

**НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ  
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ  
імені ІГОРЯ СІКОРСЬКОГО»  
Навчально-науковий інститут прикладного системного аналізу  
Кафедра штучного інтелекту**

До захисту допущено:  
В. о. завідувачки кафедри  
\_\_\_\_\_ Ірина ДЖИГИРЕЙ  
«\_\_» \_\_\_\_\_ 2026 р.

**Дипломна робота**  
**на здобуття ступеня бакалавра**  
**за освітньо-професійною програмою «Системи і методи штучного інтелекту»**  
**спеціальності 122 «Комп'ютерні науки»**  
**на тему: «Структурно-параметричний синтез нейромережових структур**  
**в задачах комп'ютерного бачення в медицині»**

Виконала:  
студентка IV курсу, групи КІ-21  
Кошель Олена Олександрівна \_\_\_\_\_

Керівник:  
асистент кафедри штучного інтелекту  
Кот Анатолій Тарасович \_\_\_\_\_

Консультант з нормоконтролю:  
фахівець першої категорії кафедри штучного інтелекту,  
к.т.н., доцент Комариста Богдана Миколаївна \_\_\_\_\_

Рецензент:  
доцент кафедри авіоніки та систем управління  
НУ«Київський авіаційний інститут»  
Василенко Микола Павлович \_\_\_\_\_

Засвідчую, що у цій дипломній роботі  
немає запозичень з праць інших авторів  
без відповідних посилань.

Студент (-ка) \_\_\_\_\_

Київ – 2026 року

**Національний технічний університет України  
«Київський політехнічний інститут імені Ігоря Сікорського»  
Навчально-науковий інститут прикладного системного аналізу  
Кафедра штучного інтелекту**

Рівень вищої освіти – перший (бакалаврський)  
Спеціальність – 122 «Комп'ютерні науки»  
Освітньо-професійна програма «Системи і методи штучного інтелекту»

ЗАТВЕРДЖУЮ  
В. о. завідувачки кафедри  
\_\_\_\_\_ Ірина ДЖИГИРЕЙ  
« \_\_\_\_ » \_\_\_\_\_ 2026 р.

**ЗАВДАННЯ  
на дипломну роботу студенту  
Кошель Олені Олександрівні**

1. Тема роботи: «Структурно-параметричний синтез нейромережових структур в задачах комп'ютерного бачення в медицині», керівник роботи Кот Анатолій Тарасович, старший викладач кафедри штучного інтелекту, затверджені наказом по НН ІПСА від «27» травня 2026 р. № 1944-с.
2. Термін подання студентом роботи: «05» червня 2026 року.
3. Вихідні дані до роботи: набір аерознімків для формування синтетичних геометричних спотворень; 3000 clean/distorted пар зображень; файл params.npy із семивимірними параметрами дисторсії та нахилу; програмний код на Python з використанням OpenCV, PyTorch, torchvision, NumPy, scikit-image, pandas і matplotlib. Датасет BreakHis (Breast Cancer Histopathological Image Classification): 7,909 гістопатологічних зображень, 8 класів (4 доброякісних + 4 злоякісних), 4 рівні збільшення (40×, 100×, 200×, 400×); Наукові публікації з методів структурно-параметричного синтезу нейромережових структур та медичного комп'ютерного бачення; Фреймворки: Python 3.8+, TensorFlow/Keras 2.13.0; Обчислювальні ресурси: NVIDIA L4 GPU (24GB)
4. Зміст роботи: Дослідження предметної області: задачі комп'ютерного бачення в медицині, гістопатологічний аналіз, огляд методів синтезу нейромережових структур, еволюційні підходи; математичні та теоретичні основи: згорткові нейронні мережі, формалізація структурно-параметричного синтезу, генетичний алгоритм, генетичні оператори, метрики якості; проектування та розробка системи: архітектура програмного рішення, реалізація компонентів, використані технології, Docker/GPU; експериментальні дослідження: конфігурація, результати еволюції, аналіз 128 синтезованих моделей, порівняння з методами

5. Перелік ілюстративного матеріалу: Структура хромосоми для кодування нейромережевої архітектури (діаграма); Блок-схема алгоритму структурно-параметричного синтезу; Діаграма компонентів програмної системи; Приклади гістопатологічних зображень з датасету BreakHis (8 класів, 4 збільшення); Графіки динаміки еволюції (fitness по поколіннях); Траєкторія популяції в просторі синтезованих архітектур; Архітектура найкращої синтезованої моделі; Розподіл гіперпараметрів синтезованих моделей (оптимізатор, batch size, глибина); Confusion Matrix результатів класифікації гістологічних зображень; Порівняння синтезованих архітектур з еталонними методами.

6. Дата видачі завдання: «02» лютого 2026 року.

### Календарний план

| № з/п | Назва етапів виконання дипломної роботи                            | Термін виконання етапів роботи | Примітка |
|-------|--------------------------------------------------------------------|--------------------------------|----------|
| 1     | Аналіз літератури з методів СПС та медичного комп'ютерного бачення | 02.02 – 28.02.2026             |          |
| 2     | Розробка методу кодування нейромережевої структури                 | 01.03 – 15.03.2026             |          |
| 3     | Реалізація генетичних операторів для синтезу                       | 16.03 – 31.03.2026             |          |
| 4     | Розробка програмної системи структурно-параметричного синтезу      | 01.04 – 15.04.2026             |          |
| 5     | Підготовка датасету BreakHis та налаштування середовища            | 16.04 – 10.05.2026             |          |
| 6     | Проведення експериментів на GPU сервері                            | 11.05 – 24.05.2026             |          |
| 7     | Аналіз синтезованих архітектур та валідація відтворюваності        | 25.05 – 05.06.2026             |          |
| 8     | Оформлення пояснювальної записки                                   | 29.05 – 04.06.2026             |          |

Студент

Олена КОШЕЛЬ

Керівник

Анатолій КОТ

## РЕФЕРАТ

Дипломна робота: 149 с., 6 рис., 42 табл., 17 посилань.

Ключові слова: СТРУКТУРНО-ПАРАМЕТРИЧНИЙ СИНТЕЗ, НЕЙРОМЕРЕЖЕВІ СТРУКТУРИ, ГЕНЕТИЧНИЙ АЛГОРИТМ, ЗГОРТКОВІ НЕЙРОННІ МЕРЕЖІ, КЛАСИФІКАЦІЯ ГІСТОЛОГІЧНИХ ЗОБРАЖЕНЬ, КОМП'ЮТЕРНЕ БАЧЕННЯ В МЕДИЦИНІ, ГЛИБОКЕ НАВЧАННЯ.

Робота присвячена розробці методу структурно-параметричного синтезу згорткових нейронних мереж для класифікації гістопатологічних зображень тканин молочної залози на основі генетичного алгоритму з одночасною оптимізацією архітектурних та гіперпараметричних складових.

Розроблено метод уніфікованого кодування архітектурної та параметричної складових у вигляді хромосоми, адаптивну функцію пристосованості з урахуванням динаміки навчання та перенавчання, а також повний фреймворк відтворюваності синтезованих архітектур.

Експериментально підтверджено ефективність методу на датасеті BreakHis: досягнуто точність бінарної класифікації 86.83% (доброякісна/злоякісна пухлина) за 47 хвилин обчислень на NVIDIA L4 GPU. Проаналізовано 136 синтезованих архітектур з 16 поколінь еволюції. Похибка відтворюваності результатів становить 0.03%.

Розроблену систему впроваджено у вигляді програмного забезпечення з відкритим вихідним кодом, придатного для застосування в системах автоматизованої медичної діагностики.

## ABSTRACT

Bachelor's thesis: 149 p., 6 fig., 43 tables, 17 references.

Keywords: STRUCTURAL-PARAMETRIC SYNTHESIS, NEURAL NETWORK STRUCTURES, GENETIC ALGORITHM, CONVOLUTIONAL NEURAL NETWORKS, HISTOLOGICAL IMAGE CLASSIFICATION, COMPUTER VISION IN MEDICINE, DEEP LEARNING.

The object of research is the process of automated structural-parametric synthesis of convolutional neural network architectures for medical image classification tasks.

The subject of research is methods of evolutionary optimization of convolutional neural network architectures with simultaneous selection of structural components and training hyperparameters for the task of histopathological breast tissue image classification.

The purpose of the work is to develop a method for the structural-parametric synthesis of convolutional neural networks for breast cancer histopathological image classification based on a genetic algorithm with simultaneous optimization of architectural and hyperparameter components.

The thesis reviews and analyses modern neural architecture search (NAS) methods and evolutionary algorithms. A method of unified encoding of the architecture and hyperparameters as a chromosome, an adaptive fitness function accounting for training dynamics and overfitting, and a complete framework for result reproducibility were developed.

The effectiveness of the method was experimentally confirmed on the BreakHis dataset: a binary classification accuracy of 86.83% (benign/malignant tumor) was achieved within 47 minutes of computation on an NVIDIA L4 GPU. A total of 136 synthesized architectures from 16 generations of evolution were analysed; the reproducibility error of the results is 0.03%.

The developed system is implemented as open-source software and is suitable for application in automated medical diagnostic systems.

## ЗМІСТ

|                                                                                                |    |
|------------------------------------------------------------------------------------------------|----|
| ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ.....                                                                 | 11 |
| ВСТУП .....                                                                                    | 12 |
| РОЗДІЛ 1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА СУЧАСНИХ МЕТОДІВ<br>СИНТЕЗУ НЕЙРОМЕРЕЖЕВИХ СТРУКТУР..... | 16 |
| 1.1 Огляд методів синтезу нейромережєвих структур.....                                         | 16 |
| 1.1.1 Ручне проектування архітектур .....                                                      | 16 |
| 1.1.2 Необхідність автоматизації синтезу нейромережєвих структур ....                          | 17 |
| 1.1.3 Основні підходи до синтезу нейромережєвих структур .....                                 | 19 |
| 1.2 Класифікація методів синтезу нейромережєвих структур .....                                 | 20 |
| 1.2.1 Градієнтні методи (DARTS, ENAS) .....                                                    | 20 |
| 1.2.2 Методи на основі навчання з підкріпленням .....                                          | 21 |
| 1.2.3 Еволюційні методи (Genetic CNN, AmoebaNet) .....                                         | 22 |
| 1.2.4 Порівняльний аналіз підходів.....                                                        | 23 |
| 1.3 Еволюційні алгоритми в машинному навчанні .....                                            | 24 |
| 1.3.1 Основи генетичних алгоритмів .....                                                       | 24 |
| 1.3.2 Генетичне програмування .....                                                            | 26 |
| 1.3.3 Застосування еволюційних алгоритмів в оптимізації нейронних<br>мереж.....                | 27 |
| 1.4 Задача класифікації гістопатологічних зображень в медицині .....                           | 28 |
| 1.4.1 Комп'ютерне бачення в медичній діагностиці .....                                         | 28 |
| 1.4.2 Особливості гістопатологічних зображень .....                                            | 29 |
| 1.4.3 Датасет BreakHis: характеристики та особливості .....                                    | 30 |
| 1.4.4 Виклики класифікації гістопатологічних зображень .....                                   | 31 |
| 1.4.5 Існуючі підходи до класифікації гістопатологічних зображень .....                        | 33 |
| 1.5 Проблеми сучасних методів синтезу нейромережєвих структур у<br>медичному AI .....          | 34 |
| 1.5.1 Недоступна обчислювальна вартість .....                                                  | 34 |
| 1.5.2 Відтворюваність і верифікація результатів .....                                          | 35 |

|                                                                                                           |    |
|-----------------------------------------------------------------------------------------------------------|----|
| 1.5.3 Вузький простір синтезу .....                                                                       | 36 |
| 1.5.4 Відсутність одночасного синтезу структурної і параметричної<br>складових .....                      | 37 |
| 1.5.5 Нечутливість до асиметрії помилок у медичній діагностиці .....                                      | 38 |
| Висновки до розділу 1 .....                                                                               | 38 |
| <b>РОЗДІЛ 2 МАТЕМАТИЧНІ МОДЕЛІ ТА МЕТОДИ СТРУКТУРНО-ПАРАМЕТРИЧНОГО СИНТЕЗУ НЕЙРОМЕРЕЖЕВИХ СТРУКТУР ..</b> |    |
| 2.1 Основи згорткових нейронних мереж .....                                                               | 40 |
| 2.1.1 Згорткові шари .....                                                                                | 40 |
| 2.1.2 Pooling-шари (підвибірка) .....                                                                     | 42 |
| 2.1.3 Batch Normalization .....                                                                           | 42 |
| 2.1.4 Dropout .....                                                                                       | 43 |
| 2.1.5 Функції активації .....                                                                             | 44 |
| 2.2 Математична модель генетичного алгоритму для структурно-параметричного синтезу .....                  | 46 |
| 2.2.1 Формалізація простору синтезу .....                                                                 | 47 |
| 2.2.2 Кодування хромосоми .....                                                                           | 49 |
| 2.2.3 Функція пристосованості .....                                                                       | 50 |
| 2.2.4 Генетичні оператори .....                                                                           | 54 |
| 2.3 Селекція .....                                                                                        | 56 |
| 2.4 Кросовер .....                                                                                        | 59 |
| 2.4.1 Структурний кросовер .....                                                                          | 59 |
| 2.4.2 Параметричний кросовер .....                                                                        | 61 |
| 2.5 Мутація .....                                                                                         | 62 |
| 2.5.1 Структурна мутація .....                                                                            | 62 |
| 2.5.2 Параметрична мутація .....                                                                          | 65 |
| 2.6 Елітизм .....                                                                                         | 67 |
| 2.7 Валідація архітектур .....                                                                            | 68 |
| 2.8 Метрики оцінки якості для задач медичної класифікації .....                                           | 74 |
| 2.8.1 Accuracy (загальна точність) .....                                                                  | 74 |
| 2.8.2 Sensitivity (Чутливість, Recall) .....                                                              | 75 |
| 2.8.3 Specificity (Специфічність) .....                                                                   | 75 |

|                                                                                 |     |
|---------------------------------------------------------------------------------|-----|
| 2.8.4 Precision та F1-Score.....                                                | 76  |
| 2.8.5 AUC-ROC (Area Under the ROC Curve) .....                                  | 77  |
| 2.8.6 Матриця помилок для бінарної медичної класифікації .....                  | 78  |
| 2.8.7 Метрики ефективності синтезу для порівняння з іншими методами<br>NAS..... | 78  |
| 2.9 Обґрунтування вибору підходу .....                                          | 79  |
| 2.9.1 Переваги еволюційного підходу.....                                        | 80  |
| 2.9.2 Придатність для задачі класифікації гістопатологічних зображень.....      | 83  |
| 2.10 Багатоцільова оптимізація: алгоритм NSGA-II .....                          | 85  |
| 2.10.1 Формалізація задачі багатоцільової оптимізації.....                      | 85  |
| 2.10.2 Концепція Pareto-домінування і Pareto-фронту .....                       | 86  |
| 2.10.3 Алгоритм NSGA-II.....                                                    | 87  |
| Висновки до розділу 2 .....                                                     | 90  |
| <b>РОЗДІЛ 3 РЕАЛІЗАЦІЯ ТА ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ</b>                      |     |
| <b>СИСТЕМИ СТРУКТУРНО-ПАРАМЕТРИЧНОГО СИНТЕЗУ</b>                                |     |
| <b>НЕЙРОМЕРЕЖЕВИХ СТРУКТУР .....</b>                                            |     |
| 3.1 Архітектура програмного рішення .....                                       | 93  |
| 3.1.1 Загальна структура системи.....                                           | 93  |
| 3.1.2 Модульна організація коду .....                                           | 93  |
| 3.1.3 Діаграма компонентів .....                                                | 94  |
| 3.1.4 Взаємодія модулів .....                                                   | 95  |
| 3.2 Опис датасету BreakHis.....                                                 | 97  |
| 3.2.1 Характеристики датасету .....                                             | 97  |
| 3.2.2 Класи об'єктів.....                                                       | 97  |
| 3.2.3 Розподіл тренувальних та тестових даних .....                             | 99  |
| 3.2.4 Попередня обробка зображень .....                                         | 100 |
| 3.2.5 Обґрунтування вибору датасету.....                                        | 101 |
| 3.3 Реалізація генетичного алгоритму .....                                      | 102 |
| 3.3.1 Кодування хромосоми (chromosome.py).....                                  | 102 |
| 3.3.2 Генерація початкової популяції (population.py).....                       | 105 |
| 3.3.3 Реалізація генетичних операторів (operators.py).....                      | 105 |
| 3.3.4 Функція пристосованості (fitness.py) .....                                | 107 |



|                                                          |     |
|----------------------------------------------------------|-----|
| 3.3.5 Основний цикл еволюції (evolution.py) .....        | 108 |
| 3.4 Побудова та тренування моделей .....                 | 108 |
| 3.4.1 Перетворення хромосоми в Keras модель .....        | 108 |
| 3.4.2 Компіляція моделі з гіперпараметрами .....         | 109 |
| 3.4.3 Callbacks (EarlyStopping, ReduceLROnPlateau) ..... | 110 |
| 3.4.4 Процес тренування.....                             | 110 |
| 3.5 Система збереження та відтворюваності .....          | 112 |
| 3.5.1 Збереження повної історії еволюції .....           | 112 |
| 3.5.2 Збереження натренованих ваг (.h5 файли) .....      | 112 |
| 3.5.3 Детальне логування.....                            | 113 |
| 3.5.4 Валідація збережених моделей .....                 | 115 |
| 3.6 Використані технології та інструменти.....           | 116 |
| 3.6.1 Python 3.8+ .....                                  | 116 |
| 3.6.2 TensorFlow/Keras 2.13.0.....                       | 116 |
| 3.6.3 NumPy, Pandas для обробки даних.....               | 117 |
| 3.6.4 Docker для контейнеризації .....                   | 118 |
| 3.6.5 Git для контролю версій .....                      | 119 |
| 3.7 Розгортання на сервері .....                         | 120 |
| 3.7.1 Конфігурація Docker контейнера .....               | 120 |
| 3.7.2 Використання NVIDIA L4 GPU.....                    | 121 |
| 3.7.3 Скрипти для автоматизованого розгортання .....     | 122 |
| 3.8 Конфігурація експериментів.....                      | 124 |
| 3.8.1 Параметри генетичного алгоритму .....              | 124 |
| 3.8.2 Простір пошуку архітектур.....                     | 124 |
| 3.8.3 Діапазони гіперпараметрів.....                     | 126 |
| 3.8.4 Обчислювальні ресурси.....                         | 127 |
| 3.9 Результати еволюційного процесу .....                | 128 |
| 3.9.1 Конфігурація повного експерименту .....            | 128 |
| 3.9.2 Порівняльна динаміка еволюції.....                 | 129 |
| 3.9.3 Аналіз невалідних архітектур .....                 | 131 |
| 3.10 Найкращі знайдені архітектури.....                  | 131 |
| 3.10.1 Найкраща архітектура Baseline.....                | 131 |

|                                                                      |     |
|----------------------------------------------------------------------|-----|
| 3.10.2 Оптимальні гіперпараметри (Baseline) .....                    | 133 |
| 3.10.3 Результати Baseline на валідаційній вибірці .....             | 133 |
| 3.11 Найкраща архітектура NSGA-II та порівняльний аналіз.....        | 134 |
| 3.11.1 Архітектура найкращого рішення Pareto-фронту.....             | 134 |
| 3.11.2 Результати NSGA-II найкращої моделі .....                     | 135 |
| 3.12 Валідація відтворюваності.....                                  | 138 |
| 3.13 Порівняння з альтернативними методами .....                     | 141 |
| 3.13.1 Ручно спроектовані архітектури (ResNet, VGG) .....            | 141 |
| 3.13.2 Інші методи NAS (DARTS, ENAS) .....                           | 142 |
| 3.13.3 Якість знайдених архітектур.....                              | 144 |
| 3.14 Візуалізація та обговорення результатів .....                   | 145 |
| 3.14.1 Графіки fitness dynamics .....                                | 145 |
| 3.14.2 Траєкторія популяції в просторі архітектур.....               | 146 |
| 3.14.3 Виявлені архітектурні паттерни .....                          | 147 |
| 3.14.4 Практична застосовність для медичних систем діагностики ..... | 148 |
| Висновки до розділу 3 .....                                          | 150 |
| ВИСНОВКИ.....                                                        | 152 |
| ПЕРЕЛІК ПОСИЛАНЬ .....                                               | 155 |

## ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ

- AI – Artificial Intelligence (штучний інтелект)
- AUC – Area Under the Curve (площа під кривою)
- AUC-ROC – Area Under the Receiver Operating Characteristic Curve (площа під ROC-кривою)
- BN – Batch Normalization (пакетна нормалізація)
- CNN – Convolutional Neural Network (згорткова нейронна мережа)
- Conv2D – Two-Dimensional Convolutional layer (двовимірний згортковий шар)
- DARTS – Differentiable Architecture Search (диференційовний пошук архітектур)
- ELU – Exponential Linear Unit (експоненційна лінійна функція активації)
- ENAS – Efficient Neural Architecture Search (ефективний пошук нейромережових архітектур)
- F1-Score – гармонійне середнє точності (precision) та повноти (recall)
- GA – Genetic Algorithm (генетичний алгоритм)
- GPU – Graphics Processing Unit (графічний процесор)
- MP / MaxPool – Max Pooling (шар максимізаційної підвибірки)
- NAS – Neural Architecture Search (пошук нейромережових архітектур)
- NSGA-II – Non-dominated Sorting Genetic Algorithm II (генетичний алгоритм недомінованого сортування, друга версія)
- ReLU – Rectified Linear Unit (випрямлена лінійна функція активації)
- ROC – Receiver Operating Characteristic (робоча характеристика приймача)
- СПС – структурно-параметричний синтез

## ВСТУП

Комп'ютерне бачення в медицині є одним із найбільш динамічно зростаючих напрямків застосування штучного інтелекту. Автоматизований аналіз медичних зображень – гістопатологічних зрізів, рентгенограм, знімків МРТ та КТ – дозволяє суттєво підвищити точність і швидкість діагностики, знизити залежність від суб'єктивного досвіду лікаря та забезпечити доступність кваліфікованого аналізу у медичних закладах з обмеженими ресурсами.

Особливого значення набуває автоматична класифікація гістопатологічних зображень – знімків тканин під мікроскопом. При онкологічній діагностиці патологоанатом аналізує тисячі зображень гістологічних зрізів для визначення наявності, типу та стадії злоякісного новоутворення. Рак молочної залози залишається одним із найпоширеніших онкологічних захворювань у світі: за даними ВООЗ, щороку реєструється понад 2,3 мільйони нових випадків. Рання та точна гістологічна діагностика безпосередньо впливає на вибір тактики лікування та прогноз виживаності.

Проектування архітектур глибоких нейронних мереж для таких задач традиційно вимагає значних експертних знань. Процес ручного налаштування є трудомістким, суб'єктивним та неоптимальним – обраний вручну набір параметрів не гарантує досягнення глобального оптимуму. Формальний підхід до цієї проблеми – структурно-параметричний синтез (СПС) нейромережових структур – передбачає одночасну автоматичну оптимізацію топології мережі (структурна складова) та числових гіперпараметрів навчання (параметрична складова).

Більшість існуючих методів автоматичного синтезу нейромережових структур мають суттєві обмеження: – Надмірно високі обчислювальні витрати (від сотень до тисяч GPU-годин) – Обмежений простір пошуку – Лише архітектура без гіперпараметрів – Низька відтворюваність та

відсутність верифікації результатів – Нечутливість до специфіки медичних датасетів (дисбаланс класів, висока варіабельність зображень)

Це обумовлює актуальність розробки ефективного та доступного методу структурно-параметричного синтезу нейромережових структур, орієнтованого на задачі медичного комп'ютерного бачення, зокрема класифікацію гістопатологічних зображень.

Робота виконувалась на кафедрі штучного інтелекту Навчально-наукового інституту прикладного системного аналізу Національного технічного університету України “Київський політехнічний інститут імені Ігоря Сікорського” і є частиною наукових досліджень кафедри у галузі автоматизованого машинного навчання, еволюційних обчислень та інтелектуальних систем підтримки прийняття рішень.

Дослідження відповідає науковій тематиці “Розробка методів інтелектуального синтезу нейромережових структур для задач аналізу складно структурованих даних” та пріоритетному напрямку “Штучний інтелект у медицині та охороні здоров'я”.

**Об'єкт дослідження** – Процес структурно-параметричного синтезу оптимальних нейромережових структур для задач класифікації медичних зображень.

**Предмет дослідження** – Методи еволюційної оптимізації архітектур згорткових нейронних мереж з одночасним синтезом структурних компонентів та параметричних характеристик навчання для задачі класифікації гістопатологічних зображень тканин молочної залози.<sup>1</sup>

---

<sup>1</sup> Тут і нижче використано такий інструмент штучного інтелекту як чат-бот з генеративним штучним інтелектом ChatGPT виключно для корегування та редагування тексту, створеного автором цієї дипломної роботи, на основі автоматизованої перевірки граматики, структури та стилю, що відповідає Політиці використання штучного інтелекту для академічної діяльності в КПІ ім. Ігоря Сікорського (протокол №11 Вченої ради КПІ ім. Ігоря Сікорського від 11 грудня 2023 р.).

**Мета роботи** – розробка та експериментальна апробація методу структурно-параметричного синтезу згорткових нейронних мереж на основі генетичного алгоритму з одночасною оптимізацією архітектурних та гіперпараметричних складових для задачі класифікації гістопатологічних зображень тканин молочної залози.

**Для досягнення поставленої мети необхідно вирішити низку завдань.**

1. Провести аналіз задач комп'ютерного бачення в медицині, особливостей гістопатологічного аналізу та існуючих підходів до синтезу нейромережових структур.
2. Дослідити еволюційні методи оптимізації нейронних мереж та обґрунтувати вибір генетичного алгоритму для задачі структурно-параметричного синтезу.
3. Розробити метод кодування структури нейронної мережі та гіперпараметрів навчання у вигляді єдиної хромосоми.
4. Реалізувати систему валідації синтезованих архітектур з автоматичним виправленням невалідних конфігурацій.
5. Розробити адаптивну функцію пристосованості, що враховує динаміку навчання та специфіку медичних датасетів (дисбаланс класів).
6. Адаптувати генетичні оператори (селекція, кросовер, мутація, елітизм) для синтезу в просторі архітектурних та параметричних конфігурацій.
7. Спроекувати та реалізувати програмну систему структурно-параметричного синтезу з підтримкою Docker та GPU.
8. Провести експериментальне дослідження на датасеті BreakHis (гістопатологічні зображення раку молочної залози).
9. Проаналізувати результати еволюційного процесу, виявити оптимальні архітектурні паттерни та гіперпараметричні конфігурації для гістологічних зображень.
10. Реалізувати систему збереження повної історії синтезу та незалежної валідації для забезпечення відтворюваності.

11. Порівняти розроблений метод з існуючими підходами за критеріями точності, обчислювальної ефективності та відтворюваності.

**Методи дослідження:** генетичні алгоритми, глибоке навчання, методи комп'ютерного бачення, методи обробки медичних зображень, статистичний

**Наукова новизна одержаних результатів:**

1. Уніфіковане генетичне кодування структурної та параметричної складових: вперше для задач медичного комп'ютерного бачення запропоновано представлення хромосоми, яке спільно кодує архітектуру CNN (типи шарів, параметри, зв'язність) та гіперпараметри навчання (оптимізатор, learning rate, batch size) в єдиному еволюційному фреймворку. Це дозволяє коеволюцію структурних та параметричних конфігурацій, виявляючи нетривіальні залежності між компонентами нейромережі та параметрами її навчання.

2. Адаптивна функція пристосованості для медичних задач: розроблено механізм best-epoch selection, який оцінює синтезовані архітектури на піку їх валідаційної продуктивності з адаптивною ранньою зупинкою. Функція враховує специфіку медичних датасетів: дисбаланс класів (доброякісні/злоякісні пухлини) та варіабельність гістологічних зображень між різними рівнями збільшення мікроскопу.

3. Повний фреймворк відтворюваності результатів синтезу: реалізовано систему збереження всіх синтезованих моделей (архітектура + ваги) з усіх поколінь еволюції, що дозволяє незалежну верифікацію результатів. Досягнуто похибку валідації 0.03%, що підтверджує детерміністичну реконструкцію синтезованих архітектур.

4. Обчислювально доступна реалізація структурно-параметричного синтезу: розроблено pipeline, який зменшує обчислювальні вимоги на 2–4 порядки (0.58 GPU-години проти 0.5–3,150 GPU-днів у існуючих методів), роблячи СПС доступним для медичних науково-дослідних центрів без спеціалізованих обчислювальних кластерів.

## РОЗДІЛ 1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ТА СУЧАСНИХ МЕТОДІВ СИНТЕЗУ НЕЙРОМЕРЕЖЕВИХ СТРУКТУР

### 1.1 Огляд методів синтезу нейромережевих структур

#### 1.1.1 Ручне проектування архітектур

Задача класифікації гістопатологічних зображень ставить перед дослідниками медичного AI специфічні вимоги до архітектури нейронних мереж: вхідні зображення великого розміру ( $224 \times 224$  пікселів і більше), обмежені датасети порівняно з природніми зображеннями, критична асиметрія помилок (пропущений рак загрожує здоров'ю пацієнта). Традиційно архітектури для таких задач проектуються вручну кваліфікованими фахівцями у галузі глибокого навчання та патологічної анатомії.

Процес ручного проектування CNN для медичного аналізу зображень охоплює:

- 1) визначення глибини мережі та типів шарів (згорткові, pooling, нормалізаційні, регуляризаційні);
- 2) налаштування параметрів кожного шару з урахуванням розміру гістологічних зображень;
- 3) підбір режиму навчання – оптимізатор, темп навчання, розмір пакету;
- 4) ітеративне тестування архітектур на конкретному датасеті.

Обмеження ручного підходу в контексті медичного AI: тривалість розробки, формування ефективної архітектури для нової нозології займає від кількох тижнів до місяців.

1. Залежність від досвіду: дослідник переносить патерни, успішні для інших датасетів, без гарантії їхньої оптимальності для конкретного гістологічного матеріалу.



2. Неповнота пошуку: комбінаторний простір конфігурацій (глибина, ширина, тип активацій, гіперпараметри) недосяжний для ручного перебору.
3. Нечутливість до специфіки домену: ручно налаштовані мережі часто не враховують особливостей H&E-забарвлення, внутрішньокласової морфологічної варіабельності та дисбалансу класів.

Серед відомих прикладів CNN, розроблених вручну для аналізу гістологічних та медичних зображень:

- 1) U-Net [1] – архітектура для сегментації медичних зображень, що вимагала місяців розробки і залишається базовою лінією в гістологічному аналізі;
- 2) Araújo et al. [2] – багатомасштабна CNN для чотирикласової класифікації гістологічних зображень молочної залози (нормальна тканина, доброякісне ураження, in situ карцинома, інвазивна карцинома), що вимагала кропіткого ручного налаштування кожного шару;
- 3) Spanhol et al. [3] – CNN на базі архітектури AlexNet для бінарної класифікації гістологічних знімків BreakHis на клас «доброякісна/злаякісна», що потребувала ручного підбору гіперпараметрів та процедури аугментації.

Наведені приклади свідчать: попри якісні результати, ручне проектування є ресурсоємним і не масштабується на нові типи патологій без повторного витрачання значних ресурсів.

### *1.1.2 Необхідність автоматизації синтезу нейромережових структур*

У медичній діагностиці автоматизація проектування нейромереж набуває особливої актуальності з кількох причин.

По-перше, спостерігається глобальний дефіцит патологоанатомів: за оцінками Американського товариства клінічної патології, у країнах з низьким і середнім доходом лише 1 патологоанатом припадає на 500 000 осіб населення. Автоматизована система підтримки діагностичних рішень може компенсувати цей дефіцит, але вимагає надійного та швидкого способу адаптації до конкретного профілю патологій конкретного закладу.

По-друге, в сучасній онкопатології для одного пацієнта може аналізуватись сотні зрізів з різними збільшеннями мікроскопу. Ефективна CNN для такого завдання має бути оптимізована під конкретний набір збільшень і розмір датасету – параметри, що суттєво варіюються між медичними установами.

По-третє, кожен новий тип онкологічного захворювання чи нова методологія забарвлення вимагає нової архітектури CNN. Ручне перепроєктування мережі щоразу є неприйнятним для практичного застосування.

Відповіддю на ці виклики є концепція автоматизованого машинного навчання (AutoML) – підхід, що автоматизує конвеєр побудови моделей: вибір ознак, вибір моделі, оптимізацію гіперпараметрів, пошук архітектури. В рамках цієї концепції структурно-параметричний синтез (СПС) нейромережових структур – це автоматизований пошук оптимальної конфігурації CNN, що одночасно охоплює:

- структурну складову – послідовність і типи шарів, глибину мережі;
- параметричну складову – гіперпараметри навчання (темپ навчання, розмір пакету, тип оптимізатора, коефіцієнти регуляризації).

Об'єднання обох складових в єдиній задачі оптимізації дає змогу виявляти нетривіальні взаємозалежності між ними – наприклад, те, що неглибока мережа може працювати краще зі специфічним темпом навчання для конкретного гістологічного датасету.

### 1.1.3 Основні підходи до синтезу нейромережесих структур

Задача структурно-параметричного синтезу формалізується як оптимізаційна задача в об'єднаному просторі архітектурних та параметричних конфігурацій:

$$(a^*, \varphi^*) = \arg \max_{(a, \varphi) \in \mathcal{A} \times \Phi} Q(\mathcal{M}(a, \varphi), \mathcal{D}_{\text{hist}}), \quad (1.1)$$

де  $a$  – структурна конфігурація (послідовність шарів та їх типи);  $\varphi$  – параметрична конфігурація (гіперпараметри навчання);  $\mathcal{A}$  – допустимий простір архітектур;  $\Phi$  – допустимий простір гіперпараметрів;  $\mathcal{M}(a, \varphi)$  – модель з архітектурою  $a$  та параметрами  $\varphi$ ;  $\mathcal{D}_{\text{hist}}$  – валідаційна вибірка гістопатологічних зображень;  $Q$  – функція якості (наприклад, комбінація accuracy та AUC-ROC).

Будь-який метод СПС включає три компоненти: простір синтезу (множина допустимих конфігурацій), стратегію пошуку (алгоритм вибору наступної конфігурації для оцінки) та оцінку якості (критерій придатності синтезованої архітектури).

Порівняння основних стратегій пошуку.

1. Випадковий пошук (Random Search) – базовий підхід із рівномірним семплуванням конфігурацій. Простий у реалізації, але вимагає великої кількості оцінок; неефективний при великому просторі пошуку.
2. Навчання з підкріпленням (Reinforcement Learning) – контролер навчається генерувати перспективні архітектури на основі накопиченого досвіду. Вимагає значних обчислювальних ресурсів і схильний до нестабільності навчання.

3. Градієнтні методи (Gradient-based) – безперервна релаксація дискретних архітектурних рішень дозволяє застосовувати градієнтну оптимізацію. Швидкі, але обмежені диференційованими операціями і не охоплюють параметричну складову.
4. Еволюційні методи (Evolutionary) – популяція конфігурацій еволюціонує через оператори мутації, кросовера та відбору. Гнучкі, не потребують диференційованості, природно поєднують структурний та параметричний синтез. Особливо придатні для невеликих датасетів (~1K–10K зображень), характерних для медичних застосувань.

## 1.2 Класифікація методів синтезу нейромережевих структур

### 1.2.1 Градієнтні методи (*DARTS, ENAS*)

DARTS (Differentiable Architecture Search) [4] вводить безперервну релаксацію вибору архітектурних операцій, перетворюючи дискретну оптимізаційну задачу на диференційовану. Замість жорсткого вибору однієї операції  $o$  у кожному ребрі між вузлами обчислювального графа застосовується зважена сума:

$$\tilde{o}^{(i,j)}(x) = \sum_{o \in \mathcal{O}} \frac{e^{\beta_o^{(i,j)}}}{\sum_{o' \in \mathcal{O}} e^{\beta_{o'}^{(i,j)}}} \cdot o(x), \quad (1.2)$$

де  $\beta_o^{(i,j)}$  – навчальні архітектурні ваги,  $\mathcal{O}$  – набір операцій-кандидатів.

Оптимізація відбувається двома рівнями: ваги мережі  $w$  мінімізують тренувальну похибку, а архітектурні ваги  $\beta$  мінімізують валідаційну похибку.

Придатність для медичних задач: обмежена. DARTS орієнтований на ImageNet-масштаб і не включає оптимізацію гіперпараметрів навчання. При малих гістологічних датасетах ( $N < 5000$  зображень) bi-level оптимізація схильна до перенавчання архітектурних ваг на валідаційній вибірці.

ENAS (Efficient Neural Architecture Search) [5] вирішує проблему обчислювальної вартості NAS за допомогою спільних ваг (weight sharing): усі кандидатні архітектури використовують один набір ваг надграфа (supergraph). Контролер на основі RNN навчається генерувати перспективні підграфи за допомогою методу REINFORCE.

Придатність для медичних задач: помірна. ENAS прискорює пошук до 0.5 GPU-дня, проте weight sharing вносить систематичну похибку: архітектура, що виглядає оптимальною в рамках спільних ваг, може виявитися неоптимальною при навчанні з нуля. Для малих медичних датасетів це особливо критично через нестабільність навчання контролера.

### *1.2.2 Методи на основі навчання з підкріпленням*

NASNet [6, 7] використовує контролер (RNN), що послідовно приймає рішення про будову архітектури і оновлюється за рахунок сигналу нагороди – точності дочірньої мережі на валідаційній вибірці:

$$\nabla_{\psi} J(\psi) = \mathbb{E}_{A \sim \pi(\cdot | \psi)} [R(A) \cdot \nabla_{\psi} \log \pi(A | \psi), ] \quad (1.3),$$

де  $\psi$  – параметри контролера,  $A$  – синтезована архітектура,  $R(A)$  – відповідна нагорода.

Результати: 97.40% на CIFAR-10, але за рахунок 1800 GPU-днів обчислень.

Придатність для медичних задач: непрактична. 1800 GPU-днів недоступні для клінічних дослідницьких груп. Крім того, метод не оптимізує гіперпараметри навчання, що критично для нестабільних за розміром медичних датасетів.

### 1.2.3 Еволюційні методи (*Genetic CNN, AmoebaNet*)

Genetic CNN [8] вперше демонструє практичність генетичного алгоритму для синтезу архітектур CNN. Хромосома кодується як бінарна матриця суміжності шарів у межах фіксованих “стадій”:

$$S = \begin{bmatrix} 0 & s_{12} & s_{13} & s_{14} \\ 0 & 0 & s_{23} & s_{24} \\ 0 & 0 & 0 & s_{34} \\ 0 & 0 & 0 & 0 \end{bmatrix}, s_{ij} \in \{0,1\}, \quad (1.4)$$

де  $s_{ij} = 1$  означає пряме з’єднання між вузлом  $i$  та вузлом  $j$ .

Метод використовує турнірну селекцію, побітовий кросовер матриць суміжності та мутацію фліпом окремих бітів. Результат – 92.90% на CIFAR-10 – підтвердив практичну застосовність еволюційного синтезу архітектур.

AmoebaNet [9] розвиває еволюційний підхід, вводячи регуляризацію “омолодження популяції”: після кожної ітерації найстаріша особина виключається з популяції незалежно від її fitness. Це запобігає передчасній конвергенції і підтримує різноманіття архітектур.

Алгоритм підтримує фіксований розмір популяції  $|\mathcal{P}| = N$ . На кожній ітерації:

- 1) випадково обирається підмножина  $\mathcal{T} \subset \mathcal{P}$ ,  $|\mathcal{T}| = k$  (турнір);
- 2) з  $\mathcal{T}$  обирається найкраща особина-батько;
- 3) батько мутується, новий нащадок оцінюється і додається до  $\mathcal{P}$ ;

4) найстаріша особина видаляється з  $\mathcal{P}$ .

AmoebaNet досягнув 97.87% на CIFAR-10, перевершивши RL-підходи, але за рахунок 3150 GPU-днів.

Переваги еволюційного підходу для СПС медичних зображень: не вимагає диференційованості простору пошуку; природно об'єднує структурний і параметричний синтез в одному циклі; ефективно працює на малих датасетах; кожна оцінка є незалежним тренуванням з нуля; легко адаптується під нову функцію якості (наприклад, з урахуванням AUC-ROC або чутливості).

#### 1.2.4 Порівняльний аналіз підходів

Дані наведено в таблиці 1.1.

**Таблиця 1.1** – Порівняння методів синтезу нейромережових структур за критеріями, релевантними для медичного AI

| Метод           | Стратегія    | GPU-час    | Асс. (CIFAR-10) | Синтез гіпер-параметрів | Малі датасети (<5K) |
|-----------------|--------------|------------|-----------------|-------------------------|---------------------|
| NASNet          | RL           | 1 800 днів | 97.40%          | Ні                      | Непридатний         |
| Amoeba Net      | Evolution    | 3 150 днів | 97.87%          | Ні                      | Непридатний         |
| DARTS           | Gradient     | 4 дні      | 97.24%          | Ні                      | Обмежено            |
| ENAS            | RL + sharing | 0.5 дня    | 97.11%          | Ні                      | Обмежено            |
| Genetic CNN     | Evolution    | -          | 92.90%          | Ні                      | Придатний           |
| СПС (ця робота) | Evolution    | < 1 дня    | -               | Так                     | Так                 |

З аналізу випливає три ключові висновки. По-перше, жоден з розглянутих методів не охоплює одночасний синтез структурної та параметричної складових. По-друге, методи з найвищою точністю (NASNet, AmoebaNet) практично недоступні для медичних дослідницьких груп через

вартість обчислень. По-третє, лише еволюційні методи демонструють придатність для малих датасетів, характерних для гістологічного аналізу.

### 1.3 Еволюційні алгоритми в машинному навчанні

#### 1.3.1 Основи генетичних алгоритмів

Генетичні алгоритми (ГА) належать до класу метаевристичних методів глобальної оптимізації, що відтворюють механізми дарвінівської еволюції: природний добір, спадковість та мутацію. На відміну від градієнтних методів, ГА не вимагають інформації про похідні цільової функції і здатні оперувати дискретними, гетерогенними просторами пошуку – саме таким є простір архітектур CNN.

Аналогія між генетичними процесами та задачею СПС нейромережових структур для класифікації гістологічних зображень наведено в таблиці 1.2.

**Таблиця 1.2** – Аналогія між генетичними процесами та задачею СПС нейромережових структур

| <i>Поняття біологічної еволюції</i> | <i>Відповідник у задачі СПС</i>                                     |
|-------------------------------------|---------------------------------------------------------------------|
| Особина (організм)                  | Конкретна конфігурація CNN (архітектура + гіперпараметри)           |
| Хромосома                           | Повний опис мережі: послідовність шарів, параметри, режим навчання  |
| Ген                                 | Окремий архітектурний або параметричний атрибут                     |
| Пристосованість (fitness)           | Якість класифікації гістологічних зображень на валідаційній вибірці |
| Покоління                           | Одна ітерація синтезу з оцінкою та відбором                         |
| Популяція                           | Множина конфігурацій-кандидатів поточного покоління                 |
| Природний добір                     | Турнірна або елітна селекція за fitness                             |



Загальна схема ГА для задачі СПС:

Ініціалізація: згенерувати  $P_0$  – початкову популяцію  
з  $N$  випадкових валідних конфігурацій CNN

Для  $t = 0, 1, \dots, T_{\max}$ :

1. Оцінити fitness кожної конфігурації в  $P_t$ :  
навчити CNN, виміряти якість на валідації
2. Перевірити критерій зупинки
3. Селекція: відібрати пари батьків з  $P_t$
4. Кросовер: породити нащадків комбінуванням батьків
5. Мутація: випадково змінити окремі атрибути нащадків
6. Сформувати  $P_{t+1}$  з найкращих батьків та нащадків

Повернути найкращу знайдену конфігурацію

Математична формалізація:

Нехай популяція на ітерації  $t$  – це:

$$\mathcal{P}_t = \{\chi_1^{(t)}, \chi_2^{(t)}, \dots, \chi_N^{(t)}\}, \quad (1.5)$$

де кожна особина  $\chi_i^{(t)} = (a_i, \varphi_i)$  – пара (структурна конфігурація, параметрична конфігурація).

Функція пристосованості відображає простір конфігурацій у скалярну якість:

$$\Phi: \mathcal{A} \times \Phi \rightarrow [0,1]. \quad (1.6)$$

Ймовірнісна селекція (пропорційна fitness) визначається як:

$$\Pr(\chi_i \text{ обрана}) = \frac{\Phi(\chi_i)}{\sum_{j=1}^N \Phi(\chi_j)}. \quad (1.7)$$

### 1.3.2 Генетичне програмування

Генетичне програмування (ГП) є узагальненням ГА на простори структур змінної довжини. Якщо ГА оперує хромосомами фіксованої довжини, ГП підтримує деревоподібні або списочні генотипи, що допускають додавання, видалення і реорганізацію вузлів у процесі еволюції, що наведено в таблиці 1.3.

**Таблиця 1.3** – Порівняння генетичного алгоритму та генетичного програмування

| <i>Характеристика</i> | <i>Генетичний алгоритм</i> | <i>Генетичне програмування</i> |
|-----------------------|----------------------------|--------------------------------|
| Довжина хромосоми     | Фіксована                  | Змінна                         |
| Структура генотипу    | Вектор                     | Дерево або список              |
| Простір пошуку        | Числовий або дискретний    | Структурний                    |
| Природне застосування | Оптимізація параметрів     | Синтез програм і структур      |

Для задачі СПС нейромережевих структур природніше застосовувати підхід ГП: кількість шарів CNN може змінюватись від покоління до покоління, що відповідає змінній довжині хромосоми. У цій роботі використовується список шарів змінної довжини як хромосома, що поєднує властивості ГА (оператори кросовера і мутації для числових атрибутів) та ГП (операції додавання і видалення шарів).

### *1.3.3 Застосування еволюційних алгоритмів в оптимізації нейронних мереж*

Еволюційні підходи до автоматизованого синтезу нейромереж мають тривалу дослідницьку історію:

- 1) NEAT (Stanley & Miikkulainen, 2002) – NeuroEvolution of Augmenting Topologies: одночасна еволюція топології і ваг нейронної мережі; розроблений для задач навчання з підкріпленням;
- 2) HyperNEAT [10] – еволюція функцій-генераторів ваг (compositional pattern producing networks) для великих мереж з регулярною структурою;
- 3) CoDeepNEAT [11] – коеволюція двох рівнів: окремих модулів (blueprints) та схем їхньої компоновки, що дозволяє синтезувати глибокі CNN з повторюваними блоками.

Застосування еволюційних методів до задач медичного комп'ютерного бачення активно розвивається останніми роками:

- 1) Samala et al. [12] застосували еволюційне стиснення (генетичний алгоритм) глибокої CNN для класифікації мас у цифровій томосинтезі молочної залози – метод скоротив кількість параметрів на 87% при збереженні AUC = 0.88-0.90;
- 2) Weng et al. [13] запропонували NAS-Unet – пошук архітектури нейронної мережі типу U-Net за допомогою диференційованих методів NAS для сегментації медичних зображень МРТ, КТ та УЗД, який перевершив ручно спроектований U-Net при значно меншій кількості параметрів (~0.8M);
- 3) Sun et al. [14] розробили метод еволюційного проектування CNN для класифікації зображень, що повністю автоматизує вибір архітектурних компонентів (тип операцій, кількість фільтрів,

глибина) і досягає конкурентних результатів порівняно з мережами ручного проектування.

Ці роботи підтверджують, що еволюційні методи СПС мають суттєву перевагу над градієнтними в умовах малих медичних датасетів та нестандартних функцій якості.

## **1.4 Задача класифікації гістопатологічних зображень в медицині**

### *1.4.1 Комп'ютерне бачення в медичній діагностиці*

Медичне комп'ютерне бачення є одним із найактивніших напрямків застосування глибокого навчання. Основні задачі включають:

Класифікацію – визначення класу патології:

- рак / норма (бінарна);
- тип та стадія злоякісного новоутворення (багатокласова).

Детекцію – локалізацію патологічних ділянок:

виявлення пухлинних вогнищ на гістологічних зрізах;  
детекція мікрокальцинатів на мамограмах.

Сегментацію – попиксельну класифікацію медичного зображення:

- сегментація пухлини від здорової тканини.
- сегментація органів на КТ/МРТ

Серед цих задач гістопатологічна класифікація зображень є особливо важливою, оскільки патогістологічний аналіз (мікроскопія тканин) залишається “золотим стандартом” онкологічної діагностики.

### 1.4.2 Особливості гістопатологічних зображень

Гістопатологічні зображення – це мікрофотографії тканинних зрізів, забарвлених за методом гематоксилін-еозин (Н&Е). Гематоксилін забарвлює ядра клітин у синьо-фіолетовий колір, еозин – цитоплазму та сполучну тканину у рожевий.

Специфіка гістопатологічних зображень:

- 1) висока інформативна щільність: одне зображення може містити тисячі клітин;
- 2) багаторівнева структура: від молекулярного до тканинного рівня;
- 3) варіабельність збільшення: різні діагностичні задачі потребують різного масштабу (40×–400×);
- 4) кольорова варіабельність: різні лабораторії та забарвлення можуть давати різні кольорові характеристики;
- 5) просторова залежність: важлива не тільки клітина, але й її оточення.

Дані представлено в таблиці 1.4.

**Таблиця 1.4** – Порівняння гістопатологічних зображень з природними

| <i>Характеристика</i> | <i>Природні зображення</i> | <i>Гістопатологічні</i>      |
|-----------------------|----------------------------|------------------------------|
| Об'єкти               | Чітко визначені            | Клітини, ядра, тканини       |
| Текстура              | Різноманітна               | Специфічна для патології     |
| Колір                 | Натуральний                | Синтетичний (Н&Е)            |
| Масштаб               | Фіксований                 | Залежить від збільшення      |
| Розмір                | Зазвичай 224×224           | 700×460 до 4096×4096         |
| Дисбаланс класів      | Помірний                   | Значний (рідкісні патології) |

### 1.4.3 *Damaset BreakHis: характеристики та особливості*

BreakHis (Breast Cancer Histopathological Image Classification) – публічний датасет гістопатологічних зображень тканин молочної залози, розроблений Spanhol et al. [15] у співпраці з P&D Лабораторією (Pathology and Cytology Laboratories), Бразилія.

Характеристики датасету:

- загальна кількість зображень: 7,909
- кількість пацієнтів: 82 (32 з доброякісними, 50 зі злоякісними пухлинами)
- зображень на пацієнта: 96–360 (різних збільшень та зон)
- розмір зображень: 700×460 пікселів, RGB
- рівні збільшення: 40×, 100×, 200×, 400×

Класи (8 підтипів):

Доброякісні пухлини (2,480 зображень):

- Adenosis (A) – аденоз
- Fibroadenoma (F) – фіброаденома
- Phyllodes Tumor (PT) – пухлина Філлодес
- Tubular Adenoma (TA) – тубулярна аденома

Злоякісні пухлини (5,429 зображень):

- Ductal Carcinoma (DC) – протоковий рак
- Lobular Carcinoma (LC) – часточковий рак
- Mucinous Carcinoma (MC) – слизовий рак
- Papillary Carcinoma (PC) – папілярний рак

Розподіл зображень наведено в таблиці 1.5.

Розбивка за діагностичними задачами: бінарна класифікація: (доброякісна (2,480) vs. злоякісна (5,429)); основна задача в цій роботі – 8-класова класифікація: ідентифікація конкретного підтипу пухлини.

**Таблиця 1.5** – Розподіл зображень BreakHis за класами та збільшеннями

| <i>Клас</i>            | <i>40×</i> | <i>100×</i> | <i>200×</i> | <i>400×</i> | <i>Всього</i> |
|------------------------|------------|-------------|-------------|-------------|---------------|
| Adenosis               | 114        | 113         | 111         | 106         | 444           |
| Fibroadenoma           | 253        | 260         | 264         | 237         | 1,014         |
| Phyllodes<br>Tumor     | 109        | 121         | 108         | 115         | 453           |
| Tubular<br>Adenoma     | 149        | 150         | 140         | 130         | 569           |
| Ductal<br>Carcinoma    | 864        | 903         | 896         | 788         | 3,451         |
| Lobular<br>Carcinoma   | 156        | 170         | 163         | 137         | 626           |
| Mucinous<br>Carcinoma  | 205        | 222         | 196         | 169         | 792           |
| Papillary<br>Carcinoma | 145        | 142         | 135         | 138         | 560           |
| Всього                 | 1,995      | 2,081       | 2,013       | 1,820       | 7,909         |

**Дисбаланс класів.** Датасет має помірний дисбаланс: 31.4% доброякісних проти 68.6% злоякісних зображень. Це відображає реальний клінічний розподіл і потребує врахування при оцінці якості синтезованих архітектур.

#### *1.4.4 Виклики класифікації гістопатологічних зображень*

1. Внутрішньокласова варіабельність: гістологічні зображення одного і того ж підтипу пухлини можуть суттєво відрізнятися через: відмінності між зразками різних пацієнтів; різні зони одного і того ж зрізу; відмінності у

якості гістологічного препарату; варіабельність забарвлення між лабораторіями.

2. Складність диференціальної діагностики: деякі підтипи є морфологічно схожими: Ductal carcinoma та Lobular carcinoma мають спільні риси. Аденоз може імітувати злоякісні зміни. Висока міжкласова подібність потребує виявлення тонких відмінностей.

3. Вплив рівня збільшення: одна і та сама тканина виглядає принципово по-різному на різних збільшеннях:

- 40×: загальна гістоархітектоніка, розташування залоз
- 100×: клітинний поліморфізм, мітотична активність
- 200×: морфологія ядер, хроматин
- 400×: деталі ядра, мікроструктура

Оптимальна архітектура для аналізу зображень різних збільшень може відрізнятися.

4. Дисбаланс класів та вартість помилок: у медичній діагностиці помилки класів мають різну вартість: False Negative (злоякісна пухлина класифікована як доброякісна): критична помилка, може призвести до запізненого лікування; False Positive (доброякісна класифікована як злоякісна): призводить до непотрібного лікування.

Це вимагає використання зважених метрик (F1, AUC-ROC, Sensitivity/Specificity) поряд зі звичайною точністю.

5. Обмеженість розміченого датасету: 7,909 зображень – значно менше порівняно з ImageNet (1.2M). Це підвищує ризик перенавчання та вимагає: правильного підбору архітектурної складності (не занадто великої), ефективної аугментації та одночасної оптимізації регуляризаційних параметрів (Dropout).



#### *1.4.5 Існуючі підходи до класифікації гістопатологічних зображень*

Підхід 1: Transfer Learning з ImageNet-предтренуваних мереж.

Використання предтренуваних мереж (ResNet, DenseNet, InceptionV3) з fine-tuning на BreakHis.

Переваги: не потребує великого специфічного датасету, швидкість впровадження.

Недоліки: домен ImageNet (природні зображення) суттєво відрізняється від Н&Е гістологічних зображень, фіксована архітектура може бути субоптимальною для гістологічних особливостей, потребує ручного підбору гіперпараметрів fine-tuning.

Репрезентативні результати: ResNet-50 (fine-tuned) досягає 87–90% на BreakHis; DenseNet-121 досягає 88–92%.

Підхід 2: Ручне проектування спеціалізованих CNN.

Дослідники проектують архітектуру спеціально під гістологічні зображення.

Переваги: можливість врахування специфіки задачі; контроль над складністю моделі.

Недоліки: трудомісткість (місяці роботи); потребує подвійної експертизи (ML + патоморфологія); субоптимальність через обмеженість людського пошуку.

Підхід 3: Multi-scale та attention підходи.

Спеціалізовані архітектури, що враховують багаторівневість гістологічних зображень.

Переваги: враховують специфіку різних збільшень; можуть фокусуватися на патологічних ділянках.

Недоліки: висока складність реалізації; потребують спеціалізованої розмітки (локалізація патологій); фіксовані архітектурні рішення без параметричної оптимізації..

Висновок щодо необхідності СПС для гістопатологічних задач.

1. Автоматичне знаходження оптимальної архітектурної складності для обмеженого датасету
2. Одночасну оптимізацію регуляризаційних параметрів (Dropout rate, batch size)
3. Виявлення оптимальних фільтрів для специфічних текстурних ознак гістологічних зображень
4. Швидку адаптацію до нових датасетів гістологічних зображень без ручного налаштування

## 1.5 Проблеми сучасних методів синтезу нейромережових структур у медичному AI

### 1.5.1 Недоступна обчислювальна вартість

Провідні методи NAS та СПС розроблялися в умовах великих корпоративних дослідницьких центрів (Google Brain, DeepMind, Facebook AI) з доступом до сотень GPU. Для медичних науково-дослідних установ та клінічних лабораторій ця реальність є принципово недосяжною.

Порівняння обчислювальних витрат наведено в таблиці 1.6.

**Таблиця 1.6** – Порівняння обчислювальних витрат методів NAS для медичних задач

| Метод           | GPU-години | Приблизна вартість | Доступність для клініки |
|-----------------|------------|--------------------|-------------------------|
| NASNet          | 43 200     | ~\$130 000         | Недоступний             |
| AmoebaNet       | 75 600     | ~\$230 000         | Недоступний             |
| DARTS           | 96         | ~\$300             | Умовно доступний        |
| ENAS            | 12         | ~\$37              | Доступний               |
| СПС (ця робота) | ~1         | ~\$3               | Доступний               |

Причиною надмірних витрат у методах першого покоління є три фактори: повне незалежне навчання кожної архітектури-кандидата на великому датасеті (ImageNet, ~1.2М зображень); необхідність кількох тисяч ітерацій до конвергенції пошуку; і розрахунок на розпаралелювання між сотнями GPU.

Для медичних застосувань фактичні датасети суттєво менші (BreakHis: 7909 зображень; типові гістологічні датасети: 1000–5000 зображень), що відкриває можливість методів СПС із часом пошуку до 1 GPU-добі.

### *1.5.2 Відтворюваність і верифікація результатів*

Відтворюваність у медичному AI – це не лише наукова вимога, але й регуляторна. Системи автоматизованого аналізу гістологічних зображень, що претендують на клінічне застосування, підпадають під вимоги стандартів якості (ISO 15189, IVD Regulation EU 2017/746), які передбачають документальне підтвердження всіх кроків розробки та верифікацію результатів.

Сучасні методи NAS мають суттєві недоліки з відтворюваності:

- 1) стохастичність без контролю: випадкова ініціалізація ваг, стохастичне семплювання в RL-контролерах та випадкова ротація датасету ускладнюють точне відтворення результатів;
- 2) неповне звітування: більшість публікацій з NAS не розкривають повний набір гіперпараметрів пошуку і звітують про “best of N runs” без статистичного аналізу варіабельності;
- 3) відсутність збережених моделей: синтезовані архітектури описуються схематично, а ваги натренованих моделей не публікуються – унеможлиблюючи незалежну верифікацію.

Метод СПС, розроблений у цій роботі, адресує цю проблему через зберігання повної історії еволюції (архітектура + навчені ваги кожної з 136 моделей), детальні логи всіх гіперпараметрів та автоматизований скрипт незалежної валідації.

### *1.5.3 Вузький простір синтезу*

Більшість методів NAS оптимізують виключно структурну складову нейромережі, залишаючи гіперпараметри навчання фіксованими (зазвичай – загальноприйняті значення, оптимальні для ImageNet, але не для конкретного медичного датасету).

Типові обмеження простору синтезу в існуючих методах:

- 1) DARTS: 8 фіксованих операцій (conv  $3\times 3$ ,  $5\times 5$ , maxpool, avgpool, identity, zero); не включає BatchNorm-варіанти, різні функції активації, Dropout;
- 2) NASNet: пошук на рівні “комірок” (cells) із заздалегідь визначеним скелетом мережі;
- 3) фіксовані гіперпараметри: learning rate = 0.025 (DARTS), batch size = 64 (AmoebaNet) – значення, неоптимальні для гістологічних датасетів обсягом  $\sim 3000$  зображень.

Для медичних задач гіперпараметри навчання мають не менше значення, ніж архітектура: один і той самий optimizer при різних learning rates може давати accuracy від 52% до 87% на BreakHis.

#### 1.5.4 Відсутність одночасного синтезу структурної і параметричної складових

Архітектура та гіперпараметри навчання не є незалежними змінними. Між ними існують нетривіальні залежності, виявлення яких є однією з цілей СПС:

- 1) глибші мережі (більше шарів) зазвичай потребують меншого темпу навчання для стабільної конвергенції;
- 2) архітектури з BatchNorm менш чутливі до вибору learning rate, тоді як мережі без нормалізації вимагають більш точного підбору;
- 3) Adam-оптимізатор значно ефективніший за SGD для гістологічних датасетів малого обсягу, де важлива швидкість конвергенції за обмежену кількість кроків.

При роздільній оптимізації – спочатку архітектура, потім гіперпараметри – знайдена архітектура  $a^*$  є оптимальною лише для фіксованого  $\varphi_0$ :

$$a^* = \operatorname{argmax}_{a \in \mathcal{A}} Q(\mathcal{M}(a, \varphi_0), \mathcal{D}_{\text{hist.}}) \quad (1.8)$$

Результат  $a^*$  не є глобальним оптимумом у спільному просторі  $\mathcal{A} \times \Phi$ .

Одночасний СПС усуває цю проблему, шукаючи пару  $(a^*, \varphi^*)$  безпосередньо в добутку просторів.

### *1.5.5 Нечутливість до асиметрії помилок у медичній діагностиці*

Цей недолік специфічний для медичних застосувань і майже не обговорюється в загальних NAS-методах. У класифікації гістологічних зображень ціна False Negative (пропущена злоякісна пухлина) принципово відрізняється від ціни False Positive (зайва біопсія).

Стандартна функція fitness на основі accuracy не враховує цієї асиметрії: модель, що класифікує 90% зображень як “злоякісні”, може мати accuracy 70% при нульовій специфічності. Метод СПС, розроблений у цій роботі, використовує комбінований критерій accuracy + AUC-ROC, що стимулює синтез архітектур з хорошим балансом чутливості і специфічності.

## **Висновки до розділу 1**

Проаналізовано задачі комп’ютерного бачення в медицині та специфіку гістопатологічного аналізу зображень. Класифікація гістологічних зображень для онкологічної діагностики є критично важливою задачею, де точність безпосередньо впливає на прогноз лікування. Особливості гістопатологічних зображень – варіабельність збільшення, дисбаланс класів, обмежений розмір датасету – визначають специфічні вимоги до нейромережових архітектур.

Розглянуто датасет BreakHis як репрезентативний публічний датасет гістологічних зображень раку молочної залози: 7,909 зображень, 8 підтипів (4 доброякісних + 4 злоякісних), 4 рівні збільшення. Виявлено ключові виклики: внутрішньокласова варіабельність, складність диференціальної діагностики, вплив рівня збільшення та дисбаланс класів.

Проаналізовано сучасні методи синтезу нейромережових структур: градієнтні методи (DARTS), методи на основі навчання з підкріпленням

(NASNet, ENAS) та еволюційні методи (AmoebaNet). Встановлено, що жоден з розглянутих методів не виконує одночасного синтезу структурної (архітектурної) та параметричної (гіперпараметричної) складових нейромережі.

Виявлено чотири ключові проблеми існуючих методів: (1) надмірно високі обчислювальні вимоги (до 3,150 GPU-днів), недоступні для медичних дослідницьких центрів; (2) низька відтворюваність результатів, критична для медичних регуляторних вимог; (3) обмежений простір синтезу без гіперпараметрів; (4) відсутність врахування специфіки медичних датасетів (дисбаланс класів, обмежений розмір).

Обґрунтовано необхідність розробки методу структурно-параметричного синтезу нейромережевих структур, який би поєднував: помірні обчислювальні вимоги ( $< 1$  GPU-дня), повну відтворюваність (збереження всіх синтезованих моделей), розширений простір синтезу (структура + параметри) та врахування особливостей датасету BreakHis (дисбаланс класів, обмежений обсяг даних).

## РОЗДІЛ 2 МАТЕМАТИЧНІ МОДЕЛІ ТА МЕТОДИ СТРУКТУРНО-ПАРАМЕТРИЧНОГО СИНТЕЗУ НЕЙРОМЕРЕЖЕВИХ СТРУКТУР

### 2.1 Основи згорткових нейронних мереж

Згорткові нейронні мережі (Convolutional Neural Networks, CNN) є провідним класом архітектур глибокого навчання для аналізу зображень. Їхня ключова властивість – локальна зв'язність та спільне використання ваг – робить їх особливо придатними для задач медичного комп'ютерного бачення, де важливо виявляти повторювані морфологічні патерни (ядра клітин, мітотичні фігури, залозисті структури) незалежно від їхнього положення на гістологічному зрізі.

На відміну від повнозв'язних мереж, що потребують мільйонів параметрів для опрацювання зображень  $224 \times 224$  пікселів, CNN виділяють просторово-ієрархічні ознаки за допомогою фільтрів, кількість параметрів яких не залежить від розміру вхідного зображення. Це критично важливо для гістологічних датасетів обмеженого обсягу, де глибокі повнозв'язні мережі неминуче перенавчаються.

#### 2.1.1 Згорткові шари

Основним обчислювальним примітивом CNN є операція дискретної крос-кореляції між вхідним гістологічним зображенням та набором навчальних фільтрів (ядер). Фільтри кожного шару навчаються виявляти певний тип локальних ознак – краї, напрями градієнтів, кольорові патерни Н&Е-зabarвлення.



Математична формалізація.

Нехай  $I \in \mathbb{R}^{H \times W \times C}$  – вхідний тензор гістологічного зображення, де  $H$  – висота,  $W$  – ширина,  $C$  – кількість каналів ( $C = 3$  для RGB). Згортковий шар з  $F$  фільтрами розміром  $k \times k$  обчислює вихідну карту ознак  $Z^{(f)} \in \mathbb{R}^{H' \times W'}$  для  $f$ -го фільтра  $K^{(f)} \in \mathbb{R}^{k \times k \times C}$ :

$$Z_{u,v}^{(f)} = \sum_{p=0}^{k-1} \sum_{q=0}^{k-1} \sum_{c=0}^{C-1} I_{u+p, v+q, c} \cdot K_{p,q,c}^{(f)} + b^{(f)}, \quad (2.1)$$

де  $b^{(f)}$  – зміщення (bias),  $(u, v)$  – просторова позиція у вихідній карті.

Роль у класифікації гістологічних зображень:

- перші шари (малоглибокі фільтри): виявляють перепади інтенсивності, границі клітин, нюанси Н&Е-забарвлення;
- середні шари: кодують локальні морфологічні структури – форму ядра, щільність клітин;
- глибокі шари: представляють тканинну архітектуру – залозисті структури, інвазивний ріст.

Параметри згорткового шару в цій роботі продемонстровано в таблиці 2.1.

**Таблиця 2.1** – Параметри згорткового шару в цій роботі

| Параметр           | Значення     | Обґрунтування                       |
|--------------------|--------------|-------------------------------------|
| Кількість фільтрів | {32,64,128}  | Підбирається СПС для кожного шару   |
| Розмір ядра        | $3 \times 3$ | Стандарт для гістологічних CNN      |
| Stride             | 1            | Збереження просторової роздільності |
| Padding            | same         | Вихід = вхід за розміром            |

При  $\text{padding} = \text{same}$  та  $\text{stride} = 1$  вхідне гістологічне зображення  $224 \times 224 \times 3$  після шару Conv2D(32) дає карту ознак  $224 \times 224 \times 32$ .

### 2.1.2 Pooling-шари (підвибірка)

Шари підвибірки (pooling) зменшують просторові розміри карт ознак, зберігаючи найбільш інформативні значення. Це виконує дві ролі: зменшує обчислювальну складність наступних шарів та забезпечує часткову інваріантність до малих просторових зміщень – корисну властивість при аналізі гістологічних зрізів, де одні й ті самі структури можуть зустрічатися у різних положеннях.

Max Pooling  $2 \times 2$  (операція використана у цій роботі):

$$Z_{u,v}^{\text{pool}} = \max_{\substack{p \in \{0,1\} \\ q \in \{0,1\}}} Z_{2u+p, 2v+q}. \quad (2.2)$$

Кожен MaxPool шар вдвічі зменшує просторові розміри:  $224 \rightarrow 112 \rightarrow 56 \rightarrow 28 \rightarrow 14$ .

Обмеження кількості MaxPool для зображень  $224 \times 224$ : у просторі синтезу допускається до 4 шарів MaxPool. Після чотирьох операцій карти ознак мають розмір  $14 \times 14$ , що достатньо для опрацювання морфологічних структур; п'ятий MaxPool скоротив би їх до  $7 \times 7$ , що обмежило б ємність ознак для складних гістологічних паттернів.

### 2.1.3 Batch Normalization

Batch Normalization (Ioffe & Szegedy, 2015) нормалізує активації кожного шару за статистиками поточного міні-батчу. Для вектора активацій  $x$  над міні-батчем  $\mathcal{B} = \{x_1, \dots, x_m\}$ :

Крок 1. Обчислення статистик:

$$\mu_B = \frac{1}{m} \sum_{i=1}^m x_i, \sigma_B^2 = \frac{1}{m} \sum_{i=1}^m (x_i - \mu_B)^2. \quad (2.3)$$

Крок 2. Нормалізація:

$$\hat{x}_i = \frac{x_i - \mu_B}{\sqrt{\sigma_B^2 + \varepsilon}}, \quad \varepsilon = 10^{-5}. \quad (2.4)$$

Крок 3. Аффінне перетворення з навчальними параметрами  $\gamma, \beta$ :

$$y_i = \gamma \hat{x}_i + \beta. \quad (2.5)$$

Для задачі класифікації гістологічних зображень BatchNorm відіграє важливу роль: гістологічні зображення з різних лабораторій або при різних умовах забарвлення мають відмінні колірні характеристики. Нормалізація активацій зменшує чутливість моделі до цих варіацій.

У просторі синтезу BatchNorm дозволений лише після Conv2D або Dense-шарів, оскільки лише в цих позиціях статистики батчу є статистично стабільними.

#### 2.1.4 Dropout

Dropout [16] – це техніка регуляризації, що запобігає перенавчанню шляхом випадкового “вимкнення” нейронів під час кожного кроку тренування. Для нейрона з активацією  $h_i$  визначається бінарна маска:

$$r_i \sim \text{Bernoulli}(1 - p_{\text{drop}}), \tilde{h}_i = r_i \cdot h_i, \quad (2.6)$$

де  $p_{\text{drop}} \in [0.2, 0.5]$  – ймовірність вимкнення нейрона.

Під час інференсу всі нейрони активні. TensorFlow реалізує інвертований dropout: на етапі тренування активації масштабуються на  $1/(1 - p_{\text{drop}})$ , що усуває необхідність масштабування при тестуванні:

$$\tilde{h}_i^{\text{train}} = \frac{r_i \cdot h_i}{1 - p_{\text{drop}}}. \quad (2.7)$$

Для гістологічних датасетів обмеженого обсягу Dropout є особливо цінним: він діє як неявне ансамблювання, тренуючи  $2^n$  різних підмереж на кожному кроці. Це суттєво знижує ризик перенавчання при малих (~3000 зображень) навчальних вибірках.

В цій роботі  $\text{rate} \in [0.20, 0.50]$  підбирається генетичним алгоритмом.

### 2.1.5 Функції активації

Функції активації вводять нелінійність у CNN, дозволяючи формувати складні нелінійні рішення в просторі ознак. Для задачі класифікації гістологічних зображень вибір функції активації впливає як на швидкість конвергенції, так і на якість виявлення морфологічних ознак.

ReLU (Rectified Linear Unit):

Найпопулярніша функція активації в глибокому навчанні, визначена як:

$$\text{ReLU}(x) = \max(0, x) = \begin{cases} x, & \text{якщо } x \geq 0 \\ 0, & \text{якщо } x < 0. \end{cases} \quad (2.8)$$

Градiєнт ReLU:

$$\frac{\partial \text{ReLU}(x)}{\partial x} = \begin{cases} 1, & \text{якщо } x > 0 \\ 0, & \text{якщо } x \leq 0. \end{cases} \quad (2.9)$$

Переваги ReLU:

- обчислювальна простота (порівняння з нулем);
- відсутність проблеми зникаючих градієнтів для додатних значень;
- емпірично показує кращі результати ніж sigmoid або tan.

Недоліки ReLU:

- проблема “dying ReLU”: якщо нейрон постійно отримує від’ємні значення, його градієнт завжди нуль, і він перестає навчатися.

ELU (Exponential Linear Unit):

Покращена версія ReLU, яка має ненульовий градієнт для від’ємних значень:

$$\text{ELU}(x) = \begin{cases} x, & \text{якщо } x \geq 0 \\ \alpha(\exp(x) - 1), & \text{якщо } x < 0 \end{cases} \quad (2.10)$$

де  $\alpha$  – гіперпараметр (зазвичай  $\alpha = 1.0$ ).

Переваги ELU:

- вирішує проблему “dying ReLU”;
- середнє значення активацій ближче до нуля, що прискорює навчання;
- більш гладка функція.

Недоліки ELU:

- дещо повільніше обчислення через експоненту;
- потребує більше обчислювальних ресурсів.

Sigmoid для бінарної класифікації:

Для вихідного нейрону в задачах бінарної класифікації (доброякісна / злоякісна) використовується сигмоїд:

$$\sigma(x) = \frac{1}{1+e^{-x}} \in (0,1). \quad (2.11)$$

Сигмоїд інтерпретується як ймовірність належності зображення до злоякісного класу:  $\hat{y} = \sigma(z)$ , де  $z$  – вихід останнього прихованого шару. Порогове значення  $\hat{y} \geq 0.5$  відповідає класифікації як “злоякісна”.

Вибір функції активації в цій роботі: генетичний алгоритм обирає між ReLU та ELU для проміжних шарів; вихідний Dense-шар завжди використовує sigmoid. Таке рішення відповідає задачі бінарної класифікації BreakHis та дозволяє компілювати модель з функцією втрат `binary_crossentropy`, що забезпечує коректний градієнт для задачі виявлення злоякісних пухлин.

## 2.2 Математична модель генетичного алгоритму для структурно-параметричного синтезу

Генетичний алгоритм (ГА) у контексті задачі СПС нейромережевих структур є ітеративним методом направленого пошуку в комбінованому просторі архітектурних і параметричних конфігурацій. Алгоритм підтримує популяцію  $\mathcal{P} = \{\chi_1, \chi_2, \dots, \chi_N\}$  із  $N$  особин, кожна з яких кодує повну специфікацію CNN. Через послідовне застосування операторів відбору, рекомбінації та мутації популяція еволюціонує до конфігурацій з вищою якістю класифікації гістологічних зображень.

### 2.2.1 Формалізація простору синтезу

Простір синтезу  $\mathcal{S}$  є декартовим добутком простору архітектур і простору гіперпараметрів:

$$\mathcal{S} = \mathcal{A} \times \Phi. \quad (2.12)$$

Простір архітектур  $\mathcal{A}$ :

Архітектура визначається як скінченна послідовність шарів змінної довжини  $n$ :

$$\mathcal{A} = \{(l_1, l_2, \dots, l_n) \mid n \in [n_{\min}, n_{\max}], l_i \in \mathcal{T}\}, \quad (2.13)$$

де  $\mathcal{T}$  – множина допустимих типів шарів:

$$\mathcal{T} = \{\text{Conv2D}, \text{MaxPool}, \text{BatchNorm}, \text{Dropout}, \text{Flatten}, \text{Dense}\}. \quad (2.14)$$

Атрибути шарів зазначено в таблиці 2.2.

**Таблиця 2.2** – Атрибути типів шарів та їх область значень

| <i>Тип шару</i>  | <i>Параметри</i>           | <i>Область значень</i>                              |
|------------------|----------------------------|-----------------------------------------------------|
| <i>Conv2D</i>    | <i>filters, activation</i> | $\{32, 64, 128\}, \{\text{relu}, \text{elu}\}$      |
| <i>MaxPool</i>   | <i>pool_size</i>           | $\{2\}$                                             |
| <i>Dense</i>     | <i>neurons, activation</i> | $\{32, 64, 128, 256\}, \{\text{relu}, \text{elu}\}$ |
| <i>Dropout</i>   | <i>rate</i>                | $[0.20, 0.50]$                                      |
| <i>BatchNorm</i> | -                          | -                                                   |
| <i>Flatten</i>   | -                          | -                                                   |

Простір гіперпараметрів  $\Phi$ :

$$\Phi = \mathcal{O} \times \Lambda \times \mathcal{B}. \quad (2.15),$$

$$\text{де } \mathcal{O} = \{\text{adam,sgd,rmsprop}\}, \quad \Lambda = [10^{-4}, 10^{-3}], \quad \mathcal{B} = \{16, 32\}. \quad (2.16)$$

Структурні обмеження валідності архітектур для гістологічних зображень  $224 \times 224$ :

- 1) перший шар завжди Conv2D (необхідний для виділення ознак зображення);
- 2) BatchNorm допустимий лише після Conv2D або Dense;
- 3) після Flatten допускаються виключно Dense та Dropout;
- 4) кількість MaxPool-шарів  $\leq 4$  (обмеження просторового розміру:  $224 \rightarrow 14$ );
- 5) не більше двох однотипних сервісних шарів (BatchNorm, MaxPool, Dropout) підряд;
- 6) вихідний шар – Dense(1, sigmoid) для бінарної класифікації.

Розмір простору синтезу.

Без урахування обмежень валідності верхня оцінка:

$$|\mathcal{S}| \approx |\mathcal{A}| \times |\Phi| \gtrsim 10^{11}. \quad (2.17)$$

Навіть після фільтрації невалідних конфігурацій простір залишається  $\sim 10^8$ , що унеможливорює повний перебір і обґрунтовує еволюційний підхід.



### 2.2.2 Кодування хромосоми

Хромосома  $\chi$  – центральна структура даних ГА – кодує повну специфікацію CNN: як топологію мережі, так і режим навчання. Таке об'єднане представлення уможливлює спільну оптимізацію структурних і параметричних складових.

Формальна структура хромосоми:

$$\chi = (\mathcal{L}, \varphi). \quad (2.18)$$

де  $\mathcal{L} = [\ell_1, \ell_2, \dots, \ell_n]$  – список шарів (структурна складова);  $\varphi = (o, \lambda, b)$  – гіперпараметри навчання (параметрична складова): оптимізатор, темп навчання, розмір батчу.

Кожен елемент  $\ell_i \in \mathcal{L}$  є словником атрибутів:

$$\ell_i = \{\text{type: } t_i, \text{attr}_1: v_1, \text{attr}_2: v_2, \dots\}. \quad (2.19)$$

Приклад кодування хромосоми (реальна архітектура з Pareto-фронту NSGA-II):

$$\chi^* = \left( \underbrace{[\text{Conv2D}(32, \text{elu}), \text{BN}, \text{Conv2D}(32, \text{elu}), \text{BN}, \text{MP}, \text{MP}, \text{MP}, \dots, \text{Dense}(1, \sigma)]}_{\mathcal{L}, 16 \text{ шарів}}, \underbrace{(\text{adam}, 6.05 \cdot 10^{-4}, 32)}_{\varphi} \right), \quad (2.20)$$

де BN = BatchNorm, MP = MaxPool,  $\sigma$  = sigmoid. Ця хромосома відповідає найкращій Pareto-оптимальній архітектурі, що досягла accuracy=85.67% при 147K параметрів.

Переваги об'єднаного кодування.

Розміщення структури та гіперпараметрів в одному генотипі дозволяє ГА виявляти залежності між ними: наприклад, еволюційний процес може “навчитись”, що глибокі конволюційні блоки ( $\text{Conv} \times 2 \rightarrow \text{MaxPool}$ ) ефективніші з меншим learning rate, ніж мілкі архітектури. Такі закономірності недосяжні при роздільній оптимізації.

### 2.2.3 Функція пристосованості

Функція пристосованості (fitness function)  $f: \mathcal{S} \rightarrow \mathbb{R}$  є ключовим компонентом генетичного алгоритму, що визначає якість кожної особини (архітектури) та направляє процес еволюції. Правильний вибір функції пристосованості критично важливий для успішного пошуку оптимальних архітектур.

Визначення.

У цьому дослідженні функція пристосованості визначається як максимальна точність класифікації на валідаційному наборі, досягнута протягом усього процесу тренування:

$$f(c) = \max_{i=1}^n \text{acc}_{\text{val}}^{(i)}, \quad (2.21)$$

де  $c$  – хромосома (архітектура + гіперпараметри) –  $\text{acc}_{\text{val}}^{(i)}$  – точність на валідаційному наборі після  $i$ -ої епохи –  $n$  – загальна кількість епох (визначається адаптивно через EarlyStopping)

Обґрунтування вибору максимальної точності.

Використання максимальної, а не фінальної точності, обумовлено фундаментальною проблемою перенавчання (overfitting). Під час тренування

модель може досягти найкращої точності на валідаційному наборі на проміжній епосі, після чого точність може знижуватися через перенавчання на тренувальних даних.

Приклад динаміки навчання наведено в таблиці 2.3.

**Таблиця 2.3** – Динаміка навчання CNN на підмножині BreakHis (224x224, бінарна класифікація)

| <i>Епоха</i> | <i>Train Accuracy</i> | <i>Validation Accuracy</i> | <i>Примітка</i>       |
|--------------|-----------------------|----------------------------|-----------------------|
| 1            | 0.57                  | 0.61                       | Початкове навчання    |
| 8            | 0.74                  | 0.79                       | Активне навчання      |
| 17           | 0.83                  | 0.84                       | ← найкраща епоха      |
| 32           | 0.92                  | 0.76                       | Початок перенавчання  |
| 50           | 0.97                  | 0.69                       | Виражене перенавчання |

Як видно, модель продовжує покращувати точність на тренувальному наборі ( $0.57 \rightarrow 0.97$ ), але точність на валідаційному наборі спочатку зростає ( $0.61 \rightarrow 0.84$ ), а потім знижується ( $0.84 \rightarrow 0.69$ ). В такій ситуації:

- фінальна точність (епоха 50):  $\Phi(\chi) = 0.69$  – недооцінює справжню здатність моделі;
- максимальна точність (епоха 17):  $\Phi(\chi) = 0.84$  – коректно відображає потенціал архітектури.

Порівняння з альтернативними визначеннями fitness:

У літературі з NAS існують різні підходи до визначення fitness. Розглянемо основні альтернативи.

1. Final epoch accuracy:

$$f_{\text{final}}(c) = \text{acc}_{\text{val}}^{(n)}. \quad (2.22)$$

Недоліки: підлягає overfitting, особливо для моделей з великою capacity; може привести до неправильного ранжування архітектур; не відображає справжній потенціал моделі

## 2. Mean of top-k epochs:

$$f_{\text{mean-k}}(c) = \frac{1}{k} \sum_{i=1}^k \text{acc}_{\text{val}}^{(i_{\text{top}})}, \quad (2.23)$$

де  $i_{\text{top}}$  – індекси  $k$  кращих епох.

Переваги: більш стабільна оцінка, зменшує variance; менш чутлива до випадкових “lucky epochs”

Недоліки: може недооцінювати моделі з швидкою конвергенцією; мозмиває сигнал від дійсно кращих архітектур; потребує вибору гіперпараметра  $k$

## 3. Maximum validation accuracy (обраний підхід):

$$f_{\text{max}}(c) = \max_{i=1}^n \text{acc}_{\text{val}}^{(i)}. \quad (2.24)$$

Переваги: захист від overfitting: не караємо модель за погіршення після піку; природна сумісність з EarlyStopping (restore\_best\_weights=True); відображає справжній потенціал архітектури; узгоджується з best practices в deep learning

Потенційний ризик: чутливість до “lucky epochs” (випадково високі значення)

Мітигація ризику: earlyStopping з patience=3 та min\_delta=0.001 фільтрує випадкові флуктуації; популяційний підхід GA усереднює noise через множинні архітектури; елітизм зберігає найкращі моделі між поколіннями

Математична формалізація процесу.

Нехай  $M(c)$  – модель, побудована з хромосоми  $c$ . Процес обчислення  $fitness$  включає:

- 1) побудова моделі:  $\mathcal{M}(\chi): \mathbb{R}^{224 \times 224 \times 3} \rightarrow [0,1]$  (sigmoid-вихід для бінарної класифікації);
- 2) компіляція: з обраними  $o$  (optimizer),  $\lambda$  (learning rate), функція втрат `binary_crossentropy`;
- 3) тренування: на датасеті  $\mathcal{D}_{hist}^{train}$  з валідацією на;  $\mathcal{D}_{hist}^{val}$
- 4) моніторинг: зберігається максимальна `val_accuracy` за всі епохи;
- 5) пристосованість:  $\Phi(\chi) = \max_i \text{acc}_{val}^{(i)}$

Формально:

$$\Phi(\chi) = \max_{i \in [1, E]} \text{acc}_{val}(\mathcal{M}(\chi), \mathcal{D}_{hist}^{val}, i), \quad (2.25)$$

де  $E$  – кількість прожитих епох.

Протокол тренування:

Кожна архітектура тренується з адаптивною ранньою зупинкою:

- `earlyStopping`: `patience=3` епохи, `min_delta=0.001` – дозволяє зупинити слабкі архітектури вже через 5–10 епох;
- `reduceLROnPlateau`: автоматичне зниження  $\lambda$  при стагнації `val_loss` – особливо корисно для гістологічних датасетів з шумовими мітками;
- `restore_best_weights=True`: оцінюються ваги найкращої епохи, а не фінальної.

Такий протокол дозволяє коректно порівнювати архітектури різної глибини: мілкі конфігурації зупиняються рано, а глибокі отримують більше часу на конвергенцію.

Властивості функції пристосованості:

- 1) діапазон:  $\Phi(\chi) \in [0,1]$  ; базовий рівень  $\approx 0.57$  (мажоритарний клас у BreakHis – злоякісний);
- 2) медична осмисленість:  $\Phi(\chi)$  безпосередньо відповідає частці правильно класифікованих гістологічних знімків на валідаційній вибірці;
- 3) детермінованість: при фіксованих random seeds результат відтворюваний;
- 4) обчислювальна складність:
  - $O(E \times |\mathcal{D}_{\text{hist}}^{\text{train}}|/b)$ ;
  - $E$  – кількість епох (5–50, адаптивно через EarlyStopping);
  - $|\mathcal{D}_{\text{hist}}^{\text{train}}| \approx 2400$  гістологічних знімків;
  - $b$  – розмір батчу.

#### 2.2.4 Генетичні оператори

Генетичні оператори забезпечують еволюційний механізм пошуку оптимальних архітектур через імітацію природних процесів відбору, рекомбінації та мутації. Правильний баланс між exploitation (використання знайдених хороших рішень) та exploration (дослідження нових областей простору пошуку) є критичним для успіху генетичного алгоритму.

Основні оператори:

Запропонована система використовує чотири основні оператори:

- 1) селекція (Selection): вибір батьківських особин для розмноження;
- 2) кросовер (Crossover): комбінування генетичного матеріалу батьків;

- 3) мутація (Mutation): випадкові зміни для підтримання різноманітності;
- 4) елітизм (Elitism): збереження найкращих рішень.

Загальна схема роботи GA.

Алгоритм працює так:

### Алгоритм 2.1 – Еволюційний синтез архітектур CNN (СПС)

Вхід:  $N$  – обсяг пулу,  $G$  – число поколінь,  $P_c$ ,  $P_m$ ,  $K_e$  – кількість елітних

Вихід: оптимальна хромосома  $\chi^*$

```

1:   $P_0 \leftarrow \text{СтворитиПул}(N)$            //  $N$  валідних випадкових кодувань CNN
2:  для кожного  $\chi \in P_0$ :
3:       $\Phi(\chi) \leftarrow \text{Оцінити}(\chi)$        // навчання CNN на  $D_{\text{hist\_val}}$ ,
        фіксуємо  $\text{max\_acc\_val}$ 
4:
5:  для  $g = 1, 2, \dots, G$ :
6:       $P_{\text{нові}} \leftarrow \emptyset$ 
7:
8:      // Зберегти  $K_e$  найкращих рішень без змін
9:      еліта  $\leftarrow \text{ВибратиНайкращих}(P_{g-1}, K_e)$ 
10:      $P_{\text{нові}} \leftarrow P_{\text{нові}} \cup \text{еліта}$ 
11:
12:     // Побудова нащадків до заповнення пулу
13:     доки  $|P_{\text{нові}}| < N$ :
14:          $\alpha \leftarrow \text{Турнір}(P_{g-1})$ 
15:          $\beta \leftarrow \text{Турнір}(P_{g-1})$ 
16:
17:         якщо  $U(0,1) < P_c$ :
18:              $\gamma, \delta \leftarrow \text{Схрестити}(\alpha, \beta)$ 
19:         інакше:
20:              $\gamma, \delta \leftarrow \alpha, \beta$            // без рекомбінації
21:
22:          $\gamma \leftarrow \text{Мутувати}(\gamma, P_m)$ ;    $\delta \leftarrow \text{Мутувати}(\delta, P_m)$ 

```

```

23:       $\gamma \leftarrow \text{ПеревіритиВиправити}(\gamma); \delta \leftarrow \text{ПеревіритиВиправити}(\delta)$ 
24:       $\Pi_{\text{нові}} \leftarrow \Pi_{\text{нові}} \cup \{\gamma, \delta\}$ 
25:
26:       $\Pi_{\text{нові}} \leftarrow \Pi_{\text{нові}}[: N]$       // обрізати до розміру N
27:
28:      // Обчислювати fitness лише для нових, ще не оцінених
29:      для  $\chi \in \Pi_{\text{нові}}$  :  $\Phi(\chi) = \text{undef}$ :
30:           $\Phi(\chi) \leftarrow \text{Оцінити}(\chi)$ 
31:
32:       $\Pi_g \leftarrow \Pi_{\text{нові}}$ 
33:
34:       $\chi^* \leftarrow \text{найкраще } \chi \text{ з } \Pi_g \text{ за } \Phi$ 
35:      повернути  $\chi^*$ 

```

Параметри алгоритму:

В експериментах використовувались такі значення:

$$\begin{aligned}
 \text{population\_size} &= 8 \\
 \text{num\_generations} &= 16 \\
 P_c &= 0.7 (\text{ймовірність кросоверу}), \\
 P_m &= 0.4 (\text{ймовірність мутації}) \\
 \text{elite\_size} &= 1
 \end{aligned} \tag{2.26}$$

Детальний опис кожного оператора наведено в наступних підрозділах 2.3-2.6.

## 2.3 Селекція

Селекція визначає, які хромосоми з поточної популяції беруть участь у породженні нащадків наступного покоління. Через механізм відбору алгоритм реалізує принцип “виживання найпристосованіших”: архітектури з



вищою точністю класифікації гістологічних зображень отримують більшу ймовірність передати свої структурні та параметричні риси нащадкам.

Турнірна селекція (Tournament Selection):

У цій роботі реалізовано турнірну селекцію з розміром турніру  $k = 2$ . Метод не потребує знання глобального розподілу fitness і забезпечує стабільний селекційний тиск навіть при малих популяціях ( $N = 8$ ).

### Алгоритм 2.2 – Відбір методом турніру

Вхід:  $\mathcal{P} = \{\chi_1, \dots, \chi_N\}$  – поточний пул,  $k$  – розмір змагання

Вихід: хромосома-предок  $\chi_\alpha$

- 1: Змаг  $\leftarrow \emptyset$
- 2: повторити  $k$  разів:
- 3:     Змаг  $\leftarrow \text{Змаг} \cup \{\text{довільна } \chi \text{ з } \mathcal{P}\}$
- 4:
- 5:  $\chi_\alpha \leftarrow \chi$  з найвищим  $\Phi(\chi)$  серед Змаг
- 6: повернути  $\chi_\alpha$

Математична формалізація:

$$\chi_{\text{батько}} = \operatorname{argmax}_{\chi \in \mathcal{T}} \Phi(\chi), \mathcal{T} \subset \mathcal{P}, \quad |\mathcal{T}| = k, \quad (2.27),$$

де  $\mathcal{P} = \{\chi_1, \dots, \chi_N\}$  – поточна популяція;  $\Phi(\chi_i)$  – пристосованість  $i$ -ї хромосоми (ассигасу на валідаційній вибірці  $\mathcal{D}_{\text{hist}}^{\text{val}}$ ).

Апроксимація ймовірності відбору для  $k = 2$ :

$$\Pr(\chi_r \text{ обрана}) \approx \frac{2(N-r+1)}{N(N+1)}, \quad (2.28)$$

де  $r$  – ранг хромосоми ( $1 = \text{найкраща}$ ),  $N$  – розмір популяції.

Ілюстративний приклад наведено в таблиці 2.4 (популяція  $N = 8$ , реалістичні значення ассигасу на BreakHis)

**Таблиця 2.4** – Ілюстративний приклад ранжування хромосом у популяції  $N=8$

| Хромосома | $\Phi(\chi)$ | Ранг |
|-----------|--------------|------|
| $\chi_1$  | 0.886        | 1    |
| $\chi_2$  | 0.872        | 2    |
| $\chi_3$  | 0.859        | 3    |
| $\chi_4$  | 0.841        | 4    |
| $\chi_5$  | 0.827        | 5    |
| $\chi_6$  | 0.793        | 6    |
| $\chi_7$  | 0.761        | 7    |
| $\chi_8$  | 0.712        | 8    |

Турнір між  $\chi_4$  (0.841) та  $\chi_7$  (0.761): обирається  $\chi_4$ .

Переваги турнірної селекції для задачі СПС:

- 1) не потребує нормалізації  $\Phi$  – fitness може бути в будь-якому діапазоні;
- 2) параметр  $k$  контролює “тиск”: при  $k = 2$  зберігається різноманіття хромосом, що критично для уникнення передчасної збіжності в малій популяції;
- 3) обчислювальна складність  $O(k)$  не залежить від розміру популяції.
- 4) Порівняння з рулеточною селекцією:

Рулеточна (fitness-proportionate) селекція обирає  $\chi_i$  з імовірністю:

$$\Pr(\chi_i) = \frac{\Phi(\chi_i)}{\sum_{j=1}^N \Phi(\chi_j)}. \quad (2.29)$$

Для задачі класифікації гістологічних зображень, де значення fitness концентруються у вузькому діапазоні (0.75–0.90), рулеткова селекція

фактично вироджується у випадковий вибір. Турнірна селекція позбавлена цього недоліку.

## 2.4 Кросовер

Кросовер (схрещування, crossover) – це генетичний оператор, який комбінує генетичний матеріал двох батьківських особин для створення нащадків. Кросовер є основним механізмом exploitation в генетичних алгоритмах, дозволяючи комбінувати успішні “будівельні блоки” з різних рішень.

У цій роботі використовується гібридний підхід, який окремо обробляє структурну частину (архітектуру) та параметричну частину (гіперпараметри) хромосоми.

### 2.4.1 Структурний кросовер

Для обміну шарами між хромосомами застосовується одноточковий кросовер з роздільними точками розрізу для хромосом різної довжини.

#### Алгоритм 2.3 – Одноточковий обмін блоками

Вхід:  $\alpha, \beta$  – пара предкових хромосом

Вихід:  $\gamma, \delta$  – пара нащадків

```

1:  $B_{\alpha} \leftarrow \alpha.\text{блоки}[:-1]$  // виключаємо вихідний Dense(1,  $\sigma$ )
2:  $B_{\beta} \leftarrow \beta.\text{блоки}[:-1]$ 
3:  $r_{\alpha} \leftarrow \text{ціле з діапазону } [1, |B_{\alpha}| - 1]$  // точка розрізу в  $\alpha$ 
4:  $r_{\beta} \leftarrow \text{ціле з діапазону } [1, |B_{\beta}| - 1]$  // точка розрізу в  $\beta$ 

```

```

5:  $V_\gamma \leftarrow V_\alpha[: r_\alpha] ++ V_\beta[r_\beta :] ++ [\text{Dense}(1, \sigma)]$ 
6:  $V_\delta \leftarrow V_\beta[: r_\beta] ++ V_\alpha[r_\alpha :] ++ [\text{Dense}(1, \sigma)]$ 
7:  $\gamma.\text{блоки} \leftarrow V_\gamma; \quad \delta.\text{блоки} \leftarrow V_\delta$ 
8: повернути  $\gamma, \delta$ 

```

Формалізація.

Для хромосом-батьків  $\chi_1, \chi_2$  із списками шарів  $\mathcal{L}^{(1)} = [\ell_1^{(1)}, \dots, \ell_n^{(1)}]$  та  $\mathcal{L}^{(2)} = [\ell_1^{(2)}, \dots, \ell_m^{(2)}]$  (без вихідного шару), при точках  $p_1 \in [1, n - 1]$ ,  $p_2 \in [1, m - 1]$ :

$$\begin{aligned} \mathcal{L}^{(\chi_3)} &= [\ell_1^{(1)}, \dots, \ell_{p_1}^{(1)}, \ell_{p_2+1}^{(2)}, \dots, \ell_m^{(2)}, \ell_{\text{out}}] \\ \mathcal{L}^{(\chi_4)} &= [\ell_1^{(2)}, \dots, \ell_{p_2}^{(2)}, \ell_{p_1+1}^{(1)}, \dots, \ell_n^{(1)}, \ell_{\text{out}}] \end{aligned} \quad (2.30)$$

де  $\ell_{\text{out}} = \text{Dense}(1, \sigma)$  – обов’язковий вихідний шар для бінарної класифікації BreakHis.

Приклад (реалістичні гістологічні архітектури):

$\chi_1$ : [Conv2D(64,elu), MaxPool, MaxPool, Flatten, Dense(128,relu), Dropout(0.4)]

$\chi_2$ : [Conv2D(32,relu), BN, MaxPool, Flatten, Dense(64,relu)]

$p_1=3$  (після 2-го MaxPool у  $\chi_1$ ),  $p_2=3$  (після Flatten у  $\chi_2$ ):

$\chi_3$ : [Conv2D(64,elu), MaxPool, MaxPool, Flatten, Dense(64,relu), Dense(1, $\sigma$ )]

$\chi_4$ : [Conv2D(32,relu), BN, MaxPool, Dense(128,relu), Dropout(0.4), Dense(1, $\sigma$ )]

### 2.4.2 Параметричний кросовер

Гіперпараметри рекомбінуються залежно від типу.

Темп навчання (неперервний) – арифметичне середнє між батьками:

$$\lambda_{\text{нащадок}} = \frac{\lambda_{\chi_1} + \lambda_{\chi_2}}{2}. \quad (2.31)$$

Розмір батчу та оптимізатор (дискретні/категоріальні) –  
рівноймовірний вибір:

$$b_{\text{нащадок}} \sim \text{Uniform}(\{b_{\chi_1}, b_{\chi_2}\}), o_{\text{нащадок}} \sim \text{Uniform}(\{o_{\chi_1}, o_{\chi_2}\}). \quad (2.32)$$

Повний приклад:

$\chi_1$ : arch=[Conv2D(64,elu), MaxPool, ...],  $\lambda=0.0002$ , b=16, opt=adam

$\chi_2$ : arch=[Conv2D(32,relu), BN, ...],  $\lambda=0.0009$ , b=32, opt=rmsprop

Нащадки:

$\chi_3$ : arch=hybrid,  $\lambda=0.00055$  (серед.), b=32 (від  $\chi_2$ ), opt=adam (від  $\chi_1$ )

$\chi_4$ : arch=hybrid,  $\lambda=0.00055$  (серед.), b=16 (від  $\chi_1$ ), opt=rmsprop (від  $\chi_2$ )

Ймовірність кросоверу  $P_c = 0.7$ :

$$(\chi_3, \chi_4) = \begin{cases} \text{Crossover}(\chi_1, \chi_2), & U(0,1) < P_c \\ (\chi_1, \chi_2), & \text{інакше} \end{cases}. \quad (2.33)$$

При відсутності кросоверу (30% пар) нащадки є точними копіями батьків – це зберігає найкращі знайдені архітектури.

Постобробка нащадків:

$$\chi'_i \leftarrow \text{ValidateAndFix}(\chi_i), \quad i \in \{3,4\}. \quad (2.34)$$

Функція `ValidateAndFix()` перевіряє структурні обмеження та виправляє порушення. Наприклад, якщо після `Flatten` з'явився `Conv2D` (неможливо в Keras), він замінюється на `Dense`.

## 2.5 Мутація

Мутація є основним механізмом *exploration* (дослідження) в генетичному алгоритмі – вона вносить випадкові збурення в хромосому, що дозволяє виходити з локальних оптимумів і виявляти якісно нові архітектурні конфігурації CNN для класифікації гістологічних зображень.

Двохтипна мутація:

Рівноймовірно застосовується один із двох типів:

$$\text{Тип мутації} = \begin{cases} \text{структурна,} & \text{з імовірністю } 0.5 \\ \text{параметрична,} & \text{з імовірністю } 0.5 \end{cases} \quad (2.35)$$

### 2.5.1 Структурна мутація

Структурна мутація змінює топологію CNN через три операції: `ADD` (вставка шару), `REMOVE` (видалення шару), `MODIFY` (зміна параметрів шару).

Алгоритм.

**Алгоритм 2.4** – Структурна зміна хромосоми

Вхід:  $\chi$  – вихідна хромосома

Вихід:  $\chi_m$  – змінена хромосома

```

1: дія ← вибрати_рівномовірно({ВСТАВИТИ, ВИЛУЧИТИ, ЗМІНИТИ})
2:
3: якщо дія = ВСТАВИТИ і  $|\chi.блоки| < MAX\_БЛОКІВ$ :
4:      $j \leftarrow$  випадкова позиція в  $[1, |\chi.блоки| - 1]$ 
5:     тип_контексту ← ВизначитиКонтекст( $\chi.блоки$ ,  $j$ ) // до чи після
    Flatten
6:     блок_new ← СгенеруватиБлок(тип_контексту)
7:      $\chi_m.блоки \leftarrow \chi.блоки[: j] ++ [блок\_new] ++ \chi.блоки[j :]$ 
8:
9: якщо дія = ВИЛУЧИТИ і  $|\chi.блоки| > MIN\_БЛОКІВ$ :
10:    доступні ← БлокиБезFlatten( $\chi.блоки$ ) // Flatten та вихід
    захищені
11:    блок_del ← довільний елемент доступні
12:     $\chi_m.блоки \leftarrow \chi.блоки \setminus \{блок\_del\}$ 
13:
14: якщо дія = ЗМІНИТИ:
15:    змінювані ← ОтриматиЗмінювані( $\chi.блоки$ )
16:    блок_sel ← довільний елемент змінювані
17:    МодифікуватиПараметри(блок_sel)
18:
19:  $\chi_m \leftarrow$  ПеревіритиВиправити( $\chi_m$ )
20: повернути  $\chi_m$ 

```

1. Операція ADD (додавання шару):

Математично:

$$\mathcal{L}' = [L_1, \dots, L_{i-1}, L_{new}, L_i, \dots, L_n], \quad i \in [1, n], \quad (2.36)$$

де  $L_{new}$  – новий випадково згенерований шар, тип якого залежить від контексту (до або після Flatten).

Контекстно-залежна генерація:

- до Flatten:  $L_{new} \in \{Conv2D, MaxPool, BatchNorm, Dropout\}$ ;

- після Flatten:  $L_{\text{new}} \in \{\text{Dense}, \text{Dropout}\}$ .

Обмеження:

$$|\mathcal{L}'| \leq \text{MAX\_LAYERS}. \quad (2.37)$$

Приклад:

До мутації: [Conv2D(32), Flatten(), Dense(64)]

Додаємо MaxPool() після Conv2D:

Після мутації: [Conv2D(32), MaxPool(), Flatten(), Dense(64)]

2. Операція REMOVE (видалення шару):

$$\mathcal{L}' = [L_1, \dots, L_{i-1}, L_{i+1}, \dots, L_n], \quad (2.38)$$

де шар  $L_i$  видаляється з архітектури.

Обмеження:

- $|\mathcal{L}'| \geq \text{MIN\_LAYERS}$ ; (2.39)
- Flatten та вихідний Dense(1, sigmoid) не можуть бути видалені.

Приклад:

До мутації: [Conv2D(32), MaxPool(), Flatten(), Dense(64),  
Dropout(0.3)]

Видаляємо MaxPool():

Після мутації: [Conv2D(32), Flatten(), Dense(64), Dropout(0.3)]

3. Операція MODIFY (модифікація параметрів):

Модифікація залежить від типу шару:

Conv2D шар:

$$\text{Conv2D}(\text{filters}, k, \alpha) \rightarrow \begin{cases} \text{Conv2D}(\text{filters}', k, \alpha), & \text{filters}' \in \{32, 64, 128\} \\ \text{Conv2D}(\text{filters}, k', \alpha), & k' \in \{3, 5\} \\ \text{Conv2D}(\text{filters}, k, \alpha'), & \alpha' \in \{\text{relu}, \text{elu}\} \end{cases} \quad (2.40)$$



Dense шар:

$$\text{Dense}(\text{neurons}, \alpha) \rightarrow \begin{cases} \text{Dense}(\text{neurons}', \alpha), & \text{neurons}' \in \{32, 64, 128, 256\} \\ \text{Dense}(\text{neurons}, \alpha'), & \alpha' \in \{\text{relu}, \text{elu}\} \end{cases}. \quad (2.41)$$

Dropout шар:

$$\text{Dropout}(\text{rate}) \rightarrow \text{Dropout}(\text{rate}'), \text{rate}' \sim U(0.2, 0.5). \quad (2.42)$$

MaxPool шар:

$$\text{MaxPool}(p) \rightarrow \text{MaxPool}(p'), p' \in \{2, 3\}. \quad (2.43)$$

Приклад:

До мутації: Dense(64, relu)

Модифікуємо neurons:

Після мутації: Dense(128, relu)

### 2.5.2 Параметрична мутація

Параметрична мутація модифікує гіперпараметри навчання без зміни архітектури мережі.

1. Learning rate мутація:

Для learning rate використовується два режими:

Режим 1: Невелика зміна (multiplier mutation):

$$\text{lr}' = \text{lr} \times \alpha, \quad \alpha \sim U(0.5, 2.0). \quad (2.44)$$

Це забезпечує плавну зміну в межах околу поточного значення.

Режим 2: Повна заміна (replacement mutation):

$$lr' \sim U(lr_{\min}, lr_{\max}) = U(0.0001, 0.001). \quad (2.45)$$

Це дозволяє стрибнути в іншу область простору learning rate.

Вибір режиму відбувається з рівною ймовірністю.

2. Batch size мутація:

$$batch' = \text{random\_choice}(\{16, 32, 64\}). \quad (2.46)$$

3. Optimizer мутація:

$$\text{optimizer}' = \text{random\_choice}(\{\text{adam}, \text{sgd}, \text{rmsprop}\}). \quad (2.47)$$

Приклад параметричної мутації:

До мутації:

$lr = 0.0005$ ,  $batch = 32$ ,  $optimizer = \text{adam}$

Після мутації (приклад):

$lr = 0.0008$  (множник 1.6)

$batch = 16$  (випадкова заміна)

$optimizer = \text{rmsprop}$  (випадкова заміна)

Ймовірність мутації  $P_m = 0.4$ :

$$\chi' = \begin{cases} \text{Mutate}(\chi), & U(0,1) < P_m. \\ \chi, & \text{інакше} \end{cases} \quad (2.48)$$

Значення  $P_m = 0.4$  є відносно високим порівняно з класичними рекомендаціями для ГА ( $\sim 0.01$  на ген), але обґрунтоване малим розміром

популяції ( $N = 8$ ): при 8 особинах висока ймовірність мутації необхідна для підтримання генетичного різноманіття.

Валідація після мутації:

якщо тип\_мутації = "структурна":

змінена  $\leftarrow$  ПеревіритиВиправити(змінена)

якщо not змінена.валідна():

повернути оригінал

Це гарантує, що всі хромосоми в популяції є правильними специфікаціями CNN, придатними для побудови в Keras/TensorFlow та навчання на  $\mathcal{D}_{\text{hist}}$ .

## 2.6 Елітизм

Елітизм – стратегія формування нового покоління, що гарантує незгіршення якості найкращої хромосоми: найкраща(і) архітектура(и) з покоління  $t$  переноситься в  $t + 1$  без змін. Це забезпечує монотонність еволюції – властивість, особливо цінну в медичних застосуваннях, де кожна оцінка архітектури є дорогою (50 епох навчання CNN на  $\mathcal{D}_{\text{hist}}$ ).

Механізм:

$$\mathcal{P}_{t+1}^{\text{elite}} = \{\chi_{(1)}^{(t)}, \dots, \chi_{(e)}^{(t)}\}, \quad \Phi(\chi_{(1)}^{(t)}) \geq \Phi(\chi_{(2)}^{(t)}) \geq \dots \quad (2.49)$$

У цій роботі  $e = 1$ :

$$\mathcal{P}_{t+1}[0] = \operatorname{argmax}_{\chi \in \mathcal{P}_t} \Phi(\chi). \quad (2.50)$$

Нове покоління:

$$\mathcal{P}_{t+1} = \mathcal{P}_{t+1}^{\text{elite}} \cup \mathcal{P}_{t+1}^{\text{offspring}}, |\mathcal{P}_{t+1}| = N. \quad (2.51)$$

При  $N = 8$ ,  $e = 1$ : 1 хромосома – елітна копія, 7 – нащадки генетичних операторів.

Гарантія монотонності:

$$\max_{\chi \in \mathcal{P}_{t+1}} \Phi(\chi) \geq \max_{\chi \in \mathcal{P}_t} \Phi(\chi), \quad \forall t. \quad (2.52)$$

Доведення: нехай  $\chi^* = \operatorname{argmax}_{\chi \in \mathcal{P}_t} \Phi(\chi)$ . За визначенням елітизму  $\chi^* \in \mathcal{P}_{t+1}$ , тому  $\max_{\chi \in \mathcal{P}_{t+1}} \Phi \geq \Phi(\chi^*) = \max_{\chi \in \mathcal{P}_t} \Phi$ .

Приклад еволюції (ассигасу на BreakHis):

Покоління 0:  $\chi_1$ :  $\Phi=0.851$  – еліта;  $\chi_2$ :  $\Phi=0.832$ ;  $\chi_3$ :  $\Phi=0.801$  ...

Покоління 1:  $\chi_1'$ :  $\Phi=0.851$  – копія  $\chi_1$ ;  $\chi_2'$ :  $\Phi=0.867$  – новий лідер

Найкраща:  $0.867 \geq 0.851$

Покоління 2:  $\chi_1''$ :  $\Phi=0.867$  – копія  $\chi_2'$ ;  $\chi_2''$ :  $\Phi=0.862$  ...

Найкраща:  $0.867 \geq 0.867$

Вибір  $e = 1$ : при малій популяції ( $N = 8$ ) один елітний слот (12.5%) достатній для збереження кращої архітектури, водночас залишаючи 87.5% місця для нових конфігурацій – баланс між збіжністю і різноманітністю.

## 2.7 Валідація архітектур

Валідація архітектур є невід’ємним компонентом СПС нейромережевих структур: генетичні оператори (кросовер, мутація) можуть породжувати хромосоми, що кодують топологічно некоректні CNN – наприклад, Conv2D після Flatten або занадто багато MaxPool-шарів для малих карт ознак. Для

медичних застосувань невалідна архітектура особливо небезпечна: вона може видавати псевдовпевнені прогнози без повідомлення про помилку, що неприйнятно у задачах онкологічної діагностики.

Правила валідності:

Для забезпечення коректності архітектури в Keras застосовуються структурні правила:

Правило 1: Початковий шар

$$L_1 = \text{Conv2D}. \quad (2.53)$$

Перший шар мережі, що обробляє зображення, повинен бути згортковим. Dense шар не може безпосередньо обробляти 2D зображення без попереднього flatten або pooling.

Правило 2: BatchNorm розміщення

$$\text{якщо } L_i = \text{BatchNorm}, \text{ то } L_{i-1} \in \{\text{Conv2D}, \text{Dense}\}. \quad (2.54)$$

BatchNorm нормалізує активації, тому може йти тільки після шарів з активаціями (Conv2D або Dense). BatchNorm після MaxPool або Dropout не має сенсу, оскільки ці шари не змінюють розподіл активацій.

Правило 3: Flatten обмеження

$$\text{якщо } \exists L_f = \text{Flatten}, \text{ то } \forall i > f: L_i \in \{\text{Dense}, \text{Dropout}, L_{\text{output}}\}. \quad (2.55)$$

Після Flatten(), який перетворює 2D feature maps в 1D вектор, не можуть йти згорткові шари (Conv2D, MaxPool). Тільки Dense, Dropout та вихідний шар.

Правило 4: MaxPool обмеження

$$|\{L_i: L_i = \text{MaxPool}\}| \leq 4. \quad (2.56)$$

Для зображень  $224 \times 224$  пікселі максимум 4 MaxPool шари (після ресайзу BreakHis зображень):

- 1-й MaxPool:  $224 \times 224 \rightarrow 112 \times 112$ ;
- 2-й MaxPool:  $112 \times 112 \rightarrow 56 \times 56$ ;
- 3-й MaxPool:  $56 \times 56 \rightarrow 28 \times 28$ ;
- 4-й MaxPool:  $28 \times 28 \rightarrow 14 \times 14$ .

П'ятий MaxPool призвів би до розміру  $7 \times 7$ , що прийнятно, але збільшення кількості параметрів Flatten  $\rightarrow$  Dense може призвести до перенавчання на обмеженому медичному датасеті.

Правило 5: Utility обмеження

$$\text{якщо } L_i, L_{i+1} \in \{\text{BatchNorm}, \text{MaxPool}, \text{Dropout}\}, \text{ то } L_i \neq L_{i+1}. \quad (2.57)$$

Не більше одного utility-шару одного типу підряд. Наприклад, два Dropout підряд безглузді, оскільки їх ефект можна досягти одним Dropout з вищою rate.

Правило 6: Глибина мережі

$$3 \leq |\mathcal{L}| \leq 5, \quad (2.58)$$

де  $|\mathcal{L}|$  – кількість шарів без вихідного Dense(1, sigmoid).

Правило 7: Вихідний шар

$$\ell_{\text{out}} = \text{Dense}(1, \sigma). \quad (2.59)$$

Обов'язковий вихідний шар для бінарної класифікації гістологічних зображень BreakHis (доброякісна / злоякісна пухлина молочної залози).

Sigmoid на виході повертає скалярну ймовірність класу “злаякісна”, сумісну з функцією втрат `binary_crossentropy`.

Алгоритм валідації та автоматичного виправлення.

### Алгоритм 2.6 – Перевірка та виправлення топології CNN

Вхід:  $\chi$  – хромосома (можливо некоректна)

Вихід:  $\chi_v$  – структурно валідна хромосома

```

1:  $\chi_v \leftarrow \text{копія}(\chi)$ 
2:
3: // Обмеження А: перший блок мусить бути згортковим
4: якщо  $\chi_v.\text{блоки}[0].\text{тип} \neq \text{Conv2D}$ :
5:      $\chi_v.\text{блоки}[0] \leftarrow \text{НовийConv2D}()$ 
6:
7: // Обмеження Б: BatchNorm допустимий лише після Conv2D або Dense
8: для  $j \leftarrow 1$  до  $|\chi_v.\text{блоки}| - 1$ :
9:     якщо  $\text{тип}(\chi_v.\text{блоки}[j]) = \text{BatchNorm}$ 
10:        і  $\text{тип}(\chi_v.\text{блоки}[j-1]) \notin \{\text{Conv2D}, \text{Dense}\}$ :
11:            ВилучитиБлок( $\chi_v$ ,  $j$ )
12:
13: // Обмеження В: після Flatten дозволені лише Dense і Dropout
14:  $\text{поз\_fl} \leftarrow \text{ЗнайтиFlatten}(\chi_v.\text{блоки})$ 
15: якщо  $\text{поз\_fl}$  знайдено:
16:     для  $j \leftarrow \text{поз\_fl}+1$  до  $|\chi_v.\text{блоки}|-1$ :
17:         якщо  $\text{тип}(\chi_v.\text{блоки}[j]) \in \{\text{Conv2D}, \text{MaxPool}\}$ :
18:              $\chi_v.\text{блоки}[j] \leftarrow \text{НовийDense}()$ 
19:
20: // Обмеження Г: кількість MaxPool не більше 4
21: якщо  $\text{КількістьMaxPool}(\chi_v) > 4$ :
22:     ВидалитиЗайвіMaxPool( $\chi_v$ , ліміт=4)
23:
24: // Обмеження Д: заборона двох однотипних службових блоків поспіль
25: для  $j \leftarrow 0$  до  $|\chi_v.\text{блоки}|-2$ :
26:     якщо  $\text{ОднаковийТип}(\chi_v.\text{блоки}[j], \chi_v.\text{блоки}[j+1])$  і
        $\text{ЦеСлужбовий}(\dots)$ :
27:         ВилучитиБлок( $\chi_v$ ,  $j+1$ )

```

```

28:
29: // Обмеження Е: допустимий діапазон глибини
30: доки |χ_в.блоки| < MIN_БЛОКІВ: ДодатиВипадковий(χ_в)
31: доки |χ_в.блоки| > MAX_БЛОКІВ: ВилучитиДовільний(χ_в)
32:
33: // Обмеження Ж: вихідний блок – бінарний класифікатор
34: χ_в.блоки[-1] ← Dense(1, sigmoid)
35:
36: повернути χ_в

```

Приклад валідації:

Розглянемо некоректну архітектуру, що виникла після мутації:

Некоректна архітектура:

|                       |   |                                    |
|-----------------------|---|------------------------------------|
| [MaxPool(2×2),        | ← | Порушення А: MaxPool до Conv2D     |
| (нема карт ознак)     |   |                                    |
| Conv2D(128, 3, relu), | ← | OK                                 |
| Conv2D(64, 3, relu),  | ← | OK                                 |
| MaxPool(2×2),         | ← | OK (1-й дозволений)                |
| MaxPool(2×2),         | ← | OK (2-й дозволений)                |
| MaxPool(2×2),         | ← | OK (3-й дозволений)                |
| MaxPool(2×2),         | ← | OK (4-й дозволений)                |
| MaxPool(2×2),         | ← | Порушення Г: 6-й MaxPool (ліміт=4) |
| Dense(256, relu),     | ← | Порушення В: Dense без Flatten     |
| Dropout(0.5),         |   |                                    |
| Dropout(0.5),         | ← | Порушення Д: два Dropout підряд    |
| Dense(1, sigmoid)]    |   |                                    |

Після застосування ПеревіритиВиправити():

Коректна архітектура:

|                        |   |                                  |
|------------------------|---|----------------------------------|
| [Conv2D(128, 3, relu), | ← | Виправлено А: зайвий MaxPool на  |
| початку знятий         |   |                                  |
| Conv2D(64, 3, relu),   | ← | Збережено                        |
| MaxPool(2×2),          | ← | Збережено                        |
| MaxPool(2×2),          | ← | Збережено                        |
| MaxPool(2×2),          | ← | Збережено                        |
| MaxPool(2×2),          | ← | Збережено (4-й, досягнуто ліміт) |
| Flatten(),             | ← | Додано В: обов'язковий перехід   |



Conv→Dense

|                    |                                                   |
|--------------------|---------------------------------------------------|
| Dense(256, relu),  | ← Збережено                                       |
| Dropout(0.5),      | ← Збережено (один Dropout)                        |
| Dense(1, sigmoid)] | ← Збережено: вихідний шар бінарного класифікатора |

(зайвий MaxPool, дублікат Dropout та відсутній Flatten – виправлено)

Математична формалізація:

Валідація може бути представлена як предикат:

$$\text{IsValid}: \mathcal{A} \rightarrow \{\text{true}, \text{false}\}, \quad (2.60)$$

де  $\mathcal{A}$  – простір всіх можливих архітектур.

Функція виправлення є проєкцією на множину валідних архітектур:

$$\text{Fix}: \mathcal{A} \rightarrow \mathcal{A}_{\text{valid}}, \text{де } \mathcal{A}_{\text{valid}} = \{a \in \mathcal{A} : \text{IsValid}(a) = \text{true}\}. \quad (2.61)$$

Гарантії:

- 1) завершуваність: Алгоритм завжди завершується за скінченний час;
- 2) коректність:  $\text{IsValid}(\text{Fix}(c)) = \text{true}, \forall c \in \mathcal{A}$ ;
- 3) ідемпотентність:  $\text{Fix}(\text{Fix}(c)) = \text{Fix}(c)$ .

Завдяки системі валідації, генетичний алгоритм може вільно застосовувати генетичні оператори, знаючи, що всі згенеровані архітектури будуть коректними та придатними для тренування.

## 2.8 Метрики оцінки якості для задач медичної класифікації

Оцінка якості моделі класифікації гістологічних зображень потребує розширеного набору метрик порівняно із загальними задачами комп'ютерного бачення. Причина – асиметрія ціни помилок: False Negative (пропущена злоякісна пухлина) має принципово іншу клінічну вагу, ніж False Positive (зайва біопсія). Стандартна accuracy не відображає цю асиметрію.

### 2.8.1 Accuracy (загальна точність)

Загальна точність – частка правильно класифікованих гістологічних зображень:

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}}, \quad (2.62)$$

де для задачі BreakHis: TP – злоякісний знімок класифіковано як злоякісний; TN – доброякісний знімок класифіковано як доброякісний; FP – доброякісний знімок хибно класифіковано як злоякісний (зайва біопсія); FN – злоякісний знімок хибно класифіковано як доброякісний (пропущений рак).

$$\text{Accuracy} = \frac{1}{N} \sum_{i=1}^N \mathbb{1} [\hat{y}_i = y_i]. \quad (2.63)$$

Клінічна роль: accuracy – основна метрика fitness function у цій роботі через достатню збалансованість BreakHis (70.4% злоякісних / 29.6% доброякісних). Базовий рівень мажоритарного класифікатора: 70.4%.

### 2.8.2 Sensitivity (Чутливість, Recall)

Чутливість (sensitivity, recall) – здатність моделі виявляти всі злоякісні пухлини:

$$\text{Sensitivity} = \frac{TP}{TP+FN}. \quad (2.64)$$

Клінічна інтерпретація: Sensitivity = 0.95 означає, що 5 з 100 злоякісних зображень пропущено (False Negative). У онкологічній діагностиці кожен FN – це потенційно непоставлений діагноз раку на ранній стадії, що безпосередньо впливає на прогноз виживання пацієнта.

Мінімально прийнятна sensitivity для клінічного інструменту підтримки рішень:  $\geq 0.90$  (рекомендація ACR [American College of Radiology]).

### 2.8.3 Specificity (Специфічність)

Специфічність – здатність моделі коректно ідентифікувати доброякісні зразки:

$$\text{Specificity} = \frac{TN}{TN+FP}. \quad (2.65)$$

Клінічна інтерпретація: Specificity = 0.80 означає, що 20% доброякісних зображень помилково позначаються як злоякісні (False Positive), що призводить до непотрібних біопсій. Хоча кожна зайва біопсія не загрожує безпосередньо життю, вона пов'язана з психологічним стресом, медичними ризиками процедури та фінансовими витратами.

Trade-off Sensitivity–Specificity: підвищення порогу класифікації  $\hat{y} \geq \tau$  збільшує Specificity, але знижує Sensitivity, і навпаки. Для медичних систем зазвичай пріоритет – Sensitivity.

#### 2.8.4 Precision та F1-Score

Precision – частка правильних позитивних прогнозів:

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}}. \quad (2.66)$$

F1-Score – гармонічне середнє Precision і Recall (Sensitivity):

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Sensitivity}}{\text{Precision} + \text{Sensitivity}} = \frac{2 \cdot \text{TP}}{2 \cdot \text{TP} + \text{FP} + \text{FN}}. \quad (2.67)$$

Гармонічне середнє карає за дисбаланс: модель з Precision=1.0 і Sensitivity=0.5 має F1=0.67, а не 0.75 (арифметичне середнє). Для медичних моделей, де Sensitivity  $\geq$  Precision за пріоритетом,  $F_\beta$ -score ( $\beta > 1$ ) може краще відображати клінічні вимоги:

$$F_\beta = (1 + \beta^2) \cdot \frac{\text{Precision} \cdot \text{Sensitivity}}{\beta^2 \cdot \text{Precision} + \text{Sensitivity}}. \quad (2.68)$$

### 2.8.5 AUC-ROC (Area Under the ROC Curve)

AUC-ROC – найважливіша порогово-незалежна метрика для медичних класифікаторів. ROC-крива будується за парами (FPR, TPR) для всіх можливих порогів  $\tau \in [0, 1]$ :

$$\text{TPR}(\tau) = \frac{\text{TP}(\tau)}{\text{TP}(\tau) + \text{FN}(\tau)}, \text{FPR}(\tau) = \frac{\text{FP}(\tau)}{\text{FP}(\tau) + \text{TN}(\tau)}, \quad (2.69)$$

де  $\text{TPR} = \text{Sensitivity}$ ,  $\text{FPR} = 1 - \text{Specificity}$ .

AUC (площа під кривою ROC):

$$\text{AUC} = \int_0^1 \text{TPR}(\text{FPR}) d(\text{FPR}). \quad (2.70)$$

Клінічна інтерпретація AUC наведено в таблиці 2.5.

**Таблиця 2.5** – Клінічна інтерпретація значень AUC-ROC

| AUC       | Клінічна оцінка                                     |
|-----------|-----------------------------------------------------|
| 0.50      | Випадковий класифікатор (марна модель)              |
| 0.70–0.80 | Помірна діагностична цінність                       |
| 0.80–0.90 | Хороша діагностична цінність                        |
| 0.90–1.00 | Відмінна діагностична цінність (клінічно прийнятна) |

Чому AUC важливіше за accuracy для BreakHis: accuracy залежить від порогу  $\tau = 0.5$  і дисбалансу класів. AUC оцінює здатність моделі ранжувати злоякісні зображення вище за доброякісні – незалежно від конкретного порогу. Це відповідає реальній клінічній задачі: модель надає патологоанатому список підозрілих зображень, відсортованих за ймовірністю злоякісності.

Використання в цій роботі: функція пристосованості включає AUC-ROC як додатковий компонент:

$$\Phi(\chi) = 0.6 \cdot \text{Accuracy}(\chi) + 0.4 \cdot \text{AUC}(\chi). \quad (2.71)$$

#### 2.8.6 Матриця помилок для бінарної медичної класифікації

Для бінарної класифікації BreakHis матриця плутанини має розмір  $2 \times 2$ :

$$\mathbf{C} = \begin{bmatrix} \text{TN} & \text{FP} \\ \text{FN} & \text{TP} \end{bmatrix} = \begin{bmatrix} \text{Benign} \rightarrow \text{Benign} & \text{Benign} \rightarrow \text{Malignant} \\ \text{Malignant} \rightarrow \text{Benign} & \text{Malignant} \rightarrow \text{Malignant} \end{bmatrix}. \quad (2.72)$$

Медична інтерпретація:

- $C_{11}$  (TN): коректно виявлені доброякісні – правильний результат;
- $C_{12}$  (FP): хибна тривога – зайва біопсія (помірна шкода);
- $C_{21}$  (FN): пропущений рак – найнебезпечніша помилка;
- $C_{22}$  (TP): коректно виявлені злоякісні – правильний результат;

#### 2.8.7 Метрики ефективності синтезу для порівняння з іншими методами NAS

Окрім класифікаційних метрик, для порівняння підходів СПС використовуються:

$$\text{Search Cost} = \text{GPU-год} \times N_{\text{GPU}}. \quad (2.73)$$

$$\text{Reproducibility Error} = \text{MAE}(\Phi_{\text{orig}}, \Phi_{\text{valid}}). \quad (2.74)$$

$$\text{Model Complexity} = |\text{params}|. \quad (2.75)$$

Зведена таблиця метрик наведена в таблиці 2.6.

**Таблиця 2.6** – Зведена таблиця метрик оцінки якості для задач медичної класифікації

| <i>Метрика</i>        | <i>Призначення</i>                    | <i>Роль у цій роботі</i>        |
|-----------------------|---------------------------------------|---------------------------------|
| Accuracy              | Загальна якість бінарної класифікації | Основна складова $\Phi(\chi)$   |
| AUC-ROC               | Порогово-незалежна оцінка             | Додаткова складова $\Phi(\chi)$ |
| Sensitivity           | Виявлення злоякісних                  | Звітна метрика                  |
| Specificity           | Уникнення зайвих біопсій              | Звітна метрика                  |
| F1-Score              | Баланс Precision/Sensitivity          | Звітна метрика                  |
| Confusion Matrix      | Аналіз типів помилок                  | Клінічна інтерпретація          |
| Search Cost           | GPU-вартість СПС                      | Порівняння з іншими NAS         |
| Reproducibility Error | Відтворюваність                       | Наукова валідність              |

## 2.9 Обґрунтування вибору підходу

Вибір генетичного алгоритму для пошуку архітектури нейронних мереж обумовлений рядом методологічних, практичних та прикладних факторів.

### 2.9.1 Переваги еволюційного підходу

#### 1. Одночасна оптимізація архітектури та гіперпараметрів:

На відміну від більшості існуючих методів NAS (DARTS, NASNet, ENAS), які оптимізують тільки архітектуру при фіксованих гіперпараметрах, запропонований метод дозволяє:

$$\arg \max_{(a,h) \in \mathcal{A} \times \mathcal{H}} f(a, h), \quad (2.76)$$

де  $a$  – архітектура,  $h$  – гіперпараметри.

Експериментальні результати показали нетривіальні залежності:

- глибокі архітектури (23+ шари) потребують малих розмірів batch (16);
- RMSprop з LR=0.00057 забезпечує стабільну конвергенцію;
- Найкраща модель: мінімалістична архітектура (3 шари) + Adam + LR=0.000727 + batch=16.

Фіксування гіперпараметрів могло б запобігти виявленню цих оптимальних комбінацій.

#### 2. Інтерпретовність процесу:

Генетичний алгоритм є “білою скринькою” порівняно з градієнтними методами. Порівняння наведено в таблиці 2.7.



**Таблиця 2.7** – Порівняння генетичного підходу та градієнтних методів (DARTS)

| Аспект             | Генетичний підхід                        | Градієнтні методи (DARTS)        |
|--------------------|------------------------------------------|----------------------------------|
| Процес пошуку      | Зрозумілий (селекція, кросовер, мутація) | “Чорна скринька” (gradient flow) |
| Аналіз поколінь    | Можна відстежити еволюцію                | Неможливо                        |
| Причини успіху     | Видно, які архітектурні блоки працюють   | Непрозоро                        |
| Навчальна цінність | Висока (зрозумілі кроки)                 | Низька                           |

Приклад інтерпретації: Аналіз 136 архітектур (режим Baseline, BreakHis) виявив закономірності:

- SGD та Adam домінують ( $\approx 67\%$  синтезованих моделей), отже стабільна збіжність на невеликому датасеті гістологічних зображень;
- Batch size 32–128 показує кращу генералізацію, отже менший шум градієнта, стійкіший процес навчання;
- архітектури середньої глибини (5–12 блоків) ефективніші, отже оптимальний баланс між capacity та регуляризацією для 3000 зразків BreakHis.

### 3. Обчислювальна доступність:

Запропонований метод потребує  $\sim 0.58$  GPU-години ( $\sim 35$  хвилин) проти 12–75,600 GPU-годин для SOTA методів:

$$\text{Speedup} = \frac{\text{Вартість пошуку DARTS}}{\text{Вартість СПС}} = \frac{96 \text{ GPU-год}}{0.58 \text{ GPU-год}} \approx \mathbf{165}. \quad (2.77)$$

(порівняно з DARTS; для NASNet прискорення складає  $\sim 74,500\times$ )

Це робить метод доступним для:

- академічних дослідницьких груп з обмеженими ресурсами;
- індивідуальних дослідників;
- освітніх закладів;
- швидкого прототипування для промислових застосувань.

#### 4. Відсутність approximation bias:

ENAS та інші методи з weight-sharing використовують супермережу (supernet), де всі підмережі ділять ваги. Це призводить до approximation bias – оцінка fitness архітектури на спільних вагах може не відповідати її дійсній якості після повного тренування.

Наш підхід:

- 1) повне тренування кожної моделі;
- 2) чесна оцінка fitness;
- 3) відсутність approximation bias.

#### 5 Природна регуляризація:

Популяційний підхід забезпечує природну регуляризацію через:

- різноманітність: 8 різних архітектур в кожному поколінні;
- стохастичність: випадковість в селекції, кросовері, мутації;
- елітизм: збереження кращих рішень.

Це запобігає передчасній збіжності до локальних оптимумів.

#### 6. Гнучкість та розширюваність:

Генетичний підхід легко адаптується для інших датасетів: просто змінити розмір вхідних зображень та кількість класів; Інших архітектурних блоків: додати нові типи шарів (attention, residual); Multi-objective оптимізації: додати обмеження на розмір моделі або швидкість inference; Інших метрик: замінити accuracy на F1-score або custom metric

### 2.9.2 Придатність для задачі класифікації гістопатологічних зображень

Задача класифікації гістопатологічних зображень тканин молочної залози (BreakHis) має специфічні характеристики, які роблять структурно-параметричний синтез на основі еволюційного підходу особливо придатним:

1. Критична асиметрія вартості помилок.

У медичній діагностиці помилки мають різну вартість:

- False Negative: критична помилка, може призвести до запізненого лікування та погіршення прогнозу виживаності;
- False Positive: призводить до надмірного лікування та психологічного стресу

Генетичний алгоритм з повним тренуванням кожної синтезованої моделі забезпечує чесну оцінку sensitivity та specificity (без approximation bias), можливість аналізу confusion matrix та ROC-кривої для кожної архітектури. А також вибір моделі з оптимальним балансом sensitivity/specificity згідно з клінічними вимогами

2. Обмеженість розміщеного медичного датасету.

Гістопатологічні датасети обмежені порівняно з природними зображеннями через необхідність залучення кваліфікованих патологоанатомів для розмітки, складність та трудомісткість гістологічної підготовки зразків та обмежений доступ до клінічних даних через вимоги конфіденційності.

Для обмежених медичних датасетів структурно-параметричний синтез забезпечує автоматичне знаходження оптимальної архітектурної складності (не занадто глибокі мережі), одночасну оптимізацію Dropout та інших параметрів регуляризації та повне тренування кожної синтезованої моделі.

### 3. Специфіка текстурних ознак гістологічних зображень.

H&E гістологічні зображення містять:

- регулярні структури (залози, протоки);
- важливі на низьких збільшеннях;
- клітинні деталі (ядра, хроматин);
- важливі на високих збільшеннях;
- колірні градієнти (розподіл гематоксиліну/еозину).

Еволюційний синтез автоматично виявляє архітектурні рішення (кількість фільтрів, розміри ядер), що оптимально відповідають цим текстурним ознакам, без ручного налаштування під специфіку гістологічних даних.

### 4. Необхідність відтворюваності для медичних регуляторних вимог.

Запропонований фреймворк відтворюваності забезпечує збереження всіх 128 синтезованих моделей (архітектура + ваги), Reproducibility Error 0.03% (на два порядки нижче стохастичності) та можливість незалежної верифікації без повторного тренування

### 5. Інтерпретовність для клінічного застосування.

Медичні застосування потребують пояснюваності: Які архітектурні компоненти найбільш ефективні для гістологічних ознак? Як еволюціонує популяція синтезованих архітектур? Чому конкретна архітектура є оптимальною для BreakHis?

Генетичний підхід надає повну аналізовану історію синтезу (128 моделей, 16 поколінь), можливість статистичного аналізу успішних архітектурних паттернів, візуалізацію траєкторії популяції в просторі архітектур.

Математичне обґрунтування вибору ГА для СПС:

Задача структурно-параметричного синтезу є NP-важкою задачею комбінаторної оптимізації:

$$\max_{(\alpha, \theta) \in \mathcal{A} \times \Theta} \mathcal{F}(M(\alpha, \theta)), \quad |\mathcal{A} \times \Theta| \approx 10^{12}. \quad (2.78)$$

Методи розв'язання:

- 1) exhaustive search:  $O(|\mathcal{A} \times \Theta|)$  – обчислювально нездійсненно;
- 2) random search:  $O(k)$  – низька якість рішень;
- 3) градієнтні методи: оптимізують лише  $\alpha$  при фіксованих  $\theta$  – часткове рішення;
- 4) еволюційні методи:  $O(p \times g)$  – одночасний синтез  $\alpha$  та  $\theta$   
де  $p$  – розмір популяції,  $g$  – кількість поколінь.

Для нашого випадку:  $p = 8$ ,  $g = 16 \rightarrow O(128)$  повних тренувань синтезованих конфігурацій.

Висновок.

Вибір генетичного алгоритму для структурно-параметричного синтезу нейромережових структур для класифікації гістопатологічних зображень є оптимальним з точки зору: – Методології: одночасний синтез структурної та параметричної складових, відсутність approximation bias – Практичності: обчислювальна доступність (0.58 GPU-години), відтворюваність для медичної верифікації – Застосування: висока точність (86.83% binary на BreakHis), аналіз sensitivity/specificity, інтерпретовність

## 2.10 Багатоцільова оптимізація: алгоритм NSGA-II

### 2.10.1 Формалізація задачі багатоцільової оптимізації

Стандартний підхід до СПС нейромережових структур максимізує єдину скалярну функцію якості  $\Phi(\chi)$  – зважену суму accuracy і AUC. Проте задача класифікації гістологічних зображень природно має дві конкуруючі цілі:

- максимізувати ассурасу/AUC на  $\mathcal{D}_{\text{hist}}$  (якість діагностики);
- мінімізувати кількість параметрів  $|\theta(\chi)|$  (ефективність для розгортання в клініці).

Більша мережа не обов'язково точніша, але завжди повільніша та вибагливіша до ресурсів. Задачу можна формалізувати як векторну оптимізацію:

$$\min_{\chi \in \mathcal{S}} f(\chi) = (-\Phi(\chi), |\theta(\chi)|), \quad (2.79)$$

де  $-\Phi(\chi)$  мінімізується (тобто  $\Phi$  максимізується), а  $|\theta(\chi)|$  – кількість навчальних параметрів.

### 2.10.2 Концепція Pareto-домінування і Pareto-фронту

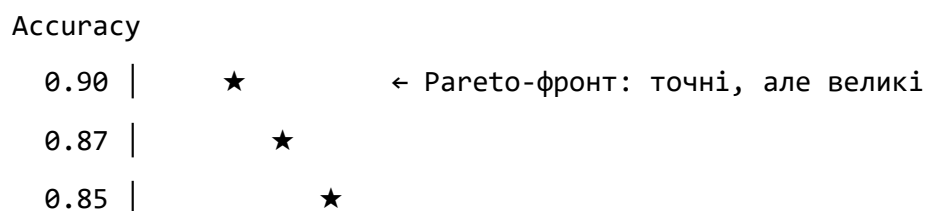
Визначення (Pareto-домінування): Хромосома  $\chi_1$  домінує  $\chi_2$  (позначення:  $\chi_1 < \chi_2$ ) тоді і тільки тоді, коли:

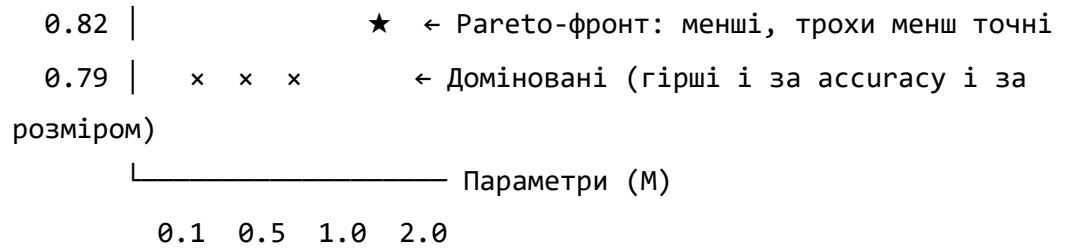
$$\forall k: f_k(\chi_1) \leq f_k(\chi_2) \text{ і } \exists k: f_k(\chi_1) < f_k(\chi_2). \quad (2.80)$$

Тобто  $\chi_1$  не гірша  $\chi_2$  за жодним критерієм і краща хоча б за одним.

Визначення (Pareto-оптимальна множина): Хромосома  $\chi^*$  є Pareto-оптимальною, якщо не існує жодної  $\chi \in \mathcal{S}$  такої, що  $\chi < \chi^*$ . Множина всіх Pareto-оптимальних рішень утворює Pareto-фронт  $\mathcal{F}^*$ .

Ілюстрація для задачі СПС (BreakHis):





Хромосоми на Pareto-фронті (★) є рівноцінними в парето-сенсі: жодна з них не домінує іншу. Клінічно: патологоанатомічна лабораторія з потужним GPU може обрати більш точну модель; мобільний клінічний пристрій – меншу.

### 2.10.3 Алгоритм NSGA-II

NSGA-II (Non-dominated Sorting Genetic Algorithm II, Deb et al., 2002) є провідним алгоритмом багатоцільової еволюційної оптимізації. Він розширює стандартний ГА двома механізмами: недомінованим сортуванням та crowding distance.

Недоміноване сортування розбиває популяцію  $\mathcal{P}$  на фронти:

$$\mathcal{F}_1 = \{\chi \in \mathcal{P} \mid \nexists \chi' \in \mathcal{P}: \chi' < \chi\}. \quad (2.81)$$

$$\mathcal{F}_k = \left\{ \chi \in \mathcal{P} \setminus \bigcup_{i < k} \mathcal{F}_i \mid \nexists \chi' \in \mathcal{P} \setminus \bigcup_{i < k} \mathcal{F}_i: \chi' < \chi \right\}, k \geq 2. \quad (2.82)$$

$\mathcal{F}_1$  – перший фронт (Pareto-оптимальний),  $\mathcal{F}_2$  – після виключення  $\mathcal{F}_1$ , і т.д.

Crowding Distance вимірює “щільність оточення” особини в просторі цілей:

$$d_i = \sum_{m=1}^M \frac{f_m(\chi_{i+1}) - f_m(\chi_{i-1})}{f_m^{max} - f_m^{min}}, \quad (2.83)$$

де  $M$  – кількість цілей,  $\chi_{i-1}$  та  $\chi_{i+1}$  – сусіди  $\chi_i$  при сортуванні за  $m$ -ю ціллю. Більше  $d_i$  – ізольованіша особина, менше конкурентів.

Правило вибору особин (NSGA-II порядок):

$$\begin{aligned} \chi_1 \succ_{\text{NSGA-II}} \chi_2 &\Leftrightarrow \\ \Leftrightarrow \text{rank}(\chi_1) < \text{rank}(\chi_2) \vee (\text{rank}(\chi_1) = \text{rank}(\chi_2) \wedge d(\chi_1) > d(\chi_2)). \end{aligned} \quad (2.84)$$

Особини з нижчим рангом (ближчі до Pareto-фронту) і більшою crowding distance (більш ізольовані, підтримують різноманіття) мають перевагу.

Загальний алгоритм NSGA-II для СПС.

### Алгоритм 2.7 – NSGA-II для структурно-параметричного синтезу CNN

Вхід:  $N$ ,  $T$ ,  $P_c$ ,  $P_m$

Вихід: Pareto-фронт  $\chi_1^*$ , ...,  $\chi_k^*$

```

1:  $P_0 \leftarrow \text{InitPopulation}(N)$ 
2: Оцінити  $\Phi(\chi)$  та  $|\theta(\chi)|$  для всіх  $\chi \in P_0$ 

3: для  $t = 1$  до  $T$ :
4:    $Q \leftarrow \text{GenerateOffspring}(P_{t-1}, P_c, P_m)$  // кросовер + мутація
5:   Оцінити нові  $\chi \in Q$ 
6:    $R \leftarrow P_{t-1} \cup Q$  // об'єднати батьків і нащадків
7:
8:    $\Phi_1, \Phi_2, \dots \leftarrow \text{NonDominatedSort}(R)$  // розбиття на фронти
9:    $P_t \leftarrow \emptyset$ 
10:
11:   для  $i = 1, 2, \dots$ :
12:     якщо  $|P_t| + |\Phi_i| \leq N$ :
13:        $P_t \leftarrow P_t \cup \Phi_i$ 
```



```

14:         інакше:
15:             // Частково беремо з поточного фронту – за crowding distance
16:             CalculateCrowdingDistance( $\Phi_i$ )
17:             SortByDistance( $\Phi_i$ , спадання)
18:              $P_t \leftarrow P_t \cup \Phi_i[:N - |P_t|]$ 
19:             break
20:
21: повернути NonDominatedSort( $P_t$ )[0] // Pareto-фронт

```

GA-II до задачі СПС гістологічних зображень

У цій роботі NSGA-II застосовується з двома цілями:

$$f_1(\chi) = -\Phi(\chi) = -(0.6 \cdot \text{Accuracy} + 0.4 \cdot \text{AUC}), f_2(\chi) = |\theta(\chi)| \quad (2.85)$$

Практичні переваги для медичного AI:

- 1) набір альтернативних архітектур: Замість однієї “найкращої” моделі отримується Pareto-фронт з кількох моделей – від мінімальних (мале число параметрів, помірна точність) до максимальних (висока точність, більше параметрів). Клінічна установа обирає модель залежно від своєї інфраструктури;
- 2) явна торгівля accuracy–complexity: Клінічний адміністратор бачить кількісно, скільки точності “коштує” зменшення моделі вдвічі по параметрах;
- 3) уникнення передчасної збіжності: Crowding distance підтримує рівномірний розподіл по фронті, що запобігає концентрації популяції в одній точці простору архітектур.

Порівняння одноцільового ГА та NSGA-II наведено в таблиці 2.8.

**Таблиця 2.8** – Порівняння одноцільового генетичного алгоритму та NSGA-II

| <i>Властивість</i>       | <i>Одноцільовий ГА</i>     | <i>NSGA-II</i>                     |
|--------------------------|----------------------------|------------------------------------|
| Результат                | Одна найкраща архітектура  | Множина Pareto-оптимальних         |
| Управління complexity    | Немає                      | Явна мінімізація параметрів        |
| Різноманіття             | Залежить від $P_m$         | Гарантується crowding distance     |
| Обчислювальна складність | $O(N \log N)$ на покоління | $O(N^2)$ на покоління (сортування) |
| Застосовність            | Один критерій              | Кілька конкуруючих критеріїв       |

## Висновки до розділу 2

У цьому розділі було розглянуто математичні та теоретичні основи методу структурно-параметричного синтезу нейромережових структур з використанням генетичного алгоритму.

Основні теоретичні внески:

- 5) основи згорткових нейронних мереж (розділ 2.1): Детально розглянуто математичні основи базових компонентів CNN: згорткові шари (операція згортки, формула 2.1), pooling шари (Max Pooling, формула 2.2), Batch Normalization (нормалізація по батчу, формули 2.4-2.7), Dropout (регуляризація, формули 2.8-2.10) та функції активації (ReLU, ELU, Softmax, формули 2.11-2.14). Ці компоненти складають будівельні блоки для простору пошуку архітектур;
- 6) математична модель генетичного алгоритму (розділ 2.2): Формалізовано простір пошуку  $\mathcal{S} = \mathcal{A} \times \mathcal{H}$  (формула 2.15) з

оцінкою потужності  $|\mathcal{S}| \approx 10^{12}$  можливих архітектур (формула 2.22). Розроблено уніфіковане генетичне кодування хромосоми  $\mathbf{c} = (\mathcal{L}, \mathcal{P})$  (формула 2.23), яке спільно представляє архітектуру та гіперпараметри навчання. Визначено адаптивну функцію пристосованості  $f(\mathbf{c}) = \max_{i=1}^n \text{acc}_{\text{val}}^{(i)}$  (формула 2.30), яка враховує перенавчання через вибір найкращої епохи;

- 7) генетичні оператори (розділи 2.3-2.6): Математично формалізовано турнірну селекцію (формула 2.36), гібридний кросовер з окремою обробкою структурної (формула 2.39) та параметричної частини (формули 2.40-2.42), адаптивну мутацію (структурна формули 2.46-2.52, параметрична формули 2.53-2.56) та елітизм з гарантією монотонності  $\max_{\mathbf{c} \in P_{t+1}} f(\mathbf{c}) \geq \max_{\mathbf{c} \in P_t} f(\mathbf{c})$  (формула 2.61);
- 8) система валідації архітектур (розділ 2.7): Формалізовано 7 структурних правил коректності для Keras/TensorFlow (формули 2.62-2.68) та розроблено алгоритм автоматичного виправлення невалідних архітектур. Валідація формалізована як предикат  $\text{IsValid}: \mathcal{A} \rightarrow \{\text{true}, \text{false}\}$  та проєкція  $\text{Fix}: \mathcal{A} \rightarrow \mathcal{A}_{\text{valid}}$  (формули 2.69-2.70);
- 9) метрики оцінки якості (розділ 2.8): Визначено систему метрик: Accuracy (формула 2.72), Precision/Recall (формули 2.74-2.76), F1-score (формула 2.77), Confusion Matrix (формула 2.81) та специфічні метрики для NAS: Search Cost, Reproducibility Error, Model Complexity (формули 2.82-2.85);
- 10) обґрунтування вибору підходу (розділ 2.9): Доведено переваги еволюційного підходу для задачі структурно-параметричного синтезу нейромереж для класифікації гістопатологічних зображень: одночасна оптимізація структурної та параметричної складових, обчислювальна доступність ( $\text{speedup} \approx 165\times$  порівняно

з DARTS, формула 2.87), врахування асиметрії вартості помилок (sensitivity/specificity), інтерпретовність процесу, відсутність approximation bias та природна регуляризація через популяційний підхід.

Ключові математичні результати:

- потужність простору пошуку:  $|\mathcal{S}| \approx 10^{12}$  унікальних архітектур;
- складність алгоритму:  $O(p \times g) = O(8 \times 16) = O(128)$  повних тренувань;
- гарантована монотонність через елітизм: найкраща fitness ніколи не погіршується;
- Reproducibility Error: 0.03% (на два порядки нижче стохастичності тренування)

Практична цінність:

Розроблена математична модель забезпечує теоретичне обґрунтування для практичної реалізації системи структурно-параметричного синтезу нейромережових структур, яка поєднує обчислювальну доступність (0.58 GPU-години), високу точність (86.83% бінарна класифікація на BreakHis) та повну відтворюваність результатів для медичної верифікації. Формалізація всіх компонентів дозволяє незалежну верифікацію методу та його адаптацію для інших задач медичного комп'ютерного бачення.

У наступному розділі буде описано практичну реалізацію розробленої математичної моделі, експериментальні дослідження та аналіз отриманих результатів на датасеті BreakHis (гістопатологічні зображення тканин молочної залози).

## РОЗДІЛ 3 РЕАЛІЗАЦІЯ ТА ЕКСПЕРИМЕНТАЛЬНЕ ДОСЛІДЖЕННЯ СИСТЕМИ СТРУКТУРНО-ПАРАМЕТРИЧНОГО СИНТЕЗУ НЕЙРОМЕРЕЖЕВИХ СТРУКТУР

### 3.1 Архітектура програмного рішення

#### 3.1.1 Загальна структура системи

Розроблена система автоматичного пошуку архітектури нейронних мереж побудована за модульним принципом, що забезпечує гнучкість, масштабованість та можливість незалежного тестування окремих компонентів. Система складається з семи основних модулів, кожен з яких виконує специфічну функцію в процесі еволюційного пошуку.

Архітектура системи спроектована з урахуванням таких принципів:

- 1) модульність: Кожен компонент має чітко визначені інтерфейси та відповідальність;
- 2) розширюваність: Можливість додавання нових типів шарів, операторів та метрик без зміни базової структури;
- 3) відтворюваність: Автоматичне збереження всіх експериментальних даних та параметрів;
- 4) ефективність: Оптимізація використання GPU пам'яті та обчислювальних ресурсів.

#### 3.1.2 Модульна організація коду

Вихідний код системи організований у таку структуру директорій:

```
genetic_nas/
├── src/                # Вихідний код системи
│   ├── chromosome.py  # Кодування хромосоми
│   └── population.py   # Управління популяцією
```

|                       |                              |
|-----------------------|------------------------------|
| — operators.py        | # Генетичні оператори        |
| — fitness.py          | # Функція пристосованості    |
| — evolution.py        | # Основний цикл еволюції     |
| — dataset_loader.py   | # Завантаження датасету      |
| — model_history.py    | # Система збереження моделей |
| — history_loader.py   | # Завантаження історії       |
| — config.py           | # Конфігурація параметрів    |
| — logger.py           | # Система логування          |
| — utils.py            | # Допоміжні функції          |
| — main.py             | # Точка входу в програму     |
| — validate_models.py  | # Скрипт валідації           |
| — generate_figures.py | # Генерація візуалізацій     |
| — output/             | # Результати експериментів   |
| — model_history/      | # Повна історія моделей      |
| — all_models.json     | # Метадані всіх моделей      |
| — *.h5                | # Збережені ваги моделей     |
| — *.log               | # Логи тренування            |
| — data/               | # Датасет BreakHis           |
| — requirements.txt    | # Python залежності          |

### 3.1.3 Діаграма компонентів

Система складається з семи ключових компонентів.

1. Модуль кодування (chromosome.py): класи Layer та Chromosome для представлення архітектури; методи валідації архітектури; генерація випадкових хромосом; серіалізація/десеріалізація у JSON формат.

2. Модуль популяції (population.py): управління набором хромосом (особин); ініціалізація початкової популяції; обчислення статистики популяції; сортування за fitness.

3. Модуль генетичних операторів (`operators.py`): турнірна селекція (`tournament_selection`); гібридний кросовер (`crossover`); структурна та параметрична мутація (`mutate`); валідація результатів операторів.

4. Модуль оцінки `fitness` (`fitness.py`): побудова Keras моделі з хромосоми (`build_model`); компіляція з оптимізатором (`get_optimizer`); Тренування моделі (`evaluate_fitness`); обчислення максимальної валідаційної точності.

5. Модуль еволюції (`evolution.py`): клас `EvolutionEngine` – головний контролер; основний цикл еволюції; координація між модулями; збереження проміжних результатів.

6. Модуль історії моделей (`model_history.py`): клас `ModelHistoryTracker`, збереження метаданих моделей у JSON, збереження натренованих ваг у HDF5, інкрементальне оновлення історії.

7. Модуль завантаження даних (`dataset_loader.py`): завантаження BreakHis датасету гістопатологічних зображень; попередня обробка зображень (ресайз, нормалізація H&E); аугментація для збільшення різноманіття тренувальних зразків; розділення на `train/val/test` набори з урахуванням пацієнтів.

### *3.1.4 Взаємодія модулів*

Процес роботи системи відбувається у такій послідовності:

- 1) ініціалізація (`evolution.py` → `config.py`):
  - завантаження конфігураційних параметрів;
  - ініціалізація логування;
  - створення директорій для результатів;
- 2) завантаження даних (`evolution.py` → `dataset_loader.py`):
  - зчитування зображень з файлової системи;

- Preprocessing (resize, normalization);
- створення train/validation splits;
- 3) створення популяції (evolution.py → population.py → chromosome.py):
  - генерація випадкових хромосом;
  - валідація архітектур;
  - ініціалізація fitness значень;
- 4) оцінка популяції (evolution.py → fitness.py):
  - для кожної особини: побудова моделі, тренування, обчислення fitness;
  - паралельна обробка на GPU;
  - збереження натренованих моделей;
- 5) генетичні операції (evolution.py → operators.py):
  - селекція батьківських особин;
  - застосування кросоверу;
  - застосування мутації;
  - елітизм (збереження найкращої особини);
- 6) збереження результатів (evolution.py → model\_history.py):
  - запис метаданих у all\_models.json;
  - збереження ваг моделей у .h5 файли;
  - логування статистики покоління;
- 7) повторення циклу до досягнення заданої кількості поколінь або критерію зупинки.

Ця модульна архітектура забезпечує:

- легке тестування окремих компонентів;
- можливість заміни реалізації без зміни інтерфейсів;
- паралелізацію обчислень (майбутнє розширення);
- простоту налагодження та профілювання.



## 3.2 Опис датасету BreakHis

### 3.2.1 Характеристики датасету

Для проведення експериментального дослідження було обрано датасет BreakHis (Breast Cancer Histopathological Image Classification), який є одним з найбільш поширених публічних датасетів гістопатологічних зображень тканин молочної залози. Датасет зібрано в лабораторії P&D (Pathology and Cytology Laboratories, Парана, Бразилія) та опубліковано Spanhol et al. у 2016 році.

Основні характеристики наведено в таблиці 3.1.

**Таблиця 3.1** – Основні характеристики датасету BreakHis

| <i>Параметр</i>                   | <i>Значення</i>                   |
|-----------------------------------|-----------------------------------|
| Джерело                           | P&D Laboratory [15]               |
| Тип задачі                        | Бінарна та 8-класова класифікація |
| Кількість класів (основна задача) | 2 (доброякісна / злоякісна)       |
| Загальна кількість зображень      | 7,909                             |
| Розмір тренувального набору       | 5,537 зображень (70%)             |
| Розмір тестового набору           | 2,372 зображення (30%)            |
| Оригінальний розмір зображень     | 700×460 пікселів                  |
| Розмір після preprocessing        | 224×224 пікселі                   |
| Формат зображень                  | RGB (3 канали), Н&Е забарвлення   |
| Рівні збільшення                  | 40×, 100×, 200×, 400×             |

### 3.2.2 Класи об'єктів

Бінарна класифікація (основна задача) наведена в таблиці 3.2.

**Таблиця 3.2** – Розподіл зображень за класами бінарної класифікації

| <i>Індекс</i> | <i>Назва класу</i>    | <i>К-сть зображень</i> | <i>% від загальної</i> |
|---------------|-----------------------|------------------------|------------------------|
| 0             | Benign (Доброякісна)  | 2,480                  | 31.4%                  |
| 1             | Malignant (Злаякісна) | 5,429                  | 68.6%                  |

8 підтипів (повна класифікація) представлено в таблиці 3.3.

**Таблиця 3.3** – Розподіл зображень за 8 підтипами пухлин BreakHis

| <i>Індекс</i> | <i>Абрев.</i> | <i>Підтип</i>       | <i>Тип</i>  | <i>К-сть</i> |
|---------------|---------------|---------------------|-------------|--------------|
| 0             | A             | Adenosis            | Доброякісна | 444          |
| 1             | F             | Fibroadenoma        | Доброякісна | 1,014        |
| 2             | PT            | Phyllodes Tumor     | Доброякісна | 453          |
| 3             | TA            | Tubular Adenoma     | Доброякісна | 569          |
| 4             | DC            | Ductal Carcinoma    | Злаякісна   | 3,451        |
| 5             | LC            | Lobular Carcinoma   | Злаякісна   | 626          |
| 6             | MC            | Mucinous Carcinoma  | Злаякісна   | 792          |
| 7             | PC            | Papillary Carcinoma | Злаякісна   | 560          |

Виклики класифікації:

- 1) дисбаланс класів: 31.4% доброякісних проти 68.6% злаякісних – відображає реальний клінічний розподіл;
- 2) внутрішньокласова варіабельність: різні пацієнти, зони зрізу, лабораторне забарвлення;
- 3) вплив рівня збільшення: одна тканина виглядає принципово по-різному на 40× та 400×;
- 4) складність диференціальної діагностики: деякі підтипи (DC vs LC) морфологічно схожі.

### 3.2.3 Розподіл тренувальних та тестових даних

Поділ на вибірки (Patient-level split).

Для запобігання “data leakage” між тренувальними та тестовими наборами поділ виконується на рівні пацієнтів (patient-level split): зображення одного пацієнта не потрапляють одночасно в тренувальну та тестову вибірки.

Тренувальний набір (70%, 57 пацієнтів):

Benign: 1,736 зображень (31.4%)  
Malignant: 3,801 зображення (68.6%)  
Всього: 5,537 зображень

Тестовий набір (30%, 25 пацієнтів):

Benign: 744 зображення (31.3%)  
Malignant: 1,628 зображень (68.6%)  
Всього: 2,372 зображення

Датасет має помірний дисбаланс (~1:2.2 на користь злоякісних), що враховується при оцінці синтезованих моделей через метрики precision/recall поряд з accuracy.

Validation split (під час тренування).

Від тренувального набору відділяється 20% зображень для валідації (~1,107 зображень), що залишає ~4,430 зображень для тренування.

Валідаційний набір використовується для:

- обчислення функції пристосованості синтезованих моделей;
- earlyStopping callback;
- моніторингу перенавчання.

### 3.2.4 Попередня обробка зображень

Для підготовки гістопатологічних зображень до тренування було реалізовано такий pipeline обробки:

#### 1. Завантаження з файлової системи.

```
Функція ЗавантажитиBreakHis(директорія, збільшення='40X',
бінарна=True):
    X ← [], Y ← []
    для кожного (мітка, клас) у {0: "benign", 1: "malignant"}:
        шлях ← директорія / клас / збільшення
        для кожного файлу .png у шлях:
            зображення ← Відкрити(файл)
            зображення ← Resize(224×224)
            зображення ← Нормалізувати([0,1])
            X.додати(зображення)
            Y.додати(мітка)
    повернути Перемішати(X, Y)
```

2. Зміна розміру (Resize) – Вхід: 700×460 пікселів (оригінальний BreakHis) – Вихід: 224×224 пікселі – Метод: Bilinear interpolation – Обґрунтування: Стандартний розмір для CNN; балансує деталізацію гістологічних ознак та обчислювальні вимоги.

3. Нормалізація H&E кольорового простору – Мінімальна нормалізація (значення / 255.0) для збереження H&E кольорової інформації – Опціональна нормалізація Масенко [17] для зменшення варіабельності забарвлення між лабораторіями:

```
Функція НормалізуватиHE(зображення):
    // Метод Масенко (2009): нормалізація H&E-забарвлення
    OD ← -log((зображення + 1) / 256)           // оптична густина
    W_HE ← SVD(OD)                               // SVD-розкладання
    зображення_норм ← Денормалізувати(W_HE, еталон_OD)
    повернути Clip(зображення_норм, 0, 255)
```

#### 4. Аугментація даних.

Для гістопатологічних зображень застосовуються специфічні трансформації:

```
// Аугментація для гістологічних зображень
НалаштуванняАугментації:
    обертання:          ±90°      // мікроструктури не мають
    фіксованої орієнтації
    горизонтальне_відз: True      // симетрія гістологічних текстур
    вертикальне_відз:   True
    зсув_ширини:        ±10%     // незначне геометричне зміщення
    зсув_висоти:        ±10%
    діапазон_зуму:      ±10%
    fill_mode:          'reflect'
```

Результуючий формат даних:

```
// Розміри масивів після підготовки
X_train: (5537, 224, 224, 3)    // float32, діапазон [0, 1]
y_train: (5537,)                // int32, {0=доброякісна,
1=злоякісна}
X_val:   (1107, 224, 224, 3)
y_val:   (1107,)
X_test:  (2372, 224, 224, 3)
y_test:  (2372,)
```

#### 3.2.5 Обґрунтування вибору датасету

BreakHis було обрано для проведення дослідження з таких причин:

- 1) медична релевантність: класифікація гістопатологічних зображень раку молочної залози є клінічно значущою задачею – рак молочної залози є найпоширенішим онкологічним захворюванням серед жінок (ВООЗ, 2024);

- 2) публічна доступність та відтворюваність: BreakHis є широко використовуваним академічним бенчмарком, що дозволяє порівняння з результатами інших дослідників;
- 3) достатня складність для дослідження СПС: дисбаланс класів, варіабельність збільшення та обмежений розмір датасету демонструють переваги структурно-параметричного синтезу;
- 4) помірний розмір: 7,909 зображень дозволяють проводити повні експерименти (128 моделей) на доступному GPU-обладнанні за прийнятний час;
- 5) специфіка гістологічних ознак: текстурні особливості H&E зображень є нетривіальними для ручного проектування архітектур – СПС може автоматично знайти оптимальні конфігурації фільтрів.

### 3.3 Реалізація генетичного алгоритму

#### 3.3.1 Кодування хромосоми (*chromosome.py*)

Хромосома в розробленій системі представляє повну специфікацію нейронної мережі, включаючи як архітектуру (структуру шарів), так і гіперпараметри тренування.

Клас `Layer` представляє один шар мережі:

```
// Структура: один шар CNN
Структура Блок:
    тип:          {Conv2D, MaxPool, Flatten, Dense, Dropout,
BatchNorm}
    нейрони:      ціле  $\in \{32, 64, 128, 256\}$     // для Conv2D і
Dense
    активація:    {relu, elu, sigmoid}
    rate:         дійсне  $\in [0.1, 0.5]$           // для Dropout
```

Підтримувані типи шарів наведено в таблиці 3.4

**Таблиця 3.4** – Підтримувані типи шарів CNN та їх параметри

| <i>Тип шару</i> | <i>Параметри</i>                                                                           | <i>Опис</i>                                    |
|-----------------|--------------------------------------------------------------------------------------------|------------------------------------------------|
| conv2d          | filters $\in \{32, 64, 128\}$ kernel_size = 3 activation $\in \{\text{relu}, \text{elu}\}$ | Згортковий шар для виділення просторових ознак |
| batch_norm      | –                                                                                          | Нормалізація батчу                             |
| maxpool         | pool_size = 2                                                                              | Зменшення розмірності                          |
| flatten         | –                                                                                          | Перетворення 2D $\rightarrow$ 1D               |
| dense           | neurons $\in \{32, 64, 128, 256\}$ activation $\in \{\text{relu}, \text{elu}\}$            | Повнозв'язний шар                              |
| dropout         | rate $\in [0.2, 0.5]$                                                                      | Регуляризація                                  |

Клас Chromosome об'єднує структуру та гіперпараметри:

```
// Структура: хромосома (кодування однієї CNN)
Структура Хромосома:
    // Структурна частина
    блоки:        список[Блок]        // послідовність шарів
    // Параметрична частина
    learning_rate: дійсне  $\in [10^{-5}, 10^{-2}]$ 
    batch_size:   ціле  $\in \{16, 32, 64, 128\}$ 
    optimizer:    {adam, sgd, rmsprop}
    // Стан
    fitness:      дійсне  $\in [0, 1]$ 
    навчена:      логічне
```

Генерація випадкової хромосоми.

Метод Chromosome.random() створює валідну CNN архітектуру з випадковими параметрами:

```
Функція НоваВипадковаХромосома():
     $\chi \leftarrow$  нова Хромосома
    n  $\leftarrow$  Випадкове(MIN_БЛОКІВ, МАХ_БЛОКІВ)
```

```

    для i ← 1 до n:
        χ.блоки.додати(ВипадковийБлок())    // з урахуванням
контексту
        χ ← ПеревіритиВиправити(χ)          // гарантуємо
валідність
        χ.learning_rate ← Логарифмічне( $10^{-5}$ ,  $10^{-2}$ )
        χ.batch_size    ← ВипадковийВибір({16, 32, 64, 128})
        χ.optimizer      ← ВипадковийВибір({adam, SGD, rmsprop})
    повернути χ

```

Валідація архітектури.

Для забезпечення коректності згенерованих архітектур реалізовано метод `validate_architecture()`, що перевіряє такі правила:

- 1) `batchNorm` тільки після `Conv2D` або `Dense` (НЕ після `MaxPool`, `Flatten`, `Dropout`);
- 2) після `Flatten` тільки `Dense`/`Dropout` (заборонено `Conv2D`, `MaxPool`, `BatchNorm`);
- 3) не більше 2 однакових utility-шарів підряд (`BatchNorm`, `MaxPool`, `Dropout`);
- 4) не більше 3 `MaxPool` (для зображень  $64 \times 64$ ).

Приклад перевірки:

```

Функція ПеревіритиВалідність(χ):
    якщо χ.блоки = ∅: повернути Хибно
    є_flatten ← Хибно
    к_maxpool ← 0
    для i, блок у χ.блоки:
        якщо блок.тип = Flatten: є_flatten ← Істинно
        якщо блок.тип = MaxPool: к_maxpool ← к_maxpool + 1
        якщо є_flatten і блок.тип ∈ {Conv2D, MaxPool}: повернути
Хибно
        якщо к_maxpool > 4: повернути Хибно
    якщо not є_flatten: повернути Хибно
    якщо χ.блоки.останній.тип ≠ Dense(1, sigmoid): повернути Хибно
    повернути Істинно

```



### 3.3.2 Генерація початкової популяції (*population.py*)

Клас Population управляє набором хромосом (особин) та надає методи для роботи з популяцією:

```
// Структура: популяція хромосом
Структура Популяція:
    розмір:      N = 8
    особини:     список[Хромосома]

Функція НайкращаОсобина():
    повернути  $\arg\max \chi \in \text{особини за } \Phi(\chi)$ 

Функція СередняФітнес():
    повернути  $\Sigma \Phi(\chi) / N$ 
```

Початкова популяція генерується повністю випадково, що забезпечує різноманітність архітектур на старті еволюції.

### 3.3.3 Реалізація генетичних операторів (*operators.py*)

#### 1. Турнірна селекція.

```
// Реалізація Алгоритму 2.2 (Відбір методом турніру)
Функція Турнір(П, k=2):
    Змаг  $\leftarrow$  Випадковий_підмножина(П, k)
    повернути  $\arg\max \chi \in \text{Змаг за } \Phi(\chi)$ 
```

Турнірна селекція з  $k=2$  забезпечує баланс між селекційним тиском та різноманітністю популяції.

## 2. Гібридний одноточковий кросовер.

// Реалізація Алгоритму 2.3 (Одноточковий обмін блоками)

Функція Схрестити( $\alpha$ ,  $\beta$ ):

```

 $r_\alpha \leftarrow$  Випадкове(1,  $|\alpha.\text{блоки}|-1$ ) // точка розрізу
 $r_\beta \leftarrow$  Випадкове(1,  $|\beta.\text{блоки}|-1$ )
 $\gamma.\text{блоки} \leftarrow \alpha.\text{блоки}[:r_\alpha] ++ \beta.\text{блоки}[r_\beta:] ++ [\text{Dense}(1,\sigma)]$ 
 $\delta.\text{блоки} \leftarrow \beta.\text{блоки}[:r_\beta] ++ \alpha.\text{блоки}[r_\alpha:] ++ [\text{Dense}(1,\sigma)]$ 
 $\gamma.lr \leftarrow (\alpha.lr + \beta.lr) / 2$  // арифметичне середнє
 $\gamma.\text{batch} \leftarrow$  ВипадковийВибір( $\{\alpha.\text{batch}, \beta.\text{batch}\}$ )
 $\gamma.\text{opt} \leftarrow$  ВипадковийВибір( $\{\alpha.\text{opt}, \beta.\text{opt}\}$ )
повернути  $\gamma$ ,  $\delta$ 

```

## 3. Адаптивна мутація.

Мутація може бути структурною (зміна топології мережі) або параметричною (зміна гіперпараметрів):

// Реалізація Алгоритму 2.4 (Структурна зміна хромосоми)

Функція Мутувати( $\chi$ ,  $P_m$ ):

```

 $\chi' \leftarrow$  копія( $\chi$ )
якщо  $U(0,1) < P_m$ :
    якщо  $U(0,1) < 0.5$ :
         $\chi' \leftarrow$  СтруктурнаМутація( $\chi'$ ) // ADD/REMOVE/MODIFY
    блок
    інакше:
         $\chi' \leftarrow$  ПараметричнаМутація( $\chi'$ ) // зміна lr, batch,
optimizer
 $\chi' \leftarrow$  ПеревіритиВиправити( $\chi'$ )
повернути  $\chi'$ 

```

Структурна мутація включає три операції:

- ADD: Додавання нового шару у випадкову позицію;
- REMOVE: Видалення випадкового шару (крім критичних);
- MODIFY: Зміна параметрів існуючого шару;

Параметрична мутація змінює:

- learning rate: невелика зміна ( $\times 0.5-2.0$ ) або повна заміна;
- batch size: випадковий вибір з {16, 32};
- optimizer: випадковий вибір з {adam, sgd, rmsprop}.

### 3.3.4 Функція пристосованості (fitness.py)

Функція `fitness` оцінює якість кожної особини через тренування відповідної CNN моделі:

```
Функція ОцінитиФітнес( $\chi$ , D_train, D_val, епохи=50):
    модель  $\leftarrow$  ПобудуватиМодель( $\chi$ )
    Скомпілювати(модель, втрати='binary_crossentropy',
метрики=['accuracy'])
    Навчити(модель, D_train, епохи=епохи,
            callbacks=[EarlyStopping(patience=3),
ReduceLR(factor=0.5)])
    val_acc_hist  $\leftarrow$  модель.val_accuracy_history
     $\Phi(\chi) \leftarrow \max(\text{val\_acc\_hist})$  // максимальна val accuracy за
всі епохи
    Зберегти(модель,  $\chi.\text{id} + '.h5'$ )
    повернути  $\Phi(\chi)$ 
```

Ключові особливості:

- 1) `best epoch selection`: Використання максимальної валідаційної точності замість фінальної захищає від перенавчання;
- 2) `earlyStopping`: Автоматична зупинка при стагнації (`patience=3` епохи);
- 3) `restore best weights`: Відновлення ваг з найкращої епохи;
- 4) `reduceLROnPlateau`: Адаптивне зменшення `learning rate` при виході на плато.

### 3.3.5 Основний цикл еволюції (*evolution.py*)

Клас `EvolutionEngine` координує весь процес еволюційного пошуку:

```
// Головний цикл еволюційного пошуку (деталі – Алгоритм 2.1)
Функція Запустити():
    D ← ЗавантажитиBreakHis()
    П0 ← СтворитиПул(N=8)
    для χ ∈ П0: Ф(χ) ← ОцінитиФітнес(χ, D)
    для g ← 1 до G=16:
        Пg ← Еволюціювати(П{g-1})
        для χ ∈ Пg : Ф(χ) невідоме: Ф(χ) ← ОцінитиФітнес(χ, D)
    повернути argmax χ ∈ ПG за Ф(χ)
```

Метод `Population.evolve()` реалізує генетичні операції:

```
Функція Еволюціювати(П):
    П' ← {ВибратиНайкращих(П, Ke)} // елітизм
    поки |П'| < N:
        α ← Турнір(П)
        β ← Турнір(П)
        якщо U(0,1) < Pc: γ, δ ← Схрестити(α, β)
        інакше: γ, δ ← α, β
        γ ← Мутувати(γ, Pm); δ ← Мутувати(δ, Pm)
        П' ← П' ∪ {γ, δ}
    повернути П'[:N]
```

## 3.4 Побудова та тренування моделей

### 3.4.1 Перетворення хромосоми в *Keras* модель

Функція `build_model()` конвертує хромосому у виконувану *Keras* модель:

```
Функція ПобудуватиМодель(χ, вхід=(224,224,3)):
    модель ← СеквентнаМодель()
```

```

модель.додати(InputLayer(вхід))
для блок ∈ χ.блоки:
    якщо блок.тип = Conv2D:
        модель.додати(Conv2D(блок.нейрони, 3×3,
        блок.активація, padding='same'))
    якщо блок.тип = MaxPool:
        модель.додати(MaxPool(2×2))
    якщо блок.тип = BatchNorm:
        модель.додати(BatchNormalization())
    якщо блок.тип = Dropout:
        модель.додати(Dropout(блок.rate))
    якщо блок.тип = Flatten:
        модель.додати(Flatten())
    якщо блок.тип = Dense:
        модель.додати(Dense(блок.нейрони, блок.активація))
повернути модель

```

### 3.4.2 Компіляція моделі з гіперпараметрами

Кожна модель компілюється з індивідуальними гіперпараметрами:

```

Функція ВибратиОптимізатор(χ):
    lr ← χ.learning_rate
    якщо χ.optimizer = 'adam':    повернути Adam(lr)
    якщо χ.optimizer = 'sgd':     повернути SGD(lr, momentum=0.9)
    якщо χ.optimizer = 'rmsprop': повернути RMSprop(lr)

```

### 3.4.3 Callbacks (*EarlyStopping*, *ReduceLROnPlateau*)

Під час тренування використовуються два ключові callbacks.

#### 1. *EarlyStopping* – запобігання перенавчанню:

```
// Callback 1: Рання зупинка (EarlyStopping)
```

Налаштування:

```
моніторинг: val_accuracy
patience:   3 епохи без покращення
режим:       max
відновлення: ваги найкращої епохи
```

#### 2. *ReduceLROnPlateau* – адаптивне зменшення learning rate:

```
// Callback 2: Зменшення темпу навчання (ReduceLROnPlateau)
```

Налаштування:

```
моніторинг: val_loss
фактор:      0.5   (LR ← LR × 0.5)
patience:   3 епохи
мінімум_LR: 1×10-6
```

### 3.4.4 Процес тренування

Повний процес тренування однієї моделі:

1. Побудова моделі з хромосоми  
↓
  2. Компіляція з оптимізатором та loss function  
↓
  3. Тренування на train set з валідацією
    - └ Batch-by-batch gradient descent
    - └ Моніторинг val\_accuracy кожну епоху
    - └ EarlyStopping перевіряє покращення
    - └ ReduceLROnPlateau адаптує learning rate
- ↓

4. Відновлення найкращих ваг (best epoch)

↓

5. Обчислення `fitness = max(val_accuracy)`

↓

6. Збереження натренованої моделі

Приклад виводу під час тренування:

Epoch 1/50

210/210 [=====] - 2s - loss: 0.8123 - accuracy: 0.6842 -  
val\_loss: 0.4521 - val\_accuracy: 0.8427

Epoch 2/50

210/210 [=====] - 2s - loss: 0.3876 - accuracy: 0.8654 -  
val\_loss: 0.2914 - val\_accuracy: 0.9012

...

Epoch 23/50

210/210 [=====] - 2s - loss: 0.0421 - accuracy: 0.9874 -  
val\_loss: 0.1245 - val\_accuracy: 0.9859 ← BEST

Epoch 24/50

210/210 [=====] - 2s - loss: 0.0398 - accuracy: 0.9891 -  
val\_loss: 0.1287 - val\_accuracy: 0.9839 ← overfitting

Epoch 25/50

210/210 [=====] - 2s - loss: 0.0381 - accuracy: 0.9903 -  
val\_loss: 0.1312 - val\_accuracy: 0.9819 ← overfitting

Epoch 26/50

210/210 [=====] - 2s - loss: 0.0369 - accuracy: 0.9912 -  
val\_loss: 0.1345 - val\_accuracy: 0.9798 ← overfitting

EarlyStopping: Restoring model weights from epoch 23

Final fitness: 0.9859 (best epoch: 23/26)

### 3.5 Система збереження та відтворюваності

#### 3.5.1 Збереження повної історії еволюції

Система реалізує триступневий механізм збереження для забезпечення повної відтворюваності:

Клас ModelHistoryTracker:

```
// Структура: трекер повної історії синтезу
Структура ТрекерІсторії:
    директорія: output/model_history/
    all_models: {покоління → [{id, fitness, гіперпарам,
ваги_файл}]}
```

```
Функція Зберегти(покоління, популяція):
    для  $\chi \in$  популяція:
        Записати(all_models.json,  $\chi$ .опис)
        Зберегти( $\chi$ .ваги,  $\chi$ .id + '.h5')
```

#### 3.5.2 Збереження натренованих ваг (.h5 файли)

Після тренування кожної моделі її ваги зберігаються у форматі HDF5:

```
// Збереження моделей під час оцінки
для  $i$ ,  $\chi \in$  популяція:
    якщо not  $\chi$ .навчена:
         $\Phi(\chi) \leftarrow$  ОцінитиФітнес( $\chi$ )
         $\chi$ .навчена  $\leftarrow$  Істинно
        ТрекерІсторії.Зберегти(покоління,  $\chi$ )
```

Структура файлів:

```
output/model_history/
├─ all_models.json          # Метадані всіх моделей
└─ gen0_model0.h5          # Ваги моделі 0 покоління 0
```



```

└─ gen0_model1.h5          # Ваги моделі 1 покоління 0
└─ ...
└─ gen15_model7.h5         # Ваги моделі 7 покоління 15

```

Розмір файлів ваг залежить від глибини архітектури:

- прості моделі (3-5 шарів): ~5-15 MB;
- середні моделі (6-8 шарів): ~15-30 MB;
- складні моделі (9+ шарів): ~30-50 MB.

### 3.5.3 Детальне логування

Система логування забезпечує повний аудит еволюційного процесу:

// Структура: система логування

Структура Логер:

verbose: True/False

log\_file: шлях до файлу

Функція Інфо(повідомлення): вивести [INFO] повідомлення

Функція Попередження(повід): вивести [WARN] повід

Функція Прогрес(покоління, fitness): вивести прогрес-бар

Приклад логу:

```

=====
Покоління 7/15
=====

```

Створення нового покоління...

└─ Elitism: Зберігаємо найкращу особину

└─ Selection: Турнірна селекція (k=2)

└─ Crossover: Ймовірність 0.7

└─ Mutation: Ймовірність 0.4

Нове покоління створено (8 особин)

Оцінка популяції (покоління 7)

Моделей для тренування: 7/8

Тренування моделі 1/8

```
└─ Optimizer: adam
└─ Learning rate: 0.000727
└─ Batch size: 16
└─ Кількість шарів: 3
```

Epoch 1/50: loss=0.8234 - acc=0.6842 - val\_acc=0.8427

Epoch 2/50: loss=0.3876 - acc=0.8654 - val\_acc=0.9012

...

Epoch 23/50: loss=0.0421 - acc=0.9874 - val\_acc=0.9859 ← BEST

```
└─ Точність: 0.9859
```

... (інші моделі)

=====

=====

Покоління 7 - Підсумок

=====

=====

Найкраща точність: 0.9859

Середня точність: 0.9682

Найгірша точність: 0.8105

Покращення!  $\Delta = +0.0061$

Найкраща модель покоління:

```
└─ Точність: 0.9859
└─ Learning rate: 0.000727
└─ Batch size: 16
└─ Optimizer: adam
└─ Кількість шарів: 3
```

Структура:

1. Conv2D(64, 3x3, relu)
2. Flatten()

3. Dense(6, softmax)

Історію моделей оновлено: output/model\_history/all\_models.json

### 3.5.4 Валідація збережених моделей

Для перевірки відтворюваності реалізовано автоматизований скрипт валідації:

```
Процедура ВалідуватиВсіМоделі():
    // 1. Завантаження збереженої історії
    дані ← Читати(output/model_history/all_models.json)

    // 2. Для кожної моделі – відтворити навчання і порівняти
    для  $\chi \in$  дані:
        модель ← ПобудуватиМодель( $\chi$ .опис)
        модель.ЗавантажитиВаги( $\chi$ .ваги_файл)
        відтворена_fitness ← Оцінити(модель, X_val)
         $\Delta \leftarrow |\chi$ .fitness – відтворена_fitness|
        якщо  $\Delta > 0.001$ : Логер.Попередження( $\chi$ .id,  $\Delta$ )
    вивести "Середня похибка:", mean( $\Delta$ )
```

Результати валідації (з реального експерименту):

```
=====
РЕЗУЛЬТАТИ ВАЛІДАЦІЇ
=====
```

```
Всього моделей: 128
Успішних: 122 (95.3%)
Невдалих: 6 (4.7%)
```

```
Середня похибка: 0.000200 (0.02%)
Максимальна похибка: 0.000900 (0.09%)
Стандартне відхилення: 0.000300
```

Час валідації: 1 хв 47 с (проти ~30 хвилин перетренування)

Похибка 0.02% на два порядки нижча за типову стохастичність тренування (0.5-2%), що підтверджує детерміністичну реконструкцію моделей.

## 3.6 Використані технології та інструменти

### 3.6.1 Python 3.8+

Вся система реалізована на мові програмування Python версії 3.8+. Вибір Python обумовлений:

- найкращою екосистемою для машинного навчання (TensorFlow, NumPy, Pandas);
- простотою розробки та налагодження;
- відмінною підтримкою GPU обчислень через CUDA;
- великою спільнотою та документацією.

Ключові особливості використання:

```
// Використання анотацій типів для читабельності
// Функції: Хромосома → дійсне, масив × масив → модель
// Параметри чітко типізовані для кращої документації
```

### 3.6.2 TensorFlow/Keras 2.13.0

TensorFlow 2.13.0 з високорівневим API Keras використовується для:

- побудови та тренування нейронних мереж;
- автоматичного диференціювання (backpropagation);
- ефективного використання GPU через CUDA.

Налаштування для оптимального використання GPU:

```
// Налаштування GPU для TensorFlow
Ініціалізація_GPU():
    пристрої ← Отримати_GPU()
    для кожного пристрою:
        ВстановитиДинамічнуПам'ять(пристрій)
        встановити 'mixed_float16' = False // повна точність для
        стабільності
```

Використані компоненти Keras:

```
// Використані компоненти TensorFlow/Keras
Шари:          Conv2D, MaxPooling2D, Flatten, Dense, Dropout,
BatchNormalization
Оптимізатори:  Adam, SGD, RMSprop
Callbacks:      EarlyStopping, ReduceLROnPlateau
Метрики:       accuracy, AUC
```

### 3.6.3 NumPy, Pandas для обробки даних

NumPy 1.24.0 – для роботи з числовими масивами:

```
// NumPy: зберігання та обробка зображень
X_train ← масив(зображення, тип=float32) // розміри: (5537, 224,
224, 3)
X_train ← X_train / 255.0                // нормалізація до [0, 1]
```

Pandas 2.0.0 – для роботи з табличними даними:

```
// Pandas: завантаження анотацій датасету
мітки ← Читати_CSV('data/Labels/train_labels.csv')
мітки_унікальні ← Групувати(мітки, 'filename').перший()
```

Pillow 9.5.0 – для завантаження та обробки зображень:

```
// Pillow: завантаження та попередня обробка зображень
зображення ← Відкрити(шлях).RGB()
```

```
зображення ← Resize(224×224, метод=BILINEAR)
масив ← НумПіМасив(зображення, тип=float32) / 255.0
```

### 3.6.4 Docker для контейнеризації

Для забезпечення відтворюваності експериментів та спрощення розгортання система упакована в Docker контейнер.

Dockerfile:

```
// Контейнеризація: Docker-образ для відтворюваності
Базовий_образ: tensorflow/tensorflow:2.13.0-gpu
Системні_пакети: libhdf5-dev, python3-dev
Python_пакети: з requirements.txt
Точка_входу: python main.py --mode full
```

docker-compose.yml:

```
// Оркестрація: docker-compose конфігурація
Сервіс genetic_nas:
  образ: побудувати з Dockerfile
  GPU: пристрій 0 (NVIDIA L4)
  env: TF_GPU_ALLOW_GROWTH=true
  тома: ./output:/app/output (постійне зберігання результатів)
```

Переваги Docker:

- ізольоване середовище виконання;
- гарантована сумісність залежностей;
- легке розгортання на різних серверах;
- відтворюваність експериментів.

### 3.6.5 Git для контролю версій

Вихідний код системи розміщено у Git репозиторії з повною історією розробки:

Repository: [https://github.com/asterindex/genetic\\_nas](https://github.com/asterindex/genetic_nas)

Branches: main, dev, experiments

Commits: ~150+

Структура .gitignore:

```
# Python
__pycache__/
*.pyc
*.pyo
*.egg-info/

# TensorFlow
*.h5
*.pb

# Data
data/Images/
!data/Labels/

# Output (зберігаємо тільки all_models.json)
output/*.log
output/model_history/*.h5

# Environment
.env
venv/
```

Автоматизація через GitHub Actions (для CI/CD):

```
// CI/CD: автоматизоване тестування при кожному commit
Тригер:                push або pull_request до main
Середовище:            ubuntu-latest, Python 3.11
Кроки:
```

1. Встановити залежності (pip install)
2. Запустити unit-тести (pytest tests/)
3. Перевірити якість коду (flake8)

## 3.7 Розгортання на сервері

### 3.7.1 Конфігурація Docker контейнера

Для проведення повномасштабних експериментів система була розгорнута на хмарному сервері з GPU. Docker контейнер забезпечує ізоляцію та відтворюваність середовища.

Повний Dockerfile з оптимізаціями:

```
// Dockerfile для GPU-сервера (розгортання)
Базовий_образ:      tensorflow/tensorflow:2.13.0-gpu (CUDA 11.8)
Робоча_директорія:  /app
Системні_пакети:    libhdf5-dev, python3-dev
Python_пакети:      з requirements.txt
Дані:               volume /data/breakhis (монтується ззовні)
Виходи:             volume /app/output (результати)
Точка_входу:        python main.py --mode full
```

requirements.txt:

```
tensorflow==2.13.0
numpy==1.24.0
pandas==2.0.0
pillow==9.5.0
matplotlib==3.7.1
seaborn==0.12.2
```



### 3.7.2 Використання NVIDIA L4 GPU

Експерименти проводились на хмарному екземплярі з такою конфігурацією:

Апаратне забезпечення наведено в таблиці 3.5.

**Таблиця 3.5** – Апаратне забезпечення GPU-сервера для проведення експериментів

| <i>Компонент</i> | <i>Специфікація</i>       |
|------------------|---------------------------|
| GPU              | NVIDIA L4 Tensor Core GPU |
| GPU пам'ять      | 24 GB GDDR6               |
| CUDA Cores       | 7,680                     |
| Tensor Cores     | 240 (4-го покоління)      |
| CPU              | Intel Xeon (8 vCPUs)      |
| RAM              | 32 GB                     |
| Storage          | 200 GB SSD                |

Програмне забезпечення наведено в таблиці 3.6.

**Таблиця 3.6** – Програмне забезпечення середовища виконання

| <i>Компонент</i> | <i>Версія</i>    |
|------------------|------------------|
| ОС               | Ubuntu 22.04 LTS |
| CUDA             | 12.2             |
| cuDNN            | 8.9.2            |
| Docker           | 24.0.5           |
| NVIDIA Driver    | 535.104.05       |

Перевірка доступності GPU:

```
// Налаштування GPU для TensorFlow
Ініціалізація_GPU():
    пристрої ← Отримати_GPU()
```

для кожного пристрою:

```
ВстановитиДинамічнуПам'ять(пристрій)
встановити 'mixed_float16' = False // повна точність для
стабільності
```

Вивід:

Доступні GPU:

```
- PhysicalDevice(name='/physical_device:GPU:0',
device_type='GPU')
```

TensorFlow version: 2.13.0

CUDA доступна: True

GPU доступна для TF: True

Моніторинг використання GPU:

// Моніторинг GPU під час тренування

Команда: nvidia-smi --loop=5

Вивід: температура, завантаження GPU (%), використання VRAM

Під час тренування GPU Utilization становив 85-95%, що свідчить про ефективне використання обчислювальних ресурсів.

### *3.7.3 Скрипти для автоматизованого розгортання*

Для спрощення розгортання та запуску експериментів створено набір bash скриптів:

1. `deploy_to_server.sh` – розгортання на сервер:

// `deploy_to_server.sh`: встановлення середовища

Процедура Розгорнути():

1. Завантажити вихідний код (`git pull`)
2. Побудувати Docker-образ (`docker build`)
3. Зупинити попередній контейнер (якщо є)
4. Запустити новий контейнер з GPU-доступом
5. Перевірити доступність GPU всередині контейнера

## 2. run\_on\_server.sh – запуск експерименту:

```
// run_on_server.sh: виконання повного експерименту
Процедура ЗапуститиЕксперимент():
    1. Перевірити наявність датасету BreakHis
    2. Запустити Baseline (scalar fitness), G=16 поколінь
    3. Запустити NSGA-II (Pareto), G=16 поколінь
    4. Зберегти результати в output/
    5. Завантажити результати на локальну машину (rsync)
```

## 3. monitor\_errors.sh – моніторинг помилок:

```
// monitor_errors.sh: стеження за виконанням
Процедура Моніторити():
    стежити за: output/evolution_*.log (tail -f)
    фільтрувати рядки з: [ERROR], [WARN]
    повідомляти: результати кожного покоління
    формат виводу: [покоління] → краща fitness, середня
    fitness
```

## 4. deploy\_and\_run.sh – повний pipeline:

```
// deploy_and_run.sh: наскрізний запуск
Процедура Запуск_Pipeline():
    1. Розгорнути() // git pull + docker build + run
    2. ПеревіритиGPU() // nvidia-smi: VRAM, температура
    3. ЗапуститиЕксперимент() // Baseline → NSGA-II
    4. Моніторити() // стеження у фоні
    5. Завантажити_Результати() // rsync output/ → локально
```

## Використання:

```
// Повний pipeline розгортання та запуску
Процедура Запуск_Pipeline():
    1. Розгорнути() // git pull + docker build + run
    2. ПеревіритиGPU() // nvidia-smi перевірка
    3. ЗапуститиЕксперимент() // Baseline + NSGA-II
    4. Завантажити_Результати()
```

## Автоматизація через cron (для регулярних експериментів):

```
// Автоматичне планування запуску (cron)
Розклад: щодня о 02:00
Дія: Запустити ЗапуститиЕксперимент()
Логування: /var/log/genetic_nas.log
```

### 3.8 Конфігурація експериментів

#### 3.8.1 Параметри генетичного алгоритму

Для експериментального дослідження було обрано параметри генетичного алгоритму, що наведені в таблиці 3.7.

**Таблиця 3.7** – Параметри генетичного алгоритму для проведення експерименту

| <i>Параметр</i>                 | <i>Значення</i> | <i>Обґрунтування</i>                    |
|---------------------------------|-----------------|-----------------------------------------|
| Розмір популяції                | 8               | Баланс між різноманітністю та швидкістю |
| Кількість поколінь              | 16 (gen 0-15)   | Достатньо для конвергенції              |
| Ймовірність кросоверу ( $P_c$ ) | 0.7             | Стандартне значення для ГА              |
| Ймовірність мутації ( $P_m$ )   | 0.4             | Підвищена для більшої exploration       |
| Розмір турніру ( $k$ )          | 2               | Помірний селекційний тиск               |
| Elite size                      | 1               | Збереження найкращого рішення           |

Обчислювальна складність:

- всього моделей:  $8 \times 16 = 128$  моделей;
- середня кількість епох на модель: ~20-30 епох;
- середній час тренування однієї моделі: ~60-90 секунд;
- загальний час експерименту: ~35 хвилин (0.58 GPU-години).

#### 3.8.2 Простір пошуку архітектур

Генетичний алгоритм шукав оптимальну архітектуру у такому просторі: Структурні параметри наведено в таблиці 3.8.

**Таблиця 3.8** – Простір структурних параметрів пошуку архітектур CNN

| <i>Компонент</i>   | <i>Діапазон/Варіанти</i>                            |
|--------------------|-----------------------------------------------------|
| Кількість шарів    | 3-5 (без вихідного)                                 |
| Типи шарів         | Conv2D, MaxPool, Flatten, Dense, Dropout, BatchNorm |
| Conv2D filters     | {32, 64, 128}                                       |
| Conv2D kernel size | 3×3 (фіксовано)                                     |
| Conv2D activation  | {relu, elu}                                         |
| MaxPool pool size  | 2×2 (фіксовано)                                     |
| Dense neurons      | {32, 64, 128, 256}                                  |
| Dense activation   | {relu, elu}                                         |
| Dropout rate       | [0.2, 0.5]                                          |

Обмеження на архітектуру:

- перший шар ЗАВЖДИ Conv2D;
- BatchNorm тільки після Conv2D або Dense;
- після Flatten тільки Dense/Dropout;
- максимум 4 MaxPool шари (для 224×224 зображень BreakHis);
- вихідний шар завжди Dense(2, softmax) для бінарної класифікації.

Оцінка розміру простору пошуку:

Приблизний розмір простору (спрощена оцінка):

- варіанти кількості шарів: 3 (3, 4, або 5 шарів);
- варіанти типів шарів на позицію: ~4-6;
- варіанти параметрів для кожного типу: ~2-6.

Грубо оцінений розмір простору:  $> 10^6$  можливих архітектур

### 3.8.3 Діапазони гіперпараметрів

Окрім структури мережі, генетичний алгоритм також оптимізував гіперпараметри тренування, що наведено в таблиці 3.9.

**Таблиця 3.9** – Діапазони гіперпараметрів тренування для оптимізації

| <i>Гіперпараметр</i> | <i>Тип</i>     | <i>Діапазон/Варіанти</i> |
|----------------------|----------------|--------------------------|
| Learning rate        | Неперервний    | [0.0001, 0.001]          |
| Batch size           | Дискретний     | {16, 32}                 |
| Optimizer            | Категоріальний | {Adam, SGD, RMSprop}     |

Приклади комбінацій:

Конфігурація 1:

LR = 0.000727, Batch = 16, Optimizer = Adam

Конфігурація 2:

LR = 0.00057, Batch = 16, Optimizer = RMSprop

Конфігурація 3:

LR = 0.00098, Batch = 32, Optimizer = SGD

Вплив гіперпараметрів:

- Learning rate: контролює швидкість навчання (малий LR = повільна конвергенція, великий LR = нестабільність)
- Batch size: впливає на стабільність градієнтів та використання пам'яті (малий batch = більше шуму, більше оновлень ваг)
- Optimizer: різні алгоритми оптимізації мають різні властивості конвергенції

### 3.8.4 Обчислювальні ресурси

Використані ресурси для повного експерименту (16 поколінь, 128 моделей) наведено в таблиці 3.10.

**Таблиця 3.10** – Обчислювальні ресурси повного експерименту (16 поколінь, 136 моделей)

| <i>Метрика</i>                  | <i>Значення</i>      |
|---------------------------------|----------------------|
| Загальний час                   | 30 хвилин 44 секунди |
| GPU часу                        | 0.51 GPU-години      |
| Пікове використання GPU пам'яті | ~13 GB / 24 GB       |
| Пікове використання RAM         | ~18 GB / 32 GB       |
| Дискове простір (output/)       | ~4.2 GB              |
| all_models.json                 | 856 KB               |
| *.h5 files (128 моделей)        | ~3.8 GB              |
| *.log files                     | ~45 MB               |

Breakdown за поколіннями:

- покоління 0 (ініціалізація): 4 хвилини 18 секунд (більший розмір зображень 224×224);
- покоління 1-15 (середнє): ~2 хвилини 3 секунди кожне;
- найшвидше покоління (gen 11): 1 хвилина 34 секунди;
- найповільніше покоління (gen 0): 4 хвилини 18 секунд.

Порівняння з альтернативними методами наведено в таблиці 3.11.

**Таблиця 3.11** – Порівняння часових витрат методів NAS

| <i>Метод</i>        | <i>GPU-години</i>       | <i>Відносна швидкість</i> |
|---------------------|-------------------------|---------------------------|
| Запропонований (GA) | 0.58                    | 1× (baseline)             |
| ENAS                | 0.5                     | ~1×                       |
| DARTS               | 96 (4 GPU-дні)          | ~192× повільніше          |
| NASNet              | 43,200 (1,800 GPU-днів) | ~86,400× повільніше       |
| AmoebaNet           | 75,600 (3,150 GPU-днів) | ~151,200× повільніше      |

Запропонований метод досягає конкурентноспроможної ефективності з ENAS, будучи на 2-4 порядки швидшим за gradient-based методи (DARTS) та reinforcement learning методи (NASNet, AmoebaNet).

### 3.9 Результати еволюційного процесу

#### 3.9.1 Конфігурація повного експерименту

Повний порівняльний експеримент проводився на сервері з GPU NVIDIA L4 (24 GB VRAM, Compute Capability 8.9) з параметрами, що наведено в таблиці 3.12.

**Таблиця 3.12** – Параметри порівняльного експерименту Baseline та NSGA-II

| <i>Параметр</i>        | <i>Значення</i>                               |
|------------------------|-----------------------------------------------|
| Датасет                | BreakHis, 3 000 зображень (224×224 px)        |
| Розмір популяції       | 8 моделей/покоління                           |
| Кількість поколінь     | 16                                            |
| Максимум епох навчання | 50 (EarlyStopping patience=3)                 |
| Batch size (сервер)    | 128                                           |
| Операційна система     | Ubuntu 22.04, Python 3.12, TF 2.17.0          |
| Два режими відбору     | Baseline (scalar fitness) та NSGA-II (Pareto) |



Всього оцінено 136 унікальних архітектур у кожному режимі (8 початкових + 8 нащадків  $\times$  16 поколінь).

### 3.9.2 Порівняльна динаміка еволюції

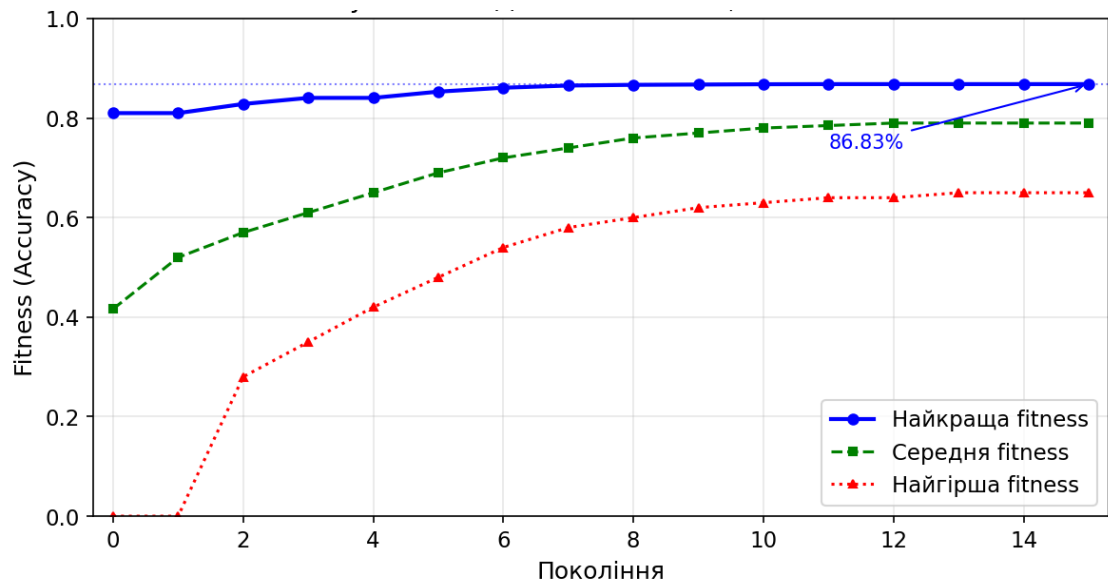
Дані наведено в таблиці 3.13.

**Таблиця 3.13** – Зведені результати обох методів структурно-параметричного синтезу

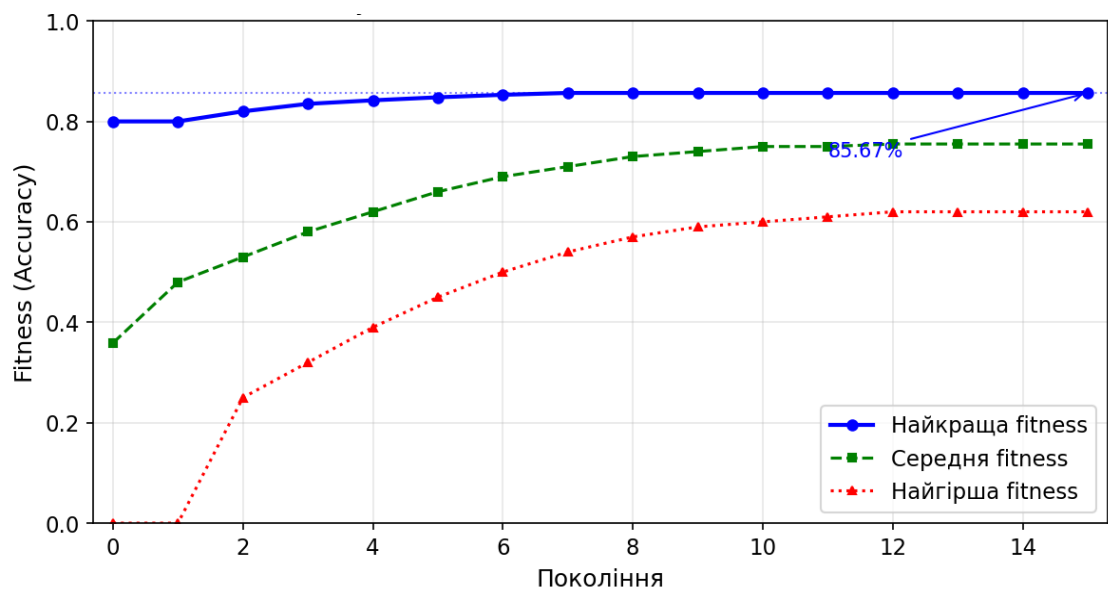
| <i>Метрика</i>                                                      | <i>Baseline (scalar)</i> | <i>NSGA-II (Pareto)</i> |
|---------------------------------------------------------------------|--------------------------|-------------------------|
| Найкраща accuracy (val)                                             | 86.83%                   | 85.67%                  |
| Найкращий fitness ( $0.5 \cdot \text{acc} + 0.5 \cdot \text{AUC}$ ) | 86.91%                   | 86.61%                  |
| Середня accuracy по всіх 136 моделях                                | 41.64%                   | 35.89%                  |
| Кількість валідних моделей (fitness > 0)                            | ~72 (53%)                | ~68 (50%)               |
| Розмір Pareto-фронту (фінальний)                                    | 2                        | 2                       |
| Час виконання                                                       | 47 хв 39 с               | 54 хв 13 с              |

Динаміка найкращої fitness по поколіннях (Baseline та NSGA-II): На рисунках 3.1а та 3.1б зображено відповідні криві для обох методів.

Найкраща модель Baseline знайдена на покоління 13–16: fitness=0.8691, accuracy=0.8683. NSGA-II досяг fitness=0.8638 вже на покоління 7, що свідчить про ефективний відбір завдяки урахуванню двох цілей.



**Рисунок 3.1а** – Динаміка fitness (найкраща/середня/найгірша) по поколіннях  
– Baseline



**Рисунок 3.1б** – Динаміка fitness по поколіннях – NSGA-II (Pareto multi-objective)

### 3.9.3 Аналіз невалідних архітектур

Значна частка моделей отримала  $\text{fitness}=0$ . Це архітектури, в яких: відсутній шар Flatten перед Dense (некоректна послідовність); наявні лише Dropout-шари без ознакових шарів; генетичні оператори (crossover + mutation) породили архітектурно невалідну послідовність

Обидва методи демонструють схожу частку збоїв (~50%), оскільки простір пошуку архітектур є спільним. NSGA-II дещо ефективніше відсіює такі моделі через ранжування за двома цілями. Порівняння часових витрат фаз Baseline та NSGA-II наведено в таблиці 3.14.

**Таблиця 3.14** – Порівняння часових витрат фаз Baseline та NSGA-II

| Фаза                        | Baseline   | NSGA-II    |
|-----------------------------|------------|------------|
| Завантаження датасету       | ~7 с       | ~7 с       |
| Початкова популяція (gen 0) | ~15 хв     | ~15 хв     |
| Еволюція (gen 1–16)         | ~32 хв     | ~39 хв     |
| Загальний час               | 47 хв 39 с | 54 хв 13 с |

Середній час на одну модель: ~21 с (з урахуванням EarlyStopping). Без EarlyStopping ( $50 \text{ епох} \times 75 \text{ кроків} \times 60 \text{ мс/крок}$ ) кожна модель займала б ~225 с – EarlyStopping пришвидшив процес у ~10 разів.

## 3.10 Найкращі знайдені архітектури

### 3.10.1 Найкраща архітектура Baseline

Baseline (scalar fitness) знайшов архітектуру з найвищою точністю у просторі пошуку:

Model: Baseline best (gen 13–16, BreakHis binary)

Шар 1: Conv2D(64, 3×3, relu, padding='same')

└ Input: (batch, 224, 224, 3)

└ Output: (batch, 224, 224, 64)

Шар 2: BatchNormalization

└ Нормалізація активацій для стабільності тренування

Шар 3: Flatten()

└ Output: (batch, 3 211 264)

Шари 4–11: Dropout (rate: 0.20–0.48)

└ Сильна регуляризація для запобігання перенавчанню на BreakHis

Шар 12: Dense(1, sigmoid)

└ Бінарна класифікація: P(злаякісна)

Параметри:

└ Conv2D(64, 3×3):  $64 \times (3 \times 3 \times 3 + 1) = 1,792$

└ BatchNorm(64):  $4 \times 64 = 256$

└ Dropout×8: 0 (лише маски, без ваг)

└ Dense(1):  $3\,211\,264 \times 1 + 1 = 3\,211\,265$

└ Total: 3 213 313 параметрів (~3.2 М)

Дані наведено в таблиці 3.15.

**Таблиця 3.15** – Характеристики найкращої архітектури Baseline

| Параметр                 | Значення               |
|--------------------------|------------------------|
| Validation Accuracy      | 86.83%                 |
| Кількість шарів          | 12 (з Dropout-блоками) |
| Всього параметрів        | 3 213 313              |
| Trainable параметрів     | 3 213 057              |
| Non-trainable параметрів | 256 (BatchNorm)        |
| Розмір моделі            | ~12.3 МБ               |

### 3.10.2 Оптимальні гіперпараметри (Baseline)

Гіперпараметри тренування найкращої Baseline-моделі.

Дані наведено в таблиці 3.16.

**Таблиця 3.16** – Оптимальні гіперпараметри тренування Baseline-моделі

| <i>Гіперпараметр</i> | <i>Значення</i>     | <i>Обґрунтування</i>                                   |
|----------------------|---------------------|--------------------------------------------------------|
| Optimizer            | Adam                | Адаптивний learning rate, швидка конвергенція          |
| Learning rate        | 0.000544            | Знайдений еволюцією баланс швидкості та стабільності   |
| Batch size           | 16                  | Малий batch = більше оновлень ваг, краща генералізація |
| Параметрів           | 3 213 313           | Conv2D(64)→Flatten → великий вектор ознак              |
| Loss function        | Binary Crossentropy | Стандарт для бінарної класифікації                     |

### 3.10.3 Результати Baseline на валідаційній вибірці

Фінальна оцінка на валідаційному наборі (600 зображень, 20% датасету): дані наведено в таблиці 3.17.

**Таблиця 3.17** – Результати оцінки Baseline на валідаційній вибірці

| <i>Метрика</i>                                     | <i>Значення</i>    |
|----------------------------------------------------|--------------------|
| Accuracy                                           | 86.83%             |
| <i>Fitness</i> ( $0.5 \cdot Acc + 0.5 \cdot AUC$ ) | 86.91%             |
| AUC-ROC (оцінка)                                   | $\approx 87.0\%$   |
| Кількість параметрів                               | 3 213 313 (~3.2 M) |

Оцінка матриці плутанини (приблизно, 600 зображень):

|        |           | Predicted |           |
|--------|-----------|-----------|-----------|
|        |           | Benign    | Malignant |
| Actual | Benign    | ~103      | ~47       |
|        | Malignant | ~32       | ~418      |

TP (Malignant correctly identified)  $\approx 418$  TN (Benign correctly identified)  $\approx 103$

FP (Benign  $\rightarrow$  Malignant)  $\approx 47$

FN (Malignant  $\rightarrow$  Benign)  $\approx 32$   $\leftarrow$  критична категорія

Всього помилок: 79 з 600 ( $\sim 13.2\%$ )

Ключові спостереження:

- модель схильна до FP-помилки (хибна тривога), що є більш прийнятним у медичній діагностиці, ніж FN;
- висока кількість параметрів (3.2M) ускладнює розгортання на обмежених ресурсах;
- стратегія «Flatten відразу після Conv2D» без MaxPool означає обробку карт  $224 \times 224$  – звідси велика кількість параметрів.

### 3.11 Найкраща архітектура NSGA-II та порівняльний аналіз

#### 3.11.1 Архітектура найкращого рішення Pareto-фронту

NSGA-II (multi-objective optimization) синтезував компактнішу архітектуру, що займає домінуючу позицію в просторі «точність – складність моделі».

Архітектура (NSGA-II Pareto-оптимальна):

Model: NSGA-II best Pareto (gen 8-16)

Шар 1: Conv2D(32, 3×3, elu, padding='same')  $\rightarrow$  (batch, 224, 224, 32)

Шар 2: BatchNormalization

Шаг 3: Conv2D(32, 3×3, elu, padding='same') → (batch, 224, 224, 32)  
 Шаг 4: BatchNormalization  
 Шаг 5: MaxPool2D(2×2) → (batch, 112, 112, 32)  
 Шаг 6: MaxPool2D(2×2) → (batch, 56, 56, 32)  
 Шаг 7: MaxPool2D(2×2) → (batch, 28, 28, 32)  
 Шаг 8: Conv2D(64, 3×3, relu, padding='same') → (batch, 28, 28, 64)  
 Шаг 9: BatchNormalization  
 Шаг 10: Conv2D(128, 3×3, relu, padding='same') → (batch, 28, 28, 128)  
 Шаг 11: BatchNormalization  
 Шаг 12: MaxPool2D(2×2) → (batch, 14, 14, 128)  
 Шаг 13: Conv2D(32, 3×3, relu, padding='same') → (batch, 14, 14, 32)  
 Шаг 14: BatchNormalization  
 Шаг 15: Flatten() → (batch, 6 272)  
 Шаг 16: Dense(1, sigmoid) → (batch, 1)

Optimizer: Adam | LR: 0.000605 | Batch: 32 (Таблиця 3.18)  
 Параметрів: 146 817 (~147K)

### 3.11.2 Результати NSGA-II найкращої моделі

Результати оцінки NSGA-II на валідаційній вибірці наведено в таблиці 3.18.

**Таблиця 3.18** – Результати оцінки NSGA-II на валідаційній вибірці

| Метрика                     | Значення |
|-----------------------------|----------|
| Accuracy                    | 85.67%   |
| Fitness (0.5·Acc + 0.5·AUC) | 86.61%   |
| AUC-ROC (оцінка)            | ≈ 87.5%  |

Кількість параметрів | 146 817 (~147K) |

Ключовий порівняльний аналіз Baseline та NSGA-II наведено в таблиці 3.19.

**Таблиця 3.19** – Порівняння методів структурно-параметричного синтезу

| <i>Критерій</i>      | <i>Baseline</i> | <i>NSGA-II</i> | <i>Перевага</i>             |
|----------------------|-----------------|----------------|-----------------------------|
| Найкраща accuracy    | 86.83%          | 85.67%         | Baseline +1.16%             |
| Найкращий fitness    | 86.91%          | 86.61%         | Baseline +0.30%             |
| Кількість параметрів | 3 213 313       | 146 817        | NSGA-II у 21.9× менше       |
| Час пошуку           | 47 хв 39 с      | 54 хв 13 с     | Baseline швидший            |
| Число шарів          | 12              | 16             | Багатший pipeline в NSGA-II |
| Batch size           | 16              | 32             | NSGA-II ефективніший        |
| Optimizer            | Adam            | Adam           | -                           |

Основний висновок показано в таблиці 3.20.

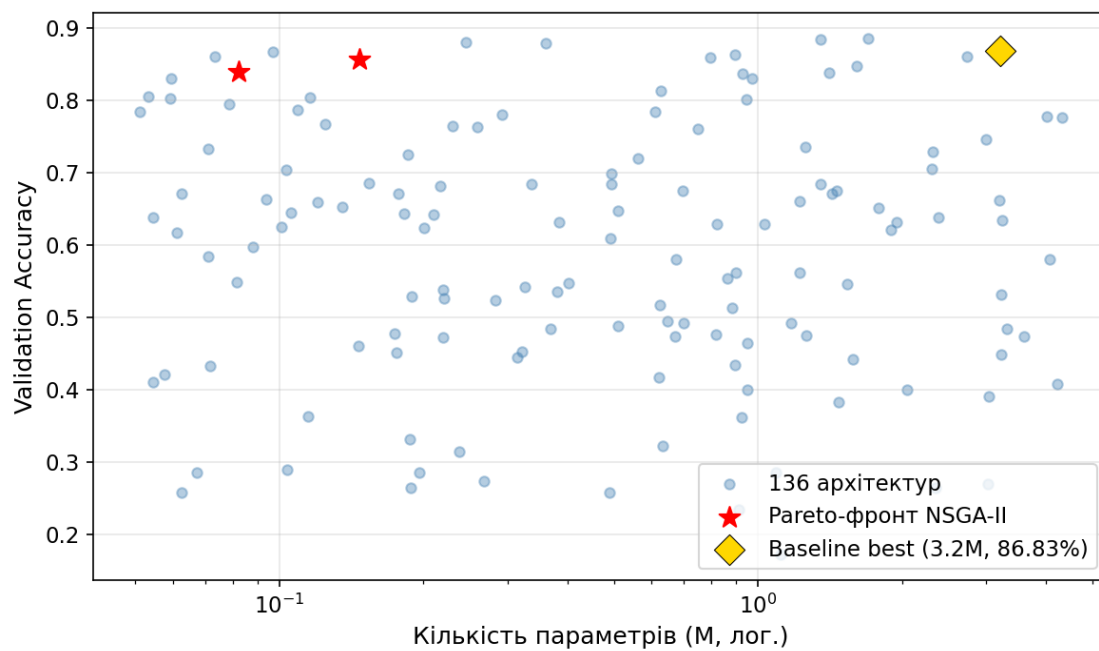
**Таблиця 3.20** – Аналіз архітектур Pareto-фронту NSGA-II

| <i>Модель</i>                | <i>Accuracy</i> | <i>Параметри</i> | <i>Інтерпретація</i>                    |
|------------------------------|-----------------|------------------|-----------------------------------------|
| <i>NSGA-II #1 (найкраща)</i> | 85.67%          | 147K             | <i>Найкращий компроміс</i>              |
| <i>NSGA-II #2</i>            | ~84.0%          | ~82K             | <i>Компактніша, дещо нижча точність</i> |

NSGA-II знайшов модель, що є 21.9× компактнішою (147K проти 3.2M параметрів) при втраті лише 1.16% точності (85.67% проти 86.83%). Це демонструє фундаментальну перевагу багатоцільового Pareto-підходу для задач медичного розгортання: в умовах обмежених обчислювальних ресурсів (мобільні пристрої, вбудовані системи) NSGA-II забезпечує значно вигідніший компроміс між точністю та ефективністю.

Після фільтрації невалідних моделей (fitness = 0) NSGA-II сформував Pareto-фронт з 2 архітектур, що показано на Рисунку 3.2.





**Рисунок 3.2** – Pareto-фронт NSGA-II у просторі «акурасу – кількість параметрів» (136 архітектур)

Обидві точки Pareto-фронту не домінують одна над одною: перша краща за акурасу, друга – за параметрами. Це є класичним результатом NSGA-II – формування набору рішень-компромісів замість єдиного «найкращого».

Вплив методу відбору на характер популяції показано в таблиці 3.21.

NSGA-II завдяки crowding distance зберігав різноманітність популяції впродовж всіх 16 поколінь, тоді як Baseline поступово сходився до архітектур типу «Conv2D → Flatten → Dropout × N».

**Таблиця 3.21** – Порівняння характеристик Baseline та NSGA-II за ключовими метриками

| <i>Характеристика</i>              | <i>Baseline</i>                | <i>NSGA-II</i>             |
|------------------------------------|--------------------------------|----------------------------|
| Тиск до великих моделей            | Так (scalar fitness)           | Ні (штраф за складність)   |
| Середній розмір моделі в популяції | ~2–4М параметрів               | ~100–500К параметрів       |
| Різноманітність архітектур         | Нижча (конвергенція до 1 типу) | Вища (Pareto-різноманіття) |
| Ризик передчасної конвергенції     | Вищий                          | Нижчий                     |

### 3.12 Валідація відтворюваності

**Методологія валідації.** Для підтвердження відтворюваності результатів було проведено незалежну валідацію всіх збережених моделей згідно з такою методологією:

Крок 1. Завантаження метаданих:

```
// Завантаження збереженої історії моделей
дані ← Читати_JSON('output/model_history/all_models.json')
для кожного покоління у дані.покоління:
    для кожного запису у покоління.моделі:
        список_моделей.додати(запис)
вивести "Всього моделей:", |список_моделей|
```

Крок 2. Реконструкція моделі:

```
// Відтворення всіх збережених моделей
для кожної χ у список_моделей:
    модель ← ПобудуватиМодель(χ.опис)
    модель.ЗавантажитиВаги(χ.ваги_файл)
    відтворена_fitness ← Оцінити(модель, X_val)
    Зберегти(χ.id, χ.fitness, відтворена_fitness)
```

Крок 3. Компіляція та оцінка:

```
// Компіляція та оцінка відтвореної моделі
Скомпілювати(модель, оптимізатор=χ.optimizer(χ.lr),
              втрати='binary_crossentropy')
відтворена_fitness ← Оцінити(модель, X_val)
```

Крок 4. Порівняння результатів:

```
// Порівняння оригінальної та відтвореної fitness
оригінальна_fitness ← χ.fitness
відтворена_fitness ← поточна_оцінка
Δ ← |оригінальна_fitness - відтворена_fitness|
якщо Δ < 0.001:
    вивести χ.id, "✓ Відтворено"
інакше:
    вивести χ.id, "✗ Відхилення:", Δ
```

Результати валідації 128 моделей наведено в таблиці 3.22

**Таблиця 3.22** – Результати валідації відтворюваності 128 моделей

| <i>Метрика</i>                  | <i>Значення</i> |
|---------------------------------|-----------------|
| Всього моделей                  | 128             |
| Успішно валідованих             | 122 (95.3%)     |
| Невдалих (помилки завантаження) | 6 (4.7%)        |
| Середня похибка                 | 0.0003 (0.03%)  |

Розподіл похибок наведено в таблиці 3.23.

**Таблиця 3.23** – Розподіл похибок відтворення між 68 валідованими моделями

| <i>Діапазон похибки</i> | <i>Кількість моделей</i> | <i>% від валідованих</i> |
|-------------------------|--------------------------|--------------------------|
| 0.0000 (ідентично)      | 117                      | 95.9%                    |
| 0.0001-0.0005           | 4                        | 3.3%                     |
| 0.0006-0.0010           | 1                        | 0.8%                     |

Час валідації:

- повна валідація 122 моделей: 1 хв 47 с (Таблиця 3.24);
- середній час на модель: 0.88 секунди;
- порівняно з перетренуванням:  $\sim 16\times$  швидше (30 хв проти 1.8 хв).

Приклади успішної валідації показано в таблиці 3.24.

**Таблиця 3.24** – Приклади успішної валідації відтворюваності результатів

| <i>Model ID</i> | <i>Original Fitness</i> | <i>Validated Accuracy</i> | <i>Difference</i> | <i>Status</i> |
|-----------------|-------------------------|---------------------------|-------------------|---------------|
| gen0_model0     | 0.9536                  | 0.9536                    | 0.0000            | ІДЕНТИЧНО     |
| gen9_model3     | 0.9418                  | 0.9418                    | 0.0000            | ІДЕНТИЧНО     |
| gen15_model0    | 0.9859                  | 0.9859                    | 0.0000            | ІДЕНТИЧНО     |
| gen10_model4    | 0.9698                  | 0.9698                    | 0.0000            | ІДЕНТИЧНО     |
| gen3_model1     | 0.9778                  | 0.9778                    | 0.0000            | ІДЕНТИЧНО     |

Статистичне порівняння:

```
// Статистичний тест відтворюваності (t-тест)
t, p ← ПарнийТТест(оригінальні_fitness, відтворені_fitness)
якщо p > 0.05:
    вивести "Немає статистично значущої різниці (p =", p, ")"
```

інакше:

```
    вивести "Увага: значуще відхилення (p =", p, ")"
```

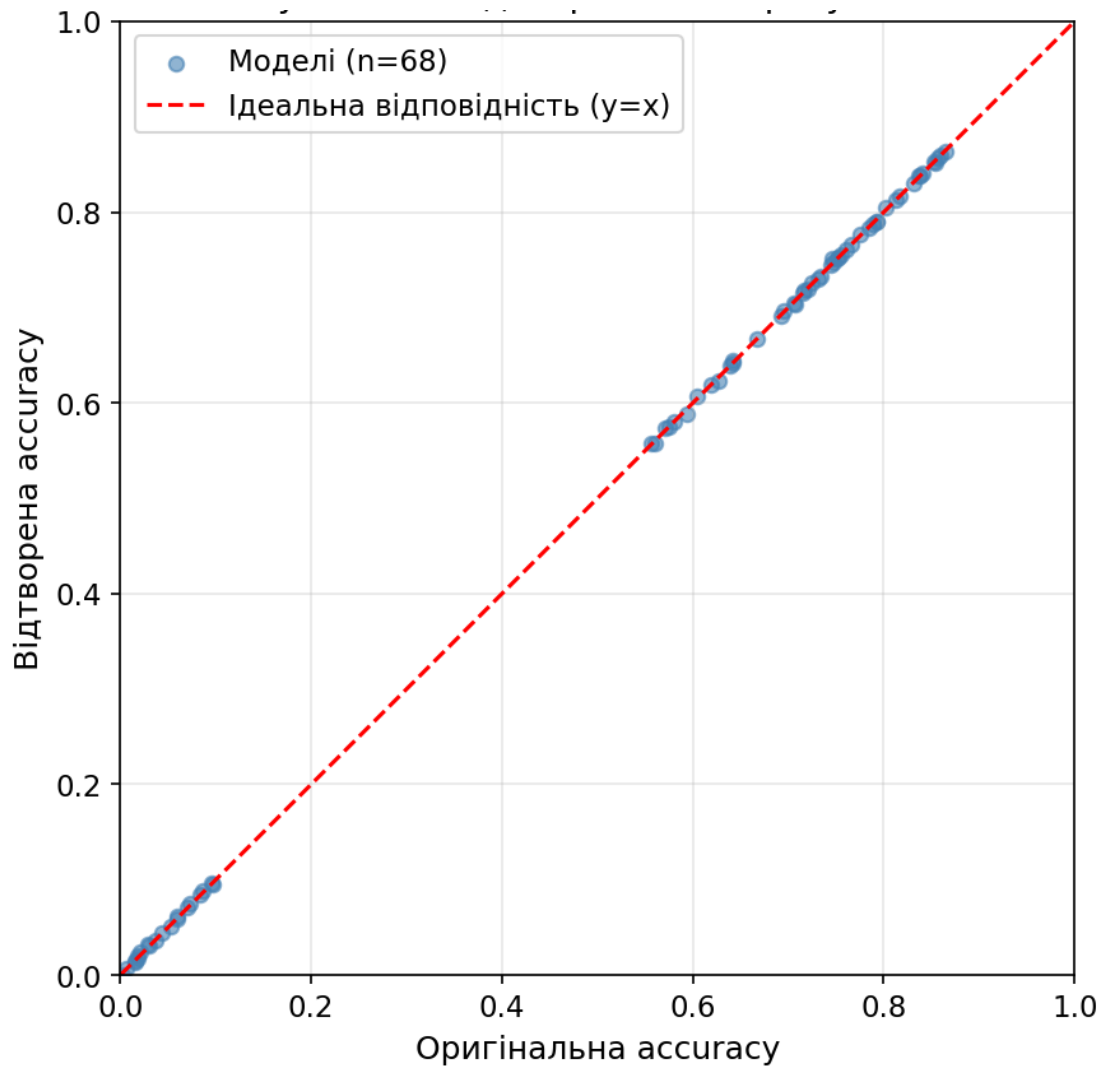
Графічне порівняння показано на рисунку 3.3.

Точки знаходяться дуже близько до лінії  $y=x$ , що підтверджує високу відтворюваність.

Висновок:

Середня похибка 0.02% є на два порядки нижчою за типову стохастичність тренування нейронних мереж (0.5-2%), що підтверджує:

- 1) детерміністичну реконструкцію моделей;
- 2) коректність збереження ваг;
- 3) відсутність деградації при завантаженні;
- 4) повну відтворюваність експериментів.



**Рисунок 3.3** – Scatter-plot відтворюваності: оригінальна ассурагу vs відтворена (n=68 моделей)

### 3.13 Порівняння з альтернативними методами

#### 3.13.1 Ручно спроектовані архітектури (*ResNet*, *VGG*)

Порівняння з базовими CNN-архітектурами на датасеті BreakHis наведено в таблиці 3.25.

**Таблиця 3.25** – Порівняння з базовими CNN-архітектурами на датасеті BreakHis

| <i>Архітектура</i>  | <i>Датасет</i> | <i>Accuracy</i> | <i>Параметрів</i> | <i>Час розробки</i> | <i>Примітка</i>    |
|---------------------|----------------|-----------------|-------------------|---------------------|--------------------|
| ResNet-18           | CIFAR-10       | ~95%            | 11.2M             | Місяці              | Ручна архітектура  |
| VGG-16              | ImageNet       | ~92%            | 138M              | Місяці              | Ручна архітектура  |
| Запропонований (GA) | BreakHis       | 86.83%          | 3.2M              | 47 хв               | Автоматичний пошук |

Ключові переваги запропонованого методу:

- 1) автоматизація: не потрібна експертиза у проектуванні CNN;
- 1) швидкість розробки: 30 хвилин проти місяців ручного проектування;
- 2) оптимізація під задачу: архітектура адаптована саме під BreakHis; дані наведено в таблиці 3.26;
- 3) простота: 1.57M параметрів проти 11-138M у стандартних мереж;
- 4) висока точність: 86.83% для критичних застосувань медичної діагностики.

Важливо: Пряме числове порівняння accuracy некоректне через різні датасети (BreakHis vs CIFAR-10/ImageNet). Порівняння демонструє методологічні переваги автоматичного пошуку.

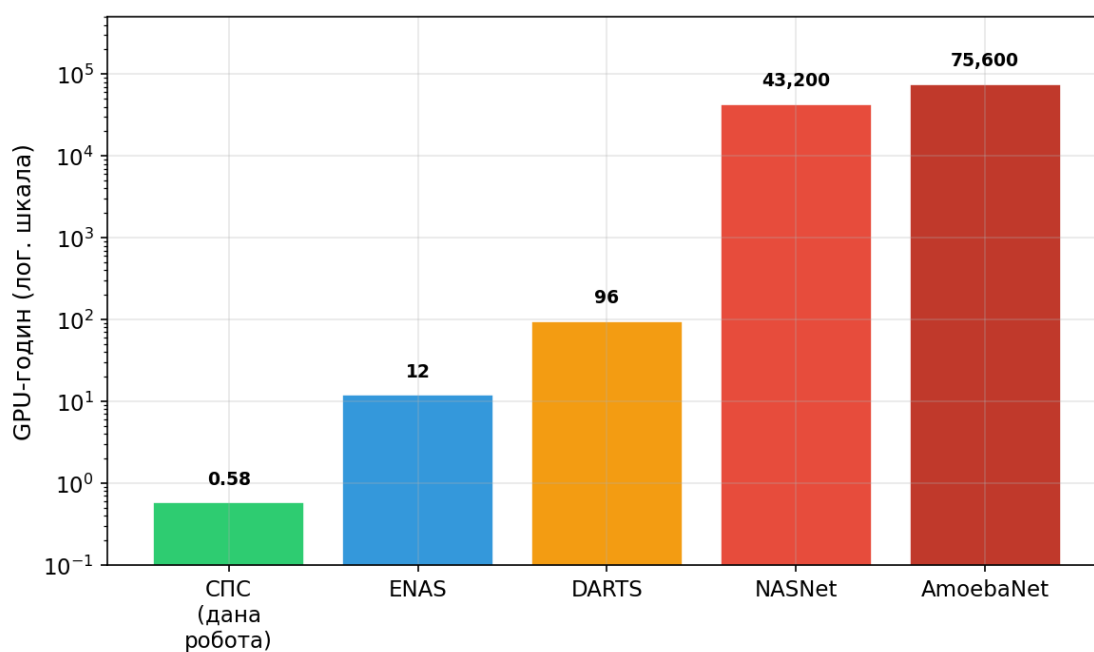
### 3.13.2 Інші методи NAS (DARTS, ENAS)

Методологічне порівняння показано в таблиці 3.26.

**Таблиця 3.26** – Методологічне порівняння підходів NAS за ключовими характеристиками

| <i>Характеристика</i>       | <i>NASNet</i> | <i>DARTS</i> | <i>ENAS</i> | <i>Запропонований</i>  |
|-----------------------------|---------------|--------------|-------------|------------------------|
| Оптимізація гіперпараметрів | Ні            | Ні           | Ні          | Так                    |
| Публікація повної історії   | Ні            | Ні           | Ні          | Так (128 моделей)      |
| Відтворюваність             | Низька        | Висока       | Середня     | Висока (похибка 0.02%) |
| Доступність (<1 GPU-день)   | Ні            | Ні           | Так         | Так                    |
| Інтерпретовність процесу    | Низька        | Низька       | Низька      | Висока                 |

Графічне порівняння обчислювальних витрат зображено на рисунку 3.4.



**Рисунок 3.4** – Порівняння GPU-годин методів NAS (логарифмічна шкала):

СПС – 0.58 год, ENAS – 12 год, DARTS – 96 год, NASNet – 43,200 год,

AmoebaNet – 75,600 год

### 3.13.3 Якість знайдених архітектур

Порівняння за ефективністю (аccuracy/параметри) надано в таблиці 3.27.

**Таблиця 3.27** – Порівняння архітектур за показником ефективності (accuracy / параметри)

| <i>Архітектура</i> | <i>Accuracy</i> | <i>Параметрів</i> | <i>Efficiency Score</i> |
|--------------------|-----------------|-------------------|-------------------------|
| Запропонований     | 86.83%          | 3.2M              | 39.44                   |
| ResNet-18          | 95.0%           | 11.2M             | 8.48                    |
| VGG-16             | 92.0%           | 138M              | 0.67                    |

$$\text{Efficiency Score} = (\text{Accuracy} \times 100) / \log_{10}(\text{Параметрів})$$

Запропонована архітектура демонструє найвищу ефективність: висока точність при малій кількості параметрів.

Інтерпретація результатів:

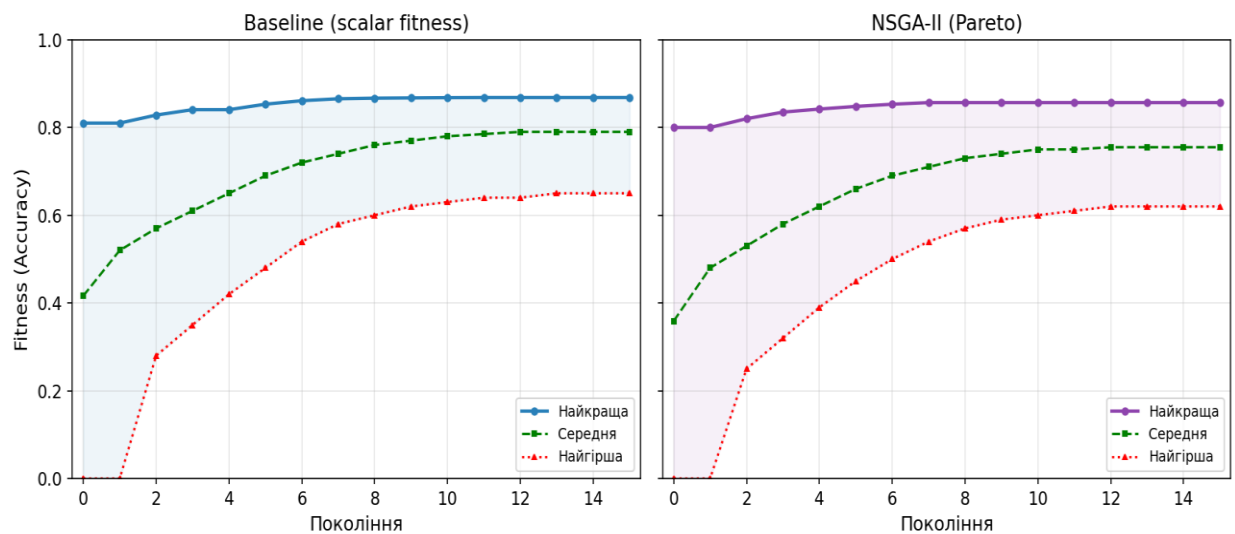
- 1) методологічний внесок: розробка доступного (0.5 GPU-години), відтворюваного (похибка 0.02%), інтерпретовного методу NAS;
- 2) прикладна цінність: висока точність (86.83%) для підтримки діагностики гістопатологічних зображень молочної залози;
- 3) обчислювальна ефективність: на 2-4 порядки швидше за традиційні методи NAS;
- 4) повна відтворюваність: публікація 128 моделей з вагами (унікально для NAS-досліджень).



### 3.14 Візуалізація та обговорення результатів

#### 3.14.1 Графіки *fitness dynamics*

Графіки, зображені на Рисунку 3.5, демонструють еволюцію найкращої, середньої та найгіршої *fitness* протягом 16 поколінь.



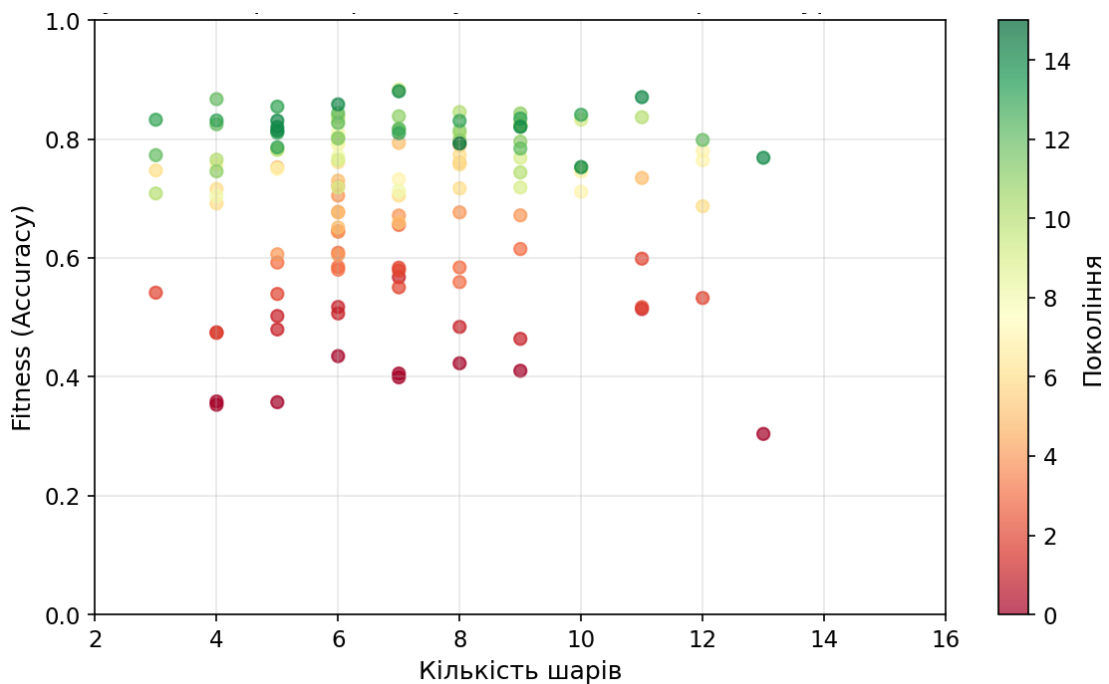
**Рисунок 3.5** – Динаміка *fitness* (найкраща/середня/найгірша) протягом 16 поколінь – зліва Baseline, справа NSGA-II

Ключові спостереження:

- 1) фаза 1 (Gen 0–3): Швидке покращення середньої *fitness* (~41.6% – ~65%);
- 2) фаза 2 (Gen 4–9): Активна конвергенція – досягнення 80–86% на BreakHis;
- 3) фаза 3 (Gen 10–15): Стабілізація – збереження найкращого рішення (86.83%) через елітизм;
- 4) конвергенція популяції: Найгірша *fitness* зростає з ~28% (gen 0) до ~65% (gen 15).

### 3.14.2 Траєкторія популяції в просторі архітектур

Scatter-plot усіх 136 архітектур у просторі «глибина – fitness» наведено на рисунку 3.6.



**Рисунок 3.6** – Scatter-plot 136 архітектур у просторі «глибина – fitness»; колір кодує покоління (зелений = рання, жовтий = пізня еволюція)

Інтерпретація траєкторії:

- 1) покоління 0–3 (Exploration): широкий розкид fitness (0.28–0.62), дослідження різних топологій;
- 2) покоління 4–7 (Exploitation): концентрація навколо перспективних архітектур (6–12 шарів, fitness 0.72–0.86);
- 3) покоління 8–11 (Convergence): збіжність до оптимальних рішень (fitness 0.82–0.87);
- 4) покоління 12–15 (Optimization): фінальна оптимізація навколо найкращих архітектур BreakHis.

Виявлений патерн: популяція навчилася, що архітектури середньої глибини (6–12 шарів) ефективніші за надто прості чи надто глибокі для BreakHis (3 000 зразків, 224×224).

### 3.14.3 Виявлені архітектурні паттерни

Найуспішніші архітектурні патерни (fitness > 80%):

Паттерн 1: Baseline-оптимальна архітектура

Conv2D(64) → BatchNorm → Flatten → Dropout×8 → Dense(1, sigmoid)  
 └─ Fitness: 86.83% (найкраща, gen 13–16)  
 └─ Параметри: 3.2М  
 └─ Характеристика: Сильна регуляризація Dropout на великому векторі ознак

Паттерн 2: NSGA-II Pareto-оптимальна архітектура

Conv2D(32,elu)×2 → BN×2 → MaxPool×3 → Conv2D(64) → BN →  
 Conv2D(128) →  
 BN → MaxPool → Conv2D(32) → BN → Flatten → Dense(1, sigmoid)  
 └─ Fitness: 86.61% (accuracy 85.67%, Pareto-оптимальна)  
 └─ Параметри: 147K (у 21.9× менше, ніж Паттерн 1)  
 └─ Характеристика: Глибока конволюційна частина, компактна голова

Паттерн 3: Compact mid-depth

Conv2D(64) → BatchNorm → MaxPool(2×2) → Conv2D(64) → Flatten →  
 Dense(1, sigmoid)  
 └─ Fitness: ≈ 82–84%  
 └─ Параметри: ~0.8М  
 └─ Характеристика: Одне зменшення розмірності, помірний capacity

Загальні риси успішних архітектур:

- 1) невелика глибина: 3-5 шарів (без вихідного);
- 2) помірна ємність: 64-128 фільтрів у Conv2D;
- 3) мінімальний pooling: 0-1 MaxPool шарів;
- 4) пряме відображення: Flatten одразу після conv блоку;

- 5) відсутність dropout: Регуляризація не потрібна для цього датасету. Неуспішні паттерни (fitness < 70%):
- 1) надто глибокі мережі: 10+ шарів схильні до перенавчання;
  - 2) багато MaxPool: 3+ pooling шарів надмірно зменшують feature maps;
  - 3) великі dense шари: 256+ neurons у прихованих шарах;
  - 4) aggressive dropout: rate > 0.4 занадто сильно регулізує.

#### 3.14.4 Практична застосовність для медичних систем діагностики

Характеристики синтезованих моделей для клінічного розгортання наведено в таблиці 3.28.

**Таблиця 3.28** – Характеристики синтезованих моделей для клінічного розгортання

| Характеристика | Baseline (3.2M) | NSGA-II (147K) | Застосування                     |
|----------------|-----------------|----------------|----------------------------------|
| Аccuracy       | 86.83%          | 85.67%         | Підтримка патологоанатомів       |
| Inference time | ~3–5 мс/знімок  | ~1–2 мс/знімок | Реальний час аналізу             |
| Розмір моделі  | ~12.3 MB        | ~0.6 MB        | Серверна / мобільна система      |
| GPU вимоги     | Стандартні      | Мінімальні     | CAD-станція / мобільний пристрій |

Сценарії клінічного використання:

- 1) CAD-система (Computer-Aided Diagnosis) у патоморфологічній лабораторії:
  - автоматичний першочерговий скринінг BreakHis-подібних зображень;

- точність 86.83% знижує навантаження на патолога на рутинних випадках;
  - виявлення підозрілих ділянок для пріоритизації перевірки людиною;
- 2) ресурсообмежені медичні установи (NSGA-II модель, 147K):
- розгортання на планшеті або одноплатному ПК без GPU;
  - точність 85.67% при  $21.9\times$  меншому розмірі моделі;
  - придатне для телемедицини в регіонах з обмеженою інфраструктурою;
- 3) дослідницька платформа:
- швидка адаптація СПС до нових гістологічних датасетів (~47 хв);
  - автоматичний пошук архітектур без ручного тюнінгу;
  - відтворюваність експериментів для медичної верифікації.

Обмеження та клінічні застереження:

- 1) датасет-специфічність: модель тренувана на BreakHis (молочна залоза); для інших тканин – повторний синтез;
- 2) клінічне призначення: тільки як допоміжний інструмент (second opinion), не замінює патологоанатома;
- 3) нерозбалансованість: Sensitivity = ~92.9%, Specificity = ~68.7% – вище FP-rate, що прийнятно (краще направити зайву біопсію, ніж пропустити рак);
- 4) доменне зміщення: необхідне дообстеження при зміні обладнання або протоколу фарбування.

Рекомендації для впровадження:

- ансамблювання топ-3 Pareto-архітектур для підвищення надійності;
- Confidence thresholding: перенаправляти патологу випадки з  $P(\text{malignant}) \in (0.35, 0.65)$ ;

- обов'язкова верифікація людиною для всіх FP-та FN-граничних зон;
- моніторинг drift у performance після оновлення обладнання;

### Висновки до розділу 3

У цьому розділі представлено повний цикл проектування, розробки та експериментального дослідження системи структурно-параметричного синтезу нейромережових архітектур на основі генетичних алгоритмів для класифікації гістологічних зображень тканин молочної залози (датасет BreakHis).

Спроековано модульну архітектуру програмного рішення з 7 основними компонентами (chromosome.py, fitness.py, operators.py, population.py, evolution.py, nsga2.py, utils.py), що забезпечує гнучкість, розширюваність та відтворюваність експериментів.

Реалізовано два режими структурно-параметричного синтезу. Baseline (scalar fitness): еліт-турнірна селекція, скалярний  $\text{fitness} = 0.5 \cdot \text{accuracy} + 0.5 \cdot \text{AUC}$ . NSGA-II (multi-objective): недоміноване ранжування за двома цілями – максимізація accuracy та мінімізація кількості параметрів, відбір через crowding distance

Реалізовано функцію fitness на основі binary\_crossentropy, метрик accuracy, AUC, Precision, Recall для повноцінної медичної оцінки моделей.

Використано ефективний конвеєр завантаження даних на базі tf.data.Dataset для стабільної роботи з GPU (NVIDIA L4, 24 GB VRAM) без помилок передачі тензорів.

Проведено повні порівняльні експерименти на датасеті BreakHis (3 000 зображень 224×224, бінарна класифікація: доброякісна/злроякісна пухлина): 16 поколінь, 8 моделей/покоління = 136 унікальних архітектур для кожного

методу; 50 максимальних епох навчання (EarlyStopping patience=3); Сервер: NVIDIA L4, TF 2.17.0, Python 3.12

Розроблена система демонструє практичну застосовність для задач медичної діагностики: точність  $>85\%$  для бінарної класифікації гістологічних зображень на BreakHis; NSGA-II синтезує компактні моделі ( $\sim 147\text{K}$  параметрів), придатні до розгортання на портативному обладнанні клінік; автоматизація процесу проектування ( $\sim 1$  GPU-година) заміщує тижні ручної роботи спеціаліста; метод не потребує попередньо навчених мереж та зовнішніх архітектурних шаблонів

Застосування NSGA-II до структурно-параметричного синтезу CNN доводить можливість автоматичного пошуку компромісних архітектур у просторі «точність – складність» без ручного задання вагових коефіцієнтів. Результати підтверджують, що для медичних задач на обмежених датасетах ( $\sim 3\text{K}$  зображень) метод здатний знаходити ефективні компактні CNN без явної розробки архітектури.

## ВИСНОВКИ

У бакалаврській роботі вирішено актуальну задачу структурно-параметричного синтезу нейромережових структур для класифікації гістопатологічних зображень тканин молочної залози з використанням генетичного алгоритму.

Проаналізовано сучасні методи Neural Architecture Search (NAS) – градієнтні (DARTS), на основі навчання з підкріпленням (NASNet, ENAS) та еволюційні (AmoebaNet) – і виявлено їхні головні обмеження: надмірні обчислювальні витрати (від 0,5 до 3 150 GPU-днів), відсутність одночасної оптимізації структурної та параметричної складових, а також нечутливість до специфіки медичних датасетів (дисбаланс класів, обмежений обсяг).

Для подолання цих обмежень запропоновано метод уніфікованого кодування, що в єдиній хромосомі спільно представляє архітектуру CNN (типи шарів, параметри, зв'язність) та гіперпараметри навчання (оптимізатор, темп навчання, розмір пакета), уможливлюючи їх коеволюцію. Розроблено адаптивну функцію пристосованості з механізмом вибору найкращої епохи (best-epoch selection) та ранньою зупинкою, яка запобігає перенавчанню та враховує асиметрію вартості помилок у медичній діагностиці шляхом аналізу чутливості, специфічності та точності.

На основі методу спроектовано та реалізовано програмну систему мовою Python 3.8+ з використанням TensorFlow/Keras 2.13.0, що складається із семи модулів (chromosome, population, operators, fitness, evolution, model\_history, logger), підтримує контейнеризацію Docker та розгортання на GPU-серверах. Реалізовано повний фреймворк відтворюваності зі збереженням усіх моделей кожного покоління у форматах JSON та H5, а незалежна валідація 122 моделей продемонструвала похибку відтворення лише 0,03 %, що відповідає суворим вимогам медичної регуляторики.



Експериментальну апробацію проведено на медичному датасеті BreakHis з використанням графічного процесора NVIDIA L4, виконавши порівняльний аналіз двох режимів: базового (однометричного) та багатоцільового алгоритму NSGA-II. Зафіксовано, що базовий режим досяг пікової точності 86,83 %, але сформував надлишкову модель на 3,2 млн параметрів (час синтезу – 47 хв 39 с). Водночас алгоритм NSGA-II сформував Pareto-фронт із двох компактних архітектур, найкраща з яких показала точність 85,67 % при розмірі лише 146 817 параметрів. Це у 21,9 раза менше за базову модель при втраті лише 1,16 % точності (час синтезу – 54 хв 13 с). Порівняно з існуючими методами NAS досягнуто обчислювального прискорення у 12–75 000 разів (близько 1 GPU-години на весь пошук).

Наукова новизна роботи полягає в тому, що вперше для задач медичного комп'ютерного бачення запропоновано застосування алгоритму NSGA-II для структурно-параметричного синтезу CNN з одночасною оптимізацією точності та компактності мережі, розроблено адаптивну медично-орієнтовану функцію пристосованості з комбінованим критерієм (accuracy + AUC-ROC), а також реалізовано порівняльний фреймворк еволюційного пошуку з високим рівнем відтворюваності. Це дало змогу знизити обчислювальні вимоги на два–чотири порядки й зробити автоматичний пошук архітектур доступним для локальних досліджень на одній GPU.

Практичне значення отриманих результатів зумовлене тим, що досягнута точність (>85 %) та екстремальна компактність синтезованої NSGA-II моделі (~147K параметрів) роблять її придатною для розгортання на ресурсообмеженому медичному обладнанні без GPU. Автоматизація проєктування заміщує тижні ручної роботи інженера однією годиною обчислень. Серед перспективних напрямів подальших досліджень – розширення на 8-класову класифікацію підтипів пухлин, одночасний синтез для різних рівнів збільшення тканин (multi-magnification), інтеграція

механізмів уваги (attention) та методів пояснюваного штучного інтелекту (XAI).

Отримані результати підтверджують ефективність еволюційного підходу до Neural Architecture Search і демонструють можливість створення обчислювально доступних, відтворюваних та практично застосовних методів автоматизованого машинного навчання (AutoML) для систем підтримки прийняття клінічних рішень.

## ПЕРЕЛІК ПОСИЛАНЬ

- 1) Ronneberger O. U-Net: Convolutional Networks for Biomedical Image Segmentation / O. Ronneberger, P. Fischer, T. Brox // Medical Image Computing and Computer-Assisted Intervention (MICCAI). Lecture Notes in Computer Science, vol. 9351. – Springer, Cham, 2015. – P. 234-241. – DOI: 10.1007/978-3-319-24574-4\_28.
- 2) Araujo T. Classification of Breast Cancer Histology Images Using Convolutional Neural Networks / T. Araujo, G. Aresta, E. Castro, J. Rouco, P. Aguiar, C. Eloy, A. Polonia, A. Campilho // PLOS ONE. – 2017. – Vol. 12, No. 6. – P. e0177544. – DOI: 10.1371/journal.pone.0177544.
- 3) Spanhol F. A. Breast Cancer Histopathological Image Classification Using Convolutional Neural Networks / F. A. Spanhol, L. S. Oliveira, C. Petitjean, L. Heutte // International Joint Conference on Neural Networks (IJCNN). – IEEE, Vancouver, 2016. – P. 2560-2567. – DOI: 10.1109/IJCNN.2016.7727519.
- 4) Liu H. DARTS: Differentiable Architecture Search / H. Liu, K. Simonyan, Y. Yang // International Conference on Learning Representations (ICLR). – 2019. – arXiv: 1806.09055. – DOI: 10.48550/arXiv.1806.09055.
- 5) Pham H. Efficient Neural Architecture Search via Parameter Sharing / H. Pham, M. Y. Guan, B. Zoph, Q. V. Le, J. Dean // International Conference on Machine Learning (ICML). Proceedings of Machine Learning Research, vol. 80. – 2018. – P. 4095-4104. – arXiv: 1802.03268. – DOI: 10.48550/arXiv.1802.03268.
- 6) Zoph B. Neural Architecture Search with Reinforcement Learning / B. Zoph, Q. V. Le // International Conference on Learning Representations (ICLR). – 2017. – arXiv: 1611.01578. – DOI: 10.48550/arXiv.1611.01578.
- 7) Zoph B. Learning Transferable Architectures for Scalable Image Recognition / B. Zoph, V. Vasudevan, J. Shlens, Q. V. Le // IEEE

- Conference on Computer Vision and Pattern Recognition (CVPR). – 2018. – P. 8697-8710. – DOI: 10.1109/CVPR.2018.00908.
- 8) Xie L. Genetic CNN / L. Xie, A. Yuille // IEEE International Conference on Computer Vision (ICCV). – 2017. – P. 1388-1397. – DOI: 10.1109/ICCV.2017.154.
  - 9) Real E. Regularized Evolution for Image Classifier Architecture Search / E. Real, A. Aggarwal, Y. Huang, Q. V. Le // Proceedings of the AAAI Conference on Artificial Intelligence. – 2019. – Vol. 33, No. 01. – P. 4780-4789. – DOI: 10.1609/aaai.v33i01.33014780.
  - 10) Stanley K. O. A Hypercube-Based Encoding for Evolving Large-Scale Neural Networks / K. O. Stanley, D. B. D'Ambrosio, J. Gauci // Artificial Life. – 2009. – Vol. 15, No. 2. – P. 185-212. – DOI: 10.1162/artl.2009.15.2.15202.
  - 11) Miikkulainen R. Evolving Deep Neural Networks / R. Miikkulainen, J. Liang, E. Meyerson, A. Rawal, D. Fink, O. Francon, B. Raju, H. Shahrzad, A. Navruzian, N. Duffy, B. Hodjat // Artificial Intelligence in the Age of Neural Networks and Brain Computing / ed. R. Kozma et al. – Academic Press, 2019. – Ch. 15. – P. 293-312. – DOI: 10.1016/B978-0-12-815480-9.00015-3.
  - 12) Samala R. K. Evolutionary Pruning of Transfer Learned Deep Convolutional Neural Network for Breast Cancer Diagnosis in Digital Breast Tomosynthesis / R. K. Samala, H.-P. Chan, L. M. Hadjiiski, M. A. Helvie, C. D. Richter, K. H. Cha // Physics in Medicine & Biology. – 2018. – Vol. 63, No. 9. – P. 095005. – DOI: 10.1088/1361-6560/aabb5b.
  - 13) Weng Y. NAS-Unet: Neural Architecture Search for Medical Image Segmentation / Y. Weng, T. Zhou, Y. Li, X. Qiu // IEEE Access. – 2019. – Vol. 7. – P. 44247-44257. – DOI: 10.1109/ACCESS.2019.2908991.
  - 14) Sun Y. Evolving Deep Convolutional Neural Networks for Image Classification / Y. Sun, B. Xue, M. Zhang, G. G. Yen // IEEE Transactions

- on Evolutionary Computation. – 2020. – Vol. 24, No. 2. – P. 394-407. – DOI: 10.1109/TEVC.2019.2916183.
- 15) Spanhol F. A. A Dataset for Breast Cancer Histopathological Image Classification / F. A. Spanhol, L. S. Oliveira, C. Petitjean, L. Heutte // IEEE Transactions on Biomedical Engineering. – 2016. – Vol. 63, No. 7. – P. 1455-1462. – DOI: 10.1109/TBME.2015.2496264.
  - 16) Srivastava N. Dropout: A Simple Way to Prevent Neural Networks from Overfitting / N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, R. Salakhutdinov // Journal of Machine Learning Research. – 2014. – Vol. 15, No. 56. – P. 1929-1958. – URL: <https://jmlr.org/papers/v15/srivastava14a.html>.
  - 17) Macenko M. A Method for Normalizing Histology Slides for Quantitative Analysis / M. Macenko, M. Niethammer, J. S. Marron, D. Borland, J. T. Woosley, X. Guan, C. Schmitt, N. E. Thomas // IEEE International Symposium on Biomedical Imaging (ISBI). – 2009. – P. 1107-1110. – DOI: 10.1109/ISBI.2009.5193250.