екрана. Це дозволяє системі зіставляти точки ознак, розраховувати геометричні

спотворення та повторно синхронізувати зображення перед вилученням водяного знака [60].

Розвиток технологій глибокого навчання дозволив розробити комплексні системи водяних знаків. В таких системах і кодер, і декодер є нейронними мережами, навченими оптимізувати тріаду водяних знаків [62]. Ці системи, такі як HiDDeN та StegaStamp, використовують архітектуру згорткової нейронної мережі (CNN) для вивчення оптимальної стратегії вбудовування для заданого набору потенційних атак [63]. Типова система нейронного водяного знаку складається з кодера, декодера та мережі-супротивника. В ній кодер вбудовує повідомлення в зображення, декодер використовується для відновлення повідомлення, мережа-супротивника для забезпечення непомітності, а шумовий шар застосовує диференційовані спотворення під час навчання, щоб забезпечити надмірне поширення сигналу та його стійкість до таких атак, як обрізання або розмиття, змушуючи кодер вбудовувати сигнал надлишково для надійності [63].

Неявні нейронні представлення (INR) переосмислюють зображення як безперервні функції просторових координат, встановлюючи парадигму цифрового водяного знаку, яка не залежить від роздільної здатності [64, 65]. Вбудовуючи корисні навантаження безпосередньо в безперервний домен INR, а не в дискретні піксельні сітки, ці фреймворки обходять традиційні обмеження пам'яті, забезпечуючи багатомасштабне водяне знакування з надвисокою роздільною здатністю з фіксованим обчислювальним обсягом [64]. Така, незалежна від координат, обробка забезпечує стійкість до серйозних геометричних спотворень, що дозволяє високоточно видобувати повідомлення навіть після екстремального обрізання або масштабування [64, 65]. Зрештою, притаманна безперервній вибірці координат масштабованість і надійність роблять архітектури INR дуже

ефективними для захисту масивних цифрових активів у великомасштабних програмах водяного знакування [66].

Традиційні нейронні водяні маркери зазвичай досягають граничної ємності 100-256 біт [58]. В судово-медичній практиці де використовуються криптографічні підписи або ланцюг збереження така ємність є недостатньою. Для вирішення цієї проблеми представили ChunkySeal, масштабована архітектура VideoSeal. Яка за рахунок великої кількості параметрів (1,8 млрд.) та геометрично високовимірного сітчастого представлення долає цю граничну ємність, досягаючи ємності 1024-біт [58]. При цьому зберігається висока стійкість до геометричних атак та стиснення зображень.

Зважаючи на проведений аналіз було сформоване детальне порівняння сучасних методів цифрового водяного знака та представлено у таблиці 1.4.

Таблиця 1.4 – Порівняння сучасних методів цифрового водяного знака

Метод	Атака на стійкість	Продуктивність
INR UHR	99.6% Обрізання (0.4% Область)	>98% точність
ASIFT-DWT-SVD	Обертання та масштабування	Високий NC (>0.95)
LSFR-DFT	50% Обрізання	BER < 0.1333
SURF-SVD-Hash	30% Обрізання	NC $\approx$ 1.00
HiDDeN	JPEG 50, Випадіння	95% Відновлення бітів
StegaStamp	Захоплення екрана	95% Відновлення бітів
ChunkySeal	Лінійні ускладнення	1024-bit Ємність

Проведений аналіз демонструє, що розглянуті методи цифрового водяного знака суттєво відрізняються за стійкістю до геометричних атак,

обсягом вбудовуваної інформації та обчислювальною складністю. Зважаючи на це вибір підходу до маркування зображень рукописного тексту потребує компромісу між стійкістю, ємністю та продуктивністю відповідно до вимог конкретної прикладної задачі.

1.6 Аналіз технологій безпечного зберігання даних

Перехід від фізичних до цифрових файлів створив значні проблеми щодо автентичності, цілісності та допустимості даних. Цифрові документи за своєю суттю є нестабільними та схильними до непомітної модифікації, тому окрім технологій маркування (як цифровий водяний знак), також важливим елементом захисту є спеціалізовані архітектури зберігання [67]. Як було сказано раніше безперервний, перевірений ланцюг зберігання є важливою безпековою вимогою і системи зберігання даних повинні також його забезпечувати, документуючи кожну дію. Сучасні архітектури все частіше використовують хмарне сховище об'єктів (Amazon S3), приватні шари, сумісні з S3 (MinIO), та децентралізовані реєстри для створення безпечних сховищ, які запобігають підробці, гарантують невідмовність та витримують складні кібератаки, такі як програми-вимагачі.

Одним з популярних методів зберігання даних є хмарне сховище об'єктів Amazon Simple Storage Service (S3) [68]. Його використання для безпечного управління файлами зосереджено на забезпеченні дотримання політик Write-Once-Read-Many (WORM), а також використання функції блокування об'єктів (Object Lock). Саме останній механізм захищає файли від несанкціонованого видалення або шифрування [68]. Коли файли або системні журнали, зберігаються у версійному бакеті з увімкненим блокуванням об'єктів, їх додатково можна перевести в «Режим відповідності». У цьому

режимі термін зберігання не можна скоротити, і навіть користувачам root-рівня заборонено видаляти версію об'єкта. Для складніших ситуацій можна застосовувати режим «Юридичні утримання». Він забезпечує безстроковий захист без фіксованої дати закінчення терміну дії, зберігаючи незмінність, доки утримання не буде видалено вручну уповноваженими особами [68]. Технічна архітектура хмарного репозиторію інформації інтегрує криптографічне хешування SHA-256 з надлишковою інфраструктурою S3. Кожен файл унікально ідентифікується своїм хешем SHA-256, який зберігається разом зі структурованими метаданими в форматі XML. Вони містять дату збору, місцезнаходження та інформацію про особу, яка його обробила. Основною перевагою цього підходу є його надзвичайна довговічність та можливість використовувати політики життєвого циклу S3 для переміщення старих, неактивних файлів на менш витратні рівні, такі як S3 Glacier Instant Retrieval, гарантуючи їхню доступність для мілісекундного пошуку [69].

Іншим провідним S3-сумісним рішенням для приватної хмари слугує MinIO [70]. Воно також може використовуватись для підтримки фізичного контролю над даними. Контроль включає юрисдикційні обмеження та обмеження конфіденційності. Модель безпеки MinIO базується на його інтеграції зі Службою шифрування ключів (KES). Вона забезпечує шифрування на стороні сервера (SSE) шляхом абстрагування взаємодії з системою керування ключами (KMS), такою як HashiCorp Vault [70]. Така архітектура гарантує, що на рівні сховища обробляються лише зашифровані потоки даних. При цьому головні криптографічні секрети залишаються ізольованими в окремому, захищеному сховищі. MinIO реалізує багаторівневий процес отримання ключів, де кожен об'єкт шифрується унікальним ключем шифрування об'єкта (ОЕК). Він сам по собі захищений

ключем шифрування ключа (KEK), отриманим з головного ключа [71]. Основною перевагою MinIO є його величезна пропускна здатність. Швидкість читання до 183 ГБ/с на стандартному обладнанні, що є важливим для швидкого отримання та обробки багатотерабайтної фізичної пам'яті або дамів жорсткого диска [71]. Також MinIO підтримує ті ж API блокування об'єктів S3, що дозволяє організаціям впроваджувати локальне сховище WORM для своїх файлів. Недоліком підходу базованого на MinIO є високі операційні витрати. Впроваджувач відповідає за управління станом обладнання, ротацією дисків та кластеризацією KMS з високою доступністю [71]. Неправильна синхронізація вузлів KES може призвести до несинхронізованих записів, що потенційно може поставити під загрозу цілісність Ланцюга зберігання під час критичної фази отримання [71].

Розробки в сфері базуються на фреймворках B-CoC (ланцюг відповідальності на основі блокчейну), які дематеріалізують журнал [72]. У цих системах фактичні файли знаходяться в Міжпланетній файловій системі (IPFS), бо вони занадто великі для зберігання в ланцюзі. При цьому їхні криптографічні хеші та журнали доступу записуються в незмінному реєстрі, такому як Ethereum або Hyperledger Fabric [72]. Це створює перевірений журнал аудиту, де будь-яка несанкціонована модифікація файлів на вузлі зберігання призводить до невідповідності хешу із записом блокчейну. Наприклад фреймворк ForensiLock використовує смарт-контракти для забезпечення контролю доступу на основі ролей (RBAC) [67]. Це гарантує, що лише уповноважені особи можуть генерувати токени, необхідні для отримання файлів з IPFS. Основною перевагою архітектури Blockchain+IPFS є усунення єдиних точок відмови та зменшення внутрішніх загроз. Наприклад скомпрометований адміністратор не може змінити історичні журнали або стерти запис про те, хто мав доступ до певного файлу. Технічним недоліком

вважається затримка отримання даних. Вона збільшується приблизно на 1,8 секунди через консенсус блокчейну та виявлення вузлів IPFS, а також складність управління закритими ключами для всіх осіб-учасників. Крім того, публічний характер деяких розподілених хеш-таблиць (DHT) в IPFS може розкрити метадані файлів, якщо їх не поєднати з надійним шифруванням на рівні проксі-сервера.

Традиційні методи захисту даних зосереджені на зберіганні та передачі. Проте під час аналізу дані стають вразливими, для вирішення цієї проблеми використовують Надійні Середовища Виконання (TEE), такі як Intel SGX або TDX. Вони забезпечують ізоляцію на апаратному рівні, створюючи анклав. Анклави являють собою зашифровані області пам'яті, недоступні навіть для операційної системи хоста або скомпрометованого гіпервізора. Це дозволяє виконувати певні операції з файлами, без їх доступу до основної оперативної пам'яті системи [73]. Такі операції включають пошук за ключовими словами у величезних наборах даних або класифікацію зображень на основі штучного інтелекту [73]. Це також актуально для збереження конфіденційності в хмарних середовищах [74]. Впровадження TEE дозволяє здійснювати віддалену атестацію. Атестація це механізм який створює криптографічний звіт підтверджуючий, що певний фрагмент коду був виконаний точно так, як передбачалося, в безпечній області [73]. Перевагою є значне зменшення бази довірених обчислень (TCB), що захищає файли від розширеного шкідливого програмного забезпечення рівня ядра. Недоліком вважається обмежена пам'ять. Наприклад TEE першого покоління від SGX може мати 128 МБ, що може призвести до суттєвого зниження продуктивності під час обробки зображень високої роздільної здатності. При цьому новіші TEE на базі віртуальних машин, такі як Intel TDX, починають зменшувати ці обмеження ємності [73].

Детальне порівняння сучасних методів безпечного зберігання представлено у таблиці 1.5.

Таблиця 1.5 – Порівняння методів безпечного зберігання

Метод	Механізм	Переваги	Недоліки
AWS S3 Object Lock	WORM	Сертифіковане хмарне сховище	Ціноутворення
MinIO + KES/KMS	Розділене керування ключами	Внутрішній контроль даних	Складність управління
B-CoC (Blockchain)	Незмінний розподілений реєстр	Високий захист від несанкціонованого доступу	Затримка та масштабованість
Intel SGX/TDX	Апаратні захищені анклави	Захист даних під час використання	Обмеження пам'яті

Проведений аналіз демонструє, що жоден із розглянутих методів безпечного зберігання не є універсальним. Кожен з методів забезпечує власний баланс між рівнем захисту, продуктивністю та операційними витратами, тому вибір конкретного рішення має визначатися вимогами до незмінності даних і безперервності ланцюга зберігання у відповідній прикладній задачі.

1.7 Постановка задач дослідження

На основі аналізу сучасного стану досліджень у сфері графологічного аналізу рукописного тексту можна зробити висновок щодо актуальності

дослідження такого аналізу та необхідності подальших досліджень з метою підвищення його ефективності.

Метою дисертаційної роботи є розробка інформаційної технології автоматизованого графологічного аналізу рукописного тексту, яка використовує сучасні технології в області ідентифікації автора та визначення особистості для створення нової технології графологічного аналізу зображень рукописного тексту. Для досягнення поставленої мети необхідно вирішити наступні завдання:

- провести аналіз сучасного стану досліджень у сфері аналізу рукописного тексту;
- удосконалити метод ідентифікації та пошуку автора на основі зображень рукописного тексту;
- розробити метод визначення психологічних характеристик особистості за рукописним текстом;
- розробити модель інформаційної технології автоматизованого аналізу рукописного тексту;
- провести експериментальну оцінку результатів дослідження.

Вирішенню поставлених завдань, що породжуються в процесі розробки моделі та методів інформаційної технології аналізу рукописного тексту, присвячені наступні розділи.

1.8 Висновки за першим розділом

У даному розділі дисертації був проведений системний аналіз та розглянуто сучасний стан досліджень у сфері графологічного аналізу рукописного тексту та підходи і технології:

1. Розглянуто основні процеси та особливості графологічного аналізу рукописного тексту, його сфери застосування та ключові проблеми. За результатами аналізу методи аналізу почерку класифіковано на чотири групи (структурні, текстурні, графемні та на основі автоматичного навчання) і встановлено, що методи глибокого навчання забезпечують найвищі показники розпізнавання.

2. Здійснено огляд та аналіз сучасних методів ідентифікації автора, визначено спільні етапи їх роботи та продемонстровано порівняння точності на публічних наборах даних. Встановлено, що найкращі результати забезпечують методи глибокого навчання, які досягають точності Top-1 понад 99%, а перспективним напрямом визначено застосування архітектури зорового трансформера.

3. Здійснено огляд та аналіз сучасних методів визначення психологічних характеристик особистості за рукописним текстом. Встановлено, що такі методи переважно належать до категорії автоматичного навчання та оцінюються на приватних наборах даних за відсутності єдиного підходу, а їхня точність варіюється в широкому діапазоні, що свідчить про суттєву залежність результату від обраної психометричної моделі та якості розмітки даних.

4. Проаналізовано психометричні моделі за критеріями психометричної стійкості та придатності до машинного навчання.

5. Розглянуто методи забезпечення цілісності цифрових зображень за допомогою цифрових водяних знаків, що поділяються на частотні, на основі інваріантних ознак та нейромережеві. Встановлено, що методи на основі інваріантних ознак забезпечують стійкість до геометричних атак, нейромережеві методи досягають близько 95% відновлення вбудованих бітів, а обґрунтованим вибором для захисту зображень рукописного тексту є невидиме активне водяне маркування.

6. Проаналізовано технології безпечного зберігання даних, що гарантують незмінність інформації та безперервний ланцюг зберігання.

7. Сформульовано мету та задачі дисертаційного дослідження. За результатами аналізу сучасного стану обґрунтовано актуальність роботи та визначено п'ять завдань: аналіз сучасного стану досліджень у сфері аналізу рукописного тексту; удосконалення методу ідентифікації та пошуку автора; розробка методу визначення психологічних характеристик особистості; розробка моделі інформаційної технології автоматизованого аналізу; експериментальна оцінка отриманих результатів.

Список використаних джерел у даному розділі наведено в повному списку використаних джерел під номерами: 1-73.

2 НАБУВ ПОДАЛЬШОГО РОЗВИТКУ МЕТОД ВИЗНАЧЕННЯ ПСИХОЛОГІЧНИХ ХАРАКТЕРИСТИК ОСОБИСТОСТІ

У другому розділі набув подальшого розвитку метод визначення психологічних характеристик особистості на основі рукописного тексту. Основну увагу приділено аналізу існуючих методів визначення особистості, формальній постановці задачі та обґрунтуванню вибору психометричної моделі «Велика п'ятірка». Запропоновано підхід до визначення особистості за допомогою методів машинного навчання, який використовує Зоровий трансформер. У процесі розробки побудовано функціональну модель визначення особистості, що включає етапи попередньої обробки вхідних даних та виділення графологічних ознак з використанням масштабно-інваріантного ознакового перетворення. Окрему увагу приділено експериментальному дослідженню методу, яке продемонструвало високу точність класифікації (до 99,48%). Запропонована модель створює основу для підвищення надійності автоматизованих систем аналізу рукописного тексту та є важливим внеском у підвищення точності визначення особистості.

2.1 Постановка задачі визначення психологічних характеристик особистості

Під час аналізу почерку для визначення особистості досліджується відмінність письмових зразків від стандартного письма і на основі цього визначається особистість особи, яка надає зразок письма [21]. Дослідження в області визначення особистості включають широкий спектр підходів оснований на різних методах.

Аналізуючи дослідження в області визначення особистості було виявлено, що в своєму дослідженні Гаврилеску [34] запропонував 3-рівневу архітектуру нейронної мережі з класифікацією ANN, SVM і KNN на 3-ому рівні, яка продемонструвала точність в 86,7%, в якості типів особистості використовувався індикатор типів Маєрс-Брігса, в якості набору даних був обраний приватний набір, який налічував 128 суб'єктів.

Об'єктивні труднощі, пов'язані з відсутністю існуючих наборів даних, складною системою вилучення ознак та низькопродуктивною нейронною мережевою системою як агрегатором ознак, призвели до того, що в [37] підхід до 3-рівневої архітектури був удосконалений наступним чином:

- базовий шар включає нормалізацію та сегментацію;
- проміжний шар генерує карту рукописного тексту за допомогою двійкового коду;
- верхній шар виконує класифікацію за допомогою прямої нейронної мережі (ПМП).

Результати цього підходу показали точність 84,4%. Як типи особистості було використано індикатор типу «Велика п'ятірка», а як набір даних було обрано приватний набір із 128 суб'єктів. Однак, питання, пов'язані з перенавчанням та необхідністю ручної кореляції графологічних ознак з ознаками індикатора типу, залишилися невирішеними. Причиною перенавчання може бути як погано розбитий набір даних, так і невідповідна модель, яка виконує класифікацію. Варіантом подолання відповідних труднощів може бути використання спеціалізованих методів класифікації.

Саме цей метод було використано в [38], де автори розробили структуру класифікації на основі SVM та використали зіставлення зі зразками для вилучення літер, що продемонструвало точність до 97%. Як типи особистості використовувався спеціальний набір рис для оцінки працевлаштування

кандидата з точки зору HR, а як набір даних було обрано приватний набір, який включав 1890 зразків зображень почерку. Однак результати свідчать про дуже вузьку спрямованість дослідження, спрямовану на визначення рівня працевлаштування на основі обмеженого набору зображень.

Однак у роботі [39] з аналогічним методом класифікації автори порівнювали сегментацію лівого, правого, верхнього та нижнього полів та використали SVM для класифікації. В результаті було продемонстровано точність 82,73%, у роботі не використовувалися типи особистості, а як набір даних було обрано приватний набір, який включав 42 суб'єкти.

У [40] було запропоновано класифікацію на основі згорткових нейронних мереж (CNN) з методами оптимізації SGD та AdaDelta, і було продемонстровано точність 82,5%. Цього було досягнуто завдяки тому, що замість типів особистості тестувалися лише графологічні ознаки, такі як поля, міжрядковий інтервал, міжсловний інтервал та нахил слів. Як набір даних було обрано приватний набір із 128 суб'єктів.

У цих двох роботах питання доцільності використання таких підходів, як CNN та SVM, для етапу вилучення ознак залишилося невирішеним. Незважаючи на те, що вони можуть впоратися з цим завданням, для навчання вони вимагають великих наборів даних. Слід також зазначити, що в цих роботах відсутній компонент ідентифікації особистості, що робить відповідні дослідження непридатними для порівняння в цій роботі.

Варіантом використання більш доцільних технологій для вилучення ознак є робота [41], в якій автори запропонували виконувати вилучення ознак за допомогою методу гістограмно-орієнтованого градієнта (HOG) та класифікувати за допомогою SVM; була продемонстрована точність 80%. Як типи особистості було використано 5 рис особистості, а саме: Енергійність, Екстраверт, Інтроверт, Неспокійність та Оптимізм, а як набір даних було

обрано приватний набір із 50 суб'єктів. Хоча метод гістограмно-орієнтованого градієнта широко використовується в обробці зображень та комп'ютерному зорі для виявлення об'єктів, він має низку недоліків, таких як чутливість до шуму та артефактів на зображеннях. Якщо розглядати цей підхід у контексті розпізнавання тексту, то до недоліків можна віднести нечутливість до незначних змін через його зосередженість на сильних градієнтах. Варіантом подолання відповідних труднощів може бути зіставлення з шаблонами, що саме й використовується в [43].

У [42] автори використали модель CNN та логістичну регресію, і обрали сім ознак для завдання класифікації, і було продемонстровано точність 65%. Як тип особистості використовувався показник типу особистості «Великої п'ятірки», а як набір даних було обрано приватний набір із 112 зразків зображень почерку, написаних десятима суб'єктами.

У [43] автори запропонували архітектуру глибокої нейронної мережі на основі CNN з довгою короткостроковою пам'яттю (LSTM) та конекціоністською часовою класифікацією (СТС); і була продемонстрована точність 97,7%. У дослідженні замість типів особистості тестувалися лише графологічні ознаки, такі як поля, міжрядковий інтервал, міжсловний інтервал, нахил слів, базова лінія, тиск пера та інші. Як набір даних було обрано публічний набір даних IAM, який включав 1539 зразків почерку, написаних 657 суб'єктами. Однак питання, пов'язані з відображенням графологічних ознак на певні риси, залишалися невирішеними. Варіантом подолання цих труднощів може бути використання індикатора типу, що є підходом, використаним у [42, 44, 46].

У [44] автори запропонували нову модель визначення рівня рис особистості (PTLDM), де для класифікації використовується мережа аналізу особистості (PAN) та PersonaNet; була продемонстрована точність 65%. Як тип

особистості використовувався показник типу особистості «Велика п'ятірка», а як набір даних обрано приватний набір із 125 суб'єктів. Результати показують, що представлені моделі глибокого навчання гірше справляються з класифікацією, ніж інші методи. Причиною цього може бути те, що моделі глибокого навчання можуть бути обчислювально дорогими та складними для навчання, а розмір та якість набору даних впливають на продуктивність моделі. Крім того, традиційні архітектури глибокого навчання, такі як CNN, можуть не оптимально підходити для виявлення тонких, але важливих особливостей почерку. У цьому контексті дослідження трансформерів зору, які завдяки механізму уваги здатні ефективніше моделювати зв'язки між різними частинами тексту, може виявитися перспективним. Іншим варіантом подолання труднощів може бути використання методів класифікації, таких як напівконтрольоване навчання з генеративно-змагальною мережею.

Саме такий підхід було використано в [46]; автори запропонували витягувати ознаки за допомогою графового підходу письмового представлення та виконувати класифікацію за допомогою напівконтрольованої генеративно-змагальної мережі (SGAN), що дозволило досягти точності 91,3%. Як типи особистості використовувався індикатор типу «Велика п'ятірка». А як набір даних було обрано приватний набір із 173 суб'єктів. Однак, витяг ознак на основі графів є проблематичним через свою обчислювальну складність та чутливість до варіацій стилю почерку, що призводить до неточних або ненадійних результатів. Крім того, створення графової структури може призвести до втрати частини інформації, присутньої в оригінальному почерку, а для певних завдань може знадобитися складне налаштування параметрів та функцій, що робить пошук альтернативи актуальним. З точки зору класифікації, питання навчання моделі залишається складним; SGAN було навчено на маркованих, немаркованих та згенерованих

генератором даних. Це може бути пов'язано з конкурентним характером процесу навчання та подвійною природою моделі. Традиційне машинне навчання може бути варіантом подолання цих проблем, де моделі, як правило, простіші, навчаються виключно на маркованих даних і часто забезпечують кращу продуктивність для завдань класифікації. Однак існують також принципово нові підходи, такі як Зоровий трансформер, який, хоча й належить до глибокого навчання, суттєво відрізняється від традиційних CNN. Зоровий трансформер має потенціал поєднати переваги глибокого навчання з ефективністю вилучення важливих ознак почерку, що робить його перспективною галуззю для подальших досліджень.

Застосування Зорового трансформеру до аналізу почерку може відкрити нові можливості для вилучення ознак та класифікації. Завдяки своїй здатності моделювати довгострокові залежності, Зоровий трансформер потенційно може виявляти складні закономірності в почерку, які не можуть виявити звичайні методи. Зоровий трансформер здатний навчатися на великих наборах даних і розглядає зображення як послідовність плям. Він також використовує механізм уваги для моделювання зв'язків між плямами, що дозволяє фіксувати як локальні, так і глобальні ознаки, що може призвести до більш точного визначення рис особистості.

На основі аналізу можна побачити, що деякі дослідження [39, 40, 43] обмежувалися виділенням графологічних ознак і не використовували у своїх дослідженнях показники типів особистості. Це свідчить про необхідність більш комплексних підходів до дослідження з використанням певних показників типів особистості.

Два дослідження [34, 41], які використовували показник типу Майєрса-Бріггса та 5 рис особистості (Енергійність, Екстравертність, Інтровертність,

Сумлінність та Оптимізм), призвели до посередньої точності в діапазоні від 80% до 87%.

Якщо розглянути дослідження [37, 42, 44, 46], які використовували індикатор типу особистості «Велика п'ятірка», можна побачити, що ця область все ще має потенціал для вдосконалення. Про це свідчить той факт, що отримана точність ідентифікації особистості все ще нижче 92%.

З точки зору виявлення ознак, можна побачити різноманітність вибору графологічних ознак для проведення досліджень та неоднозначність порівняння цих ознак з показниками типів особистості, що породжується в результаті суб'єктивного вибору графологічних ознак.

Таким чином, поєднання цих проблем та недоліків вказує на необхідність додаткових досліджень у галузі визначення психологічних характеристик особистості на основі рукописного тексту. Все це дає підстави стверджувати, що доцільно проводити дослідження визначення особистості з використанням певної метричної моделі особистості, методів машинного навчання та автоматизованого виявлення ознак.

Тому було запропоновано підхід до визначення особистості на основі рукописного тексту в якому використовується публічний набір даних, використовується індикатор типів особистості Велика п'ятірка, використовується підхід попередньої обробки з масштабно-інваріантним ознаковим перетворенням, а також використовується метод машинного навчання Зоровий трансформер.

Метою даної частини дослідження є розробка підходу до визначення психологічних характеристик особистості на основі рукописного тексту за допомогою методів машинного навчання. Підвищення точності визначення психологічних характеристик особистості та автоматизація виявлення ознак дасть можливість виконувати більш точний аналіз рукописного тексту при

визначенні особистості. В свою чергу, це підвищить надійність автоматизованих систем аналізу рукописного тексту в області визначення особистості, що допоможе графологам та експертам з почерку під час їх роботи.

Для досягнення мети були поставлені наступні завдання:

- розробити функціональну модель визначення психологічних характеристик особистості за допомогою методів машинного навчання;
- провести експериментальне дослідження, щодо запропонованого підходу.

2.2 Підхід визначення психологічних характеристик

Для комплексного підходу до визначення психологічних характеристик особистості важливим залишається вибір психометричних моделей особистості. У цьому дослідженні було використано індикатор типу особистості «Велика п'ятірка». Дослідження [75], яке порівнювало психологічні тести для оцінки особистості, показало, що, незважаючи на те, що алгоритми, навчені на інших тестах, можуть мати кращу ефективність, риси «Великої п'ятірки» є набагато інформативнішими та мають більшу варіабельність у продуктивності залежно від алгоритму. Також в [76] показано, що «Велика п'ятірка» виявилася надійним провісником критичних життєвих результатів (продуктивність праці, академічні досягнення та фізичне здоров'я). Крім того, було підтверджено численними незалежними дослідницькими групами та масштабними метааналізами, що «Велика п'ятірка» демонструє стабільність у різних демографічних групах та контекстах, на відміну від багатьох типологій особистості, яким бракує послідовних факторних структур [76].

В основі «Великої п'ятірки» лежать п'ять вимірів. Кожен з цих вимірів налічує два підвиміри. Розглянемо кожен з вимірів окремо.

Відкритість до досвіду стосується здатності та схильності людей досліджувати та створювати новий досвід, проявляючи допитливість, уяву, креативність, цінування мистецтва та нестандартних ідей. Два підвиміри це Інтелект та Відкритість. Інтелект відображає схильність до участі в абстрактній та інтелектуальній інформації, тоді як Відкритість відображає схильність до участі в естетичній та сенсорній інформації (сприйняття та уява). За цією шкалою люди оцінюються відповідно до дихотомії: винахідливі/допитливі (вищий бал) та послідовні/обережні (нижчий бал).

Сумлінність стосується людини, яка є обережною, вдумливою, відповідальною, організованою, працювитою, успішною та наполегливою. Вторинні виміри сумлінності це Працьовитість та Упорядкованість. Працьовитість стосується здатності до безперервної, цілеспрямованої діяльності, що пов'язано з продуктивністю та професійною етикою, тоді як упорядкованість стосується схильності до впорядкованості, організації та систематизації, що пов'язано з такими якостями, як охайність та старанність. За цією шкалою люди оцінюються відповідно до дихотомії: ефективні/організовані (вищі бали) проти легких на підйом/недбайливих (нижчі бали).

Екстраверсія стосується людей, які легко виражають позитивні емоції, є товариськими, рішучими, балакучими та активними. Два підвиміри це ентузіазм та асертивність. Ентузіазм означає дружелюбність, комунікабельність та схильність до позитивного впливу, тоді як асертивність відображає схильність до активності, рушійної сили та соціального домінування. За цією шкалою люди оцінюються відповідно до дихотомії: товариські/енергійні (вищі бали) проти самотніх/замкнутих (нижчі бали).

Співробітництво стосується людини, яка схильна бути співчутливою, довірливою та готовою допомогти, гнучкою та толерантною. Підвимірами співробітництва є Співчуття та Ввічливість. Співчуття стосується схильності емоційно піклуватися про інших, тоді як Ввічливість стосується схильності виявляти гарні манери, дотримуватися соціальних норм та уникати агресії. За цією шкалою люди оцінюються відповідно до дихотомії: доброзичливі/співчутливі (вищі бали) проти звинувачень/відстороненості (нижчі бали).

Нейротизм стосується людей, яким бракує стабільності та контролю над своїми емоціями, які легко емоційні, тривожні, депресивні, гнівні, стурбовані та невпевнені. Підвимірами нейротизму є Замкнутість та Нестабільність. Замкнутість стосується схильності до негативних внутрішніх впливів (гальмування), тоді як нестабільність стосується емоційної нестабільності, труднощів у контролі емоційних імпульсів та чутливості до негативних зовнішніх впливів (розгальмування).

Описи рис особистості «Великої п'ятірки» та пов'язаних з ними графологічних особливостей підсумовано в наступних абзацах.

Особи з високим рівнем Відкритості до досвіду часто демонструють висхідну базову лінію, що свідчить про оптимізм та амбіції. Їхній стиль письма може мати нахил, що свідчить про незалежність та оригінальність, а також широкий міжрядковий інтервал та широкий міжсловний інтервал, що свідчить про креативність, потребу в просторі та незалежне мислення. Крім того, їхній натиск пера часто легкий, що свідчить про чутливість, а розмір літер великий або змінний, що свідчить про креативність. Приклад рукописного тексту з набору CVL, який було визначено як Відкритість до досвіду зображено на рис. 2.1.

Imagine a vast sheet of paper on which straight Lines, Triangles, Squares Pentagons, Hexagons and other figures, instead of remaining fixed in their places, move freely about, on or in the surface, but without the power of rising above or sinking below it, very much like shadows - only hard and with luminous edges - and you will then have a pretty correct notion of my country and countrymen. Alas, a few years ago. I should have said "my universe": but now my mind has been opened to higher views of things.

Рисунок 2.1 – Зображення тексту з набору позначене як Відкритість

Сумлінність підтверджується високорегульованими рисами почерку. Пряма базова лінія вказує на самодисципліну та організованість, тоді як прямий нахил свідчить про практичність та контроль. Ця риса пов'язана з нормальними характеристиками за міжрядковим інтервалом, міжсловним інтервалом, натиском пера та розміром літер, що свідчить про впорядкованість, дотримання норм, баланс, контроль та практичність. Приклад рукописного тексту з набору CVL, який було визначено як Сумлінність зображено на рис. 2.2.

Imagine a vast sheet of paper on which straight lines, Triangles, Squares, Pentagons, Hexagons, and other figures, instead of remaining fixed in their places, move freely about, on or in the surfaces, but without the power of rising above or sinking below it, very much like shadows - only hard and with luminous edges - and you will then have a pretty correct notion of my country and countrymen. Alas, ~~at~~ a few years ago, I should have said "my universe": but now my mind has been opened to higher views of things

Рисунок 2.2 – Зображення тексту з набору позначене як Сумлінність

Особливості почерку, пов'язані з екстраверсією, включають висхідну базову лінію, що свідчить про ентузіазм та впевненість. Нахил передбачає відкритість та товариськість. Міжрядковий інтервал зазвичай нормальний або широкий, що свідчить про товариську натуру, а нормальний міжрядковий інтервал свідчить про легкість спілкування. Ця риса часто включає сильний натиск пера, що свідчить про наполегливість, та великий розмір літер, що свідчить про виразність. Приклад рукописного тексту з набору CVL, який було визначено як Екстраверсія зображено на рис. 2.3.

Imagine a vast sheet of paper on which straight lines, Triangles, Squares, Pentagons, Hexagons, and other figures, instead of remaining fixed in their places, more freely ~~about~~ ^{about}, on or in the surface, but without the power of rising above or sinking below it, very much like shadows – only hard and with luminous edges – and you will then have a pretty correct notion of my country and countrymen. Also, a few years ago, I should have said "my universe": but now my mind has been opened to higher views of things.

Рисунок 2.3 – Зображення тексту з набору позначене як Екстраверсія

На Співробітництво вказує пряма або злегка висхідна базова лінія, що свідчить про щирість, та прямий або злегка похилий нахил, що свідчить про емпатію. Нормальний міжрядковий інтервал свідчить про співпрацю, а нормальний або злегка широкий міжрядковий інтервал свідчить про толерантність. Тиск пера часто легкий або нормальний, що свідчить про ніжність, а розмір літер нормальний або злегка закруглений, що свідчить про теплоту. Приклад рукописного тексту з набору CVL, який було визначено як Співробітництво зображено на рис. 2.4.

Ознаки нейротизму включають спадну базову лінію, що свідчить про песимізм або втому, та нерівний нахил, що свідчить про емоційну нестабільність.

Imagine a vast sheet of paper on which straight Lines, Triangles, Squares, Pentagons, Hexagons, and other figures, instead of remaining fixed in their place, move freely about, on or in the surface, but without the power of rising above or sinking below it, very much like shadows - only hard and with luminous edges - and you will then have a pretty correct notion of my country and countrymen. Alas, a few years ago, I should have said "my universe": but now my mind has been opened to higher views of things.

Рисунок 2.4 – Зображення тексту з набору позначене як Співробітництво

Письменник може використовувати вузький міжрядковий інтервал, що свідчить про тривогу або напругу, та демонструвати нерівний міжрядковий інтервал, що свідчить про перепади настрою. Тиск пера часто змінний, що свідчить про коливання емоцій, а розмір літер маленький або тісний, що свідчить про невпевненість. Приклад рукописного тексту з набору CVL, який було визначено як Нейротизм зображено на рис. 2.5.

Imagine a vast sheet of paper on which straight
 lines, Triangles, Squares, Pentagons, Hexagons, and other
 figures, instead of remaining fixed in their places,
 move freely about, on or in the surface, but without
 the power of rising above or sinking below it,
 very much like shadows, — only hard and with
 luminous edges — and you will then have a pretty
 correct notion of my country and countrymen.
 And, a few years ago, I should have said
 "my universe": but now my mind has been
 opened to higher views of things.

Рисунок 2.5 – Зображення тексту з набору позначене як Нейротизм

Основні графологічні ознаки для кожної з п'яти рис особистості за моделлю «Великої п'ятірки» узагальнено у скороченому форматі та представлені в таблиці 2.1. В цій таблиці, для кожної риси з індикатора типу особистості наведено характерні значення базової лінії, нахилу письма, міжрядкового й міжсловного інтервалів, натиску пера та розміру літер. Це дозволяє звести описані вище співвідношення до єдиної структури ознак і використовувати її як основу для формування розмітки навчального набору даних.

Таблиця 2.1 – Відношення психологічних рис та графологічних ознак

Риса	Базова лінія	Нахил	Міжрядковий інтервал	Міжсловний інтервал	Натиск пера	Розмір літер
Відкритість	Висхідна	Нахилене	Широкий	Широкий	Легкий	Великий або змінний
Сумлінність	Пряма	Пряме	Норм.	Норм.	Норм.	Норм.
Екстраверсія	Висхідна	Похиле	Норм. або Широкий	Норм.	Сильний	Великий
Співробітництво	Пряма або Злегка Висхідна	Прямий або Злегка похиле	Норм.	Норм. або Злегка Широкий	Легкий або Норм.	Норм. або Злегка закругленим
Нейротизм	Спадова	Нерівне	Вузький	Нерівний	Змінний	Мал. або тісний

Вибір психометричної моделі особистості та формування співвідношення формує основу для подальшої побудови методу автоматизованого визначення психологічних характеристик особистості за рукописним текстом.

2.3 Метод визначення психологічних характеристик особистості

Побудова системи аналізу рукописного тексту для визначення психологічних характеристик особистості ґрунтується на тому, що під час аналізу почерку для визначення особистості досліджується відмінність письмових зразків від стандартного письма і на основі цього визначається особистість особи, яка надає зразок письма.

Проаналізувавши підходи та методи аналізу рукописного тексту для визначення психологічних характеристик особистості використані в літературі, можна зробити висновок, що підхід до побудови моделі визначення особистості за допомогою методу машинного навчання повинен включати в себе наступні кроки:

1. Попередня обробка вхідних даних;
2. Виділення графологічних ознак рукописного тексту;
3. Визначення особистості на основі ознак письмових зразків рукописного тексту;

Загальна схема підходу до визначення особистості на основі рукописного тексту за допомогою методу машинного навчання побудована на цих кроках зображена на рис. 2.6.

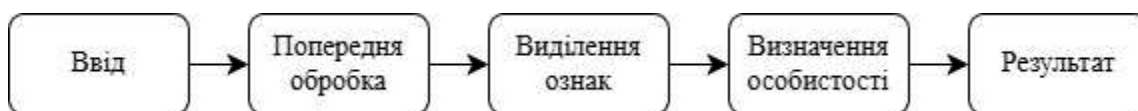


Рисунок 2.6 – Загальна схема підходу визначення особистості

На базі цієї загальної схеми була побудована функціональна модель визначення особистості за допомогою методів машинного навчання. Ця

модель включає в себе крок попередньої обробки та об'єднані кроки виділення графологічних ознак та визначення особистості.

Як джерело вхідних даних було використано публічний набір зображень рукописного тексту. Набір даних Computer Vision Lab (CVL) складається із зображень рукописних текстів німецькою та англійською мовами, відібраних з літературних творів [77]. Набір даних складається з 7 різних рукописних текстів для кожної теми, і загалом у наборі взяли участь 310 авторів [77]. З всього набору була використана лише частина рукописного тексту, що склала загалом 404 зображення та 75 суб'єктів.

Набір даних CVL був створений та призначений для ідентифікації автора, тому, окрім унікального ідентифікатора, який визначає автора, необхідно створити додаткові позначки, щоб можна було використовувати набір даних для визначення рис особистості. Рукописний текст аналізувався вручну та за допомогою графологічних ознак, риси особистості визначалися за моделлю «Великої п'ятірки», а кожному зображенню в наборі даних було присвоєно відповідний ідентифікатор. Кореляція графологічних ознак з рисами особистості була сформована згідно з таблицею 2.1. Набір даних було розділено на п'ять класів відповідно до нових ідентифікаторів.

Для роботи з моделями машинного навчання набір даних поділено на такі частини: навчання, валідація та тестування. Недоліком цього підходу є те, що точність може змінюватися залежно від тестового набору. Для подолання цієї проблеми було використано підхід подвійної валідації. Цей підхід передбачає виконання валідації з використанням підходів пошуку та класифікації. Підхід до пошуку зображень базується на обчисленні відстаней або подібності між зображеннями. Пошукові системи обчислюють відстань або подібність запитуваного зображення до інших зображень у наборі даних. Найбільш схожі зображення (тобто ті, що мають найменші відстані)

повертаються як результати пошуку. Підхід до класифікації зображень базується на призначенні заздалегідь визначених категорій/міток на основі їхнього візуального вмісту. Механізми класифікації навчаються на категоризованих/мічених наборах даних для розпізнавання закономірностей та класифікації нових немаркованих даних.

Крок попередньої обробки вхідних даних це важлива складова підходу визначення особистості на основі рукописного тексту, бо він виконує одразу декілька завдань.

Враховуючи, що вхідними даними є зображення, і воно може бути різнокольоровим, робота з таким зображенням є складною та також вимагає додаткових обчислювальних можливостей. Для полегшення роботи із зображенням першим кроком попередньої обробки є перетворення зображення з кольорового на градації сірого.

Вхідні зображення можуть мати різні розміри, наприклад, у наборі даних CVL деякі зображення мають середній розмір 2500×1500 . Це дуже великий розмір для моделей машинного навчання, оскільки обробка такого зображення призведе до величезної кількості токенів, і, наприклад, обчислювальна складність самоуваги в зорових трансформерах масштабується квадратично з кількістю токенів. Щоб уникнути обчислювальної складності обробки таких зображень, одним із кроків попередньої обробки є крок вилучення ділянок. Зображення розділяється на ділянки фіксованого розміру, розмір ділянки було обрано рівним 32×32 , оскільки цей розмір часто використовується для зображень дуже високої роздільної здатності, і це значно зменшує кількість токенів, які обробляє модель.

Зображення в наборі даних CVL мають досить великі внутрішні поля, в яких немає відповідної інформації, тому вилучення ділянок слід виконувати з урахуванням цього. Для вилучення відповідних ділянок зображення було

додано крок попередньої обробки, на якому використовувався підхід перетворення ознак інваріантних до масштабу [78]. Цей підхід перетворює дані зображення в координати, інваріантні до масштабу, відносно локальних об'єктів. В результаті формується велика кількість ознак, які щільно покривають зображення у всьому діапазоні масштабів та розташувань. Координати цих ознак використовуються як центральна точка латки шириною 32 та висотою 32.

Під час попередньої обробки вхідних даних також важливо відокремити об'єкти від їхнього фону. Для цього було додано крок попередньої обробки на якому використано алгоритм автоматичного визначення порогу для бінаризації зображень у градаціях сірого [79]. Оптимальний поріг вибирається відповідно до дискримінантного критерію таким чином, щоб максимізувати міжкласову дисперсію рівнів сірого. Зображення бінаризується шляхом встановлення всіх пікселів зі значенням яскравості нижче оптимального порогу на 0 (чорний), а всіх пікселів зі значенням вище порогу на 1 (білий).

Модель попередньої обробки вхідних даних рукописного тексту зображена на рис. 2.7.

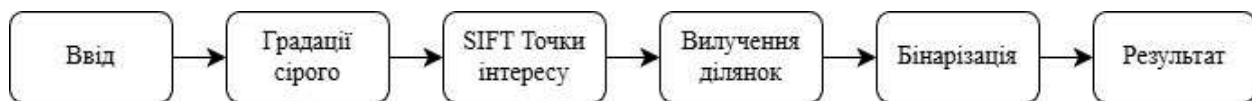


Рисунок 2.7 – Модель попередньої обробки вхідних даних

Як модель машинного навчання було обрано Зоровий трансформер (ViT) [80]. Ця модель використовує увагу як основний механізм навчання. Архітектуру моделі показано на рис. 2.8. Основою Зорового трансформера є Transformer Encoder, який складається з двох основних шарів: багатоголової самоуваги та нейронної мережі прямого зв'язку. Transformer

Encoder відповідає за обробку вхідних токенів за допомогою двох основних шарів для створення контекстно-залежних представлень. Багатоголова самоувага виконує такі функції: фіксує зв'язки між кількома патчами вхідної послідовності, обчислює зважену суму вбудовування патчів, щоб зосередитися на важливих патчах, враховуючи як глобальний, так і локальний контекст. Функція нейронної мережі прямого зв'язку зводиться до формування нелінійності та здатності вивчати складні зв'язки між ділянками. Результат роботи кожного з цих шарів передається на підрівень нормалізації та залишкових зв'язків, який відповідає за стабілізацію та прискорення навчання шляхом нормалізації вхідних даних для кожного з шарів та допомагає пом'якшити проблему зникнення градієнта.

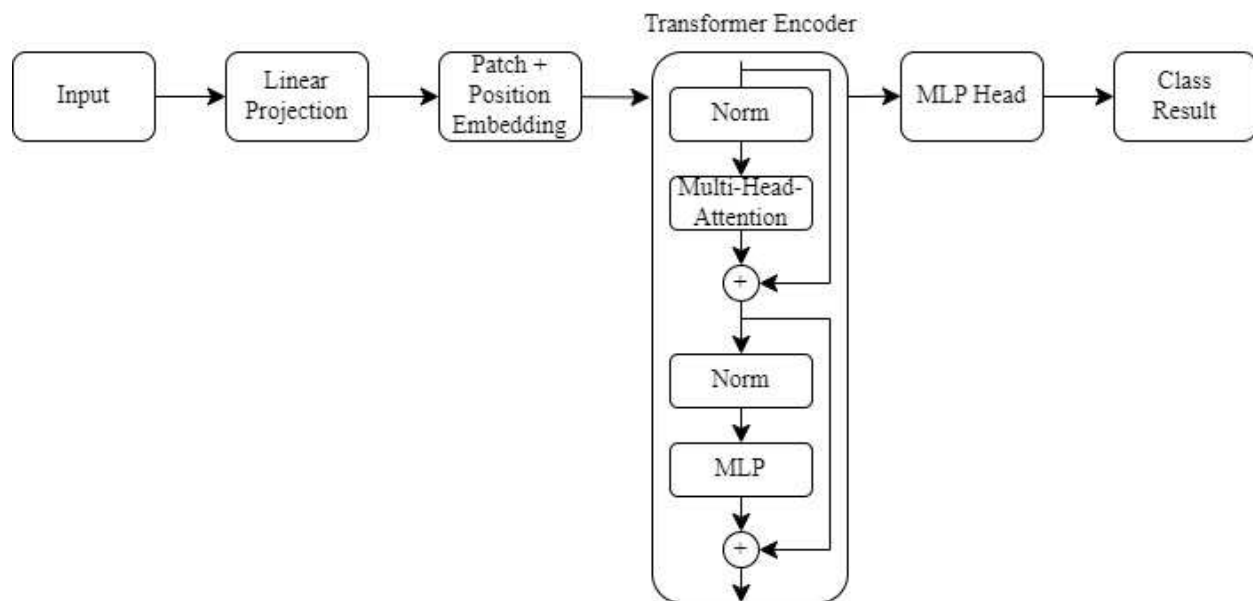


Рисунок 2.8 – Архітектура моделі Зоровий трансформер

Запропонована функціональна модель визначення особистості за допомогою методів машинного навчання зображена на рис. 2.9.

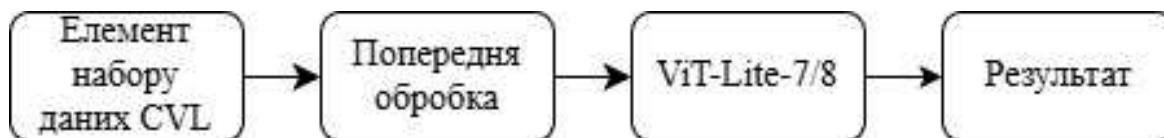


Рисунок 2.9 – Запропонована функціональна модель визначення особистості

Після попередньої обробки вхідні дані подаються в комп'ютерний зір трансформера. Під час використання зображень з набору даних передбачається, що зображення розбите на ділянки, оскільки Зоровий трансформер працює з ділянками зображення та спирається на шари перетворення для навчання відповідних функцій. Кожна ділянка лінійно проєктується у вектор, до якого додається вектор позиційного кодування, що містить інформацію про розташування ділянки на зображенні. Потім послідовність цих векторів надходить на вхід трансформера. Механізм уваги дозволяє моделі встановлювати зв'язки між різними ділянками зображення, враховуючи їх взаємне розташування та вміст. Після проходження через шари трансформера вихідний вектор категоризується за допомогою повнозв'язного шару. Результатом Зорового трансформера є ймовірність належності зображення до певного класу. Була використана модель навчання малого масштабу [81]. Ця модель підходить для невеликих наборів даних і дозволяє навчатися на невеликих наборах даних, таких як CVL, і менших, ніж ті, що традиційно використовуються в комп'ютерному зорі, такі як ImageNet. Використана модель ViT-Lite-7/8 включає 7 шарів Transformer Encoder та 8 головок уваги, що разом з класифікатором та патчами 32×32 забезпечує понад 3,74 мільйона параметрів.

Запропонований графологічний аналіз для визначення особистості реалізовано як метод, що складається з послідовності етапів перетворення вхідного зображення рукописного тексту у вектор ймовірностей належності до

класів рис особистості за моделлю «Велика п'ятірка». Формалізація методу здійснюється шляхом побудови математичної моделі, яка описує весь процес.

Нехай X_{in} – вхідне зображення з набору даних CVL. Тоді весь процес зображений на рис 2.9 можна описати функцією F , яка відображає вхідне зображення у вектор ймовірностей класів Y :

$$Y = F(X_{in}) = (C_{ViT} \circ H)(X_{in}) \quad (2.1)$$

де:

H – функція попередньої обробки;

C_{ViT} – модель класифікатора Зоровий трансформер.

В свою чергу функція попередньої обробки представлена на рис 2.7 описується як послідовність кроків H :

$$H = h_5 \circ h_4 \circ h_3 \circ h_2 \circ h_1 \quad (2.2)$$

де:

h_5 – формування тензора

h_4 – бінаризація

h_3 – розбиття на фрагменти

h_2 – виділення ознак

h_1 – градації сірого

Градації сірого (h_1) – перетворення кольорового зображення в інтенсивність яскравості:

$$h_1 = 0.299R + 0.587G + 0.114B \quad (2.3)$$

Виділення ознак (h_2) – метод масштабно-інваріантного ознакового перетворення для виявлення ключових точок, інваріантних до масштабу:

$$h_2 = SIFT(h_1) \quad (2.4)$$

Вилучення фрагменту (h_3) – розбиття зображення на фрагменти (патчі - x_p). Розмір фрагмента становить 32x32 пікселі:

$$h_3 = \{x_{i,j}\}, \text{ де } x_p \in R^{32 \times 32} \quad (2.5)$$

Бінаризація (h_4) – перетворення кожного патча в бінарний вид (0 або 1) за порогом T :

$$h_4 = \begin{cases} 255, \text{ якщо } x_p > T^* \\ 0, \text{ якщо } x_p \leq T^* \end{cases} \quad (2.6)$$

Формування результату (h_5) – формування тензора H , який подається на вхід нейромережі.

Після попередньої обробки дані потрапляють до ViT-Lite-7/8. Обробку цих даних можна виразити через механізм Self-Attention:

Кожен патч x_p перетворюється у вектор через матрицю ваг E з додаванням позиційного кодування E_{pos} :

$$z_0 = [x_{class}; x_p^1 E; x_p^2 E; \dots x_p^n E] + E_{pos} \quad (2.7)$$

Результат, z_0 , є першим вхідним сигналом для стекових трансформерних кодерів. L стекових кодерів представляють другий компонент архітектури. В

ньому $l = 1 \dots L$. Кожен трансформер приймає на вхід ознаки, представлені у вигляді тензора $(n+1, d)$, і видає вихідний сигнал тієї ж розмірності.

На другому кроці мережа вивчає більш абстрактні функції з вбудованих патчів, використовуючи стек L трансформерних кодерів. Це відповідає рівнянням (2.8) та (2.9)

$$z'_l = MSA(LN(z_{l-1})) + z_{l-1} \quad (2.8)$$

$$z_l = MLP(LN(z'_l)) + z'_l \quad (2.9)$$

де:

MSA – багатоголова самоувага;

LN – нормалізація шарів.

Компонент кодера містить механізм багаторівневої уваги (МНА) та дворівневий MLP з нормалізацією шарів та залишковими зв'язками між ними.

Вихідний вектор класифікаційного токена нормалізується та подається у повнозв'язний шар; ймовірності належності до п'яти класів рис особистості за моделлю «Велика п'ятірка» обчислюються функцією активації Softmax на виході голови класифікації (MLP Head):

$$Y = softmax(MLP(z_L^0)) \quad (2.10)$$

Де z_L^0 – стан класифікаційного токена на останньому шарі.

2.4 Експериментальне дослідження методу визначення особистості

Для перевірки запропонованого підходу та проведення експерименту було навчено кілька моделей, при цьому архітектура моделі, параметри та інше - однакові, ключова різниця полягає лише в перевірці та в композиції набору даних. Набір даних було скомпоновано наступним чином:

1. Набір даних який ділиться на половину, де перша половина використовується для навчання та валідації, а друга половина для тестування;
2. Набір даних в якому з 7 зразків рукописного тексту від кожного суб'єкта виділяється 2 зразка для навчання, 1 для валідації та 4 зразка для тестування.

На рис. 2.10 продемонстровано композицію набору даних для класифікаційного підходу.

0001-1 train	0001-2 val	0001-3 train	0001-4 test	0001-5 test	0001-6 test	0001-7 test	0001-8 test	0002-1 train	0002-2 val	0002-3 train	0002-4 test	0002-5 test	0002-6 test	0002-7 test	0002-8 test	0003-1 train	0003-2 val	0003-3 train	0003-4 test	0003-5 test	0003-6 test	0003-7 test	0003-8 test
0004-1 train	0004-2 val	0004-3 train	0004-4 test	0004-5 test	0004-6 test	0004-7 test	0004-8 test	0005-1 train	0005-2 val	0005-3 train	0005-4 test	0005-5 test	0005-6 test	0005-7 test	0005-8 test	0006-1 train	0006-2 val	0006-3 train	0006-4 test	0006-5 test	0006-6 test	0006-7 test	0006-8 test
0012-1 train	0012-2 val	0012-3 train	0012-4 test	0012-5 test	0012-6 test	0012-7 test	0012-8 test	0013-1 train	0013-2 val	0013-3 train	0013-4 test	0013-5 test	0013-6 test	0013-7 test	0013-8 test	0014-1 train	0014-2 val	0014-3 train	0014-4 test	0014-5 test	0014-6 test	0014-7 test	0014-8 test
0015-1 train	0015-2 val	0015-3 train	0015-4 test	0015-5 test	0015-6 test	0015-7 test	0015-8 test	0016-1 train	0016-2 val	0016-3 train	0016-4 test	0016-5 test	0016-6 test	0016-7 test	0016-8 test	0017-1 train	0017-2 val	0017-3 train	0017-4 test	0017-5 test	0017-6 test	0017-7 test	0017-8 test
0018-1 train	0018-2 val	0018-3 train	0018-4 test	0018-5 test	0018-6 test	0018-7 test	0018-8 test	0020-1 train	0020-2 val	0020-3 train	0020-4 test	0020-5 test	0020-6 test	0020-7 test	0020-8 test	0021-1 train	0021-2 val	0021-3 train	0021-4 test	0021-5 test	0021-6 test	0021-7 test	0021-8 test
0022-1 train	0022-2 val	0022-3 train	0022-4 test	0022-5 test	0022-6 test	0022-7 test	0022-8 test	0023-1 train	0023-2 val	0023-3 train	0023-4 test	0023-5 test	0023-6 test	0023-7 test	0023-8 test	0024-1 train	0024-2 val	0024-3 train	0024-4 test	0024-5 test	0024-6 test	0024-7 test	0024-8 test
0025-1 train	0025-2 val	0025-3 train	0025-4 test	0025-5 test	0025-6 test	0025-7 test	0025-8 test	0026-1 train	0026-2 val	0026-3 train	0026-4 test	0026-5 test	0026-6 test	0026-7 test	0026-8 test	0027-1 train	0027-2 val	0027-3 train	0027-4 test	0027-5 test	0027-6 test	0027-7 test	0027-8 test
0028-1 train	0028-2 val	0028-3 train	0028-4 test	0028-5 test	0028-6 test	0028-7 test	0028-8 test	0029-1 train	0029-2 val	0029-3 train	0029-4 test	0029-5 test	0029-6 test	0029-7 test	0029-8 test	0041-1 train	0041-2 val	0041-3 train	0041-4 test	0041-5 test	0041-6 test	0041-7 test	0041-8 test
0042-1 train	0042-2 val	0042-3 train	0042-4 test	0042-5 test	0042-6 test	0042-7 test	0042-8 test	0047-1 train	0047-2 val	0047-3 train	0047-4 test	0047-5 test	0047-6 test	0047-7 test	0047-8 test	0050-1 train	0050-2 val	0050-3 train	0050-4 test	0050-5 test	0050-6 test	0050-7 test	0050-8 test
0052-1 train	0052-2 val	0052-3 train	0052-4 test	0052-5 test	0052-6 test	0053-1 train	0053-2 val	0053-3 train	0053-4 test	0053-5 test	0053-6 test	0057-1 train	0057-2 val	0057-3 train	0057-4 test	0057-5 test	0058-1 train	0058-2 val	0058-3 train	0058-4 test	0058-5 test	0058-6 test	0059-1 train
0059-2 val	0059-3 train	0059-4 test	0059-5 test	0073-1 train	0073-2 val	0073-3 train	0073-4 test	0073-5 test	0073-6 test	0074-1 train	0074-2 val	0074-3 train	0074-4 test	0074-5 test	0074-6 test	0075-1 train	0075-2 val	0075-3 train	0075-4 test	0075-5 test	0075-6 test	0076-1 train	0076-2 val
0076-3 train	0076-4 test	0076-5 test	0077-1 train	0077-2 val	0077-3 train	0077-4 test	0077-5 test	0078-1 train	0078-2 val	0078-3 train	0078-4 test	0078-5 test	0078-6 test	0078-7 test	0078-8 test	0079-1 train	0079-2 val	0079-3 train	0079-4 test	0079-5 test	0080-1 train	0080-2 val	0080-3 train
0080-4 test	0080-5 test	0081-1 train	0081-2 val	0081-3 train	0081-4 test	0081-5 test	0083-1 train	0083-2 val	0083-3 train	0083-4 test	0083-5 test	0084-1 train	0084-2 val	0084-3 train	0084-4 test	0084-5 test	0095-1 train	0095-2 val	0095-3 train	0095-4 test	0095-5 test	0095-6 test	0095-7 test
0095-8 test	0096-1 train	0096-2 val	0096-3 train	0096-4 test	0096-5 test	0097-1 train	0097-2 val	0097-3 train	0097-4 test	0097-5 test	0098-1 train	0098-2 val	0098-3 train	0098-4 test	0098-5 test	0098-6 test	0152-1 train	0152-2 val	0152-3 train	0152-4 test	0152-5 test	0152-6 test	0152-7 test
0153-1 train	0153-2 val	0153-3 train	0153-4 test	0153-5 test	0155-1 train	0155-2 val	0155-3 train	0155-4 test	0155-5 test	0156-1 train	0156-2 val	0156-3 train	0156-4 test	0156-5 test	0157-1 train	0157-2 val	0157-3 train	0157-4 test	0157-5 test	0157-6 test	0161-1 train	0161-2 val	0161-3 train
0161-4 test	0161-5 test	0161-6 test	0162-1 train	0162-2 val	0162-3 train	0162-4 test	0162-5 test	0163-1 train	0163-2 val	0163-3 train	0163-4 test	0163-5 test	0164-1 train	0164-2 val	0164-3 train	0164-4 test	0164-5 test	0164-6 test	0165-1 train	0165-2 val	0165-3 train	0165-4 test	0165-5 test
0165-6 test	0165-7 test	0166-1 train	0166-2 val	0166-3 train	0166-4 test	0166-5 test	0167-1 train	0167-2 val	0167-3 train	0167-4 test	0167-5 test	0168-1 train	0168-2 val	0168-3 train	0168-4 test	0168-5 test	0169-1 train	0169-2 val	0169-3 train	0169-4 test	0169-5 test	0169-6 test	0169-7 test
0169-8 test	0170-1 train	0170-2 val	0170-3 train	0170-4 test	0170-5 test	0171-1 train	0171-2 val	0171-3 train	0171-4 test	0171-5 test	0172-1 train	0172-2 val	0172-3 train	0172-4 test	0172-5 test	0173-1 train	0173-2 val	0173-3 train	0173-4 test	0173-5 test	0173-6 test	0173-7 test	0173-8 test
0173-9 test	0175-1 train	0175-2 val	0175-3 train	0175-4 test	0175-5 test	0176-1 train	0176-2 val	0176-3 train	0176-4 test	0176-5 test	0177-1 train	0177-2 val	0177-3 train	0177-4 test	0177-5 test	0178-1 train	0178-2 val	0178-3 train	0178-4 test	0178-5 test	0178-6 test	0178-7 test	0178-8 test
0179-1 train	0179-2 val	0179-3 train	0179-4 test	0179-5 test	0179-6 test																		

Рисунок 2.10 – Композиція набору даних для класифікаційного підходу.

На рис. 2.11 продемонстровано композицію набору даних для пошукового підходу.

0001-1 train	0001-2 val	0001-3 train	0001-4 val	0001-6 train	0001-7 train	0001-8 train	0002-1 train	0002-2 val	0002-3 train	0002-4 val	0002-6 train	0002-7 train	0002-8 train	0003-1 train	0003-2 val	0003-3 train	0003-4 val	0003-6 train	0003-7 train	0003-8 train
0004-1 train	0004-2 val	0004-3 train	0004-4 val	0004-6 train	0004-7 train	0004-8 train	0005-1 train	0005-2 val	0005-3 train	0005-4 val	0005-6 train	0005-7 train	0005-8 train	0006-1 train	0006-2 val	0006-3 train	0006-4 val	0006-6 train	0006-7 train	0006-8 train
0012-1 train	0012-2 val	0012-3 train	0012-4 val	0012-6 train	0012-7 train	0012-8 train	0013-1 train	0013-2 val	0013-3 train	0013-4 val	0013-6 train	0013-7 train	0013-8 train	0014-1 train	0014-2 val	0014-3 train	0014-4 val	0014-6 train	0014-7 train	0014-8 train
0015-1 train	0015-2 val	0015-3 train	0015-4 val	0015-6 train	0015-7 train	0015-8 train	0016-1 train	0016-2 val	0016-3 train	0016-4 val	0016-6 train	0016-7 train	0016-8 train	0017-1 train	0017-2 val	0017-3 train	0017-4 val	0017-6 train	0017-7 train	0017-8 train
0018-1 train	0018-2 val	0018-3 train	0018-4 val	0018-6 train	0018-7 train	0018-8 train	0020-1 train	0020-2 val	0020-3 train	0020-4 val	0020-6 train	0020-7 train	0020-8 train	0021-1 train	0021-2 val	0021-3 train	0021-4 val	0021-6 train	0021-7 train	0021-8 train
0022-1 train	0022-2 val	0022-3 train	0022-4 val	0022-6 train	0022-7 train	0022-8 train	0023-1 train	0023-2 val	0023-3 train	0023-4 val	0023-6 train	0023-7 train	0023-8 train	0024-1 train	0024-2 val	0024-3 train	0024-4 val	0024-6 train	0024-7 train	0024-8 train
0025-1 train	0025-2 val	0025-3 train	0025-4 val	0025-6 train	0025-7 train	0025-8 train	0026-1 train	0026-2 val	0026-3 train	0026-4 val	0026-6 train	0026-7 train	0026-8 train	0027-1 train	0027-2 val	0027-3 train	0027-4 val	0027-6 train	0027-7 train	0027-8 train
0028-1 train	0028-2 val	0028-3 train	0028-4 val	0028-6 train	0028-7 train	0028-8 train	0029-1 train	0029-2 val	0029-3 train	0029-4 val	0029-6 train	0029-7 train	0029-8 train	0041-1 train	0041-2 val	0041-3 train	0041-4 val	0041-6 train	0041-7 train	0041-8 train
0042-1 train	0042-2 val	0042-3 train	0042-4 val	0042-6 train	0042-7 train	0042-8 train	0047-1 train	0047-2 val	0047-3 train	0047-4 val	0047-6 train	0047-7 train	0047-8 train	0050-1 train	0050-2 val	0050-3 train	0050-4 val	0050-6 train	0050-7 train	0050-8 train
0052-1 test	0052-2 test	0052-3 test	0052-4 test	0052-6 test	0053-1 test	0053-2 test	0053-3 test	0053-4 test	0053-6 test	0057-1 test	0057-2 test	0057-3 test	0057-4 test	0057-6 test	0058-1 test	0058-2 test	0058-3 test	0058-4 test	0058-6 test	0059-1 test
0059-2 test	0059-3 test	0059-4 test	0059-6 test	0073-1 test	0073-2 test	0073-3 test	0073-4 test	0073-6 test	0074-1 test	0074-2 test	0074-3 test	0074-4 test	0074-6 test	0075-1 test	0075-2 test	0075-3 test	0075-4 test	0075-6 test	0076-1 test	0076-2 test
0076-3 test	0076-4 test	0076-6 test	0077-1 test	0077-2 test	0077-3 test	0077-4 test	0077-6 test	0078-1 test	0078-2 test	0078-3 test	0078-4 test	0078-6 test	0079-1 test	0079-2 test	0079-3 test	0079-4 test	0079-6 test	0080-1 test	0080-2 test	0080-3 test
0080-4 test	0080-6 test	0081-1 test	0081-2 test	0081-3 test	0081-4 test	0081-6 test	0083-1 test	0083-2 test	0083-3 test	0083-4 test	0083-6 test	0084-1 test	0084-2 test	0084-3 test	0084-4 test	0084-6 test	0095-1 test	0095-2 test	0095-3 test	0095-4 test
0095-6 test	0096-1 test	0096-2 test	0096-3 test	0096-4 test	0096-6 test	0097-1 test	0097-2 test	0097-3 test	0097-4 test	0097-6 test	0098-1 test	0098-2 test	0098-3 test	0098-4 test	0098-6 test	0152-1 test	0152-2 test	0152-3 test	0152-4 test	0152-6 test
0153-1 test	0153-2 test	0153-3 test	0153-4 test	0153-6 test	0155-1 test	0155-2 test	0155-3 test	0155-4 test	0155-6 test	0156-1 test	0156-2 test	0156-3 test	0156-4 test	0156-6 test	0157-1 test	0157-2 test	0157-3 test	0157-4 test	0157-6 test	0161-1 test
0161-2 test	0161-3 test	0161-4 test	0161-6 test	0162-1 test	0162-2 test	0162-3 test	0162-4 test	0162-6 test	0163-1 test	0163-2 test	0163-3 test	0163-4 test	0163-6 test	0164-1 test	0164-2 test	0164-3 test	0164-4 test	0164-6 test	0165-1 test	0165-2 test
0165-3 test	0165-4 test	0165-6 test	0166-1 test	0166-2 test	0166-3 test	0166-4 test	0166-6 test	0167-1 test	0167-2 test	0167-3 test	0167-4 test	0167-6 test	0168-1 test	0168-2 test	0168-3 test	0168-4 test	0168-6 test	0169-1 test	0169-2 test	0169-3 test
0169-4 test	0169-6 test	0170-1 test	0170-2 test	0170-3 test	0170-4 test	0170-6 test	0171-1 test	0171-2 test	0171-3 test	0171-4 test	0171-6 test	0172-1 test	0172-2 test	0172-3 test	0172-4 test	0172-6 test	0173-1 test	0173-2 test	0173-3 test	0173-4 test
0173-6 test	0175-1 test	0175-2 test	0175-3 test	0175-4 test	0175-6 test	0176-1 test	0176-2 test	0176-3 test	0176-4 test	0176-6 test	0177-1 test	0177-2 test	0177-3 test	0177-4 test	0177-6 test	0178-1 test	0178-2 test	0178-3 test	0178-4 test	0178-6 test
0179-1 test	0179-2 test	0179-3 test	0179-4 test	0179-6 test																

Рисунок 2.11 – Композиція набору даних для пошукового підходу.

Також при проведенні експерименту було використано random seed, він використовувався в усіх модулях моделі. Це необхідно для відтворюваності результату та для того щоб пересвідчитись, що забезпечується однаковий порядок перемішування набору даних під час кожного запуску експерименту.

Для уникнення перенавчання було обрано оптимальну кількість ітерацій з можливістю дострокового завершення навчання. Навчання достроково завершується, якщо значення середньої втрати навчання за партію не зменшується протягом 10 епох.

Параметри, які були обрані та використані при налаштуванні моделі представлені в табл. 2.2.

Параметри, які були обрані та використані при налаштуванні масштабно-інваріантного ознакового перетворення представлені в табл. 2.3.

Таблиця 2.2 – Параметри налаштування моделі

Параметр	Значення	Опис
lr	0.0005	класичний learning rate;
num-epochs-warmup	5	кількість епох розігріву темпу навчання
batch-size	128	розмір пакета
num-workers	4	кількість PyTorch працівників
num-epochs-patience	10	кількість епох, після котрих навчання буде зупинено, якщо валідація втрат більше не покращується;
num-epochs	60	кількість епох

Таблиця 2.3 – Параметри масштабно-інваріантного ознакового перетворення

Параметр	Значення
nfeatures	0
nOctaveLayers	3
contrastThreshold	0.04
edgeThreshold	10
sigma	3.75

Під час навчання моделей було проведено валідацію навчання, у вигляді збирання та контролю таких метрик як loss_val, loss_train, accuracy_val,

accuracy_train. Вони збиралися для підтвердження коректного процесу навчання.

На рис. 2.12-2.15 продемонстровані графіки які було отримано під час навчання моделей. Графіки на рис. 2.12-2.13 демонструють метрики зібрані під час навчання моделі для класифікаційного підходу, графіки на рис. 2.14-2.15 демонструють метрики зібрані під час навчання моделі для пошукового підходу.

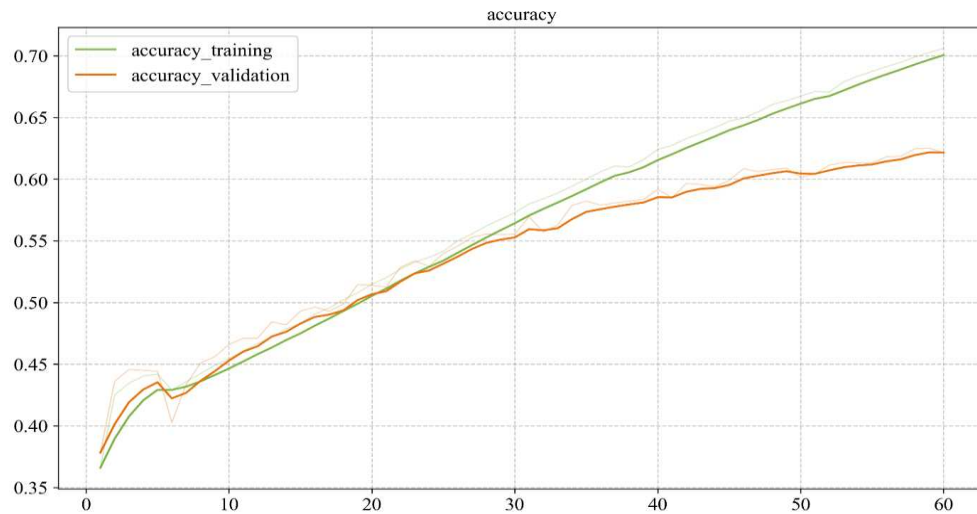


Рисунок 2.12 – Точність валідації та навчання моделі на основі ViT на тестовому наборі даних для класифікаційного підходу.

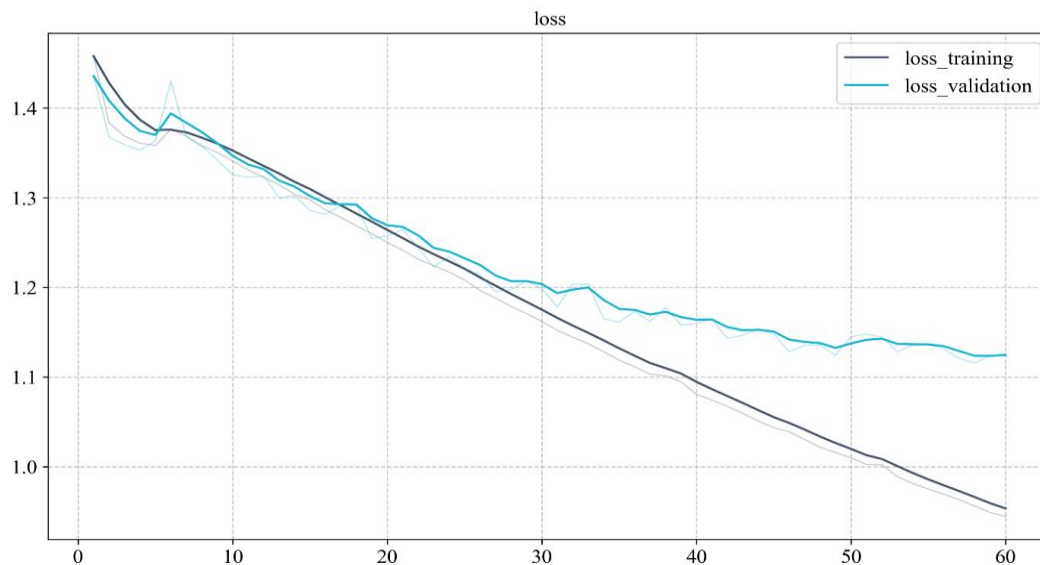


Рисунок 2.13 – Втрати валідації та навчання моделі на основі ViT на тестовому наборі даних для класифікаційного підходу.

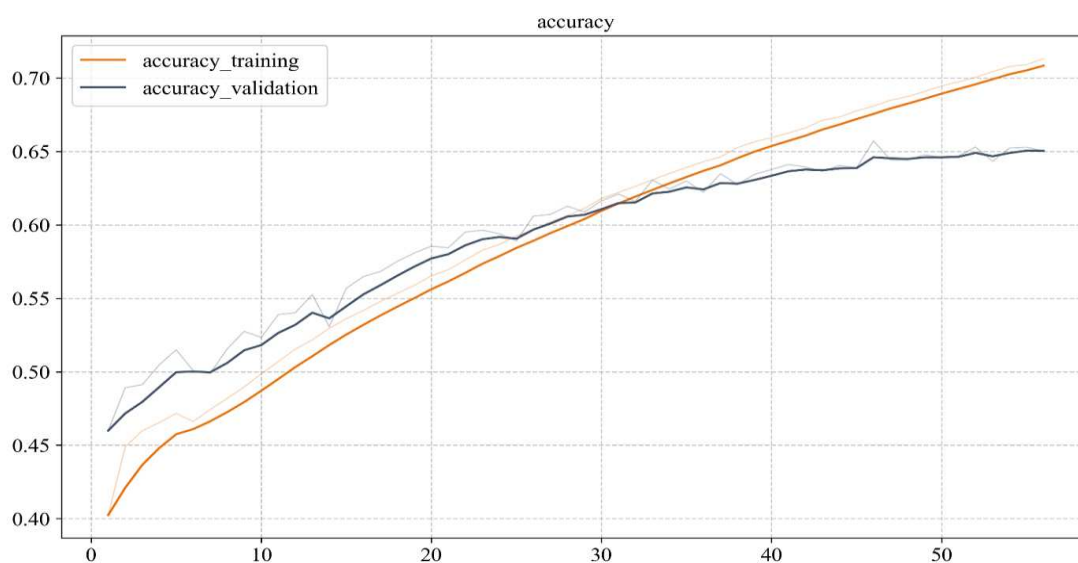


Рисунок 2.14 – Точність валідації та навчання моделі на основі ViT на тестовому наборі даних для пошукового підходу.

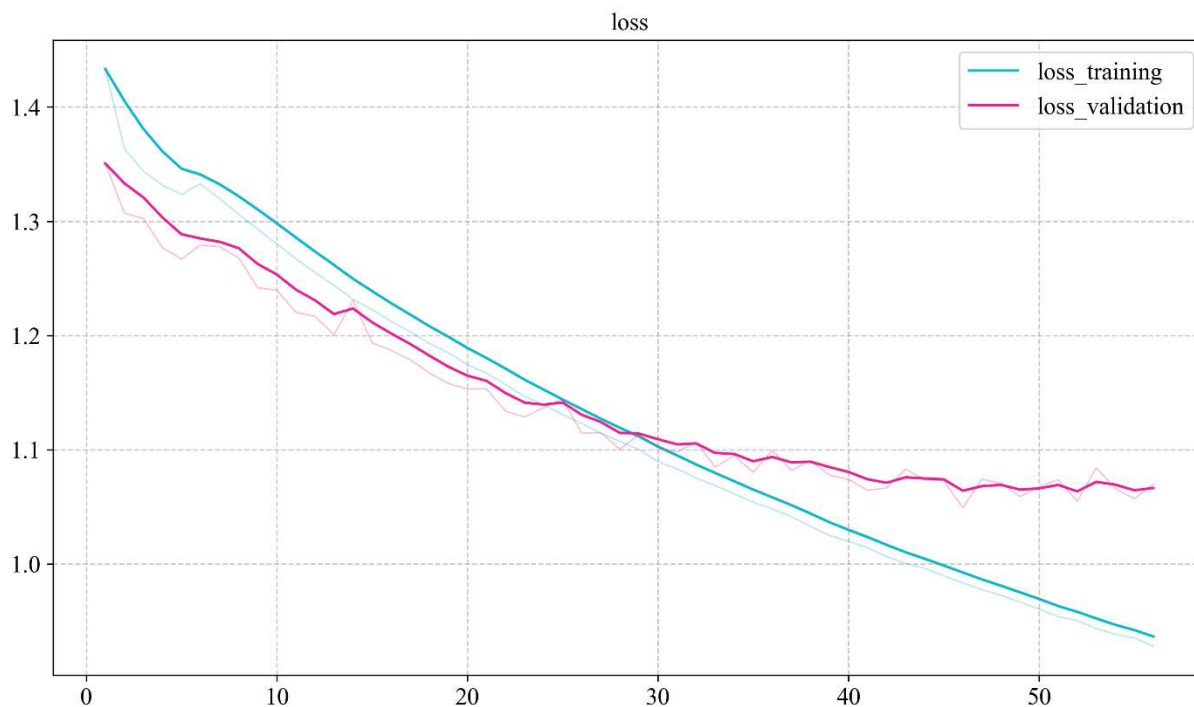


Рисунок 2.15 – Втрати валідації та навчання моделі на основі ViT на тестовому наборі даних для пошукового підходу.

На рис. 2.16 представлено матрицю невідповідностей для підходу класифікації.

Матриця невідповідностей моделі відображає правильно та неправильно класифіковані зразки. На рис. 2.16 вісь X – це справжня мітка, вісь Y – передбачувана мітка, а перетин осей показує кількість правильних передбачень.

Для основного експерименту був проведений аналіз точності моделі для підходу класифікації та точність пошуку відповідних ознак на основі евклідової відстані їх векторів ознак для пошукового підходу. Додатково для обох підходів було враховано й інші базові метрики, корисні для оцінки успішності підходу визначення особистості на основі рукописного тексту за допомогою методів машинного навчання.

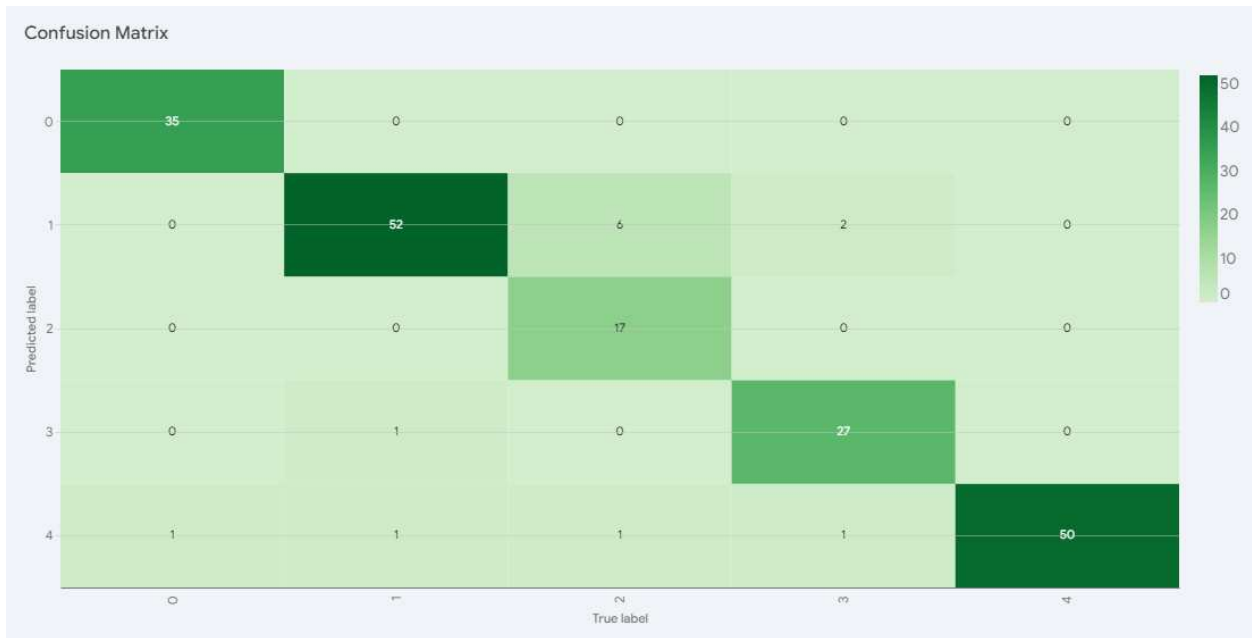


Рисунок 2.16 – Матриця невідповідностей моделі на основі ViT на тестовому наборі даних для підходу класифікації

Метрики, які були використані для оцінки ефективності пошуку: м'який топ k, суворий топ k, повнота на k (2.11), влучність на k (2.12), середнє значення точності (2.13).

$$Recall@k = \frac{numRelItemsInTopK}{numRelItemsPerLabel} , \quad (2.11)$$

де $numRelItemsInTopK$ – кількість релевантних елементів у перших k рекомендаціях,

$numRelItemsPerLabel$ – загальна кількість релевантних елементів у всьому наборі даних.

$$Precision@k = \frac{numRelItemsInTopK}{k} , \quad (2.12)$$

де $numRelItemsInTopK$ – кількість релевантних елементів у перших k рекомендаціях,

k – кількість рекомендацій.

$$mAP = \frac{1}{n} \sum_{k=1}^{k=n} AvgP_k, \quad (2.13)$$

де n – кількість документів,

$AvgP$ – середня влучність документа.

Метрики, які були використані для оцінки ефективності класифікації: точність (2.14), влучність (2.15), повнота (2.16), оцінка F1 (2.17), істинно позитивний (TP), істинно негативний (TN), та помилки першого та другого роду хибнопозитивний (FP) та хибнонегативний (FN) результат.

$$Accuracy = \frac{TN + TP}{TP + FP + TN + FN}, \quad (2.14)$$

де TN – істинно негативний,

TP – істинно позитивний,

FP – хибний позитивний,

FN – хибний негативний.

$$Precision = \frac{TP}{TP + FP}, \quad (2.15)$$

де TP – істинно позитивний,

FP – хибний позитивний.

$$Recall = \frac{TP}{TP + FN}, \quad (2.16)$$

де TP – істинно позитивний,

FN – хибний негативний.

$$F1Score = \frac{2 \times Precision \times Recall}{Precision + Recall}, \quad (2.17)$$

де $Precision$ – влучність,

$Recall$ – повнота.

Метрики оцінки продуктивності отримані в результаті проведення тестів для підходу класифікації та для пошукового підходу представлені в табл. 2.4, 2.5.

Таблиця 2.4 – Метрики оцінки продуктивності ViT для підходу класифікації

Метрика	Характеристика	Значення
Accuracy	Сумлінність	99.48 %
	Співробітництво	94.84 %
	Екстраверсія	96.39 %
	Відкритість	97.93 %
	Нейротизм	97.93 %
F1 Score	Сумлінність	0.98
	Співробітництво	0.91
	Екстраверсія	0.82
	Відкритість	0.93
	Нейротизм	0.96

Продовження таблиці 2.4

Precision	Сумлінність	1.0
	Співробітництво	0.86
	Екстраверсія	1.0
	Відкритість	0.96
	Нейротизм	0.92
Recall	Сумлінність	0.97
	Співробітництво	0.96
	Екстраверсія	0.70
	Відкритість	0.89
	Нейротизм	1.0

Результати дослідження показали, що Зоровий трансформер може значно покращити визначення особистості на основі рукописного тексту, оскільки демонструє більшу точність.

Таблиця 2.5 – Метрики оцінки продуктивності ViT для пошукового підходу з використанням евклідової відстані

Метрика	k	Значення
Hard Top K	1	0.84
	2	0.65
Soft Top K	1	0.84
	2	0.91
	3	0.94
	4	0.97
	5	0.98
Precision at K	1	0.84
	2	0.78
	3	0.75
	4	0.73
	5	0.68

Продовження таблиці 2.5

Recall at K	1	0.02
	2	0.03
	3	0.05
	4	0.07
	5	0.08
mAP		0.45

На відміну від роботи [37], в якій модель значною мірою залежить від етапів попередньої обробки, таких як нормалізація, сегментація, вилучення ознак та формування проміжного представлення, Зоровий трансформер працює з ділянками зображення та спирається на шари трансформації для вивчення відповідних ознак, ця відмінність вирішує проблему неоднозначності у зіставленні ознак з показниками типів особистості.

Водночас, порівнюючи значення точності для кожної окремої риси, можна побачити, що для екстраверсії їхня модель продемонструвала 87,1%, а запропонований підхід 96,39%, для нейротизму – 86,3%, а запропонований підхід 97,93%, для співробітництва – 79,6%, а запропонований підхід 94,84%, для сумлінності – 80,4%, а запропонований підхід 99,48%, для відкритості до досвіду – 88,6%, а запропонований підхід 97,93%.

На відміну від роботи [44], в якій для кожної окремої риси особистості навчалася окрема модель, у цій роботі навчалася єдина модель, яка розрізняє п'ять рис особистості. Таке розмежування дозволяє зменшити обчислювальну складність.

Це стає можливим завдяки механізму самоуваги, який дозволяє моделі звертати увагу на різні частини вхідних даних на основі їх відповідності до завдання. Водночас, порівнюючи значення точності для кожної окремої риси, можна побачити, що для екстраверсії їхня модель продемонструвала 85%, а

запропонований підхід 96,39%, для нейротизму – 70%, а запропонований підхід 97,93%, для співробітництва – 75%, а запропонований підхід 94,84%, для сумлінності – 65%, а запропонований підхід 99,48%, для відкритості до досвіду – 65%, а запропонований підхід 97,93%.

На відміну від роботи [46], в якій модель навчалася з використанням маркованих, немаркованих та згенерованих генератором даних, у цій роботі навчання проводилося лише з використанням маркованих даних, ця різниця дозволяє спростити процес навчання, зменшити обчислювальні витрати та краще контролювати якість навчання та точність результатів. Це можливо завдяки тому, що навчання Зорового трансформера, на відміну від SGAN, базується виключно на маркованих даних. Водночас, порівнюючи значення точності, можна побачити, що їхня модель продемонструвала в середньому 91,3%, а запропонований підхід в середньому 97%.

Таким чином, підхід до визначення особистості на основі рукописного тексту з використанням методів машинного навчання має великий практичний потенціал для графологів та експертів з почерку, дозволяючи їм ефективніше виконувати аналіз під час роботи. Цей підхід може прискорити та автоматизувати процес визначення психологічного портрета автора та позитивно вплинути на розвиток автоматизованих систем графологічного аналізу. Такий автоматизований аналіз почерку дозволяє швидко та об'єктивно оцінити особисті якості людини, такі як рівень екстраверсії, нейротизму, співробітництва, сумлінності та відкритості до досвіду.

Очікувані ефекти від впровадження цього підходу включають підвищення точності та об'єктивності аналізу почерку, скорочення часу, необхідного для експертизи, зменшення впливу людського фактору та суб'єктивних оцінок, а також можливість аналізу великих обсягів рукописних даних. Це, у свою чергу, призведе до більш ефективної роботи експертів,

підвищення якості підбору персоналу та глибшого розуміння психологічного стану людини в різних сферах діяльності.

Використання публічного набору даних позитивно впливає на відтворюваність дослідження. Також передбачається, що перед використанням набору даних буде проведено графологічний аналіз, і кожен зразок рукописного тексту буде позначено однією з п'яти оцінок згідно з моделлю «Великої п'ятірки» (екстраверсія, нейротизм, співробітництво, сумлінність, відкритість до досвіду). Під час відтворення результатів для методу пошуку важливо використовувати евклідову відстань; під час використання інших методів розрахунку відстані результати можуть відрізнятися.

2.5 Висновки за другим розділом

У цьому розділі набув подальшого розвитку та був експериментально досліджений метод визначення особистості за рукописним текстом, спрямований на підвищення точності та автоматизацію аналізу графологічних ознак. Як психометрична основа була обґрунтована та використана модель «Велика п'ятірка», а також встановлено кореляцію її рис із відповідними графологічними особливостями почерку.

1. За результатами аналізу сучасних методів визначення особистості за рукописним текстом виявлено їхні основні недоліки – схильність до перенавчання, потребу в ручній кореляції графологічних ознак із рисами особистості та суб'єктивність вибору ознак, через що точність методів на основі моделі «Велика п'ятірка» не перевищувала 92%. На цій основі обґрунтовано та сформульовано задачу побудови підходу з використанням публічного набору даних CVL, індикатора «Велика п'ятірка», попередньої

обробки на основі масштабно-інваріантного ознакового перетворення та методу машинного навчання Зоровий трансформер.

2. Обґрунтовано вибір психометричної моделі «Велика п'ятірка» як найбільш інформативної та стабільної в різних демографічних групах та сформовано таблицю відповідності п'яти рис особистості (Відкритість, Сумлінність, Екстраверсія, Співробітництво, Нейротизм) графологічним ознакам почерку (базова лінія, нахил, міжрядковий і міжсловний інтервали, натиск пера, розмір літер), яка стала основою для розмітки набору даних.

3. Було запропоновано метод визначення психологічних характеристик особистості, що включає етап попередньої обробки (переведення в градації сірого, вилучення інформативних ділянок за допомогою масштабно-інваріантного ознакового перетворення, розбиття на фрагменти 32×32 та бінаризацію) і класифікацію на основі архітектури Зоровий трансформер ViT-Lite-7/8, було виконано математичну формалізацію методу та реалізовано підхід подвійної валідації – класифікаційний і пошуковий.

4. Проведено експериментальне дослідження на розміченому наборі CVL, яке підтвердило високу ефективність підходу. У класифікаційному підході досягнуто точності 99,48% для риси Сумлінність, 97,93% для Відкритості та Нейротизму, 96,39% для Екстраверсії та 94,84% для Співробітництва, а в пошуковому підході з показником Soft Top K до 0,98. Також запропонований метод показав кращі результати, в порівнянні з іншими підходами. Це підтверджує успішне вирішення поставлених завдань і створює надійну основу для подальшого розвитку автоматизованих систем графологічного аналізу в області визначення особистості.

Перелік використаних джерел у даному розділі наведено у повному переліку використаних джерел під номерами: 21–81.

3 УДОСКОНАЛЕНИЙ МЕТОД ІДЕНТИФІКАЦІЇ ТА ПОШУКУ АВТОРА

У третьому розділі розглянуто удосконалений метод ідентифікації та пошуку автора на основі рукописного тексту. Основну увагу приділено аналізу існуючих методів ідентифікації та пошуку автора, постановці задачі та обґрунтуванню вибору методу виявлення ознак. Запропоновано підхід до ідентифікації автора за допомогою методів машинного навчання, який використовує Зоровий трансформер та масштабно-інваріантний центрально-окружний детектор (CenSurE). У процесі розробки побудовано функціональну модель ідентифікації автора та проведена її формалізація шляхом побудови математичної моделі. Окрему увагу приділено експериментальному дослідженню методу, яке продемонструвало покращення показників пошуку, зокрема Hard Top K до 0.96 та Soft Top K до 0.96, а також збільшення кількості виявлених ознак порівняно з базовими методами. Запропонована модель створює основу для підвищення надійності автоматизованих систем аналізу рукописного тексту та є важливим внеском у підвищення точності ідентифікації автора.

3.1 Постановка задачі ідентифікації та пошуку автора

Технологічний прогрес людства безпосередньо впливає на повсякденну діяльність людини. Незважаючи на це, рукописний текст продовжує використовуватися в різних сферах людського життя. Унікальність почерку ґрунтується на стані здоров'я, віці, здібностях, рисах особистості та характері людини, що створює теоретичну основу для його аналізу. Аналіз почерку має три основні галузі: графологію, почеркознавчу експертизу та судову

лінгвістику [21]. Графологія вивчає поведінку людини, яку можна розпізнати за почерком [4], а її надійність характеризується тим, що вимірювані змінні в почерку надзвичайно послідовні [21].

Ідентифікація автора під час аналізу почерку використовується в різних сферах діяльності, включаючи медицину, криміналістику та судочинство [4]. Одним із прикладів використання аналізу почерку є його використання в криміналістиці під час розслідування різних злочинів [82]. Аналіз почерку використовується для відтворення емоційного стану людини на момент злочину, оскільки почерк може розкрити емоційний стан людини. Це може бути одним із ключових факторів у визначенні обставин злочину [83].

Існує чимало публікацій, присвячених аналізу почерку в різних сферах людської діяльності. Однак широке використання таких методів все ще обмежується недостатньою формалізацією характеристик почерку, суб'єктивною експертною оцінкою, обмеженим апаратним забезпеченням та недосконалими алгоритмами. Усі ці фактори перешкоджають підвищенню точності та надійності результатів.

Сучасні підходи до аналізу почерку часто не враховують мінливість почерку та потребують модифікації для підвищення точності результатів аналізу. Більшість підходів не враховують фізичний та/або емоційний стан людини, тип інструменту, що використовується для написання тексту, швидкість письма та мовні характеристики. Усі ці фактори обмежують використання аналізу почерку в таких критичних галузях, як криміналістика, дослідження документів та інші.

Сучасні розробки в оптичних технологіях у поєднанні з інформаційними технологіями розширюють можливості аналізу почерку завдяки створенню камер високої роздільної здатності. Це дозволяє виявляти дрібні деталі в характеристиках почерку авторів, що в поєднанні із сучасними методами

обробка візуальних даних на основі штучного інтелекту відкриває нові можливості для аналізу [4].

Використання штучного інтелекту, зокрема методів машинного навчання, дозволить створити більш ефективне та надійне програмне забезпечення для ідентифікації авторів почерку. Це програмне забезпечення не лише ідентифікуватиме автора, але й надасть експертам обґрунтовані характеристики почерку, що підтверджують авторство. Такі програмні інструменти дозволять експертам робити більш об'єктивні висновки під час аналізу почерку, що, у свою чергу, розширить практичне застосування аналізу почерку.

Тому, враховуючи все вищезазначене, необхідні подальші дослідження для вдосконалення підходів, які подолають перелічені вище обмеження та забезпечать краще розуміння механізмів формування та ідентифікації почерку. Одним із ключових завдань є розробка єдиних стандартів аналізу почерку. Зокрема, це передбачає інтеграцію сучасних автоматизованих статистичних методів, що підвищить об'єктивність та валідність результатів.

Отже, актуальність розробки та вдосконалення підходів до ідентифікації автора на основі рукописного тексту є суттєво важливою, оскільки це не тільки подолає існуючі обмеження, але й відкриє нові можливості для практичного застосування в різних галузях.

Аналіз досліджень у галузі ідентифікації автора показав, що існує два підходи до ідентифікації автора за рукописним текстом: глибоке навчання та звичайні алгоритми ручного створення [83].

Автори в [22] запропонували використовувати щільно обчислювані дескриптори RootSIFT для фіксації локальних властивостей рукописних документів, використовувати супервектори GMM як метод кодування та використовувати методи опорних векторів Exemplar (Exemplar-SVM) для

вимірювання подібності. Цей підхід спирається на ручно створені ознаки та складний конвеєр, що включає адаптацію GMM та Exemplar-SVM, що робить його чутливим до налаштування параметрів та обмежує його здатність узагальнюватися до невидимих варіацій рукописного тексту. Трансформер зору потенційно може запропонувати більш оптимізоване та адаптоване рішення, що усуває залежність від ручно створених ознак.

У статті [23] автори запропонували метод навчання ознак без учителя з використанням активацій згорткової нейронної мережі (CNN), де сурогатні класи створюються шляхом кластеризації для навчання глибокої мережі залишків, а активації з передостаннього шару служать ознаками для наступних завдань класифікації, таких як ідентифікація автора. Запропонований підхід сильно залежить від якості кластеризації для створення набору класів ознак почерку. Але це, у свою чергу, значно впливає на повноту виявлення ознак почерку, на відміну від методів машинного навчання з учителем. Очікується, що цю проблему можна подолати за допомогою моделі ViT, яка дозволить виявляти внутрішні закономірності та зв'язки в почерку.

У публікації [24] автори використали метод зваженої регуляризації згладжування міток (WLSR) для розширення навчального набору немаркованими даними та збільшення роздільної здатності мережі CNN (ResNet-50). Цей метод показав хороші результати завдяки використанню немаркованих даних. Але таке використання мережі CNN обмежує здатність усього підходу враховувати глобальні залежності у стилі письма, а продуктивність підходу також залежить від ефективності WLSR. Модель ViT може навчатися більш повним та надійним представленням без необхідності явного доповнення даних, як у випадку WLSR.

У [30] було запропоновано метод глибокого адаптивного навчання для ідентифікації автора на основі зображень окремих слів з використанням багатозадачного навчання. До процесу навчання додається допоміжне завдання, щоб забезпечити появу ознак повторного використання. Запропонований метод переносить переваги вивчених ознак згорткової нейронної мережі з допоміжного завдання, такого як явне розпізнавання контенту, на основне завдання ідентифікації автора в рамках однієї процедури. Однак, такий метод сильно залежить від відповідності вибраних допоміжних завдань та ефективності адаптивного згорткового шару мережі у виявленні узагальнюючих ознак. ViT може покращити цей підхід, навчаючись більш повним та надійним представленням безпосередньо з образів слів, без необхідності попередньо визначених допоміжних завдань.

У статті [25] було представлено згорткову нейронну мережу (ЗНМ) для вирішення проблеми ідентифікації автора. Також оцінюється продуктивність ознак, отриманих з ЗНМ, при подачі їх до класифікатора на основі машини опорних векторів (SVM), враховуючи традиційний підхід до класифікації. Цей підхід може бути чутливим до певних параметрів процесу генерації текстури та архітектури ЗНМ. На противагу цьому, модель ViT здатна безпосередньо витягувати відповідні деталі з оригінальних зображень рукопису та формувати глибші представлення стилю письма.

У [26] використовується комбінація глибоких дескрипторів та дескрипторів рукописного введення для вивчення шаблонів з рукописних зображень. Спочатку з рукописних зображень витягуються локальні області. Потім ці області одночасно подаються до глибоких дескрипторів та дескрипторів рукописного введення для генерації локальних описів. Витягнуті локальні ознаки потім агрегуються для створення цілої матриці опису. Далі, застосовуючи кодування векторних локально агрегованих дескрипторів

(VLAD) до матриці опису, отримується одновимірний вектор ознак для представлення авторського шаблону. До недоліків цього підходу можна віднести той факт, що використання лише локальних фрагментів без урахування їх просторових зв'язків може обмежити можливість виявлення глибших стилістичних шаблонів. Також поєднання ручно створених ознак та глибоких описових ознак може негативно вплинути на ефективність усього підходу, оскільки він дуже залежить від вибору конкретних дескрипторів та їхнього впливу на кінцевий результат.

На противагу цьому, модель ViT може навчатися більш надійним представленням, фіксуючи глобальні зв'язки між латками. Це дозволяє їй побудувати більш повне та цілісне розуміння стилю письма людини.

У публікації [27] автори представили модель без сегментації, засновану на CNN, з використанням слабо контрольованого механізму вибору регіонів для характеристики автора заданої вибірки. Хоча цей підхід є ефективним, він має обмеження. Його ефективність залежить від використання фіксованого відсотка верхніх регіонів для аналізу та чутлива до якості карти ймовірностей CNN. На противагу цьому, модель ViT пропонує більш гнучкий метод, оскільки вона може динамічно вибирати найбільш інформативні регіони, вивчаючи їхню важливість з контексту.

Автори в [31] запропонували глибоку нейронну мережу FragNet з двома шляхами: пірамідою ознак для вилучення карт ознак та шляхом фрагментів для прогнозування автора. Зосередження FragNet на ізольованих зображеннях слів обмежує її здатність захоплювати ширшу контекстну інформацію та робить її чутливою до варіацій якості зображення, що потенційно призводить до перенавчання та труднощів з обробкою різноманітних стилів почерку. Хоча ViT вимагає обчислювальних ресурсів, вона може покращити ідентифікацію автора, використовуючи свій механізм самоуваги для захоплення

довгострокових залежностей у більших текстових зразках та забезпечуючи більшу стійкість до варіацій зображень.

У роботі [84], на відміну від підходів на основі CNN, представлено метод на основі ViT з використанням алгоритму SIFT для виявлення ключових точок, який демонструє високу продуктивність на кількох наборах даних, включаючи судово-медичні дані, з високою точністю top-1 та аналізом впливу сценарію та стилю письма. Однак використання SIFT може бути неоптимальним, оскільки воно може не виявити достатню кількість надійних ключових точок для аналізу ViT.

У публікації [29] автори запропонували використовувати Resnet-34, навчений для виявлення глибоких ознак з невеликих ділянок, ділянки витягуються за допомогою алгоритму FAST та детектора Harris Corner (HC), тоді як машинно-навчені ознаки кодується за допомогою VLAD та тріангуляційного вбудовування. Хоча це ефективно, використання детектора ключових точок FAST та детектора HC може пропустити важливу стилістичну інформацію, присутню в інших частинах рукописного тексту, оскільки вони покладаються на прості локалізовані ознаки, такі як кути. Використання детекторів ключових точок SIFT або CenSurE може ще більше покращити вилучення ділянок, оскільки вони враховують ширший контекст навколо кожної ключової точки та є більш стійкими до варіацій масштабу, обертання та освітлення.

Автори в [32] запропонували глобально-контекстну залишкову рекурентну нейронну мережу (GR-RNN), де просторовий зв'язок між послідовністю фрагментів моделюється рекурентною нейронною мережею (RNN) для посилення розрізнявальної здатності локальних ознак фрагментів, а також використовується додаткова інформація між глобально-контекстними та локальними фрагментами. Однак цей метод потребує сегментації

зображення слів або підслів, що вимагає додаткових попередніх кроків для його застосування до інших документів.

У статті [85] автори запропонували використовувати сіамську мережу для генерації закодованого представлення та вилучення ідентифікаційної мітки, вхідні дані сегментуються на два входи, а потім обробляються екстрактором ознак Xception. Хоча цей підхід демонструє вражаючу точність, найбільшим недоліком є час навчання, який триває від п'яти до семи днів.

Аналіз наукових публікацій [22-32, 84, 85] показав, що, незважаючи на значний прогрес, існуючі методи ідентифікації автора мають суттєві обмеження. З одного боку, підходи на основі CNN часто не здатні ефективно враховувати глобальний контекст та довгострокові залежності в почерку. З іншого боку, новіші архітектури, такі як глибокі згорткові мережі або рекурентні нейронні мережі, хоча й демонструють високий потенціал для аналізу складних шаблонів, їхня ефективність значною мірою компенсується високими обчислювальними витратами на навчання та критичною чутливістю до методів підготовки вхідних даних.

Таким чином, поєднання цих проблем та недоліків вказує на необхідність додаткових досліджень у галузі ідентифікації автора для створення оптимізованого та комплексного підходу до ідентифікації автора.

Отже, метою даної частини дослідження є розробка методу ідентифікації автора, який використовуватиме машинне навчання, аналізуватиме та використовуватиме методи виявлення ознак і методи попередньої обробки зображень. Це призведе до підвищення надійності автоматизованих систем ідентифікації автора за почерком, що посилить аналітичні можливості фахівців у галузі графології та експертизи почерку.

Для досягнення мети були поставлені такі завдання:

- побудувати функціональну модель ідентифікації автора за допомогою методу машинного навчання;
- провести експериментальне дослідження запропонованого підходу.

3.2 Метод виявлення ознак зображень

Під час аналізу зображень рукописного тексту ідентифікація точок інтересу (ТІ), які іноді також називають ключовими точками або визначними ознаками є важливим етапом для надійного вилучення ознак. Такі точки представляють локалізовані області на зображенні, які демонструють високу інформаційну ентропію. Прикладами таких областей є з'єднання штрихів, кінці та максимуми локальної кривизни. Виявлення таких координат є важливим, оскільки рукописний текст характеризується високою внутрішньоавторською мінливістю та стохастичним шумом. Зосереджуючи аналіз на дискретних ТІ, а не на всій піксельній сітці, система ефективно зменшує розмірність даних, зберігаючи при цьому основні топологічні властивості письма. Виявлення ТІ у зображеннях рукописного тексту є важливим також тому, що ТІ здатні забезпечувати геометричну інваріантність. Почерк підлягає афінним перетворенням. Це проявляється через зміни нахилу, масштабу та обертання. Виявлення стабільних ТІ дозволяє побудувати структурний шаблон тексту, який не змінюється навіть під час використання інших інструментів або письма з іншою швидкістю. Така структурна точність є невід'ємною складовою графологічного аналізу. Вона дозволяє виявляти мікроознаки, такі як оцінки нахилу пера та вузли, чутливі до тиску, які є важливими як для ідентифікації автора, так і для високоточного оптичного розпізнавання символів.

З огляду на ці вимоги, важливим етапом є вибір ефективного алгоритму виявлення здатного надійно ідентифікувати ТІ. Оскільки рукописний текст містить багато дрібних деталей важливим критерієм вибору детектора є його здатність зберігати геометричну інваріантність без деградації якості локалізації на різних масштабах. Саме тому класичні алгоритми потребують адаптації або заміни на більш сучасні аналоги, що не використовують апроксимацію зі зниженням роздільної здатності.

Метод масштабно-інваріантних центрально-окружних детекторів (Center Surround Extrema або CenSurE) [86, 87] забезпечує справжній багатомасштабний дескриптор. Вектор ознак створюється з використанням повної просторової роздільної здатності на всіх масштабах піраміди, на відміну від SIFT та SURF, які знаходять екстремуми на субдискретизованих пікселях, що знижує точність на більших масштабах. CenSurE подібний до SIFT та SURF, але деякі ключові відмінності наведено в таблиці 3.1.

Таблиця 3.1 – Порівняльна таблиця основних відмінностей характеристик методів виявлення ознак

	CenSurE	SIFT	SURF
Роздільна здатність	Кожен піксель	Пірамідальна субдискретизація	Пірамідальна субдискретизація
Метод фільтрації країв	Гарріс	Гессіан	Гессіан
Метод екстремумів масштабового простору	Лапласа, центральний оточуючий	Лапласа, DOG	Гессіан, DOB
Обертальна інваріантність	Приблизна	Так	Ні
Просторова роздільна здатність у масштабі	Повна	Субдискретизація	Субдискретизація

CenSurE зберігає зображення з його початковою роздільною здатністю. В той же час SIFT будує обчислювально дорогу піраміду зображення за рахунок розмивання та багаторазової зміни розміру зображення, а SURF покладається на вейвлет-апроксимації Хаара. CenSurE знаходить особливості, масштабуючи самі фільтри. Зазвичай використовуються дворівневі центрально-оточуючі фільтри у формі восьмикутників або зірок.

Бібліотека програмного забезпечення OpenCV внесла зміни до оригінального алгоритму CenSurE і додала його під назвою дескриптор STAR. Хоча бібліотека внесла певні структурні модифікації, основне алгоритмічне перетворення все ще найкраще зобразити через його формальне визначення.

Формально процес виявлення ознак можна виразити як математичне відображення. Якщо ми визначимо наше початкове зображення як I , а результуючий набір ключових точок як K , алгоритм виглядатиме так:

$$F_{CenSuRe}: I \rightarrow K \quad (3.1)$$

Нехай вхідними даними є вхідне зображення з набору даних CVL представлене як двовимірна матриця зображення $I(x,y)$.

Нехай S – дискретна множина шкал фільтрів.

$$S = \{s_1, s_2, \dots, s_n\} \quad (3.2)$$

Нехай V буде тривимірним об'ємом відгуків дворівневих фільтрів:

$$V(x, y, s) = (I * f_s)(x, y) \quad (3.3)$$

де:

$V(x,y,s)$ – значення скалярної відповіді в піксельних координатах (x, y) та масштабі S ;

f_s – дворівневий центральо-окружний фільтр (восьмикутник), масштабований до розміру s .

Ми можемо виразити процес виявлення як вилучення прогресивно менших підмножин точок (x, y, s) з об'єму V .

Перший набір кандидатів C_I створюється зберігаючи лише координати з сильною базовою відповіддю (де абсолютна відповідь фільтра відповідає `responseThreshold` (τ_{res})):

$$C_1 = \{(x, y, s) \mid |V(x, y, s)| \geq \tau_{res}\} \quad (3.4)$$

де:

$|V(x,y,s)|$ – абсолютна величина відгуку фільтра в цій точці;

τ_{res} – параметр `responseThreshold` з OpenCV, що діє як мінімально необхідний контраст.

Далі фільтрується C_I , щоб переконатися, що кожна точка є локальним екстремумом.

Нехай $N_{x,y,s}$ представляє локальну околицю меж, визначених як $N = \text{suppressNonmaxSize}$. Створюється другий набір кандидатів C_2 :

$$C_2 = \{(x, y, s) \in C_1 \mid V(x, y, s) \text{ is a strict extremum in } N_{x,y,s}\} \quad (3.5)$$

де:

$N_{x,y,s}$ – локальна тривимірна околиця точки (x, y) у масштабі s ;

«strict extremum» означає, що значення $V(x,y,s)$ має бути суворо більшим за всі інші точки $N_{x,y,s}$ (локальний максимум) або суворо меншим за всі інші точки (локальний мінімум).

Далі застосовуємо метод відхилення ребер Гарріса/Гессіана. Для кожної точки в C_2 обчислюється матриця Гессіана H .

Нехай $\gamma(H)$ буде обчисленням визначальної траєкторії коефіцієнта кривини:

$$\gamma(H) = \frac{Tr(H)^2}{Det(H)} \quad (3.6)$$

де:

H – 2x2 матриця Гессіана обчислена в точці (x,y,s) ;

$Tr(H)$ – це слід матриці Гессіана (сума незмішаних похідних другого порядку, $L_{xx} + L_{yy}$);

$Det(H)$ – є визначником матриці Гессіана ($L_{xx} L_{yy} - L_{xy}^2$).

Нехай T буде абсолютним порогом для відхилення ребра:

$$T_{line} = \frac{(r_{th}+1)^2}{r_{th}} \quad (3.7)$$

де:

r_{th} – максимально допустиме співвідношення між головними кривинами (власними значеннями H), безпосередньо відображається з параметрів OpenCV (*lineThresholdProjected* та *lineThresholdBinarized*).

Кінцевий вихідний набір ключових точок K визначається як:

$$K = \{(x, y, s) \in C_2 \mid \gamma(H(x, y, s)) < T_{line}\} \quad (3.8)$$

Структурна схема методу наведена на рис 3.1.



Рисунок 3.1 – Структурна схема методу CenSurE

Для методів виявлення ознак важливими характеристиками є: обчислювальна ефективність, ефективне використання пам'яті, висока продуктивність та точність. Перевагами цього методу можна вважати оптимізований підхід до знаходження екстремумів, спочатку використовуючи

Лапласіан на всіх масштабах, а потім здійснюючи крок фільтрації за допомогою методу Гарріса для відкидання кутів зі слабкими відгуками. Він також вважається неймовірно швидким і високоточним тому що знаходить особливості, масштабуючи самі фільтри.

3.3 Метод ідентифікації автора

Для ідентифікації автора аналіз почерку використовує метод індивідуальної особистої ідентифікації. Суть методу полягає у всебічному дослідженні окремих характеристик документа, що досліджується, та порівнянні їх з відомими проявами цих характеристик в інших документах (порівняльний матеріал) [3].

Аналізуючи підходи, що використовуються в літературі для ідентифікації автора на основі аналізу рукописного тексту, можна виділити спільні частини, що використовуються при використанні технологій машинного навчання. Спільні частини включають:

- 1) попередню обробку даних;
- 2) вилучення графологічних ознак;
- 3) ідентифікацію автора.

Відповідно до цих частин побудовано загальну схему підходу до ідентифікації автора на основі зображення рукописного тексту з використанням машинного навчання, яка показана на рис. 3.2.

Першим кроком моделі йде вхідне зображення. Як джерело вхідних даних було використано загальнодоступний набір даних CVL [77]. Цей набір даних містить зображення рукописних текстів німецькою та англійською мовами, взяті з різних літературних творів. Він містить матеріали 310 різних

авторів, кожен з яких надав сім унікальних зразків тексту. Весь набір даних, що охоплює всіх 310 суб'єктів, був використаний для аналізу та навчання.



Рисунок 3.2 – Загальна схема підходу до ідентифікації автора

Моделі машинного навчання вимагають розділення набору даних на такі частини: навчання, валідація та тестування. Для перевірки здатності моделі передбачати автора заданого тексту було використано підхід подвійної валідації. Цей підхід передбачає виконання валідації за допомогою методів пошуку та класифікації.

Підхід пошуку зображень базується на обчисленні відстаней або подібності між зображеннями та на використанні інформативних інваріантних швидко обчислюваних ознак. Методи пошуку отримують запитуване зображення та обчислюють відстань (подібність) до інших зображень у наборі даних. Зображення з найближчими відстанями вважаються схожими та повертаються як результати.

Підхід класифікації зображень базується на призначенні заздалегідь визначених категорій/міток на основі їхнього візуального вмісту [2]. Методи класифікації класифікують немарковані дані та розпізнають шаблони завдяки навчанню на категоризованих/маркованих.

Крок попередньої обробки є важливою частиною підходу до ідентифікації автора, оскільки він виконує кілька цілей одночасно.

Обробка візуальних даних різного кольору має притаманну складність і вимагає значних обчислювальних ресурсів. Щоб пом'якшити ці проблеми, фундаментальний крок попередньої обробки включає перетворення зображення у градації сірого.

Вхідні зображення в наборах даних, таких як CVL, часто мають значні розміри, деякі приклади мають середній розмір 2500×1500 пікселів. Така висока роздільна здатність створює значний виклик для моделей машинного навчання через велику кількість згенерованих токенів. Це особливо проблематично для архітектур, таких як трансформери зору, де обчислювальна складність самоуваги масштабується квадратично з кількістю токенів. Щоб зменшити цю обчислювальну складність, вводиться крок попередньої обробки вилучення ділянок. Це включає поділ зображення на латки фіксованого розміру, зазвичай 32×32 пікселі. Цей розмір латки зазвичай використовується для зображень з дуже високою роздільною здатністю, оскільки він ефективно зменшує загальну кількість токенів, які модель повинна обробити, тим самим керуючи обчислювальною складністю.

Враховуючи, що поля з порожнім простором присутні в зображеннях набору даних CVL, необхідна цільова стратегія вилучення ділянок. Щоб отримати найбільш релевантну інформацію з зображення, було додано додатковий крок попередньої обробки вилучення ознак з використанням CENSURE (також відомий як OpenCV STAR) [87]. Цей крок перетворює дані

зображення на ознаки на екстремумах фільтрів центрального оточення в кількох масштабах, в результаті чого отримуються локальні координати об'єкта. Ці локальні координати об'єкта використовуються як центральні точки для вилучення ділянок розміром 32×32 пікселі.

Для ефективного відділення об'єктів від їхнього фону під час попередньої обробки вхідних даних було використано алгоритм автоматичного встановлення порогу [79]. Цей алгоритм бінаризує зображення у градаціях сірого, вибираючи оптимальний поріг, який максимізує міжкласову дисперсію рівнів сірого. Пікселі нижче цього порогу встановлюються на 0 (чорний), а вище - на 1 (білий), чітко сегментуючи передній план від фону.

Модель попередньої обробки рукописних текстових даних показано на рис. 3.3.

Для кроку ідентифікації автора базовою моделлю машинного навчання було обрано Зоровий трансформер [80], який використовує механізм уваги як основну парадигму навчання.

По суті, ViT побудовано на трансформерному кодувальнику, архітектурі, що складається з двох критичних шарів: багатоголової самоуваги та нейронної мережі прямого зв'язку. Цей кодувальник розроблений для обробки вхідних токенів, генеруючи контекстно-залежні представлення через взаємодію цих шарів. Зокрема, механізм багатоголової самоуваги виконує кілька ключових функцій: він ефективно фіксує складні зв'язки між різними патчами у вхідній послідовності та обчислює зважену суму вбудовування патчів. Це зважування стратегічно підкреслює важливі патчі, інтегруючи як глобальну, так і локальну контекстну інформацію. Доповнюючи це, нейронна мережа прямого зв'язку вводить нелінійність, дозволяючи моделі вивчати складні, нелінійні зв'язки між патчами. Після кожного з цих обчислювальних

етапів виходи проходять через підрівень нормалізації та залишкових зв'язків. Цей підрівень має вирішальне значення для стабілізації та прискорення процесу навчання шляхом нормалізації вхідних даних для кожного шару та зменшення добре відомої проблеми зникнення градієнтів.



Рисунок 3.3 – Модель попередньої обробки рукописних текстових даних

Відповідно до високорівневої моделі було побудовано функціональну модель для ідентифікації автора за допомогою технології машинного

навчання. Ця модель базується на описаних кроках, а саме попередньої обробки та кроках вилучення графологічних ознак та ідентифікації автора.

Запропонована функціональна модель для ідентифікації автора за допомогою машинного навчання показана на рис. 3.4.

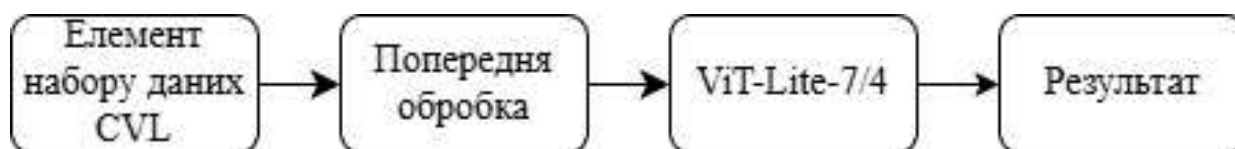


Рисунок 3.4 – Запропонована функціональна модель для ідентифікації автора

Попередньо оброблені вхідні дані подаються в Зоровий трансформер, зокрема в модель навчання малого масштабу, детально описану в [81].

Формалізація роботи запропонованої функціональної моделі здійснюється шляхом побудови математичної моделі, яка описує весь процес.

Нехай X_{in} – вхідне зображення з набору даних CVL. Тоді весь процес зображений на рис 3.4 можна описати функцією F , яка відображає вхідне зображення у вектор ймовірностей класів Y :

$$Y = F(X_{in}) = (M_{ViT} \circ P)(X_{in}) \quad (3.9)$$

де:

P – функція попередньої обробки;

M_{ViT} – модель класифікатора Зоровий трансформер.

В свою чергу функція попередньої обробки представлена на рис 3.3 описується як послідовність кроків P :

$$P = p_5 \circ p_4 \circ p_3 \circ p_2 \circ p_1 \quad (3.10)$$

де:

p_5 – формування тензора

p_4 – бінаризація

p_3 – розбиття на фрагменти

p_2 – виділення ознак

p_1 – градації сірого

Градації сірого (p_1) – перетворення кольорового зображення в інтенсивність яскравості:

$$p_1 = 0.299R + 0.587G + 0.114B \quad (3.11)$$

Виділення ознак (p_2) – метод масштабно-інваріантних центрально-окружних детекторів для виділення ознак:

$$p_2 = CenSurE(p_1) \quad (3.12)$$

Вилучення фрагменту (p_3) – розбиття зображення на фрагменти (патчі). Розмір фрагмента становить 32x32 пікселі:

$$p_3 = \{x_{i,j}\}, \text{ де } x_p \in R^{32 \times 32} \quad (3.13)$$

Бінаризація (p_4) – перетворення кожного патча в бінарний вид (0 або 1) за порогом T :

$$p_4 = \begin{cases} 255, \text{ якщо } x_p > T^* \\ 0, \text{ якщо } x_p \leq T^* \end{cases} \quad (3.14)$$

Формування результату (p_5) – формування тензора P , який подається на вхід нейромережі.

Після попередньої обробки дані потрапляють до ViT-Lite-7/4. Обробку цих даних можна виразити через механізм Self-Attention:

Кожен патч x_p перетворюється у вектор через матрицю ваг E з додаванням позиційного кодування E_{pos} :

$$z_0 = [x_{class}; x_p^1 E; x_p^2 E; \dots x_p^n E] + E_{pos} \quad (3.15)$$

Результат, z_0 , є першим вхідним сигналом для стекових трансформерних кодерів. L стекових кодерів представляють другий компонент архітектури. В ньому $l = 1 \dots L$. Кожен трансформер приймає на вхід ознаки, представлені у вигляді тензора $(n+1, d)$, і видає вихідний сигнал тієї ж розмірності.

На другому кроці мережа вивчає більш абстрактні функції з вбудованих патчів, використовуючи стек L трансформерних кодерів. Це відповідає рівнянням (3.16) та (3.17)

$$z'_l = MSA(LN(z_{l-1})) + z_{l-1} \quad (3.16)$$

$$z_l = MLP(LN(z'_l)) + z'_l \quad (3.17)$$

де:

MSA – багатоголова самоувага;

LN – нормалізація шарів.

Компонент кодера містить механізм багаторівневої уваги (МНА) та дворівневий MLP з нормалізацією шарів та залишковими зв'язками між ними.

Останній крок – активація Softmax на виході голови класифікації (MLP Head):

$$Y = \text{softmax}(MLP(z_L^0)) \quad (3.18)$$

де z_L^0 – стан класифікаційного токена на останньому шарі.

3.4 Експериментальна перевірка метода ідентифікації автора

Експериментальна частина, проведена для перевірки запропонованого підходу, включає кілька частин. Перша частина пов'язана з методами виявлення ознак. Друга частина пов'язана з ідентифікацією автора.

Методи виявлення ознак були застосовані до зображень з бази даних CVL для вилучення важливих ознак. Щоб продемонструвати різницю між базовим методом та запропонованим методом, виявлені області зображень були виділені.

Рис. 3.5 демонструє ключові точки, визначені алгоритмом виявлення ознак SIFT на зразку зображення з набору даних. Виділені області відповідають ознакам, які детектор вважає найбільш значущими.

Рис. 3.6 демонструє ключові точки, визначені алгоритмом виявлення ознак CenSurE на зразку зображення з набору даних. Виділені області відповідають ознакам, які детектор вважає найбільш значущими.

Wird 'ich zum Augenblicke sagen:
 Verweile doch! du bist so schön!
 Dann magst du mich in Fesseln schlagen,
 Dann will ich gern zu Grunde gehn!
 Dann mag die Todtenglocke schallen,
 Dann bist du meines Dienstes frey,
 Die Uhr mag stehn, der Seiger fallen,
 Es sey die Zeit für mich verby!

Рисунок 3.5 – Демонстрація виявлення ознак за допомогою базового методу

Wird 'ich zum Augenblicke sagen:
 Verweile doch! du bist so schön!
 Dann magst du mich in Fesseln schlagen,
 Dann will ich gern zu Grunde gehn!
 Dann mag die Todtenglocke schallen,
 Dann bist du meines Dienstes frey,
 Die Uhr mag stehn, der Seiger fallen,
 Es sey die Zeit für mich verby!

Рисунок 3.6 – Демонстрація виявлення ознак за допомогою запропонованого методу

Для кількісної оцінки ефективності різних методів виявлення ознак було проведено порівняльний аналіз, зосереджений на обчислювальній ефективності та виході ознак. Кожен метод був протестований на всьому наборі даних для обробки кожного елемента за однакових умов.

Під час експерименту вимірювалися ключові показники ефективності: середній час визначення ключових точок на елемент, загальний час, необхідний для повного набору даних, та загальна кількість виправлень, вилучених з виявлених ключових точок. Результати цієї оцінки представлені в таблиці 3.2. Ці дані забезпечують пряме порівняння, підкреслюючи компроміси між швидкістю обробки та щільністю виявлення ознак кожного методу.

Таблиця 3.2 – Порівняння методів виявлення ознак

Метод	Час/елемент, с	Загальний час, с	Заг. кіл. патчів
Baseline (SIFT)	0.8665	1390.93	2555409
SURF	0.4633	744.50	3209245
ORB	0.1662	267.28	801000
Proposed (CenSurE)	0.52	848.44	4368819

Друга частина, пов'язана з ідентифікацією автора, включає навчання кількох моделей. Дві моделі відрізняються складом набору даних та підходом до верифікації. Склад наборів даних був наступним:

1. Набір даних був розділений на дві рівні половини. Першу половину було виділено для навчання та валідації моделі, тоді як другу половину було зарезервовано виключно як резервний тестовий набір.

2. Було реалізовано стратегію розподілу для кожного суб'єкта. Для кожного з 310 суб'єктів сім зразків рукописного тексту було розподілено наступним чином: два зразки для навчання, один для валідації та решта чотири для тестування.

Щоб забезпечити відтворюваність результатів та рівномірне перетасування даних по всіх модулях моделі, було використано фіксоване випадкове початкове значення. Щоб уникнути перенавчання, було реалізовано механізм ранньої зупинки, який зупиняє навчання, якщо функція втрат не покращується протягом 10 епох.

Параметри, вибрані під час налаштування моделі, представлені в таблиці 3.3.

Таблиця 3.3 – Параметри конфігурації моделі

Параметр	Знач	Опис
lr	0.0005	швидкість навчання
num-epochs-warmup	5	епокси розігріву
batch-size	128	розмір пакета
num-workers	4	працівники PyTorch
num-epochs-patience	10	кількість епох, після яких навчання зупиниться, якщо втрата валідації більше не покращиться
num-epochs	60	кількість епох

Щоб переконатися, що моделі були навчені правильно, показники втрат (`loss_val`, `loss_train`) та точності (`accuracy_val`, `accuracy_train`) контролювалися на навчальних та валідаційних вибірках.

Рис. 3.7–3.10 ілюструють результати розробки моделей для підходів класифікації (рис. 3.7, 3.8) та пошуку (рис. 3.9, 3.10).

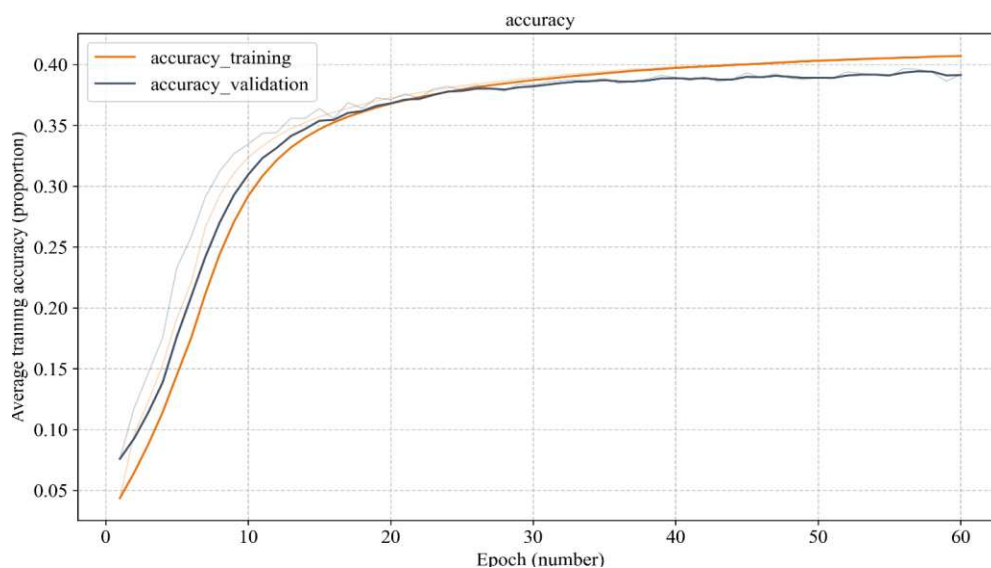


Рисунок 3.7 – Точність валідації та навчання запропонованого підходу на тестовому наборі даних для класифікаційного підходу

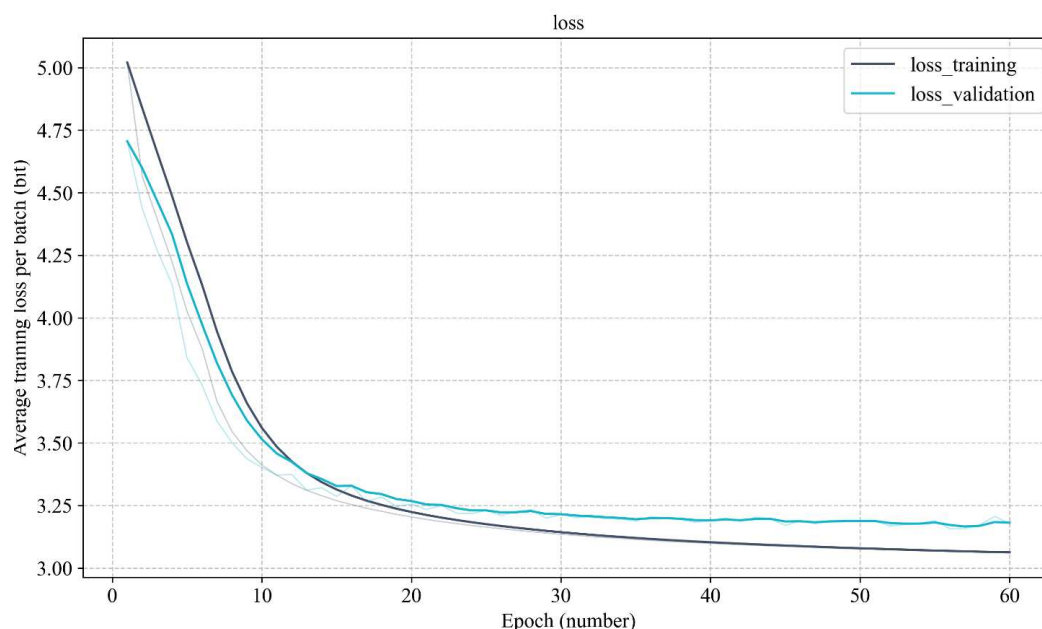


Рисунок 3.8 – Втрати валідації та навчання запропонованого підходу на тестовому наборі даних для класифікаційного підходу

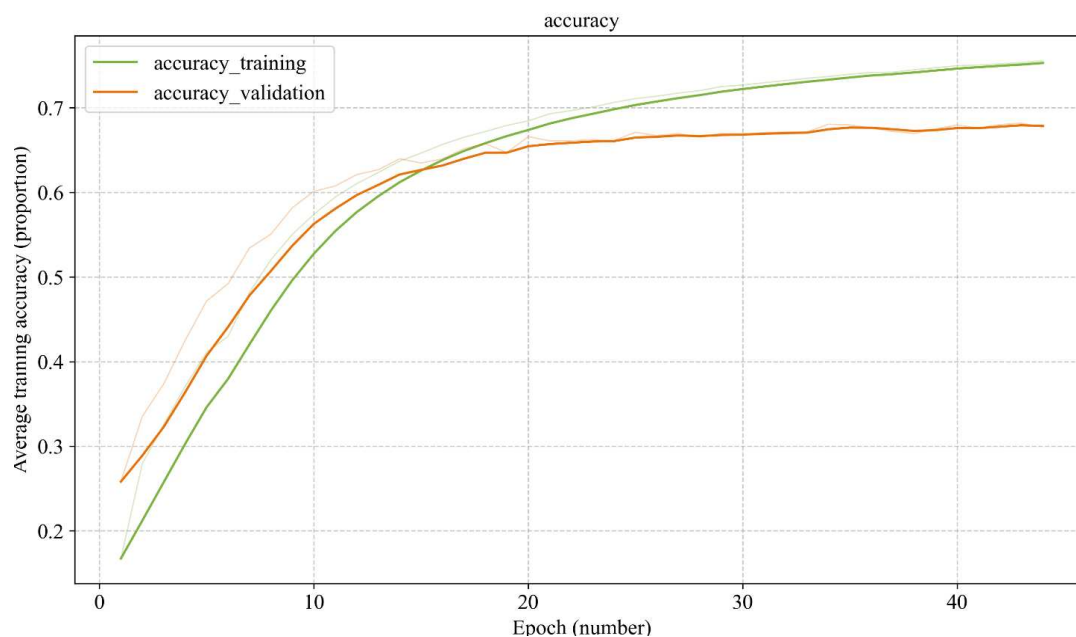


Рисунок 3.9 – Точність валідації та навчання запропонованого підходу на тестовому наборі даних для підходу пошуку

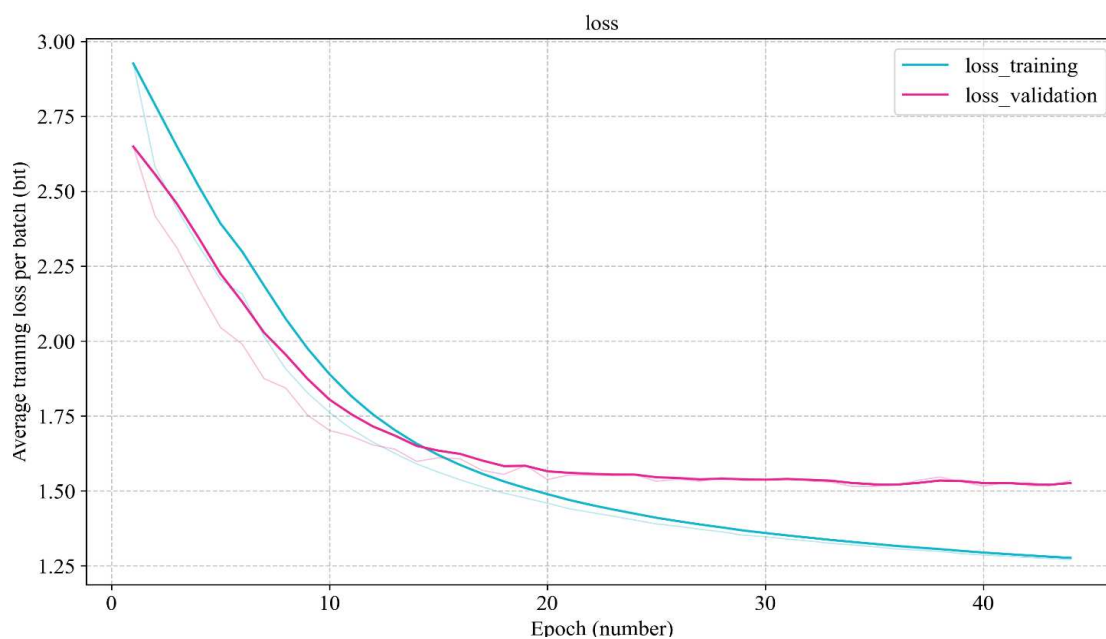


Рисунок 3.10 – Валідація та втрати навчання запропонованого підходу на тестовому наборі даних для підходу пошуку

В основному експерименті оцінювалася ефективність моделей для двох підходів до ідентифікації автора. Для класифікаційного підходу ключовим показником була точність моделі, тоді як для підходу пошуку була точність ідентифікації на основі квадрата евклідової відстані між векторами ознак. З метою комплексної валідації також було розраховано низку додаткових метрик, що стосуються оцінки моделей машинного навчання.

Результати тестування підходу класифікації, а саме, основні показники ефективності моделі представлені в таблиці 3.4.

Таблиця 3.4 – Показники ефективності класифікаційного підходу

Метрика	k	Базовий	Запропонований
Accuracy	1	0.98	0.98
	2	0.99	0.99
	3	0.99	0.99
	4	0.99	0.99
	5	0.99	0.99
	6	0.99	0.99
	7	0.99	0.99
	8	0.99	0.99
	9	0.99	0.99
	10	0.99	0.99

Результати тестування підходу пошуку, а саме, основні показники ефективності моделі представлені в таблиці 3.5.

Таблиця 3.5 – Показники ефективності пошуку даних з використанням квадрата евклідової відстані

Метрика	k	Базовий	Запропонований
Hard Top K	1	0.95	0.96
	2	0.89	0.91
Soft Top K	1	0.95	0.96
	2	0.96	0.97
	3	0.96	0.97
	4	0.97	0.97
	5	0.97	0.98
mAP		0.87	0.89

Показники продуктивності методів виявлення ознак наведені в таблиці 3.2 вказують, що запропонований підхід на основі CenSurE досягає найвищої кількості виявлених ознак, зі значенням понад 4,3 мільйона. Це значно більше, ніж SIFT (~2,56 млн), SURF (~3,21 млн) та ORB (~0,80 млн). Хоча ORB є найшвидшим методом, виявляючи ознаки в середньому за 0,17с на елемент, його значно менша кількість патчів свідчить про те, що він фіксує набагато менше деталей почерку. Запропонований метод є збалансованим варіантом. Він обробляє кожен елемент в середньому за 0,52с (приблизно на 40% швидше, ніж базовий показник SIFT (0,87 с)), при цьому виявляючи приблизно на 71% більше інформативних ключових точок. Загалом, ці результати свідчать про те, що CenSurE забезпечує більш детальне виявлення ознак для завдань ідентифікації автора без надмірних обчислювальних витрат.

Показники продуктивності запропонованого підходу ідентифікації автора, наведені в таблицях 3.4 та 3.5, вказують на покращення ідентифікації

автора як у завданнях класифікації, так і в завданнях пошуку. Основне покращення спостерігається в завданнях пошуку. Такі показники, як Hard Top K, Soft Top K та середня середня точність (mAP), демонструють збільшення порівняно з базовим підходом.

Переваги запропонованого підходу також можна побачити в порівнянні з іншими відомими підходами. На відміну від методів, які залежать від складних, багатоетапних конвеєрів та ручних функцій, таких як дескриптори RootSIFT [22], або комбінації глибоких та ручних функцій [26], запропонований підхід використовує більш спрощений та обчислювально ефективний метод вилучення ознак. Використання таких методів вилучення ознак також допомагає уникнути чутливого налаштування параметрів та усуває необхідність балансування та комбінування ознак.

Крім того, запропонований підхід усуває відомі недоліки інших архітектур, таких як ті, що базуються на CNN. Відомі недоліки моделей ResNet-50 [24] та FragNet [31] включають обмежену здатність захоплювати глобальний контекст та труднощі з обробкою різноманітних стилів почерку. Оскільки запропонований підхід базується на ViT, він використовує механізм самоуваги для фіксації довгострокових залежностей та зв'язків між латками. Ці характеристики не лише усувають відомі недоліки архітектур CNN, але й слугують перевагами над методами на основі ділянок [26, 27].

Нарешті, оскільки ViT працює безпосередньо із зображеннями, запропонований підхід показує чудові результати без використання допоміжних завдань [30] або складної попередньої обробки сегментації [32], таким чином уникаючи потенційних ускладнень, пов'язаних з використанням додаткових підготовчих кроків.

Важливо також зазначити, що запропонований підхід підтримує високі значення точності для класифікаційного підходу. Це, разом із покращеними

значеннями для підходу пошуку, дозволяє зробити висновок, що запропонований підхід демонструє високу ефективність у завданнях ідентифікації автора.

Отже, запропонований підхід до ідентифікації автора на основі рукописного тексту сприяє підвищенню ефективності автоматизованих систем аналізу почерку. Він також служить цінним інструментом для графологів та судово-медичних експертів з документів шляхом оптимізації їхньої професійної діяльності.

3.5 Висновки за третім розділом

У цьому розділі було розроблено та експериментально досліджено удосконалений метод ідентифікації автора за рукописним текстом, зосереджений на підвищенні точності та обчислювальної ефективності. Як метод виявлення ознак було обґрунтовано та використано метод масштабно-інваріантних центрально-окружних детекторів, а також сформоване його формальне та структурне визначення.

1. За результатами аналізу існуючих методів ідентифікації автора виявлено, що підходи на основі згорткових нейронних мереж недостатньо враховують глобальний контекст і довгострокові залежності в зображенні, є чутливими до якості попередньої обробки та потребують значних обчислювальних ресурсів. На цій основі обґрунтовано та сформульовано задачу побудови удосконаленого методу, що поєднує детектор ключових точок для вилучення інформативних ділянок із архітектурою Зоровий трансформер.

2. Обґрунтовано та обрано метод виявлення ознак на основі масштабно-інваріантних центрально-окружних детекторів (CenSurE), який, на відміну від

SIFT та SURF, зберігає повну просторову роздільну здатність на всіх масштабах і забезпечує геометричну інваріантність; для методу виконано формальне та структурне визначення. Експериментальне порівняння показало, що CenSurE вилучає на 71% більше інформативних ділянок та скорочує середній час попередньої обробки одного зображення на 40%.

3. Розроблено удосконалений метод ідентифікації автора, що інтегрує детектор CenSurE для вилучення ознак та архітектуру глибокого навчання Зоровий трансформер, побудовано функціональну модель ідентифікації автора та проведено її математичну формалізацію. Передбачено два підходи до ідентифікації, класифікаційний та пошуковий, що дозволяє подвійну валідацію результатів.

4. Проведено експериментальну перевірку на наборі даних CVL, яка підтвердила ефективність підходу. У класифікаційному підході досягнуто високої точності, а в пошуковому досягнуто підвищення показників ідентифікації Hard Top K та Soft Top K до 0,96 та середньої точності mAP до 0,89 порівняно з базовим методом. Отримані результати підтверджують успішне вирішення поставлених завдань щодо підвищення точності та обчислювальної ефективності ідентифікації автора.

Ці результати свідчать про успішне вирішення поставлених завдань, що створює надійну основу для подальшого розвитку автоматизованих систем графологічного аналізу в області ідентифікації автора.

Перелік використаних джерел у даному розділі наведено у повному переліку використаних джерел під номерами: 4–87.

4 МОДЕЛЬ ІНФОРМАЦІЙНОЇ ТЕХНОЛОГІЇ АВТОМАТИЗОВАНОГО АНАЛІЗУ РУКОПИСНОГО ТЕКСТУ

У четвертому розділі розглянуто модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту. Основну увагу приділено аналізу існуючих систем аналізу почерку та розробці цілісної архітектури, яка поєднує удосконалені методи ідентифікації автора та визначення особистості, для забезпечення комплексного графологічного аналізу. У процесі роботи побудовано функціональну та структурну моделі інформаційної технології автоматизованого графологічного аналізу рукописного тексту, а також архітектуру інформаційної системи, що реалізує цю технологію. Окрему увагу приділено формалізації задачі та опису складових частин інформаційної технології. Запропонована модель інформаційної технології є основою для подальшого розвитку систем поведінкової біометрії та забезпечує підвищення ефективності автоматизованого графологічного аналізу рукописного тексту.

4.1 Постановка задачі автоматизованого аналізу рукописного тексту

Галузь програмного забезпечення для аналізу почерку чітко розділена на дві групи, це професійні судово-медичні системи (Forensic Grade) та комерційні/наукові інструменти для психологічного профілювання. Пріоритетним напрямком розвитку галузі є впровадження автоматизованої системи розпізнавання почерку [14]. Інтеграція спеціалізованого програмного забезпечення дозволяє мінімізувати вплив людського фактору та забезпечує перехід до математично обґрунтованої перевірки автентичності реквізитів.

Такі системи забезпечують високу точність ідентифікації автора в режимі реального часу, дозволяють оперативно перевіряти великі обсяги документації та ефективно виявляти підробки, зроблені за допомогою сучасних технічних засобів, таких як плоттери або авторучки.

Ще одне критичне застосування полягає в обробці історичних архівів, де ручний аналіз рукописів залишається суттєвим вузьким місцем. Рукописи можуть бути дуже різними типами текстів: листи відомих особистостей, родичів чи предків, романи чи вірші, щоденники чи юридичні звіти [88]. Робота з історичними документами включає вирішення проблем транскрибування тексту, а також атрибуції анонімних рукописів, ідентифікації авторів листів тощо. Крім того, застосування аналізу почерку поширюється на поведінкові науки. Сектор психології та управління персоналом зацікавлений у методах швидкої оцінки психологічного портрета кандидатів. Аналіз вважається інструментом для виявлення м'яких навичок, оцінки стійкості до стресу, схильності до лідерства або навіть деструктивної поведінки. Інтеграція автоматизованого аналізу почерку в цій галузі може забезпечити оцінку та зворотний зв'язок у режимі реального часу, розширюючи його доступність та зручність використання в різних умовах [89].

На міжнародному рівні існують системи, що включають як функції автоматичного розпізнавання образів, так і інтерактивні інструменти, що допомагають судово-медичним експертам. FISH вважається предком автоматизованих систем аналізу почерку. Розроблена Федеральним управлінням кримінальної поліції Німеччини у 1990-х роках, вона була адаптована та широко використовувалася Секретною службою США для аналізу листів із погрозами президентам та високопосадовцям [90]. Як системи підтримки рішень, вони не роблять остаточного висновку, а лише звужує пошук експерта і працює як рекомендаційна система. У 2025 році Секретна

служба США оголосила про модернізацію FISH з використанням найновіших алгоритмів штучного інтелекту для пришвидшення пошуку [91]. Незважаючи на потенційну модернізацію, FISH залишається інструментом підтримки рішень, обмеженим своєю оригінальною архітектурою, зосереджуючись виключно на зіставленні образів без розуміння поведінкового наміру, що стоїть за написанням.

Як еволюція FISH була розроблена нова система аналізу почерку та ідентифікації автора під назвою WANDA [92]. Клієнт-серверна платформа WANDA це відкрита система для судово-медичної експертизи почерку та підпису, яка підтримує інтеграцію сторонніх модулів та різних систем письма через архітектуру плагінів. Програма дозволяє автоматизувати збір зразків, розпізнавання алографів, вилучення ознак та ідентифікацію особи на основі оцифрованих даних, використовуючи стандарт WandaML XML для опису документів. Використовуючи фільтри імпорту зображень (зокрема, через програмне забезпечення IBIS) та зручний графічний інтерфейс для експертних нотаток, платформа поєднує традиційні судово-медичні методи з інноваційними розробками незалежних дослідників в єдиному робочому середовищі.

Центр передового досвіду в аналізі документів (CEDAR) при Університеті Буффало розробляє систему судово-медичного аналізу документів, відому як CEDAR-FOX [93]. CEDAR-FOX реалізує повний цикл аналізу: від видалення фону до розрахунку коефіцієнта правдоподібності (LLR). Унікальною особливістю є можливість роботи як з повними сторінками тексту, так і з короткими записами. Система аналізує як макроознаки (ентропія, поріг), так і мікроознаки (структура символів). Вона також включає інструменти для перевірки підпису та ідентифікації автора за обмеженою кількістю тексту. Хоча CEDAR-FOX чудово працює для аналізу мікроознак та

розрахунків коефіцієнта правдоподібності, він є суто детерміністичним інструментом, який розглядає почерк як фізичний слід без урахування психологічної проєкції. Центр статистики та застосувань у судових доказах (CSAFE) розробляє пакет з відкритим кодом для мови R. Він використовує розкладання графів райдужних трикутників та стохастичну кластеризацію (k-середніх) для створення груп гліфів. Ймовірність авторства розраховується на основі частоти потрапляння елементів почерку в ці кластери. Цей підхід значною мірою спирається на статистичну частоту «кластерів», яку можна легко спотворити через невеликі розміри вибірки або спроби автора замаскувати свій почерк.

В Україні судова експертиза почерку регулюється Кримінально-процесуальним кодексом та інструкціями Міністерства юстиції. Основними суб'єктами є державні спеціалізовані установи, такі як Київський науково-дослідний інститут судових експертиз (КНДІСЕ). Хоча основа експертизи, що проводиться українськими експертами, залишається традиційною (порівняльний аналіз загальних та індивідуальних ознак експертом), інструменти активно впроваджуються. Зокрема, використання відеоспектральних компараторів та програмного забезпечення Regula Forensic Studio є стандартом для технічної експертизи документів та аналізу підписів на наявність технічної підробки. З іншого боку, використання порівняльного аналізу загальних та індивідуальних ознак є трудомістким та не має прогностичної сили.

Що стосується комерційних/наукових інструментів для психологічного профілювання, то існує система AvgMISC, яка використовує гібрид SVM та CNN для аналізу офлайн-рукописного письма різними мовами (англійською, арабською, китайською) [94]. Автори стверджують про кореляцію результатів з психологічними тестами, але ці дослідження часто мають обмежені вибірки.

Також існує платформа Nyskel, яка пропонує API для класифікації зображень, включаючи модель ідентифікатора особистості рукописного письма. Вона позиціонується як інструмент для HR, маркетингу та освіти, що забезпечує оцінку впевненості моделі на наявність певної риси.

Беручи до уваги поточний стан, стає зрозуміло, що сфера розділена між детермінованими судово-медичними інструментами, які розглядають рукопис виключно як фізичний слід, та гнучкими, але ненадійними психологічними моделями. В цій роботі пропонується інформаційна технологія для автоматизованого аналізу рукописного письма, яка поєднує судово-медичну експертизу з поведінковими висновками.

Метою роботи є розробка моделі інформаційної технології для автоматизованого графологічного аналізу рукописного тексту, яка забезпечує комплексну ідентифікацію автора та його психологічне профілювання.

Відповідно до поставленої мети були поставлені такі завдання:

1. Побудувати функціональну модель інформаційної технології для автоматизованого графологічного аналізу рукописного тексту.
2. Виконати формалізацію роботи шляхом побудови математичної моделі.
3. Побудувати архітектуру інформаційної системи, що реалізує інформаційну технологію для автоматизованого графологічного аналізу рукописного тексту.

4.2 Інформаційна технологія автоматизованого графологічного аналізу рукописного тексту

Запропоновано наступну загальну схему інформаційної технології для автоматизованого аналізу почерку, рис. 4.1. У центрі системи знаходиться

блок обробки, який взаємодіє з кількома потоками даних: вхідні параметри включають дані ідентифікації користувача та зображення рукописного тексту; управління та методологічна підтримка здійснюється за допомогою використання методів ідентифікації автора, визначення психологічних характеристик особистості та введення цифрових водяних знаків для захисту інформації. Ресурсна база системи ґрунтується на безпечному файловому сховищі та сервері даних користувачів. Результатом роботи технології є вихідні продукти, а саме: детальний звіт аналізу, список виявлених характеристик почерку та масив вибраних ознак для подальшої обробки.

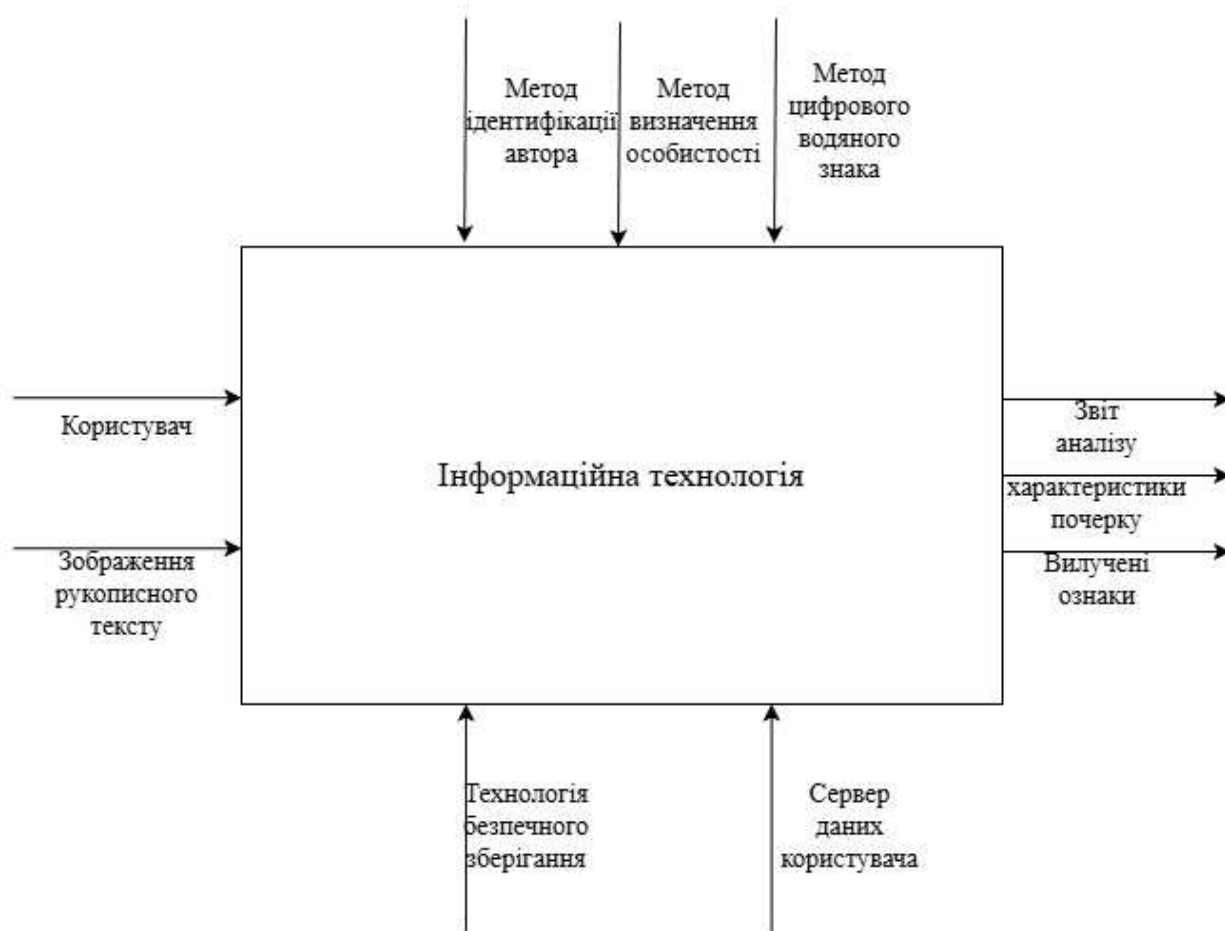


Рисунок 4.1 – Загальна схема інформаційної технології

За характером результатів запропонована інформаційна технологія функціонує як рекомендаційна система.

Процес автоматизованого аналізу почерку пропонується здійснювати за етапами, представленими на схемі декомпозиції інформаційної технології рис. 4.2. Процес починається з автентифікації користувача та завантаження графічних даних, після чого йде етап безпечного зберігання з використанням стеганографічних методів (цифрових водяних знаків). Ключова аналітична обробка зображення поділяється на два паралельні вектори: визначення авторства документа та виявлення особистих характеристик особи на основі аналізу її почерку. Заключний етап включає агрегацію отриманих результатів та формування вихідних звітів на основі вилучених ознак та характеристик.

Доцільно детально розглянути етапи, показані на рис. 4.2, які лежать в основі функціонування технології:

1) Автентифікація/Авторизація користувача: Початковий етап перевірки прав доступу суб'єкта до системи, який здійснюється за допомогою даних з сервера користувача.

2) Зберігання та нанесення водяних знаків: Отримані зображення рукописного тексту проходять процедуру захисту шляхом нанесення цифрових водяних знаків та розміщуються в захищеному файловому сховищі.

3) Визначення авторства документа: Використання методів ідентифікації автора для встановлення або підтвердження особи автора поданого рукопису. Базовий метод, що використовується для ідентифікації автора, детально описаний у роботі [2] та використовує Зоровий трансформер.

4) Визначення рис характеру людини: Психологічний та графологічний аналіз тексту за допомогою методів визначення особистості для формування психологічного портрета. Базовий метод, що використовується для виявлення

особистості, детально описаний у роботі [4] та використовує модель рис особистості Велика п'ятірка.

5) Формування звіту: Заключний етап, на якому синтезуються дані з усіх модулів та формуються кінцеві результати: аналітичний звіт, список характеристик почерку та масив вилучених ознак.

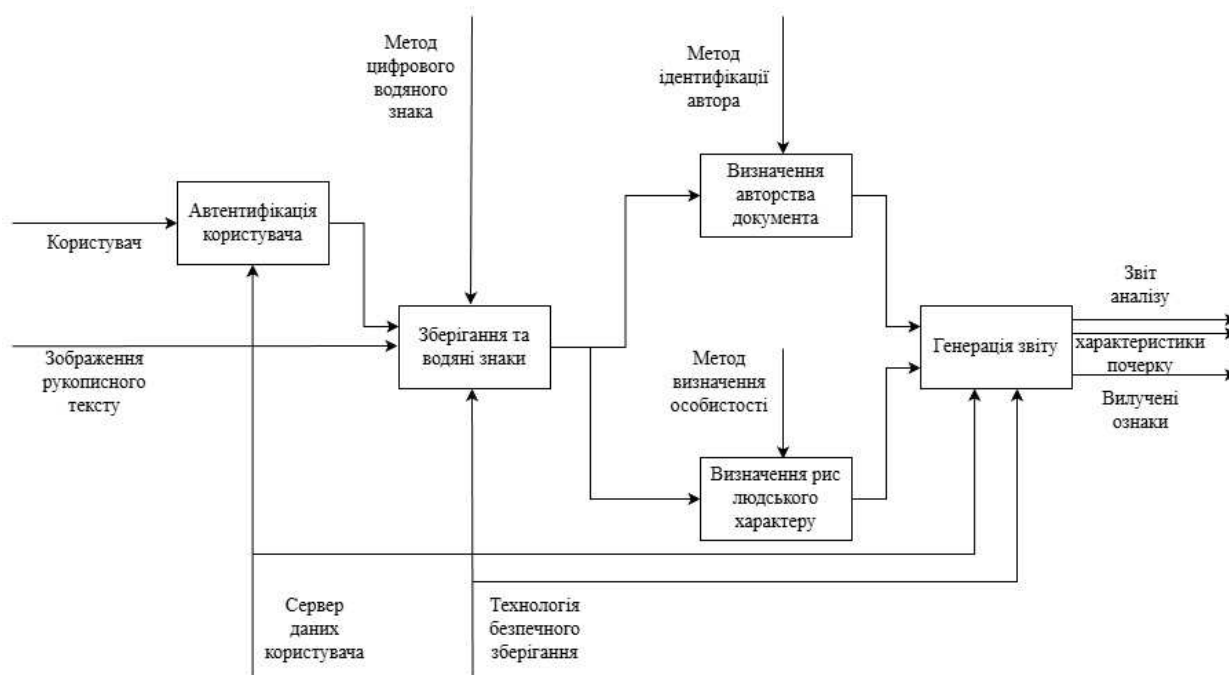


Рисунок 4.2 – Схема декомпозиції інформаційної технології

Формалізація роботи запропонованої інформаційної технології здійснюється шляхом побудови математичної моделі, яка описує весь процес – від збору та зберігання вхідних зображень до формування звітів на основі алгоритмів аналізу рукописного тексту.

Систему S можна представити як кортеж наборів вхідних даних, вихідних даних, механізмів керування та ресурсів:

$$S = \{X, Y, M, R, F\} \quad (4.1)$$

де:

X – набір вхідних даних;

Y – набір вихідних результатів;

M – набір методів (керуючих впливів);

R – набір ресурсів;

F – набір функцій перетворення.

Представимо вхідні дані (X) як:

$$X = \{x_{user}, x_{img}\} \quad (4.2)$$

де:

x_{user} – дані ідентифікації користувача;

x_{img} – цифрові зображення рукописного тексту.

Представимо методи та алгоритми (M) як:

$$M = \{m_{wm}, m_{id}, m_{pi}\} \quad (4.3)$$

де:

m_{wm} – метод цифрового водяного знаку;

m_{id} – метод ідентифікації автора;

m_{pi} – метод визначення особистості.

Представимо ресурси (R) як:

$$R = \{r_{srv}, r_{store}\} \quad (4.4)$$

де:

r_{srv} – сервер даних користувача;

r_{store} – безпечне сховище файлів.

Представимо вихідні результати (Y) як:

$$Y = \{y_{rep}, y_{char}, y_{feat}\} \quad (4.5)$$

де:

y_{rep} – звіт аналізу;

y_{char} – характеристики почерку;

y_{feat} – вилучені ознаки.

Процес перетворення вхідних даних у вихідні описується композицією функцій $f_i \in F$ що відповідають блокам на діаграмі декомпозиції.

Функція автентифікації (f_1) – перевіряє користувача за допомогою даних сервера:

$$a = f_1(x_{user}, r_{srv}) \quad (4.6)$$

де:

a – токен доступу.

Функція зберігання та маркування (f_2) – обробляє зображення за допомогою водяних знаків та зберігає їх:

$$d_{sec} = f_2(x_{imgr}, a, m_{wm}, r_{store}) \quad (4.7)$$

де:

d_{sec} – захищені дані.

Функція визначення авторства (f_3) – аналізує захищені дані для визначення автора:

$$res_{id} = f_3(d_{sec}, m_{id}) \quad (4.8)$$

де:

res_{id} – результат ідентифікації.

Функція ідентифікації рис характеру (f_4) – аналізує захищені дані для визначення психотипу:

$$res_{traits} = f_4(d_{sec}, m_{pi}) \quad (4.9)$$

де:

res_{traits} – виявлені психологічні риси.

Функція генерації звітів (f_5) – агрегує результати аналізу та отримує доступ до ресурсів для генерації кінцевих вихідних даних.

$$Y = f_5(res_{id}, res_{traits}, r_{srv}, r_{store}) \quad (4.10)$$

Підставляючи проміжні функції, загальну роботу інформаційної технології можна описати оператором Φ :

$$Y = \Phi(X, M, R) = f_5(f_3(\Omega, m_{id}), f_4(\Omega, m_{pi}), r_{srv}, r_{store}) \quad (4.11)$$

де Ω - результат захисту даних:

$$\Omega = f_2(x_{imgr}, f_1(x_{user}, r_{srv}), m_{wm}, r_{store}) \quad (4.12)$$

На основі запропонованої концептуальної та математичної моделей було розроблено інформаційну систему для автоматизованого графологічного аналізу почерку. Система реалізує повний цикл обробки даних та роботи: від збору та зберігання вхідних зображень до формування звітів на основі алгоритмів аналізу почерку.

4.3 Інформаційна система автоматизованого графологічного аналізу рукописного тексту

Розробка архітектури для автоматизованого аналізу вимагає балансу між високопродуктивними обчисленнями та суворим дотриманням нормативних вимог. Спираючись на проведений і представлений в таблиці 1.5 аналіз технологій безпечного зберігання було вирішено будувати архітектуру системи на базі технології хмарних обчислень. В якості провайдера хмарних обчислень був обраний Amazon Web Services (AWS). Основними перевагами цього провайдера є велика кількість послуг і технологій, безпека та відповідність вимогам і сертифікаціям (особливо в сфері зберігання даних), надійна безперебійна робота.

На основі загальної схеми інформаційної технології для автоматизованого графологічного аналізу рукописного тексту зображеної на рис. 4.1 була розроблена схема інформаційної системи автоматизованого графологічного аналізу рукописного тексту з використанням хмарного провайдера AWS. Загальна схема такої системи представлена на рис. 4.3.

Система налічує такі складові як Представлення та ідентифікація, API та оркестрування, Обробка, Зберігання та метадані.

Перше, що бачить перед собою користувач інформаційної системи це Інтерфейс користувача. AWS дозволяє розгортати веб додатки з використанням таких сучасних технологій як React або TypeScript в хмарі. На схемі рис. 4.3 це відображено за допомогою рівня Представлення та ідентифікації. Він відповідає за розгортання за допомогою технології AWS Amplify або комбінації Amazon S3 та CloudFront. Ці технології дозволяють розповсюджувати інтерфейс користувача через глобальну мережу доставки контенту (CDN). Цей підхід гарантує, що конфіденційні дані завжди передаються через HTTPS, зберігаючи при цьому низьку затримку. Центральним елементом цього рівня є Amazon Cognito. Він відповідає за авторизацію та автентифікацію користувачів за допомогою багатфакторної автентифікації (MFA). Функціональність Cognito включає такі обов'язки як вхід; керування пулами ідентифікації та пулами користувачів для видачі JSON Web Tokens (JWT). Ці токени необхідні для авторизації кожного запиту до серверних служб. Це гарантує, що лише перевірені користувачі можуть взаємодіяти з базовими структурами даних. За рахунок цього також виконуються вимоги контролю доступу, включаючи вимоги таких відомих стандартів безпеки як Закон про перенесення та підзвітність медичного страхування (HIPAA) або Інформаційні служби кримінального правосуддя (CJIS).

Запити з інтерфейсу користувача в клієнт-серверних додатках зазвичай йдуть на API. На схемі це відображено за допомогою рівня API та оркестрація. Оркестрація на цьому рівні відіграє головну роль. Вона здійснюється за допомогою AWS Step Functions. Вони виступають в ролі основного механізму логіки для всієї системи. Такий підхід дозволяє використовувати роз'єднану

архітектуру, що відрізняється від традиційних монолітних програм. Цей підхід називається безсерверний кінцевий автомат. В такій архітектурі кожна фаза аналізу, від початкового завантаження до остаточного звіту, розглядається як окремий крок.

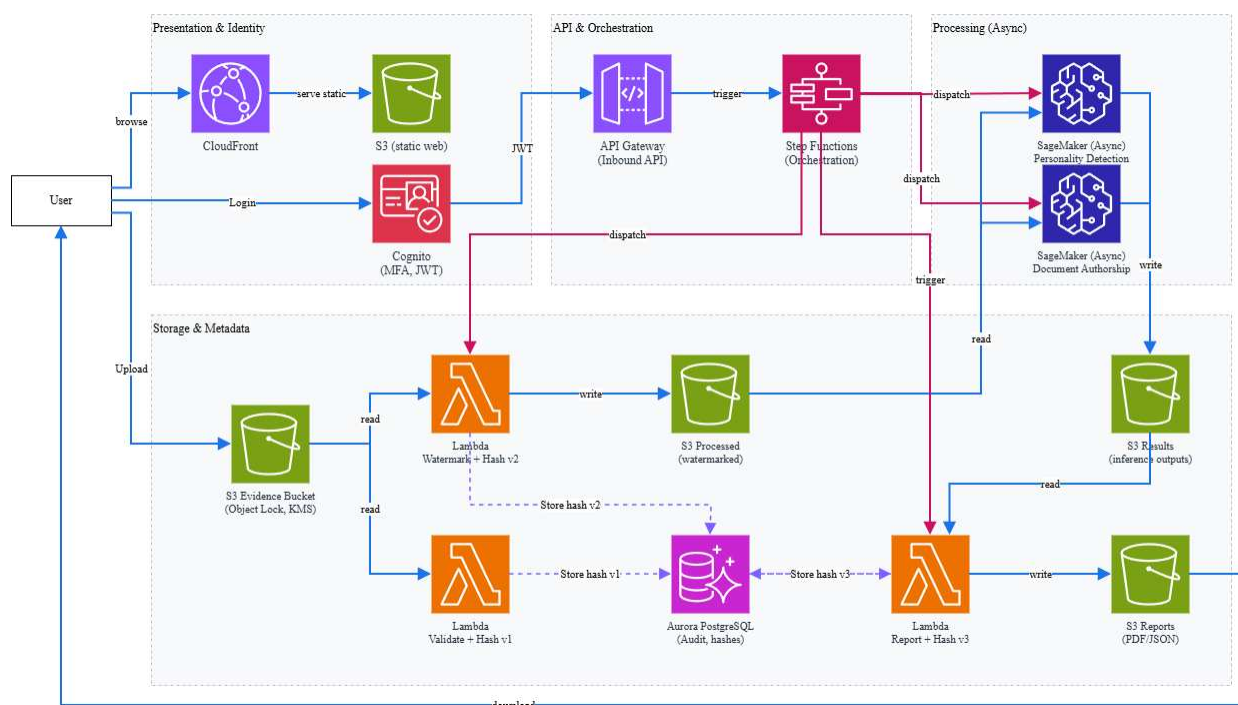


Рисунок 4.3 – Загальна схема побудови інформаційної системи

Коли користувач завантажує документ або зображення, Step Functions запускає послідовність подій. Ця послідовність починається з процесу нанесення цифрового водяного знака. Він виконується функцією AWS Lambda, що працює в середовищі Python. Такий підхід дозволяє реалізувати та використовувати один або декілька методів визначених у проведеному і представленому в таблиці 1.4 аналізі. Функція вбудовує надійний цифровий підпис або водяний знак у зображення, щоб запобігти несанкціонованому розповсюдженню та позначити файл як прийнятий.

Для підтримки ланцюга зберігання система обчислює криптографічні хеші (зазвичай з використанням алгоритму SHA-256). Наприклад хеш генерується як для оригінального файлу одразу після отримання так і на файл після нанесення водяного знака. Ця стратегія подвійного хешування гарантує, що цілісність файлів можна перевірити в будь-який момент життєвого циклу. Схема організації ланцюга зберігання доказів для інформаційної системи представлена на рис. 4.4.

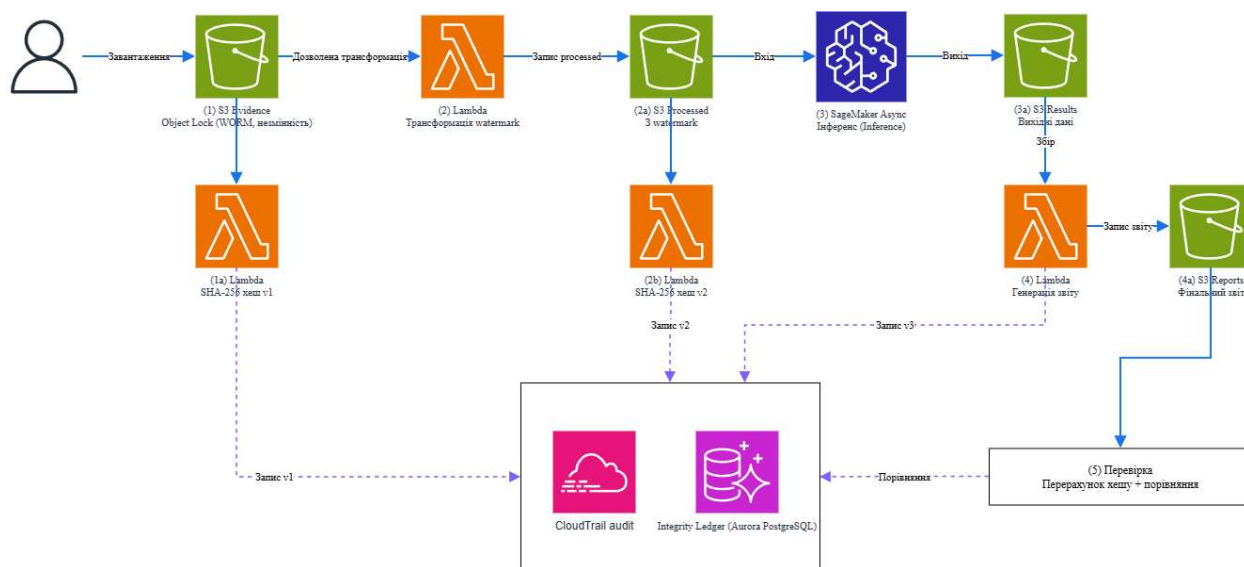


Рисунок 4.4 – Схема організації ланцюга зберігання доказів

Окрім таких процесів як нанесення водяного знака та обчислення хешів завантажені файли мають бути збережені. На схемі рис 4.3 це відображено за допомогою рівня Зберігання та метадані. Стратегія зберігання розроблена на основі принципу незмінності. Amazon S3 служить основним сховищем для об'єктів. Додатково доступна технологія S3 Object Lock з режимом відповідності. Ця функція запобігає від видалення або перезаписування даних протягом заздалегідь визначеного періоду зберігання. Заборона поширюється на всіх користувачів, включаючи тих, хто має права адміністратора. Збережені

дані додатково захищені службою AWS Key Management Service (KMS). Вона використовує ключі для шифрування, що гарантує нечитабельність сховища даних без певного криптографічного матеріалу. Паралельно зі сховищем об'єктів використовується реляційний сервер метаданих Amazon Aurora (PostgreSQL). Особливістю Aurora є здатність обробляти складні журнали аудиту та реляційні зв'язки між діями користувача, метаданими зображень та результатами висновків машинного навчання. Ця реляційна структура допомагає зі звітністю і може включати реконструкцію часової шкали подій на основі різноманітних точок даних.

Після завантаження зображення користувач може запустити процес аналізу зображення. Оркестратор відповідає за запуск цього процесу через запит. На схемі рис 4.3 це відображено за допомогою рівня Обробка. Цей рівень відповідає за графологічний аналіз, для якого використовуються методи на основі машинного навчання (рис. 2.9, рис. 3.4). Моделі на основі машинного навчання є обчислювально вимогливими і цей аспект потрібно враховувати в архітектурі системи. Для роботи з моделями машинного навчання Amazon надає технологію SageMaker Asynchronous Inference. Вона чудово підходить для задач які потребують значного часу, такі як обробка зображень високої роздільної здатності, навчання або донавчання моделей машинного навчання. Такий підхід дозволяє одночасно розгортати вже навчену модель визначення особистості та донавчати і розгортати модель ідентифікації автора. SageMaker автоматично керує масштабуванням надаючи необхідні ресурси графічного або центрального процесора на вимогу, що особливо корисно при паралельних операціях аналізу. Детальна схема організації роботи цього рівня з використанням SageMaker відображена на рис 4.5. З точки зору безпеки, моделі машинного навчання залишаються ізольованими в приватних

підмережах віртуальної приватної хмари (VPC). Такий підхід гарантує, що логіка висновку ніколи не буде доступна поза межами системи.

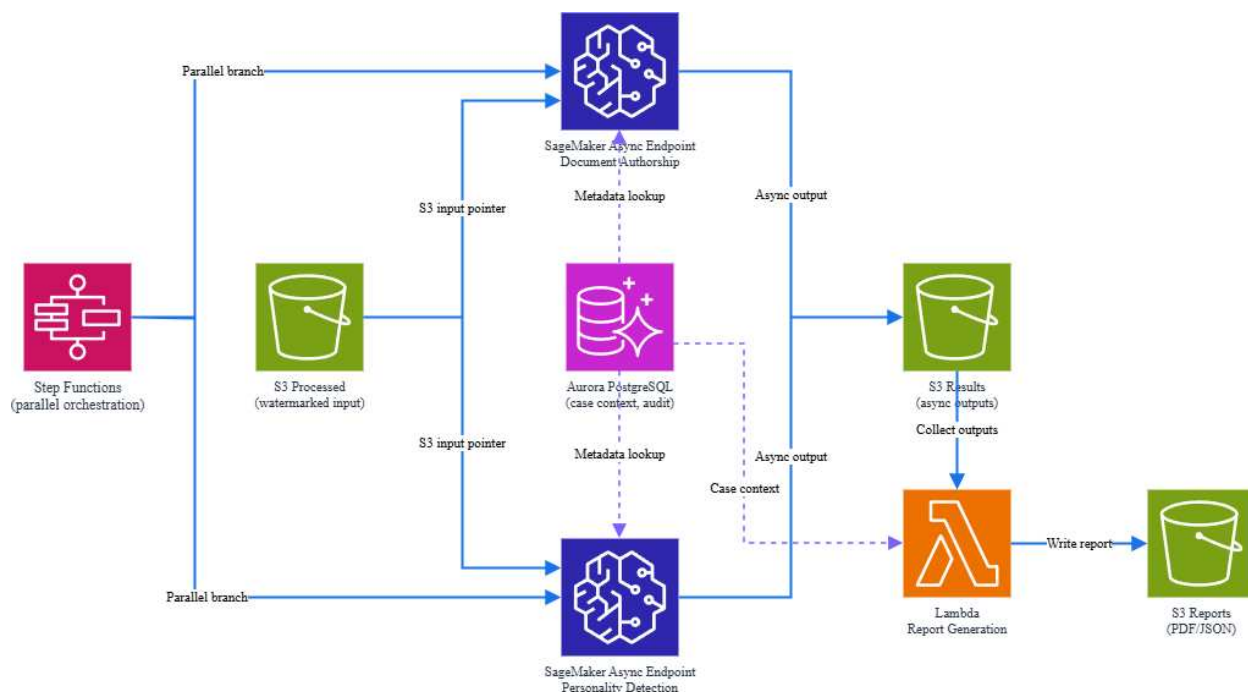


Рисунок 4.5 – Схема організації роботи моделей машинного навчання

Після завершення аналізу з використанням модулів ідентифікації автора та визначення особистості користувач має можливість сформулювати звіт. Оркестратор запускає процес генерації звіту. За агрегацію результатів аналізу та формування звіту відповідає функція AWS Lambda. Ця функція отримує результати висновку з SageMaker та поєднує їх зі специфічними метаданими та журналами аудиту, що зберігаються в базі даних Aurora. Використовуючи такі бібліотеки, як ReportLab або Pandoc, система генерує структурований звіт у форматі PDF або JSON, який синтезує всі результати в один документ. Така автоматизована звітність гарантує, що дані залишаються узгодженими. Це також убезпечує від людських помилок на етапі транскрипції результатів.

Сформований звіт слугує основою системи підтримки рішень надаючи структуровану інформацію для обґрунтованого прийняття рішень.

Загальна безпека системи та відповідність вимогам досягаються також на рівні мережі. Це можливо через використання AWS PrivateLink та кінцеві точки VPC. Такий підхід гарантує, що весь трафік між сховищами S3, моделями SageMaker та функціями Lambda на сервері залишається в межах приватної мережі Amazon. Повна схема організації безпеки доступу та інформації зображена на рис. 4.6.

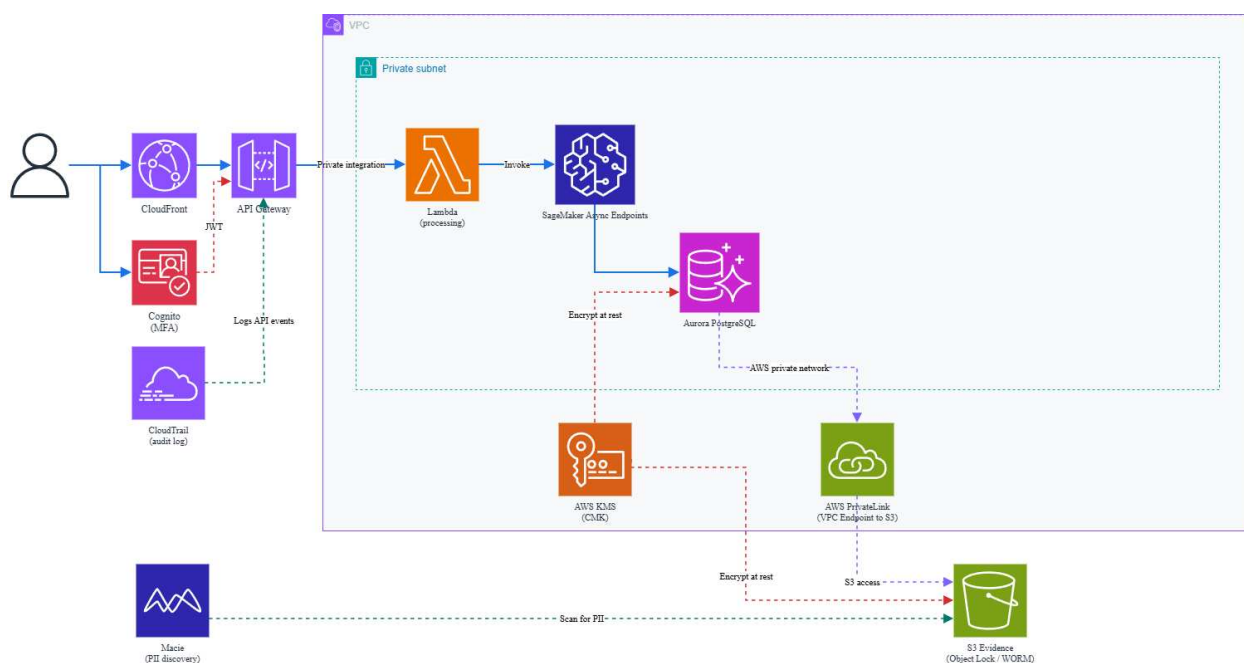


Рисунок 4.6 – Схема організації безпеки доступу та інформації

Додатково для забезпечення захисту персональної інформації (PII) використовується Amazon Macie. Його задача автоматично сканувати сховища S3 на предмет безпеки та контролю доступу, знаходити та класифікувати конфіденційні дані та автоматизувати пошук. Крім того, AWS CloudTrail надає повний, незмінний журнал кожного виклику API, здійсненого користувачем.

Цей журнал аудиту дозволяє реконструювати історію, щоб точно визначити, який користувач отримував доступ до якого зображення, які перетворення були виконані та коли було створено остаточний звіт.

4.4 Висновки за четвертим розділом

У цьому розділі була розроблена та представлена модель інформаційної технології автоматизованого графологічного аналізу рукописного тексту. Запропонована модель інформаційної технології об'єднує удосконалений метод визначення особистості та удосконалений метод ідентифікації автора для забезпечення автоматизованого аналізу, який включає як ідентифікацію автора, так і його психологічне профілювання.

1. За результатами аналізу існуючого програмного забезпечення для аналізу почерку виявлено, що системи діляться на детерміновані судово-експертні системи та гнучкі, але ненадійні комерційні засоби психологічного профілювання. На базі цього аналізу обґрунтовано та сформульовано задачу побудови інформаційної технології, що поєднує ідентифікацію автора з його психологічним профілюванням, та визначено три завдання, побудова функціональної моделі, математична формалізація та розробка архітектури інформаційної системи.

2. Розроблено інформаційну технологію автоматизованого графологічного аналізу. Побудовано функціональну та структурну моделі інформаційної технології, здійснено математичну формалізацію задачі. Охоплюється повний цикл обробки даних, починаючи від автентифікації користувача та завантаження зображення рукописного тексту, його аналізу для ідентифікації автора та визначення рис особистості і закінчуючи агрегацією результатів та формуванням комплексного звіту.

3. Розроблено архітектуру інформаційної системи на базі хмарного провайдера AWS, що реалізує запропоновану технологію. Система складається з рівнів представлення та ідентифікації (Cognito, MFA/JWT), API й оркестрування (Step Functions, Lambda), обробки (SageMaker у межах ізольованої VPC) та зберігання й метаданих (S3 Object Lock, KMS, Aurora). Цілісність і незмінність даних забезпечено криптографічним хешуванням SHA-256, незмінним сховищем та безперервним ланцюгом зберігання доказів із журналюванням (CloudTrail, Macie), що відповідає вимогам стандартів HIPAA та CJIS.

Ці результати свідчать про успішне вирішення поставлених завдань, розроблена модель інформаційної технології автоматизованого аналізу рукописного тексту створює надійне підґрунтя для подальшого розвитку систем поведінкової біометрії, успішно вирішуючи завдання щодо підвищення автоматизації, об'єктивності та точності графологічного аналізу.

Перелік використаних джерел у даному розділі наведено у повному переліку використаних джерел під номерами: 2–95.

ВИСНОВКИ

У дисертаційній роботі розв'язано актуальне науково-практичне завдання підвищення точності та автоматизації аналізу зображень рукописного тексту шляхом розробки моделі та удосконалення методів інформаційної технології автоматизованого графологічного аналізу, що забезпечують комплексну ідентифікацію автора та його психологічне профілювання. При цьому були отримані такі наукові та практичні результати.

Проведено аналіз сучасного стану досліджень у сфері графологічного аналізу рукописного тексту. Розглянуто основні процеси аналізу почерку, здійснено огляд та аналіз сучасних методів ідентифікації автора та визначення особистості. Проаналізовано психометричні моделі особистості, методи забезпечення цілісності цифрових зображень з використанням цифрових водяних знаків та технології безпечного зберігання даних. На основі проведеного аналізу сформульовано мету та задачі дослідження.

Отримала подальший розвиток модель інформаційної технології автоматизованого аналізу рукописного тексту, яка, на відміну від існуючих, об'єднує удосконалений метод визначення особистості та удосконалений метод ідентифікації автора для забезпечення автоматизованого аналізу, що включає як ідентифікацію автора, так і його психологічне профілювання. Побудовано функціональну та структурну моделі інформаційної технології, здійснено математичну формалізацію задачі та розроблено архітектуру інформаційної системи на базі хмарних технологій, яка охоплює повний цикл обробки даних.

Набув подальшого розвитку метод визначення психологічних характеристик особистості за рукописним текстом, який, на відміну від

існуючих, використовує архітектуру глибокого навчання Зоровий трансформер, що завдяки механізму уваги ефективно фіксує як локальні, так і глобальні ознаки почерку. Як психометричну основу обґрунтовано та використано модель «Велика п'ятірка», а для ефективного вилучення інформативних ділянок із зображень застосовано масштабно-інваріантне ознакове перетворення. Експериментальне дослідження на наборі даних CVL підтвердило високу ефективність запропонованого підходу, показавши максимальну точність класифікації до 99,48% для риси Сумлінність та високу точність у пошуковому підході з показником Soft Top K до 0,98.

Удосконалено метод ідентифікації та пошуку автора за рукописним текстом, в якому, на відміну від існуючих, за рахунок інтеграції масштабно-інваріантних центрально-окружних детекторів (CenSurE) отримується більша кількість ознак із зображення рукописного тексту. Побудовано функціональну модель ідентифікації автора та проведено її формалізацію шляхом побудови математичної моделі. Запропонований метод забезпечив збільшення загальної кількості вилучених ознак на 71%, скорочення середнього часу попередньої обробки на 40%, а також підвищення показників ідентифікації, зокрема зростання Hard Top K до 0,96, Soft Top K до 0,96 та середньої точності (mAP) до 0,89 порівняно з базовим підходом.

Проведено апробацію та впровадження результатів розроблених моделі та методів. Розроблена модель інформаційної технології автоматизованого аналізу рукописного тексту є підґрунтям для подальшого розвитку систем поведінкової біометрії. Запропоновані моделі та методи дозволяють збільшити точність визначення психологічних характеристик при визначенні особистості, підвищити точність ідентифікації та пошуку автора за рукописним текстом, збільшити кількість виділених ознак та зменшити середній час попередньої обробки зображень рукописного тексту.

Результати дисертаційної роботи підтверджено експериментальними дослідженнями, що засвідчили їхню достовірність та відповідність теоретичним положенням. Отримані методи та модель становлять науково-практичну основу для розвитку та впровадження автоматизованих систем графологічного аналізу у сферах криміналістики, судової експертизи, поведінкової біометрії, управління персоналом та обробки історичних архівів.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. Shupyliuk M., Martovytskyi V. Model of information technology for automated handwriting analysis. Scientific notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences. 2026. P. 442–447. DOI: <https://doi.org/10.32782/2663-5941/2026.1.2/55>
2. Shupyliuk M., Martovytskyi V., Romanenkov Y. Enhancing writer identification and writer retrieval with CenSurE and Vision Transformers. Technology Audit and Production Reserves. 2025. Vol. 6, No. 2(86). P. 6–14. DOI: <https://doi.org/10.15587/2706-5448.2025.343943>
3. Shupyliuk M., Martovytskyi V. Analysis of personality detection and writer identification methods. Control, Navigation and Communication Systems. 2025. Vol. 79, No. 1. P. 138–142. DOI: <https://doi.org/10.26906/SUNZ.2025.1.138-142>
4. Shupyliuk M., Martovytskyi V., Bolohova N., Romanenkov Y., Osiievskyi S., Liashenko S., Nesmiian O., Nikiforov I., Sukhoteplyi V., Lapchenkov Y. Devising an approach to personality identification based on handwritten text using a vision transformer. Eastern-European Journal of Enterprise Technologies. 2025. Vol. 1, No. 2(133). P. 53–65. DOI: <https://doi.org/10.15587/1729-4061.2025.322726>
5. Smelyakov K., Shupyliuk M., Martovytskyi V., Tovchyrechko D., Ponomarenko O. Efficiency of image convolution. 2019 IEEE 8th International Conference on Advanced Optoelectronics and Lasers (CAOL). 2019. P. 578–583. DOI: <https://doi.org/10.1109/CAOL46282.2019.9019450>
6. Shupyliuk M., Martovytskyi V. Novel Vision Transformer based method for Personality detection from handwriting. Seventh International Scientific

and Technical Conference COMPUTER AND INFORMATION SYSTEMS AND TECHNOLOGIES. 2024. P. 5–6.

7. Шупилюк М. В., Мартовицький В. О. Дослідження методів попередньої обробки зображень з рукописним текстом. Матеріали дванадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2(4), с. 111, 2024. DOI: <https://doi.org/10.32620/PI.24.t2>

8. Шупилюк М. В., Мартовицький В. О. Analysis of keypoint detection methods for images with handwritten text. Матеріали п'ятнадцятої міжнародної науково–технічної конференції Сучасні напрями розвитку інформаційно–комунікаційних технологій та засобів управління, том 2(2), с. 61, 2025. DOI: <https://doi.org/10.32620/ICT.25.t2>

9. Шупилюк М. В., Мартовицький В. О. Підхід до ідентифікації автора за допомогою censure та зорового трансформера. Матеріали тринадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 3(4), с. 39, 2025. DOI: <https://doi.org/10.32620/PI.25.t3>

10. Шупилюк М. В., Мартовицький В. О. Model of information technology for automated handwriting analysis. Матеріали шістнадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2, с. 132, 2026. DOI: <https://doi.org/10.32620/ICT.26.t2>

11. Joshi P., Agarwal A., Dhavale A., Suryavanshi R., Kodollikar S. Handwriting Analysis for Detection of Personality Traits using Machine Learning Approach. International Journal of Computer Applications. 2015. Vol. 130. P. 40–45. DOI: <https://doi.org/10.5120/ijca2015907189>

12. Malik M. I., Liwicki M., Dengel A., Found B. Man vs. Machine: A Comparative Analysis for Forensic Signature Verification. Journal of Forensic Document Examination. 2014. Vol. 24. P. 21–35. DOI: <https://doi.org/10.31974/jfde24-21-35>

13. Baranoski F. L. Writer Identification Based on Forensic Science Approach. CLEI Electronic Journal. 2007. P. 1–10.
14. Hicklin R. A., Eisenhart L., Richetelli N., Miller M. D., Belcastro P., Burkes T. M., Parks C. L., Smith M. A., Buscaglia J., Peters E. M., Perlman R. S., Abonamah J. V., Eckenrode B. A. Accuracy and reliability of forensic handwriting comparisons. Proceedings of the National Academy of Sciences. 2022. Vol. 119, No. 32. DOI: <https://doi.org/10.1073/pnas.2119944119>
15. Lombardi F., Marinai S. Deep Learning for Historical Document Analysis and Recognition - A Survey. Journal of Imaging. 2020. Vol. 6, No. 10. P. 110. DOI: <https://doi.org/10.3390/jimaging6100110>
16. Piccinini G. A Robust Framework for Forensic Offline Signature Verification With Enhanced Detection, Noise Removal, and Multi. IEEE Access. 2025. Vol. 13. P. 172457. DOI: <https://doi.org/10.1109/ACCESS.2025.3617221>
17. Mazzolini D., Mignone P., Pavan P., Vessio G. An easy-to-explain decision support framework for forensic analysis of dynamic signatures. Forensic Science International: Digital Investigation. 2021. Vol. 38. P. 301216. DOI: <https://doi.org/10.1016/j.fsidi.2021.301216>
18. Afzali P., Rezapour A., Rezaee Jordehi A. Leveraging deep feature learning for handwriting biometric authentication. International Journal of Research in Industrial Engineering. 2024. Vol. 13, No. 1. P. 88–103. DOI: <https://doi.org/10.22105/riej.2024.432510.1412>
19. Rassul Y. H., Ahmed A. M., Fattah P., Hassan B. A., Abdulkareem A. W., Rashid T. A., Lu J. Advancing Offline Handwritten Text Recognition: A Systematic Review of Data Augmentation and Generation Techniques. arXiv. 2025. DOI: <https://doi.org/10.48550/arxiv.2507.06275>

20. Ahmed B. Q., Hassan Y. F., Elsayed A. S. Offline text-independent writer identification using a codebook with structural features. PLOS ONE 18. 2023. Vol. 18(4). April. P. 1–31. DOI: <https://doi.org/10.1371/journal.pone.0284680>
21. Hengl M. Comparison of the Branches of Handwriting Analysis. Часопис Національного університету Острозька академія. Серія: Право. 2014. № 2(10).
22. Christlein V., Bernecker D., Hönig F., Maier A., Angelopoulou E. Writer Identification Using GMM Supervectors and Exemplar-SVMs. Pattern Recognition. 2017. Vol. 63. March. P. 258–267. DOI: <https://doi.org/10.1016/j.patcog.2016.10.005>
23. Christlein V., Gropp M., Fiel S., Maier A. Unsupervised Feature Learning for Writer Identification and Writer Retrieval. 14th IAPR International Conference on Document Analysis and Recognition (ICDAR). 2017. Vol. 13. 9-15 November. DOI: <https://doi.org/10.1109/ICDAR.2017.165>
24. Chen S., Wang Y., Lin C., Ding W., Cao Z. Semi-supervised Feature Learning For Improving Writer Identification. Information Sciences. 2019. Vol. 482. May. P. 156–170. DOI: <https://doi.org/10.1016/j.ins.2019.01.024>
25. Helal L. G., et al. Representation Learning and Dissimilarity for Writer Identification. 2019 International Conference on Systems, Signals and Image Processing (IWSSIP). 2019. June. P. 63–68. DOI: <https://doi.org/10.1109/IWSSIP.2019.8787293>
26. Sulaiman A., Omar K., Nasrudin M. F., Arram A. Length Independent Writer Identification Based on the Fusion of Deep and Hand-Crafted Descriptors. IEEE Access. 2019. Vol. 7. June. P. 91772–91784. DOI: <https://doi.org/10.1109/ACCESS.2019.2927286>

27. Kumar P., Sharma A. Segmentation-free writer identification based on convolutional neural network. *Computers & Electrical Engineering*. 2020. Vol. 85. June. DOI: <https://doi.org/10.1016/j.compeleceng.2020.106707>
28. Koepf M., Kleber F., Sablatnig R. Writer Identification and Writer Retrieval Using Vision Transformer for Forensic Documents. *Document Analysis Systems*. 2022. May. P. 352–366. DOI: https://doi.org/10.1007/978-3-031-06555-2_24
29. Semma A., Hannad Y., Siddiqi I., Djeddi C., El Youssfi El Kettani M. Writer Identification using Deep Learning with FAST Keypoints and Harris corner detector. *Expert Systems with Applications*. 2021. Vol. 184. Dec. P. 115473. DOI: <https://doi.org/10.1016/j.eswa.2021.115473>
30. He S., Schomaker L. Deep adaptive learning for writer identification based on single handwritten word images. *Pattern Recognition*. 2019. Vol. 88. April. P. 64–74. DOI: <https://doi.org/10.1016/j.patcog.2018.11.003>
31. He S., Schomaker L. FragNet: Writer Identification Using Deep Fragment Networks. *IEEE Transactions on Information Forensics and Security*. 2020. Vol. 15. March. P. 3013–3022. DOI: <https://doi.org/10.1109/TIFS.2020.2981236>
32. He S., Schomaker L. GR-RNN: Global-Context Residual Recurrent Neural Networks for Writer Identification. *Pattern Recognition*. 2021. Vol. 117. Apr. DOI: <https://doi.org/10.48550/arXiv.2104.05036>
33. Wirmanto S., Agustini D.A.R., Atmanto D.A. Offline Handwriting Writer Identification using Depth-wise Separable Convolution with Siamese Network. *International Journal On Informatics Visualization*. 2024. March. P. 535–541. DOI: <https://dx.doi.org/10.62527/joiv.8.1.2148>
34. Gavrilescu M. Study on determining the Myers-Briggs personality type based on individual's handwriting. *Proceedings of the 5th IEEE International*

Conference on E-Health and Bioengineering. 2015. Nov. P. 1–6. DOI: <https://doi.org/10.1109/EHB.2015.7391603>

35. Połap D., Wozniak M. Flexible neural network architecture for handwritten signatures recognition. *International Journal of Electronics and Telecommunications*. 2016. Vol. 62(2). April. P. 197–202. DOI: <https://doi.org/10.1515/eletel-2016-0027>

36. Topaloglu M., Ekmekci S. Gender detection and identifying one's handwriting with handwriting analysis. *ExpertSystems with Applications*. 2017. Vol. 79. March. P. 236–243. DOI: <https://doi.org/10.1016/j.eswa.2017.03.001>

37. Gavrilescu M., Vizireanu N. Predicting the big five personality traits from handwriting. *EURASIP Journal on Image and Video Processing*. 2018. Vol. 2018(1). July. DOI: <https://doi.org/10.1186/s13640-018-0297-3>

38. Joshi P., Ghaskadbi P., Tendulkar S. A machine learning approach to employability evaluation using handwriting analysis. *Proceedings of the Communications in Computer and Information Science ICAICR 2018*. 2018. June. P. 253–263. DOI: https://doi.org/10.1007/978-981-13-3140-4_23

39. Wijaya W., Tolle H., Utaminigrum F. Personality analysis through handwriting detection using android based mobile device. *International Journal of Information Technology and Computer Science*. 2018. Vol. 2(2). Feb. P. 114–128. DOI: <https://doi.org/10.25126/jitecs.20172237>

40. Fatimah S. H., Djamal E. C., Ilyas R., Renaldi F. Personality features identification from handwriting using convolutional neural networks. *Proceedings of the 4th International Conference on Information Technology, Information Systems and Electrical Engineering, ICITISEE*. 2019. Nov. P. 119–124. DOI: <https://doi.org/10.1109/ICITISEE48480.2019.9003855>

41. Chitlangia A., Malathi G. Handwriting analysis based on histogram of oriented gradient for predicting personality traits using SVM. *Procedia Computer*

Science. 2019. Vol. 165. Jan. P. 384–390. DOI: <https://doi.org/10.1016/j.procs.2020.01.034>

42. Thomas S., Goel M., Agrawal D. A framework for analyzing financial behavior using machine learning classification of personality through handwriting analysis. *Journal of Behavioral and Experimental Finance*. 2020. Vol. 26(2). March. DOI: <https://doi.org/10.1016/j.jbef.2020.100315>

43. Pathak A. R., Raut A., Pawar S., Nangare M., Abbott H. S., Chandak P. Personality analysis through handwriting recognition. *Journal of Discrete Mathematical Sciences and Cryptography*. 2020. Vol. 23(1). Jan. P. 19–33. DOI: <https://doi.org/10.1080/09720529.2020.1721856>

44. Elngar A. A., et al. A deep learning based analysis of the big five personality traits from handwriting samples using image processing. *Journal of Information Technology Management*. 2021. Vol. 12. P. 3–35. DOI: <http://doi.org/10.22059/jitm.2020.78884>

45. Bernardo L. S., Damasevicius R., De Albuquerque V. H. C., Maskeliunas R. A hybrid two-stage SqueezeNet and support vector machine system for Parkinson's disease detection based on handwritten spiral patterns. *International Journal of Applied Mathematics and Computer Science*. 2021. Vol. 31(4). Dec. P. 549–561. DOI: <https://doi.org/10.34768/amcs-2021-0037>

46. Rahman A. U., Halim Z. Predicting the big five personality traits from hand-written text features through semi-supervised learning. *Multimedia Tools and Applications*. 2022. Vol. 81(23). Sep. P. 1–17. DOI: <https://doi.org/10.1007/s11042-022-13114-5>

47. Samsuryadi S. A Framework for Determining the Big Five Personality Traits Using Machine Learning Classification through Graphology. *Journal of Electrical and Computer Engineering*. 2023. Vol. 2023(1). Jan. P. 1–15. DOI: <https://doi.org/10.1155/2023/1249004>

48. McCrae R. R., Costa P. T. The Five-Factor Theory of Personality. *Handbook of Personality: Theory and Research*. New York: Guilford Press. 1999. P. 139–151.
49. Alamsyah D., Samsuryadi S., Widhiarso W., Hasan S. Handwriting Analysis for Personality Trait Features Identification using CNN. *2022 International Conference on Data Science and Its Applications (ICoDSA)*. 2022. P. 232–238. DOI: <https://doi.org/10.1109/icodsa55874.2022.9862910>
50. Furnham A. The big five versus the big four: the relationship between the Myers-Briggs Type Indicator (MBTI) and the five-factor model of personality. *Personality and Individual Differences*. 1996. DOI: [https://doi.org/10.1016/0191-8869\(96\)00033-5](https://doi.org/10.1016/0191-8869(96)00033-5)
51. Ashton M. C., Lee K. Empirical, theoretical, and practical advantages of the HEXACO model of personality structure. *Personality and Social Psychology Review*. 2007. DOI: <https://doi.org/10.1177/1088868306294907>
52. NERIS Analytics Limited. NERIS Type Explorer® Assessment Framework: Mapping Jungian Nomenclature to Big Five (FFM) Trait Methodology. Technical Research and Validation Manual. 2020. URL: <https://www.16personalities.com/articles/our-theory> (дата звернення: 12.05.2026).
53. Gardner W. L., Martinko M. J. Using the Myers-Briggs Type Indicator in Management and Leadership: A Review of the Literature and Research Agenda. *Journal of Management*. 1996. DOI: [https://doi.org/10.1016/S0149-2063\(96\)90012-4](https://doi.org/10.1016/S0149-2063(96)90012-4)
54. Hook J. N., Hall T. W., Davis D. E., Van Tongeren D. R., Conner M. The Enneagram: A systematic review of the literature and directions for future research. *Journal of Clinical Psychology*. 2021. Vol. 77, No. 4. P. 865–883. DOI: <https://doi.org/10.1002/jclp.23097>

55. Wadhera S., Kamra A., Rajpal A., Jain A. A Comprehensive Review on Digital Image Watermarking. arXiv. 2022. DOI: <https://doi.org/10.48550/arXiv.2207.06909>
56. Review of Image Forensic Techniques Based on Deep Learning. Mathematics. 2023. Vol. 11, No. 14. P. 3134. DOI: <https://doi.org/10.3390/math11143134>
57. Bistroń M., Żurada J. M., Piotrowski Z. Deep Learning for Image Watermarking: A Comprehensive Review and Analysis of Techniques, Challenges, and Applications. Sensors. 2026. Vol. 26, No. 2. P. 444. DOI: <https://doi.org/10.3390/s26020444>
58. Petrov A., Fernandez P., Souček T., Elsahar H. We Can Hide More Bits: The Unused Watermarking Capacity in Theory and in Practice. arXiv. 2025. DOI: <https://doi.org/10.48550/arXiv.2510.12812>
59. Taj R., Tao F., Kanwal S. et al. A reversible-zero watermarking scheme for medical images. Scientific Reports. 2024. Vol. 14. P. 17320. DOI: <https://doi.org/10.1038/s41598-024-67672-9>
60. Wang C., Zhang Y., Zhou X. Robust Image Watermarking Algorithm Based on ASIFT against Geometric Attacks. Applied Sciences. 2018. Vol. 8, No. 3. P. 410. DOI: <https://doi.org/10.3390/app8030410>
61. Chen W., Ren N., Zhu C., Zhou Q., Seppänen T., Keskinarkaus A. Screen-Cam Robust Image Watermarking with Feature-Based Synchronization. Applied Sciences. 2020. Vol. 10, No. 21. P. 7494. DOI: <https://doi.org/10.3390/app10217494>
62. Zhong X., Das A., Alrasheedi F., Tanvir A. A Brief, In-Depth Survey of Deep Learning-Based Image Watermarking. Applied Sciences. 2023. Vol. 13, No. 21. P. 11852. DOI: <https://doi.org/10.3390/app132111852>

63. Mareen H., Antchougov L., Van Wallendael G., Lambert P. Blind Deep-Learning-Based Image Watermarking Robust Against Geometric Transformations. arXiv. 2024. DOI: <https://doi.org/10.48550/arXiv.2402.09062>
64. Sun N., Yuan L., Fang H., Lu Y., Ling H., Xie S., Zhao C. Ultra-high Resolution Watermarking Framework Resistant to Extreme Cropping and Scaling. 39th Conference on Neural Information Processing Systems (NeurIPS 2025). URL: <https://openreview.net/forum?id=JyxsfFEiua> (дата звернення: 12.05.2026).
65. Wang Y., Zhu X., Ye G., Zhang S. Achieving Resolution-Agnostic DNN-based Image Watermarking: A Novel Perspective of Implicit Neural Representation. Proceedings of the 32nd ACM International Conference on Multimedia (ACM MM). 2024. P. 10354–10362. DOI: <https://doi.org/10.1145/3664647.3681138>
66. Sun W., Liu J., Chen L., Dong W., Di F. MarkINeRV: A Robust Watermarking Scheme for Neural Representation for Videos Based on Invertible Neural Networks. Computers, Materials and Continua. 2024. Vol. 80, No. 3. P. 4031–4046. DOI: <https://doi.org/10.32604/cmc.2024.052745>
67. Nath S., Summers K., Baek J., Ahn G.-J. Digital Evidence Chain of Custody: Navigating New Realities of Digital Forensics. Proceedings IEEE 6th International Conference on Trust, Privacy and Security in Intelligent Systems, and Applications (TPS-ISA). 2024. P. 11–20. DOI: <https://doi.org/10.1109/TPS-ISA62245.2024.00012>
68. S3 object lock: immutability and WORM. Solved Magazine, Scalify. URL: <https://www.solved.scalify.com/how-aws-s3-object-lock-can-help-your-data-stay-safe/> (дата звернення: 12.05.2026).
69. AWS S3 plugin – Cost optimization for medical image storage? Google Groups. URL: https://groups.google.com/g/orthanc-users/c/VBB8AK_t4AM (дата звернення: 12.05.2026).

70. MinIO Best Practices – Security and Access Control. MinIO Blog. URL: <https://www.min.io/blog/s3-security-access-control> (дата звернення: 12.05.2026).

71. Kumar T. K. Encrypted Minio Storage with KMS Setup. URL: <https://krishnakumar4a4.github.io/post/encrypted-minio-storage-with-kms/> (дата звернення: 12.05.2026).

72. Bonomi S., Casini M., Ciccotelli C. B-CoC: A Blockchain-Based Chain of Custody for Evidences Management in Digital Forensics. International Conference on Blockchain Economics, Security and Protocols (Tokenomics 2019). Open Access Series in Informatics (OASICS). 2020. Vol. 71. P. 12:1–12:15. DOI: <https://doi.org/10.4230/OASICS.Tokenomics.2019.12>

73. Casella B. A performance analysis of VM-based Trusted Execution Environments for Confidential Federated Learning. arXiv. 2025. DOI: <https://doi.org/10.48550/arXiv.2501.11558>

74. Amaithi Rajan A., V V., M A. Secure and fault tolerant cloud based framework for medical image storage and retrieval in a distributed environment. Scientific Reports. 2025. Vol. 15. P. 32965. DOI: <https://doi.org/10.1038/s41598-025-16903-8>

75. Celli F., Lepri B. Is Big Five better than MBTI? Proceedings of the Fifth Italian Conference on Computational Linguistics CLiC-It. 2018. P. 93–98. DOI: <https://doi.org/10.4000/books.aaccademia.3147>

76. Mammadov S. Big Five personality traits and academic performance: A meta-analysis. Journal of Personality. 2022. Vol. 90. P. 222–255. DOI: <https://doi.org/10.1111/jopy.12663>

77. Kleber F., Fiel S., Diem M., Sablatnig R. CVL-DataBase: An Off-Line Database for Writer Retrieval, Writer Identification and Word Spotting. 2013 12th

International Conference on Document Analysis and Recognition. 2013. P. 560–564. DOI: <https://doi.org/10.1109/icdar.2013.117>

78. Lowe D. G. Distinctive Image Features from Scale-Invariant Keypoints. *International Journal of Computer Vision*. 2004. Vol. 60, No. 2. P. 91–110. DOI: <https://doi.org/10.1023/b:visi.0000029664.99615.94>

79. Otsu N. A Threshold Selection Method from Gray-Level Histograms. *IEEE Transactions on Systems, Man, and Cybernetics*. 1979. Vol. 9, No. 1. P. 62–66. DOI: <https://doi.org/10.1109/tsmc.1979.4310076>

80. Dosovitskiy A., Beyer L., Kolesnikov A., Weissenborn D., Zhai X., Unter-thiner T. et al. An Image is Worth 16 x 16 Words: Transformers for Image Recognition at Scale. *arXiv*. 2021. DOI: <https://doi.org/10.48550/arXiv.2010.11929>

81. Hassani A., Walton S., Shah N., Abuduweili A., Li J., Shi H. Escaping the Big Data Paradigm with Compact Transformers. *arXiv*. 2021. DOI: <https://doi.org/10.48550/arXiv.2104.05704>

82. Aliyev E. Forensic Handwriting Analysis to Determine the Psychophysiological Traits. *International Journal of Religion*. 2024. Vol. 5, No. 6. P. 511–530. DOI: <https://doi.org/10.61707/2r6bmr11>

83. Pandey N., Singh B., Singh S. Review on handwriting examination on unusual surface. *IP International Journal of Forensic Medicine and Toxicological Sciences*. 2024. Vol. 8, No. 4. P. 125–131. DOI: <https://doi.org/10.18231/j.ijfmts.2023.028>

84. Koepf M., Kleber F., Sablatnig R. Writer identification and writer retrieval using Vision Transformer for forensic documents. *Document Analysis Systems*. Cham: Springer. 2022. P. 352–366. DOI: https://doi.org/10.1007/978-3-031-06555-2_24

85. Sutddy W., Agustini D. A. R., Atmanto D. A. Offline Handwriting Writer Identification using Depth-wise Separable Convolution with Siamese

Network . JOIV: International Journal on Informatics Visualization. 2024. Vol. 8, No. 1. P. 535–541. DOI: <https://doi.org/10.62527/joiv.8.1.2148>

86. Gauglitz S., Höllerer T., Turk M. Evaluation of Interest Point Detectors and Feature Descriptors for Visual Tracking. *International Journal of Computer Vision*. 2011. Vol. 94. DOI: <https://doi.org/10.1007/s11263-011-0431-5>

87. Agrawal M., Konolige L., Blas M. CenSurE: Center Surround Extremas for Realtime Feature Detection and Matching. *10th European Conference on Computer Vision*. Berlin, Heidelberg: Springer. 2008. P. 102–115. DOI: https://doi.org/10.1007/978-3-540-88693-8_8

88. Tikhonov A., Rabus A. Handwritten Text Recognition of Ukrainian Manuscripts in the 21st Century: Possibilities, Challenges, and the Future of the First Generic AI-based Model. *Kyiv-Mohyla Humanities Journal*. 2024. (11). P. 226–247. DOI: <https://doi.org/10.18523/2313-4895.11.2024.226-247>

89. Deore S., Kalokhe P., Patil S., Nagarkar S. et al. Identifying the personality traits using handwriting recognition in a real-time environment. *Ingénierie des Systèmes d’Information*. 2025. Vol. 30, No. 1. P. 203–211. DOI: <https://doi.org/10.18280/isi.300117>

90. Dikevych K. G. Application of the method of cybernetic modeling in foreign practice during the writing of forensic handwriting examinations. *Kriminalistyka i sudova ekspertyza*. 2021. Vol. 66. P. 766–771. DOI: <https://doi.org/10.33994/kndise.2021.66.72>

91. You’ve Been Gone: FSD Forensics Advancements. Behind the Shades. 2025. URL: <https://www.secretservice.gov/newsroom/behind-the-shades/2025/02/youve-been-gone-fsd-forensics-advancements> (дата звернення: 10.01.2026)

92. Erp M., Vuurpijl L., Franke K., Schomaker L. The WANDA Measurement Tool for Forensic Document Examination. *J Forensic Doc Exam.* 2003. Vol. 16
93. Srihari S., Huang C., Srinivasan H. Search engine for handwritten documents. *Proc. SPIE 5676, Document Recognition and Retrieval XII.* 2005. P. 66–75. DOI: <https://doi.org/10.1117/12.585883>
94. Garoot A., Suen C. Y. Measuring the Big Five Factors from Handwriting Using Ensemble Learning Model AvgMiSC. *Intertwining Graphonomics with Human Movements: materials IGS 2022.* Cham: Springer, 2022. Vol. 13424. DOI: https://doi.org/10.1007/978-3-031-19745-1_12

ДОДАТОК А

Список публікацій здобувача

1. Shupyliuk M., Martovytskyi V. Model of information technology for automated handwriting analysis. Scientific notes of Taurida National V.I. Vernadsky University. Series: Technical Sciences. 2026. P. 442–447. DOI: <https://doi.org/10.32782/2663-5941/2026.1.2/55> (Фахове видання категорії Б).
2. Shupyliuk M., Martovytskyi V., Romanenkov Y. Enhancing writer identification and writer retrieval with CenSurE and Vision Transformers. Technology Audit and Production Reserves. 2025. Vol. 6, No. 2(86). P. 6–14. DOI: <https://doi.org/10.15587/2706-5448.2025.343943> (Входить до міжнародної наукометричної бази Scopus).
3. Shupyliuk M., Martovytskyi V. Analysis of personality detection and writer identification methods. Control, Navigation and Communication Systems. 2025. Vol. 79, No. 1. P. 138–142. DOI: <https://doi.org/10.26906/SUNZ.2025.1.138-142> (Фахове видання категорії Б).
4. Shupyliuk M., Martovytskyi V., Bolohova N., Romanenkov Y., Osiievskyi S., Liashenko S., Nesmiiian O., Nikiforov I., Sukhoteplyi V., Lapchenkov Y. Devising an approach to personality identification based on handwritten text using a vision transformer. Eastern-European Journal of Enterprise Technologies. 2025. Vol. 1, No. 2(133). P. 53–65. DOI: <https://doi.org/10.15587/1729-4061.2025.322726> (Входить до міжнародної наукометричної бази Scopus, 3 квартиль).
5. Smelyakov K., Shupyliuk M., Martovytskyi V., Tovchyrechko D., Ponomarenko O. Efficiency of image convolution. 2019 IEEE 8th International Conference on Advanced Optoelectronics and Lasers (CAOL). 2019. P. 578–583.

DOI: <https://doi.org/10.1109/CAOL46282.2019.9019450> (Входить до міжнародної наукометричної бази Scopus).

6. Shupyliuk M., Martovytskyi V. Novel Vision Transformer based method for Personality detection from handwriting. Seventh International Scientific and Technical Conference COMPUTER AND INFORMATION SYSTEMS AND TECHNOLOGIES. 2024. P. 5–6.

7. Шупилюк М. В., Мартовицький В. О. Дослідження методів попередньої обробки зображень з рукописним текстом. Матеріали дванадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2(4), с. 111, 2024. DOI: <https://doi.org/10.32620/PI.24.t2>

8. Шупилюк М. В., Мартовицький В. О. Analysis of keypoint detection methods for images with handwritten text. Матеріали п'ятнадцятої міжнародної науково–технічної конференції Сучасні напрями розвитку інформаційно–комунікаційних технологій та засобів управління, том 2(2), с. 61, 2025. DOI: <https://doi.org/10.32620/ICT.25.t2>

9. Шупилюк М. В., Мартовицький В. О. Підхід до ідентифікації автора за допомогою censure та зорового трансформера. Матеріали тринадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 3(4), с. 39, 2025. DOI: <https://doi.org/10.32620/PI.25.t3>

10. Шупилюк М. В., Мартовицький В. О. Model of information technology for automated handwriting analysis. Матеріали шістнадцятої міжнародної науково–технічної конференції Проблеми інформатизації, том 2, с. 132, 2026. DOI: <https://doi.org/10.32620/ICT.26.t2>

ДОДАТОК Б

Акт впровадження

ЗАТВЕРДЖУЮ
Генеральний директор
ПРАТ «Бізнес Автоматика»
Андрій Пилипенко
20 лютого 2026 р.

АКТ

Впровадження результатів дисертаційної роботи

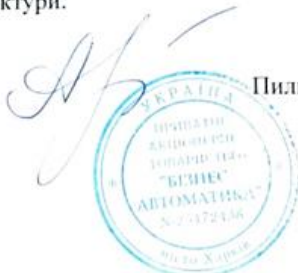
ШУПИЛЮК Микити Володимировича

Підтверджую, що наступні результати наукових досліджень ШУПИЛЮК Микити Володимировича, а саме:

- метод визначення особистості за рахунок застосування зорового трансформера, заснованого на глибинних нейронних мережах, що дозволяє підвищити точність визначення особистості;
- модель інформаційної технології автоматизованого аналізу рукописного тексту, яка базується на вдосконалених методах ідентифікації та визначення особистості та дозволяє проводити комплексну ідентифікацію автора та його психологічне профілювання.

Використання методу забезпечує збільшення точності визначення психологічних характеристик при визначенні особистості та дозволяє проводити комплексне психологічне профілювання автора. А модель інформаційної технології формалізує повний цикл обробки даних від автентифікації користувача та завантаження зображень рукописного тексту до агрегації результатів і формування вихідних звітів. Модель забезпечує інтеграцію визначення авторства документа та виявлення особистих характеристик для повноцінної ідентифікації та психологічного профілювання автора в межах єдиної функціональної архітектури.

20 лютого 2026



Пилипенко А.Г.