

НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ КОРАБЛЕБУДУВАННЯ
ІМЕНІ АДМІРАЛА МАКАРОВА
МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ

Кваліфікаційна наукова
праця на правах рукопису

Трухов Артем Сергійович

УДК 004.93:519.237

ДИСЕРТАЦІЯ

НЕГАУСІВСЬКІ ЙМОВІРНІСНІ МОДЕЛІ ТА ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ДЛЯ
РОЗПІЗНАВАННЯ ОБРАЗІВ

Спеціальність 122 – Комп'ютерні науки

Галузь знань 12 – Інформаційні технології

Подається на здобуття наукового ступеня доктора філософії (PhD)

Дисертація містить результати власних досліджень. Використання ідей, результатів і текстів інших авторів мають посилання на відповідне джерело

 _____ А.С. Трухов

Науковий керівник Приходько Сергій Борисович, доктор технічних наук, професор



Миколаїв - 2024

АНОТАЦІЯ

Трухов А.С. Негаусівські ймовірнісні моделі та інформаційні технології для розпізнавання образів. – Кваліфікаційна наукова праця на правах рукопису.

Дисертація на здобуття наукового ступеня доктора філософії (PhD) за спеціальністю 122 "Комп'ютерні науки" (галузь 12 – Інформаційні технології). – Національний університет кораблебудування імені адмірала Макарова, Міністерство освіти і науки України, Миколаїв, 2024.

Дисертаційна робота присвячена вирішенню важливого науково-практичного завдання підвищення ймовірності та точності розпізнавання образів шляхом побудови ймовірнісних моделей, зокрема, еліпсоїдів прогнозування для нормалізованих даних та створенню на їх основі інструментарію інформаційної технології (ІТ) обробки інформації для розпізнавання облич та клавіатурного почерку.

Актуальність роботи полягає в тому, що розпізнавання образів є важливою складовою багатьох сучасних технологій, таких як біометричні системи, системи відеоспостереження, автоматизовані системи контролю доступу, медичні дослідження та інші. Тому недостатня ймовірність розпізнавання образів може мати серйозні наслідки. У зв'язку з цим на сьогодні існує потреба в розробці та вдосконаленні методів розпізнавання, що забезпечують високу точність та надійність в умовах змінних і багатовимірних даних.

Метою дисертаційної роботи є підвищення ймовірності та точності розпізнавання образів, зокрема за рахунок побудови негаусівських ймовірнісних моделей для нормалізованих даних.

Робочою науковою гіпотезою дисертаційного дослідження є твердження, що підвищення ймовірності та точності розпізнавання образів досягається за рахунок використання негаусівських ймовірнісних моделей, які дозволяють враховувати відмінність розподілу даних від нормального.

Для побудови вказаних негаусівських ймовірнісних моделей пропонується використовувати відповідний метод на основі багатовимірних нормалізуючих перетворень, які дозволяють враховувати кореляцію між змінними та наблизити розподіл до нормального, що підвищує ймовірність та точність розпізнавання образів.

Для досягнення поставленої мети необхідно вирішити такі завдання:

- проаналізувати існуючі математичні моделі, що використовуються для розпізнавання образів, та оцінити їх точність при роботі з негаусівськими даними;
- дослідити існуючі нормалізуючі перетворення та методи для оцінювання їх параметрів; обрати перетворення для нормалізації десятивимірних векторів характеристик облич та дев'ятивимірних векторів характеристик клавіатурного почерку;
- побудувати ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, який побудовано за допомогою десятивимірного нормалізуючого перетворення;
- побудувати ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, який побудовано за допомогою дев'ятивимірного нормалізуючого перетворення;
- удосконалити ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, який побудовано за допомогою десятивимірного нормалізуючого перетворення;
- удосконалити ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, який побудовано за допомогою дев'ятивимірного нормалізуючого перетворення;

- створити на основі удосконалених ймовірнісних моделей інформаційні технології для розпізнавання облич і клавіатурного почерку.

Наукова новизна одержаних результатів полягає у наступному.

1) Вперше побудована ймовірнісна модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірного перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу, що дозволяє підвищити точність і ймовірність розпізнавання облич.

2) Вперше побудована ймовірнісна модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірного перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу, що дозволяє підвищити точність і ймовірність розпізнавання клавіатурного почерку.

3) Удосконалено ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірного перетворення Бокса-Кокса, що дозволяє підвищити точність і ймовірність розпізнавання облич.

4) Удосконалено ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірного перетворення Бокса-Кокса, що дозволяє підвищити точність і ймовірність розпізнавання клавіатурного почерку.

Практичне значення одержаних результатів полягає у наступному. Розроблено інструментарій ІТ для розпізнавання облич та клавіатурного почерку на основі вектору характеристик, використовуючи мову програмування Python, що забезпечує її високу гнучкість та можливість інтеграції з сучасними платформами.

У вступі дисертації розкрито сутність та значущість науково-практичного завдання, обґрунтовано потребу в проведенні даного дослідження. Наведено загальний опис роботи в такій структурі: актуальність обраної теми, зв'язок дослідження з науковими програмами, планами і темами; мета і завдання дослідження; наукова новизна і практична цінність отриманих результатів; особистий внесок здобувача; апробація результатів дослідження та публікації.

У першому розділі дисертації виконано аналіз існуючих методів і моделей для розпізнавання образів та здійснено обґрунтування необхідності проведення досліджень за обраною темою.

У другому розділі дисертації розглянуто існуючі взаємо-зворотні нормалізуючі перетворення, здійснено вибір перетворення для нормалізації дев'ятивимірних та десятивимірних векторів ознак клавіатурного почерку та облич відповідно. Проведено аналіз і видалення викидів з даних клавіатурного почерку за допомогою відстані Махаланобіса, використовуючи нормалізовані дані, отримані за допомогою багатовимірного перетворення Бокса-Кокса.

У третьому розділі було побудовано негаусівські ймовірнісні моделі у вигляді дев'ятивимірного еліпсоїда прогнозування для нормалізованих векторів характеристик клавіатурного почерку, та десятивимірного еліпсоїда прогнозування для нормалізованих векторів характеристик облич, з використанням багатовимірного перетворення Бокса-Кокса і квантилів Хі-квадрат та F -розподілу.

У четвертому розділі дисертації запропоновано ІТ для розпізнавання облич та клавіатурного почерку. Розроблено інструментарії ІТ на основі побудованих негаусівських ймовірнісних моделей для розпізнавання облич за десятивимірним вектором ознак та для розпізнавання клавіатурного почерку за дев'ятивимірним

вектором ознак. Було розроблено відповідне ПЗ, використовуючи мову програмування Python, та інструкції користувача для розпізнавання облич та клавіатурного почерку, які розраховані на використання зазначеного ПЗ.

Ключові слова: розпізнавання образів, квадрат відстані Махаланобіса, еліпсоїд прогнозування, нормалізуюче перетворення, негаусівські дані, інформаційна технологія, машинне навчання, нейронні мережі, класифікація, прийняття рішень, людський фактор, Python

Список публікацій здобувача за темою дисертації:

1) в яких опубліковані основні наукові результати дисертації:

- *статті у наукових виданнях, включених на дату опублікування до переліку наукових фахових видань України:*

1. Приходько, С., & Трухов, А. (2023). Побудова правил прийняття рішень для розпізнавання облич на основі квадрата відстані Махаланобіса для нормалізованих даних. *Information Technology: Computer Science, Software Engineering and Cyber Security*, (2), 50-58. DOI: 10.32782/IT/2023-2-6.

2. Prykhodko, S. B., & Trukhov, A. S. (2024). Face recognition using ten-variate prediction ellipsoids for normalized data with different quantiles. *Applied Aspects of Information Technology*, 7(2), 151-161. DOI: 10.15276/aait.07.2024.11.

- *статті у періодичних наукових виданнях, проіндексованих у базах даних Web of Science Core Collection та/або Scopus (крім видань держави, визнаної Верховною Радою України державою-агресором):*

3. Prykhodko, S. B., & Trukhov, A. S. (2024). Face recognition using the ten-variate prediction ellipsoid for normalized data based on the Box-Cox transformation. *Radio Electronics, Computer Science, Control*, (2), 82-89. DOI: 10.15588/1607-3274-2024-2-9 (індексується наукометричною базою Web of Science).

4. Prykhodko, S., & Trukhov, A. (2024). Application of a Ten-Variate Prediction Ellipsoid for Normalized Data and Machine Learning Algorithms for Face

Recognition. *International Workshop on Computer Modeling and Intelligent Systems*, Zaporizhzhia, Ukraine, May 3, 2024. CEUR Workshop Proceedings, Vol.3702, pp. 362-375, 2024. <https://ceur-ws.org/Vol-3702/paper30.pdf> (Aachen, Germany, ISSN 1613-0073. Стаття включена до міжнародних наукометричних баз Scopus та DBLP).

2) які засвідчують апробацію матеріалів дисертації:

5. **Трухов, А.С., & Приходько, С.Б.** (26-28 жовтня 2021 р.). Аналіз існуючих математичних моделей для розпізнавання образів. *Матеріали II Всеукраїнської науково-практичної інтернет конференції*. – Миколаїв: НУК імені адмірала Макарова, 2021. – с. 8-10.

6. **Трухов, А.С., & Приходько, С.Б.** (2022). Математичні моделі та методи для ідентифікації та виділення образів. *Матеріали XXII Всеукраїнської науково-технічної конференції молодих вчених, аспірантів та студентів*. Одеса, 21-22 квітня 2022 р. - Одеса, Видавництво ОНТУ, 2022 р. – с. 44-46.

7. **Трухов, А.С., & Приходько, С.Б.** (2022). Застосування квадрата відстані махаланобіса в задачі розпізнавання облич. *Матеріали III Всеукраїнської науково-практичної інтернет конференції* (26-28 жовтня 2022 р.). – Миколаїв: НУК імені адмірала Макарова, 2022. – с. 35-37.

8. **Prykhodko, S., & Trukhov, A.** (2023). Free software for object identification and feature extraction. *Матеріали XIV-ої Міжнародної науково-практичної конференції «Free and Open Source Software»* (14-16 лютого 2023 р.). – Харків: Харківський національний економічний університет імені Семена Кузнеця, 2023.–110 с. 18-19.

9. **Трухов, А.С., & Приходько, С.Б.** (2023). Побудова правил прийняття рішень для нормалізованих даних у задачі розпізнавання образів. *Матеріали IV Всеукраїнської науково-практичної інтернет конференції* (30-31 жовтня 2023 р.). – Миколаїв: НУК імені адмірала Макарова, 2023. – с. 7-9.

10. **Трухов, А.С., & Приходько, С.Б. (2024).** Розпізнавання клавіатурного почерку за допомогою дев'ятивимірних еліпсоїдів прогнозування для нормалізованих даних з різними квантилями. *Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2024)* - (Миколаїв, 30-31 жовтня 2024 р.) - с. 45 47.

ABSTRACT

Trukhov A.S. Non-Gaussian probabilistic models and information technologies for pattern recognition. – Manuscript of the qualification scientific work.

Thesis for the degree of philosophy doctor in specialty 122 "Computer Science" (field 12 – Information Technologies). – Admiral Makarov National University of Shipbuilding, Ministry of Education and Science of Ukraine, Mykolaiv, 2024.

The dissertation is dedicated to solving the critical scientific and practical problem of improving the probability and accuracy of pattern recognition by developing probabilistic models, specifically prediction ellipsoids for normalized data, and creating information technology (IT) tools for face recognition and keystroke dynamics recognition based on these models.

The relevance of the study lies in the fact that pattern recognition is a vital component of many modern technologies, such as biometric systems, video surveillance systems, automated access control systems, medical research, and others. Therefore, insufficient recognition accuracy can lead to serious consequences. In this regard, there is a growing need to develop and improve recognition methods that ensure high accuracy and reliability under conditions of variable and multivariate data.

The dissertation aim is to improve the probability and accuracy of pattern recognition, specifically through the development of non-Gaussian probabilistic models for normalized data.

The working scientific hypothesis of the dissertation is that improving the probability and accuracy of pattern recognition can be achieved through the use of non-Gaussian probabilistic models, which account for deviations in data distribution from normality.

To construct these non-Gaussian probabilistic models, the proposed approach utilizes multivariate normalizing transformations that account for variable correlation and approximate the data distribution to normality, thereby enhancing the probability and accuracy of pattern recognition.

To achieve this aim we need to solve the following tasks:

- Analyze existing mathematical models used for pattern recognition and evaluate their accuracy when working with non-Gaussian data.
- Investigate existing normalizing transformations and methods for evaluating their parameters; choose a transformation for normalizing ten-variate feature vectors for face recognition and nine-variate feature vectors for keystroke dynamics.
- Build a probabilistic model for face recognition in the form of a prediction ellipsoid with a chi-square quantile for a normalized feature vector, constructed using a ten-variate normalizing transformation.
- Build a probabilistic model for keystroke dynamics recognition in the form of a prediction ellipsoid with an F -distribution quantile for a normalized feature vector, constructed using a nine-variate normalizing transformation.
- Improve the probabilistic model for face recognition in the form of a prediction ellipsoid with a chi-square quantile for a normalized feature vector, constructed using a ten-variate normalizing transformation.
- Improve the probabilistic model for keystroke dynamics recognition in the form of a prediction ellipsoid with a chi-square quantile for a normalized feature vector, constructed using a nine-variate normalizing transformation.
- Develop information technologies for face and keystroke dynamics recognition based on the improved probabilistic models.

The scientific novelty of the obtained results is as follows.

- 1) For the first time, a probabilistic model for face recognition has been constructed in the form of a prediction ellipsoid with an F -distribution quantile for a normalized feature vector, which, unlike existing models, takes into account the deviation of the data distribution from Gaussian by applying the ten-variate Box-Cox transformation and the number of data points by using the F -distribution quantile, which allows to increase the accuracy and probability of face recognition.
- 2) For the first time, a probabilistic model for keyboard handwriting recognition has been constructed in the form of a prediction ellipsoid with an F -distribution quantile for a normalized feature vector, which, unlike existing models, takes into account the deviation of the data distribution from Gaussian by applying the

nine-variate Box-Cox transformation and the number of data points by using the F -distribution quantile, which allows to increase the accuracy and probability of keyboard handwriting recognition.

3) The probabilistic model for face recognition in the form of a prediction ellipsoid with a Chi-square quantile for a normalized feature vector has been improved, which, unlike existing models, takes into account the deviation of the data distribution from Gaussian by applying the ten-variate Box-Cox transformation, which allows to increase the accuracy and probability of face recognition.

4) The probabilistic model for keyboard handwriting recognition in the form of a prediction ellipsoid with a Chi-square quantile for a normalized feature vector has been improved, which, unlike existing models, takes into account the deviation of the data distribution from Gaussian by applying the nine-variate Box-Cox transformation, which allows to increase the accuracy and probability of keyboard handwriting recognition.

The practical significance of the obtained results is as follows. The IT tools developed for face and keystroke dynamics recognition are based on feature vectors and implemented using Python programming language, ensuring high flexibility and integration with modern platforms.

The introduction outlines the essence and importance of the scientific and practical problem, substantiates the need for this research, and provides an overview of the study in the following structure: relevance of the topic, connection to scientific programs and plans, research objectives, scientific novelty, practical value of the results, personal contribution of the author, and dissemination of results and publications.

In Section 1, the analysis of existing methods and models for pattern recognition is presented. The section identifies key limitations in handling non-Gaussian data distributions and emphasizes the need for further research in developing advanced probabilistic models to address these challenges.

In Section 2, existing invertible normalizing transformations are reviewed. The section describes the selection of appropriate transformations for normalizing nine-

variate and ten-variate feature vectors for keystroke dynamics and face data, respectively. It also details the analysis and removal of outliers from keystroke dynamics data using the Mahalanobis distance, applied to data normalized via the multivariate Box-Cox transformation.

In Section 3, non-Gaussian probabilistic models are constructed. These models include the nine-variate prediction ellipsoid for normalized keystroke dynamics feature vectors and the ten-variate prediction ellipsoid for normalized face feature vectors. The section explains the use of the multivariate Box-Cox transformation and the incorporation of Chi-square and F-distribution quantiles in model construction.

In Section 4, information technology (IT) tools for face and keystroke dynamics recognition are proposed. The section discusses the development of IT tools based on the constructed non-Gaussian probabilistic models, tailored for ten-variate feature vectors for face recognition and nine-variate feature vectors for keystroke dynamics recognition. It also outlines the implementation of the corresponding software in Python and provides a methodology for utilizing the software in practical applications.

Keywords: pattern recognition, squared Mahalanobis distance, prediction ellipsoid, normalizing transformation, non-Gaussian data, information technology, machine learning, neural networks, classification, decision making, human factor, Python

ЗМІСТ

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ, СКОРОЧЕНЬ І ТЕРМІНІВ	15
ВСТУП.....	16
РОЗДІЛ 1 АНАЛІЗ ІСНУЮЧИХ МАТЕМАТИЧНИХ МОДЕЛЕЙ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ ТА ОБҐРУНТУВАННЯ НЕОБХІДНОСТІ ПРОВЕДЕННЯ ДОСЛІДЖЕННЯ ЗА ОБРАНОЮ ТЕМОЮ	23
1.1 Аналіз існуючих математичних моделей для розпізнавання образів	23
1.2 Формування набору даних для розпізнавання облич	31
1.3 Формування набору даних для розпізнавання клавіатурного почерку.....	42
1.4 Обґрунтування необхідності проведення досліджень за обраною темою	45
1.5 Висновки до розділу 1	47
РОЗДІЛ 2 ВИБІР НОРМАЛІЗУЮЧОГО ПЕРЕТВОРЕННЯ ДЛЯ ПОБУДОВИ НЕГАУСІВСЬКИХ ЙМОВІРНІСНИХ МОДЕЛЕЙ	50
2.1 Одновимірні взаємо-зворотні нормалізуючі перетворення	51
2.2 Багатовимірне перетворення Бокса-Кокса.....	53
2.3 Вибір перетворення для нормалізації десятидимірних характеристик облич	54
2.4 Вибір перетворення для нормалізації та дев'ятидимірних характеристик клавіатурного почерку	59
2.5 Висновки за розділом 2.....	68
РОЗДІЛ 3 ПОБУДОВА НЕГАУСІВСЬКИХ ЙМОВІРНІСНИХ МОДЕЛЕЙ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ	70

3.1 Побудова десятивимірних еліпсоїдів прогнозування для розпізнавання облич з різними квантилями розподілів.....	70
3.2 Побудова дев'ятивимірних еліпсоїдів прогнозування для розпізнавання клавіатурного почерку з різними квантилями розподілів	76
3.3 Порівняння точності та ймовірності розпізнавання образів за побудованими еліпсоїдами прогнозування для нормалізованих даних та за алгоритмами машинного навчання	83
3.4 Висновки за розділом 3	96
РОЗДІЛ 4 ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ	100
4.1 Створення інформаційної технології для розпізнавання облич	101
4.2 Створення інформаційної технології для розпізнавання клавіатурного почерку	106
4.3 Інструкція користувача для ІТ розпізнавання облич	110
4.4 Інструкція користувача для ІТ розпізнавання клавіатурного почерку.....	120
4.5 Висновки за розділом 4.....	132
ВИСНОВКИ.....	135
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	140
ДОДАТОК А. АКТ ВПРОВАДЖЕННЯ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЙНОЇ РОБОТИ	156
ДОДАТОК Б. АКТ ВПРОВАДЖЕННЯ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЙНОЇ РОБОТИ В НАВЧАЛЬНИЙ ПРОЦЕС	158
ДОДАТОК В. СПИСОК ПУБЛІКАЦІЙ ЗДОБУВАЧА ЗА ТЕМОЮ ДИСЕРТАЦІЇ ТА ВІДОМОСТІ ПРО АПРОБАЦІЮ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЇ.....	160

**ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СИМВОЛІВ, ОДИНИЦЬ,
СКОРОЧЕНЬ І ТЕРМІНІВ**

БКП – перетворення Бокса-Кокса;

ІТ – інформаційна технологія;

КВМ – квадрат відстані Махаланобіса;

ММП – метод максимальної правдоподібності;

ПЗ – програмне забезпечення;

АЕ - autoencoder;

ІF - isolation forest;

ОCSVM - one-class support vector machine;

ВСТУП

Обґрунтування вибору теми дослідження. Розпізнавання образів — це процес ідентифікації або класифікації об'єктів на основі характерних ознак або властивостей [1]. Це складний міждисциплінарний напрямок, що охоплює комп'ютерні науки, машинне навчання, штучний інтелект, математику та статистику. Основна мета розпізнавання образів полягає в тому, щоб навчити комп'ютерні системи розпізнавати шаблони [2]: за допомогою різноманітних алгоритмів та математичних моделей системи здатні аналізувати величезні масиви даних та автоматично ідентифікувати об'єкти або виявляти закономірності.

Ця технологія широко використовується у різних сферах: для контролю якості та автоматизації виробничих процесів; для аналізу медичних зображень [3, 4], діагностики хвороб та моніторингу стану пацієнтів; для розробки систем безпілотних автомобілів [5]; для ідентифікації людей [6], моніторингу відеозаписів та контролю доступу [7]; у фінансових установах розпізнавання образів допомагає виявляти шахрайські транзакції [8], аналізуючи незвичайні моделі поведінки в діяльності облікового запису; в електронній комерції воно допомагає персоналізувати роботу користувачів шляхом розпізнавання моделей покупок і вподобань.

Одне з найважливіших застосувань розпізнавання образів сьогодні — у системах автентифікації користувачів, яке відбувається, зокрема, за допомогою розпізнавання обличчя та клавіатурного почерку [9, 10]. Ці системи аналізують особливі характеристики, унікальні для особи, і використовують їх для перевірки. Виходячи з цього, недостатня ймовірність розпізнавання образів може мати серйозні наслідки в таких сферах, як безпека, автоматизація, персоналізація та ін. Тому на сьогодні існує потреба в розробці та вдосконаленні методів розпізнавання образів.

В задачах розпізнавання облич і клавіатурного почерку для мультикласового розпізнавання часто використовують лінійний дискримінантний аналіз [11-14], метод опорних векторів [15-19], k-найближчих сусідів [20-22], нейронні мережі [23-28] та інші.

Для розпізнавання облич та клавіатурного почерку характерне використання однокласового розпізнавання, коли навчання моделі відбувається лише на об'єктах цільового класу, без представників інших класів [29-30]. Це дозволяє побудувати модель для окремого представника та проводити розпізнавання у контексті, чи належать дані ознаки конкретній особі. Однокласове розпізнавання особливо корисна для виявлення аномалій та нетипових подій, оскільки вона фокусується на виявленні відхилень від цільового стану. До найбільш популярних методів однокласового розпізнавання належить використання еліпсоїда прогнозування [31-33] та алгоритмів машинного навчання, такі як однокласова опорно-векторна машина [34, 35], ізоляційний ліс [36] та автокодувальник [37, 38].

Основним недоліком сучасних підходів до розпізнавання образів є те, що вони часто ґрунтуються на припущенні про багатовимірний нормальний розподіл даних [39-41]. Однак у реальному світі ця умова не завжди виконується [42], що може призвести до зниження точності та ймовірності розпізнавання для негаусівських даних. Одним з вирішення цього обмеження є використання нормалізуючих перетворень, які допомагають коригувати розподіл даних, наближаючи його до нормального [43-46]. Нормалізуючі перетворення бувають одновимірні та багатовимірні. Одновимірні перетворення застосовуються до кожної ознаки окремо, тоді як багатовимірні перетворення враховують взаємозв'язки між ознаками.

Таким чином, підвищення ймовірності та точності розпізнавання образів досягається за рахунок використання негаусівських ймовірнісних моделей, які дозволяють враховувати відмінність розподілу даних від нормального. Для побудови вказаних негаусівських ймовірнісних моделей пропонується

використовувати відповідний метод на основі нормалізуючих перетворень, які дозволяють наблизити розподіл до нормального.

Зв'язок роботи з науковими програмами, планами, темами. Дисертаційну роботу виконано у Національному університеті кораблебудування імені адмірала Макарова (НУК) відповідно до планів НДР з ініціативної теми «Негаусівські ймовірнісні моделі для розпізнавання облич» (Державний реєстраційний номер: 0124U004650).

Мета і завдання дослідження. Метою роботи є підвищення ймовірності та точності розпізнавання образів, зокрема за рахунок побудови негаусівських ймовірнісних моделей для нормалізованих даних.

Для досягнення поставленої мети необхідно вирішити такі завдання:

- проаналізувати існуючі математичні моделі, що використовуються для розпізнавання образів, та оцінити їх точність при роботі з негаусівськими даними;
- дослідити існуючі нормалізуючі перетворення та методи для оцінювання їх параметрів; обрати перетворення для нормалізації десятивимірних векторів характеристик облич та дев'ятивимірних векторів характеристик клавіатурного почерку;
- побудувати ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, який побудовано за допомогою десятивимірного нормалізуючого перетворення;
- побудувати ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, який побудовано за допомогою дев'ятивимірного нормалізуючого перетворення;
- удосконалити ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, який побудовано за допомогою десятивимірного нормалізуючого перетворення;

- удосконалити ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, який побудовано за допомогою дев'ятивимірною нормалізуючого перетворення;

- створити на основі удосконалених ймовірнісних моделей інформаційні технології для розпізнавання облич і клавіатурного почерку.

Об'єктом дослідження є процес розпізнавання образів.

Предметом дослідження є ймовірнісні моделі для розпізнавання образів.

Методи дослідження. Для вирішення поставлених завдань було застосовано методи теорії ймовірності, математичної статистики та багатовимірною статистичного аналізу для оцінки характеристик випадкових величин, перевірки статистичних гіпотез про нормальність розподілів даних, для оцінювання параметрів нормалізуючих перетворень, а також для побудови ймовірнісних моделей.

Наукова новизна одержаних результатів полягає у наступному.

1) Вперше побудована ймовірнісна модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірною перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу, що дозволяє підвищити точність і ймовірність розпізнавання облич.

2) Вперше побудована ймовірнісна модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірною перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу, що дозволяє підвищити точність і ймовірність розпізнавання клавіатурного почерку.

3) Удосконалено ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем Хі-квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірного перетворення Бокса-Кокса, що дозволяє підвищити точність і ймовірність розпізнавання облич.

4) Удосконалено ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем Хі-квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірного перетворення Бокса-Кокса, що дозволяє підвищити точність і ймовірність розпізнавання клавіатурного почерку.

Практичне значення одержаних результатів полягає у наступному. Розроблено інструментарій ІТ для розпізнавання облич та клавіатурного почерку на основі вектору характеристик, використовуючи мову програмування Python, що забезпечує її високу гнучкість та можливість інтеграції з сучасними платформами.

Результати дисертаційної роботи, висновки та рекомендації, що містяться у роботі, схвалені та використовуються в наступних установах: ТОВ НВЦ «Діагностика та контроль». Завдяки застосуванню інформаційної технології для розпізнавання облич та інформаційної технології для розпізнавання клавіатурного почерку, створених з використанням розроблених ймовірнісних моделей, які враховують негаусівський розподіл даних, підвищилась точність та ймовірність розпізнавання облич та клавіатурного почерку.

Результати дисертаційної роботи також використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та управління проектами (ННІКНУП) Національного університету кораблебудування ім. адмірала Макарова наступних дисциплін: «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»;

«Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки».

Акти про впровадження результатів дисертаційної роботи наведені у додатках А і Б.

Особистий внесок здобувача. Усі наукові результати отримано здобувачем самостійно. У працях, опублікованих у співавторстві, дисертанту належать: в [47] – десятивимірний еліпсоїд прогнозування для нормалізованих за допомогою десяткового логарифму даних для розпізнавання облич; в [48] – десятивимірні еліпсоїди прогнозування для нормалізованих за допомогою одновимірного та багатовимірного перетворення Бокса-Кокса даних для розпізнавання облич; в [49] – десятивимірні еліпсоїди прогнозування для нормалізованих даних з різними квантилями розподілів Хі-квадрат та F -розподілу для розпізнавання облич; в [50] – реалізовані алгоритми машинного навчання, а саме однокласова опорна векторна машина, ізоляційний ліс та автокодувальник для порівняння з побудованим десятивимірним еліпсоїдом для нормалізованих даних для розпізнавання облич.

Апробація результатів досліджень. Основні результати дисертації доповідалися та обговорювалися на 8 науково-технічних конференціях, зокрема:

- II Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2021) (Миколаїв, 26-28 жовтня 2021 р.).

- XXII Всеукраїнській науково-технічній конференції молодих вчених, аспірантів та студентів «Стан, досягнення та перспективи інформаційних систем і технологій» (Одеса, 21-22 квітня 2022 р.).

- III Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2022) (Миколаїв, 26-28 жовтня 2022 р.).

- XIV Міжнародній науково-практичній конференції «Free and Open Source Software» (Харків, 14-16 лютого 2023 р.).

- IV Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2023) (Миколаїв, 30-31 жовтня 2023 р.).
- V Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2024) (Миколаїв, 30-31 жовтня 2024 р.).
- XI International Conference «Information Technology and Implementation» (IT&I-2024) (Kyiv, November 20-21, 2024).
- Seventh International Workshop on Computer Modeling and Intelligent Systems (CMIS-2024) (Zaporizhzhia, May 3, 2024).

Публікації. Результати роботи викладені у 10 публікаціях. Основні наукові результати дисертації опубліковано у 4 статтях, з яких дві статті у наукових виданнях, включених на дату опублікування до переліку наукових фахових видань України, а дві статті у періодичних наукових виданнях, проіндексованих у базах даних Web of Science Core Collection та/або Scopus (крім видань держави, визнаної Верховною Радою України державою-агресором). Ще 6 публікацій засвідчують апробацію матеріалів дисертації.

Структура і обсяг дисертації. Дисертаційна робота складається з вступу, чотирьох розділів, висновків, списку використаних джерел з 129 найменувань і 3 додатків. Загальний обсяг дисертації становить 161 сторінок, обсяг основного тексту складає 139 сторінок. Дисертаційна робота містить 50 рисунків та 45 таблиць. У додатках наведені документи про впровадження результатів роботи, список публікацій здобувача за темою дисертації та відомості про апробацію результатів дисертації.

РОЗДІЛ 1

АНАЛІЗ ІСНУЮЧИХ МАТЕМАТИЧНИХ МОДЕЛЕЙ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ ТА ОБҐРУНТУВАННЯ НЕОБХІДНОСТІ ПРОВЕДЕННЯ ДОСЛІДЖЕННЯ ЗА ОБРАНОЮ ТЕМОЮ

1.1 Аналіз існуючих математичних моделей для розпізнавання образів

Розпізнавання образів — це технічна дисципліна, яка фокусується на розробці методів та систем для визначення класу об'єкта з набору можливих класів. Об'єкт, який аналізується, описується набором характеристик. Вибір правила прийняття рішень повинен базуватися на використанні обмеженого числа ознак, які найкраще відрізняють один образ від іншого. Правило прийняття рішень (вирішальне правило або класифікатор) — це математичний вираз або алгоритм, який визначає належність образу до певного класу [51].

Розпізнавання образів знаходить широке застосування у різних галузях. У медицині вони використовуються для аналізу медичних зображень [3, 4], включаючи рентгенівські знімки та томограми. У сфері безпеки — для біометричної автентифікації, наприклад, під час аналізу відбитків пальців чи розпізнавання осіб [6, 7]. У промисловості такі алгоритми застосовуються для автоматизованого контролю за якістю продукції, а фінансовому секторі — задля унеможливлення шахрайства [8]. Тому на сьогодні існує потреба підвищення точності та ймовірності розпізнавання образів.

Система розпізнавання образів, аналогічно біонічним системам, зазвичай працює в двох режимах: навчання та тестування. Навчання є ключовим етапом процесу розпізнавання, метою якого є створення еталонних описів класів. Форма цих описів визначається тим, як вони будуть використовуватися в правилах прийняття рішень. Вибірка — це набір об'єктів, представлених значеннями ознак, належність яких до певного класу відома. Зазвичай вибірка ділиться на дві частини: навчальну та тестову [52]. Навчальна вибірка використовується для

створення правил прийняття рішень, які дозволяють системі розпізнавати об'єкти за їх ознаками. Тестова вибірка, в свою чергу, служить для оцінки якості цих правил, тобто їх здатності правильно визначати клас об'єктів. Важливо, щоб тестова вибірка не перетиналася з навчальною вибіркою і щоб обидві вибірки були репрезентативними для загальної популяції [53]. Це забезпечує точність і надійність системи розпізнавання образів.

Зазвичай процес розпізнавання складається з декількох етапів [54]:

1) Збір даних - це основа будь-якої системи розпізнавання образів. Дані можуть надходити з різних джерел, таких як зображення, текст, сенсорні зчитування або біометричні показники. Якість і кількість даних значно впливають на продуктивність системи.

2) Попередня обробка даних - це етап, на якому початкові дані піддаються різним операціям, таким як фільтрація, нормалізація, сегментація, виділення об'єкту на зображенні, видалення фону та інше, з метою підвищити якість і зменшити обсяг даних для наступного аналізу.

3) Отримання ознак - це етап, на якому визначаються важливі характеристики об'єктів, які показують їх властивості та відмінності. Ознаки можуть належати до різних типів і можуть мати форму векторів, матриць, графів та ін.

4) Побудова правил прийняття рішень для розпізнавання - це етап, на якому створюється ймовірнісна модель, яка може розрізняти класи образів за їх ознаками. Класифікатор може ґрунтуватися на різних методах, таких як статистичні, використання нейронних мереж, геометричні та ін., і може бути навчений на певному наборі даних з відомими мітками класів.

5) Прийняття рішення - це етап, на якому здійснюється класифікація нових даних за допомогою побудованої моделі. Результатом цього етапу є присвоєння мітки класу або ймовірності належності до кожного класу для кожного об'єкта.

Розпізнавання облич та клавіатурного почерку є важливими завданнями в галузі біометричного розпізнавання [9, 10], засновані на аналізі та інтерпретації унікальних характеристик людини.

Розпізнавання облич включає аналіз геометричних, текстурних та інших візуальних особливостей зображення обличчя, таких як відстань між очима, форма носа або структура шкіри [55]. Це завдання актуальне у системах відеоспостереження, біометричної аутентифікації, а також для доступу до електронних пристроїв.

Розпізнавання клавіатурного почерку пов'язане з аналізом динамічних характеристик процесу введення тексту, таких як часові інтервали між натисканнями клавіш та тривалість їх утримання [56]. Ці параметри є унікальними для кожного користувача та використовуються в системах аутентифікації, особливо в онлайн-сервісах, де потрібний додатковий захист який не видимий користувачу [57].

В задачах розпізнавання облич та клавіатурного почерку використовується велика кількість моделей та методів для розпізнавання образів, як статистичні підходи, так і алгоритми глибокого навчання. Одними з найбільш широко використовуваних статистичних методів є метод головних компонент (PCA) та лінійний дискримінантний аналіз (LDA) [11-14, 58, 59]. PCA — це метод зменшення розмірності, який перетворює характеристики в новий набір некорельованих ознак, які називаються головними компонентами. У розпізнаванні облич PCA часто використовується для створення "власних облич" (eigenfaces), які є набором стандартизованих компонентів обличчя, що захоплюють найбільш значущі варіації в даних [60]. У задачах розпізнавання клавіатурного почерку PCA може бути використаний для виділення основних характеристик часу натискання та відпускання клавіш. LDA, з іншого боку, використовується для знаходження лінійної комбінації ознак, яка найкраще розділяє два або більше класи об'єктів або подій. Він максимізує відношення міжкласової дисперсії до внутрішньокласової дисперсії, роблячи його ефективним для завдань класифікації. Поєднання PCA та LDA має назву Fisherfaces, який спочатку використовує PCA для зменшення розмірності, а потім застосовує LDA для максимізації роздільності класів [61].

Іншими широко використовуваними методами є метод опорних векторів (SVM) [15-19], k-найближчих сусідів (k-NN) [20-22] та нейронні мережі [23-28]. SVM — це модель навчання з вчителем, яка використовується для завдань класифікації та регресії. Вона знаходить гіперплощину, яка найкраще розділяє дані на різні класи, роблячи його дуже ефективним для розпізнавання облич. Метод k-NN — це простий алгоритм навчання на основі екземплярів, який класифікує об'єкт на основі більшості голосів його k найближчих сусідів. Він може бути використаний для розпізнавання облич та клавіатурного почерку шляхом порівняння схожості нових реалізацій з базою даних відомих класів.

Нейронні мережі, особливо згорткові нейронні мережі (CNN), революціонізували розпізнавання образів. CNN використовують згорткові шари для автоматичного навчання просторових ієрархій ознак, роблячи їх дуже ефективними для обробки структурованих сіткових даних, таких як зображення [62]. Інші методи глибокого навчання, такі як рекурентні нейронні мережі (RNN) [63, 64] та генеративно-змагальні мережі (GAN) [65, 66] також можуть бути використані для розпізнавання облич та клавіатурного почерку, захоплюючи тимчасові залежності та генеруючи синтетичні дані для навчання та доповнення відповідно.

Методи, описані вище, зазвичай використовуються для мультикласової класифікації, але в контексті розпізнавання облич та клавіатурного почерку часто використовується однокласова класифікація [29-30]. Однокласова класифікація — це тип навчання, де модель навчається на даних, які належать лише одному, цільовому, класу, і потім використовується для виявлення відхилень або аномалій. Основна мета однокласової класифікації — визначити межі цільового класу і виявити будь-які відхилення від цих меж. У контексті розпізнавання облич та динаміки натискання клавіш однокласова класифікація використовується для виявлення несанкціонованого доступу або шахрайства. Наприклад, система може бути навчена розпізнавати обличчя або патерни натискання клавіш конкретного користувача і виявляти будь-які відхилення від цільового стану. Це дозволяє системам безпеки швидко реагувати на підозрілу діяльність і запобігати

несанкціонованому доступу. Однією з головних переваг однокласової класифікації є її здатність працювати з незбалансованими даними, де цільовий клас представлений значно більшою кількістю об'єктів, ніж представники інших класів. Крім того, однокласова класифікація може бути більш стійкою до змін у даних, оскільки вона фокусується на виявленні відхилень від цільового класу, а не на класифікації інших.

Для оцінювання точності та ймовірності розпізнавання у сфері однокласового розпізнавання використовують різні метрики. Найпоширенішими серед них є accuracy, precision, recall (sensitivity), specificity та F1 score [68, 69].

Метрики оцінювання базуються на даних матриці невідповідності, яка містить чотири основні показники. True Positive (TP) представляє випадки, коли модель правильно ідентифікує аномалію. False Positive (FP) відображає ситуації, коли модель хибно приймає цільову точку як аномалію. True Negative (TN) характеризує випадки, коли модель коректно визначає цільову точку. False Negative (FN) фіксує ситуації, коли модель неправильно приймає аномалію як цільову точку.

Accuracy оцінює загальну точність розпізнавання моделі, та обчислюється як відношення правильно розпізнаних точок до загальної кількості об'єктів у вибірці:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}.$$

Precision дозволяє визначити частку правильно виявлених аномалій серед усіх точок, розпізнаних як аномалії:

$$Precision = \frac{TP}{TP + FP}.$$

Recall, або sensitivity, показує частку аномалій, правильно ідентифікованих моделлю, серед усіх аномалій у вибірці:

$$Recall = \frac{TP}{TP + FN}.$$

Specificity надає додаткову інформацію, обчислюючи частку правильно ідентифікованих цільових точок серед усіх цільових точок:

$$Specificity = \frac{TN}{TN + FP}.$$

F1 score є гармонійним середнім між precision та recall, надаючи збалансовану оцінку:

$$F1\ score = 2 * \frac{Precision * Recall}{Precision + Recall}.$$

Кожна з цих метрик виконує певну функцію, надаючи детальне уявлення про загальну точність моделі та її здатність розрізняти аномалії і цільові об'єкти в різних аспектах. Precision та Recall надають детальне уявлення про точність моделі при роботі з даними, що не належать до цільового образу, тоді як F1 Score служить підсумовуючою оцінкою, що враховує обидва аспекти. Specificity, в свою чергу, дає розуміння, як модель працює з цільовими образами. Accuracy дозволяє швидко отримати загальну точність, враховуючі всі об'єкти, при тому ймовірність розпізнавання визначається сукупністю Recall та Specificity.

Популярні методи однокласової класифікації включають використання еліпсоїда прогнозування [31-33] та алгоритмів машинного навчання, такі як однокласова опорно-векторна машина (OCSVM) [34, 35], ізоляційний ліс (IF) [36] та автокодувальник (AE) [37, 38].

Однокласова опорно-векторна машина — це варіант методу опорних векторів, який навчається на даних цільового класу і використовує гіперплощину для визначення його меж. Будь-які точки, що знаходяться поза цими межами, вважаються аномальними. Ізоляційний ліс— це ансамблевий метод, який використовує випадкові дерева для ізоляції аномальних точок, та вимірює, наскільки легко можна ізолювати точку, і точки, які легко ізолюються, вважаються аномальними. Автокодувальник – це тип нейронних мереж, які навчаються стискати вхідні дані в представлення нижчої розмірності, а потім відновлювати їх. Аномалії виявляються шляхом порівняння вхідних даних з відновленими даними; великі відхилення вказують на наявність аномалій.

Еліпсоїд прогнозування — це геометричне представлення, яке застосовується в багатовимірному аналізі для визначення області, в якій, з

високою ймовірністю, перебувають більшість точок даних. Це ефективний інструмент для виявлення викидів, аналізу структури даних та перевірки припущень статистичних моделей. Еліпсоїд прогнозування будується на основі середнього вектора даних і коваріаційної матриці, яка визначає розміри та орієнтацію еліпсоїда. Коваріаційна матриця враховує кореляції між різними змінними, що дозволяє точніше описати розподіл даних у багатовимірному просторі.

Рисунок 1.1 ілюструє приклад еліпсоїда прогнозування з 3 змінними. Об'єкти в його межах класифікуються як цільовий клас, тоді як ті, що знаходяться за його межами, ідентифікуються як аномалії та належать до іншого класу.

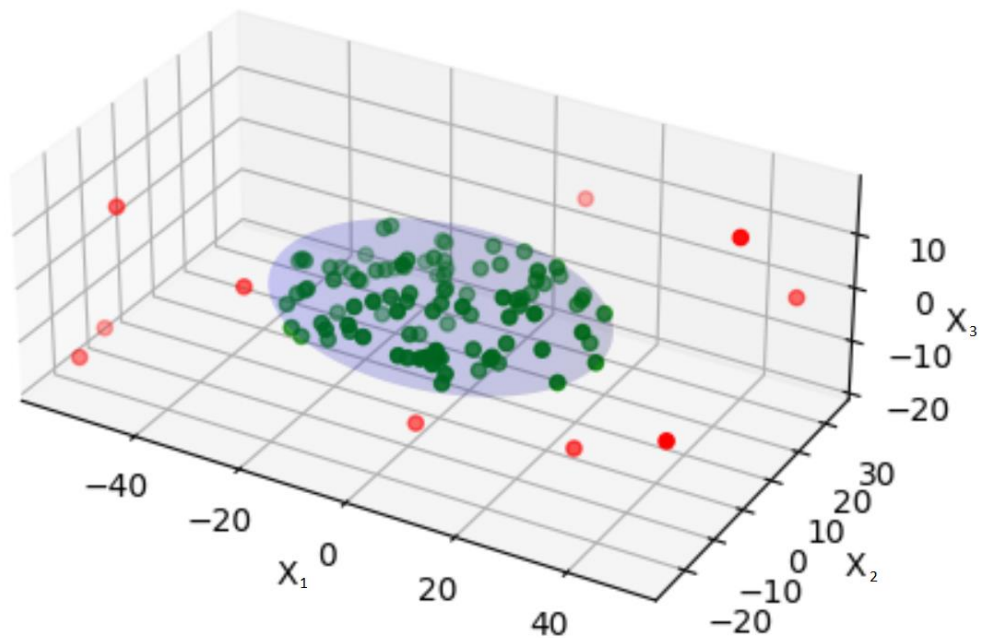


Рисунок 1.1 – Приклад тривимірного еліпсоїда прогнозування

Ліва частина еліпсоїда прогнозування представляє собою квадрат відстані Махаланобіса (КВМ) значення якого для i -ої точки даних обраховується за наступною формулою:

$$D^2(X_i) = (X_i - \bar{X})^T S_X^{-1} (X_i - \bar{X}), \quad (1.1)$$

де, $\bar{X}$ – вектор вибірових середніх; X_i – вектор даних в i -ій точці; S_X – вибіркова коваріаційна матриця даних, яка обчислюється як $S_X = \frac{1}{N} \sum_{i=1}^N (X_i - \bar{X})(X_i - \bar{X})^T$.

КВМ визначає, наскільки далеко точка знаходиться від центру розподілу даних, тим самим описуючи область, у якій точки даних задовольняють умову $D^2(X_i) < C$, де C — критичне значення, яке визначається на основі Хі-квадрат або F -розподілу для бажаного рівня значущості [70]. Таким чином усі спостереження одного класу лежать в межах еліпсоїда, а представники інших класів — поза межами.

Значення КВМ відповідає розподілу Хі-квадрат, у цьому випадку еліпсоїд набуває форму [71]:

$$(X - \bar{X})^T S_X^{-1} (X - \bar{X}) = \chi_{k,\alpha}^2, \quad (1.2)$$

де $\chi_{k,\alpha}^2$ - значення квантиля розподілу Хі-квадрат для k ступенів свободи та рівня значущості α , k дорівнює кількості характеристик.

F -розподіл є важливим інструментом багатовимірної аналізу, що використовується для визначення критичного значення еліпсоїда прогнозування та визначається як співвідношення двох масштабованих Хі-квадрат розподілів [72]. В такому випадку еліпсоїд прогнозування набуває вигляду [73]:

$$(X - \bar{X})^T S_X^{-1} (X - \bar{X}) = \frac{k(N^2-1)}{N(N-k)} F_{k,N-k,\alpha}, \quad (1.3)$$

де $F_{k,N-k,\alpha}$ - квантиль F -розподілу з k та $N - k$ ступенями свободи та рівнем значущості α . Значення квантиля F -розподілу, на відміну від Хі-квадрат, залежить не лише від кількості змінних, а і від розміру вибірки.

Зазвичай, для вибірок із більш ніж 20 спостереженнями результати застосування квантиля Хі-квадрат або F -розподілу, дають майже ідентичні результати [74], проте на практиці це не завжди підтверджується, тому є сенс побудувати відповідні моделі і порівняти їх точність.

Для визначення викидів у багатовимірних даних та однокласової класифікації, як правило, беруть рівень значущості, що дорівнює 0,005 [75, 76].

Побудова еліпсоїда прогнозування базується на припущенні, що дані мають багатовимірний нормальний розподіл даних [39], що призводить до зниження точності розпізнавання, якщо ця умова не виконується. Побудова моделей у вигляді еліпсоїда прогнозування відбувається на основі квантиля Хі-

квадрат [71], проте ще одним важливим чинником, який впливає на точність та ймовірність розпізнавання є врахування кількості точок, тобто використання квантиля F -розподілу.

Існуючі методи часто базуються на припущенні, що дані мають багатовимірний нормальний розподіл [39-41], однак, це не завжди виконується в реальному світі [42], що в результаті, призводить до зниження точності та ймовірності розпізнавання образів для негаусівських даних. Одним з ефективних способів подолання обмежень, пов'язаних з розподілом даних, є використання нормалізуючих перетворень, які допомагають коригувати розподіл даних, наближаючи його до нормального [43-46]. Нормалізуючі перетворення можуть бути одновимірними та багатовимірними [67]. Одновимірні перетворення застосовуються до кожної ознаки окремо та фокусуються на індивідуальних характеристиках, наближаючи їх до нормального розподілу. Багатовимірні нормалізуючі перетворення, на відміну від одновимірних, враховують взаємозв'язки між ознаками [75]. Вони розглядають дані як цілісну систему, де кожна ознака може впливати на інші.

Тому існує потреба у використанні ймовірнісних моделей, для побудови яких використовуються нормалізуючі перетворення. Ці моделі враховують відхилення від нормального розподілу, що призводить до підвищення точності та ймовірності розпізнавання образів.

1.2 Формування набору даних для розпізнавання облич

Вибір набору даних має вирішальне значення для побудови правил прийняття рішень. Для дослідження було обрано відомий набір даних Pins Face Recognition [77], який широко використовується при будові вирішальних правил для розпізнавання облич [78]. З метою забезпечення надійності та точності дослідження, було проведено фільтрацію набору даних, після чого здійснено додаткове завантаження фотографій з відкритих ресурсів для розширення набору.

Всього було отримано 400 високоякісних фотографій, по 200 для кожної з двох осіб.

1.2.1 Вибір інструменту для ідентифікації та вилучення характеристик обличчя на зображенні

Процес ідентифікації та вилучення характеристик об'єктів є важливим кроком у вирішенні проблем розпізнавання образів. Результат усього розпізнавання залежить від якості та реалізації обраного методу, оскільки на основі вибраних даних здійснюється подальша обробка ключових точок та розробка правил прийняття рішень.

Для вирішення цієї задачі використовуються багато різних моделей, які базуються на згорткових нейронних мережах, гістограмах орієнтованих градієнтів, каскадах Хаара та інших. Dlib, MTCNN та OpenCV - це найвідоміші та найбільш використовувані відкриті інструменти, які можна використовувати у проектах для пошуку обличчя на зображенні та формування вектора ознак для подальшої обробки [79, 80].

Метод Віюлі-Джонса реалізований у бібліотеці комп'ютерного зору OpenCV мови програмування Python, яка використовує попередньо навчені класифікатори для виявлення обличчя на зображенні. Метод використовує інтегральне представлення для обчислення яскравості прямокутної області зображення. Інтегральне представлення образу - це матриця, яка відповідає розміру оригінального зображення. Кожен її елемент зберігає суму інтенсивностей усіх пікселів зліва та вгорі від цього елемента [81]. Класичний метод Віюлі-Джонса використовує прямокутні ознаки, які називаються примітивами Хаара. Елементи матриці розраховуються за формулою:

$$I(x, y) = \sum_{x' \leq x, y' \leq y} i(x', y'),$$

де $I(x, y)$ – значення точки (x, y) інтегрального зображення; $i(x', y')$ – значення інтенсивності початкового зображення. Кожен елемент матриці $I(x, y)$ являє собою суму пікселів у прямокутнику від $i(0, 0)$ до $i(x, y)$. Розрахунок матриці

займає лінійний час, пропорційний числу пікселів у зображенні і його можна обчислити за формулою:

$$I(x, y) = i(x, y) - I(x - 1, y - 1) + I(x, y - 1) + I(x - 1, y).$$

У класичному методі Віоли-Джонса використовуються стандартні прямокутні ознаки (рис. 1.2.a) [82]. У розширеному методі Віоли-Джонса, який представлений у бібліотеці OpenCV, використовується розширений набір ознак (рис. 1.2.b) [83].

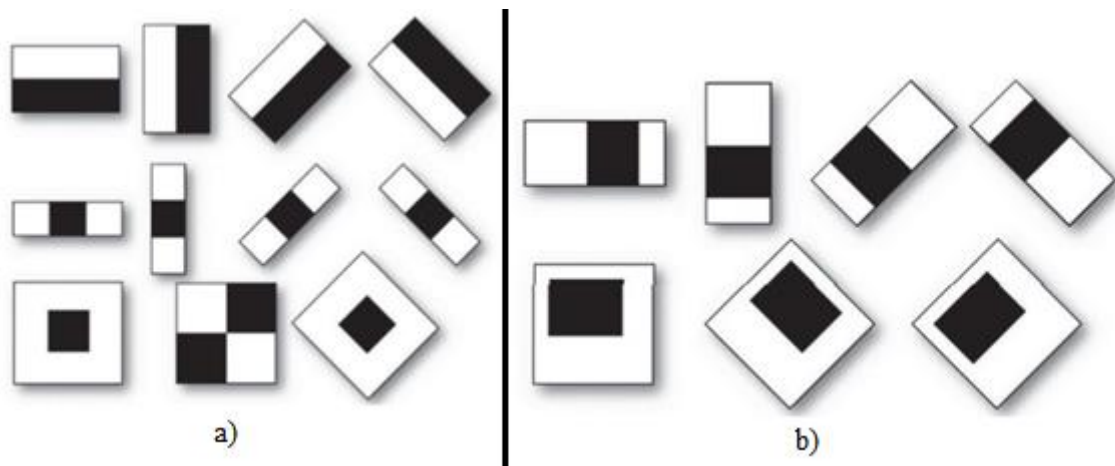


Рисунок 1.2 – Примітиви Хаара - а) основні ознаки - б) додаткові ознаки

Ознаки обчислюються за формулою:

$$F = U - V,$$

де U – сума значень яскравостей точок, що закриваються світлою частиною ознаки, а V – сума значень яскравостей точок, що закриваються темною частиною ознаки. Для обчислення використовується поняття інтегрального зображення. Хаарподібні ознаки описують значення перепаду яскравості по осі X та Y зображення відповідно.

Слід зазначити, що цей детектор має вкрай низьку ймовірність помилкового виявлення обличчя на зображеннях високої якості. Метод добре працює і виявляє риси обличчя навіть при спостереженні об'єкта під невеликим кутом приблизно до 30° . При куті нахилу більше 30° ймовірність виявлення обличчя різко падає, що значно ускладнює або унеможлиблює використання алгоритму в сучасних виробничих системах з урахуванням їхніх потреб.

MTCNN - це скорочення від Multi-task Cascaded Convolutional Neural Networks, що означає багатозадачні каскадні згорткові нейронні мережі. Цей метод складається з трьох етапів (рис. 1.3) [84], кожен з яких використовує окрему згорткову мережу (рис. 1.4) [85], які здатні розпізнавати обличчя та положення ключових точок, таких як очі, ніс і рот. На етапі підготовки метод масштабує вхідне зображення, таким чином, щоб отримати набір різних масштабів, що зветься пірамідою вхідних зображень, яка дозволяє MTCNN виявляти обличчя різних розмірів та положень на зображенні.

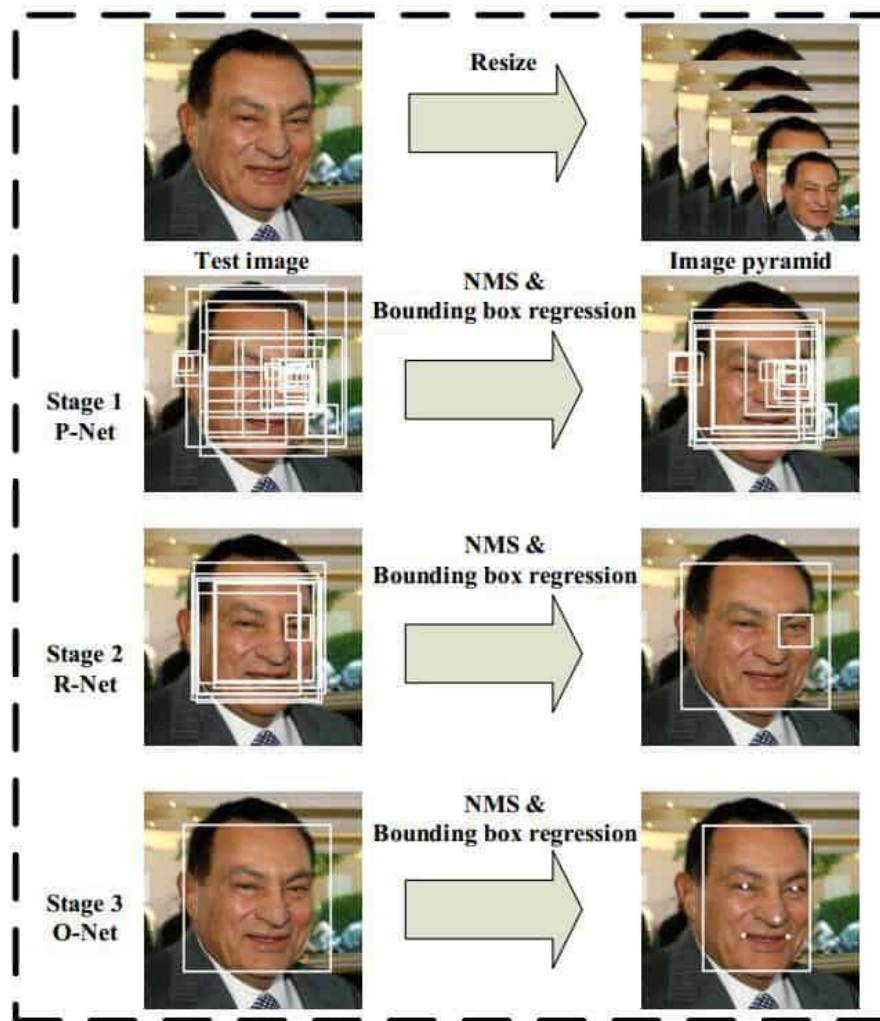


Рисунок 1.3 – Етапи роботи MTCNN [84]

На першому етапі використовується мережа пропозицій (P-Net). Ця мережа швидко переглядає зображення і виділяє всі можливі місця, де може бути обличчя. Вона також коригує форму і розмір цих областей, щоб вони були

максимально схожими на образ. Після цього вона об'єднує всі місця, які перекриваються одне з одним, і залишає тільки ті, які мають найбільшу ймовірність бути обличчями.

На другому етапі використовується мережа уточнення (R-Net). Ця мережа більш складна, ніж попередня, і вона приймає на вхід результат роботи P-Net. R-Net відкидає ті області, які не схожі на обличчя, і знову коригує форму і розмір тих, які залишилися. Вона також використовує спеціальний метод, який називається немаксимальне придушення (NMS) [86], щоб зменшити кількість областей на зображенні, які перекриваються.

На третьому етапі використовується мережа виводу (O-Net). Ця мережа найскладніша з усіх і вона знову перевіряє всі області, які обробила R-Net. Вона ще раз відкидає ті, які не схожі на обличчя, і ще раз коригує форму і розмір тих, які залишилися, але O-Net також знаходить і відмічає 5 ключових точок обличчя, а саме, де знаходяться очі, ніс і рот. Це допомагає в розробці систем розпізнавання, а також для вирівнювання обличчя.

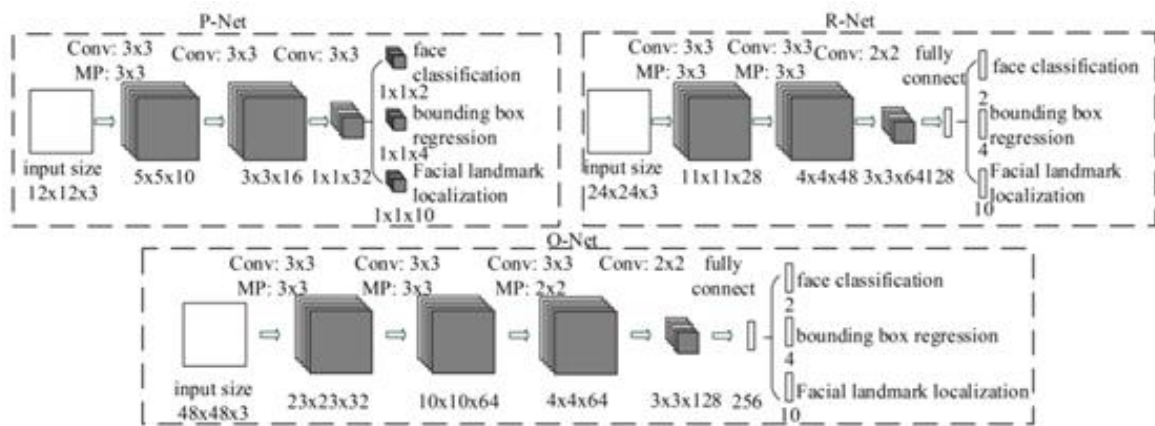


Рисунок 1.4 – Структура MTCNN [85]

MTCNN має високу точність знаходження об'єкту на зображенні, так як використовує каскадні згорткові мережі, та здатна виявляти обличчя різних розмірів, орієнтацій, виразів та освітлення, а також ті, які частково закриті іншими об'єктами або повернуті на невеликий кут.

Dlib - це бібліотека загального призначення для крос-платформного програмування розроблена мовою C++. Вона містить алгоритми машинного навчання та інструменти для створення складного програмного забезпечення. Однією з функцій Dlib є виявлення облич та ключових точок обличчя, які використовуються для розпізнавання, вирівнювання, аналізу емоцій та інших застосувань [87].

Dlib використовує HOG + Linear SVM для виявлення облич. HOG - це скорочення від Histogram of Oriented Gradients, що означає гістограма орієнтованих градієнтів. Це метод отримання ознак, який описує форму і текстуру об'єкта за допомогою статистики напрямків градієнтів пікселів. Linear SVM - це скорочення від Linear Support Vector Machine, що означає лінійна машина опорних векторів [88].

Dlib HOG + Linear SVM працює наступним чином, спочатку зображення перетворюється в сірі тони і масштабується до меншого розміру, щоб зменшити обчислювальну складність. Потім воно розбивається на дрібні клітинки, для кожної з яких обчислюється гістограма орієнтованих градієнтів (HOG). Гістограма описує напрямки і силу зміни яскравості пікселів в клітинці. Далі зображення розбивається на більші блоки, розміром по декілька клітинок з попереднього кроку. Для кожного блоку обчислюється вектор ознак, який складається з гістограм усіх клітинок в блоці. Вектор ознак нормалізується, щоб зменшити вплив зміни освітлення. Таким чином, отримується набір векторів ознак для всього зображення, який називається HOG дескриптором. Створюється піраміда, яка складається з декількох копій зображення з різними розмірами, це дозволяє виявляти обличчя різних масштабів на оригінальному зображенні, які не були знайдені раніше. По кожному зображенню в піраміді проводиться ковзаюче вікно, яке має фіксований розмір. Для кожного вікна обчислюється HOG дескриптор і подається на вхід лінійної машини опорних векторів (SVM), яка є класифікатором, що навчений розрізняти обличчя і не обличчя. Якщо SVM видає позитивний результат, то область вважається кандидатом на обличчя. На шостому кроці, всі кандидати на обличчя, які були виявлені на різних рівнях

піраміди, переводяться в координати оригінального зображення і об'єднуються за допомогою методу немаксимального придушення (NMS), який відкидає області, що перекриваються, та залишає ті, які з високою ймовірністю є обличчями.

Після виявлення облич, Dlib використовує алгоритм Ensemble of Regression Trees (ERT) для визначення положення 68 ключових точок обличчя, таких як очі, ніс, рот, брови та щелепа [89] (рис. 1.5).

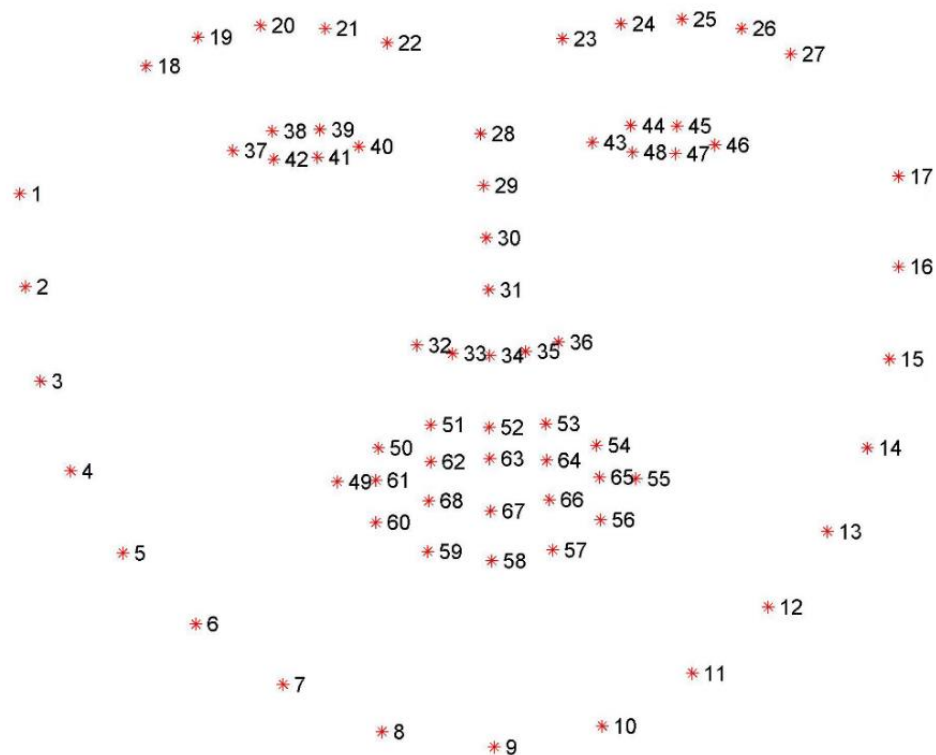


Рисунок 1.5 – Ключові точки обличчя, які розпізнає Dlib [90]

На якісних зображеннях усі методи працюють досить точно, але при спотвореннях метод Віоли-Джонса показав найгірший результат. MTCNN і Dlib працюють майже однаково, але MTCNN краще розпізнає маленькі зображення та за слабкого освітлення, і має вищу точність, ніж дескриптор функції Dlib. Він також добре справляється з обличчями, що дивляться не тільки прямо, але і під непростими кутами. Однак, MTCNN потребує більше ресурсів, оскільки працює за допомогою нейронних мереж, і знаходить лише 5 точок, а саме очі, ніс і краї рота. OpenCV не має вбудованого механізму отримання координат ключових

точок, тому найзручнішим в цьому плані є Dlib, оскільки він має функціонал для отримання масиву координат із 68 точок, і працює швидше, ніж MTCNN.

Загалом MTCNN вважається найточнішим детектором обличчя, а Dlib відомий своєю швидкістю та ефективністю. При цьому Dlib є найбільш інформативним, оскільки дозволяє отримати найбільшу кількість ключових точок обличчя. Залежно від застосування, один може бути більш придатним, ніж інший.

На основі тестових даних було перевірено швидкість роботи інструментів. Виявилося, що найшвидше працює метод Віюлі-Джонса з результатом 11,25 кадрів в секунду, потім Dlib - 9,41 кадрів в секунду і MTCNN - 4,92 кадрів в секунду.

Враховуючи високу швидкість, хорошу точність знаходження об'єкту на зображенні і кількість ключових точок, які можна розпізнати, Dlib було обрано найоптимальнішою бібліотекою для розробки програми мовою Python для пошуку та отримання ключових точок обличчя на зображенні.

1.2.2 Формування векторів характеристик облич

Для створення векторів характеристик був розроблений застосунок мовою Python з використанням бібліотеки Dlib. Алгоритм роботи програми складається з кількох кроків, які виконуються послідовно для досягнення бажаних результатів. На першому етапі програма знаходить обличчя на вхідному зображенні за допомогою алгоритму пошуку Dlib. На другому кроці видаляється задній фон, це зменшує розмиття та шум у зображенні, що можуть заважати розпізнаванню обличчя. Потім обличчя вирівнюється таким чином, щоб очі були на одному рівні, це допомагає створити однорідне та пропорційне обличчя на всьому зображенні, що полегшує визначення ключових точок, а також дозволяє уникнути помилок, пов'язаних з положенням обличчя на вхідному зображенні [91, 92]. На останньому кроці програма отримує масив ключових точок та обчислює відстані між ними для кожного зображення.

Важливість обирання правильного набору ключових точок полягає в тому, що вони впливають на якість та точність розпізнавання обличчя. Якщо набір ключових точок недостатній або неправильний, це може призвести до помилок, спотворень або невідповідностей у результатах. Також, якщо характеристики не стійкі до змін освітлення, виразу обличчя або окулярів, це може збільшити шум або варіацію у даних. Різні автори використовують різні набори ключових точок для розпізнавання обличчя, в залежності від їхньої мети та методу [93-95]. Але важливо включати незмінні об'єкти, як ніс, очі, підборіддя та рот, оскільки вони є найбільш виразними та розпізнавальними рисами обличчя. Ці об'єкти не змінюються з часом або виразом обличчя, тому вони можуть служити як надійні орієнтири для розпізнавання обличчя.

Формування вектору характеристик відбувається на основі отриманих відстаней у декілька кроків. Проаналізувавши наявні роботи, було вирішено зосередитися на 17 найбільш інформативних ключових точках, що використовуються для вимірювання 16 різних відстаней обличчя (рис. 1.6).

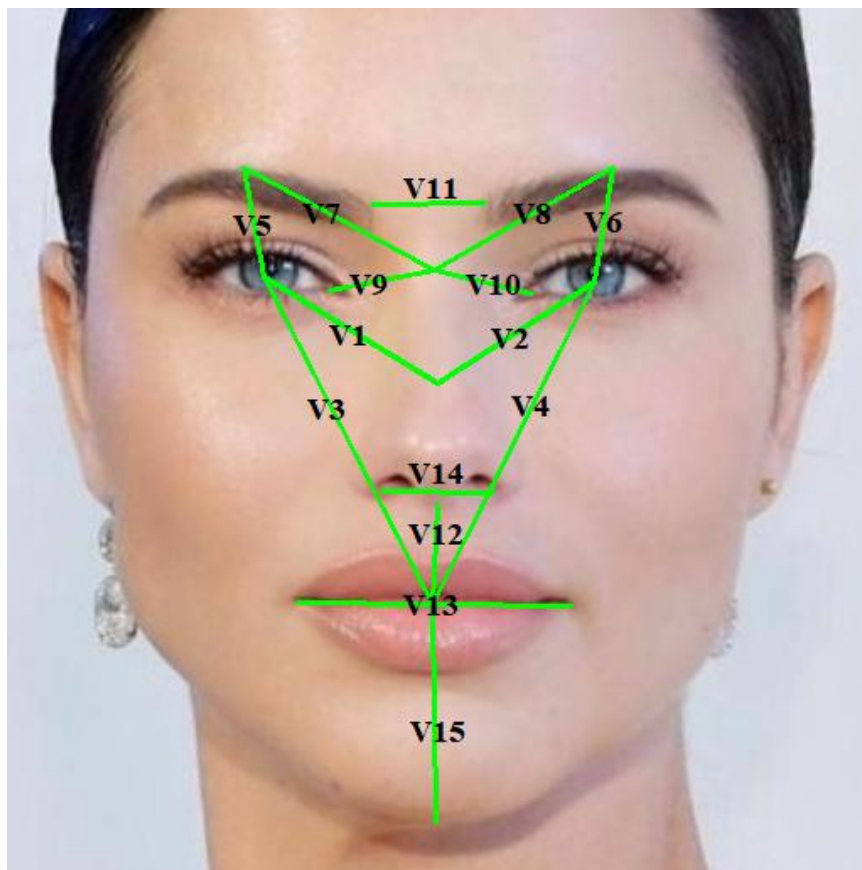


Рисунок 1.6 – Відстані обличчя для формування вектору характеристик

В першу чергу симетричні характеристики усереднюються, вираховуються середні значення відстаней між симетричними рисами обличчя [96], наприклад від очей до носа, від кутів рота до підборіддя або від брів до очей. Ці середні значення використовуються як характеристики для розпізнавання обличчя, замість відстаней з кожної сторони окремо. Це дозволяє отримати більш стабільні значення, які не залежать від асиметрії обличчя та мінімізують вплив повороту голови. Наступним важливим кроком є ділення всіх відстаней між ключовими точками на обличчі на певну сталу одиницю. Це дозволяє отримати незалежні від масштабу характеристики, які можуть бути порівняні між різними зображеннями, також це допомагає усунути вплив зміни положення голови або кута зору на відстані між рисами обличчя [97, 98]. Зазвичай використовують нормування характеристик за допомогою ділення всіх відстаней на відстань між очима. В результаті вектор характеристик складається з 10 змінних (таблиця 1.1).

Таблиця 1.1 – Опис вектору характеристик

Номер характеристики	Формула	Опис
1	$(v1 + v2)/2d$	Усереднена відстань від очей до середини носа
2	$(v3 + v4)/2d$	Усереднена відстань від очей до центру рота
3	$(v5 + v6)/2d$	Усереднена відстань від очей до центру брів
4	$(v7 + v8)/2d$	Усереднена відстань від брів до верху носа
5	$(v9 + v10)/2d$	Усереднена відстань від куточків очей до верху носа
6	$v11/d$	Відстань між бровами
7	$v12/d$	Відстань між носом і серединою рота
8	$v13/d$	Ширина рота
9	$v14/d$	Ширина носу
10	$v15/d$	Відстань від рота до підборіддя

В результаті обробки наявних зображень розробленим застосунком, було створено по 200 векторів характеристик для кожної з двох осіб.

1.2.3 Тестові характеристики отриманих вибірок

Вектор вибірових середніх значень вибірки особи 1 дорівнює $\bar{X} = \{0,63311, 1,16902; 0,30584; 0,62898; 0,30033; 0,36469; 0,30606; 0,80355; 0,36186; 0,63804\}$. Коваріаційна матриця вибірки в таблиці 1.2, діапазони характеристик – в таблиці 1.3.

Таблиця 1.2 – Коваріаційна матриця вибірки особи 1

0,0 ² 12	0,0 ² 11	-0,0 ³ 76	-0,0 ⁴ 94	-0,0 ³ 14	0,0 ⁵ 38	-0,0 ³ 58	0,0 ³ 15	0,0 ³ 46	-0,0 ² 12
0,0 ² 11	0,0 ² 24	-0,0 ⁴ 45	0,0 ³ 24	-0,0 ³ 18	0,0 ⁴ 89	0,0 ³ 5	0,0 ³ 13	0,0 ³ 65	0,0 ³ 39
-0,0 ³ 76	-0,0 ⁴ 45	0,0 ² 14	0,0 ³ 56	0,0 ⁴ 3	0,0 ³ 19	0,0 ³ 85	-0,0 ³ 19	-0,0 ³ 17	0,0 ² 15
-0,0 ⁴ 94	0,0 ³ 24	0,0 ³ 56	0,0 ³ 42	-0,0 ⁴ 2	0,0 ³ 2	0,0 ³ 31	-0,0 ³ 18	-0,0 ⁴ 49	0,0 ³ 54
-0,0 ³ 14	-0,0 ³ 18	0,0 ⁴ 3	-0,0 ⁴ 2	0,0 ⁴ 7	0,0 ⁵ 18	0,0 ⁴ 2	0,0 ⁴ 21	-0,0 ³ 1	0,0 ⁴ 7
0,0 ⁵ 38	0,0 ⁴ 89	0,0 ³ 19	0,0 ³ 2	0,0 ⁵ 18	0,0 ³ 69	0,0 ³ 15	0,0 ³ 11	0,0 ⁴ 46	0,0 ⁴ 64
-0,0 ³ 58	0,0 ³ 5	0,0 ³ 85	0,0 ³ 31	0,0 ⁴ 2	0,0 ³ 15	0,0 ² 12	-0,0 ³ 15	-0,0 ⁴ 69	0,0 ² 16
0,0 ³ 15	0,0 ³ 13	-0,0 ³ 19	-0,0 ³ 18	0,0 ⁴ 21	0,0 ³ 11	-0,0 ³ 15	0,0 ² 33	0,0 ² 11	0,0 ² 13
0,0 ³ 46	0,0 ³ 65	-0,0 ³ 17	-0,0 ⁴ 49	-0,0 ³ 1	0,0 ⁴ 46	-0,0 ⁴ 69	0,0 ² 11	0,0 ² 1	0,0 ³ 21
-0,0 ² 12	0,0 ³ 39	0,0 ² 15	0,0 ³ 54	0,0 ⁴ 7	0,0 ⁴ 64	0,0 ² 16	0,0 ² 13	0,0 ³ 21	0,0 ² 45

Таблиця 1.3 – Діапазони характеристик вибірки особи 1

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Min	0,5489	1,0703	0,2289	0,5778	0,2729	0,3109	0,2138	0,6987	0,2933	0,4681
Max	0,7085	1,3166	0,4232	0,6935	0,3217	0,4316	0,3996	1,0575	0,4608	0,8095

Вектор вибірових середніх значень вибірки особи 2 дорівнює $\bar{X} = \{0,62727; 1,24348; 0,31716; 0,67790; 0,30163; 0,40401; 0,35468; 0,87551; 0,41770; 0,80282\}$. Коваріаційна матриця вибірки в таблиці 1.4, діапазони характеристик – в таблиці 1.5.

Таблиця 1.4 – Коваріаційна матриця вибірки особи 2

0,0 ² 15	0,0 ² 14	-0,0 ³ 53	0,0 ⁴ 73	-0,0 ³ 1	0,0 ³ 36	-0,0 ³ 54	-0,0 ³ 36	0,0 ³ 19	-0,0 ² 14
0,0 ² 14	0,0 ² 26	-0,0 ³ 3	0,0 ³ 11	-0,0 ³ 12	-0,0 ⁴ 48	0,0 ³ 46	-0,0 ² 13	0,0 ³ 11	-0,0 ³ 32
-0,0 ³ 53	-0,0 ³ 3	0,0 ² 1	0,0 ³ 6	-0,0 ⁴ 18	0,0 ³ 73	0,0 ³ 27	0,0 ³ 7	-0,0 ⁴ 31	0,0 ³ 75
0,0 ⁴ 73	0,0 ³ 11	0,0 ³ 6	0,0 ³ 63	-0,0 ⁴ 42	0,0 ² 11	0,0 ⁵ 43	0,0 ³ 8	-0,0 ⁴ 13	0,0 ³ 39
-0,0 ³ 1	-0,0 ³ 12	-0,0 ⁴ 18	-0,0 ⁴ 42	0,0 ⁴ 65	-0,0 ⁴ 17	0,0 ⁴ 23	-0,0 ³ 1	-0,0 ⁴ 57	0,0 ⁴ 91
0,0 ³ 36	-0,0 ⁴ 48	0,0 ³ 73	0,0 ² 11	-0,0 ⁴ 17	0,0 ² 28	-0,0 ³ 38	0,0 ² 13	-0,0 ³ 24	0,0 ³ 16
-0,0 ³ 54	0,0 ³ 46	0,0 ³ 27	0,0 ⁵ 43	0,0 ⁴ 23	-0,0 ³ 38	0,0 ³ 98	-0,0 ³ 65	-0,0 ³ 21	0,0 ² 1
-0,0 ³ 36	-0,0 ² 13	0,0 ³ 7	0,0 ³ 8	-0,0 ³ 1	0,0 ² 13	-0,0 ³ 65	0,0 ² 75	0,0 ² 12	0,0 ² 22
0,0 ³ 19	0,0 ³ 11	-0,0 ⁴ 31	-0,0 ⁴ 13	-0,0 ⁴ 57	-0,0 ³ 24	-0,0 ³ 21	0,0 ² 12	0,0 ³ 72	0,0 ⁴ 9
-0,0 ² 14	-0,0 ³ 32	0,0 ³ 75	0,0 ³ 39	0,0 ⁴ 91	0,0 ³ 16	0,0 ² 1	0,0 ² 22	0,0 ⁴ 9	0,0 ² 45

Таблиця 1.5 – Діапазони характеристик вибірки особи 2

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Min	0,5467	1,0959	0,235	0,6123	0,28	0,2733	0,2772	0,7142	0,3606	0,6209
Max	0,7448	1,4012	0,4268	0,7411	0,3271	0,5519	0,4868	1,165	0,5025	0,922

В майбутньому вибірка особи 1 буде розділена на навчальну та тренувальну, в той час як вибірка особи 2 буде використовуватися як тестова для перевірки точності та ймовірності розпізнавання обличчя іншої особи створеною моделлю.

1.3 Формування набору даних для розпізнавання клавіатурного почерку

1.3.1 Формування векторів характеристик

У контексті розпізнавання клавіатурного почерку якість набору даних має вирішальне значення для ефективності та точності моделі. Такий набір зазвичай містить детальні записи про поведінку користувача при наборі тексту, фіксуючи різні часові параметри, як-от тривалість натискання клавіші і інтервали між

натисканням і відпусканням клавіш. Розмір та різноманітність патернів, представлених у наборі даних, суттєво впливають на якість роботи алгоритмів розпізнавання. Більш комплексні набори даних дають змогу проводити глибший аналіз і краще навчати моделі, так як мінімізують людський фактор, що дозволяє чітко розрізняти цільові та аномальні патерни набору, підвищуючи надійність системи розпізнавання.

Набір даних динаміки клавіатурного почерку CMU [99], використаний у цій роботі, надає детальні записи набору тексту для 51 респондента, кожен із яких вводив статичний пароль: “.tie5Roanl”. Збір даних передбачав проведення восьми сесій для кожного респондента, при цьому між сесіями була щонайменше доба. Під час кожної сесії респонденти вводили пароль 50 разів, що дало 400 повторень на респондента. У підсумку було зібрано 20,400 зразків для всіх учасників, що забезпечує достатній обсяг даних для аналізу та мінімізує людський фактор у вигляді самопочуття, настрою, тощо.

Структура набору включає ідентифікатор респондента, номер сесії, номер повторення та загалом 31 часову характеристику. У наборі використано стандартизовані позначення для різних часових характеристик клавіатурної динаміки. Колонки, позначені як H.key, відображають тривалість натискання конкретної клавіші (час від початку натискання до відпускання клавіші). Колонки DD.key1.key2 представляють інтервали між натисканням двох послідовних клавіш (keydown-keydown час), тоді як колонки UD.key1.key2 вказують на keyup-keydown час, який вимірює інтервал між відпусканням однієї клавіші та натисканням наступної. Важливо зазначити, що UD часи можуть іноді бути негативними, а сума H часу та UD часу для певного діграма дорівнює відповідному DD часу.

Для спрощення моделі було вирішено зосередитися на 9 ключових характеристиках, тому вектор ознак приймає форму: $X = \{ H.t, H.i, H.e, H.5, H.R, H.o, H.a, H.n, H.l \}$.

1.3.2 Тестові характеристики отриманих вибірок

Зі всього набору випадковим чином було вибрано дані з ідентифікатором s015 для побудови моделі в якості цільового образу. Водночас s004 буде використано для формування тестової вибірки, щоб перевірити розпізнавання клавіатурного почерку іншого користувача.

Вектор вибірових середніх значень вибірки s015 дорівнює $\bar{X} = \{0,07549; 0,07038; 0,07824; 0,06275; 0,06923; 0,08839; 0,08639; 0,07513; 0,07500\}$. Коваріаційна матриця вибірки в таблиці 1.6, діапазони характеристик – в таблиці 1.7.

Таблиця 1.6 – Коваріаційна матриця вибірки s015

0,0 ³ 22	0, 0 ⁴ 34	0, 0 ⁴ 51	0, 0 ⁴ 12	0, 0 ⁵ 76	0, 0 ⁴ 13	0, 0 ⁴ 31	-0, 0 ⁵ 37	-0, 0 ⁴ 24
0, 0 ⁴ 34	0, 0 ³ 21	-0, 0 ⁴ 17	-0, 0 ⁵ 11	0, 0 ⁴ 1	0, 0 ⁴ 18	0, 0 ⁴ 42	-0, 0 ⁴ 14	0, 0 ⁵ 12
0, 0 ⁴ 51	-0, 0 ⁴ 17	0,0 ³ 25	-0, 0 ⁵ 19	0, 0 ⁵ 41	0, 0 ⁵ 12	-0, 0 ⁴ 25	0, 0 ⁴ 21	0, 0 ⁴ 34
0, 0 ⁴ 12	-0, 0 ⁵ 11	-0, 0 ⁵ 19	0,0 ³ 2	0, 0 ⁴ 11	-0, 0 ⁵ 28	0, 0 ⁴ 23	-0, 0 ⁴ 15	-0, 0 ⁵ 72
0, 0 ⁵ 76	0, 0 ⁴ 1	0, 0 ⁵ 41	0, 0 ⁴ 11	0,0 ³ 14	0, 0 ⁴ 16	-0, 0 ⁵ 51	0, 0 ⁵ 47	0, 0 ⁵ 4
0, 0 ⁴ 13	0, 0 ⁴ 18	0, 0 ⁵ 12	-0, 0 ⁵ 28	0, 0 ⁴ 16	0, 0 ³ 14	-0, 0 ⁴ 11	0, 0 ⁵ 39	0, 0 ⁵ 13
0, 0 ⁴ 31	0, 0 ⁴ 42	-0, 0 ⁴ 25	0, 0 ⁴ 23	-0, 0 ⁵ 51	-0, 0 ⁴ 11	0, 0 ³ 35	-0, 0 ⁴ 39	0, 0 ⁵ 35
-0, 0 ⁵ 37	-0, 0 ⁴ 14	0, 0 ⁴ 21	-0, 0 ⁴ 15	0, 0 ⁵ 47	0, 0 ⁵ 39	-0, 0 ⁴ 39	0, 0 ³ 28	0, 0 ⁵ 79
-0, 0 ⁴ 24	0, 0 ⁵ 12	0, 0 ⁵ 34	-0, 0 ⁵ 72	0, 0 ⁵ 4	0, 0 ⁵ 13	0, 0 ⁵ 35	0, 0 ⁵ 79	0, 0 ³ 21

Таблиця 1.7 – Діапазони характеристик вибірки s015

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
Min	0,018	0,0138	0,0304	0,0061	0,0164	0,0357	0,0238	0,0123	0,0077
Max	0,1526	0,109	0,127	0,0985	0,0961	0,146	0,1302	0,1465	0,1109

Вектор вибірових середніх значень вибірки s004 дорівнює $\bar{X} = \{0,10411; 0,10615; 0,08058; 0,08757; 0,13740; 0,10035; 0,12935; 0,08527; 0,09369\}$.

Коваріаційна матриця вибірки в таблиці 1.8, діапазони характеристик – в таблиці 1.9.

Таблиця 1.8 – Коваріаційна матриця вибірки s004

0,0 ³ 39	0, 0 ⁴ 75	0, 0 ⁴ 79	0, 0 ³ 24	0, 0 ³ 18	0, 0 ⁴ 86	0, 0 ³ 11	0, 0 ⁴ 12	0, 0 ⁵ 21
0, 0 ⁴ 75	0, 0 ³ 14	0, 0 ⁴ 29	0, 0 ⁴ 82	0, 0 ⁴ 94	0, 0 ⁴ 52	0, 0 ⁴ 55	0, 0 ⁴ 11	0, 0 ⁵ 77
0, 0 ⁴ 79	0, 0 ⁴ 29	0, 0 ³ 17	0, 0 ⁴ 36	0, 0 ⁴ 43	0, 0 ⁴ 13	0, 0 ⁴ 3	-0, 0 ⁶ 6	-0, 0 ⁵ 87
0, 0 ³ 24	0, 0 ⁴ 82	0, 0 ⁴ 36	0, 0 ³ 43	0, 0 ³ 25	0, 0 ³ 11	0, 0 ³ 12	0, 0 ⁴ 18	0, 0 ⁴ 14
0, 0 ³ 18	0, 0 ⁴ 94	0, 0 ⁴ 43	0, 0 ³ 25	0, 0 ³ 53	0, 0 ⁴ 72	0, 0 ³ 12	0, 0 ⁴ 31	0, 0 ⁵ 38
0, 0 ⁴ 86	0, 0 ⁴ 52	0, 0 ⁴ 13	0, 0 ³ 11	0, 0 ⁴ 72	0, 0 ³ 2	0, 0 ⁴ 66	0, 0 ⁵ 7	0, 0 ⁴ 13
0, 0 ³ 11	0, 0 ⁴ 55	0, 0 ⁴ 3	0, 0 ³ 12	0, 0 ³ 12	0, 0 ⁴ 66	0, 0 ³ 34	0, 0 ⁵ 39	0, 0 ⁵ 65
0, 0 ⁴ 12	0, 0 ⁴ 11	-0, 0 ⁶ 6	0, 0 ⁴ 18	0, 0 ⁴ 31	0, 0 ⁵ 7	0, 0 ⁵ 39	0, 0 ³ 12	0, 0 ⁴ 24
0, 0 ⁵ 21	0, 0 ⁵ 77	-0, 0 ⁵ 87	0, 0 ⁴ 14	0, 0 ⁵ 38	0, 0 ⁴ 13	0, 0 ⁵ 65	0, 0 ⁴ 24	0, 0 ³ 14

Таблиця 1.9 – Діапазони характеристик вибірки s004

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
Min	0,051	0,0732	0,0082	0,0301	0,0557	0,0605	0,0457	0,0515	0,0325
Max	0,1664	0,1452	0,151	0,1584	0,2207	0,1502	0,1954	0,1286	0,1288

Всього було відібрано по 400 векторів характеристик для кожного з двох наборів з ідентифікаторами s004 та s015.

1.4 Обґрунтування необхідності проведення досліджень за обраною темою

Необхідність проведення цього дослідження обумовлена низкою важливих чинників, пов'язаних із недостатньою точністю існуючих методів для розпізнавання образів, оскільки, відомо, що більшість статистичних методів, які використовуються для розпізнавання образів передбачають, що дані мають багатовимірний нормальний розподіл.

Для оцінювання відхилення багатовимірної розподілу даних від нормального використано тест Мардіа [100]. Він заснований на аналізі багатовимірних асиметрії β_1 та ексцесу β_2 даних, які є показниками того, наскільки дані відхиляються від нормального розподілу та обчислюються за наступними формулами:

$$\beta_1 = \frac{1}{N^2} \sum_{i=1}^N \sum_{j=1}^N \{(\mathbf{X}_i - \bar{\mathbf{X}})^T \mathbf{S}_X^{-1} (\mathbf{X}_j - \bar{\mathbf{X}})\}^3; \quad (1.4)$$

$$\beta_2 = \frac{1}{N} \sum_{j=1}^N \{(\mathbf{X}_j - \bar{\mathbf{X}})^T \mathbf{S}_X^{-1} (\mathbf{X}_j - \bar{\mathbf{X}})\}^2. \quad (1.5)$$

де, $\bar{\mathbf{X}}$ – вектор вибірових середніх; $\mathbf{X}_i$ – вектор даних в i -ій точці; $\mathbf{X}_j$ – вектор даних в j -ій точці; $\mathbf{S}_X$ – вибіркова коваріаційна матриця даних, яка обчислюється як $\mathbf{S}_X = \frac{1}{N} \sum_{i=1}^N (\mathbf{X}_i - \bar{\mathbf{X}})(\mathbf{X}_i - \bar{\mathbf{X}})^T$.

Величина асиметрії, розрахована за допомогою формули (1.4), помножена на $N/6$, розподілена за законом Хі-квадрат з $p(p+1)(p+2)/6$ ступенями свободи, де p – це кількість змінних, а величина ексцесу, розрахована за допомогою формули (1.5), розподілена згідно із законом нормального розподілу із математичним сподіванням $p(p+2)$ і дисперсією $8p(p+2)/N$.

Відповідно до тесту Мардіа багатовимірний розподіл вибірки для розпізнавання обличчя особи 1 відхиляється від нормального закону, тому що значення тестової статистики розподілу даної вибірки для багатовимірної асиметрії $\beta_1 N/6$, яка становить 403,88, перевищує значення квантиля розподілу Хі-квадрат, що становить 277,77 для 220 ступенів свободи та рівня значущості 0,005. Натомість, тестова статистика розподілу даної вибірки для багатовимірної ексцесу β_2 , яка становить 121,41, не перевищує значення квантиля нормального розподілу, яке становить 127,97 для математичного сподівання 120 і дисперсії 9,6 та рівня значущості 0,005.

Тест Мардіа показав, що набір даних клавіатурного почерку з ідентифікатором s015 демонструє відхилення від багатовимірної нормальності, оскільки тестова статистика для багатовимірної асиметрії $\beta_1 N/6$, що дорівнює 391,54, перевищує значення квантиля Хі-квадрат 215,53 для 165 ступенів свободи

та рівня значущості 0,005. Подібним чином тестова статистика для багатовимірного ексцесу β_2 зі значенням 113,32 перевищує значення квантиля нормального розподілу 102,62 із середнім значенням 99, дисперсією 1,98 і рівнем значущості 0,005.

Ще один важливий фактор, що може впливати на точність розпізнавання, — це наявність викидів, які спотворюють статистичні оцінки. Викиди впливають на результати аналізу, тому їх усунення є важливим етапом перед проведенням розпізнавання. Для виявлення аномальних точок застосовується квадрат відстані Махаланобіса (1.1). Водночас, цей метод базується на припущенні, що розподіл даних є нормальним.

Вибірки клавіатурного почерку і облич мають розподіл, що відхиляється від багатовимірного нормального, що напряду впливає на можливість застосування підходів до виявлення аномалій та точність результуючої моделі, побудованої на негаусівських даних. Таким чином, виникає потреба у проведенні досліджень за обраною темою.

1.5 Висновки до розділу 1

У даному розділі наведено огляд існуючих методів та моделей для розпізнавання образів, здійснено обґрунтування необхідності проведення досліджень за обраною темою.

1) Для мультикласового розпізнавання використовують лінійний дискримінантний аналіз, метод опорних векторів, k-найближчих сусідів та нейронні мережі. Для розпізнавання облич та клавіатурного почерку характерне використання однокласової класифікації, де модель навчається на даних, які належать лише одному, цільовому, класу, і потім використовується для виявлення відхилень або аномалій. Популярні методи однокласової класифікації включають використання еліпсоїда прогнозування та алгоритмів машинного навчання, такі як однокласова опорно-векторна машина, ізоляційний ліс та автокодувальник.

2) Існуючі методи часто базуються на припущенні, що дані мають багатовимірний нормальний розподіл, однак, це не завжди виконується в реальному світі, що в результаті, призводить до зниження точності та ймовірності розпізнавання образів для негаусівських даних. Одним з ефективних способів подолання обмежень, пов'язаних з розподілом даних, є використання нормалізуючих перетворень, які допомагають коригувати розподіл даних, наближаючи його до нормального. Нормалізуючі перетворення можуть бути одновимірними та багатовимірними. Одновимірні перетворення застосовуються до кожної ознаки окремо та фокусуються на індивідуальних характеристиках, наближаючи їх до нормального розподілу. Багатовимірні нормалізуючі перетворення, на відміну від одновимірних, враховують взаємозв'язки між ознаками. Вони розглядають дані як цілісну систему, де кожна ознака може впливати на інші.

3) Було досліджено сучасні методи виявлення та отримання вектору характеристик обличчя з зображення. На якісних зображеннях усі методи працюють досить точно, але при спотвореннях метод Віоли-Джонса показав найгірший результат. MTCNN і Dlib працюють майже однаково, але MTCNN краще розпізнає маленькі зображення та за слабкого освітлення, і має вищу точність, ніж дескриптор функції Dlib. Він також добре справляється з обличчями, що дивляться не тільки прямо, але і під непростими кутами. Однак, MTCNN потребує більше ресурсів, оскільки працює за допомогою нейронних мереж, і знаходить лише 5 точок, а саме очі, ніс і краї рота. OpenCV не має вбудованого механізму отримання координат ключових точок, тому найзручнішим в є Dlib, оскільки він має функціонал для отримання масиву координат із 68 точок, і працює швидше, ніж MTCNN.

4) З використанням бібліотеки Dlib було розроблено застосунок для пошуку та вилучення характеристик облич з зображення. Було сформовано 400 векторів ознак, по двісті для кожної з двох осіб.

5) Визначені метрики для всебічного оцінювання розроблених моделей, що включають в себе accuracy, precision, recall, specificity та F1 score. Кожна з цих

метрик виконує певну функцію, надаючи детальне уявлення про загальну точність моделі та її здатність розрізняти аномалії і цільові об'єкти в різних аспектах. Precision та Recall надають детальне уявлення про точність моделі при роботі з даними, що не належать до цільового образу, тоді як F1 Score служить підсумовуючою оцінкою, що враховує обидва аспекти. Specificity, в свою чергу, дає розуміння, як модель працює з цільовими образами. Accuracy дозволяє швидко отримати загальну точність, враховуючі всі об'єкти.

6) Проведено аналіз розподілів даних для розпізнавання облич та клавіатурного почерку. На основі критерію Мардіа, було визначено, що набори даних відхиляються від багатовимірного нормального розподілу, що підтверджує необхідність проведення дослідження за обраною темою.

Результати досліджень поточного розділу автором опубліковано в [47, 101-104].

РОЗДІЛ 2

ВИБІР НОРМАЛІЗУЮЧОГО ПЕРЕТВОРЕННЯ ДЛЯ ПОБУДОВИ НЕГАУСІВСЬКИХ ЙМОВІРНІСНИХ МОДЕЛЕЙ

У другому розділі досліджено існуючі нормалізуючі перетворення. Обрано перетворення для нормалізації десятивимірних векторів характеристик облич та дев'ятивимірних векторів характеристик клавіатурного почерку з метою подальшого застосування для побудови негаусівських ймовірнісних моделей.

Для цього було застосовано нормалізуючі перетворення до вибірки для розпізнавання облич особи 1 та до вибірки для розпізнавання клавіатурного почерку з ідентифікатором s015 для визначення нормалізації, яка найкраще підходить до наявних даних. Цей крок є необхідним, оскільки подальший аналіз, зокрема перевірка на наявність викидів за допомогою квадрату відстані Махаланобіса (1.1), базується на припущенні про нормальність розподілу даних та застосовується до повної нормалізованої вибірки. Виявлення та видалення викидів перед розподілом даних на тестову та тренувальну вибірки є критично важливим, оскільки це запобігає перенесенню спотворень у подальші етапи аналізу. Таким чином, нормалізація всього набору даних слугує основою для попередньої обробки, дозволяючи уникнути впливу викидів на якість моделей і забезпечити об'єктивність оцінки ефективності перетворень згідно тесту Мардіа.

Нормалізація є важливим кроком у попередній обробці даних, спрямованим на перетворення даних для наближення розподілу до нормального. Нормалізуючі перетворення поділяються на дві основні групи: одновимірні та багатовимірні перетворення.

Одновимірні зосереджені на незалежному перетворенні окремих змінних. Вони простіші в реалізації, тому ідеально підходять, коли характеристики вважаються незалежними або аналізуються ізольовано. До них зокрема належать перетворення десятичного логарифму та Бокса-Кокса.

Багатовимірні перетворення призначені для нормалізації наборів даних, що містять кілька взаємозалежних змінних, забезпечуючи збереження зв'язків між ними. На відміну від одновимірних перетворень, які розглядають кожну характеристику незалежно, багатовимірні перетворення працюють із набором даних у цілому, завдяки чому створюється узгоджене представлення набору даних, полегшуючи розуміння закономірностей і тенденцій, які можуть бути приховані під час незалежного аналізу змінних. Одним з недоліків є те, що такі перетворення складніші в реалізації.

Основна відмінність полягає в масштабі перетворення: одновимірні перетворення розглядають змінні як ізольовані сутності, тоді як багатовимірні трансформують набори даних цілісно, зберігаючи залежності між змінними. Таким чином, багатовимірні перетворення є складнішими за своєю суттю, але пропонують переваги у збереженні структурної цілісності набору даних.

2.1 Одновимірні взаємо-зворотні нормалізуючі перетворення

Одновимірні перетворення працюють з окремими змінними незалежно. Їх основною метою є зменшення асиметрії, стабілізація дисперсії та наближення змінної до гаусівського розподілу. Ці трансформації особливо ефективні для наборів даних, де змінні аналізуються окремо або вважаються такими, що мають незначну взаємозалежність.

2.1.1 Перетворення десяткового логарифму

Нормалізація даних за допомогою перетворення десяткового логарифму є однією з найпростіших та найпоширеніших технік. Вона особливо ефективна для даних із довгим правим хвостом, коли екстремальні значення непропорційно впливають на аналіз і виконується наступним чином [105, 106]:

$$Z = \lg(X), \quad (2.1)$$

де Z - вектор нормалізованих величин.

Зворотнє перетворення має вигляд:

$$X = 10^Z. \quad (2.2)$$

Однак перетворення десяткового логарифму вимагає строго позитивних вхідних значень. Якщо в наборі даних є нульові або негативні значення, можна додати константу, щоб змістити дані в позитивну область, тоді перетворення набуває виду:

$$Z = \lg(C + X), \quad (2.3)$$

де C – константа.

2.1.2 Одновимірне перетворення Бокса-Кокса

Це перетворення має лише один параметр лямбда та виконується за наступною формулою [107]:

$$x(\lambda) = \begin{cases} (X^\lambda - 1)/\lambda, & \lambda \neq 0; \\ \ln(X), & \lambda = 0, \end{cases} \quad (2.4)$$

де, X - випадкова величина, яка нормалізується; $x(\lambda)$ – нормалізована нормально розподілена величина; λ – параметр розподілу.

Перетворення (2.4) придатне для роботи лише для додатних значень випадкової величини X , тому було запропоновано розширений варіант:

$$x(\lambda) = \begin{cases} \frac{(X+\lambda_2)^{\lambda_1}-1}{\lambda_1}, & \lambda \neq 0; \\ \ln(X + \lambda_2), & \lambda = 0. \end{cases} \quad (2.5)$$

В (2.5) $\lambda = \{\lambda_1, \lambda_2\}^T$, де λ_2 обирається так, щоб $X + \lambda_2 > 0$ для всіх X .

У контексті оригінального перетворення Бокса-Кокса виконується одновимірне перетворення з одним параметром λ , застосованим поелементно до вектора. У випадку багатовимірних даних це перетворення зазвичай застосовується k разів як одновимірне перетворення до кожного стовпця, кожен зі своїм унікальним значенням k із k -змінного вектора $\theta = \{\lambda_1, \lambda_2, \dots, \lambda_k\}$. Визначення оптимального λ є критичним кроком, який має на меті наблизити розподіл перетвореного вихідної величини якомога ближче до нормального розподілу.

Найпопулярнішим методом знаходження оптимального значення параметру лямбда є метод максимальної правдоподібності (ММП) [108]:

$$l(\lambda) = C - \frac{N}{2} \ln \sum_{i=1}^N \frac{(x(\lambda)_i - \overline{x(\lambda)})^2}{n} + (\lambda - 1) \sum_{i=1}^N \ln(x_i), \quad (2.6)$$

де C – константа; $\overline{x(\lambda)} = \frac{1}{n} \sum_{i=1}^N x(\lambda)_i$.

2.2 Багатовимірне перетворення Бокса-Кокса.

Існує багатовимірне взаємо-зворотне нормалізуюче перетворення $\psi = \{\psi_1, \psi_2, \dots, \psi_N\}^T$, яке перетворює негаусівський випадковий вектор $\mathbf{X} = \{X_1, X_2, \dots, X_N\}^T$ у гаусівський випадковий вектор $\mathbf{Z} = \{Z_1, Z_2, \dots, Z_N\}^T$:

$$\mathbf{Z} = \psi(\mathbf{X}), \quad (2.7)$$

яке має зворотне перетворення

$$\mathbf{X} = \psi^{-1}(\mathbf{Z}). \quad (2.8)$$

У разі використання багатовимірного перетворення Бокса-Кокса компоненти вектору $\mathbf{Z}$ багатовимірного нормалізуючого перетворення визначаються як

$$Z_j = x(\lambda_j) = \begin{cases} (X_j^{\lambda_j} - 1)/\lambda_j, & \lambda_j \neq 0; \\ \ln(X_j), & \lambda_j = 0, \end{cases} \quad (2.9)$$

де, j змінюється від 1 до кількості змінних; X_j – випадкова величина, яка нормалізується; Z_j – нормалізована нормально розподілена величина; λ_j – параметр перетворення Бокса-Кокса.

Для багатовимірного перетворення Бокса-Кокса логарифмічну функцію правдоподібності можна записати як [76, 109]:

$$l(\mathbf{X}, \boldsymbol{\theta}) = \sum_{j=1}^k (\lambda_j - 1) \sum_{i=1}^N \ln(x_{ji}) - \frac{N}{2} \ln[\det(\mathbf{S}_Z)], \quad (2.10)$$

де k – кількість змінних; $\boldsymbol{\theta} = \{\lambda_1, \lambda_2, \dots, \lambda_k\}^T$; $\mathbf{S}_Z$ – це вибіркова коваріаційна матриця для компонентів вектору $\mathbf{Z}$:

$$\mathbf{S}_Z = \frac{1}{N} \sum_{i=1}^N (\mathbf{Z}_i - \bar{\mathbf{Z}})(\mathbf{Z}_i - \bar{\mathbf{Z}})^T, \quad (2.11)$$

яка описує розподіл даних у багатовимірному просторі. У контексті перетворення Бокса-Кокса враховується, що компоненти вектора $\mathbf{Z}$ можуть мати ненульові математичні сподівання. Для цього в рівнянні застосовується центрування за допомогою середнього значення $\bar{\mathbf{Z}}$.

2.3 Вибір перетворення для нормалізації десятивимірних характеристик облич

2.3.1 Нормалізація характеристик облич перетворенням десяткового логарифму

За допомогою перетворення (2.1) виконується нормалізація векторів характеристик особи 1. Після застосування перетворення вектор середніх нормалізованої вибірки становить $\bar{\mathbf{Z}} = \{-0,19917; 0,06745; -0,51762; -0,20160; -0,52258; -0,43920; -0,51702; -0,09604; -0,44314; -0,19760\}$. Коваріаційна матриця отриманої вибірки відображена в таблиці 2.1, діапазони характеристик – в таблиці 2.2.

Таблиця 2.1 – Коваріаційна матриця нормалізованої перетворенням десяткового логарифму вибірки особи 1

0,0 ³ 57	0,0 ³ 28	-0,0 ³ 75	-0,0 ⁴ 46	-0,0 ³ 13	0,0 ⁵ 21	-0,0 ³ 57	0,0 ⁴ 59	0,0 ³ 37	-0,0 ³ 56
0,0 ³ 28	0,0 ³ 32	-0,0 ⁴ 28	0,0 ⁴ 59	-0,0 ³ 1	0,0 ⁴ 39	0,0 ³ 26	0,0 ⁴ 27	0,0 ³ 28	0,0 ⁴ 86
-0,0 ³ 75	-0,0 ⁴ 28	0,0 ² 27	0,0 ³ 53	0,0 ⁴ 59	0,0 ³ 29	0,0 ² 17	-0,0 ³ 14	-0,0 ³ 31	0,0 ² 15
-0,0 ⁴ 46	0,0 ⁴ 59	0,0 ³ 53	0,0 ³ 2	-0,0 ⁴ 2	0,0 ³ 16	0,0 ³ 31	-0,0 ⁴ 69	-0,0 ⁴ 46	0,0 ³ 25
-0,0 ³ 13	-0,0 ⁴ 99	0,0 ⁴ 59	-0,0 ⁴ 2	0,0 ³ 15	0,0 ⁵ 53	0,0 ⁴ 4	0,0 ⁴ 15	-0,0 ³ 18	0,0 ⁴ 76
0,0 ⁵ 21	0,0 ⁴ 39	0,0 ³ 29	0,0 ³ 16	0,0 ⁵ 53	0,0 ³ 98	0,0 ³ 25	0,0 ⁴ 69	0,0 ⁴ 56	0,0 ⁴ 45
-0,0 ³ 57	0,0 ³ 26	0,0 ² 17	0,0 ³ 31	0,0 ⁴ 4	0,0 ³ 25	0,0 ² 25	-0,0 ³ 13	-0,0 ³ 15	0,0 ² 16
0,0 ⁴ 59	0,0 ⁴ 27	-0,0 ³ 14	-0,0 ⁴ 69	0,0 ⁴ 15	0,0 ⁴ 68	-0,0 ³ 13	0,0 ³ 89	0,0 ³ 67	0,0 ³ 44
0,0 ³ 37	0,0 ³ 28	-0,0 ³ 31	-0,0 ⁴ 46	-0,0 ³ 18	0,0 ⁴ 57	-0,0 ³ 15	0,0 ³ 67	0,0 ³ 14	0,0 ³ 13
-0,0 ³ 56	0,0 ⁴ 86	0,0 ² 15	0,0 ³ 25	0,0 ⁴ 76	0,0 ⁴ 45	0,0 ² 16	0,0 ³ 44	0,0 ³ 13	0,0 ² 22

Таблиця 2.2 – Діапазони характеристик нормалізованої перетворенням десяткового логарифму вибірки особи 1

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9	Z_{10}
Min	-0,2605	0,0295	-0,6404	-0,2382	-0,564	-0,5074	-0,67	-0,1557	-0,5327	-0,3297
Max	-0,1497	0,1195	-0,3734	-0,159	-0,4926	-0,365	-0,3984	0,0243	-0,3365	-0,0918

В результаті застосування перетворення (2.1) нормалізовані дані відхиляються від багатовимірного нормального розподілу, адже значення тестової статистики розподілу вибірки першої особи для багатовимірної асиметрії $B_1N/6$, яка дорівнює 381,91, перевищує значення квантиля розподілу Хі-квадрат, що становить 277,77 для 220 ступенів свободи та рівня значущості 0,005. Натомість, тестова статистика розподілу даної вибірки для багатовимірного ексцесу β_2 , яка становить 122,35, не перевищує значення квантиля нормального розподілу, яке становить 127,97 для математичного сподівання 120 і дисперсії 9,6 та рівня значущості 0,005.

2.3.2 Нормалізація характеристик облич одновимірним перетворенням Бокса-Кокса

Отримання оцінок параметрів відбувається k разів в залежності від кількості змінних, в даному випадку 10, що призводить до додаткового навантаження при застосуванні перетворення, при тому кожна змінна нормалізується не залежно від інших, тобто не враховується кореляція між ознаками.

В результаті вирішення задачі за методом максимальної правдоподібності логарифмічної функції (2.6) було отримано наступні оцінки параметрів: $\hat{\lambda}_1 = 2,4283$, $\hat{\lambda}_2 = -4,1501$, $\hat{\lambda}_3 = -0,4214$, $\hat{\lambda}_4 = -3,6769$, $\hat{\lambda}_5 = 6,6011$, $\hat{\lambda}_6 = -0,4094$, $\hat{\lambda}_7 = 0,3126$, $\hat{\lambda}_8 = -4,2148$, $\hat{\lambda}_9 = -1,3849$, $\hat{\lambda}_{10} = 1,0563$.

Унаслідок застосування одновимірного перетворення Бокса-Кокса з компонентами (2.4) вектор середніх нормалізованої вибірки особи 1 становить $\bar{\mathbf{Z}} = \{-0,2754; 0,1127; -1,5532; -1,2375; -0,1514; -1,2544; -0,9927; -0,3888; -2,2663; -0,3576\}$. Коваріаційна матриця нормалізованої вибірки наведена в таблиці 2.3, діапазони характеристик – в таблиці 2.4.

Таблиця 2.3 – Коваріаційна матриця нормалізованої одновимірним перетворенням Бокса-Кокса вибірки особи 1

0,0 ³ 32	0,0 ³ 25	-0,0 ² 21	-0,0 ³ 42	0,0 ⁷ 83	0,0 ⁵ 29	-0,0 ³ 66	0,0 ³ 25	0,0 ² 27	-0,0 ³ 57
0,0 ³ 25	0,0 ³ 45	-0,0 ³ 24	0,0 ³ 79	0,0 ⁷ 92	0,0 ³ 15	0,0 ³ 46	0,0 ³ 22	0,0 ² 3	0,0 ³ 11
-0,0 ² 21	-0,0 ³ 24	0,039	0,025	0,0 ⁶ 22	0,0 ² 37	0,01	-0,0 ² 34	-0,012	0,0 ² 8
-0,0 ³ 42	0,0 ³ 79	0,025	0,032	0,0 ⁶ 16	0,0 ² 72	0,0 ² 61	-0,0 ² 56	-0,0 ² 64	0,0 ² 46
0,0 ⁷ 83	0,0 ⁷ 92	0,0 ⁶ 22	0,0 ⁶ 16	0,0 ¹⁰ 95	0,0 ⁸ 47	0,0 ⁷ 56	0,0 ⁷ 38	0,0 ⁵ 14	0,0 ⁷ 87
0,0 ⁵ 29	0,0 ³ 15	0,0 ² 37	0,0 ² 72	0,0 ⁸ 47	0,012	0,0 ² 13	0,0 ³ 95	0,0 ² 14	0,0 ³ 28
-0,0 ³ 66	0,0 ³ 46	0,01	0,0 ² 61	0,0 ⁷ 56	0,0 ² 14	0,0 ² 62	-0,0 ² 11	-0,0 ² 23	0,0 ² 35
0,0 ³ 25	0,0 ³ 22	-0,0 ² 34	-0,0 ² 56	0,0 ⁷ 38	0,0 ³ 95	-0,0 ² 11	0,026	0,032	0,0 ² 34
0,0 ² 27	0,0 ² 3	-0,012	-0,0 ² 64	-0,0 ⁵ 14	0,0 ² 14	-0,0 ² 23	0,032	0,13	0,0 ² 2
-0,0 ³ 57	0,0 ³ 11	0,0 ² 8	0,0 ² 46	0,0 ⁷ 87	0,0 ³ 28	0,0 ² 35	0,0 ² 34	0,0 ² 2	0,0 ² 43

Таблиця 2.4 – Діапазони характеристик нормалізованої одновимірним перетворенням Бокса-Кокса вибірки особи 1

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9	Z_{10}
Min	-0,3159	0,0592	-2,0443	-1,7719	-0,1515	-1,4982	-1,224	-0,838	-3,225	-0,5221
Max	-0,2335	0,164	-1,0363	-0,7727	-0,1514	-1,003	-0,7976	0,0498	-1,3894	-0,1894

В результаті застосування одновимірного перетворення Бокса-Кокса з компонентами (2.4), де кожен елемент вектору $\mathbf{Z}$ обчислюється незалежно від інших, згідно тесту Мардіа нормалізовані дані відхиляються від багатовимірного нормального розподілу, тому що значення тестової статистики розподілу вибірки для багатовимірної асиметрії $B_1 N/6$, що становить 376,73, більше значення квантиля розподілу Хі-квадрат, яке дорівнює 277,77 для 220 ступенів свободи та

рівня значущості 0,005. Проте, тестова статистика розподілу даної вибірки для багатовимірного ексцесу β_2 , яка дорівнює 119,4, не перевищує значення квантиля нормального розподілу, що становить 127,97 для математичного сподівання 120 і дисперсії 9,6 та рівня значущості 0,005.

2.3.3 Нормалізація характеристик облич десятивимірним перетворенням Бокса-Кокса

В результаті вирішення задачі за методом максимальної правдоподібності логарифмічної функції (2.9) було отримано наступні оцінки параметрів: $\hat{\lambda}_1 = -0,37747$, $\hat{\lambda}_2 = 0,65858$, $\hat{\lambda}_3 = 0,42217$, $\hat{\lambda}_4 = 1,21823$, $\hat{\lambda}_5 = 5,90503$, $\hat{\lambda}_6 = -0,33303$, $\hat{\lambda}_7 = 0,42365$, $\hat{\lambda}_8 = -2,37054$, $\hat{\lambda}_9 = -0,91670$, $\hat{\lambda}_{10} = 1,20418$.

Після застосування десятивимірного перетворення Бокса-Кокса з компонентами (2.10) вектор середніх нормалізованої вибірки особи 1 становить $\bar{\mathbf{Z}} = \{-0,50137; 0,16415; -0,93473; -0,35418; -0,16921; -1,20364; -0,93330; -0,29971; -1,69772; -0,34637\}$. Коваріаційна матриця вибірки наведена в таблиці 2.5, діапазони характеристик – в таблиці 2.6.

Таблиця 2.5 – Коваріаційна матриця нормалізованої десятивимірним перетворенням Бокса-Кокса вибірки особи 1

0,0 ² 43	0,0 ² 2	-0,0 ² 29	-0,0 ³ 17	0,0 ⁶ 69	-0,0 ⁴ 16	-0,0 ² 22	0,0 ³ 64	0,0 ² 59	-0,0 ² 2
0,0 ² 2	0,0 ² 21	-0,0 ⁴ 89	0,0 ³ 2	0,0 ⁶ 47	0,0 ³ 32	0,0 ³ 93	0,0 ³ 32	0,0 ² 418	0,0 ³ 33
-0,0 ² 29	-0,0 ⁴ 89	0,0 ² 53	0,0 ³ 99	0,0 ⁶ 18	0,0 ² 13	0,0 ² 33	-0,0 ³ 81	-0,0 ² 26	0,0 ² 28
-0,0 ³ 17	0,0 ³ 2	0,0 ³ 99	0,0 ³ 35	0,0 ⁷ 43	0,0 ³ 69	0,0 ³ 56	-0,0 ³ 37	-0,0 ³ 39	0,0 ³ 45
0,0 ⁶ 69	0,0 ⁶ 47	0,0 ⁶ 18	0,0 ⁷ 43	0,0 ⁹ 51	0,0 ⁷ 12	0,0 ⁶ 11	0,0 ⁷ 69	0,0 ⁵ 2	0,0 ⁶ 18
-0,0 ⁴ 16	0,0 ³ 32	0,0 ² 13	0,0 ³ 69	0,0 ⁷ 12	0,01	0,0 ² 11	0,0 ³ 69	0,0 ³ 91	0,0 ³ 25
-0,0 ² 22	0,0 ³ 93	0,0 ² 33	0,0 ³ 56	0,0 ⁶ 11	0,0 ² 11	0,0 ² 48	-0,0 ³ 63	-0,0 ² 12	0,0 ² 29
0,0 ³ 64	0,0 ³ 32	-0,0 ³ 81	-0,0 ³ 37	0,0 ⁷ 69	0,0 ³ 69	-0,0 ³ 63	0,012	0,014	0,0 ² 22
0,0 ² 59	0,0 ² 42	-0,0 ² 26	-0,0 ³ 39	0,0 ⁵ 2	0,0 ³ 91	-0,0 ² 12	0,014	0,049	0,0 ² 12
-0,0 ² 2	0,0 ³ 33	0,0 ² 28	0,0 ³ 45	0,0 ⁶ 18	0,0 ³ 25	0,0 ² 29	0,0 ² 22	0,0 ² 12	0,0 ² 38

Таблиця 2.6 – Діапазони характеристик нормалізованої десятивимірним перетворенням Бокса-Кокса вибірки особи 1

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9	Z_{10}
Min	-0,6732	0,0695	-1,0977	-0,4001	-0,1693	-1,4283	-1,1326	-0,5651	-2,267	-0,4975
Max	-0,3681	0,3015	-0,7211	-0,2953	-0,1691	-0,9697	-0,7601	0,0524	-1,1285	-0,1866

В результаті застосування десятивимірного перетворення Бокса-Кокса з компонентами (2.10) нормалізовані дані все ще відхиляються від багатовимірного нормального розподілу, адже значення тестової статистики розподілу вибірки першої особи для багатовимірної асиметрії $B_1N/6$, яка дорівнює 332,73, перевищує значення квантиля розподілу Хі-квадрат, що становить 277,77 для 220 ступенів свободи та рівня значущості 0,005. Але, тестова статистика розподілу даної вибірки для багатовимірного ексцесу β_2 , яка становить 118,8, не перевищує значення квантиля нормального розподілу, що дорівнює 127,97 для математичного сподівання 120 і дисперсії 9,6 та рівня значущості 0,005.

Жодне з нормалізуючих перетворень не змогло задовільно нормалізувати вибірку для особи 1, тим не менш нормалізуючі перетворення допомогли значно наблизити розподіл даних до нормального. Найефективнішим, за результатами тесту Мардіа, виявилось десятивимірне перетворення Бокса-Кокса, тому воно буде використовуватися при подальшому аналізі.

В майбутньому вибірка буде поділена на тренувальну та тестову, саме до тренувальної вибірки буде застосовуватися обране перетворення при побудові негаусівських ймовірнісних моделей для розпізнавання облич. Важливо зазначити, що обране багатовимірне перетворення враховує кореляцію між змінними, що дозволяє краще моделювати складні залежності між ознаками.

Після застосування нормалізуючого перетворення вибірка була перевірена на наявність викидів за допомогою KBM. Аналіз підтвердив, що у наборі даних викидів немає.

2.4 Вибір перетворення для нормалізації та дев'ятивимірних характеристик клавіатурного почерку

2.4.1 Нормалізація характеристик клавіатурного почерку перетворенням десяткового логарифму

За допомогою перетворення (2.1) виконується нормалізація векторів характеристик з ідентифікатором s015. Після застосування перетворення вектор середніх нормалізованої вибірки з ідентифікатором s015 становить $\bar{\mathbf{Z}} = \{-1,13142; -1,16298; -1,11647; -1,21719; -1,16735; -1,05761; -1,07567; -1,13738; -1,13558\}$. Коваріаційна матриця вибірки наведена в таблиці 2.7, діапазони характеристик – в таблиці 2.8.

Таблиця 2.7 – Коваріаційна матриця нормалізованої перетворенням десяткового логарифму вибірки з ідентифікатором s015

0,0 ² 89	0,0 ² 14	0,0 ² 16	0,0 ³ 37	0,0 ³ 22	0,0 ³ 32	0,0 ³ 75	0,0 ³ 37	-0,0 ² 1
0,0 ² 14	0,0 ² 99	-0,0 ³ 60	-0,0 ³ 34	0,0 ³ 32	0,0 ³ 61	0,0 ² 13	-0,0 ³ 69	-0,0 ⁴ 27
0,0 ² 16	-0,0 ³ 60	0,0 ² 92	-0,0 ³ 1	0,0 ³ 17	-0,0 ⁴ 48	-0,0 ³ 68	0,0 ² 12	0,0 ⁴ 35
0,0 ³ 37	-0,0 ³ 34	-0,0 ³ 1	0,016	0,0 ³ 54	0,0 ⁴ 13	0,0 ³ 38	-0,0 ² 13	-0,0 ³ 5
0,0 ³ 22	0,0 ³ 32	0,0 ³ 17	0,0 ³ 54	0,0 ² 76	0,0 ³ 52	-0,0 ³ 36	0,0 ⁴ 5	0,0 ³ 38
0,0 ³ 32	0,0 ³ 61	-0,0 ⁴ 48	0,0 ⁴ 13	0,0 ³ 52	0,0 ² 36	-0,0 ³ 26	0,0 ⁴ 23	0,0 ⁴ 15
0,0 ³ 75	0,0 ² 13	-0,0 ³ 68	0,0 ³ 38	-0,0 ³ 36	-0,0 ³ 26	0,012	-0,0 ² 15	-0,0 ⁴ 51
0,0 ³ 37	-0,0 ³ 69	0,0 ² 12	-0,0 ² 13	0,0 ⁴ 5	0,0 ⁴ 23	-0,0 ² 15	0,013	0,0 ⁴ 31
-0,0 ² 1	-0,0 ⁴ 27	0,0 ⁴ 35	-0,0 ³ 5	0,0 ³ 38	0,0 ⁴ 15	-0,0 ⁴ 51	0,0 ⁴ 31	0,011

Таблиця 2.8 – Діапазони характеристик нормалізованої перетворенням десяткового логарифму вибірки з ідентифікатором s015

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9
Min	-1,7447	-1,8601	-1,5171	-2,2147	-1,7852	-1,4473	-1,6234	-1,9101	-2,1135
Max	-0,8165	-0,9626	-0,8962	-1,0066	-1,0173	-0,8357	-0,8854	-0,8342	-0,9551

В результаті застосування перетворення (2.1) нормалізовані дані відхиляються від багатовимірного нормального розподілу, адже значення тестової статистики розподілу вибірки з ідентифікатором s015 для багатовимірної асиметрії $B_1N/6$, яка дорівнює 2770,22, перевищує значення квантиля Хі-квадрат 215,53 для 165 ступенів свободи та рівня значущості 0,005. Подібним чином тестова статистика для багатовимірного ексцесу β_2 зі значенням 176,52 перевищує значення квантиля нормального розподілу 102,62 із середнім значенням 99, дисперсією 1,98 і рівнем значущості 0,005.

Як видно з результатів, перетворення десяткового логарифму збільшило відхилення даних від багатовимірного нормального розподілу. Це свідчить про те, що застосоване перетворення не дозволило досягти бажаного рівня нормалізації даних.

Такі результати вказують на необхідність розгляду альтернативних методів нормалізації, які можуть краще впоратися з даною вибіркою.

2.4.2 Нормалізація характеристик клавіатурного почерку одновимірним перетворенням Бокса-Кокса

Унаслідок вирішення задачі за методом максимальної правдоподібності логарифмічної функції (2.6) було отримано наступні оцінки параметрів: $\hat{\lambda}_1 = 0,9204$, $\hat{\lambda}_2 = 1,3617$, $\hat{\lambda}_3 = 1,2207$, $\hat{\lambda}_4 = 1,7079$, $\hat{\lambda}_5 = 2,2530$, $\hat{\lambda}_6 = 1,0741$, $\hat{\lambda}_7 = 1,5590$, $\hat{\lambda}_8 = 1,3660$, $\hat{\lambda}_9 = 2,0702$.

В результаті застосування одновимірного перетворення Бокса-Кокса з компонентами (2.4) вектор середніх нормалізованої вибірки з ідентифікатором s015 становить $\bar{\mathbf{Z}} = \{-0,9859; -0,7144; -0,7825; -0,5802; -0,4427; -0,8622; -0,6271; -0,7105; -0,4807\}$. Коваріаційна матриця вибірки зображена в таблиці 2.9, а діапазони характеристик – в таблиці 2.10.

Таблиця 2.9 – Коваріаційна матриця нормалізованої одновимірним перетворенням Бокса-Кокса вибірки з ідентифікатором s015

0,0 ³ 34	0, 0 ⁴ 15	0, 0 ⁴ 35	0, 0 ⁵ 23	0, 0 ⁶ 39	0, 0 ⁴ 13	0, 0 ⁵ 98	-0, 0 ⁵ 18	-0, 0 ⁵ 16
0,0 ⁴ 15	0,0 ⁴ 3	-0, 0 ⁵ 38	0, 0 ⁷ 58	0, 0 ⁶ 17	0, 0 ⁵ 71	0, 0 ⁵ 41	-0, 0 ⁵ 2	0, 0 ⁷ 55
0,0 ⁴ 35	-0,0 ⁵ 38	0,0 ⁴ 8	-0, 0 ⁶ 21	0, 0 ⁷ 38	0, 0 ⁶ 8	-0, 0 ⁵ 37	0, 0 ⁵ 41	0, 0 ⁶ 11
0,0 ⁵ 23	0,0 ⁷ 58	-0,0 ⁶ 21	0,0 ⁵ 36	0, 0 ⁷ 53	-0, 0 ⁶ 32	0, 0 ⁵ 1	-0, 0 ⁶ 66	-0, 0 ⁷ 35
0,0 ⁶ 39	0,0 ⁶ 17	0,0 ⁷ 38	0,0 ⁷ 53	0,0 ⁶ 15	0, 0 ⁶ 5	-0, 0 ⁸ 95	0, 0 ⁷ 76	0, 0 ⁸ 42
0,0 ⁴ 13	0,0 ⁵ 71	0,0 ⁶ 8	-0,0 ⁶ 32	0,0 ⁶ 5	0,0 ⁴ 95	-0, 0 ⁵ 24	0, 0 ⁵ 14	0, 0 ⁶ 13
0,0 ⁵ 98	0,0 ⁵ 41	-0,0 ⁵ 37	0,0 ⁵ 1	-0,0 ⁸ 95	-0,0 ⁵ 24	0,0 ⁴ 21	-0, 0 ⁵ 37	0, 0 ⁷ 79
-0,0 ⁵ 18	-0,0 ⁵ 2	0,0 ⁵ 41	-0,0 ⁶ 66	0,0 ⁷ 76	0,0 ⁵ 14	-0,0 ⁵ 37	0,0 ⁴ 41	0, 0 ⁶ 24
-0,0 ⁵ 16	0,0 ⁷ 55	0,0 ⁶ 11	-0,0 ⁷ 35	0,0 ⁸ 42	0,0 ⁶ 13	0,0 ⁷ 79	0,0 ⁶ 24	0,0 ⁶ 69

Таблиця 2.10 – Діапазони характеристик нормалізованої одновимірним перетворенням Бокса-Кокса вибірки з ідентифікатором s015

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9
Min	-1,0596	-0,7322	-0,8077	-0,5854	-0,4438	-0,905	-0,6396	-0,7303	-0,483
Max	-0,8939	-0,6985	-0,7532	-0,5743	-0,4416	-0,8131	-0,6147	-0,679	-0,478

В результаті застосування одновимірного перетворення Бокса-Кокса з компонентами (2.4), де кожен елемент вектору $\mathbf{Z}$ обчислюється незалежно від інших, нормалізовані дані відхиляються від багатовимірного нормального розподілу. Хоч значення тестової статистики розподілу для багатовимірної асиметрії $B_1N/6$, яка дорівнює 211,97, не перевищує значення квантиля Хі-квадрат, яке становить 215,53, для 165 ступенів свободи та рівня значущості 0,005. Проте тестова статистика для багатовимірного ексцесу β_2 зі значенням 109,04 перевищує значення квантиля нормального розподілу 102,62 із середнім значенням 99, дисперсією 1,98 і рівнем значущості 0,005.

2.4.3 Нормалізація характеристик клавіатурного почерку дев'ятивимірним перетворенням Бокса-Кокса

Після вирішення задачі за методом максимальної правдоподібності логарифмічної функції (2.9) було отримано наступні оцінки параметрів: $\hat{\lambda}_1 = 0,99387$, $\hat{\lambda}_2 = 1,36053$, $\hat{\lambda}_3 = 1,22020$, $\hat{\lambda}_4 = 1,75211$, $\hat{\lambda}_5 = 2,29656$, $\hat{\lambda}_6 = 1,04469$, $\hat{\lambda}_7 = 1,64663$, $\hat{\lambda}_8 = 1,35121$, $\hat{\lambda}_9 = 2,05989$.

В результаті застосування дев'ятивимірного перетворення Бокса-Кокса з компонентами (2.10) вектор середніх нормалізованої вибірки з ідентифікатором s015 становить $\bar{\mathbf{Z}} = \{-0,92901; -0,71493; -0,78274; -0,56612; -0,43445; -0,88128; -0,59626; -0,71740; -0,48303\}$. Коваріаційна матриця вибірки наведена в таблиці 2.11, діапазони характеристик – в таблиці 2.12.

Таблиця 2.11 – Коваріаційна матриця нормалізованої дев'ятивимірним перетворенням Бокса-Кокса вибірки з ідентифікатором s015

0,0 ³ 23	0,0 ⁴ 13	0,0 ⁴ 29	0,0 ⁵ 17	0,0 ⁶ 29	0,0 ⁴ 12	0,0 ⁵ 66	-0,0 ⁵ 17	-0,0 ⁵ 13
0,0 ⁴ 13	0,0 ⁴ 3	-0,0 ⁵ 38	0,0 ⁷ 57	0,0 ⁶ 15	0,0 ⁵ 62	0,0 ⁵ 34	-0,0 ⁵ 2	0,0 ⁷ 56
0,0 ⁴ 29	-0,0 ⁵ 38	0,0 ⁴ 8	-0,0 ⁶ 19	0,0 ⁷ 32	0,0 ⁶ 85	-0,0 ⁵ 3	0,0 ⁵ 42	0,0 ⁶ 1
0,0 ⁵ 17	0,0 ⁷ 57	-0,0 ⁶ 19	0,0 ⁵ 28	0,0 ⁷ 41	-0,0 ⁶ 3	0,0 ⁶ 75	-0,0 ⁶ 61	-0,0 ⁷ 33
0,0 ⁶ 29	0,0 ⁶ 15	0,0 ⁷ 32	0,0 ⁷ 41	0,0 ⁶ 12	0,0 ⁶ 48	-0,0 ⁸ 57	0,0 ⁷ 7	0,0 ⁸ 38
0,0 ⁴ 12	0,0 ⁵ 62	0,0 ⁶ 85	-0,0 ⁶ 3	0,0 ⁶ 48	0,0 ³ 1	-0,0 ⁵ 2	0,0 ⁵ 16	0,0 ⁶ 14
0,0 ⁵ 66	0,0 ⁵ 34	-0,0 ⁵ 3	0,0 ⁶ 75	-0,0 ⁸ 57	-0,0 ⁵ 2	0,0 ⁴ 14	-0,0 ⁵ 31	0,0 ⁷ 64
-0,0 ⁵ 17	-0,0 ⁵ 2	0,0 ⁵ 42	-0,0 ⁶ 61	0,0 ⁷ 7	0,0 ⁵ 16	-0,0 ⁵ 31	0,0 ⁴ 44	0,0 ⁶ 26
-0,0 ⁵ 13	0,0 ⁷ 56	0,0 ⁶ 1	-0,0 ⁷ 33	0,0 ⁸ 38	0,0 ⁶ 14	0,0 ⁷ 64	0,0 ⁶ 26	0,0 ⁶ 73

Таблиця 2.12 – Діапазони характеристик нормалізованої дев'ятивимірним перетворенням Бокса-Кокса вибірки з ідентифікатором s015

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9
Min	-0,9876	-0,7328	-0,808	-0,5707	-0,4354	-0,9278	-0,606	-0,7381	-0,4854
Max	-0,8509	-0,699	-0,7535	-0,5609	-0,4334	-0,829	-0,5861	-0,6849	-0,4802

Унаслідок застосування дев'ятивимірної перетворення Бокса-Кокса з компонентами (2.10) нормалізовані дані все ще відхиляються від багатовимірної нормального розподілу, тому що тестова статистика багатовимірної ексцесу β_2 зі значенням 109,01 перевищує значення квантиля нормального розподілу 102,62 із середнім значенням 99, дисперсією 1,98 і рівнем значущості 0,005. Хоча значення тестової статистики розподілу вибірки для багатовимірної асиметрії $B_1 N/6$, яка дорівнює 212,07, менше значення квантиля Хі-квадрат 215,53 для 165 ступенів свободи та рівня значущості 0,005.

Вибірка клавіатурного почерку все ще відхиляється від нормального закону розподілу після застосування нормалізуючих перетворень. Найменш ефективним виявилось перетворення десяткового логарифму, яке згідно тесту Мардіа збільшило відхилення від нормального розподілу. Нормалізовані вибірки за допомогою одновимірної та дев'ятивимірної перетворення Бокса-Кокса мають майже однакові показники тестових статистик асиметрії та ексцесу, тим не менш їх розподіл все ще відхиляється від нормального. Такі результати можуть свідчити про наявність аномальних точок в даних, тому наступним етапом буде пошук та видалення викидів в характеристиках клавіатурного почерку.

2.4.4 Пошук та видалення викидів у характеристиках клавіатурного почерку

Один із найпоширеніших методів виявлення та видалення викидів базується на квадраті відстані Махаланобіса (КВМ) (1.1). Цей метод обчислює відстань кожної точки даних від багатовимірної середньої, враховуючи кореляції між змінними. Основною передумовою цього методу є те, що набір даних має відповідати багатовимірному нормальному розподілу. Так як згідно тесту Мардіа багатовимірний розподіл характеристик клавіатурного почерку відхиляється від нормального, необхідно виконати нормалізацію характеристик.

Для нормалізації вибірки було обрано дев'ятивимірне перетворення Бокса-Кокса. В такому випадку KBM для нормалізованих даних для i -ої точки даних обраховується за наступною формулою:

$$D^2(\mathbf{Z}_i) = (\mathbf{Z}_i - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z}_i - \bar{\mathbf{Z}}), \quad (2.12)$$

де, $\bar{\mathbf{Z}}$ – вектор нормалізованих вибірових середніх, $\mathbf{S}_Z$ – вибіркова коваріаційна матриця, яка визначається як (2.11), $\mathbf{Z}_i$ – вектор даних в i -ій точці.

Так як навіть після нормалізації характеристики клавіатурного почерку демонструють відхилення від нормального розподілу, що, ймовірно, пов'язано з наявністю викидів, які, як правило, суттєво зміщують середнє значення та впливають на коваріаційну матрицю, що ускладнює використання статистичних методів. Незважаючи на ці виклики, нормалізація даних залишається важливим кроком у попередній обробці, оскільки вона зменшує асиметрію, стабілізує дисперсію та наближає розподіл до багатовимірної нормальності. Це, у свою чергу, значно покращує точність і надійність методів, зокрема тих, що базуються на відстані Махаланобіса.

Для кожного вектору ознак у наборі даних обчислюється квадрат відстані Махаланобіса. Значення KBM порівнюються з критичним значенням розподілу Хі-квадрат, яке залежить від кількості вимірів (ступенів свободи) та обраного рівня значущості. Для характеристик клавіатурного почерку 9 вимірів і рівня значущості 0,005 критичне значення становить 23,59. Будь-який вектор ознак із KBM, що перевищує це порогове значення, класифікується як викид.

Це ітеративний процес, що включає кілька послідовних етапів: нормалізацію вхідних даних, обчислення KBM для кожної точки нормалізованого набору даних, порівняння значень KBM з критичним значенням, ідентифікацію точки з найвищим значенням KBM та її видалення.

Під час першої ітерації вектор з індексом 295, який має найбільше значення KBM 37,44, визначається як викид і видаляється. Цей ітераційний процес обчислення SMD, виявлення найбільшого викиду та його видалення повторюється доти, доки в наборі даних не залишиться суттєвих викидів. Після

видалення шести викидів, згідно тесту Мардіа, тестова статистика багатовимірного ексцесу β_2 , яка дорівнює 102,47, менше ніж значення квантиля нормального розподілу 102,65.

Всього було виконано 11 ітерацій та видалено 10 викидів. В таблиці 2.13 відображено оптимальні значення параметрів λ дев'ятивимірного перетворення Бокса-Кокса для кожної ітерації. Таблиця 2.14 містить детальну інформацію про значення KBM та відповідні індекси для кожного виявленого і видаленого викиду.

Таблиця 2.13 – Оптимальні значення параметрів λ дев'ятивимірного перетворення Бокса-Кокса

№	λ_1	λ_2	λ_3	λ_4	λ_5	λ_6	λ_7	λ_8	λ_9
2	1,2249	1,3553	1,2321	1,7562	2,3126	1,0402	1,6400	1,3525	2,0647
3	1,2290	1,3566	1,2342	1,7569	2,3126	1,4365	1,6464	1,3502	2,0405
4	1,2286	1,3669	1,2394	1,7514	2,3087	1,4293	1,6426	1,5744	2,0337
5	1,2995	1,3627	1,2488	1,7582	2,3145	1,4377	1,6351	1,5792	2,0404
6	1,3005	1,3684	1,2480	1,7420	2,3098	1,0635	1,6409	1,5759	2,0477
7	1,3123	1,3611	1,2410	1,7085	2,3483	1,0791	1,6351	1,5772	2,0520
8	1,3113	1,3848	1,2335	1,7121	2,3429	1,0692	1,6314	1,5822	2,1392
9	1,1595	1,3913	1,2248	1,7099	2,3378	1,0781	1,6302	1,5547	2,1344
10	1,1558	1,3926	1,2235	1,7042	2,3448	0,9186	1,6256	1,5550	2,1494
11	1,1205	1,3860	1,2235	1,7012	2,3553	0,9879	1,6221	1,5511	2,1495

Таблиця 2.14 – Видалені викиди

№	SMD	Index	№	SMD	Index
1	37,44	295	6	26,963	323
2	36,962	160	7	26,868	45
3	30,742	306	8	25,776	263
4	28,833	388	9	24,515	294
5	28,662	214	10	23,972	204

В результаті ітеративного видалення викидів був отриманий набір даних з вектором вибірових середніх $\bar{X} = \{0,07525; 0,07022; 0,07823; 0,063; 0,06911; 0,08829; 0,08605; 0,07505; 0,0751\}$. Коваріаційна матриця вибірки в таблиці 2.15, діапазони характеристик – в таблиці 2.16.

Таблиця 2.15 – Коваріаційна матриця вибірки після видалення викидів

0,0 ³ 19	0,0 ⁴ 26	0,0 ⁴ 54	0,0 ⁴ 13	0,0 ⁵ 49	0,0 ⁴ 12	0,0 ⁴ 26	-0,0 ⁵ 6	-0,0 ⁴ 17
0,0 ⁴ 26	0,0 ³ 21	-0,0 ⁴ 19	0,0 ⁵ 51	0,0 ⁵ 82	0,0 ⁴ 14	0,0 ⁴ 4	-0,0 ⁴ 1	0,0 ⁵ 14
0,0 ⁴ 50	-0,0 ⁴ 19	0,0 ³ 25	0,0 ⁶ 1	0,0 ⁵ 35	0,0 ⁵ 11	-0,0 ⁴ 26	0,0 ⁴ 27	0,0 ⁶ 1
0,0 ⁴ 13	0,0 ⁵ 51	0,0 ⁶ 1	0,0 ³ 19	0,0 ⁴ 16	-0,0 ⁵ 52	0,0 ⁴ 31	-0,0 ⁴ 18	-0,0 ⁵ 57
0,0 ⁵ 49	0,0 ⁵ 82	0,0 ⁵ 35	0,0 ⁴ 16	0,0 ³ 13	0,0 ⁴ 14	-0,0 ⁵ 62	0,0 ⁵ 56	0,0 ⁵ 48
0,0 ⁴ 12	0,0 ⁴ 14	0,0 ⁵ 11	-0,0 ⁵ 52	0,0 ⁴ 14	0,0 ³ 12	-0,0 ⁵ 36	0,0 ⁵ 36	0,0 ⁵ 86
0,0 ⁴ 26	0,0 ⁴ 4	-0,0 ⁴ 26	0,0 ⁴ 31	-0,0 ⁵ 62	-0,0 ⁵ 36	0,0 ³ 35	-0,0 ⁴ 38	0,0 ⁵ 19
-0,0 ⁵ 6	-0,0 ⁴ 1	0,0 ⁴ 27	-0,0 ⁴ 18	0,0 ⁵ 56	0,0 ⁵ 36	-0,0 ⁴ 38	0,0 ³ 27	0,0 ⁵ 7
-0,0 ⁴ 17	0,0 ⁵ 14	0,0 ⁶ 1	-0,0 ⁵ 57	0,0 ⁵ 48	0,0 ⁵ 86	0,0 ⁵ 19	0,0 ⁵ 7	0,0 ³ 2

Таблиця 2.16 – Діапазони характеристик вибірки після видалення викидів

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
Min	0,018	0,0138	0,0304	0,0061	0,0164	0,052	0,0238	0,0123	0,0077
Max	0,117	0,109	0,127	0,0985	0,0961	0,1235	0,1302	0,1246	0,1103

В результаті застосування дев'ятивимірного перетворення Бокса-Кокса з компонентами (2.10) до отриманої вибірки після видалення викидів вектор середніх нормалізованої вибірки з ідентифікатором s015 становить $\bar{Z} = \{-0,84318; -0,70311; -0,78096; -0,58235; -0,42376; -0,92021; -0,60467; -0,63286; -0,46337\}$.

Коваріаційна матриця вибірки наведена в таблиці 2.17, діапазони характеристик – в таблиці 2.18.

Таблиця 2.17 – Коваріаційна матриця нормалізованої дев'ятивимірним перетворенням Бокса-Кокса після видалення викидів вибірки s015

0,0 ³ 1	0,0 ⁵ 64	0,0 ⁴ 23	0,0 ⁵ 13	0,0 ⁶ 11	0,0 ⁵ 92	0,0 ⁵ 44	-0,0 ⁶ 97	-0,0 ⁶ 59
0,0 ⁵ 64	0,0 ⁴ 26	-0,0 ⁵ 38	0,0 ⁶ 33	0,0 ⁷ 99	0,0 ⁵ 52	0,0 ⁵ 32	-0,0 ⁶ 74	0,0 ⁷ 26
0,0 ⁴ 23	-0,0 ⁵ 38	0,0 ⁴ 8	-0,0 ⁷ 7	0,0 ⁷ 14	0,0 ⁶ 87	-0,0 ⁵ 34	0,0 ⁵ 31	-0,0 ⁷ 12
0,0 ⁵ 13	0,0 ⁶ 33	-0,0 ⁷ 7	0,0 ⁵ 36	0,0 ⁷ 6	-0,0 ⁶ 8	0,0 ⁵ 11	-0,0 ⁶ 47	-0,0 ⁷ 16
0,0 ⁶ 11	0,0 ⁷ 99	0,0 ⁷ 14	0,0 ⁷ 6	0,0 ⁷ 82	0,0 ⁶ 42	-0,0 ⁷ 12	0,0 ⁷ 4	0,0 ⁸ 31
0,0 ⁵ 92	0,0 ⁵ 52	0,0 ⁶ 87	-0,0 ⁶ 8	0,0 ⁶ 42	0,0 ³ 12	-0,0 ⁶ 84	0,0 ⁵ 11	0,0 ⁶ 41
0,0 ⁵ 44	0,0 ⁵ 32	-0,0 ⁵ 34	0,0 ⁵ 11	-0,0 ⁷ 12	-0,0 ⁶ 84	0,416	-0,0 ⁵ 19	0,0 ⁷ 43
-0,0 ⁶ 97	-0,0 ⁶ 74	0,0 ⁵ 31	-0,0 ⁶ 47	0,0 ⁷ 4	0,0 ⁵ 11	-0,0 ⁵ 19	0,0 ⁴ 14	0,0 ⁶ 11
-0,0 ⁶ 59	0,0 ⁷ 26	-0,0 ⁷ 12	-0,0 ⁷ 16	0,0 ⁸ 31	0,0 ⁶ 41	0,0 ⁷ 43	0,0 ⁶ 11	0,0 ⁶ 43

Таблиця 2.18 – Діапазони характеристик нормалізованої дев'ятивимірним перетворенням Бокса-Кокса після видалення викидів вибірки s015

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9
Min	-0,8826	-0,7196	-0,8060	-0,5877	-0,4246	-0,9577	-0,6151	-0,6440	-0,4652
Max	-0,8118	-0,6881	-0,7519	-0,5764	-0,4229	-0,8840	-0,5939	-0,6192	-0,4612

Після видалення викидів, розподіл нормалізованої вибірки за допомогою дев'ятивимірного перетворення Бокса-Кокса згідно тесту Мардіа наближено є гаусівським, оскільки значення тестової статистики багатовимірної асиметрії $B_1N/6$, яка дорівнює 167,59, не перевищує значення квантиля Хі-квадрат 215,53 для 165 ступенів свободи та рівня значущості 0,005, а тестова статистика багатовимірного ексцесу β_2 зі значенням 100,78 менше значення квантиля нормального розподілу 102,67 із середнім значенням 99, дисперсією 2,03 і рівнем значущості 0,005.

2.5 Висновки за розділом 2

У даному розділі було проаналізовано нормалізуючі перетворення та методи для оцінювання їх параметрів. Розглянуто перетворення десяткового логарифму, одновимірне та багатовимірне перетворення Бокса-Кокса. Здійснено нормалізацію наборів для розпізнавання облич особи 1 та розпізнавання клавіатурного почерку з ідентифікатором s015 за допомогою всіх розглянутих перетворень, проведено ефективність оцінювання за допомогою тесту Мардіа.

1) Для нормалізації десятидимірних векторів характеристик облич було обрано багатовимірне перетворення Бокса-Кокса, за допомогою якого вдалося максимально наблизити розподіл даної вибірки до нормального. Ще одною перевагою даного перетворення є здатність враховувати кореляцію між ознаками.

2) Для дев'ятидимірних векторів характеристик клавіатурного почерку найменш ефективним виявилось перетворення десяткового логарифму, яке збільшило відхилення від нормального розподілу. Нормалізовані вибірки за допомогою одновимірного та дев'ятидимірного перетворення Бокса-Кокса мають майже однакові показники тестових статистик асиметрії та ексцесу, тим не менш було обрано багатовимірне перетворення Бокса-Кокса завдяки своїй здатності враховувати взаємозв'язки між змінними.

3) Оцінювання параметрів перетворення Бокса-Кокса здійснювалося за допомогою методу максимальної правдоподібності. Використання МПП у порівнянні з іншими методами, такими як метод найменших квадратів, дозволяє отримувати ефективні оцінки, в тому числі для даних з розподілом, що відхиляється від нормального.

4) Після застосування нормалізуючих перетворень виконано перевірку вибірок на наявність викидів. У дев'ятидимірній вибірці для розпізнавання клавіатурного почерку з ідентифікатором s015, було ідентифіковано та видалено 10 викидів. Натомість у наборі даних для розпізнавання облич викидів виявлено не було.

Результати досліджень поточного розділу автором опубліковано в [47-49, 103, 110].

РОЗДІЛ 3

ПОБУДОВА НЕГАУСІВСЬКИХ ЙМОВІРНІСНИХ МОДЕЛЕЙ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ

У третьому розділі побудовано негаусівські ймовірнісні моделі у вигляді дев'ятивимірного еліпсоїда прогнозування для нормалізованих векторів характеристик клавіатурного почерку, та десятивимірного еліпсоїда прогнозування для нормалізованих векторів характеристик облич, з використанням багатовимірного перетворення Бокса-Кокса і квантилів Хі-квадрат та F -розподілу.

3.1 Побудова десятивимірних еліпсоїдів прогнозування для розпізнавання облич з різними квантилями розподілів

Для побудови правил прийняття рішень у вигляді десятивимірних еліпсоїдів прогнозування для розпізнавання облич треба мати відповідний набір даних, що містить в собі вибірку для побудови моделі та тестування. Отриманий набір даних розпізнавання облич містить по 200 векторів характеристик для кожної з 2 осіб та був розділений наступним чином: 100 – навчальна вибірка для побудови десятивимірних еліпсоїдів прогнозування, ще 100 – тестова вибірка цільового класу. Інші 200 – тестова вибірка особи 2, яка буде використовуватися для тестування та оцінювання точності розпізнавання іншої особи. Характеристики тестової вибірки особи 2 наведені в таблицях 1.4 та 1.5.

Вектор вибірових середніх значень навчальної вибірки для розпізнавання облич дорівнює $\bar{X} = \{0,633; 1,16289; 0,30095; 0,62589; 0,29937; 0,36696; 0,3015; 0,80108; 0,35845; 0,62901\}$. Коваріаційна матриця навчальної вибірки відображена в таблиці 3.1, а діапазони характеристик – в таблиці 3.2.

Таблиця 3.1 – Коваріаційна матриця навчальної вибірки для розпізнавання облич

0,0 ² 1	0,0 ³ 89	-0,0 ³ 6	-0,0 ⁴ 63	-0,0 ³ 12	-0,0 ³ 11	-0,0 ³ 52	0,0 ⁴ 1	0,0 ³ 34	-0,0 ² 11
0,0 ³ 89	0,0 ² 2	0,0 ⁵ 39	0,0 ³ 26	-0,0 ³ 16	0,0 ⁴ 63	0,0 ³ 52	-0,0 ⁴ 36	0,0 ³ 46	0,0 ³ 24
-0,0 ³ 6	0,0 ⁵ 39	0,0 ² 11	0,0 ³ 47	0,0 ⁵ 77	0,0 ³ 19	0,0 ³ 71	-0,0 ⁴ 92	-0,0 ⁴ 94	0,0 ² 12
-0,0 ⁴ 63	0,0 ³ 26	0,0 ³ 47	0,0 ³ 39	-0,0 ⁴ 28	0,0 ³ 17	0,0 ³ 31	-0,0 ³ 27	-0,0 ⁴ 58	0,0 ³ 34
-0,0 ³ 12	-0,0 ³ 16	0,0 ⁵ 77	-0,0 ⁴ 28	0,0 ⁴ 8	0,0 ⁴ 22	0,0 ⁴ 27	-0,0 ⁵ 46	-0,0 ⁴ 99	0,0 ⁴ 76
-0,0 ³ 11	0,0 ⁴ 63	0,0 ³ 19	0,0 ³ 17	0,0 ⁴ 22	0,0 ³ 73	0,0 ³ 27	0,0 ⁴ 29	-0,0 ⁴ 2	0,0 ⁴ 17
-0,0 ³ 52	0,0 ³ 52	0,0 ³ 71	0,0 ³ 31	0,0 ⁴ 27	0,0 ³ 27	0,0 ² 12	-0,0 ³ 15	-0,0 ⁴ 73	0,0 ² 15
0,0 ⁴ 1	-0,0 ⁴ 36	-0,0 ⁴ 92	-0,0 ³ 27	-0,0 ⁵ 46	0,0 ⁴ 29	-0,0 ³ 15	0,0 ² 28	0,0 ² 1	0,0 ² 13
0,0 ³ 34	0,0 ³ 46	-0,0 ⁴ 94	-0,0 ⁴ 58	-0,0 ⁴ 99	-0,0 ⁴ 2	-0,0 ⁴ 73	0,0 ² 1	0,0 ² 1	0,0 ³ 16
-0,0 ² 11	0,0 ³ 24	0,0 ² 12	0,0 ³ 34	0,0 ⁴ 76	0,0 ⁴ 17	0,0 ² 15	0,0 ² 13	0,0 ³ 16	0,0 ² 42

Таблиця 3.2 – Діапазони характеристик навчальної вибірки для розпізнавання облич

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Min	0,5573	1,0703	0,2400	0,5778	0,2729	0,3109	0,2261	0,6987	0,2933	0,4681
Max	0,7085	1,2873	0,3958	0,6698	0,3217	0,4316	0,3996	1,0007	0,4604	0,8095

Вектор вибірових середніх значень тестової вибірки особи 1 для розпізнавання облич становить $\bar{X} = \{0,63322; 1,17515; 0,31072; 0,63206; 0,30129; 0,36242; 0,31062; 0,80601; 0,36526; 0,64708\}$. Діапазони характеристик тестової вибірки особи 1 зображені в таблиці 3.3, коваріаційна матриця тестової вибірки особи 1 відображена в таблиці 3.4.

На основі результатів тесту Мардіа, досліджено, що багатовимірний розподіл навчальної вибірки відхиляється від нормального. Про це свідчить тестова статистика для багатовимірної асиметрії $N\beta_1/6$, яка дорівнює 289,20, що перевищує значення квантиля Хі-квадрат, яке становить 277,77, для 220 ступенів свободи і рівня значущості 0,005. Натомість, тестова статистика для багатовимірного ексцесу β_2 зі значенням 122,35 не перевищує квантиль нормального розподілу, який становить 127,97, для середнього значення 120, дисперсії 9,6 і рівня значущості 0,005. Отже, необхідно застосувати

десятивимірне перетворення Бокса-Кокса до негаусового випадкового вектора $\mathbf{X} = \{X_1, X_2, \dots, X_{10}\}^T$, щоб перетворити його на гаусівський випадковий вектор $\mathbf{Z} = \{Z_1, Z_2, \dots, Z_{10}\}^T$.

Таблиця 3.3 – Діапазони характеристик тестової вибірки особи 1

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9	X_{10}
Min	0,5489	1,0780	0,2289	0,5919	0,2843	0,3131	0,2138	0,6994	0,2940	0,4741
Max	0,6972	1,3166	0,4232	0,6935	0,3171	0,4265	0,3792	1,0575	0,4608	0,7998

Таблиця 3.4 – Коваріаційна матриця тестової вибірки особи 1

0,0 ² 14	0,0 ² 13	-0,0 ³ 92	-0,0 ³ 13	-0,0 ³ 15	0,0 ³ 11	-0,0 ³ 64	0,0 ³ 29	0,0 ³ 57	-0,0 ² 12
0,0 ² 13	0,0 ² 27	-0,0 ³ 15	0,0 ³ 18	-0,0 ³ 22	0,0 ³ 14	0,0 ³ 42	0,0 ³ 27	0,0 ³ 79	0,0 ³ 43
-0,0 ³ 92	-0,0 ³ 15	0,0 ² 16	0,0 ³ 62	0,0 ⁴ 43	0,0 ³ 21	0,0 ³ 95	-0,0 ³ 32	-0,0 ³ 28	0,0 ² 18
-0,0 ³ 13	0,0 ³ 18	0,0 ³ 62	0,0 ³ 43	-0,0 ⁴ 19	0,0 ³ 25	0,0 ³ 28	-0,0 ³ 12	-0,0 ⁴ 61	0,0 ³ 69
-0,0 ³ 15	-0,0 ³ 22	0,0 ⁴ 43	-0,0 ⁴ 19	0,0 ⁴ 59	-0,0 ⁴ 14	0,0 ⁵ 36	0,0 ⁴ 41	-0,0 ³ 11	0,0 ⁴ 48
0,0 ³ 11	0,0 ³ 14	0,0 ³ 21	0,0 ³ 25	-0,0 ⁴ 14	0,0 ³ 75	0,0 ⁴ 61	0,0 ³ 21	0,0 ³ 13	0,0 ³ 15
-0,0 ³ 64	0,0 ³ 42	0,0 ³ 95	0,0 ³ 28	0,0 ⁵ 36	0,0 ⁴ 61	0,0 ² 12	-0,0 ³ 16	-0,0 ⁴ 97	0,0 ² 16
0,0 ³ 29	0,0 ³ 27	-0,0 ³ 32	-0,0 ³ 12	0,0 ⁴ 41	0,0 ³ 21	-0,0 ³ 16	0,0 ² 38	0,0 ² 12	0,0 ² 12
0,0 ³ 57	0,0 ³ 79	-0,0 ³ 28	-0,0 ⁴ 61	-0,0 ³ 11	0,0 ³ 13	-0,0 ⁴ 97	0,0 ² 12	0,0 ² 1	0,0 ³ 19
-0,0 ² 12	0,0 ³ 43	0,0 ² 18	0,0 ³ 69	0,0 ⁴ 48	0,0 ³ 15	0,0 ² 16	0,0 ² 12	0,0 ³ 19	0,0 ² 47

Наступним кроком виконується нормалізація навчальної вибірки десятивимірним перетворенням Бокса-Кокса. Оптимальні оцінки параметрів перетворення знаходяться за методом максимальної правдоподібності логарифмічної функції (2.9): $\hat{\lambda}_1 = -0,57999$, $\hat{\lambda}_2 = -0,97321$, $\hat{\lambda}_3 = 0,48717$, $\hat{\lambda}_4 = 2,88804$, $\hat{\lambda}_5 = 4,07142$, $\hat{\lambda}_6 = -0,23103$, $\hat{\lambda}_7 = 0,08710$, $\hat{\lambda}_8 = -1,29733$, $\hat{\lambda}_9 = -1,18668$, $\hat{\lambda}_{10} = 1,33216$.

Після застосування десятивимірного перетворення Бокса-Кокса з компонентами (2.10) вектор середніх нормалізованої вибірки особи 1 становить $\bar{\mathbf{Z}} = \{-0,52627; 0,13911; -0,91083; -0,25654; -0,24379; -1,13224; -1,14383; -0,26329;$

$-2,03304; -0,34494\}$. Коваріаційна матриця вибірки наведена в таблиці 3.5, діапазони характеристик – в таблиці 3.6.

Таблиця 3.5 – Коваріаційна матриця нормалізованої навчальної вибірки

0,0 ³ 44	0,0 ² 13	-0,0 ² 23	-0,0 ⁴ 55	-0,0 ⁵ 63	-0,0 ³ 8	-0,0 ² 32	0,0 ⁴ 7	0,0 ² 61	-0,0 ² 2
0,0 ² 13	0,0 ² 1	-0,0 ⁴ 19	0,0 ⁴ 77	-0,0 ⁵ 28	0,0 ³ 17	0,0 ² 11	-0,0 ⁴ 41	0,0 ² 28	0,0 ³ 13
-0,0 ² 23	-0,0 ⁴ 19	0,0 ² 37	0,0 ³ 36	0,0 ⁶ 51	0,0 ² 12	0,0 ² 38	-0,0 ³ 35	-0,0 ² 2	0,0 ² 19
-0,0 ⁴ 55	0,0 ⁴ 77	0,0 ³ 36	0,0 ⁴ 67	-0,0 ⁶ 27	0,0 ³ 23	0,0 ³ 38	-0,0 ³ 18	-0,0 ³ 27	0,0 ³ 12
-0,0 ⁵ 63	-0,0 ⁵ 28	0,0 ⁶ 51	-0,0 ⁶ 27	0,0 ⁷ 48	0,0 ⁵ 17	0,0 ⁵ 21	-0,0 ⁶ 37	-0,0 ⁴ 24	0,0 ⁵ 17
-0,0 ³ 8	0,0 ³ 17	0,0 ² 12	0,0 ³ 23	0,0 ⁵ 17	0,0 ² 85	0,0 ² 27	0,0 ³ 12	-0,0 ² 12	0,0 ⁴ 89
-0,0 ² 32	0,0 ² 11	0,0 ² 38	0,0 ³ 38	0,0 ⁵ 21	0,0 ² 27	0,01	-0,0 ³ 9	-0,0 ² 28	0,0 ² 38
0,0 ⁴ 7	-0,0 ⁴ 41	-0,0 ³ 35	-0,0 ³ 18	-0,0 ⁶ 37	0,0 ³ 12	-0,0 ³ 9	0,0 ² 71	0,014	0,0 ² 17
0,0 ² 61	0,0 ² 28	-0,0 ² 2	-0,0 ³ 27	-0,0 ⁴ 24	-0,0 ² 12	-0,0 ² 28	0,014	0,086	0,0 ² 12
-0,0 ² 2	0,0 ³ 13	0,0 ² 19	0,0 ³ 12	0,0 ⁵ 17	0,0 ⁴ 89	0,0 ² 38	0,0 ² 17	0,0 ² 12	0,0 ² 31

Таблиця 3.6 – Діапазони характеристик нормалізованої навчальної вибірки

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9	Z_{10}
Min	-0,6959	0,0658	-1,0285	-0,2752	-0,2444	-1,3412	-1,3945	-0,4565	-2,7695	-0,4776
Max	-0,3815	0,2239	-0,7458	-0,2374	-0,2432	-0,9275	-0,8817	0,0007	-1,2728	-0,1842

Нормалізована вибірка не відхиляється від багатовимірного нормального розподілу. Це підтверджується тестом Мардіа: тестова статистика для багатовимірної асиметрії $N\beta_1/6$, яка дорівнює 265,59, не перевищує значення квантиля χ^2 -квадрат, яке становить 277,77, для 220 ступенів свободи і рівня значущості 0,005. Тестова статистика для багатовимірного ексцесу β_2 зі значенням 121,52 не перевищує квантиль нормального розподілу, який становить 127,97, для середнього значення 120, дисперсії 9,6 і рівня значущості 0,005. Це дозволяє перейти до побудови негаусівської ймовірнісної моделі на основі еліпсоїда прогнозування для нормалізованих даних.

На основі (1.2) побудовано ймовірнісну модель у вигляді десятивимірної еліпсоїда прогнозування для нормалізованих характеристик облич з квантилем розподілу Хі-квадрат:

$$(\mathbf{Z} - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z} - \bar{\mathbf{Z}}) = \chi_{10, 0,005}^2. \quad (3.1)$$

Використовуючи (3.1), правило прийняття рішень для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем Хі-квадрат для нормалізованого вектору характеристик набуває вигляду:

$$(\mathbf{Z}_i - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z}_i - \bar{\mathbf{Z}}) \leq \chi_{10, 0,005}^2, \quad (3.2)$$

де, $\mathbf{Z}_i$ – нормалізований вектор даних в i -ій точці, $\mathbf{S}_Z^{-1}$ – обернена коваріаційна матриця нормалізованої навчальної вибірки особи 1, що зображена в таблиці 3.7. Значення квантиля розподілу Хі-квадрат для 10 ступенів свободи та рівня значущості 0,005 становить 25,19.

Таблиця 3.7 – Обернена коваріаційна матриця нормалізованої навчальної вибірки

6268,86	-9276,87	1954,51	-9819,11	37083,29	146,09	1939,33	-367,92	-18,65	1511,72
-9276,87	15781,23	-2169,39	9868,96	18524,66	-96,93	-3350,47	427,63	-16,33	-1957,80
1954,51	-2169,39	1708,65	-8458,20	23170,00	102,91	372,36	-160,66	-14,66	271,38
-9819,11	9868,96	-8458,20	67339,70	53240,44	-1253,80	-1793,76	1960,44	65,97	-3199,75
37083,29	18524,66	23170,00	53240,44	31160969,71	-4638,47	-5260,58	-12071,28	7707,78	788,73
146,09	-96,93	102,91	-1253,80	-4638,47	179,96	-78,15	-112,80	4,81	243,04
1939,33	-3350,47	372,36	-1793,76	-5260,58	-78,15	1025,00	100,07	-11,73	-49,94
-367,92	427,63	-160,66	1960,44	-12071,28	-112,80	100,07	457,04	-52,80	-595,29
-18,65	-16,33	-14,66	65,97	7707,78	4,81	-11,73	-52,80	23,44	26,38
1511,72	-1957,80	271,38	-3199,75	788,73	243,04	-49,94	-595,29	26,38	1758,48

На основі (1.3) побудовано ймовірнісну модель у вигляді десятивимірної еліпсоїда прогнозування для нормалізованих характеристик облич з квантилем F -розподілу:

$$(\mathbf{Z} - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z} - \bar{\mathbf{Z}}) = \frac{k(N^2-1)}{N(N-k)} F_{10, 90, 0,005}. \quad (3.3)$$

Використовуючи (3.3), побудовано вирішальне правило для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик:

$$(\mathbf{Z}_i - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z}_i - \bar{\mathbf{Z}}) \leq \frac{k(N^2-1)}{N(N-k)} F_{10, 90, 0,005}. \quad (3.4)$$

Значення квантиля F -розподілу для рівня значущості 0,005, зі ступенями свободи 10 і 90 дорівнює 2,77. Вся права частина рівняння дорівнює 30,77, що перевищує критичне значення еліпсоїда прогнозування з квантилем Хі-квадрат 25,19. Таблиця 3.8 містить порівняння за визначеними метриками побудованих еліпсоїдів прогнозування для негаусівських даних (PENGD) та еліпсоїдів прогнозування для нормалізованих даних (PEND) (3.2, 3.4) з різними квантилями розподілів в задачі розпізнавання облич.

Таблиця 3.8 – Порівняння застосування побудованих моделей

Quantiles	Models	Specificity	Recall	Precision	F1 score	Accuracy
Chi-square	PENGD	0,9200	0,9550	0,9598	0,9574	0,9433
	PEND	0,9700	0,9750	0,9848	0,9793	0,9733
F-distribution	PENGD	0,99	0,9350	0,9946	0,9639	0,9533
	PEND	1	0,9350	1	0,9664	0,9567

Не залежно від квантиля, нормалізація за допомогою десятивимірного перетворення Бокса-Кокса призводить до покращення точності розпізнавання моделі в порівнянні з моделями еліпсоїдів прогнозування для не нормалізованих даних. Точність еліпсоїда прогнозування для нормалізованих даних з квантилем розподілу Хі-квадрат перевершує аналогічну модель з квантилем F -розподілу, пропонуючи більш збалансоване рішення для розпізнавання цільового образу і аномалій. Використання квантиля F -розподілу може бути доцільним в задачах, де опізнати цільовий образ важливіше ніж відсіяти аномалії.

Еліпсоїд прогнозування, побудований на основі квантиля Хі-квадрат, є популярним інструментом для вирішення задач у галузі розпізнавання образів та

аналізу даних [31, 33, 71]. Хоча автори зазначають припущення про необхідність багатовимірного нормального розподілу даних, у науковій літературі недостатньо конкретних рішень для негаусівських даних. Одним із варіантів вирішення проблеми є використання еліпсоїда прогнозування для нормалізованих даних [45]. У даній роботі було удосконалено еліпсоїд прогнозування для нормалізованих даних з квантилем χ^2 -квадрат за рахунок більшої кількості характеристик у порівнянні з попередніми дослідженнями. Це дозволило покращити точність (Accuracy) розпізнавання облич з 0,9433 до 0,9733, а також ймовірність розпізнавання цільового класу (Specificity) з 0,9200 до 0,9700, та ймовірність розпізнавання представників інших класів (Recall) з 0,9550 до 0,9750.

У літературі не було виявлено застосувань еліпсоїда прогнозування з квантилем F -розподілу для даних, що відхиляються від багатовимірного нормального розподілу. У цій роботі вперше побудована ймовірнісна модель у вигляді еліпсоїда прогнозування для нормалізованих даних з квантилем F -розподілу, що враховує відхилення від нормального розподілу та кількість точок даних. Експериментальні дослідження показали, що її застосування дозволило підвищити точність розпізнавання (Accuracy) облич з 0,9533 до 0,9567, а також ймовірність розпізнавання цільового класу (Specificity) з 0,99 до 1, при тому ймовірність розпізнавання представників інших класів (Recall) залишилася 0,9350.

3.2 Побудова дев'ятивимірних еліпсоїдів прогнозування для розпізнавання клавіатурного почерку з різними квантилями розподілів

У початковому наборі даних було видалено викиди, після чого вибірку було перетасовано випадковим чином, щоб забезпечити рівномірний розподіл точок між навчальними та тестовими даними, зменшуючи будь-яке потенційне зміщення через порядок даних. Набір даних було розділено на навчальну, яка

використовується для побудови правил прийняття рішень і тестову вибірки, кожна з яких містить 50% даних, що дорівнює 195 векторам.

Додатково сформовано тестовий набір іншої особи з ідентифікатором s004, щоб перевірити як модель розпізнає об'єкти, що не належать цільовому образу. Характеристики цього набору наведені в таблицях 1.6 та 1.7.

Вектор вибірових середніх значень навчальної вибірки для розпізнавання клавіатурного почерку дорівнює $\bar{X} = \{0,07635; 0,07052; 0,07875; 0,06254; 0,06955; 0,08806; 0,08752; 0,07447; 0,07495\}$. Коваріаційна матриця відображена в таблиці 3.9, а діапазони характеристик навчальної вибірки - в таблиці 3.10.

Таблиця 3.9 – Коваріаційна матриця навчальної вибірки для розпізнавання клавіатурного почерку

0,0 ³ 2	0,0 ⁴ 16	0,0 ⁴ 62	0,0 ⁴ 42	0,0 ⁵ 53	0,0 ⁵ 71	0,0 ⁴ 38	0,0 ⁵ 14	-0,0 ⁵ 75
0,0 ⁴ 16	0,0 ³ 21	-0,0 ⁴ 19	-0,0 ⁵ 26	-0,0 ⁶ 35	0,0 ⁴ 14	0,0 ⁴ 62	-0,0 ⁴ 1	0,0 ⁵ 15
0,0 ⁴ 62	-0,0 ⁴ 19	0,0 ³ 26	-0,0 ⁶ 84	0,0 ⁴ 11	0,0 ⁵ 71	-0,0 ⁴ 51	0,0 ⁵ 99	0,0 ⁴ 17
0,0 ⁴ 42	-0,0 ⁵ 26	-0,0 ⁶ 84	0,0 ³ 21	0,0 ⁴ 24	-0,0 ⁶ 75	0,0 ⁴ 37	-0,0 ⁴ 28	-0,0 ⁵ 87
0,0 ⁵ 53	-0,0 ⁶ 35	0,0 ⁴ 11	0,0 ⁴ 24	0,0 ³ 13	0,0 ⁴ 12	-0,0 ⁵ 36	-0,0 ⁵ 53	0,0 ⁵ 72
0,0 ⁵ 71	0,0 ⁴ 14	0,0 ⁵ 71	-0,0 ⁶ 75	0,0 ⁴ 12	0,0 ³ 12	0,0 ⁴ 1	0,0 ⁵ 16	0,0 ⁴ 18
0,0 ⁴ 38	0,0 ⁴ 62	-0,0 ⁴ 51	0,0 ⁴ 37	-0,0 ⁵ 36	0,0 ⁴ 1	0,0 ³ 36	-0,0 ⁴ 61	0,0 ⁶ 69
0,0 ⁵ 14	-0,0 ⁴ 1	0,0 ⁵ 99	-0,0 ⁴ 28	-0,0 ⁵ 53	0,0 ⁵ 16	-0,0 ⁴ 61	0,0 ³ 27	0,0 ⁴ 23
-0,0 ⁵ 74	0,0 ⁵ 15	0,0 ⁴ 17	-0,0 ⁵ 87	0,0 ⁵ 72	0,0 ⁴ 18	0,0 ⁶ 69	0,0 ⁴ 23	0,0 ³ 22

Таблиця 3.10 – Діапазони характеристик навчальної вибірки для розпізнавання клавіатурного почерку

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
Min	0,018	0,0138	0,0304	0,0061	0,0188	0,052	0,0288	0,0322	0,0077
Max	0,1127	0,109	0,127	0,0903	0,0961	0,1235	0,1302	0,1246	0,1103

Вектор вибірових середніх значень тестової вибірки з ідентифікатором s015, що буде використовуватися для перевірки точності моделі при роботі з

цільовим класом, для розпізнавання клавіатурного почерку становить $\bar{X} = \{0,07416; 0,06992; 0,07770; 0,06347; 0,06866; 0,08851; 0,08458; 0,07564; 0,07525\}$. Діапазони характеристик тестової вибірки цільового класу відображені в таблиці 3.11, а коваріаційна матриця – в таблиці 3.12.

Таблиця 3.11 – Діапазони характеристик тестової вибірки цільового класу

	X_1	X_2	X_3	X_4	X_5	X_6	X_7	X_8	X_9
Min	0,0285	0,0335	0,0315	0,0306	0,0164	0,0531	0,0238	0,0123	0,023
Max	0,117	0,1064	0,1151	0,0985	0,0914	0,1196	0,1243	0,1064	0,1008

Таблиця 3.12 – Коваріаційна матриця тестової вибірки цільового класу

0,0 ³ 18	0,0 ⁴ 35	0,0 ⁴ 45	-0,0 ⁴ 16	0,0 ⁵ 36	0,0 ⁴ 18	0,0 ⁴ 11	-0,0 ⁴ 12	-0,0 ⁴ 26
0,0 ⁴ 35	0,0 ³ 2	-0,0 ⁴ 19	0,0 ⁴ 13	0,0 ⁴ 16	0,0 ⁴ 15	0,0 ⁴ 18	-0,0 ⁵ 93	0,0 ⁵ 14
0,0 ⁴ 45	-0,0 ⁴ 19	0,0 ³ 25	0,0 ⁵ 16	-0,0 ⁵ 45	-0,0 ⁵ 46	-0,0 ⁵ 24	0,0 ⁴ 44	-0,0 ⁴ 17
-0,0 ⁴ 16	0,0 ⁴ 13	0,0 ⁵ 16	0,0 ³ 16	0,0 ⁵ 93	-0,0 ⁵ 99	0,0 ⁴ 27	-0,0 ⁵ 85	-0,0 ⁵ 28
0,0 ⁵ 36	0,0 ⁴ 16	-0,0 ⁵ 45	0,0 ⁵ 93	0,0 ³ 13	0,0 ⁴ 17	-0,0 ⁴ 1	0,0 ⁴ 17	0,0 ⁵ 25
0,0 ⁴ 18	0,0 ⁴ 15	-0,0 ⁵ 46	-0,0 ⁵ 99	0,0 ⁴ 17	0,0 ³ 12	-0,0 ⁴ 17	0,0 ⁵ 54	-0,0 ⁶ 45
0,0 ⁴ 11	0,0 ⁴ 18	-0,0 ⁵ 24	0,0 ⁴ 27	-0,0 ⁴ 1	-0,0 ⁴ 17	0,0 ³ 34	-0,0 ⁴ 14	0,0 ⁵ 36
-0,0 ⁴ 12	-0,0 ⁵ 93	0,0 ⁴ 44	-0,0 ⁵ 85	0,0 ⁴ 17	0,0 ⁵ 54	-0,0 ⁴ 14	0,0 ³ 26	-0,0 ⁵ 89
-0,0 ⁴ 26	0,0 ⁵ 14	-0,0 ⁴ 17	-0,0 ⁵ 28	0,0 ⁵ 25	-0,0 ⁶ 45	0,0 ⁵ 36	-0,0 ⁵ 89	0,0 ³ 17

На основі результатів тесту Мардіа, досліджено, що багатовимірний розподіл навчальної вибірки відхиляється від нормального. Про це свідчить тестова статистика для багатовимірної асиметрії $N\beta_1/6$, яка дорівнює 286,99, що перевищує значення квантиля Хі-квадрат, яке становить 215,53, для 165 ступенів свободи і рівня значущості 0,005. Також, тестова статистика для багатовимірного ексцесу β_2 зі значенням 105,43 перевищує квантиль нормального розподілу, який становить 104,19, для середнього значення 99, дисперсії 4,06 і рівня значущості 0,005.

Отже, необхідно застосувати дев'ятивимірне перетворення Бокса-Кокса до негаусового випадкового вектора $\mathbf{X} = \{X_1, X_2, \dots, X_9\}^T$, щоб перетворити його на гаусівський випадковий вектор $\mathbf{Z} = \{Z_1, Z_2, \dots, Z_9\}^T$. Оптимальні оцінки параметрів перетворення знаходяться за методом максимальної правдоподібності логарифмічної функції (2.9): $\hat{\lambda}_1 = 1,36763$, $\hat{\lambda}_2 = 1,48069$, $\hat{\lambda}_3 = 1,07797$, $\hat{\lambda}_4 = 1,73925$, $\hat{\lambda}_5 = 2,10044$, $\hat{\lambda}_6 = 1,14981$, $\hat{\lambda}_7 = 1,56607$, $\hat{\lambda}_8 = 1,16855$, $\hat{\lambda}_9 = 2,11457$.

Після застосування дев'ятивимірного перетворення Бокса-Кокса з компонентами (2.10) вектор середніх нормалізованої навчальної вибірки становить $\bar{\mathbf{Z}} = \{-0,70932; -0,66184; -0,86764; -0,57016; -0,47427; -0,81642; -0,62417; -0,81443; -0,47084\}$. Діапазони характеристик вибірки наведені в таблиці 3.13, а коваріаційна матриця в таблиці 3.14.

Таблиця 3.13 – Діапазони характеристик нормалізованої навчальної вибірки

	Z_1	Z_2	Z_3	Z_4	Z_5	Z_6	Z_7	Z_8	Z_9
Min	-0,7282	-0,6742	-0,9062	-0,5749	-0,4760	-0,8407	-0,6361	-0,8403	-0,4729
Max	-0,6943	-0,6500	-0,8274	-0,5662	-0,4726	-0,7912	-0,6123	-0,7807	-0,4684

Таблиця 3.14 – Коваріаційна матриця нормалізованої навчальної вибірки

0,0 ⁴ 29	0,0 ⁵ 14	0,0 ⁴ 2	0,0 ⁵ 2	0,0 ⁶ 13	0,0 ⁵ 21	0,0 ⁵ 39	0,0 ⁷ 73	-0,0 ⁷ 68
0,0 ⁵ 14	0,0 ⁴ 16	-0,0 ⁵ 48	-0,0 ⁷ 21	0,0 ⁷ 71	0,0 ⁵ 27	0,0 ⁵ 43	-0,0 ⁵ 16	0,0 ⁷ 31
0,0 ⁴ 2	-0,0 ⁵ 48	0,0 ³ 18	-0,0 ⁶ 42	0,0 ⁶ 43	0,0 ⁵ 42	-0,0 ⁴ 11	0,0 ⁵ 51	0,0 ⁶ 65
0,0 ⁵ 2	-0,0 ⁷ 21	-0,0 ⁶ 42	0,0 ⁵ 31	0,0 ⁶ 15	-0,0 ⁶ 15	0,0 ⁵ 14	-0,0 ⁵ 2	-0,0 ⁷ 32
0,0 ⁶ 13	0,0 ⁷ 71	0,0 ⁶ 43	0,0 ⁶ 15	0,0 ⁶ 33	0,0 ⁶ 58	0,0 ⁸ 8	-0,0 ⁶ 17	0,0 ⁶ 13
0,0 ⁵ 21	0,0 ⁵ 27	0,0 ⁵ 42	-0,0 ⁶ 15	0,0 ⁶ 58	0,0 ⁴ 55	0,0 ⁵ 15	0,0 ⁶ 92	0,0 ⁶ 72
0,0 ⁵ 39	0,0 ⁵ 43	-0,0 ⁴ 11	0,0 ⁵ 14	0,0 ⁸ 8	0,0 ⁵ 15	0,0 ⁴ 22	-0,0 ⁴ 1	0,0 ⁷ 12
0,0 ⁷ 73	-0,0 ⁵ 16	0,0 ⁵ 51	-0,0 ⁵ 2	-0,0 ⁶ 17	0,0 ⁶ 92	-0,0 ⁴ 1	0,0 ³ 11	0,0 ⁶ 89
-0,0 ⁷ 68	0,0 ⁷ 31	0,0 ⁶ 66	-0,0 ⁷ 32	0,0 ⁷ 13	0,0 ⁶ 72	0,0 ⁷ 12	0,0 ⁶ 89	0,0 ⁶ 57

Нормалізована вибірка не відхиляється від багатовимірного нормального розподілу. Це підтверджується тестом Мардіа: тестова статистика для багатовимірної асиметрії $N\beta_1/6$, яка дорівнює 175,47, не перевищує значення квантиля Хі-квадрат, яке становить 215,53, для 165 ступенів свободи і рівня значущості 0,005. Також, тестова статистика для багатовимірного ексцесу β_2 зі значенням 99,76 не перевищує квантиль нормального розподілу, який становить 104,19, для середнього значення 99, дисперсії 4,06 і рівня значущості 0,005. Це дозволяє перейти до побудови негаусівської ймовірнісної моделі на основі еліпсоїда прогнозування для нормалізованих даних.

На основі (1.2) ймовірнісна модель у вигляді дев'ятивимірного еліпсоїда прогнозування для нормалізованих характеристик клавіатурного почерку з квантилем розподілу Хі-квадрат набуває вигляду:

$$(\mathbf{Z} - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z} - \bar{\mathbf{Z}}) = \chi_{9, 0,005}^2. \quad (3.5)$$

Використовуючи (3.5), побудовано вирішальне правило для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем Хі-квадрат для нормалізованого вектору характеристик:

$$(\mathbf{Z}_i - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z}_i - \bar{\mathbf{Z}}) \leq \chi_{9, 0,005}^2, \quad (3.6)$$

де, $\mathbf{Z}_i$ – нормалізований вектор даних в i -ій точці, $\mathbf{S}_Z^{-1}$ – обернена коваріаційна матриця нормалізованої навчальної вибірки s015, що зображена в таблиці 3.15. Значення квантиля розподілу Хі-квадрат для 9 ступенів свободи та рівня значущості 0,005 становить 23,59.

На основі (1.3) побудовано ймовірнісну модель у вигляді дев'ятивимірного еліпсоїда прогнозування для нормалізованих характеристик клавіатурного почерку з квантилем F -розподілу:

$$(\mathbf{Z} - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z} - \bar{\mathbf{Z}}) = \frac{k(N^2-1)}{N(N-k)} F_{9, 186, 0,005}. \quad (3.7)$$

Використовуючи (3.7), правило прийняття рішень для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик набуває вигляду:

$$(\mathbf{Z}_i - \bar{\mathbf{Z}})^T \mathbf{S}_Z^{-1} (\mathbf{Z}_i - \bar{\mathbf{Z}}) \leq \frac{k(N^2-1)}{N(N-k)} F_{9, 186, 0,005}. \quad (3.8)$$

Таблиця 3.15 – Обернена коваріаційна матриця нормалізованої навчальної вибірки

41430,72	-3064,57	-5343,78	-24880,03	3610,23	-1053,36	-8351,18	-1140,79	13186,17
-3064,57	67103,69	1541,35	8500,16	-13663,89	-2765,86	-12304,99	-75,67	-1134,91
-5343,78	1541,35	6654,65	3080,33	-7291,29	-290,20	3841,95	187,68	-8124,73
-24880,03	8500,16	3080,33	355493,15	-155308,21	3154,51	-15817,32	4676,50	5436,30
3610,23	-13663,89	-7291,29	-155308,21	3131842,06	-32258,30	11012,78	3605,18	-35231,55
-1053,36	-2765,86	-290,20	3154,51	-32258,30	18959,85	-928,60	-75,98	-22894,43
-8351,18	-12304,99	3841,95	-15817,32	11012,78	-928,60	54745,70	4558,32	-13134,08
-1140,79	-75,67	187,68	4676,50	3605,18	-75,98	4558,32	9589,50	-15353,47
13186,17	-1134,91	-8124,73	5436,30	-35231,55	-22894,43	-13134,08	-15353,47	1836658,75

Значення квантиля F -розподілу для рівня значущості 0,005, зі ступенями свободи 9 і 186 дорівнює 2,74, а значення коефіцієнта 9,44. Вся права частина рівняння становить 25,86, що перевищує критичне значення еліпсоїда прогнозування з квантилем Хі-квадрат 23,59.

Таблиця 3.16 містить порівняння за визначеними метриками побудованих еліпсоїдів прогнозування для негаусівських даних (PENGD) та еліпсоїдів прогнозування для нормалізованих даних (PEND) (3.6, 3.8) з різними квантилями розподілів в задачі розпізнавання клавіатурного почерку.

Таблиця 3.16 – Порівняння побудованих ймовірнісних моделей для розпізнавання клавіатурного почерку

Quantiles	Models	Specificity	Recall	Precision	F1 score	Accuracy
Chi-square	PENGD	0,9795	0,9225	0,9893	0,9547	0,9412
	PEND	0,9949	0,9700	0,9974	0,9835	0,9782
F-distribution	PENGD	0,9949	0,9025	0,9972	0,9475	0,9328
	PEND	1	0,9650	1	0,9822	0,9765

Результати свідчать, що використання квантиля χ^2 -квадрат забезпечує вищу загальну точність та кращий баланс між Recall і Specificity порівняно з F -розподілом. Водночас еліпсоїд прогнозування з квантилем F -розподілу більше орієнтований на точне визначення цільових точок, демонструючи дещо меншу чутливість до виявлення аномалій.

Нормалізація даних із застосуванням дев'ятивимірної перетворення Бокса-Кокса суттєво покращила всі метрики для обох типів еліпсоїдів, підкреслюючи її роль як ключового етапу в створенні ефективних моделей розпізнавання.

Дані клавіатурного почерку, як правило, суттєво відхиляються від нормального розподілу через індивідуальні особливості поведінкових характеристик користувачів [42], що обмежує використання еліпсоїда прогнозування. У роботі [45] було запропоновано тривимірний еліпсоїд прогнозування для нормалізованих даних, що дозволяє виконати припущення про багатовимірний нормальний розподіл даних. У даній роботі ймовірнісну модель у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованих даних було удосконалено завдяки використанню дев'яти характеристик клавіатурного почерку, що дозволило покращити точність (Accuracy) з 0,9412 до 0,9782, а також ймовірність розпізнавання представників інших класів (Recall) з 0,9225 до 0,9700 та ймовірність розпізнавання цільового образу (Specificity) з 0,9795 до 0,9949.

Важливим кроком стало також використання еліпсоїда прогнозування на основі квантиля F -розподілу. Такий підхід раніше не знаходив застосування у задачах розпізнавання клавіатурного почерку для даних, що відхиляються від нормального розподілу. У цій роботі вперше розроблено ймовірнісну модель еліпсоїда прогнозування для нормалізованих даних з квантилем F -розподілу. Її використання забезпечило покращення точності (Accuracy) з 0,9328 до 0,9765, ймовірності розпізнавання представників інших класів (Recall) з 0,9025 до 0,9650 та ймовірність розпізнавання цільового образу (Specificity) з 0,9949 до 1.

3.3 Порівняння точності та ймовірності розпізнавання образів за побудованими еліпсоїдами прогнозування для нормалізованих даних та за алгоритмами машинного навчання

Серед підходів до однокласової класифікації значну популярність здобули методи машинного навчання. Серед найпоширеніших прикладів можна виділити однокласову опорну векторну машину, ізоляційний ліс та автокодувальник. OCSVM — це варіант методу опорних векторів, який навчається виключно на даних цільового класу. Його мета полягає у побудові гіперплощини, що відокремлює цільові точки даних від решти простору, забезпечуючи максимальний відступ між ними. IF — це ансамблевий метод, який виконує ізоляцію аномалій через випадковий вибір ознак та поступове розділення точок даних. Цей процес триває доти, доки аномалії не будуть ізольовані у невеликих розділах, причому для цього потрібно менше розділень, ніж для виокремлення цільових точок. AE — це нейронна мережа, призначена для побудови ефективних уявлень про цільові дані. Вона виконує кодування вхідних даних у простір нижчої розмірності (кодувальник), а потім відновлює їх у початковій формі (де кодувальник). Аномалії виявляються шляхом аналізу помилки реконструкції: чим більша помилка, тим вища ймовірність, що точка є аномальною.

3.3.1 Однокласова опорна векторна машина

OCSVM будує функцію рішень або межу, яка ефективно відокремлює цільові дані від решти простору ознак. Ця межа встановлюється шляхом знаходження гіперплощини з максимальним відступом та оптимізації відстані між гіперплощиною та початком координат у багатовимірному просторі ознак. В OCSVM використовується неявна функція перетворення $\varphi(\cdot)$, яка здійснює нелінійну проєкцію, оцінену через ядрову функцію. Ядрова функція слугує відображенням із початкового простору ознак у простір вищої розмірності: $k(x, y) = \varphi(x) \cdot \varphi(y)$ [111].

Найпоширеніші ядрові функції включають:

1. лінійне ядро обчислює скалярний добуток у початковому просторі ознак і підходить для лінійно роздільних даних;
2. поліноміальне ядро моделює нелінійні зв'язки між точками даних шляхом піднесення скалярного добутку до заданого степеня, дозволяючи створювати складні межі рішень;
3. радіально-базисна функція, що використовує Гауссову функцію, ефективно захоплює складні зв'язки між точками даних, особливо у випадках, коли дані не є лінійно роздільними;
4. сигмоїдне ядро, засноване на функції гіперболічного тангенса, добре працює з нелінійними структурами у даних і підходить для моделювання складних взаємозв'язків між ознаками та класами.

Алгоритм будує межу рішень під час навчання, яка відокремлює більшість даних від простору ознак та визначається як [112]:

$$g(x) = \omega^T \varphi(x) - \rho,$$

де ω – це нормальний вектор гіперплощини, ρ – зміщення.

Формулювання OCSVM як задачі квадратичної оптимізації має на меті мінімізувати ваговий вектор, одночасно максимізуючи відступ, з урахуванням певних обмежень:

$$\min_{\omega, \xi, \rho} \frac{\|\omega\|^2}{2} - \rho + \frac{1}{vN} \sum_{i=1}^N \xi_i,$$

$$\text{за умови: } \omega^T \varphi(x_i) \geq \rho - \xi_i, \xi_i \geq 0,$$

де ξ_i – змінні відступу, які використовуються для моделювання помилок розділення; $v \in (0,1]$ – параметр регуляризації, який визначає розв'язок, встановлюючи верхню межу для частки викидів, і нижню межу для кількості тренувальних прикладів, що використовуються як опорні вектори [113].

Задача оптимізації зазвичай розв'язується у її дуальній формі, що приводить до функції рішень, яка може класифікувати нові точки даних як такі,

що належать до цільового класу, або як аномалії. Функція повертає додатне значення для цільових точок і від'ємне — для інших:

$$f(x) = \text{sgn}(g(x)).$$

Багато сучасних мов програмування містять реалізацію алгоритму. У Python побудова OCSVM зазвичай передбачає створення об'єкта OneClassSVM з модуля SVM у бібліотеці scikit-learn.

3.3.2 Ізоляційний ліс

Відходячи від традиційних підходів, що базуються на виокремленні цільових точок даних, метод IF пропонує унікальний підхід, орієнтуючись безпосередньо на ізоляцію аномалій. Цей метод ґрунтується на побудові ізоляційних дерев — бінарних дерев, де внутрішні вузли представляють ознаки та значення поділу, а листові вузли відповідають окремим точкам даних [114].

Процес побудови ізоляційних дерев починається з випадкового вибору ознаки та значення поділу в її діапазоні. Цей процес повторюється, доки кожна точка даних не буде ізолювана у власному листовому вузлі або не досягнуто заданої максимальної глибини дерева. Завдяки цьому стохастичному вибору ознак і значень поділу метод ефективно ізолює аномалії без попередніх знань про розподіл даних. Аномалії, які зазвичай легше ізолювати, ніж цільові точки даних, розташовуються у більш розріджених регіонах простору ознак. Відповідно, для їх ізоляції потрібно менше поділів на шляхах від кореня до листових вузлів. Таким чином, для кожної точки даних обчислюється середня довжина шляху від кореня до листових вузлів через усі дерева у лісі [115].

Далі для кожної точки даних розраховується оцінка аномальності на основі її середньої довжини шляху [116]:

$$s(x, n) = 2^{-\frac{E(h(x))}{c(n)}},$$

де $E(h(x)) = \frac{\sum_{i=1}^t h_i(x)}{t}$ - це середня довжина шляху x через t ізоляційних дерев;
 $c(n)$ - середня тривалість невдалого пошуку в двійковому дереві:

$$c(n) = 2H(n - 1) - (2(n - 1)/n),$$

де $H(i) = \ln(i) + \gamma$, γ - стала Ейлера.

Точки даних із коротшими шляхами, ближчими до кореня, вважаються більш імовірними аномаліями, тоді як ті, що мають довші шляхи, вважаються більш типовими. Для класифікації точок даних як аномалій або нормальних встановлюється порогове значення оцінки аномальності: точки вище порогу класифікуються як аномалії, а точки нижче — як нормальні.

Isolation Forest використовує ансамбль таких ізоляційних дерев, де кожне дерево працює незалежно, вносячи свій вклад у фінальну аномальну оцінку для заданої точки даних. Такий ансамблевий підхід підвищує стійкість і точність алгоритму, роблячи його більш ефективним у роботі з шумом та варіативністю даних.

Алгоритм IsolationForest реалізовано у широко використовуваній бібліотеці машинного навчання Python scikit-learn, яка пропонує різноманітні параметри для оптимізації його роботи. Серед них особливо важливим є параметр `contamination`, який визначає поріг у функції рішень, що дозволяє відокремлювати нові точки даних як цільові або аномальні [117]. Інші важливі параметри моделі IsolationForest включають: `n_estimators` — кількість дерев у лісі; `max_samples` — максимальна кількість зразків, що використовуються для навчання кожного дерева; `max_features` — максимальна кількість ознак, що використовуються для поділу кожного вузла.

3.3.3 Автокодувальник

Автокодувальник (автоенкодер чи автокодер) — це тип штучної нейронної мережі, призначений для задач, що потребують ефективного вивчення про особливості даних, зменшення розмірності та виявлення аномалій. Як метод навчання без вчителя, він складається з двох основних компонентів: кодувальника та декодувальника [118].

Основною метою автокодера є навчитися створювати стиснене та змістовне уявлення вхідних даних, яке часто називається латентним простором. Це уявлення захоплює найважливіші патерни та ознаки даних, одночасно відкидаючи непотрібну або надмірну інформацію. Кодувальник відповідає за перетворення вхідних даних у цей компресований латентний простір. Це досягається за допомогою серії нелінійних перетворень через шари, поступово зменшуючи розмірність вхідних даних. Ці перетворення допомагають виявляти важливі риси та структури, що важливі для розуміння даних. У результаті, латентний простір стає компактним, абстрактним уявленням вхідних даних, зберігаючи їх найважливіші аспекти.

З іншого боку, декодувальник бере це компресоване уявлення з латентного простору та відновлює оригінальні вхідні дані. Архітектура декодувальника зазвичай є дзеркальним відображенням кодувальника, але в зворотному напрямку, і його шари розширюють дані назад до їх початкового виміру. Так само, як кодувальник застосовує нелінійні перетворення, декодувальник застосовує зворотні перетворення, щоб забезпечити якомога точніше відновлення даних. Здатність відновлювати дані з мінімальною помилкою вказує на те, що кодувальник навчився ефективно уявляти вхідні дані.

Під час навчання автокодер намагається мінімізувати помилку реконструкції, що є різницею між оригінальними вхідними даними та відновленими вихідними. Помилка реконструкції є важливим показником здатності моделі навчитися релевантним ознакам даних. Це зазвичай досягається оптимізацією функції втрат, такої як середньоквадратична помилка або бінарна крос-ентропія. Ці функції втрат кількісно визначають різницю між оригінальними та відновленими даними, а оптимізація здійснюється за допомогою методів, заснованих на градієнтах, таких як зворотне поширення.

У задачах виявлення аномалій або однокласової класифікації автокодер тренується лише на екземплярах цільового класу. Таким чином, модель навчається розпізнавати регулярні патерни та структури, що визначають цільові дані. Коли автокодер стикається з новими даними, він відновлює їх з низькою

помилкою, якщо нові дані належать до цільового класу. Проте, якщо вхідні дані є аномалією, тобто, вони відрізняються від вивчених патернів, помилка реконструкції буде вищою, оскільки автокодер не здатний точно відновити незнайомі дані [119]. Встановивши поріг для помилки реконструкції, можна класифікувати точки даних як нормальні або аномальні залежно від того, чи перевищує помилка поріг, чи ні.

При розробці моделей автокодерів важливим аспектом є вибір відповідних фреймворків машинного навчання. Серед найбільш популярних фреймворків для створення та тренування автокодерів є TensorFlow та PyTorch [120].

TensorFlow, розроблений компанією Google, пропонує багатий екосистему інструментів та бібліотек, що робить його дуже зручним для створення та розгортання моделей машинного навчання в масштабах. Статичний граф обчислень у TensorFlow забезпечує ефективність обчислень, що робить його ідеальним для великих моделей і для розподілених обчислень. TensorFlow широко використовується як у дослідженнях, так і в промислових застосунках.

PyTorch, розроблений лабораторією досліджень штучного інтелекту компанії Facebook, славиться своєю гнучкістю та простотою у використанні, особливо для досліджень і прототипування. На відміну від статичного графа TensorFlow, PyTorch використовує динамічний граф обчислень, що дозволяє гнучко працювати з моделями та зручно налагоджувати код. Інтуїтивно зрозумілий API робить PyTorch привабливим для дослідників і розробників, які хочуть швидко проводити експерименти.

У розробці моделей автокодеру, було вирішено зупинитися на поєднанні TensorFlow та Keras. Keras діє як високорівневий API для нейронних мереж, що працює поверх TensorFlow, спрощуючи процес проєктування та тренування моделей, а також надає зручний інтерфейс для створення архітектур нейронних мереж і є ідеальним для тих, хто хоче швидко реалізувати складні моделі без необхідності працювати з низькорівневими аспектами. TensorFlow надає обчислювальну інфраструктуру, що забезпечує ефективні операції для навчання, оптимізації та виводу.

3.3.4 Реалізація алгоритмів машинного навчання для розпізнавання облич

Цей розділ описує реалізацію алгоритмів машинного навчання OCSVM, IF та AE для розпізнавання характеристик обличчя.

OCSVM реалізований мовою Python за допомогою бібліотеки `scikit-learn`. Ця реалізація дозволяє налаштовувати кілька критичних параметрів, включаючи функцію ядра та параметр регуляризації ν . Було встановлено значення ν на 0,15, що визначає допустиму частку помилок навчання та встановлює верхню межу для частки аномальних даних. Для моделювання нелінійних зв'язків між точками даних було обрано радіально-базисна функцію. Параметр `gamma` встановлений на "auto", що дозволяє його значенню автоматично обчислюватися на основі оберненого значення кількості ознак, що впливає на діапазон для кожного навчального прикладу; менші значення збільшують вплив, а більші значення локалізують його.

Алгоритм IF, також реалізований за допомогою `scikit-learn` [121], що надає можливість налаштування декількох параметрів для оптимізації навчання. Ключовим параметром є `contamination`, який визначає поріг для категоризації нових точок даних як цільові або аномальні, значення якого було встановлено до 0,08, що ефективно балансує виявлення справжніх аномалій проти хибних спрацьовувань. Додатково `n_estimators` встановлено на 100, `max_samples` встановлено на 256 і `max_features` встановлено на 1.0. Вибір порогового значення залежить від застосування; в даному аналізі оптимальним значенням порогу є 0.

Для моделі AE було використано TensorFlow та Keras. Перед тим як передати дані в нейронну мережу, застосовується нормалізація `min-max` для кожної ознаки окремо, масштабуючи всі ознаки в діапазон $[0, 1]$. Ця техніка стандартизує ознаки, сприяючи стабільним та ефективним процесам навчання.

Архітектура AE складається з вхідного шару, налаштованого для прийому дев'ятивимірної представлення даних. Модель включає повнозв'язні шари для операцій кодування та декодування. Під час етапу кодування вхідні дані

стискаються в низькорозмірне уявлення, поступово зменшуючи розмірність з 10 до 8, а потім до 6, створюючи латентний шар в структурі мережі. Цей шар змушує модель захоплювати основні ознаки вхідних даних, одночасно мінімізуючи надлишковість [122].

Кожен шар кодування використовує функцію активації rectified linear unit (ReLU), яка вводить нелінійність, що сприяє виділенню складних ознак. Етап декодування інвертує цей процес, розширюючи розмірність назад до 8 і в кінцевому підсумку до оригінальних 10 вимірів, використовуючи функції активації ReLU для збереження вивчених нелінійних взаємозв'язків. Останній шар використовує функцію активації sigmoid, щоб обмежити вихідні значення в діапазоні $[0, 1]$, що є стандартним вибором для задач реконструкції та бінарної класифікації, які потребують плавних та інтерпретованих результатів. Структура АЕ показана на Рисунок 3.1.

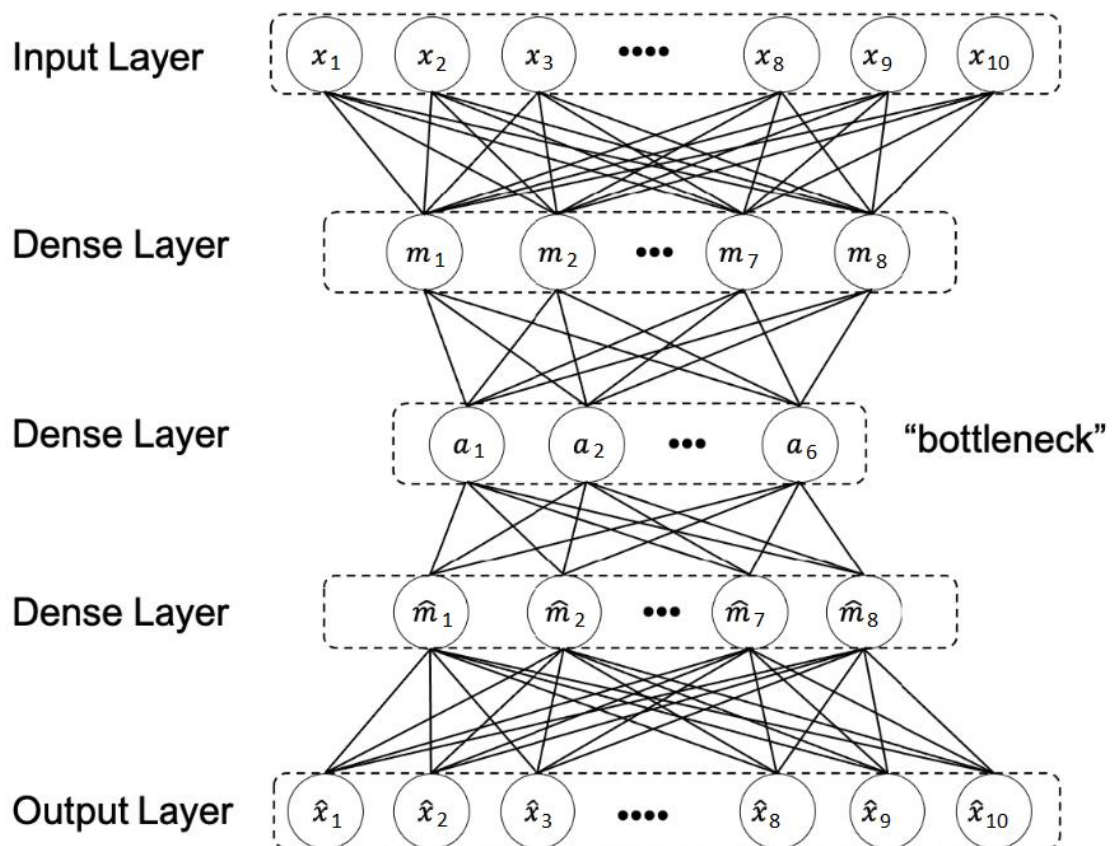


Рисунок 3.1 – Структура автоенкодера для розпізнавання характеристик облич

Для навчання моделі був використаний оптимізатор Adam у поєднанні з функцією втрат бінарна крос-ентропія, що є стандартною метрикою для задач реконструкції, спрямованих на мінімізацію розбіжностей між оригінальними та відновленими даними. Оптимізатор Adam поєднує переваги AdaGrad та RMSProp [123], динамічно коригуючи швидкість навчання під час тренування для швидшої збіжності та покращення ефективності. Функція втрат бінарна крос-ентропія ефективно вимірює різницю між вхідними та відновленими виходами, що робить її підходящою для задач бінарної класифікації. Процес навчання охоплює 100 епох, з розміром батчу 8 екземплярів.

Таблиця 3.17 містить порівняння результатів розпізнавання реалізованих алгоритмів машинного навчання, а також еліпсоїда прогнозування для негаусівських даних (PENGD-Chi) і еліпсоїда прогнозування для нормалізованих даних (PEND-Chi) з квантилем розподілу χ^2 та еліпсоїда прогнозування для негаусівських даних (PENGD-F) і еліпсоїда прогнозування для нормалізованих даних (PEND-F) з квантилем F -розподілу.

Таблиця 3.17 – Порівняння побудованих моделей для розпізнавання облич

Model	Specificity	Recall	Precision	F1 score	Accuracy
PENGD-Chi	0,9200	0,9550	0,9598	0,9574	0,9433
PEND-Chi	0,9700	0,9750	0,9848	0,9793	0,9733
PENGD-F	0,9900	0,9350	0,9946	0,9639	0,9533
PEND-F	1	0,9350	1	0,9664	0,9567
OCSVM	0,7700	0,9100	0,8878	0,8988	0,8633
IF	0,8800	0,9300	0,9393	0,9347	0,9133
AE	0,8700	0,9600	0,9366	0,9481	0,9300

Загалом, всі розроблені моделі демонструють хорошу точність. OCSVM показав найнижчі результати серед усіх оцінюваних моделей за всіма метриками, такими як precision, recall, specificity, F1 score та accuracy. IF продемонстрував

конкурентоспроможні результати за більшістю метрик, займаючи проміжне місце між результатами OCSVM та іншими моделями.

АЕ показав результат, порівнянний з PENGD-Chi, однак, завдяки застосуванню нормалізуючого перетворення PEND-Chi досяг значних покращень порівняно з АЕ, зокрема точність розпізнавання (Accuracy) 0,9733 проти 0,93, ймовірність розпізнавання цільового образу (Specificity) 0,97 проти 0,87 та ймовірність розпізнавання представників інших класів (Recall) 0,975 проти 0,96. Найвищу ймовірність розпізнавання цільового образу (Specificity) мають PENGD-F зі значенням 0,99 та PEND-F зі значенням 1, нормалізація допомогла покращити Specificity і загальну точність (Accuracy) до 0,9567, при цьому ймовірність розпізнавання представників інших класів (Recall) нижче 0,9350, ніж у PENGD-Chi, PEND-Chi та АЕ.

Результати підкреслюють важливість застосування нормалізуючих перетворень для покращення моделей еліпсоїда прогнозування для задач розпізнавання, особливо в умовах наявності негаусівських розподілів даних.

3.3.5 Реалізація алгоритмів машинного навчання для розпізнавання клавіатурного почерку

У цьому розділі описано реалізацію кількох алгоритмів машинного навчання, застосованих для розпізнавання клавіатурного почерку. Для цієї мети були обрані моделі OCSVM, IF та АЕ, кожна з яких володіє специфічними характеристиками, що дозволяють ефективно виявляти аномалії та виконувати однокласову класифікацію.

Модель OCSVM була розроблена з використанням Python та бібліотеки scikit-learn. Важливим налаштуванням є параметр ν , який визначає допустиму частку помилок під час навчання, він встановлений до 0,03. Для побудови нелінійних зв'язків між даними було використано радіально-базисне ядро, а параметр gamma було налаштовано на значення "auto", що дозволяє автоматично підлаштовувати його залежно від кількості ознак у наборі.

Алгоритм Isolation Forest також реалізований за допомогою бібліотеки `scikit-learn` і має низку налаштувань для поліпшення результатів. Один з найважливіших параметрів — це `contamination`, що визначає поріг для класифікації даних як нормальних або аномальних. У даному випадку значення `contamination` було встановлено на 0,05, що забезпечує оптимальний баланс між точністю виявлення аномалій і мінімізацією хибних спрацьовувань. Додаткові параметри включають `n_estimators` - кількість дерев у лісі, встановлено значення 100; `max_samples` - максимальна кількість вибірок на кожному дереві встановлено значення 256; `max_features` - максимальна кількість ознак для кожного дерева, встановлено 1,0 для врахування всіх ознак.

Модель автоенкодера була реалізована за допомогою TensorFlow та Keras, що дозволяє зручно будувати та тренувати нейронні мережі. Для попередньої обробки даних використовувався метод нормалізації `min-max`, що масштабував значення кожної ознаки до діапазону $[0, 1]$. Така нормалізація допомагає забезпечити стабільність і швидкість навчання, оскільки зменшується ймовірність того, що одна ознака переважає інші через різницю в масштабах.

Архітектура автоенкодера складається з кількох етапів. Вхідний шар приймає дев'ятивимірне представлення даних, що відповідає кількості ознак у наборі. Кодувальник складається з двох шарів, які по черзі зменшують розмірність даних — спочатку до 8, а потім до 6 вимірів. Для кожного з цих шарів використовуються функції активації ReLU, що допомагає моделі навчитися нелінійних залежностей між ознаками.

Етап декодування є зворотнім процесом, на якому розмірність даних збільшується до початкових значень. Спочатку відновлюється 8 вимірів, а потім — до оригінальних 9. Для цього також використовуються функції активації ReLU, що сприяє збереженню складних нелінійних взаємозв'язків, вивчених під час кодування. В кінці мережі розташований останній шар, який використовує функцію активації `sigmoid`, щоб обмежити вихідні значення в діапазоні $[0, 1]$. Це є стандартним підходом для задач відновлення даних або бінарної класифікації, оскільки дозволяє зберігати чіткість і інтерпретованість результатів.

Такий підхід дозволяє моделі ефективно відновлювати інформацію, стискаючи її до найважливіших ознак і потім відновлюючи до вихідного вигляду. Структура автоенкодера зображена на Рисунку 3.2.

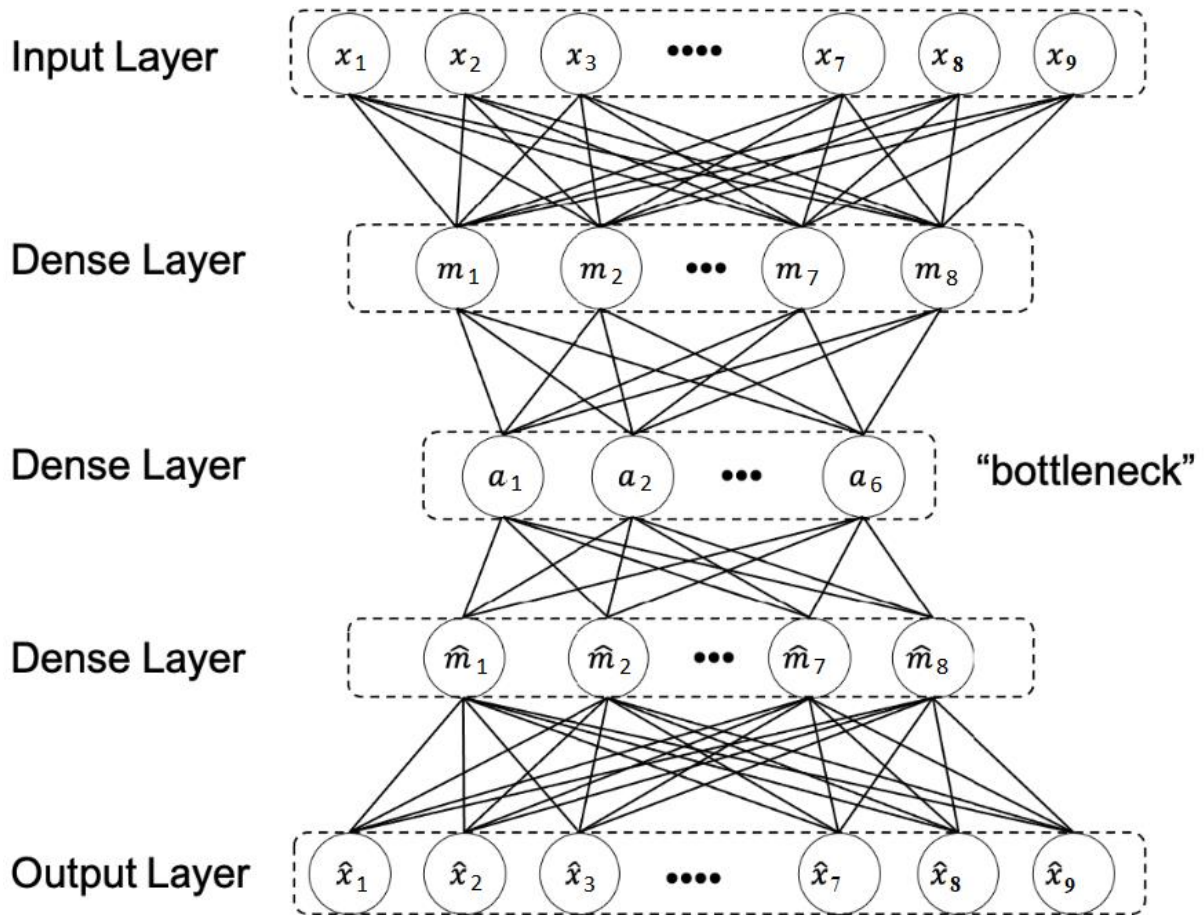


Рисунок 3.2 – Структура автоенкодера для розпізнавання характеристик клавіатурного почерку

Для навчання моделі був використаний оптимізатор Adam у поєднанні з функцією втрат бінарна крос-ентропія. Процес навчання охоплює 25 епох, з розміром батчу 16 екземплярів.

Таблиця 3.18 містить порівняння результатів розпізнавання реалізованих алгоритмів машинного навчання, а також еліпсоїда прогнозування для негаусівських даних (PENGD-Chi) і еліпсоїда прогнозування для нормалізованих даних (PEND-Chi) з квантилем розподілу χ^2 -квадрат та еліпсоїда прогнозування для негаусівських даних (PENGD-F) і еліпсоїда прогнозування для нормалізованих даних (PEND-F) з квантилем F -розподілу.

Усі моделі демонструють високу точність у задачі розпізнавання клавіатурного почерку. Однак PENG-D-Chi та PENG-D-F виділяються найнижчою точністю серед оцінюваних моделей, що вказує на вплив негаусівського розподілу даних. З іншого боку, як OCSVM, так і AE демонструють дуже хороші показники за кількома метриками, що відображає їх здатність з високою точністю виявляти справжні аномалії. Ці моделі ефективно використовують свої архітектури для захоплення складних взаємозв'язків у даних, що сприяє їхній стійкій роботі. У той же час IF показав найнижчу точність серед моделей машинного навчання.

Таблиця 3.18 – Порівняння побудованих моделей для розпізнавання клавіатурного почерку

Model	Specificity	Recall	Precision	F1 score	Accuracy
PENG-D-Chi	0,9795	0,9225	0,9893	0,9547	0,9412
PEND-Chi	0,9949	0,9700	0,9974	0,9835	0,9782
PENG-D-F	0,9949	0,9025	0,9972	0,9475	0,9328
PEND-F	1	0,9650	1	0,9822	0,9765
OCSVM	0,9744	0,9675	0,9872	0,9773	0,9697
IF	0,9333	0,9500	0,9669	0,9584	0,9445
AE	0,9641	0,9625	0,9821	0,9722	0,9630

Зрештою, PEND-Chi має найвищу точність (Accuracy) 0,9782, на другому місці PEND-F 0,9765, далі йдуть 0,9697 (OCSVM) і 0,9630 (AE), ймовірність розпізнавання цільового образу (Specificity) становить 0,9949 (PEND-Chi), найвищий показник 1 в PEND-F проти 0,9744 (OCSVM) і 0,9641 (AE) та ймовірність розпізнавання представників інших класів (Recall) становить 0,97 (PEND-Chi) і 0,9650 (PEND-F) проти 0,9675 (OCSVM) і 0,9625 (AE). Це підкреслює важливість нормалізуючих перетворень для покращення ймовірнісних моделей в задачах розпізнавання, особливо в ситуаціях, що передбачається негаусівський розподіл даних.

3.4 Висновки за розділом 3

У третьому розділі було побудовано ймовірнісні моделі у вигляді еліпсоїдів прогнозування для нормалізованих за допомогою десятивимірного перетворення Бокса-Кокса даних з різними квантилями розподілів для розпізнавання облич. Було побудовано ймовірнісні моделі у вигляді еліпсоїдів прогнозування для нормалізованих за допомогою дев'ятивимірного перетворення Бокса-Кокса даних з різними квантилями розподілів для розпізнавання клавіатурного почерку.

1) Вперше побудована ймовірнісна модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірного перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу. Використання побудованої ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9533 до 0,9567, ймовірність розпізнавання цільового образу (Specificity) з 0,99 до 1, при тому ймовірність розпізнавання представників інших класів (Recall) залишилася 0,9350.

2) Вперше побудована ймовірнісна модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірного перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу. Використання побудованої ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9328 до 0,9765, ймовірність розпізнавання цільового образу (Specificity) з 0,9949 до 1 та ймовірність розпізнавання представників інших класів (Recall) з 0,9025 до 0,9650.

3) Удосконалено ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого

вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірною перетворення Бокса-Кокса. Використання створеної ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9433 до 0,9733, ймовірність розпізнавання цільового образу (Specificity) з 0,9200 до 0,9700 та ймовірність розпізнавання представників інших класів (Recall) з 0,9550 до 0,9750.

4) Удосконалено ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірною перетворення Бокса-Кокса. Використання створеної ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9412 до 0,9782, ймовірність розпізнавання цільового образу (Specificity) з 0,9795 до 0,9949 та ймовірність розпізнавання представників інших класів (Recall) з 0,9225 до 0,9700.

5) Розглянуто та реалізовано алгоритми машинного навчання для однокласового розпізнавання, такі як однокласова опорно векторна машина, ізоляційний ліс та автокодувальник для розпізнавання облич та клавіатурного почерку.

6) Не залежно від квантиля, нормалізація за допомогою десятивимірною перетворення Бокса-Кокса призводить до покращення точності розпізнавання моделей для розпізнавання облич та клавіатурного почерку в порівнянні з моделями еліпсоїдів прогнозування для не нормалізованих даних. Це вказує на необхідність використання багатовимірних перетворень для побудови ймовірнісних моделей у задачах розпізнавання образів.

7) Результати свідчать, що використання квантиля χ^2 -квадрат забезпечує вищу загальну точність та кращий баланс між Recall і Specificity порівняно з F -розподілом. Водночас еліпсоїди прогнозування з квантилем F -

розподілу більше орієнтовані на точне визначення цільових точок, демонструючи дещо меншу чутливість до виявлення аномалій.

8) Виконано порівняння побудованих ймовірнісних моделей для розпізнавання облич з реалізованими алгоритмами машинного навчання. Однокласова опорно векторна машина та ізоляційний ліс показали найнижчу точність. Автокодувальник (АЕ) показав результат, порівнянний з еліпсоїдом прогнозування для негаусівських даних з квантилем розподілу Хі-квадрат (PENGD-Chi), однак, завдяки застосуванню нормалізуючого перетворення еліпсоїд прогнозування для нормалізованих даних з квантилем Хі-квадрат (PEND-Chi) досяг значних покращень порівняно з АЕ, зокрема точність розпізнавання (Accuracy) 0,9733 проти 0,93, ймовірність розпізнавання цільового образу (Specificity) 0,97 проти 0,87 та ймовірність розпізнавання представників інших класів (Recall) 0,975 проти 0,96. Найвищу ймовірність розпізнавання цільового образу (Specificity) мають еліпсоїд прогнозування для негаусівських даних з квантилем F -розподілу (PENGD-F) зі значенням 0,99 та еліпсоїд прогнозування для нормалізованих даних з квантилем F -розподілу (PEND-F) зі значенням 1, нормалізація допомогла покращити Specificity і загальну точність (Accuracy) до 0,9567, при цьому ймовірність розпізнавання представників інших класів (Recall) нижче 0,9350, ніж у PENGD-Chi, PEND-Chi та АЕ. Це підкреслює важливість нормалізуючих перетворень для покращення моделей еліпсоїда прогнозування для задач розпізнавання облич при роботі з негаусівськими даними.

9) Виконано порівняння розроблених ймовірнісних моделей для розпізнавання клавіатурного почерку з реалізованими алгоритмами машинного навчання. Еліпсоїди прогнозування для негаусівських даних виділяються найнижчою точністю серед оцінюваних моделей. ІF показав найнижчу точність серед моделей машинного навчання. OCSVM і АЕ демонструють майже однакову точність та ймовірність розпізнавання. Еліпсоїд прогнозування для нормалізованих даних з квантилем розподілу Хі-квадрат виявився (PEND-Chi) найточнішою моделлю (Accuracy) 0,9782, на другому місці еліпсоїд

прогнозування для нормалізованих даних з квантилем F -розподілу PEND-F 0,9765, далі йдуть 0,9697 (OCSVM) і 0,9630 (AE), ймовірність розпізнавання цільового образу (Specificity) становить 0,9949 (PEND-Chi), найвищий показник 1 в PEND-F проти 0,9744 (OCSVM) і 0,9641 (AE) та ймовірність розпізнавання представників інших класів (Recall) становить 0,97 (PEND-Chi) і 0,9650 (PEND-F) проти 0,9675 (OCSVM) і 0,9625 (AE). Це підкреслює важливість нормалізуючих перетворень для покращення ймовірнісних моделей в задачах розпізнавання клавіатурного почерку.

Результати досліджень поточного розділу автором опубліковано в [47-50, 103, 110, 124].

РОЗДІЛ 4

ІНФОРМАЦІЙНІ ТЕХНОЛОГІЇ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ

У четвертому розділі розроблено інформаційні технології для розпізнавання облич за 10 характеристиками та розпізнавання клавіатурного почерку за 9 характеристиками на основі розроблених негаусівських ймовірнісних моделей. Для цього було використано мову програмування Python та Qt Creator 14.0.2 Community. Також запропоновано інструкції користувача для розпізнавання облич та клавіатурного почерку з використанням відповідного програмного забезпечення.

Інформаційні технології – це широкий клас дисциплін та сфер діяльності, що належать до технологій управління та обробки даних, у тому числі із застосуванням обчислювальної техніки [125]. Згідно з визначенням, прийнятим ЮНЕСКО, ІТ - це комплекс взаємопов'язаних наукових, технологічних, інженерних наук, які вивчають методи ефективної організації праці людей, зайнятих обробкою та зберіганням інформації за допомогою обчислювальної техніки та методи організації та взаємодії з людьми та виробничим обладнанням, їх практичне застосування, а також пов'язані з усім цим соціальні, економічні та культурні проблеми [126].

В даний час під інформаційними технологіями найчастіше розуміють комп'ютерні технології. Зокрема, ІТ мають справу з використанням комп'ютерів та програмного забезпечення для зберігання, перетворення, захисту, обробки, передачі та отримання інформації. Фахівців з комп'ютерної техніки та програмування часто називають ІТ-фахівцями.

Тому кінцевим визначенням можна вважати: ІТ – це сукупність методів, виробничих та програмно-технологічних засобів, об'єднаних у технологічний ланцюжок, що забезпечує збирання, зберігання, обробку, виведення та

поширення інформації. Інформаційні технології призначені для зниження трудомісткості процесів використання інформаційних ресурсів.

4.1 Створення інформаційної технології для розпізнавання облич

Метою створення ІТ для розпізнавання облич є підвищення точності та ймовірності розпізнавання облич за рахунок використання відповідних ймовірнісних моделей, які наведені в розділі 3 дисертації.

ІТ для розпізнавання облич призначена для вирішення наступних задач:

- виділення ознак та підготовка даних: ІТ надає змогу завантажити зображення з обличчям або зробити фотографії за допомогою камери, після чого автоматично виділяє ключові відстані між рисами обличчя для розпізнавання;
- нормалізація даних за допомогою завантажених ймовірнісних моделей;
- побудова негаусівських ймовірнісних моделей у вигляді еліпсоїда прогнозування на основі коваріаційних матриць та векторів середніх, отриманих з завантажених моделей;
- обчислення квадрата відстані Махаланобіса для вимірювання схожості між біометричними даними особи та попередньо визначеним набором моделей;
- ідентифікація та класифікація: ІТ використовує перетворені відстані та статистичні пороги для визначення відповідності особи наявному профілю;
- ідентифікація нерозпізнаних осіб і розпізнавання їх як невідомих для подальшої обробки;
- відображення детальних результатів для огляду, підтримуючи ручне або автоматизоване прийняття рішень;
- збереження отриманих результатів у файлах з табличним представленням даних.

Математичне забезпечення інформаційної технології для розпізнавання за рисами обличчя включає математичні моделі та алгоритми для обробки та аналізу даних. До них належать: одновимірне та багатовимірне перетворення Бокса-Кокса для нормалізації даних, яке дозволяє зменшити асиметрію та

підвищити точність статистичного аналізу; побудова негаусівських ймовірнісних моделей у вигляді еліпсоїда прогнозування які враховують відхилення розподілу даних від нормального; обчислення квадрата відстані Махаланобіса для оцінювання схожості рис обличчя з наявним набором даних; застосування квантилів розподілів для розпізнавання, чи відповідають характеристики обличчя особи заданому профілю.

Інформаційне забезпечення ІТ для розпізнавання обличчя включає попередньо визначені моделі даних для збереження коваріаційних матриць, векторів середніх значень та параметрів перетворення; файли для збереження результатів розпізнавання.

Програмне забезпечення інформаційної технології для розпізнавання обличчя — це кросплатформний застосунок, розроблений з використанням Python, Qt Creator та бібліотеки Dlib, який призначений для аналізу особливостей обличчя та зіставлення їх з попередньо навченими моделями для аутентифікації та ідентифікації користувачів.

Python — це високорівнева багатофункціональна мова програмування, яка вирізняється простотою та зручністю читання коду. Створена для забезпечення чистого та ефективного написання програм, Python широко використовується в розробці вебдодатків, аналізі даних, штучному інтелекті та наукових обчисленнях. Завдяки великій екосистемі бібліотек Python надає готові інструменти для численних задач, що значно скорочує час розробки. Інтерпретована природа мови забезпечує її незалежність від платформи, дозволяючи працювати на Windows, macOS і Linux [127].

Для застосунків розпізнавання обличчя основні переваги Python включають наявність потужних бібліотек, таких як Dlib, TensorFlow і OpenCV [128], які спрощують реалізацію складних алгоритмів обробки зображень та машинного навчання. Завдяки зручному синтаксису Python, розробники можуть зосередитися на логіці й вирішенні задач замість роботи з деталями коду, що робить мову ідеальною для швидкого прототипування та ітеративної розробки.

Величезна спільнота Python є ще однією значною перевагою. Розробники можуть швидко знайти ресурси, навчальні матеріали й підтримку для будь-якої задачі, що робить мову особливо корисною для академічних і дослідницьких проєктів, де важливі гнучкість і швидкість експериментів.

Qt Creator — це потужне середовище розробки для створення кросплатформних застосунків. Розроблене творцями фреймворку Qt, воно підтримує дизайн і реалізацію інтерфейсів користувача поряд із надійною серверною логікою. Qt Creator відзначається можливістю створення адаптивних і привабливих графічних інтерфейсів, що є важливим для систем розпізнавання облич, які потребують інтерактивності.

Головна перевага Qt Creator — його кросплатформність. Застосунки, створені з використанням Qt, можуть працювати на Windows, macOS, Linux і навіть на вбудованих системах без значних змін у коді. Це забезпечує ефективне розгортання застосунків на різних платформах.

Для розробників на Python Qt надає відмінні інструменти інтеграції через PyQt та PySide, що дозволяє поєднати простоту Python із потужністю Qt для розробки графічних інтерфейсів. Крім того, Qt Creator має вдосконалений інструмент дизайну, який спрощує створення інтерфейсів, а інтегровані засоби налагодження та профілювання пришвидшують розробку [129].

Поєднання Python та Qt Creator забезпечує створення потужного, надійного і ефективного застосунку. Python спрощує інтеграцію складних математичних моделей і функцій, а Qt Creator гарантує якісний інтерфейс. Dlib є ключовим елементом застосунку, що забезпечує надійне розпізнавання облич. Її попередньо навчені моделі та сучасні алгоритми дозволяють точно локалізувати орієнтири обличчя навіть у складних умовах.

Дане програмне забезпечення об'єднує сучасні математичні методи, можливості машинного навчання та інтуїтивно зрозумілий інтерфейс для забезпечення надійного розпізнавання облич. Серед можливих сфер використання: системи безпечного доступу, персоналізовані користувацькі інтерфейси та управління відвідуваністю.

Обрані технології забезпечують сумісність із операційними системами Windows, macOS та Linux, що дозволяє розгортати застосунок у різноманітних середовищах без змін у коді. Це знижує витрати на впровадження та підвищує гнучкість застосунку в умовах різних технічних інфраструктур. Графічне вікно програми наведено на рис. 4.1.

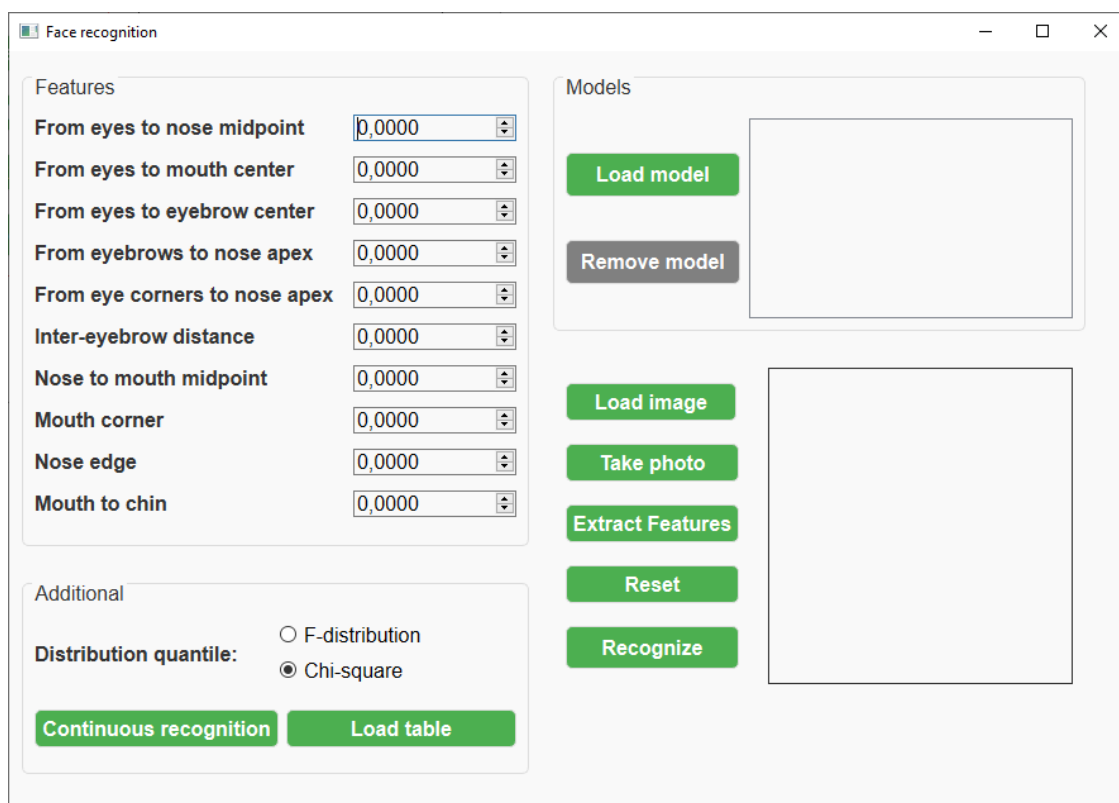


Рисунок 4.1 – Графічне вікно програми для розпізнавання облич

Діаграма станів ПЗ «Face recognition» показана на рис. 4.2. Основним станом застосунку є відображення головного вікна, під час якого відбувається очікування дій користувача.

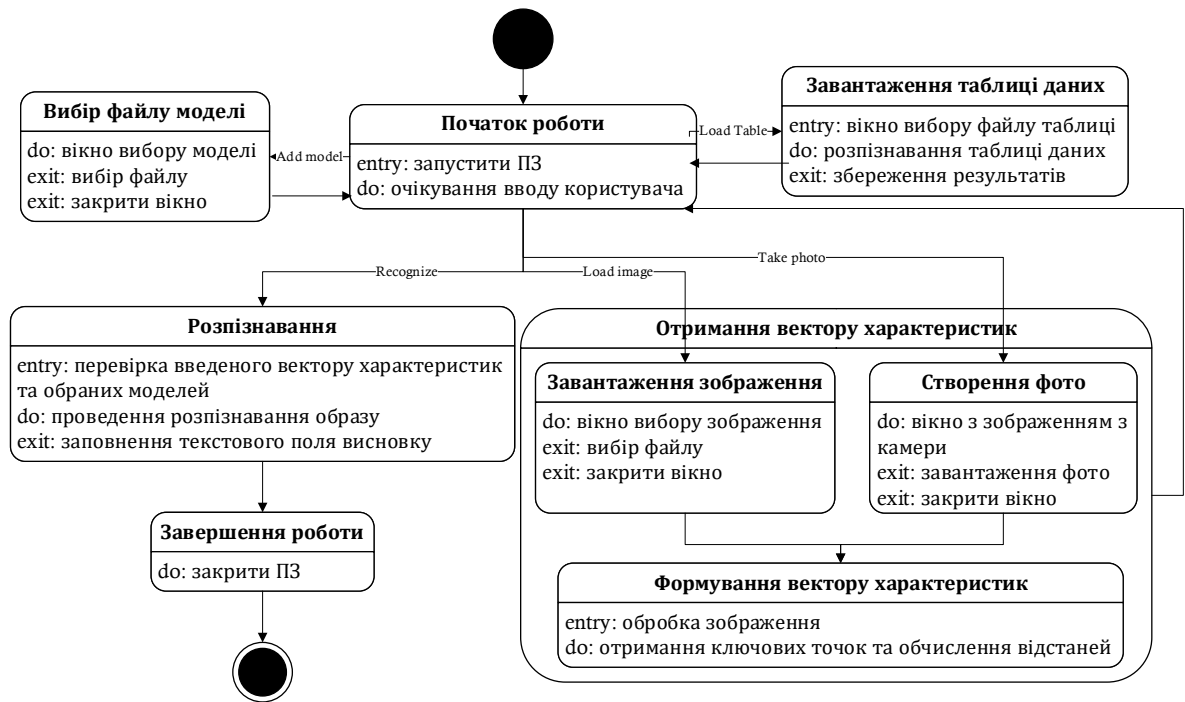


Рисунок 4.2 – Діаграма станів застосунку для розпізнавання облич

Технічне забезпечення інформаційної технології. Для ефективної роботи застосунку потрібні такі мінімальні апаратні характеристики:

Процесор: процесор не нижче Pentium IV або еквівалентний.

Оперативна пам'ять: мінімум 256 МБ ОЗП (рекомендовано 1 ГБ для оптимальної продуктивності).

Місце на диску: щонайменше 200 МБ вільного місця для встановлення.

Ці характеристики забезпечують стабільну роботу застосунку на стандартному обладнанні. Процес пошуку ключових точок та обрахування відстаней між ними для отримання вектору характеристик потребує найбільше ресурсів.

Методичне забезпечення інформаційної технології. Методичним забезпеченням інформаційної технології для розпізнавання облич є інструкція користувача для розпізнавання облич з використанням відповідного програмного забезпечення – комп'ютерної програми, розробленої в рамках даної роботи.

Впровадження інформаційної технології. Інформаційну технологію для розпізнавання облич було впроваджено в виробничі процеси ТОВ НВЦ

«Діагностика та контроль». Впровадження запропонованої ІТ для розпізнавання облич дозволило підвищити точність та ймовірність розпізнавання облич. Також, математичне забезпечення та інструментарій інформаційної технології для розпізнавання облич використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та управління проектами (ННІКНУП) Національного університету кораблебудування ім. адмірала Макарова наступних дисциплін: «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»; «Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки». Акти впровадження результатів дисертаційної роботи наведені в додатках А і Б.

4.2 Створення інформаційної технології для розпізнавання клавіатурного почерку

Метою створення ІТ для розпізнавання клавіатурного почерку є підвищення точності та ймовірності розпізнавання клавіатурного почерку за рахунок використання відповідних ймовірнісних моделей, які наведені в розділі 3 дисертації.

ІТ для розпізнавання клавіатурного почерку призначена для вирішення наступних задач:

- отримання вектору характеристик клавіатурного почерку та підготовка даних до розпізнавання;
- нормалізація даних за допомогою завантажених ймовірнісних моделей;
- побудова негаусівських ймовірнісних моделей у вигляді еліпсоїда прогнозування на основі коваріаційних матриць та векторів середніх, отриманих з завантажених моделей для розпізнавання клавіатурного почерку;

- обчислення квадрата відстані Махаланобіса для вимірювання схожості між образом клавіатурного почерку та попередньо визначеним набором моделей;
- ідентифікація та класифікація: ІТ використовує перетворені характеристики та статистичні пороги для визначення відповідності об'єкту наявним класам;
- ідентифікація нерозпізнаних об'єктів і розпізнавання їх як невідомих для подальшої обробки;
- відображення детальних результатів для огляду, підтримуючи ручне або автоматизоване прийняття рішень;
- збереження отриманих результатів у файлах з табличним представленням даних.

Математичне забезпечення інформаційної технології для розпізнавання клавіатурного почерку за вектором характеристик включає математичні моделі та алгоритми для обробки та аналізу даних. До них належать: одновимірне та багатовимірне перетворення Бокса-Кокса для нормалізації даних, яке дозволяє зменшити асиметрію та підвищити точність статистичного аналізу; побудова негаусівських ймовірнісних моделей у вигляді еліпсоїда прогнозування які враховують відхилення розподілу даних від нормального; обчислення квадрата відстані Махаланобіса для оцінювання схожості клавіатурного почерку з наявним набором даних; застосування квантилів розподілів для розпізнавання, чи відповідають характеристики клавіатурного почерку заданому профілю.

Інформаційне забезпечення ІТ для розпізнавання клавіатурного почерку включає попередньо визначені моделі даних для збереження коваріаційних матриць, векторів середніх значень та параметрів перетворення; файли для збереження результатів розпізнавання.

Програмне забезпечення інформаційної технології для розпізнавання клавіатурного почерку розроблений з використанням потужних технологій Python [127] та Qt Creator [129]. Python забезпечує простоту розробки і широкий набір інструментів для реалізації складних функцій, тоді як Qt Creator дозволяє створювати сучасний і зручний інтерфейс користувача.

Цей застосунок є кросплатформним і може бути розгорнутий на операційних системах Windows, macOS та Linux, що забезпечує його універсальність та адаптивність. Завдяки такій реалізації, програмне забезпечення не потребує значних ресурсів для роботи, що робить його ефективним у різних середовищах.

Для розпізнавання клавіатурного почерку застосунок використовує створені ймовірнісні моделі, які аналізують час утримання дев'яти символів при наборі тексту. Графічне вікно застосунку наведено на рис. 4.3.

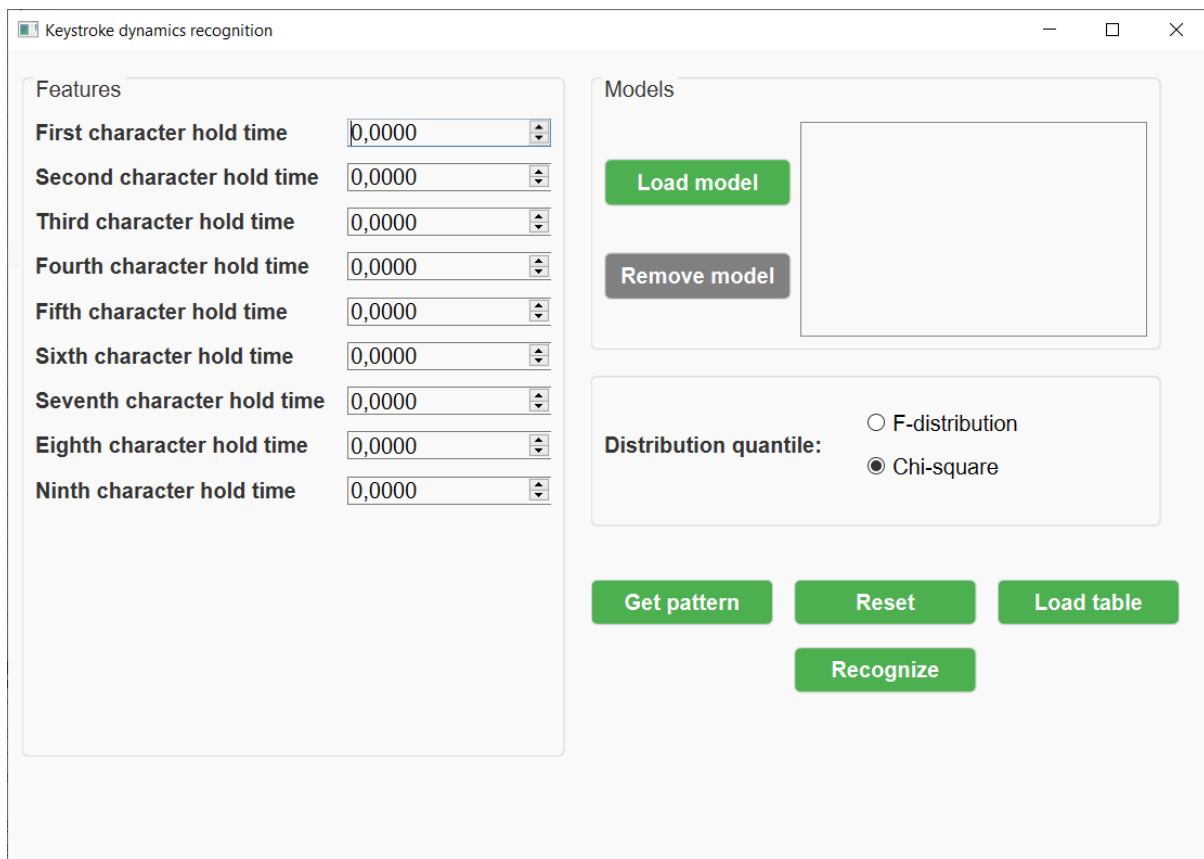


Рисунок 4.3 – Графічне вікно застосунку для розпізнавання клавіатурного почерку

Діаграма станів застосунку для розпізнавання клавіатурного почерку показана на рис. 4.4. Основним станом є відображення головного вікна, під час якого відбувається очікування дій користувача.

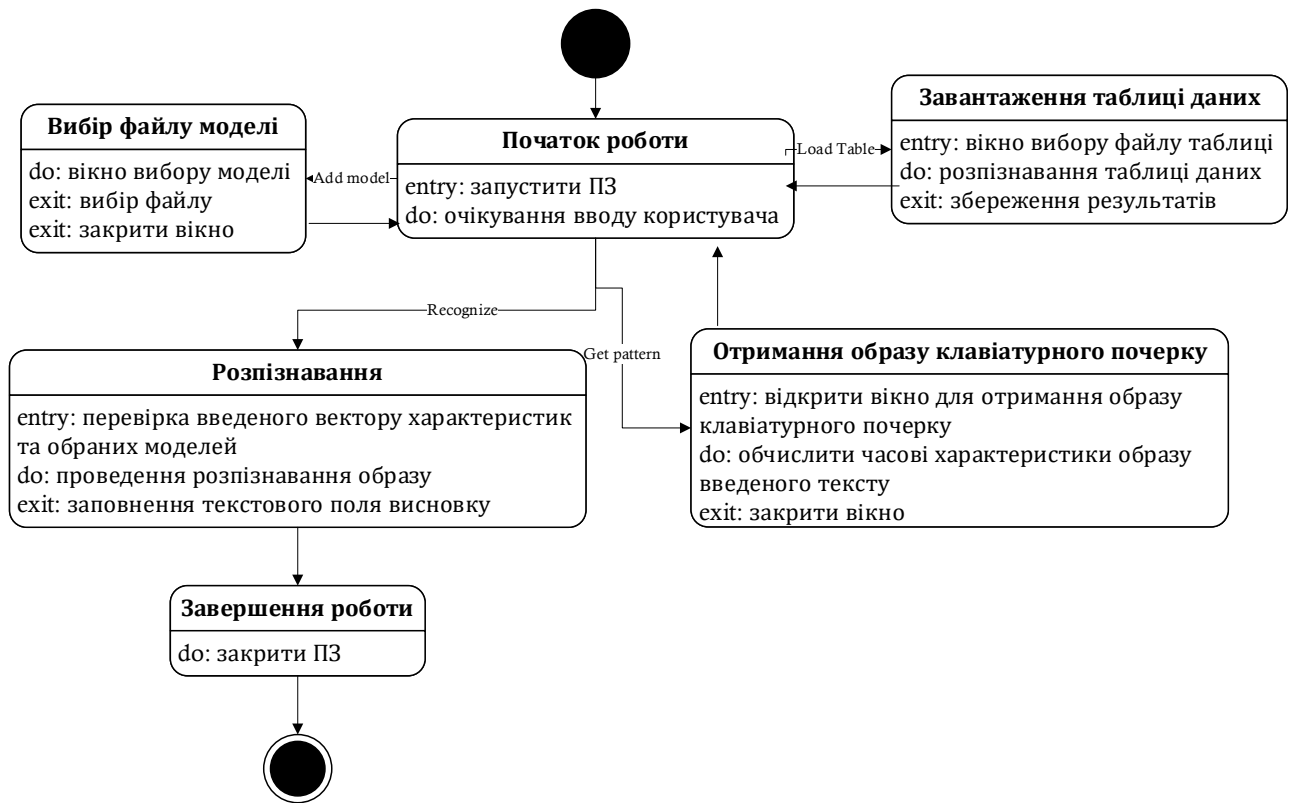


Рисунок 4.4 – Діаграма станів застосунку для розпізнавання клавіатурного почерку

Технічне забезпечення інформаційної технології. Для ефективної роботи застосунку потрібні такі мінімальні апаратні характеристики:

Процесор: процесор не нижче Pentium IV або еквівалентний.

Оперативна пам'ять: мінімум 128 МБ ОЗП (рекомендовано 1 ГБ для оптимальної продуктивності).

Місце на диску: щонайменше 100 МБ вільного місця для встановлення.

Ці характеристики забезпечують стабільну роботу застосунку на стандартному обладнанні.

Методичне забезпечення інформаційної технології. Методичним забезпеченням інформаційної технології для розпізнавання клавіатурного почерку є інструкція користувача для розпізнавання клавіатурного почерку з використанням відповідного програмного забезпечення – комп'ютерної програми, розробленої в рамках даної роботи.

Впровадження інформаційної технології. Інформаційну технологію для розпізнавання клавіатурного почерку було впроваджено в виробничі процеси ТОВ НВЦ «Діагностика та контроль». Впровадження запропонованої ІТ для розпізнавання клавіатурного почерку дозволило підвищити точність та ймовірність розпізнавання клавіатурного почерку. Також, математичне забезпечення та інструментарій інформаційної технології для розпізнавання клавіатурного почерку використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та управління проектами (ННІКНУП) Національного університету кораблебудування ім. адмірала Макарова наступних дисциплін: «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»; «Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки». Акти впровадження результатів дисертаційної роботи наведені в додатках А і Б.

4.3 Інструкція користувача для ІТ розпізнавання облич

Інструкція користувача для розпізнавання облич розрахована на використання відповідного програмного забезпечення – комп'ютерної програми з графічним інтерфейсом «Face recognition», розробленої в рамках даної роботи. Після запуску програми відкривається її графічне вікно, яке наведено на рис. 4.1.

В першу чергу необхідно завантажити файли моделі, які будуть використовуватися для розпізнавання. Для кожної завантаженої моделі буде побудовано еліпсоїд прогнозування для нормалізованих даних. Для цього, необхідно натиснути кнопку Load model, після чого відкриється вікно для вибору файлу, характерне для конкретної операційної системи, в даному випадку Windows (рис. 4.5). Важливо зазначити, що застосунок очікує файл з розширенням json.

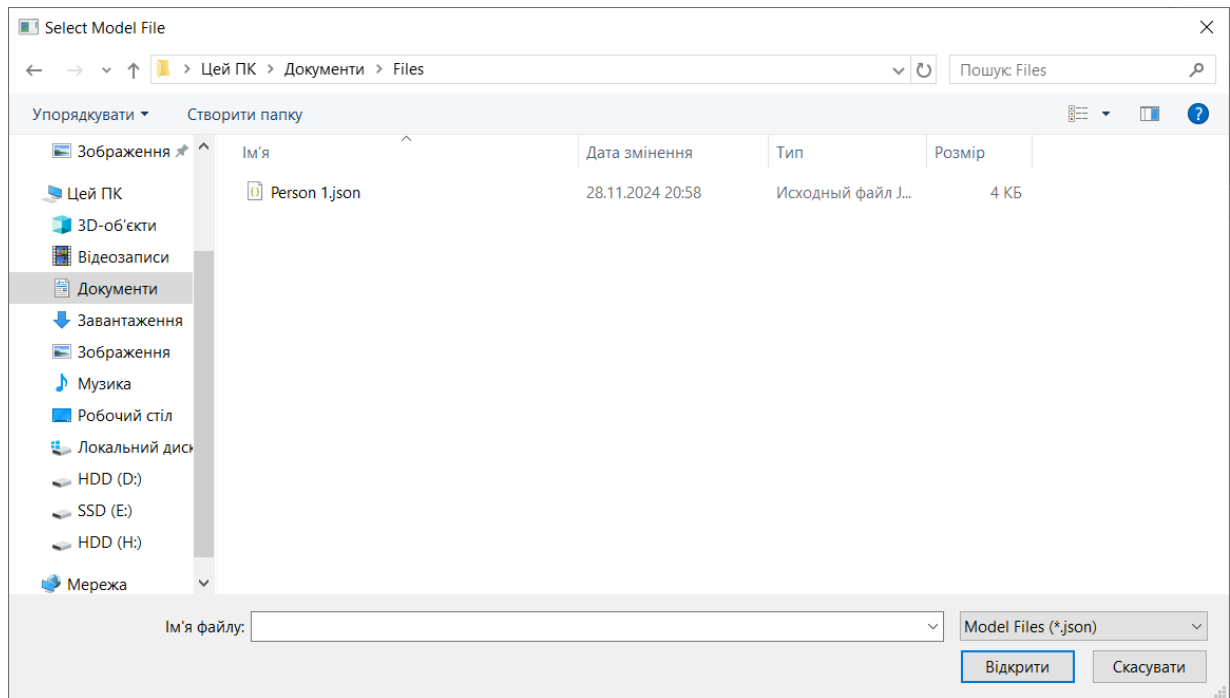


Рисунок 4.5 – Вікно для вибору файлу моделі

Файл моделі містить всі дані для побудови ймовірнісної моделі у вигляді еліпсоїда прогнозування для нормалізованих даних, і включає в себе одновимірні масиви з 10 елементів `TransformationParameters` і `MeansVector` та двовимірний масив розмірністю 10 на 10 `CovarianceMatrix`, де `TransformationParameters` - це оцінки параметрів перетворення Бокса-Кокса, `MeansVector` – вектор середніх нормалізованої навчальної вибірки, `CovarianceMatrix` – коваріаційна матриця нормалізованої навчальної вибірки.

У випадку, якщо файл не відповідає визначеній структурі, то з'явиться помилка `Invalid Format` (рис. 4.6).

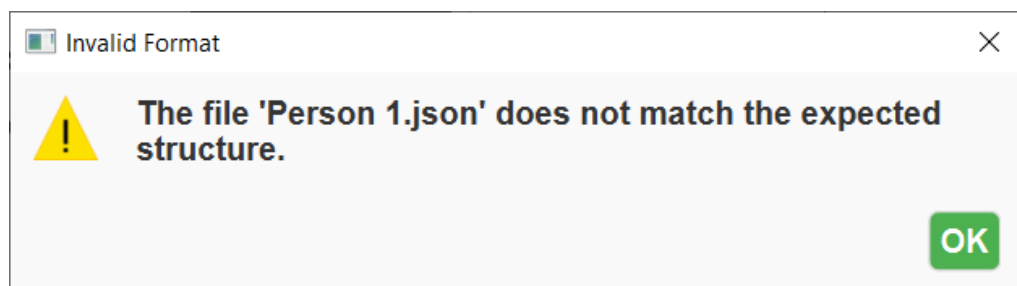


Рисунок 4.6 – Помилка Invalid Format

Застосунок підтримує одночасне додавання декількох моделей для розпізнавання (рис. 4.7).

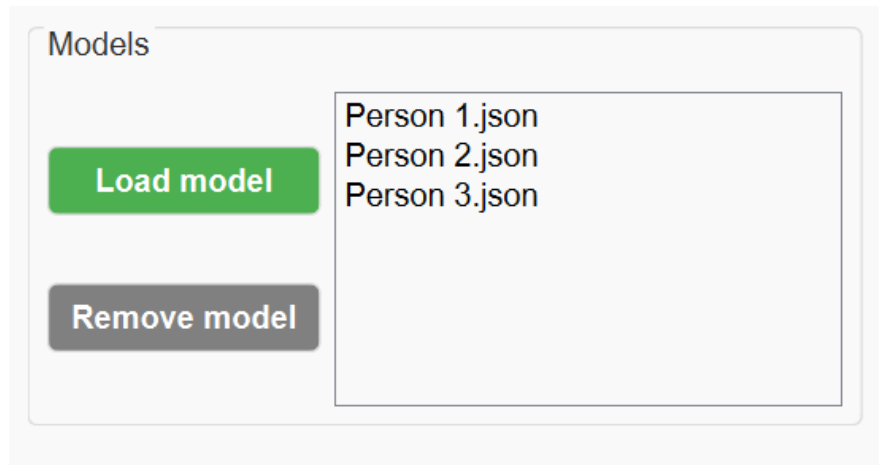


Рисунок 4.7 – Одночасне завантаження декількох моделей

У випадку додавання моделі з таким самим ім'ям з'явиться помилка Duplicate Entry.

Якщо необхідно видалити вже додану модель, її необхідно вибрати в списку. При цьому кнопка Remove model стане активною, а після її натискання виділена модель буде видалена зі списку і не використовуватиметься при розпізнаванні (рис. 4.8).

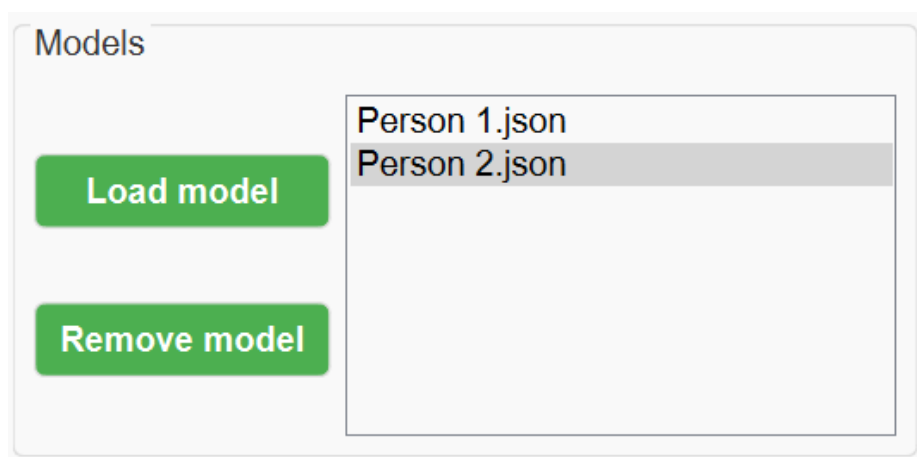


Рисунок 4.8 – Видалення файлу моделі

Наступним етапом є отримання вектору характеристик, який складається з 10 елементів: усереднена відстань від очей до середини носа; усереднена відстань від очей до центру рота; усереднена відстань від очей до центру брів;

усереднена відстань від брів до верху носа; усереднена відстань від куточків очей до верху носа; відстань між бровами; відстань між носом і серединою рота; ширина рота; ширина носу; відстань від рота до підборіддя. Кожна з відстаней ділиться на відстань між очима, щоб знизити вплив спотворень, викликаними відстанню обличчя від камери. Для цього є декілька варіантів, першим і найпростішим є завантаження зображення з пам'яті пристрою, що відбувається за допомогою кнопки Load image. При цьому відкривається вікно вибору файлу, специфічне для операційної системи з фільтром розширення файлів до найбільш популярних форматів зображення *.png *.jpg *.jpeg *.bmp (рис. 4.9).

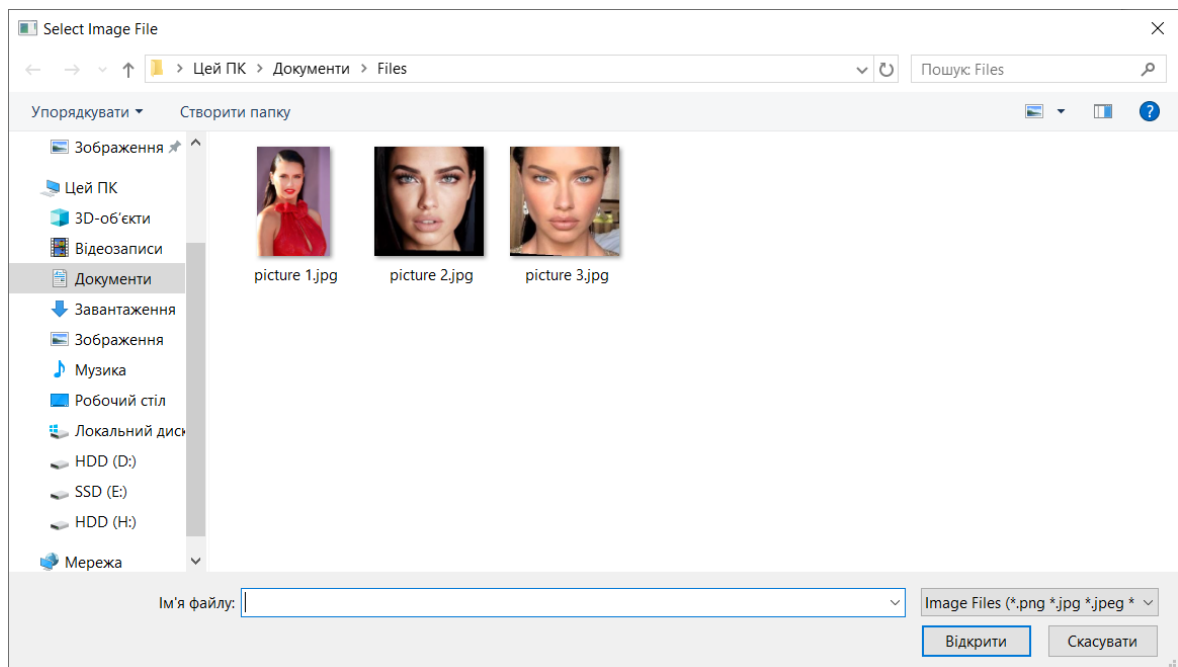


Рисунок 4.9 – Вікно вибору файлу для розпізнавання облич

Після вибору файлу, він відображається у головному вікні програми (рис. 4.10). За допомогою кнопки Extract Features відбувається пошук обличчя на зображенні. Воно відділяється від іншого зображення та вирівнюється, після чого відбувається розрахунок всіх необхідних відстаней. Таким чином, вектор характеристик заповнюється автоматично, а користувачу відображаються всі відстані, які беруть участь у розрахунку вектору характеристик, зеленим кольором, а відстань між очима – червоним (рис. 4.11).

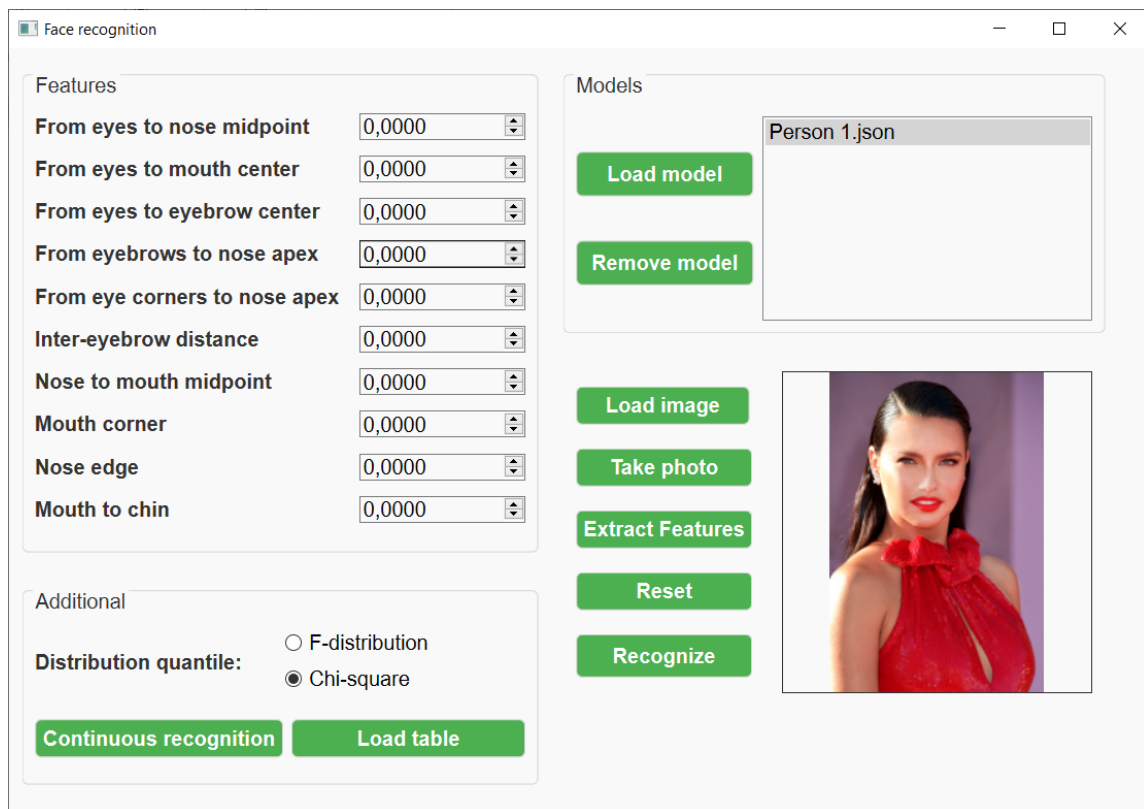


Рисунок 4.10 – Відображення вибраного файлу

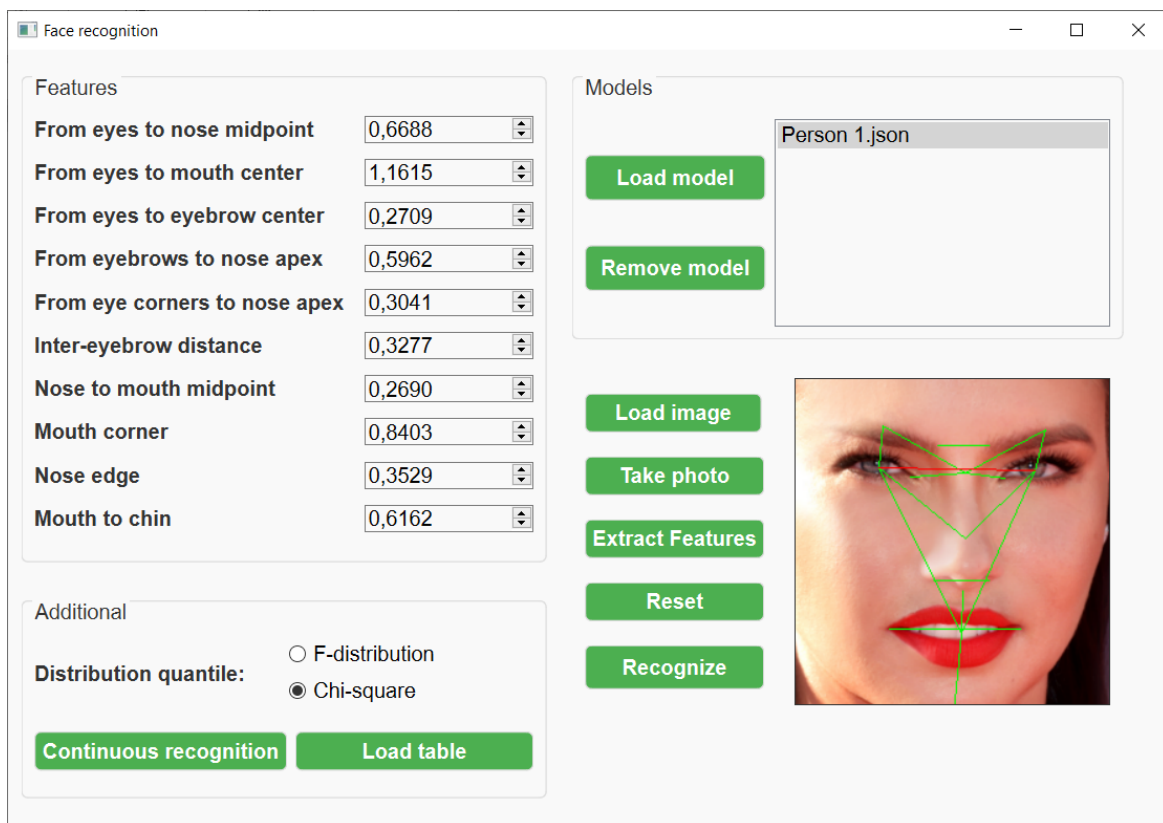


Рисунок 4.11 – Результат натискання кнопки Extract Features

Іншим, аналогічним способом, є отримання зображення з камери пристрою, на якому працює застосунок. Це відбувається за допомогою кнопки Take photo. При цьому з'являється окреме вікно, на якому відображається зображення з камери та кнопка для захоплення фотографії (рис. 4.12).

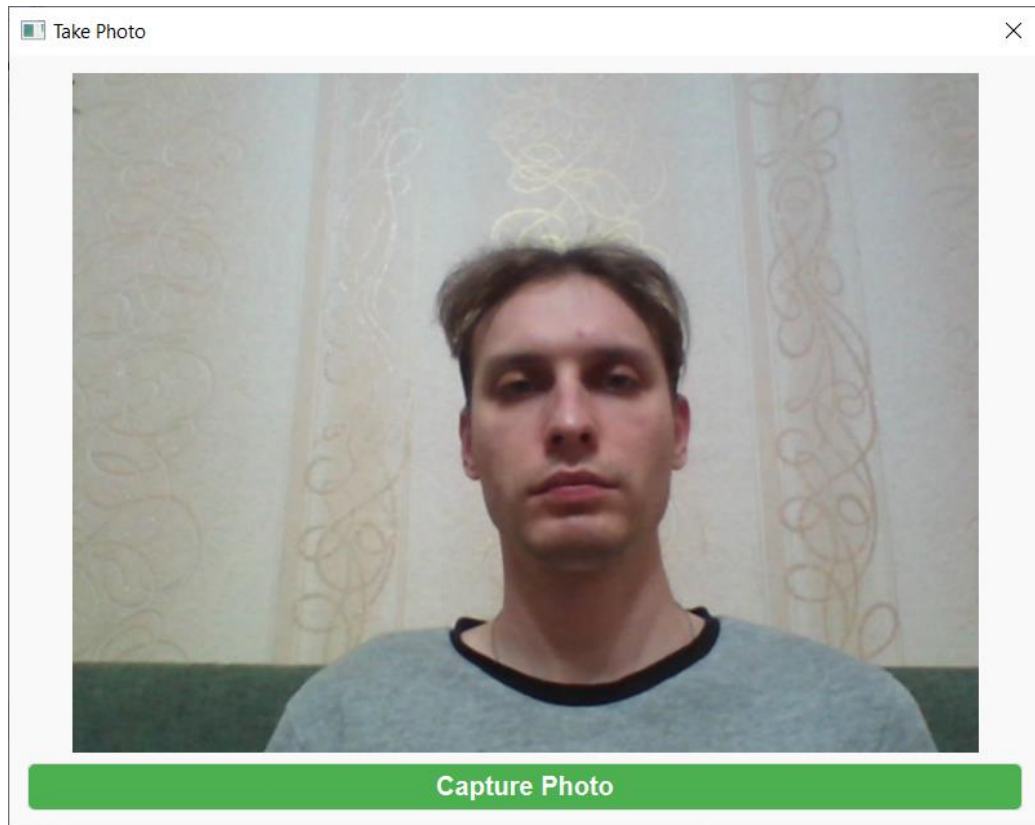


Рисунок 4.12 – Вікно Take Photo

Логіка роботи застосунку при натисканні Extract Features залишається незмінною (рис. 4.13), на фото знаходиться обличчя, яке відділяється від іншого зображення, та вирівнюється. Поверх відображаються вектори, які використовуються для подальшого обчислення відстаней та формування векторів характеристик для розпізнавання обличчя.

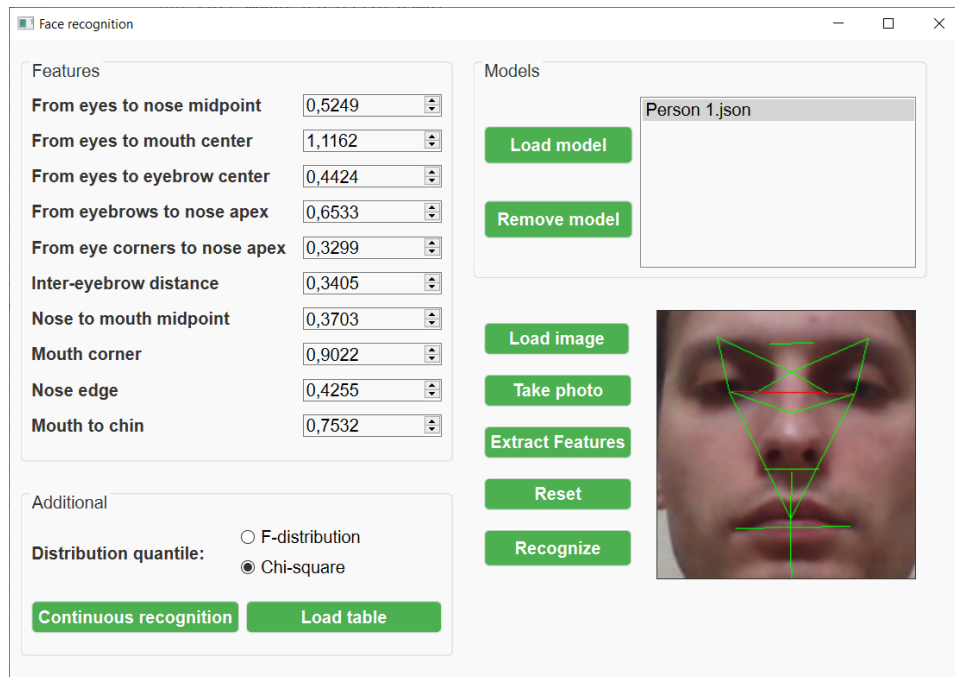


Рисунок 4.13 – Отримання відстаней між ключовими точками обличчя

У випадку, якщо зображення не містить жодного обличчя, або містить більше ніж 1 обличчя, виникає відповідна помилка (рис. 4.14).

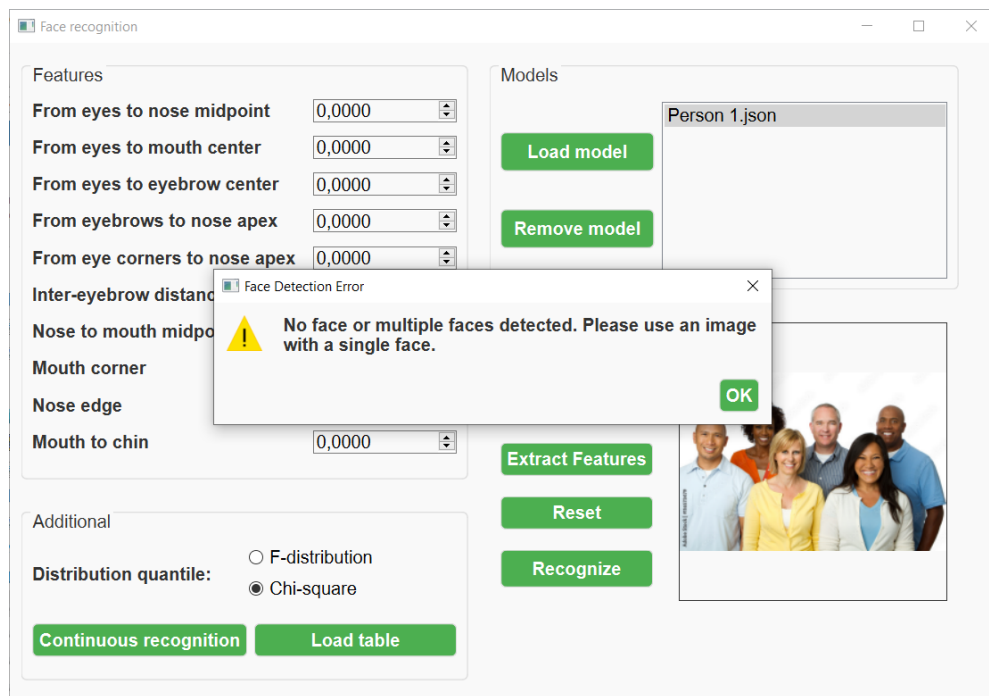


Рисунок 4.14 – Помилка при пошуку обличчя на зображенні

Якщо помилково натиснути кнопку Extract Features до того, як завантажити зображення чи зробити фото, виникає помилка (рис. 4.15).

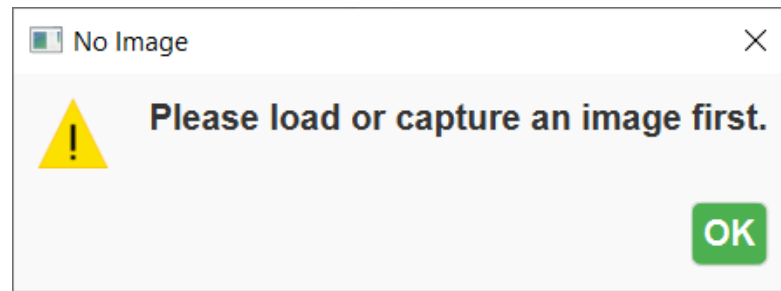


Рисунок 4.15 – Помилка No Image

У випадку, якщо жоден із запропонованих способів не підходить, то дані можуть бути введені вручну, хоча це не рекомендований варіант, так як велика ймовірність помилки при ручному вводі даних може призвести до не правильного розпізнавання обличчя.

Після отримання вектору характеристик, відбувається розпізнавання характеристик обличчя за допомогою кнопки Recognize. У випадку, якщо користувач завантажив декілька моделей, то відбувається поступове розпізнавання введеного користувачем вектору характеристик на приналежність до доданих моделей. У випадку знайденої відповідності з'являється повідомлення (рис. 4.16).

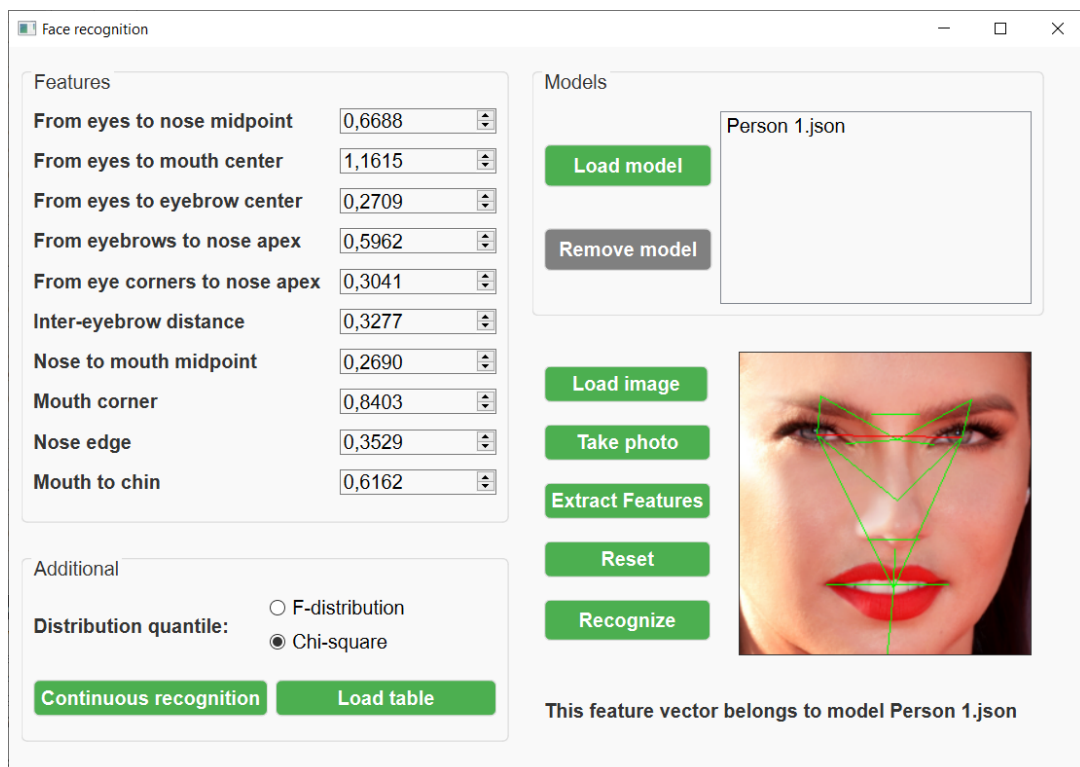


Рисунок 4.16 – Вивід результату роботи застосунку

Якщо введений вектор характеристик не належить до жодної з завантажених моделей, то виникає відповідне повідомлення (рис. 4.17). Це особливо корисно в ситуаціях, коли важливо не лише правильно ідентифікувати особу серед завчасно визначених моделей облич, а й надати інформацію про неналежність образу до жодного із завчасно визначеного списку класів, який визначається за допомогою додавання та видалення файлів моделей.

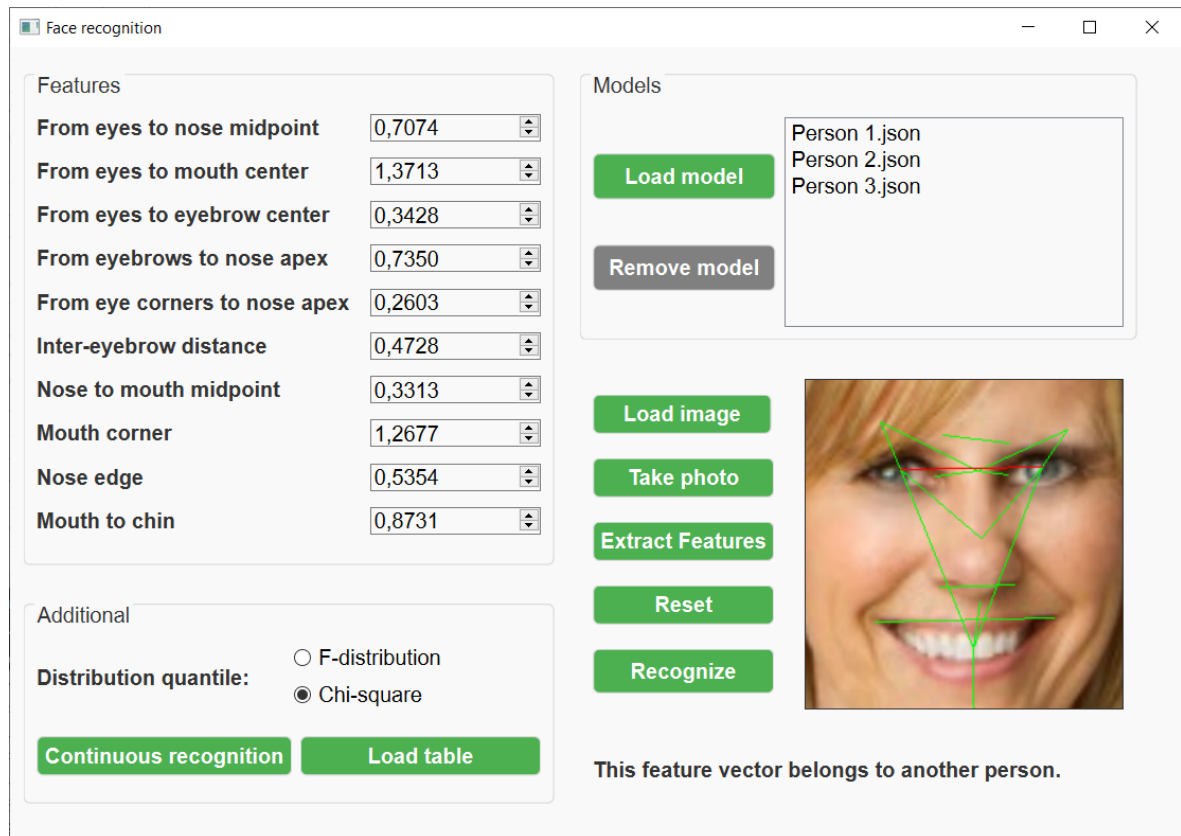


Рисунок 4.17 – Вивід результату роботи застосунку

Якщо ж користувач не завантажив модель, або не заповнив один з елементів вектору характеристик, то виникають відповідні помилки (рис. 4.18-4.19).

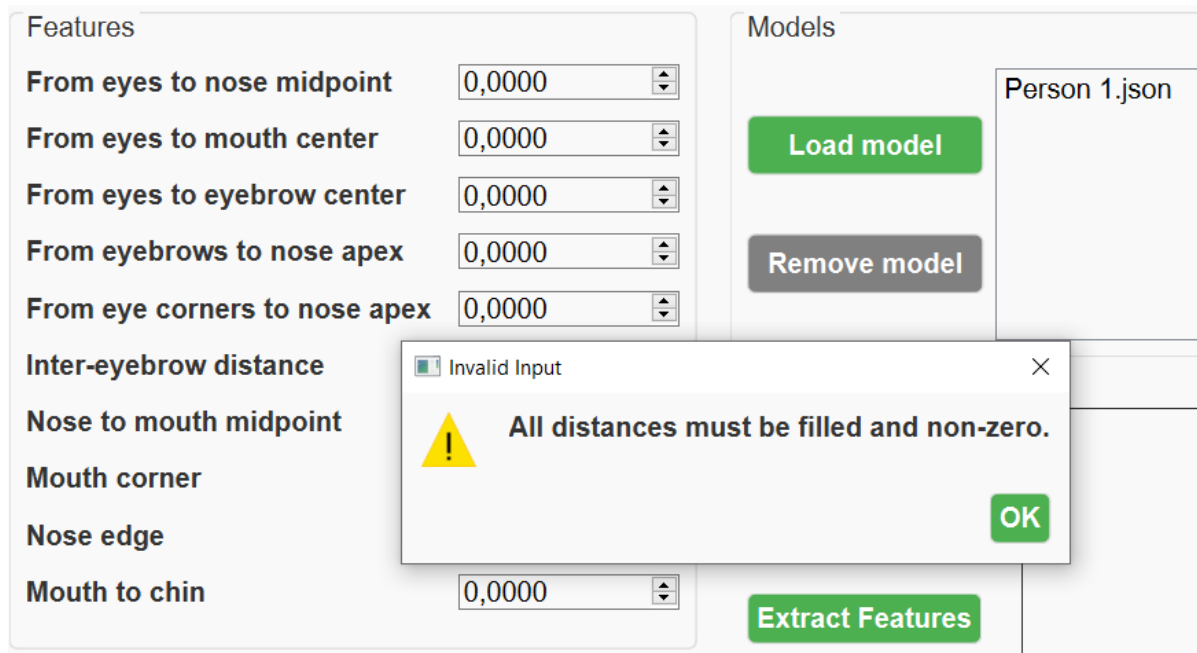


Рисунок 4.18 – Помилка Invalid Input

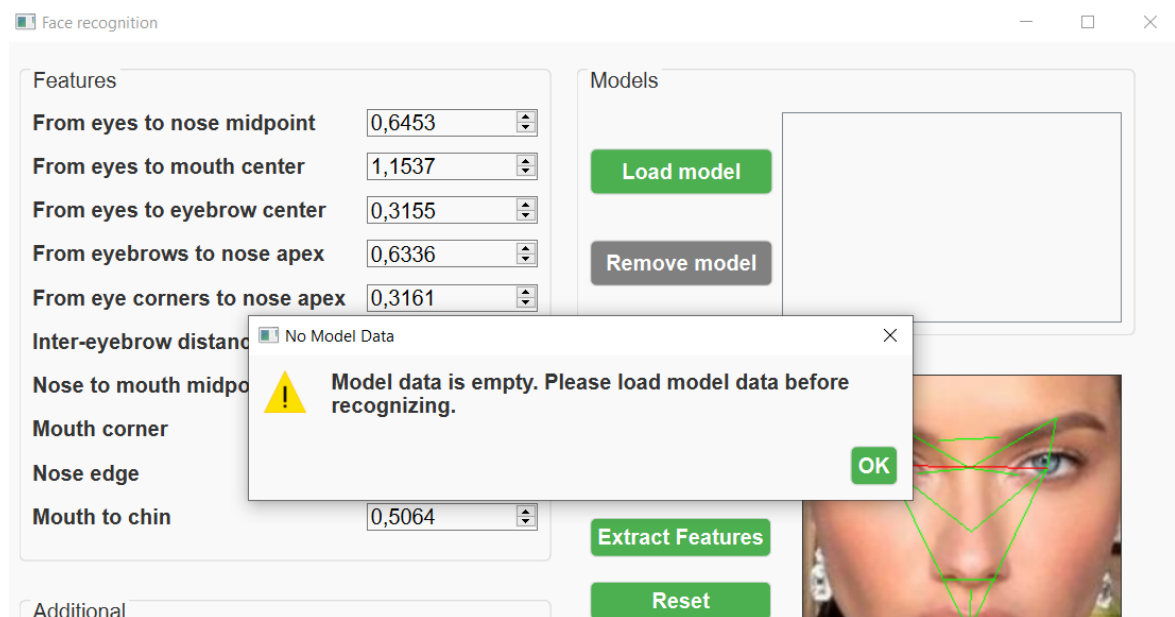


Рисунок 4.19 – Помилка No Model Data

Якщо треба виконати повторне розпізнавання чи просто очистити всі введені характеристики, зображення та висновок, можна використати кнопку **Reset**.

Додатково в групі **Additional** можна обрати необхідний квантиль розподілу. Так як було визначено, що для характеристик облич, використання квантиля

розподілу Хі-квадрат призводить до більш точного розпізнавання, то він обраний як значення за замовчуванням.

Проте це можна змінити: значення квантиля розподілу Хі-квадрат для 10 ступенів свободи та рівня значущості 0,005 становить 25,19, що менше ніж значення квантиля F -розподілу для рівня значущості 0,005, зі ступенями свободи 10 і 90 дорівнює 30,77.

Continuous recognition – це спеціальний режим роботи, який автоматично раз на секунду робить знімок, визначає вектор характеристик і проводить розпізнавання на основі завантажених моделей, що дуже корисно, наприклад, для організації доступу до деяких приміщень за фотографією людини. Цей процес відбувається безперервно, поки не буде розпізнано 1 фото.

Якщо необхідно виконати розпізнавання великої кількості даних, це можна зробити за допомогою кнопки Load table. При цьому відкривається вікно для завантаження файлів формату xls або csv. Ці файли мають містити десятивимірні вектори характеристик та заголовки колонок з назвою в форматі Distance{i}, де i змінюється від 1 до 10. При тому застосунок автоматично створить копію файлу, в якому буде додана колонка з назвою моделі, до якої належить кожен вектор та відкриється вікно для збереження файлу.

4.4 Інструкція користувача для ІТ розпізнавання клавіатурного почерку

Інструкція користувача для розпізнавання клавіатурного почерку розрахована на використання відповідного програмного забезпечення – комп'ютерної програми з графічним інтерфейсом «Keystroke dynamics recognition», розробленої в рамках даної роботи. Після запуску програми відкривається її графічне вікно, яке наведено на рис. 4.3.

Для початку необхідно завантажити файли моделей, які будуть використовуватися для розпізнавання. Для кожної завантаженої моделі створюється еліпсоїд прогнозування на основі нормалізованих даних. Для цього

слід натиснути кнопку Load model, після чого відкриється стандартне вікно вибору файлу, характерне для вашої операційної системи, у цьому випадку Windows (рис. 4.20). Важливо, що застосунок очікує файл із розширенням .json.

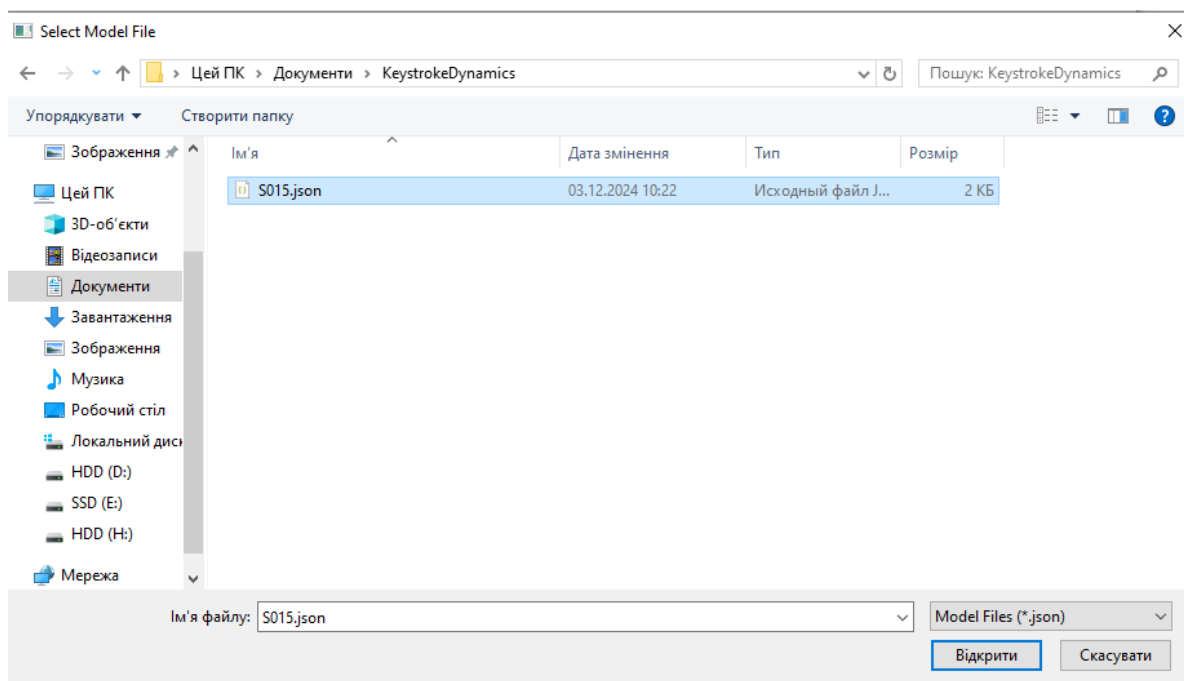


Рисунок 4.20 – Вікно вибору файлу моделі

Файл моделі містить усі необхідні дані для побудови ймовірнісної моделі у вигляді еліпсоїда прогнозування для нормалізованих даних. Зокрема, він включає:

- одновимірний масив з 9 елементів TransformationParameters, який містить оцінки параметрів перетворення Бокса-Кокса;
- одновимірний масив з 9 елементів MeansVector, який містить вектор середніх нормалізованої навчальної вибірки;
- двовимірний масив розмірністю 9x9 CovarianceMatrix, який містить коваріаційну матрицю нормалізованої навчальної вибірки.

Структура файлу моделі зображена на рис. 4.21.

```
{
  "TransformationParameters": [ "1,367629143", "1,48069292", "1,077970148", "1,739252703", "2,10044014", "1,149810219", "1,566069236", "1,168545425", "2,114572145" ],
  "CovarianceMatrix": [
    [ "2,88668E-05", "1,43542E-06", "1,97994E-05", "2,02509E-06", "1,26595E-07", "2,1049E-06", "3,9097E-06", "7,2795E-08", "-6,75351E-08" ],
    [ "1,43542E-06", "1,59957E-05", "-4,75432E-06", "-2,12275E-08", "7,09222E-08", "2,70712E-06", "4,31461E-06", "-1,6059E-06", "3,11449E-08" ],
    [ "1,97994E-05", "-4,75432E-06", "0,000174968", "-4,23071E-07", "4,24927E-07", "4,1652E-06", "-1,07337E-05", "5,12602E-06", "6,56328E-07" ],
    [ "2,02509E-06", "-2,12275E-08", "-4,23071E-07", "3,11192E-06", "1,46595E-07", "-1,4522E-07", "1,36181E-06", "-2,02308E-06", "-3,18062E-08" ],
    [ "1,26595E-07", "7,09222E-08", "4,24927E-07", "1,46595E-07", "3,34042E-07", "5,83113E-07", "7,57234E-09", "-1,68085E-07", "1,29062E-08" ],
    [ "2,1049E-06", "2,70712E-06", "4,1652E-06", "-1,4522E-07", "5,83113E-07", "5,52873E-05", "1,51291E-06", "9,20386E-07", "7,24285E-07" ],
    [ "3,9097E-06", "4,31461E-06", "-1,07337E-05", "1,36181E-06", "7,57234E-09", "1,51291E-06", "2,18655E-05", "-1,03199E-05", "1,21811E-08" ],
    [ "7,2795E-08", "-1,6059E-06", "5,12602E-06", "-2,02308E-06", "-1,68085E-07", "9,20386E-07", "-1,03199E-05", "0,00011571", "8,9427E-07" ],
    [ "-6,75351E-08", "3,11449E-08", "6,56328E-07", "-3,18062E-08", "1,29062E-08", "7,24285E-07", "1,21811E-08", "8,9427E-07", "5,64807E-07" ]
  ],
  "MeansVector": [ "-0,709323482", "-0,661840192", "-0,867639078", "-0,570160548", "-0,474273371", "-0,81641845", "-0,624166619", "-0,814426139", "-0,47084439" ]
}
```

Рисунок 4.21 – Структура файлу моделі

У випадку, якщо файл не відповідає визначеній структурі, то з'явиться помилка Invalid Format (рис. 4.22).

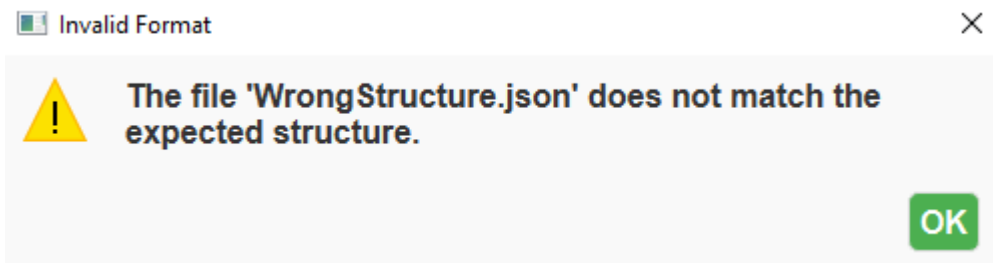


Рисунок 4.22 – Помилка Invalid Format

У випадку, якщо файл моделі представляє собою не правильно сформований json файл, то під час спроби відкрити його, з'явиться помилка, яка вказує на позицію, де виникли проблеми зі зчитуванням файлу (рис. 4.23).

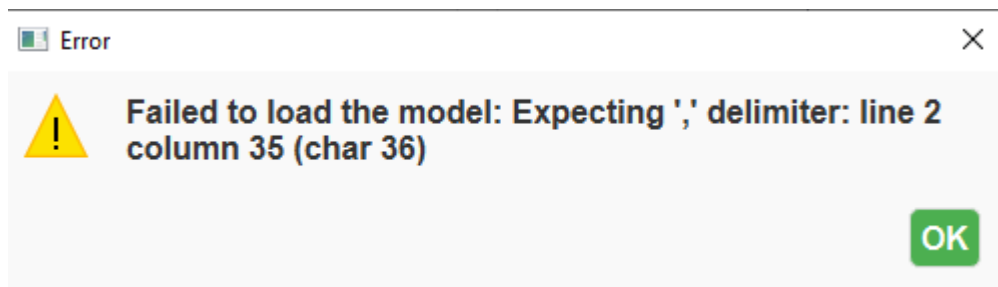


Рисунок 4.23 – Помилка при зчитуванні не правильного файлу моделі

Застосунок підтримує одночасне додавання декількох моделей для розпізнавання (рис. 4.24).

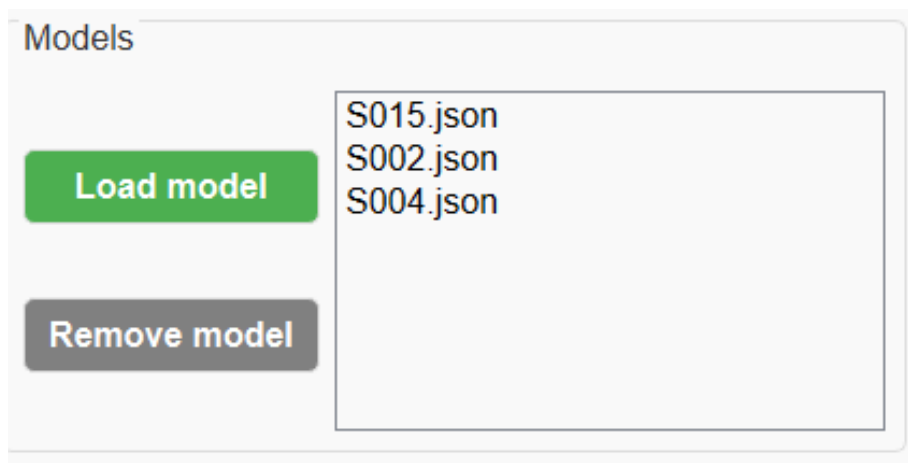


Рисунок 4.24 – Одночасне завантаження декількох моделей

Якщо необхідно видалити вже додану модель, її необхідно вибрати в списку. При цьому кнопка Remove model стане активною, а після її натискання виділена модель буде видалена зі списку і не використовуватиметься при розпізнаванні (рис. 4.25).

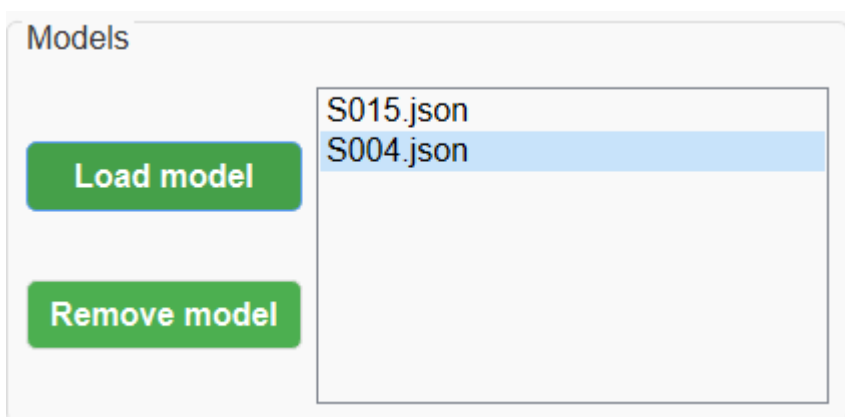


Рисунок 4.25 – Видалення файлу моделі

Наступним етапом є отримання вектору характеристик, який складається з 9 елементів, а саме часу утримання в секундах 9 символів при друку, в даному випадку «t», «i», «e», «5», «R», «o», «a», «n», «l» та натиснути кнопку Recognize.

Для початку розпізнавання є 3 основні способи, перший – це ввести вектор характеристик вручну. Проте, не завжди є можливість отримати ці дані, а також велика ймовірність допустити помилку при введенні, тому цей спосіб не є

рекомендованим, натомість пропонується використовувати автоматичне заповнення вектору характеристик.

Для автоматичного заповнення вектору характеристик, необхідно натиснути кнопку Get pattern, при цьому відкриється відповідне вікно (рис. 4.26).

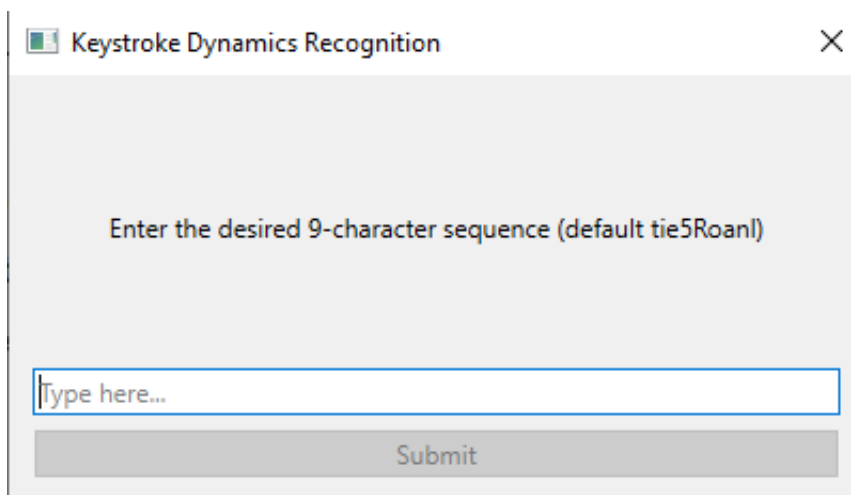


Рисунок 4.26 – Вікно Get pattern

Дане вікно вимірює час утримання символу при друці, час між натисканням клавіш – ігнорується, так як ці метрики не враховуються в розроблених моделях. При тому кнопка Submit залишається не активною, поки не буде введено потрібну кількість символів (рис. 4.27).

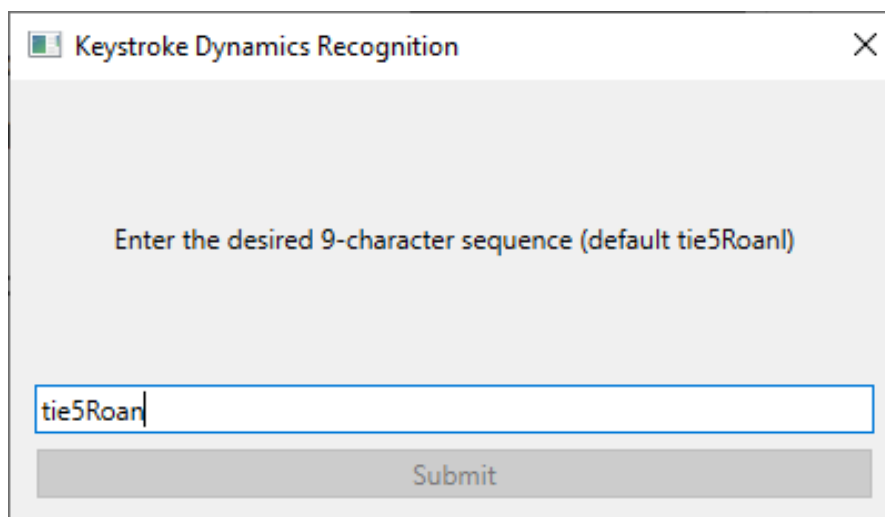


Рисунок 4.27 – Кнопка не активується при недостатній кількості введених символів

Після вводу всіх даних та заміру часових метрик, кнопка Submit активується (рис. 4.28) та дозволяє перейти до головного вікна з заповненими часовими характеристиками (рис. 4.29).

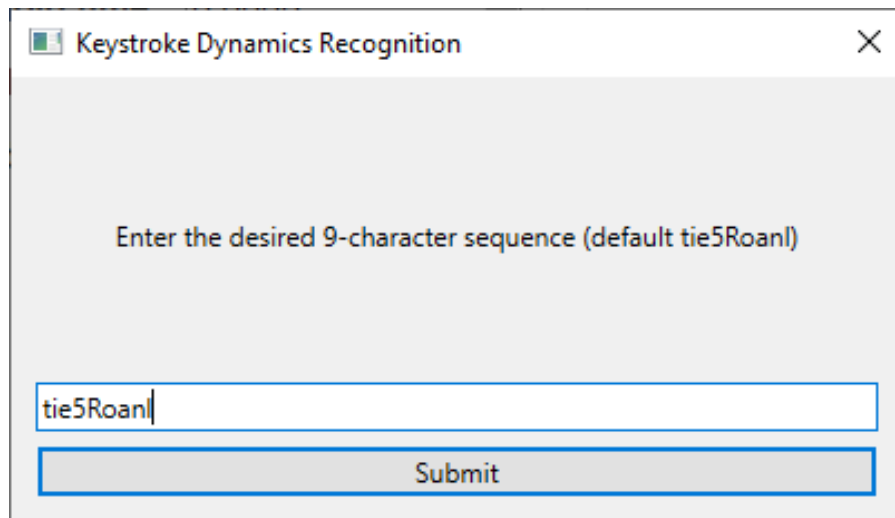


Рисунок 4.28 – Завершення зчитування часових метрик

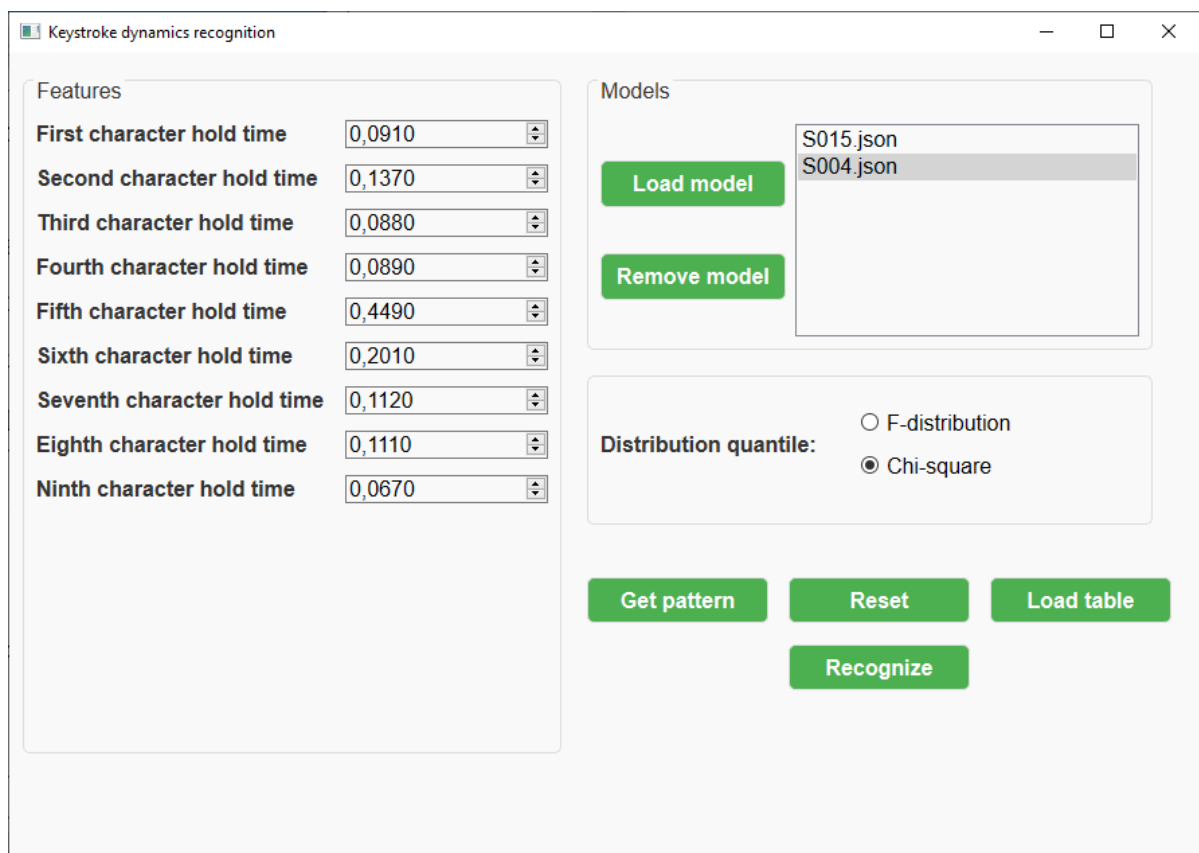


Рисунок 4.29 – Автоматичне заповнення часових метрик

Після отримання вектору характеристик, відбувається розпізнавання характеристик клавіатурного почерку за допомогою кнопки Recognize. У випадку, якщо користувач завантажив декілька моделей, то відбувається поступове розпізнавання введеного користувачем вектору характеристик на приналежність до доданих моделей. У випадку знайденої відповідності з'являється повідомлення (рис. 4.31).

Рисунок 4.31 – Вивід результату роботи застосунку

Якщо введений вектор характеристик не належить до жодної з завантажених моделей, то виникає відповідне повідомлення (рис. 4.32). Це особливо корисно в ситуаціях, коли важливо не лише правильно ідентифікувати особу серед завчасно визначених моделей клавіатурного почерку, а й надати інформацію про неналежність образу до жодного із завчасно визначеного списку класів, який визначається за допомогою додавання та видалення файлів моделей.

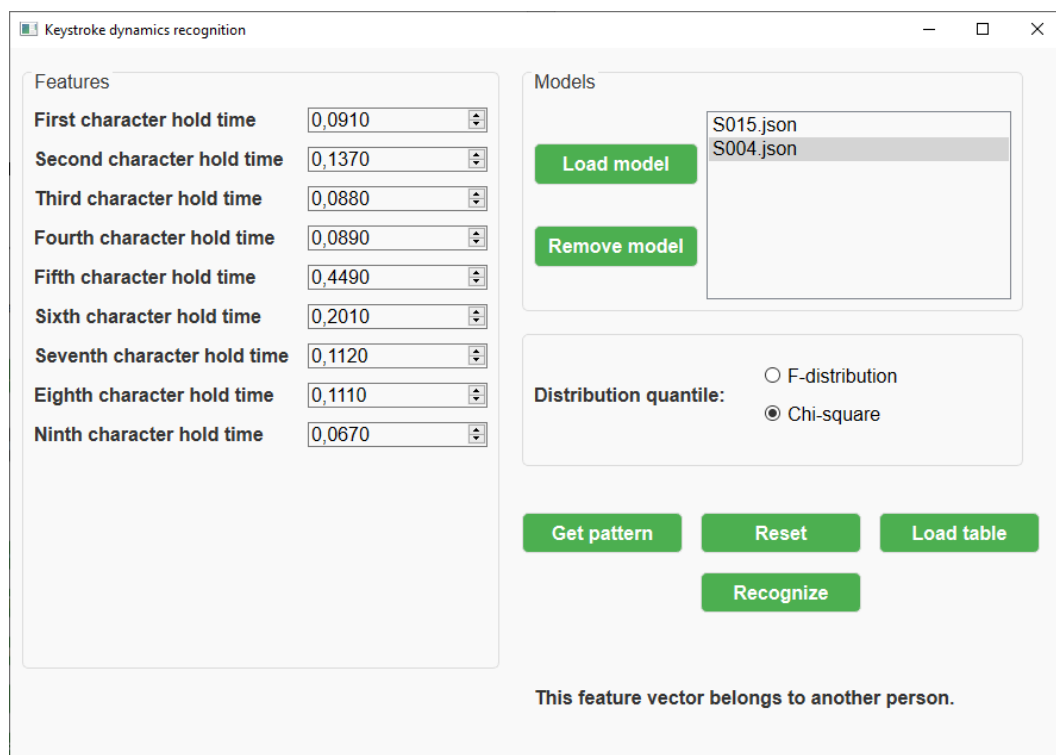


Рисунок 4.32 – Вивід результату роботи застосунку

Якщо ж користувач не завантажив модель, або не заповнив один з елементів вектору характеристик, то виникають відповідні помилки (рис. 4.33-4.34).

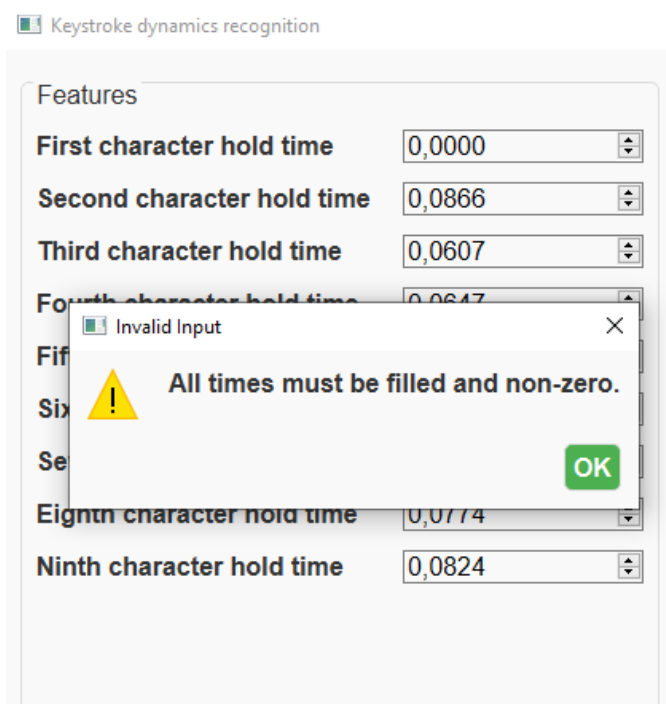


Рисунок 4.33 – Помилка Invalid Input

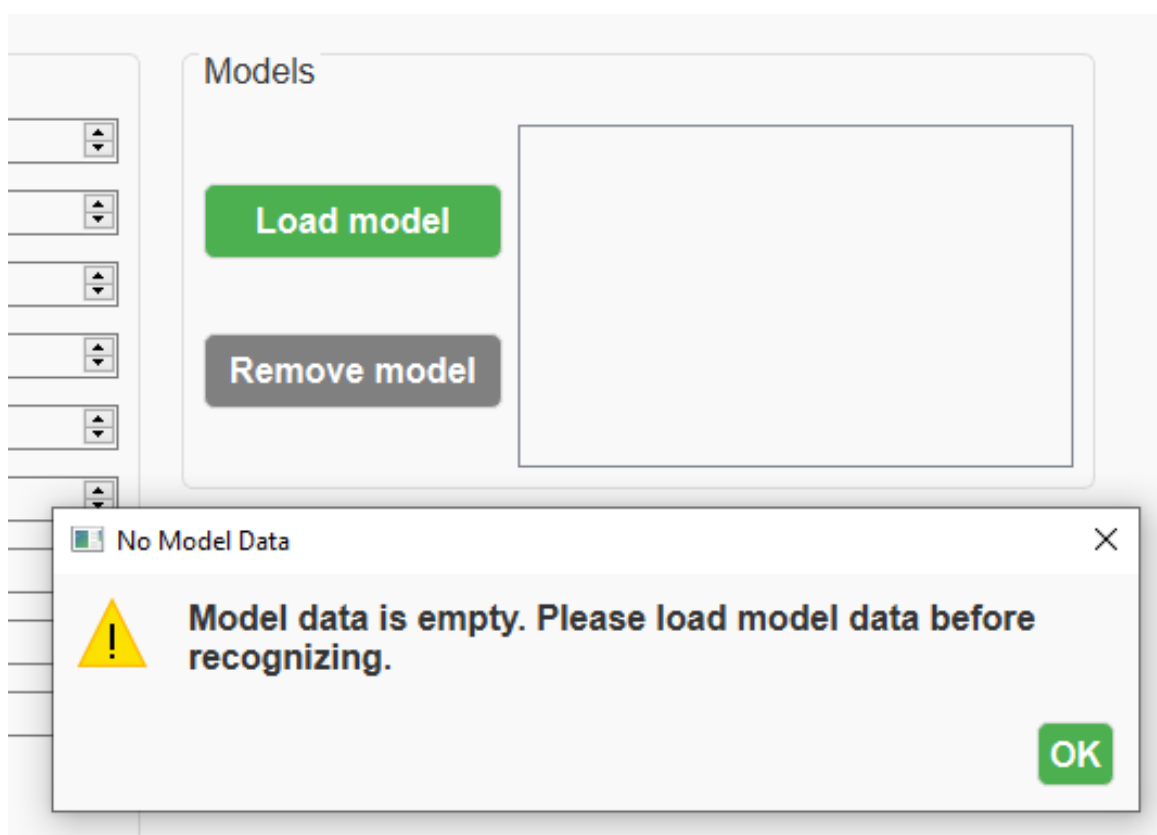


Рисунок 4.34 – Помилка No Model Data

Якщо треба виконати повторне розпізнавання чи просто очистити всі введені характеристики та висновок, можна використати кнопку Reset.

Додатково перед розпізнаванням можна обрати необхідний квантиль розподілу (рис. 4.35). Так як було визначено, що для характеристик клавіатурного почерку, використання квантиля розподілу Хі-квадрат призводить до більш точного розпізнавання, то він обраний як значення за замовчуванням.

Проте це можна змінити; значення квантиля розподілу Хі-квадрат для 9 ступенів свободи та рівня значущості 0,005 становить 23,59, що менше ніж значення квантиля F -розподілу для рівня значущості 0,005, зі ступенями свободи 9 і 186 дорівнює 25,86.

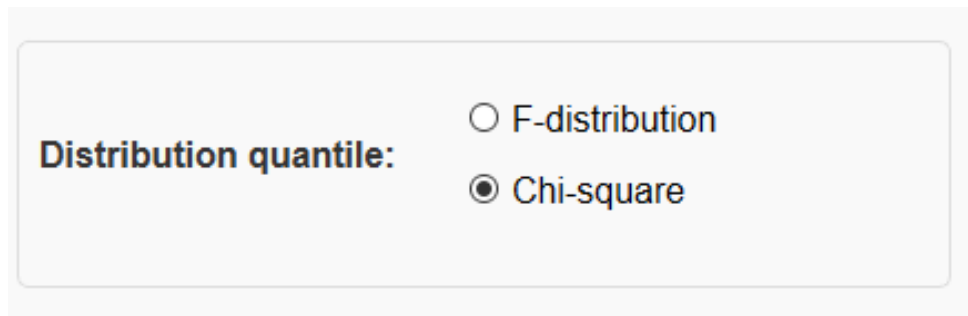


Рисунок 4.35 – Група Additional

Якщо необхідно виконати розпізнавання великої кількості даних, це можна зробити за допомогою кнопки Load table. В такому випадку необхідно завантажити моделі, у випадку їх відсутності виникне помилка (рис. 4.36).

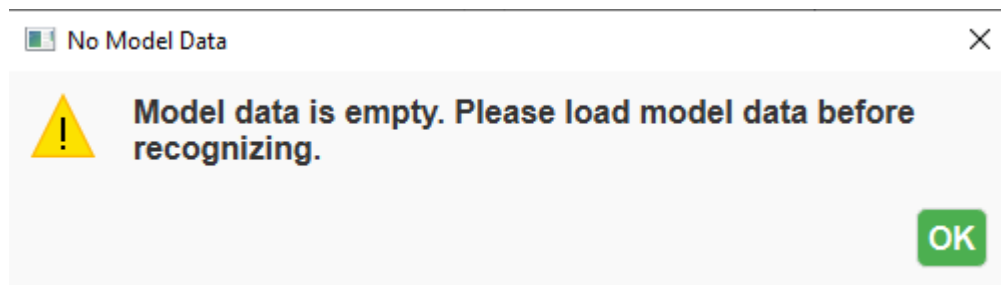


Рисунок 4.36 – Помилка No Model Data

При цьому відкривається вікно для завантаження файлів формату xls або csv (рис. 4.37).

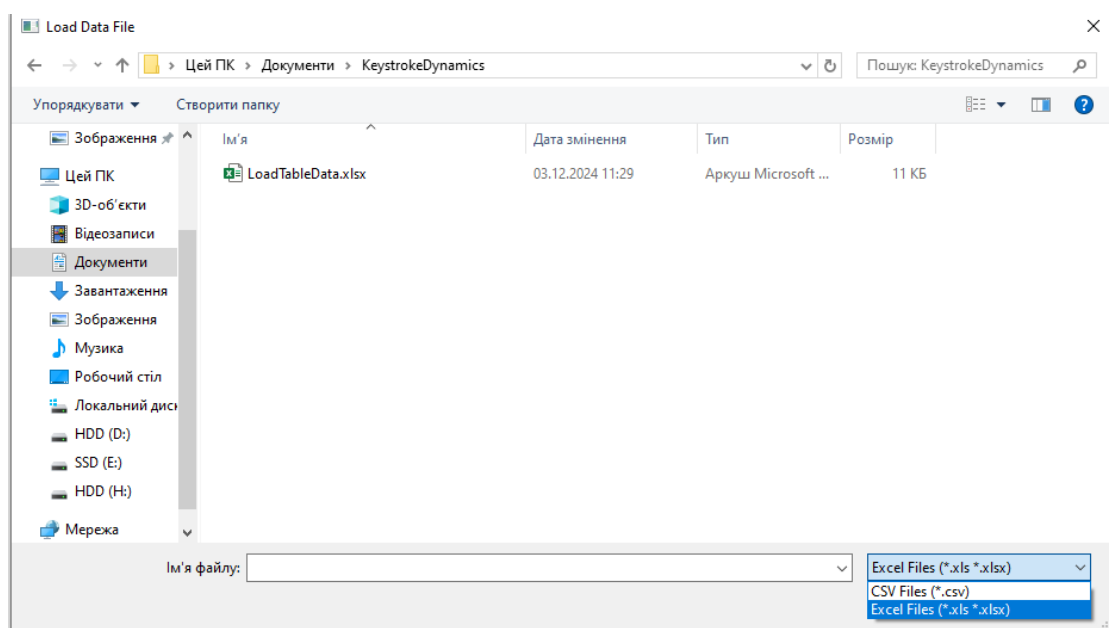


Рисунок 4.37 – Завантаження таблиці даних

Ці файли мають містити дев'ятивимірні вектори характеристик та заголовки колонок з назвою в форматі Time{i}, де i змінюється від 1 до 9 (рис. 4.38). В тестовому файлі об'єкти, що належать до класу з ідентифікатором S015, були заповнені в рядки з 2 по 12. З 13 по 22 – знаходяться представники інших класів.

	A	B	C	D	E	F	G	H	I
1	Time1	Time2	Time3	Time4	Time5	Time6	Time7	Time8	Time9
2	0,0547	0,0694	0,0657	0,0576	0,0681	0,0839	0,0824	0,0602	0,0684
3	0,0518	0,0359	0,08	0,0747	0,0756	0,0784	0,0888	0,0817	0,0628
4	0,065	0,079	0,0713	0,0061	0,0592	0,0958	0,1022	0,0846	0,0657
5	0,0689	0,0766	0,0681	0,0776	0,0668	0,1029	0,1022	0,0848	0,0805
6	0,0504	0,0887	0,065	0,036	0,081	0,0971	0,1088	0,0489	0,0681
7	0,1127	0,0671	0,0916	0,0715	0,0824	0,0882	0,1069	0,0655	0,0652
8	0,08	0,0758	0,0785	0,0512	0,0471	0,0948	0,0961	0,0747	0,0676
9	0,0528	0,0858	0,0813	0,047	0,061	0,0979	0,1016	0,0924	0,089
10	0,0744	0,0744	0,0589	0,0879	0,0755	0,0766	0,0927	0,038	0,0528
11	0,0383	0,0715	0,0351	0,0562	0,0729	0,084	0,0719	0,0869	0,0553
12	0,0892	0,0618	0,0618	0,0644	0,0732	0,1154	0,1115	0,0368	0,0908
13	0,1119	0,1452	0,085	0,062	0,1278	0,095	0,0755	0,0911	0,0964
14	0,0626	0,0998	0,0919	0,0911	0,1067	0,1003	0,1228	0,089	0,1082
15	0,0977	0,1132	0,0837	0,089	0,1341	0,0985	0,1183	0,0816	0,0795
16	0,0721	0,109	0,0657	0,0723	0,1333	0,1095	0,1156	0,0934	0,0974
17	0,0657	0,0971	0,0692	0,085	0,1346	0,094	0,1053	0,0937	0,0942
18	0,051	0,1003	0,0787	0,0824	0,1402	0,0958	0,1135	0,0826	0,0905
19	0,0805	0,1088	0,0818	0,0647	0,1276	0,0879	0,1349	0,1014	0,109
20	0,0837	0,0985	0,0858	0,0647	0,1317	0,0824	0,1206	0,0816	0,0927
21	0,085	0,1077	0,0824	0,0623	0,1586	0,0884	0,1299	0,0948	0,0971
22	0,0895	0,1	0,0813	0,0729	0,1185	0,1077	0,1143	0,0847	0,095

Рисунок 4.38 – Таблиці даних для розпізнавання

У випадку, якщо колонки таблиці не відповідають вимогами, виникне помилка (рис. 4.39).

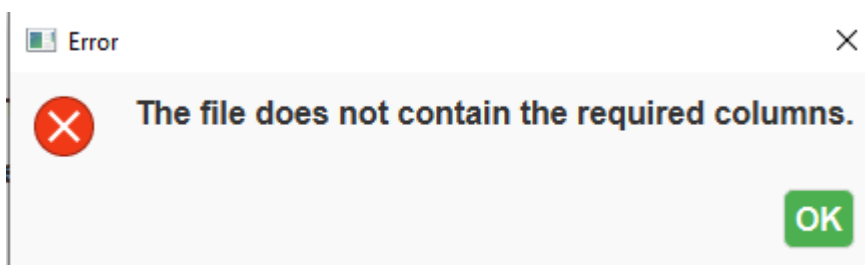


Рисунок 4.39 – Помилка обробки таблиці не вірного формату

Після вибору таблиці, застосунок автоматично проведе розпізнавання векторів характеристик клавіатурного почерку, створить копію та відкриє вікно для збереження отриманого файлу (рис. 4.40). При тому застосунок покаже повідомлення про успішне завершення розпізнавання (рис. 4.41).

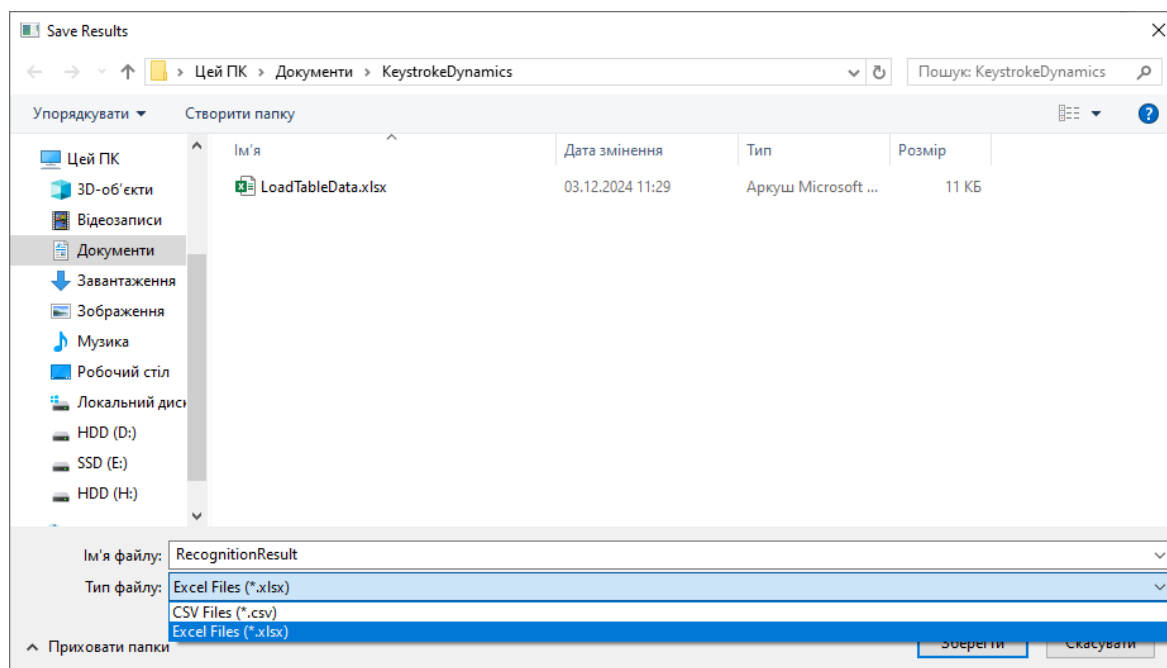


Рисунок 4.40 – Вікно збереження таблиці після розпізнавання



Рисунок 4.41 – Успішне завершення розпізнавання

В файлі з результатами розпізнавання буде додана колонка, значення якої заповнюється назвою моделі, до якої належить кожен вектор, у випадку, якщо об'єкт не розпізнано жодною з доданих моделей, в колонці Model буде стояти значення None (рис. 4.42).

	A	B	C	D	E	F	G	H	I	J
1	Time1	Time2	Time3	Time4	Time5	Time6	Time7	Time8	Time9	Model
2	0,0547	0,0694	0,0657	0,0576	0,0681	0,0839	0,0824	0,0602	0,0684	S015.json
3	0,0518	0,0359	0,08	0,0747	0,0756	0,0784	0,0888	0,0817	0,0628	S015.json
4	0,065	0,079	0,0713	0,0061	0,0592	0,0958	0,1022	0,0846	0,0657	S015.json
5	0,0689	0,0766	0,0681	0,0776	0,0668	0,1029	0,1022	0,0848	0,0805	S015.json
6	0,0504	0,0887	0,065	0,036	0,081	0,0971	0,1088	0,0489	0,0681	S015.json
7	0,1127	0,0671	0,0916	0,0715	0,0824	0,0882	0,1069	0,0655	0,0652	S015.json
8	0,08	0,0758	0,0785	0,0512	0,0471	0,0948	0,0961	0,0747	0,0676	S015.json
9	0,0528	0,0858	0,0813	0,047	0,061	0,0979	0,1016	0,0924	0,089	S015.json
10	0,0744	0,0744	0,0589	0,0879	0,0755	0,0766	0,0927	0,038	0,0528	S015.json
11	0,0383	0,0715	0,0351	0,0562	0,0729	0,084	0,0719	0,0869	0,0553	S015.json
12	0,0892	0,0618	0,0618	0,0644	0,0732	0,1154	0,1115	0,0368	0,0908	S015.json
13	0,1119	0,1452	0,085	0,062	0,1278	0,095	0,0755	0,0911	0,0964	None
14	0,0626	0,0998	0,0919	0,0911	0,1067	0,1003	0,1228	0,089	0,1082	None
15	0,0977	0,1132	0,0837	0,089	0,1341	0,0985	0,1183	0,0816	0,0795	None
16	0,0721	0,109	0,0657	0,0723	0,1333	0,1095	0,1156	0,0934	0,0974	None
17	0,0657	0,0971	0,0692	0,085	0,1346	0,094	0,1053	0,0937	0,0942	None
18	0,051	0,1003	0,0787	0,0824	0,1402	0,0958	0,1135	0,0826	0,0905	None
19	0,0805	0,1088	0,0818	0,0647	0,1276	0,0879	0,1349	0,1014	0,109	None
20	0,0837	0,0985	0,0858	0,0647	0,1317	0,0824	0,1206	0,0816	0,0927	None
21	0,085	0,1077	0,0824	0,0623	0,1586	0,0884	0,1299	0,0948	0,0971	None
22	0,0895	0,1	0,0813	0,0729	0,1185	0,1077	0,1143	0,0847	0,095	None

Рисунок 4.42 – Результати розпізнавання таблиці даних

4.5 Висновки за розділом 4

В четвертому розділі було створено ІТ для розпізнавання облич за десятивимірним вектором характеристик із застосуванням ймовірнісної моделі побудованої з використанням десятивимірного перетворення Бокса-Кокса. Також було створено ІТ для розпізнавання клавіатурного почерку за дев'ятивимірним вектором характеристик із застосуванням ймовірнісної моделі побудованої з використанням дев'ятивимірного перетворення Бокса-Кокса.

1) Створена інформаційна технологія розпізнавання облич за десятивимірним вектором характеристик із застосуванням ймовірнісної моделі побудованої з використанням десятивимірного перетворення Бокса-Кокса дозволяє підвищити точність та ймовірність розпізнавання облич.

2) Розроблено інструкцію користувача для розпізнавання облич, яка розрахована на використання відповідного інструментарію ІТ – програмного забезпечення, яке було створено.

3) Створена інформаційна технологія розпізнавання клавіатурного почерку за дев'ятивимірним вектором характеристик із застосуванням ймовірнісної моделі побудованої з використанням дев'ятивимірного перетворення Бокса-Кокса дозволяє підвищити точність та ймовірність розпізнавання клавіатурного почерку.

4) Розроблено інструкцію користувача для розпізнавання клавіатурного почерку, яка розрахована на використання відповідного інструментарію ІТ – програмного забезпечення, яке було створено.

5) Розроблену інформаційну технологію для розпізнавання облич було впроваджено в виробничі процеси ТОВ НВЦ «Діагностика та контроль». Впровадження запропонованої ІТ для розпізнавання облич дозволило підвищити точність та ймовірність розпізнавання облич. Також, математичне забезпечення та інструментарій інформаційної технології для розпізнавання облич використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та управління проектами (ННІКНУП) Національного університету кораблебудування ім. адмірала Макарова наступних дисциплін: «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»; «Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки». Акти впровадження результатів дисертаційної роботи наведені в додатках А і Б.

6) Інформаційну технологію для розпізнавання клавіатурного почерку було впроваджено в виробничі процеси ТОВ НВЦ «Діагностика та контроль». Впровадження запропонованої ІТ для розпізнавання клавіатурного почерку дозволило підвищити точність та ймовірність розпізнавання клавіатурного почерку. Також, математичне забезпечення та інструментарій інформаційної технології для розпізнавання клавіатурного почерку використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та

управління проектами (ННІКНУП) Національного університету кораблебудування ім. адмірала Макарова наступних дисциплін: «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»; «Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки». Акти впровадження результатів дисертаційної роботи наведені в додатках А і Б.

ВИСНОВКИ

В дисертаційній роботі вирішено важливе науково-практичне завдання підвищення точності та ймовірності розпізнавання образів, а саме облич та клавіатурного почерку, за рахунок використання відповідних ймовірнісних моделей у вигляді еліпсоїдів прогнозування, що побудовані з використанням багатовимірною нормалізуючого перетворення Бокса-Кокса, та на їх основі створено інструментарії інформаційних технологій для розпізнавання облич та клавіатурного почерку за відповідними векторами характеристик.

1) Проаналізовано існуючі методи і моделі для розпізнавання образів. Визначено, що існуючі методи в основному базуються на припущенні, що дані мають багатовимірний нормальний розподіл, однак, це не завжди виконується в реальному світі, що в результаті, призводить до зниження точності та ймовірності розпізнавання образів для негаусівських даних. Тому виникає потреба у розробці відповідних ймовірнісних моделей для розпізнавання образів, в тому числі облич і клавіатурного почерку, що враховують відхилення розподілу даних від нормального.

2) Одним з ефективних способів подолання обмежень, пов'язаних з розподілом даних, є використання нормалізуючих перетворень, які допомагають коригувати розподіл даних, наближаючи його до нормального. Одновимірні перетворення застосовуються до кожної ознаки окремо та фокусуються на індивідуальних характеристиках, наближаючи їх до нормального розподілу, проте не завжди призводять до задовільної нормалізації даних. Багатовимірні нормалізуючі перетворення є більш доцільними, тому що на відміну від одновимірних, враховують взаємозв'язки між ознаками. Десятивимірне перетворення Бокса-Кокса було обрано у якості нормалізуючого перетворення для нормалізації десятидимірних характеристик облич. У порівнянні з одновимірними перетвореннями десятичного логарифму та Бокса-Кокса, це

перетворення дозволило краще нормалізувати десятивимірні характеристики облич з врахуванням кореляції між ознаками. Дев'ятивимірне перетворення Бокса-Кокса було обрано у якості нормалізуючого перетворення для нормалізації дев'ятивимірних характеристик клавіатурного почерку. У порівнянні з одновимірними перетвореннями десяткового логарифму та Бокса-Кокса, це перетворення дозволило краще нормалізувати дев'ятивимірні характеристики клавіатурного почерку з врахуванням кореляції між ознаками.

3) Вперше побудована ймовірнісна модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню десятивимірного перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу. Використання побудованої ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9533 до 0,9567, ймовірність розпізнавання цільового образу (Specificity) з 0,99 до 1, при тому ймовірність розпізнавання представників інших класів (Recall) залишилася 0,9350.

4) Вперше побудована ймовірнісна модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем F -розподілу для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірного перетворення Бокса-Кокса та кількість точок даних завдяки використанню квантиля F -розподілу. Використання побудованої ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9328 до 0,9765, ймовірність розпізнавання цільового образу (Specificity) з 0,9949 до 1 та ймовірність розпізнавання представників інших класів (Recall) з 0,9025 до 0,9650.

5) Удосконалено ймовірнісну модель для розпізнавання облич у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню

десятивимірною перетворення Бокса-Кокса. Використання створеної ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9433 до 0,9733, ймовірність розпізнавання цільового образу (Specificity) з 0,9200 до 0,9700 та ймовірність розпізнавання представників інших класів (Recall) з 0,9550 до 0,9750.

6) Удосконалено ймовірнісну модель для розпізнавання клавіатурного почерку у вигляді еліпсоїда прогнозування з квантилем χ^2 -квадрат для нормалізованого вектору характеристик, яка на відміну від існуючих моделей враховує відхилення розподілу даних від гаусівського завдяки застосуванню дев'ятивимірною перетворення Бокса-Кокса. Використання створеної ймовірнісної моделі дозволило підвищити точність розпізнавання (Accuracy) з 0,9412 до 0,9782, ймовірність розпізнавання цільового образу (Specificity) з 0,9795 до 0,9949 та ймовірність розпізнавання представників інших класів (Recall) з 0,9225 до 0,9700.

7) Виконано порівняння побудованих ймовірнісних моделей для розпізнавання облич з реалізованими алгоритмами машинного навчання. Однокласова опорно векторна машина та ізоляційний ліс показали найнижчу точність. Автокодувальник (AE) показав результат, порівнянний з еліпсоїдом прогнозування для негаусівських даних з квантилем розподілу χ^2 -квадрат (PENGD- χ^2), однак, завдяки застосуванню нормалізуючого перетворення еліпсоїд прогнозування для нормалізованих даних з квантилем χ^2 -квадрат (PEND- χ^2) досяг значних покращень порівняно з AE, зокрема точність розпізнавання (Accuracy) 0,9733 проти 0,93, ймовірність розпізнавання цільового образу (Specificity) 0,97 проти 0,87 та ймовірність розпізнавання представників інших класів (Recall) 0,975 проти 0,96. Найвищу ймовірність розпізнавання цільового образу (Specificity) мають еліпсоїд прогнозування для негаусівських даних з квантилем F -розподілу (PENGD- F) зі значенням 0,99 та еліпсоїд прогнозування для нормалізованих даних з квантилем F -розподілу (PEND- F) зі значенням 1, нормалізація допомогла покращити Specificity і загальну точність (Accuracy) до 0,9567, при цьому ймовірність розпізнавання представників інших

класів (Recall) нижче 0,9350, ніж у PENGD-Chi, PEND-Chi та AE. Це підкреслює важливість нормалізуючих перетворень для покращення моделей еліпсоїда прогнозування для задач розпізнавання облич при роботі з негаусівськими даними.

8) Виконано порівняння розроблених ймовірнісних моделей для розпізнавання клавіатурного почерку з реалізованими алгоритмами машинного навчання. Еліпсоїди прогнозування для негаусівських даних виділяються найнижчою точністю серед оцінюваних моделей. IF показав найнижчу точність серед моделей машинного навчання. OCSVM і AE демонструють майже однакову точність та ймовірність розпізнавання. Еліпсоїд прогнозування для нормалізованих даних з квантилем розподілу χ^2 -квадрат виявився (PEND-Chi) найточнішою моделлю (Accuracy) 0,9782, на другому місці еліпсоїд прогнозування для нормалізованих даних з квантилем F -розподілу PEND-F 0,9765, далі йдуть 0,9697 (OCSVM) і 0,9630 (AE), ймовірність розпізнавання цільового образу (Specificity) становить 0,9949 (PEND-Chi), найвищий показник 1 в PEND-F проти 0,9744 (OCSVM) і 0,9641 (AE) та ймовірність розпізнавання представників інших класів (Recall) становить 0,97 (PEND-Chi) і 0,9650 (PEND-F) проти 0,9675 (OCSVM) і 0,9625 (AE). Це підкреслює важливість нормалізуючих перетворень для покращення ймовірнісних моделей в задачах розпізнавання клавіатурного почерку.

9) Створено ІТ для розпізнавання облич за десятивимірним вектором характеристик, розроблено відповідний інструментарій ІТ. Було розроблено відповідне ПЗ мовою Python і Qt Creator та інструкцію користувача для розпізнавання облич, яка розрахована на використання зазначеного ПЗ.

10) Створено ІТ для розпізнавання клавіатурного почерку за дев'ятивимірним вектором характеристик, розроблено відповідний інструментарій ІТ. Було розроблено відповідне ПЗ мовою Python і Qt Creator та інструкцію користувача для розпізнавання клавіатурного почерку, яка розрахована на використання зазначеного ПЗ.

11) Результати досліджень впроваджено в виробничі процеси ТОВ НВЦ «Діагностика та контроль», що підтверджується відповідними актами, які наведені у додатку А. Також результати дослідження використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та управління проектами (ННІКНУП) Національного університету кораблебудування ім. адмірала Макарова наступних дисциплін: «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»; «Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки», що підтверджується відповідним актом, який наведений у додатку Б.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Abiodun, O. I., Jantan, A., Omolara, A. E., Dada, K. V., Umar, A. M., Linus, O. U., ... Kiru, M. U. (2019). Comprehensive review of artificial neural network applications to pattern recognition. *IEEE Access*, 7, 158820-158846. DOI: 10.1109/ACCESS.2019.2945545.
2. Patil, A., & Rane, M. (2021). Convolutional neural networks: an overview and its applications in pattern recognition. *Information and Communication Technology for Intelligent Systems: Proceedings of ICTIS 2020*, 1, 21-30. DOI: 10.1007/978-981-15-7078-0_3.
3. Fatima, M., & Pasha, M. (2017). Survey of machine learning algorithms for disease diagnostic. *Journal of Intelligent Learning Systems and Applications*, 9(01), 1-16. DOI: 10.4236/jilsa.2017.91001.
4. Miller, D. D., & Brown, E. W. (2018). Artificial intelligence in medical practice: the question to the answer?. *The American journal of medicine*, 131(2), 129-133. DOI: 10.1016/j.amjmed.2017.10.035.
5. Gupta, A., Anpalagan, A., Guan, L., & Khwaja, A. S. (2021). Deep learning for object detection and scene perception in self-driving cars: Survey, challenges, and open issues. *Array*, 10, 100057. DOI: 10.1016/J.ARRAY.2021.100057.
6. Li, L., Mu, X., Li, S., & Peng, H. (2020). A review of face recognition technology. *IEEE Access*, 8, 139110-139120. DOI:10.1109/ACCESS.2020.3011028.
7. Bukola, A. C., Owolawi, P. A., Du, C., & Van Wyk, E. (2024). A Systematic Review and Comparative Analysis Approach to Boom Gate Access Using Plate Number Recognition. *Computers*, 13(11), 286. DOI: 10.3390/computers13110286.
8. Karale, A., Tiwari, A., Wadkar, A., Patil, A., & Waghulde, D. (2022). Online Transaction Security Using Face Recognition: A Review. *International Journal for Research in Applied Science & Engineering Technology*, 10(6), 2709-2717. DOI: 10.22214/ijraset.2022.44214.

9. Chenchov, I., Aleksieva-Petrova, A., & Petrov, M. (2021). Authentication mechanisms and classification: a literature survey. In *Intelligent Computing: Proceedings of the 2021 Computing Conference*, 3, 1051-1070. Springer International Publishing. DOI: 10.1007/978-3-030-80129-8_69.
10. Raul, N., Shankarmani, R., & Joshi, P. (2020). A comprehensive review of keystroke dynamics-based authentication mechanism. In *International Conference on Innovative Computing and Communications: Proceedings of ICICC 2019*, 2, 149-162. Springer Singapore. DOI: 10.1007/978-981-15-0324-5_13.
11. Bhattacharyya, Suman & Rahul, Kumar. (2013). Face recognition by linear discriminant analysis. *International Journal of Communication Network Security*. 2. 31-35. DOI: 10.47893/IJCNS.2014.1087.
12. Lu, J., Plataniotis, K. N., & Venetsanopoulos, A. N. (2003). Face recognition using LDA-based algorithms. *IEEE Transactions on Neural networks*, 14(1), 195-200. DOI: 10.1109/TNN.2002.806647.
13. Zainuddin, Zahir & Laswi, A. (2017). Implementation of The LDA Algorithm for Online Validation Based on Face Recognition. *Journal of Physics: Conference Series*. 801, 012047. DOI: 10.1088/1742-6596/801/1/012047.
14. Lu, J., Plataniotis, K. N., & Venetsanopoulos, A. N. (2003). Face recognition using LDA-based algorithms. *IEEE Transactions on Neural networks*, 14(1), 195-200. DOI: 10.1109/TNN.2002.806647.
15. Fierrez, J., Pozo, A., Martinez-Diaz, M., Galbally, J., & Morales, A. (2018). Benchmarking touchscreen biometrics for mobile authentication. *IEEE transactions on information forensics and security*, 13(11), 2720-2733. DOI: 10.1109/TIFS.2018.2833042.
16. Li, Q., & Chen, H. (2019, February). CDAS: a continuous dynamic authentication system. In *Proceedings of the 2019 8th International Conference on Software and Computer Applications*, pp. 447-452. DOI: 10.1145/3316615.3316691.
17. Chen, L., Zhong, Y., Ai, W., & Zhang, D. (2019, June). Continuous authentication based on user interaction behavior. In *2019 7th International*

Symposium on Digital Forensics and Security (ISDFS), 1-6. IEEE. DOI: 10.1109/ISDFS.2019.8757539.

18. Pal, M., & Foody, G. (May 2010) Feature Selection for Classification of Hyperspectral Data by SVM. In *IEEE Transactions on Geoscience and Remote Sensing*, vol. 48, no. 5, 2297-2307. DOI: 10.1109/TGRS.2009.2039484.

19. Jain, S., Chaudhari, V., Chuadhari, R., Chavan, T., & Shahane, P. (2022). A Survey on Face Recognition Techniques in Machine Learning. *International Journal of Scientific Research in Computer Science, Engineering and Information Technology (IJSRCSEIT)*, 8(6), 50-66. DOI: 10.32628/CSEIT228558.

20. Dhanabal, S., & Chandramathi, S. J. I. J. C. A. (2011). A review of various k-nearest neighbor query processing techniques. *International Journal of Computer Applications*, 31(7), 14-22. DOI: 10.5120/3836-5332.

21. Bansal, M., Goyal, A., & Choudhary, A. (2022). A comparative analysis of K-nearest neighbor, genetic, support vector machine, decision tree, and long short term memory algorithms in machine learning. *Decision Analytics Journal*, 3, 100071. DOI: 10.1016/j.dajour.2022.100071.

22. Alghamdi, Shatha & Elrefaei, Lamiaa. (2017). Dynamic Authentication of Smartphone Users Based on Touchscreen Gestures. *Arabian Journal for Science and Engineering*. 43. DOI: 10.1007/s13369-017-2758-x.

23. Baynath, P., Soyjaudah, K. S., & Khan, M. H. M. (2017, August). Keystroke recognition using neural network. In *2017 5th International Symposium on Computational and Business Intelligence (ISCBI)*, pp. 86-90. IEEE. DOI: 10.1109/ISCBI.2017.8053550.

24. Andrean, A., Jayabalan, M., & Thiruchelvam, V. (2020). Keystroke dynamics based user authentication using deep multilayer perceptron. *International Journal of Machine Learning and Computing*, 10(1), 134-139. DOI: 10.18178/ijmlc.2020.10.1.910.

25. Sharma, A., Jureček, M., & Stamp, M. (2023). Keystroke Dynamics for User Identification. *arXiv preprint arXiv:2307.05529*. DOI: 10.48550/arXiv.2307.05529.

26. Kamencay, P., Benco, M., Mizdos, T., & Radil, R. (2017). A new method for face recognition using convolutional neural network. *Advances in Electrical and Electronic Engineering*, 15(4), 663-672. DOI:10.15598/aeee.v15i4.2389.
27. Guo, S., Chen, S., & Li, Y. (2016). Face recognition based on convolutional neural network and support vector machine. In *2016 IEEE International conference on Information and Automation (ICIA)*, pp. 1787-1792. IEEE. DOI: 10.1109/ICInfA.2016.7832107.
28. Hu, G., Yang, Y., Yi, D., Kittler, J., Christmas, W., Li, S. Z., & Hospedales, T. (2015). When face recognition meets with deep learning: an evaluation of convolutional neural networks for face recognition. In *Proceedings of the IEEE international conference on computer vision workshops*, pp. 142-150. DOI: 10.1109/ICCVW.2015.58.
29. Li, Z., Cai, R., Li, H., Lam, K. Y., Hu, Y., & Kot, A. C. (2022). One-class knowledge distillation for face presentation attack detection. *IEEE Transactions on Information Forensics and Security*, 17, 2137-2150. DOI: 10.48550/arXiv.2205.03792.
30. Raul, N., Shankarmani, R., & Joshi, P. (2020). A comprehensive review of keystroke dynamics-based authentication mechanism. In *International Conference on Innovative Computing and Communications: Proceedings of ICICC 2019*, 2, 149-162. Springer Singapore. DOI: 10.1007/978-981-15-0324-5_13.
31. Ghorbani, H. (2019). Mahalanobis distance and its application for detecting multivariate outliers. *Facta Universitatis, Series: Mathematics and Informatics*, 583-595. DOI: 10.22190/FUMI1903583G.
32. Bezerra, V. H., da Costa, V. G. T., Barbon Junior, S., Miani, R. S., & Zarpelão, B. B. (2019). IoTDS: A one-class classification approach to detect botnets in Internet of Things devices. *Sensors*, 19(14), 3188. DOI: 10.3390/s19143188.
33. Kim, S. G., Park, D., & Jung, J. Y. (2021). Evaluation of one-class classifiers for fault detection: Mahalanobis classifiers and the mahalanobis–taguchi system. *Processes*, 9(8), 1450. DOI: 10.3390/pr9081450.

34. Seliya, N., Abdollah Zadeh, A., & Khoshgoftaar, T. M. (2021). A literature review on one-class classification and its potential applications in big data. *Journal of Big Data*, 8, 1-31. DOI: 10.1186/s40537-021-00514-x.
35. Damer, N., Grebe, J. H., Zienert, S., Kirchbuchner, F., & Kuijper, A. (2019, September). On the generalization of detecting face morphing attacks as anomalies: Novelty vs. outlier detection. In *2019 IEEE 10th International Conference on Biometrics Theory, Applications and Systems (BTAS)*, pp. 1-5. IEEE. DOI: 10.1109/BTAS46853.2019.9185995.
36. Xu, H., Pang, G., Wang, Y., & Wang, Y. (2023). Deep isolation forest for anomaly detection. *IEEE Transactions on Knowledge and Data Engineering*, 35(12), 12591-12604. DOI: 10.1109/TKDE.2023.3270293.
37. Vaiyapuri, T., & Binbusayyis, A. (2020). Application of deep autoencoder as an one-class classifier for unsupervised network intrusion detection: a comparative evaluation. *PeerJ Computer Science*, 6, 327. DOI: 10.7717/peerj-cs.327.
38. Oza, P., & Patel, V. M. (2019, May). Active authentication using an autoencoder regularized cnn-based one-class classifier. In *2019 14th IEEE International Conference on Automatic Face & Gesture Recognition (FG 2019)*, pp. 1-8. IEEE. DOI: 10.1109/FG.2019.8756525.
39. Gallego, G., Cuevas, C., Mohedano, R., & Garcia, N. (2013). On the Mahalanobis distance classification criterion for multidimensional normal distributions. *IEEE Transactions on Signal Processing*, 61(17), 4387-4396. DOI: 10.1109/TSP.2013.2269047.
40. Sen, S., Diawara, N., & Iftexharuddin, K. M. (2015). Statistical pattern recognition using Gaussian copula. *Journal of Statistical Theory and Practice*, 9, 768-777. DOI: 10.1080/15598608.2015.1008607.
41. Andrew R. Webb, & Keith D. Copsey (2003). *Statistical pattern recognition*. John Wiley & Sons.
42. Oyebola, O. (2023). Examining the Distribution of Keystroke Dynamics Features on Computer, Tablet and Mobile Phone Platforms. In *Mobile Computing and*

Sustainable Informatics: Proceedings of ICMCSI 2023, pp. 613-620. Singapore: Springer Nature Singapore. DOI: 10.1007/978-981-99-0835-6_43.

43. Lam, K. H., Meijer, K. A., Loonstra, F. C., Coerver, E. M. E., Twose, J., Redeman, E., ... Killestein, J. (2021). Real-world keystroke dynamics are a potentially valid biomarker for clinical disability in multiple sclerosis. *Multiple Sclerosis Journal*, 27(9), 1421-1431. DOI: 10.1177/1352458520968797.

44. Prykhodko, S., Prykhodko, A., & Shutko, I. (2021, September). Estimating the size of web apps created using the CakePHP framework by nonlinear regression models with three predictors. In *2021 IEEE 16th International Conference on Computer Sciences and Information Technologies (CSIT)* (Vol. 1, pp. 333-336). IEEE. DOI: 10.1109/CSIT52700.2021.9648680.

45. Prykhodko, S., Makarova, L., Prykhodko, K., & Pukhalevych, A. (2020, February). Application of transformed prediction ellipsoids for outlier detection in multivariate non-Gaussian data. In *2020 IEEE 15th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET)*, pp. 359-362. IEEE. DOI: 10.1109/TCSET49122.2020.235454.

46. Meshkani, S. A., Mehrabi, B., Yaghubpur, A., & Alghalandis, Y. F. (2011). The application of geochemical pattern recognition to regional prospecting: A case study of the Sanandaj–Sirjan metallogenic zone, Iran. *journal of geochemical exploration*, 108(3), 183-195. DOI: 10.1016/j.gexplo.2011.01.006.

47. Приходько, С., & Трухов, А. (2023). Побудова правил прийняття рішень для розпізнавання облич на основі квадрата відстані Махаланобіса для нормалізованих даних. *Information Technology: Computer Science, Software Engineering and Cyber Security*, (2), 50-58. DOI: 10.32782/IT/2023-2-6.

48. Prykhodko, S. B., & Trukhov, A. S. (2024). Face recognition using the ten-variate prediction ellipsoid for normalized data based on the Box-Cox transformation. *Radio Electronics, Computer Science, Control*, (2), 82-89. DOI: 10.15588/1607-3274-2024-2-9.

49. Prykhodko, S. B., & Trukhov, A. S. (2024). Face recognition using ten-variate prediction ellipsoids for normalized data with different quantiles. *Applied Aspects of Information Technology*, 7(2), 151-161. DOI: 10.15276/aait.07.2024.11.

50. Prykhodko, S., & Trukhov, A. (2024). Application of a Ten-Variate Prediction Ellipsoid for Normalized Data and Machine Learning Algorithms for Face Recognition. *International Workshop on Computer Modeling and Intelligent Systems*, 3702, 362-375.

51. Довбиш А. С. Основи теорії розпізнавання образів. У 2-х ч. Ч. 1 [Електронний ресурс] : навч. посіб. / А. С. Довбиш, І. В. Шелехов. – Суми : Сумський держ. ун-т, 2015. – 109 с. Режим доступу: <http://kist.ntu.edu.ua/textPhD/tro2.pdf> (дата звернення: 30.03.2024).

52. He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770-778. DOI: 10.48550/arXiv.1512.03385.

53. Spanhol, F. A., Oliveira, L. S., Petitjean, C., & Heutte, L. (2015). A dataset for breast cancer histopathological image classification. *IEEE transactions on biomedical engineering*, 63(7), 1455-1462. DOI: 10.1109/TBME.2015.2496264.

54. Khosrowabadi, R., Quek, H. C., Wahab, A., & Ang, K. K. (2010, August). EEG-based emotion recognition using self-organizing map for boundary detection. In *2010 20th international conference on pattern recognition*, pp. 4242-4245. IEEE. DOI: 10.1109/ICPR.2010.1031.

55. Beham, M. P., & Roomi, S. M. M. (2013). A review of face recognition methods. *International Journal of Pattern Recognition and Artificial Intelligence*, 27(04), 1356005. DOI:10.1142/S0218001413560053.

56. Dias, T., Vitorino, J., Maia, E., Sousa, O., & Praça, I. (2023). KeyRecs: A keystroke dynamics and typing pattern recognition dataset. *Data in Brief*, 50, 109509. DOI: 10.1016/j.dib.2023.109509.

57. Alshehri, A., Coenen, F., & Bollegala, D. (2019). Behavioural biometric continuous user authentication using multivariate keystroke streams in the spectral domain. In *Knowledge Discovery, Knowledge Engineering and Knowledge*

Management: 9th International Joint Conference, IC3K 2017, Funchal, Madeira, Portugal, November 1-3, 2017, Revised Selected Papers 9, pp. 43-66. Springer International Publishing. DOI: 10.1007/978-3-030-15640-4_3.

58. Oh, S. K., Yoo, S. H., & Pedrycz, W. (2013). Design of face recognition algorithm using PCA-LDA combined for hybrid data pre-processing and polynomial-based RBF neural networks: Design and its application. *Expert Systems with Applications*, 40(5), 1451-1466. DOI: 10.1016/j.eswa.2012.08.046.

59. Singh, S., Sharma, M., & Rao, N. S. (2012). Accurate face recognition using pca and lda. *International Journal of Managment, IT and Engineering*, 2(4), 101-121. DOI: 10.1109/ACCT.2012.52.

60. Zafaruddin, G. M., & Fadewar, H. S. (2019). Face recognition using eigenfaces. In *Computing, Communication and Signal Processing: Proceedings of ICCASP 2018*, pp. 855-864. Springer Singapore. DOI: 10.1007/978-981-13-1513-8_87.

61. Taneti, M. (2022). Secure Face Recognition System Based on Eigenface and Fisherface Techniques. *International Journal of Multidisciplinary Educational Research*, 11.

62. Chauhan, R., Ghanshala, K. K., & Joshi, R. C. (2018, December). Convolutional neural network (CNN) for image detection and recognition. In *2018 first international conference on secure cyber computing and communication (ICSCCC)*, pp. 278-282. IEEE. DOI:10.1109/ICSCCC.2018.8703316.

63. Lu, X., Zhang, S., Hui, P., & Lio, P. (2020). Continuous authentication by free-text keystroke based on CNN and RNN. *Computers & Security*, 96, 101861. DOI: 10.1016/j.cose.2020.101861.

64. Xiaofeng, L., Shengfei, Z., & Shengwei, Y. (2019). Continuous authentication by free-text keystroke based on CNN plus RNN. *Procedia computer science*, 147, 314-318. DOI: 10.1016/j.procs.2019.01.270.

65. Peral, I. E., Gil, L. S., & Zola, F. (2022). Conditional Generative Adversarial Network for keystroke presentation attack. In *Investigación en*

Ciberseguridad Actas de las VII Jornadas Nacionales (7º. 2022. Bilbao), pp. 228-231. Fundación Tecnalia Research and Innovation. DOI: 10.48550/arXiv.2212.08445.

66. Luo, M., Cao, J., Ma, X., Zhang, X., & He, R. (2021). FA-GAN: Face augmentation GAN for deformation-invariant face recognition. *IEEE Transactions on Information Forensics and Security*, 16, 2341-2355. DOI: 10.1109/TIFS.2021.3053460.

67. Hawkins, D. M., & Weisberg, S. (2017). Combining the box-cox power and generalised log transformations to accommodate nonpositive responses in linear and mixed-effects linear models. *South African Statistical Journal*, 51(2), 317-328. DOI: 10.37920/sasj.2017.51.2.5.

68. Hilal, W., Gadsden, S. A., & Yawney, J. (2022). Financial fraud: a review of anomaly detection techniques and recent advances. *Expert systems With applications*, 193, 116429. DOI: 10.1016/j.eswa.2021.116429.

69. Nurmaini, S., Darmawahyuni, A., Sakti Mukti, A. N., Rachmatullah, M. N., Firdaus, F., & Tutuko, B. (2020). Deep learning-based stacked denoising and autoencoder for ECG heartbeat classification. *Electronics*, 9(1), 135. DOI: 10.3390/electronics9010135.

70. Prykhodko, S., Prykhodko, N., Makarova, L., Kudin, O., Smykodub, T. & Prykhodko, A. (2017) A. Detecting bivariate outliers on the basis of normalizing transformations for non-Gaussian data. *Advanced Information Systems and Technologies, Proceedings of the V International Scientific Conference*. 17-19.

71. Etherington, T. R. (2021). Mahalanobis distances for ecological niche modelling and outlier detection: implications of sample size, error, and bias for selecting and parameterising a multivariate location and scatter method. *PeerJ*, 9, e11436. DOI: 10.7717/peerj.11436

72. McNeish, D. (2020). Should we use F-tests for model fit instead of chi-square in overidentified structural equation models?. *Organizational Research Methods*, 23(3), 487-510. DOI: 10.1177/1094428118809495

73. Sun, Y., Jiang, L., Pan, J., Sheng, S., & Hao, L. (2023). A satellite imagery smoke detection framework based on the Mahalanobis distance for early fire

identification and positioning. *International Journal of Applied Earth Observation and Geoinformation*, 118, 103257. DOI: 10.1016/j.jag.2023.103257.

74. Gorsuch, R. L., & Lehmann, C. (2017). Chi-square and F Ratio: Which should be used when? *Journal of Methods and Measurement in the Social Sciences*, 8(2), 58-71. DOI: 10.2458/v8i2.22990.

75. Prykhodko, S., Prykhodko, N., Makarova, L., & Pugachenko, K. (2017). Detecting outliers in multivariate non-Gaussian data on the basis of normalizing transformations. In *2017 IEEE First Ukraine Conference on Electrical and Computer Engineering (UKRCON)*, pp. 846-849. IEEE. DOI: 10.1109/UKRCON.2017.8100366.

76. Johnson, R. A., & Wichern, D. W. (2002). Applied multivariate statistical analysis. – Pearson Prentice Hall, 2007. – 800 p.

77. Burak. Dataset Pins Face Recognition. 2019. URL: <https://www.kaggle.com/datasets/hereisburak/pins-face-recognition> (дата звернення: 08.02.2023).

78. Saib, Y. M., & Pudaruth, S. (2021). Is face recognition with masks possible?. *International Journal of Advanced Computer Science and Applications*, 12(7). DOI: 10.14569/IJACSA.2021.0120706.

79. Majeed, K. A., Abbas, Z., Bakhtyar, M., Durrani, Z., Baber, J., & Ullah, I. (2021). Face Detectors Evaluation to Select the Fastest among DLIB, HAAR Cascade, and MTCNN. *Pakistan Journal of Emerging Science and Technologies*, 2, 1-13. DOI: 10.5281/zenodo.5089390.

80. Singh, N., & Brisilla, R. M. (2021, November). Comparison analysis of different face detecting techniques. In *2021 Innovations in Power and Advanced Computing Technologies (i-PACT)*, pp. 1-6. IEEE. DOI: 10.1109/i-PACT52855.2021.9696583.

81. Vikram, K., & Padmavathi, S. (2017, January). Facial parts detection using Viola Jones algorithm. In *2017 4th international conference on advanced computing and communication systems (ICACCS)*, pp. 1-4. IEEE. DOI: 10.1109/ICACCS.2017.8014636.

82. Tripathy, R., & Daschoudhury, R. (2014). Real-time face detection and tracking using haar classifier on SOC. *International Journal of Electronics and Computer Science Engineering*, 3(2), 175-184.
83. Kurukulasooriya, A., & Dharmarathne, A. T. (2011). Image Searching with Eigenfaces and Facial Characteristics. In *International Conference on Signal Processing, Image Processing, and Pattern Recognition*, pp. 215-224. Berlin, Heidelberg: Springer Berlin Heidelberg. DOI: 10.1007/978-3-642-27183-0_23.
84. Xiang, W. (2022). The Research and Analysis of Different Face Recognition Algorithms. In *Journal of Physics: Conference Series* (Vol. 2386, No. 1, p. 012036). IOP Publishing. DOI: 10.1088/1742-6596/2386/1/012036.
85. Ku, H., & Dong, W. (2020). Face recognition based on mtcnn and convolutional neural network. *Frontiers in Signal Processing*, 4(1), 37-42. DOI: 10.22606/fsp.2020.41006.
86. Du, J. (2020). High-precision portrait classification based on mtcnn and its application on similarity judgement. In *Journal of Physics: Conference Series* (Vol. 1518, No. 1, p. 012066). IOP Publishing. DOI: 10.1088/1742-6596/1518/1/012066.
87. Zhou, H., Chen, P., & Shen, W. (2018). A multi-view face recognition system based on cascade face detector and improved Dlib. In *MIPPR 2017: Pattern Recognition and Computer Vision* (Vol. 10609, pp. 37-42). SPIE. DOI: 10.1117/12.2282829.
88. Preethi, B. M., Sunitha, C., Parameshvar, M., Dharshini, B., & Gokul, S. (2023). Emotion Detection using HOG for Crime Detection. *Indian Journal of Science and Technology*, 16(41), 3617-3626. DOI: 10.17485/IJST/v16i41.1580.
89. Munasinghe, M. I. N. P. (2018, June). Facial expression recognition using facial landmarks and random forest classifier. In *2018 IEEE/ACIS 17th International Conference on Computer and Information Science (ICIS)*, pp. 423-427. IEEE. DOI: 10.1109/ICIS.2018.8466510.
90. Kopp, P., Bradley, D., Beeler, T., & Gross, M. (2019). Analysis and improvement of facial landmark detection. *Swiss Federal Institute of Technology Zurich, Zurich*. DOI: 10.13140/RG.2.2.10980.42886.

91. Haghpanah, M. A., Saeedizade, E., Masouleh, M. T., & Kalhor, A. (2022). Real-time facial expression recognition using facial landmarks and neural networks. In *2022 International Conference on Machine Vision and Image Processing (MVIP)*, pp. 1-7. IEEE. DOI: 10.1109/MVIP53647.2022.9738754.
92. Amato, G., Falchi, F., Gennaro, C., & Vairo, C. (2018). A comparison of face verification with facial landmarks and deep features. In *10th International Conference on Advances in Multimedia (MMEDIA)*, pp. 1-6.
93. Gaber, A., Taher, M. F., Wahed, M. A., & Shalaby, N. M. (2021). SVM classification of facial functions based on facial landmarks and animation Units. *Biomedical Physics & Engineering Express*, 7(5), 055008. DOI: 10.1088/2057-1976/ac107c.
94. Juhong, A., & Pintavirooj, C. (2017). Face recognition based on facial landmark detection. In *2017 10th Biomedical Engineering International Conference (BMEiCON)*, pp. 1-4. IEEE. DOI: 10.1109/BMEiCON.2017.8229173.
95. Bijarnia, S., & Singh, P. (2016). Age invariant face recognition using minimal geometrical facial features. In *Advanced Computing and Communication Technologies: Proceedings of the 9th ICACCT, 2015*, pp. 71-77. Springer Singapore. DOI: 10.1007/978-981-10-1023-1_7.
96. Nestor, A., Vettel, J. M., & Tarr, M. J. (2013). Internal representations for face detection: An application of noise-based image classification to BOLD responses. *Human brain mapping*, 34(11), 3101-3115. DOI: 10.1002/hbm.22128.
97. Johnston, B., & Chazal, P. D. (2018). A review of image-based automatic facial landmark identification techniques. *EURASIP Journal on Image and Video Processing*, 2018(1), 86. DOI: 10.1186/s13640-018-0324-4.
98. Wu, Y., & Ji, Q. (2019). Facial landmark detection: A literature survey. *International Journal of Computer Vision*, 127(2), 115-142. DOI: 10.1007/s11263-018-1097-z.
99. Killourhy, K. S., & Maxion, R. A. (2009). Comparing anomaly-detection algorithms for keystroke dynamics. In *2009 IEEE/IFIP international conference on*

dependable systems & networks, pp. 125-134. IEEE. DOI: 10.1109/DSN.2009.5270346.

100. Mardia, K. V. (1970). Measures of multivariate skewness and kurtosis with applications. *Biometrika*, 57(3), 519-530. DOI: 10.1093/biomet/57.3.519.

101. Трухов, А.С., & Приходько, С.Б. (26-28 жовтня 2021 р.). Аналіз існуючих математичних моделей для розпізнавання образів. *Матеріали II Всеукраїнської науково-практичної інтернет конференції*. – Миколаїв: НУК імені адмірала Макарова, 2021. – с. 8-10.

102. Трухов, А.С., & Приходько, С.Б. (2022). Математичні моделі та методи для ідентифікації та виділення образів. *Матеріали XXII Всеукраїнської науково-технічної конференції молодих вчених, аспірантів та студентів*. Одеса, 21-22 квітня 2022 р. - Одеса, Видавництво ОНТУ, 2022 р. – с. 44-46.

103. Трухов, А.С., & Приходько, С.Б. (2022). Застосування квадрата відстані махаланобіса в задачі розпізнавання облич. *Матеріали III Всеукраїнської науково-практичної інтернет конференції* (26-28 жовтня 2022 р.). – Миколаїв: НУК імені адмірала Макарова, 2022. – с. 35-37.

104. Prykhodko, S., & Trukhov, A. (2023). Free software for object identification and feature extraction. *Матеріали XIV-ої Міжнародної науково-практичної конференції «Free and Open Source Software»* (14-16 лютого 2023 р.). – Харків: Харківський національний економічний університет імені Семена Кузнеця, 2023.–110 с. 18-19.

105. Bagenal, M. (1955). A note on the relations of certain parameters following a logarithmic transformation. *Journal of the Marine Biological Association of the United Kingdom*, 34(2), 289-296.

106. Benoit, K. (2011). Linear regression models with logarithmic transformations. *London School of Economics, London*, 22(1), 23-36.

107. Blum, L., Elgendi, M., & Menon, C. (2022). Impact of box-cox transformation on machine-learning algorithms. *Frontiers in Artificial Intelligence*, 5, 877569. DOI: 10.3389/frai.2022.877569.

108. Osborne, J. (2010). Improving your data transformations: Applying the Box-Cox transformation. *Practical Assessment, Research, and Evaluation*, 15(1), 12. DOI: 10.7275/qbpc-gk17.
109. Lindsey, C., & Sheather, S. (2010). Power transformation via multivariate Box–Cox. *The Stata Journal*, 10(1), 69-81. DOI: 10.1177/1536867X1001000108.
110. Трухов, А.С., & Приходько, С.Б. (2023). Побудова правил прийняття рішень для нормалізованих даних у задачі розпізнавання образів. *Матеріали IV Всеукраїнської науково-практичної інтернет конференції* (30-31 жовтня 2023 р.). – Миколаїв: НУК імені адмірала Макарова, 2023. – с. 7-9.
111. Ghiasi, R., Khan, M. A., Sorrentino, D., Diaine, C., & Malekjafarian, A. (2024). An unsupervised anomaly detection framework for onboard monitoring of railway track geometrical defects using one-class support vector machine. *Engineering Applications of Artificial Intelligence*, 133, 108167. DOI: 10.1016/j.engappai.2024.108167.
112. Todkar, S. S., Baltazart, V., Ihamouten, A., Dérobert, X., & Guilbert, D. (2021). One-class SVM based outlier detection strategy to detect thin interlayer debondings within pavement structures using Ground Penetrating Radar data. *Journal of Applied Geophysics*, 192, 104392. DOI: 10.1016/j.jappgeo.2021.104392.
113. Nguyen, T. B. T., Liao, T. L., & Vu, T. A. (2019). Anomaly detection using one-class SVM for logs of juniper router devices. In *Industrial Networks and Intelligent Systems: 5th EAI International Conference, INISCOM 2019, Ho Chi Minh City, Vietnam, August 19, 2019, Proceedings*, pp. 302-312. Springer International Publishing. DOI: 10.1007/978-3-030-30149-1_24.
114. Lesouple, J., Baudoin, C., Spigai, M., & Tourneret, J. Y. (2021). Generalized isolation forest for anomaly detection. *Pattern Recognition Letters*, 149, 109-119. DOI: 10.1016/j.patrec.2021.05.022.
115. Chabchoub, Y., Togbe, M. U., Boly, A., & Chiky, R. (2022). An in-depth study and improvement of Isolation Forest. *IEEE Access*, 10, 10219-10237. DOI: 10.1109/ACCESS.2022.3144425.

116. Imashev, S. A. (2021, November). Extended isolation forest—application to outlier detection in geomagnetic data. In *IOP Conference Series: Earth and Environmental Science* (Vol. 929, No. 1, p. 012022). IOP Publishing. DOI: 10.1088/1755-1315/929/1/012022.
117. Togbe, M. U., Barry, M., Boly, A., Chabchoub, Y., Chiky, R., Montiel, J., & Tran, V. T. (2020). Anomaly detection for data streams based on isolation forest using scikit-multiflow. In *Computational Science and Its Applications—ICCSA 2020: 20th International Conference, Cagliari, Italy, July 1–4, 2020, Proceedings, Part IV* 20, pp. 15-30. Springer International Publishing. DOI: 10.1007/978-3-030-58811-3_2.
118. Xu, W., Jang-Jaccard, J., Singh, A., Wei, Y., & Sabrina, F. (2021). Improving performance of autoencoder-based network anomaly detection on NSL-KDD dataset. *IEEE Access*, 9, 140136-140146. DOI: 10.1109/ACCESS.2021.3116612.
119. Wang, Z., & Cha, Y. J. (2021). Unsupervised deep learning approach using a deep auto-encoder with a one-class support vector machine to detect damage. *Structural Health Monitoring*, 20(1), 406-425. DOI: 10.1177/1475921720934051.
120. Novac, O. C., Chirodea, M. C., Novac, C. M., Bizon, N., Oproescu, M., Stan, O. P., & Gordan, C. E. (2022). Analysis of the application efficiency of TensorFlow and PyTorch in convolutional neural network. *Sensors*, 22(22), 8872. DOI: 10.3390/s22228872.
121. Togbe, M. U., Barry, M., Boly, A., Chabchoub, Y., Chiky, R., Montiel, J., & Tran, V. T. (2020). Anomaly detection for data streams based on isolation forest using scikit-multiflow. In *Computational Science and Its Applications—ICCSA 2020: 20th International Conference, Cagliari, Italy, July 1–4, 2020, Proceedings, Part IV* 20, pp. 15-30. Springer International Publishing. DOI: 10.1007/978-3-030-58811-3_2.
122. Sewak, M., Sahay, S. K., & Rathore, H. (2020). An overview of deep learning architecture of deep neural networks and autoencoders. *Journal of Computational and Theoretical Nanoscience*, 17(1), 182-188. DOI: 10.1166/jctn.2020.8648.

123. Kartowisastro, I. H., & Latupapua, J. (2023). A Comparison of Adaptive Moment Estimation (Adam) and RMSProp Optimisation Techniques for Wildlife Animal Classification Using Convolutional Neural Networks. *Revue d'Intelligence Artificielle*, 37(4). DOI: 10.18280/ria.370424.

124. Трухов, А.С., & Приходько, С.Б. (2024). Розпізнавання клавіатурного почерку за допомогою дев'ятивимірних еліпсоїдів прогнозування для нормалізованих даних з різними квантилями. *Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2024)* - (Миколаїв, 30-31 жовтня 2024 р.) - с. 45-47.

125. Близнюк, М. (2022). Інформаційні технології в технологічній освіті. *Перспективи та інновації науки*, (9 (14)). DOI: 10.52058/2786-4952-2022-9(14)-43-52.

126. Біловус, Л. І., & Яблонська, Н. М. (2021). *Інформаційні технології у системі управління організацією* (Doctoral dissertation, Документ в інформаційному просторі: традиції минулого та виклики сучасності [Електронний ресурс]: матеріали VIII Всеукраїнської науково-практичної інтернет-конференції (Львів, 1–2 грудня 2021 року). Українська академія друкарства, кафедра інформаційної, бібліотечної та архівної справи. Львів: УАД, 2021. 112 с.. 17–22.

127. Rawat, A. (2020). A Review on Python Programming. *International Journal of Research in Engineering, Science and Management*, 3(12), 8-11.

128. Hasan, R. T., & Sallow, A. B. (2021). Face detection and recognition using opencv. *Journal of Soft Computing and Data Mining*, 2(2), 86-97. DOI: 10.30880/jscdm.2021.02.02.008.

129. Fitzpatrick, M. (2019). Create Simple GUI Applications with Python & Qt5. *Leanpub book*.

**ДОДАТОК А. АКТ ВПРОВАДЖЕННЯ РЕЗУЛЬТАТІВ
ДИСЕРТАЦІЙНОЇ РОБОТИ**



Вих.№3010 від 25.11.2024р.

АКТ

**впровадження результатів дисертаційної роботи на здобуття наукового ступеня доктора
філософії (PhD) в галузі знань 12 «Інформаційні технології»
за спеціальністю 122 «Комп'ютерні науки»
Трухова Артема Сергійовича**

Я, що нижче підписався, директор ТОВ НВЦ «Діагностика та контроль» Ярковець Олександр Чеславович, цим Актом засвідчую, що результати дисертації Трухова Артема Сергійовича, а саме: інструментарій інформаційної технології для розпізнавання облич та інструментарій інформаційної технології для розпізнавання клавіатурного почерку було впроваджено в систему контролю доступу до виробничих приміщень.

Впровадження зазначених вище результатів дозволило підвищити точність та ймовірність розпізнавання облич та клавіатурного почерку.

Директор ТОВ НВЦ
«Діагностика та контроль»



О.Ч. Ярковець

**ДОДАТОК Б. АКТ ВПРОВАДЖЕННЯ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЙНОЇ
РОБОТИ В НАВЧАЛЬНИЙ ПРОЦЕС**



МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
**НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ КОРАБЛЕБУДУВАННЯ
 ІМЕНІ АДМІРАЛА МАКАРОВА**

54007, м. Миколаїв, пр. Героїв України 9, тел./факс +38(0512) 42-42-80, e-mail: university@nuos.edu.ua



«ЗАТВЕРДЖУЮ»

Ректор

проф. Євген ТРУШЛЯКОВ

«29» листопада 2024 р.

АКТ

впровадження результатів дисертаційної роботи на здобуття наукового ступеня
 доктора філософії (PhD) в галузі знань 12 «Інформаційні технології» за
 спеціальністю 122 «Комп'ютерні науки»
 Трухова Артема Сергійовича

Результати дисертаційної роботи Трухова А.С. використовуються при викладанні студентам Навчально-наукового інституту комп'ютерних наук та управління проектами (ННІКНУП) наступних дисциплін:

- «Методи аналізу результатів експериментальних досліджень» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма «Комп'ютерні науки»;
- «Застосування методів багатомірного статистичного аналізу в інформаційних технологіях» спеціальності 122 «Комп'ютерні науки», галузь знань 12 «Інформаційні технології», освітня програма 122 «Комп'ютерні науки».

Директор ННІКНУП,
 к.т.н., проф. НУК

Тетяна ФАРІОНОВА

Виконавець: Н.О.Кіскіна (096)200-70-56

nuos.edu.ua

ДОДАТОК В. СПИСОК ПУБЛІКАЦІЙ ЗДОБУВАЧА ЗА ТЕМОЮ ДИСЕРТАЦІЇ ТА ВІДОМОСТІ ПРО АПРОБАЦІЮ РЕЗУЛЬТАТІВ ДИСЕРТАЦІЇ

Список публікацій здобувача за темою дисертації:

1) в яких опубліковані основні наукові результати дисертації:

- статті у наукових виданнях, включених на дату опублікування до переліку наукових фахових видань України:

1. Приходько, С., & Трухов, А. (2023). Побудова правил прийняття рішень для розпізнавання облич на основі квадрата відстані Махаланобіса для нормалізованих даних. *Information Technology: Computer Science, Software Engineering and Cyber Security*, (2), 50-58. DOI: 10.32782/IT/2023-2-6.

2. Prykhodko, S. B., & Trukhov, A. S. (2024). Face recognition using ten-variate prediction ellipsoids for normalized data with different quantiles. *Applied Aspects of Information Technology*, 7(2), 151-161. DOI: 10.15276/aait.07.2024.11.

- статті у періодичних наукових виданнях, проіндексованих у базах даних *Web of Science Core Collection* та/або *Scopus* (крім видань держави, визнаної Верховною Радою України державою-агресором):

3. Prykhodko, S. B., & Trukhov, A. S. (2024). Face recognition using the ten-variate prediction ellipsoid for normalized data based on the Box-Cox transformation. *Radio Electronics, Computer Science, Control*, (2), 82-89. DOI: 10.15588/1607-3274-2024-2-9 (індексується наукометричною базою *Web of Science*).

4. Prykhodko, S., & Trukhov, A. (2024). Application of a Ten-Variate Prediction Ellipsoid for Normalized Data and Machine Learning Algorithms for Face Recognition. *International Workshop on Computer Modeling and Intelligent Systems*, Zaporizhzhia, Ukraine, May 3, 2024. *CEUR Workshop Proceedings*, Vol.3702, pp. 362-375, 2024. <https://ceur-ws.org/Vol-3702/paper30.pdf> (Aachen, Germany, ISSN

1613-0073. Стаття включена до міжнародних наукометричних баз Scopus та DBLP).

2) які засвідчують апробацію матеріалів дисертації:

5. Трухов, А.С., & Приходько, С.Б. (26-28 жовтня 2021 р.). Аналіз існуючих математичних моделей для розпізнавання образів. *Матеріали II Всеукраїнської науково-практичної інтернет конференції*. – Миколаїв: НУК імені адмірала Макарова, 2021. – с. 8-10.

6. Трухов, А.С., & Приходько, С.Б. (2022). Математичні моделі та методи для ідентифікації та виділення образів. *Матеріали XXII Всеукраїнської науково-технічної конференції молодих вчених, аспірантів та студентів*. Одеса, 21-22 квітня 2022 р. - Одеса, Видавництво ОНТУ, 2022 р. – с. 44-46.

7. Трухов, А.С., & Приходько, С.Б. (2022). Застосування квадрата відстані махаланобіса в задачі розпізнавання облич. *Матеріали III Всеукраїнської науково-практичної інтернет конференції* (26-28 жовтня 2022 р.). – Миколаїв: НУК імені адмірала Макарова, 2022. – с. 35-37.

8. Prykhodko, S., & Trukhov, A. (2023). Free software for object identification and feature extraction. *Матеріали XIV-ої Міжнародної науково-практичної конференції «Free and Open Source Software»* (14-16 лютого 2023 р.). – Харків: Харківський національний економічний університет імені Семена Кузнеця, 2023.–110 с. 18-19.

9. Трухов, А.С., & Приходько, С.Б. (2023). Побудова правил прийняття рішень для нормалізованих даних у задачі розпізнавання образів. *Матеріали IV Всеукраїнської науково-практичної інтернет конференції* (30-31 жовтня 2023 р.). – Миколаїв: НУК імені адмірала Макарова, 2023. – с. 7-9.

10. Трухов, А.С., & Приходько, С.Б. (2024). Розпізнавання клавіатурного почерку за допомогою дев'ятивимірних еліпсоїдів прогнозування для нормалізованих даних з різними квантилями. *Всеукраїнській науково-практичній інтернет конференції «Інформаційні технології: моделі, алгоритми, системи» (ITMAS – 2024)* - (Миколаїв, 30-31 жовтня 2024 р.) - с. 45 47.