

**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

«__» _____ 20__ р.

Дипломна робота

на здобуття ступеня бакалавра

за освітньо-професійною програмою «Системи і методи штучного інтелекту»

спеціальності 122 «Комп'ютерні науки»

на тему: «Інтелектуальна веб-система автоматизованої перевірки коду»

Виконав:

студент IV курсу, групи КІ-21

Шульга Олексій Олексійович _____

Керівник:

доцент кафедри ШІ, к.т.н., доцент,

Тимошенко Юрій Олександрович _____

Консультант з нормоконтролю:

фахівець першої категорії кафедри штучного інтелекту, к.т.н., доцент,

Комариста Богдана Миколаївна _____

Рецензент:

Доцент кафедри Інформаційної безпеки, к.т.н., доцент

Гальчинський Леонід Юрійович _____

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

Студент _____

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

«15» січня 2026 р.

ЗАВДАННЯ

на дипломну роботу студенту

Шульзі Олексію Олексійовичу

1. Тема роботи «Інтелектуальна веб-система автоматизованої перевірки коду», керівник роботи Тимошенко Юрій Олександрович, доцент, кандидат технічних наук, 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: технічна документація та специфікації Google Gemini API, стандарти обробки синтаксичних дерев (AST) мови Python, офіційні специфікації FastAPI та React.js, тестові набори вихідного тексту студентських програм.

4. Зміст роботи: аналіз сучасного стану та проблем автоматизації перевірки коду, обґрунтування технологічного стеку та клієнт-серверної архітектури, розробка алгоритмів AST-нормалізації та дворангового семантичного кешування запитів, програмна реалізація веб-системи,

експериментальний аналіз часової та економічної ефективності гібридного конвеєра.

5. Перелік ілюстративного матеріалу: презентація до захисту

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

Календарний план

№ з/п	Назва етапів виконання дипломної роботи	Термін виконання етапів роботи	Примітка
1	Аналіз предметної області та існуючих рішень.	20.04.2026	Виконано
2	Проектування архітектури веб-системи.	27.04.2026	Виконано
3	Розробка алгоритмів AST-аналізу та кешування.	04.05.2026	Виконано
4	Програмна реалізація серверної частини.	11.05.2026	Виконано
5	Програмна реалізація клієнтської частини.	18.05.2026	Виконано
6	Тестування системи та проведення експериментів.	25.05.2026	Виконано
7	Оформлення пояснювальної записки.	01.06.2026	Виконано
8	Підготовка презентації та матеріалів до захисту.	05.06.2026	Виконано

Студент

Олексій ШУЛЬГА

Керівник

Юрій ТИМОШЕНКО

РЕФЕРАТ

Дипломна робота: 85 с., 9 рис., 5 табл., 34 посилань, додаток.

АБСТРАКТНЕ СИНТАКСИЧНЕ ДЕРЕВО (AST), ВЕЛИКІ МОВНІ МОДЕЛІ (LLM), СЕМАНТИЧНЕ КЕШУВАННЯ, FASTAPI, REACT.JS, ВЕРИФІКАЦІЯ КОДУ, БЕНЧМАРКІНГ.

Пряме використання хмарних моделей ШІ для перевірки коду студентів викликає надмірну латентність та високу вартість хмарних токенів, що потребує архітектурної оптимізації.

Об'єкт дослідження – процес автоматизованої верифікації та технічного рецензування програмного коду студентів.

Предмет дослідження – методи аналізу коду на основі абстрактних синтаксичних дерев (AST) та технології багаторівневого семантичного кешування.

Мета роботи – проєктування та програмна реалізація веб-системи перевірки коду з гібридним конвеєром обробки для підвищення швидкодії та мінімізації витрат на хмарні API.

У роботі оглянуто платформи перевірки коду та проаналізовано проблему латентності ШІ. Запропоновано архітектуру трирівневого конвеєра: локальний аналіз, семантичний кеш (L2) та Gemini API. Розроблено метод шаблонізації AST-дерев для пошуку логічних клонів і розроблено комплекс на FastAPI та React.js. Виконано генерацію мутацій коду для тестування. Отримано метрики швидкодії конвеєра. Встановлено, що кешування зменшує кількість звернень до ШІ на 65% і знижує латентність обробки еквівалентного коду з 4,2 с до 0,04 с.

ABSTRACT

Bachelor's thesis: 85 p., 9 figures, 5 tables, 34 references, appendix.

ABSTRACT SYNTAX TREE (AST), LARGE LANGUAGE MODELS (LLM), SEMANTIC CACHING, FASTAPI, REACT.JS, CODE VERIFICATION, BENCHMARKING

Direct cloud LLM integration for student code verification causes excessive processing latency and high token costs, requiring advanced architectural optimization.

The object of the study is the process of automated verification and technical reviewing of student source code.

The subject of research is AST-based code analysis methods and multi-level semantic caching technologies.

The purpose of the work is to design and implement an intelligent web system for code verification using a hybrid analysis pipeline to improve performance and minimize cloud API costs.

Existing verification platforms were reviewed, and the high latency of cloud AI models was analyzed. An architecture of a hybrid three-tier pipeline (local AST, L2 cache, Gemini API) was proposed. An AST templating method for logical clone detection was developed. A complete client-server complex was developed using FastAPI and React.js. Synthetic data generation and automated benchmarking were executed. Performance metrics were obtained. It was established that multi-level caching reduces direct AI API calls by 65% and decreases the average processing latency of equivalent solutions from 4,2 s to 0,04 s

ЗМІСТ

ПЕРЕЛІК ПРИЙНЯТИХ СКОРОЧЕНЬ	8
ВСТУП.....	9
РОЗДІЛ 1 АНАЛІЗ СУЧАСНОГО СТАНУ ТА ПРОБЛЕМ АВТОМАТИЗАЦІЇ ПЕРЕВІРКИ ПРОГРАМНОГО КОДУ	11
1.1 Проблематика оцінювання коду та обмеження існуючих методів верифікації.....	11
1.2 Огляд та порівняльна характеристика існуючих інструментальних засобів аналізу коду.....	14
1.3 Обґрунтування використання великих мовних моделей в освітньому процесі	19
1.4 Формальна постановка задачі інтелектуальної перевірки та верифікації коду	21
Висновки до розділу 1	28
РОЗДІЛ 2 ПРОЄКТУВАННЯ АРХІТЕКТУРИ ТА СТРУКТУРИ ІНТЕЛЕКТУАЛЬНОЇ ВЕБ-СИСТЕМИ	30
2.1 Аналіз функціональних та нефункціональних вимог до веб- системи.....	30
2.2 Вибір архітектурного патерну та обґрунтування технологічного стеку.....	32
2.3 Проєктування логічної структури та моделі бази даних веб- системи.....	39
Висновки до розділу 2.....	43
РОЗДІЛ 3 АЛГОРИТМІЧНЕ ТА ІНТЕЛЕКТУАЛЬНЕ ЗАБЕЗПЕЧЕННЯ СИСТЕМИ.....	45

3.1 Методи представлення та парсингу програмного коду на основі синтаксичного аналізу.....	45
3.2 Розробка методології пром프트-інжинірингу для формування навчального фідбеку.....	48
3.3 Алгоритм інтелектуальної обробки та категоризації помилок за допомогою Gemini API.....	51
3.4 Проєктування логіки пост-процесингу та валідації відповідей нейромережевої моделі	56
Висновки до розділу 3	60
РОЗДІЛ 4 ПРОГРАМНА РЕАЛІЗАЦІЯ ТА ТЕСТУВАННЯ ВЕБ-СИСТЕМИ.....	63
4.1 Реалізація серверної частини та API-інтерфейсів взаємодії...	63
4.2 Розробка клієнтського інтерфейсу та інтеграція вебредактора коду.....	66
4.3 Налаштування інтеграції з сервісами Google Gemini та обробка виключних ситуацій	68
4.4 Тестування системи та аналіз результатів інтелектуальної перевірки	70
Висновки до розділу 4.....	76
ВИСНОВКИ	78
ПЕРЕЛІК ПОСИЛАНЬ	81
ДОДАТОК А ЛІСТИНГ ПРОГРАМИ.....	85

ПЕРЕЛІК ПРИЙНЯТИХ СКОРОЧЕНЬ

API – інтерфейс програмування застосунків (Application Programming Interface)

AST – абстрактне синтаксичне дерево (Abstract Syntax Tree)

LLM – велика мовна модель (Large Language Model)

JSON – об'єктна нотація JavaScript (JavaScript Object Notation)

ORM – об'єктно-реляційне відображення (Object-Relational Mapping)

REST – передача репрезентативного стану (Representational State Transfer)

СКБД – система керування базами даних

ШІ – штучний інтелект

ВСТУП

Навчання¹ сучасним технологіям програмування є складним процесом, який вимагає постійного балансу між кількістю практичних завдань та швидкістю надання якісного зворотного зв'язку. У багатьох навчальних закладах щорічне зростання кількості студентів комп'ютерних спеціальностей створює критичне навантаження на професорсько-викладацький склад, що неминуче призводить до суттєвих затримок у рецензуванні вихідного коду (code review). Навіть незначні недоліки в організації цього процесу можуть спричинити втрату мотивації здобувачів освіти та зниження загального рівня інженерної підготовки. З огляду на це, освітні установи все більше зацікавлені в застосуванні автоматизованих систем, здатних миттєво верифікувати алгоритмічні рішення та підтримувати безперервний процес навчання.

Один із передових підходів до вирішення цієї проблеми базується на інтеграції інструментів штучного інтелекту – зокрема генеративних великих мовних моделей (LLM), які здатні виконувати глибокий семантичний аналіз програмного тексту та надавати рекомендації в режимі реального часу. Проте в умовах потокової перевірки студентських робіт пряме використання комерційних хмарних API стикається з серйозними викликами: високою затримкою відповідей (latency) та стрімким зростанням фінансових витрат через обробку семантично однакових фрагментів коду. У цьому контексті дедалі важливішу роль відіграють оптимізовані гібридні архітектури, здатні адаптуватися до специфіки масових запитів. Поєднання методів локального детермінованого парсингу на основі абстрактних синтаксичних дерев (AST) із

¹ Тут і нижче використано такий(і) інструмент(и) 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.).

технологіями багаторівневого семантичного кешування дозволяє ефективно дедуплікувати еквівалентні інженерні рішення.

Розробка спеціалізованої системи, що поєднує нормалізацію коду на локальному сервері з інтелектуальним ядром хмарної моделі, дає змогу суттєво підвищувати швидкодію аудиту, знижувати операційні витрати на обчислювальні токени і підтримувати високу гнучкість у формуванні структурованого технічного фідбеку. Значущість дослідження полягає в його потенціалі радикально оптимізувати конвеєр обробки даних, що має стратегічне значення не лише для університетської освіти, а й для корпоративних платформ підготовки ІТ-фахівців із подібними викликами.

РОЗДІЛ 1 АНАЛІЗ СУЧАСНОГО СТАНУ ТА ПРОБЛЕМ АВТОМАТИЗАЦІЇ ПЕРЕВІРКИ ПРОГРАМНОГО КОДУ

1.1 Проблематика оцінювання коду та обмеження існуючих методів верифікації

Сучасна індустрія розробки програмного забезпечення характеризується стрімким зростанням складності архітектурних рішень та постійним підвищенням вимог до надійності, безпеки та ефективності систем. У зв'язку з цим контроль якості вихідного тексту програм перетворився з допоміжного етапу на невід'ємну частину життєвого циклу розробки (Software Development Life Cycle, SDLC), без якої неможливо гарантувати стабільність та підтримуваність програмного продукту [4]. Особливої гостроти це питання набуває в освітньому середовищі комп'ютерних спеціальностей, де перевірка коду має на меті не лише формальну верифікацію працездатності, а й формування глибоких професійних інженерних компетенцій майбутніх фахівців.

Процес підготовки кадрів у галузі Computer Science у провідних вищих навчальних закладах, зокрема у Національному технічному університеті України «Київський політехнічний інститут імені Ігоря Сікорського», сьогодні стикається з серйозними об'єктивними викликами. Кількість студентів на IT-спеціальностях неухильно зростає, що створює безпрецедентне навантаження на професорсько-викладацький склад. Це призводить до виникнення кількох критичних проблем у навчальному процесі, першою з яких є значна затримка зворотного зв'язку (Feedback Delay). Студенти часто виконують лабораторні роботи та здають їх на перевірку, але отримують розгорнуту рецензію лише через кілька тижнів, безпосередньо під час захисту. Дослідження в галузі інженерної педагогіки доводять, що ефективність навчання прямо залежить від швидкості отримання фідбеку: якщо студент бачить свої помилки відразу після написання коду, він значно краще засвоює правильні патерни розробки

[5]. Крім того, через надмірний обсяг робіт викладачі не завжди спроможні приділити однакову увагу кожному рядку програмного тексту, що вносить елемент суб'єктивності та втому в процес оцінювання. Наявні традиційні системи автоматизації (наприклад, контест-системи типу e-judge) перевіряють код виключно за принципом «чорної скриньки» шляхом проведення динамічного тестування: чи видає програма правильний результат на заданому наборі тестів. Проте код, який працює функціонально правильно, може бути написаний вкрай неефективно, мати вразливості у безпеці або грубо порушувати архітектурні принципи, що є неприпустимим для якісної професійної підготовки.

Історично склалося так, що найбільш якісним та надійним методом перевірки вважається ручне рецензування (Code Review) досвідченим ментором [6]. Такий аналіз дає змогу виявити складні логічні пастки, оцінити читабельність, чистоту коду та відповідність специфічним архітектурним патернам. Проте зазначений підхід має суттєві обмеження, які унеможливають його використання у ролі єдиного інструменту в сучасній освіті. Насамперед, спостерігається надзвичайно висока трудомісткість, адже рецензування лише однієї студентської роботи середньої складності може займати від 15 до 40 хвилин чистого часу фахівця. Також значний вплив має людський чинник, оскільки експерт може пропустити критичну помилку через втому або зниження концентрації уваги. У поєднанні зі зростанням чисельності студентських груп це створює складну проблему масштабованості, оскільки ручний підхід фізично неможливо застосувати для великих потоків без критичної втрати якості та деталізації аналізу.

Для автоматизації рутинних перевірок вихідного тексту програм були розроблені засоби статичного аналізу (SAST), зокрема класичні лінтери (Pylint, Flake8). Ці інструменти базуються на аналізі абстрактних синтаксичних дерев (AST) та наборі жорстко прописаних детермінованих правил і регулярних виразів [7]. Незважаючи на високу швидкість роботи, подібні інструменти мають суттєву семантичну обмеженість, оскільки у них

повністю відсутнє розуміння наміру (Intent) автора програмного тексту. Літерспроможний зафіксувати невідповідність довжини назви змінної або порушення відступів, проте він не здатний усвідомити факт того, що студент переплутав два схожих за логікою, але різних за суттю алгоритми, або обрав неоптимальну структуру даних для вирішення конкретного завдання. Додатковою проблемою є велика кількість хибних спрацьовувань (False Positives), що зрештою змушує студента ігнорувати повідомлення системи через виникнення ефекту «втоми від сповіщень». При цьому такі засоби мають низьку освітню цінність, адже сухі технічні коди помилок без контекстних роз'яснень природною мовою не вчать здобувача освіти способам уникнення та виправлення подібних дефектів у майбутньому.

Більш просунуті комплексні інструменти аналізу якості (наприкладі платформи SonarQube) використовують складні метрики обчислення цикломатичної та когнітивної складності для інтегральної оцінки технічного боргу [8]. Проте їх впровадження у навчальний процес вищої школи є ускладненим через громіздкість інфраструктури, потребу у виділених серверах, СКБД і тривалому конфігуруванні під кожен конкретний проект, оскільки вони орієнтовані на велику промислову розробку. Аналітичні звіти таких платформ розраховані на досвідчених інженерів чи керівників проектів, а не на студентів, які тільки вивчають ази алгоритмізації та опановують базові парадигми програмування.

Проведений аналіз сучасного стану демонструє існування глибокого «семантичного розриву» між можливостями наявних автоматизованих засобів верифікації та реальними потребами вищої ІТ-освіти. Студенту в процесі навчання потрібен не просто лічильник синтаксичних відхилень чи лог компілятора, а повноцінний інтелектуальний помічник, здатний до глибокого семантичного аналізу коду – розуміння логіки, архітектурного контексту та алгоритмічної суті рішення. Використання великих LLM, зокрема хмарного інструментарію Google Gemini API, відкриває принципово нові можливості для подолання цих обмежень [12]. Завдяки здатності сучасних

нейромережових архітектур до інтерпретації програмного тексту та коду як єдиної семантичної структури, стає можливим створення систем, які об'єднують швидкість автоматизації з глибиною людського рецензування, забезпечуючи розуміння намірів програміста та надання адаптивних навчальних рекомендацій. Проте пряма інтеграція таких хмарних ШІ-сервісів у навчальний процес викликає нові інженерні виклики: високу вартість обчислювальних токенів, накладні часові витрати на очікування відповідей та ризики блокування запитів через Rate Limits [18]. Це зумовлює необхідність побудови гібридних систем, де можливості великих мовних моделей оптимізуються за допомогою локальних детермінованих алгоритмів аналізу структури коду та багаторівневого семантичного кешування результатів перевірки.

1.2 Огляд та порівняльна характеристика існуючих інструментальних засобів аналізу коду

Для розробки ефективної інтелектуальної системи автоматизованої верифікації необхідно провести глибокий критичний аналіз існуючих технологічних рішень, які використовуються для контролю якості програмного забезпечення [14]. Сучасний ринок інструментарію пропонує широкий спектр засобів, які класифікуються за методом аналізу, рівнем заглиблення в структуру вихідного тексту та цільовим призначенням.

Статичні аналізатори коду, або літери, є першим ешеленом захисту якості у сучасних пайплайнах розробки. Найбільш репрезентативними представниками цієї категорії є інструменти Pylint для мови програмування Python та ESLint для екосистеми JavaScript. Алгоритмічна база літерів ґрунтується на детермінованому парсингу вихідного тексту програм та побудові абстрактних синтаксичних дерев [15]. Парсер послідовно

перетворює лінійний текст програми у розгалужену деревоподібну структуру, де кожен вузол відповідає певній конструкції граматики мови, наприклад, оголошенню змінної, тілу циклу чи сигнатурі функції. Після побудови графа система застосовує набір жорстко прописаних детермінованих правил для пошуку дефектів. Основним аспектом роботи лінтерів є виявлення ознак слабого проєктування або так званих «запахів коду» (Code Smells). Вони ефективно знаходять невикористані ідентифікатори, занадто довгі тіла функцій, а також фіксують порушення офіційних стилістичних стандартів іменування та оформлення, таких як специфікація керівництва з написання коду Python Enhancement Proposal 8 (PEP 8). Крім того, лінтери забезпечують синтаксичну валідацію, гарантуючи відповідність програмного тексту формальній граматиці, що дозволяє заблокувати критичні помилки типу `SyntaxError` до етапу безпосереднього виконання програми. Проте до фундаментальних недоліків лінтерів належить повна відсутність аналізу потоку даних (Data Flow Analysis) та ігнорування семантики тексту. Статичний лінтер принципово неспроможний визначити логічну коректність алгоритму або передбачити динамічний стан системи під час виконання.

Наступним рівнем автоматизації є спеціалізовані платформи комплексного аналізу якості та оцінки безпеки рівня підприємства (Enterprise Solutions), найбільш відомими представниками яких є системи SonarQube та DeepSource. Дані рішення призначені для комплексного моніторингу технічного боргу масштабних програмних систем [16]. Методологія їхнього аналізу базується на комбінації кількох підходів. По-перше, вони здійснюють автоматизоване обчислення складних структурних метрик, зокрема цикломатичної складності за Маккейбом та когнітивної складності алгоритмів. По-друге, ці платформи використовують спеціалізовані алгоритми індексування та порівняння хеш-функцій для виявлення дублювання фрагментів коду (Copy-Paste Detection). По-третє, вони виконують сканування на наявність потенційних вразливостей та зон ризику (Security Hotspots) на основі галузевих баз даних, таких як відкритий проєкт із безпеки вебдодатків

(Open Worldwide Application Security Project, OWASP), блокуючи можливості проведення атак типу SQL-ін'єкцій або переповнення буфера. Незважаючи на високу промислову ефективність, впровадження платформ класу SonarQube у реальний навчальний процес університетських кафедр стикається з проблемою високого порогу входження для здобувачів освіти. Студенту молодших курсів вкрай важко інтерпретувати узагальнені графіки надійності (reliability) чи супроводжуваності (maintainability) та перераховувати абстрактний технічний борг у хвилинах без контекстного, точкового роз'яснення кожної конкретної помилки.

З появою та розвитком великих мовних моделей сформувався принципово новий клас інструментів інтелектуального аналізу коду нового покоління, таких як GitHub Copilot або загальні інтерфейси типу ChatGPT [17]. Вони функціонують на базі нейромережевої архітектури Transformer, яка пройшла етап попереднього навчання на терабайтах відкритих репозиторіїв вихідного коду. До ключових переваг ШІ-асистентів належить глибока контекстуальна обізнаність. На відміну від класичних статичних аналізаторів, мовні моделі здатні інтерпретувати код не просто як жорсткі граматичні конструкції, а як семантичну послідовність токенів з урахуванням складних нелінійних логічних зв'язків, що дозволяє виявляти тонкі логічні помилки в алгоритмах. Крім того, критично важливим для освітньої мети є використання природної мови: ШІ генерує детальний розбір дефектів людською мовою, адаптуючи термінологію під контекст поточного запиту.

Разом з тим, пряме використання універсальних ШІ-асистентів у вищій школі виявило серйозні деструктивні фактори [18]. Першим є так званий «ефект милиці», коли інструменти на кшталт GitHub Copilot автоматично дописують логічні блоки або цілі функції за студента. Це повністю блокує розвиток самостійного алгоритмічного мислення та веде до швидкої деградації інженерних навичок програмування. Другим фактором є проблема нейромережових галюцинацій, за яких модель може впевнено згенерувати посилання на неіснуючі бібліотеки, методи або видати синтаксично

помилкову конструкцію за правильну. Третім фактором виступає закрита комерційна модель більшості якісних інструментів, що потребують дорогої передплати та суттєво обмежує масове впровадження в освітній процес державних університетів.

Для проведення об’єктивного системного аналізу та зіставлення характеристик розглянутих інструментів розроблено порівняльну матрицю. Усі засоби оцінено за ключовими критеріями, які є визначальними для побудови ефективної системи автоматизованої перевірки студентських робіт. Порівняльні характеристики наведено в таблиці 1.1.

Таблиця 1.1 – Порівняльна характеристика інструментальних засобів аналізу коду

<i>Критерій порівняння</i>	<i>Лінтери (Pylint)</i>	<i>SonarQube</i>	<i>LLM-асистенти</i>	<i>Пропонована система</i>
Об’єкт аналізу	Формальний синтаксис та стилістика коду	Структурні метрики, дублікати та безпека	Семантика, логіка та контекст рішення	Структура, семантика та алгоритмічна логіка
Метод перевірки	Детерміновані правила (Regex / AST)	Комбінований (метрики + статичні правила)	Нейромережевий аналіз тексту програми	Гібридний (Дворанговий AST-аналіз + Gemini API)
Пояснення помилок	Мінімальне (сухі коди помилок та попереджень)	Шаблонні звіти та оцінка технічного боргу	Високе (розлогі відповіді природною мовою)	Адаптивне навчання (структурований JSON з дефектами)
Розуміння логіки	Повністю відсутнє	Обмежене (на основі розрахунку складності)	Високе (але з ризиком галюцинацій)	Високе (з валідацією через канонічний AST-фільтр)

Кінець табл. 1.1

<i>Критерій порівняння</i>	<i>Лінтери (Pylint)</i>	<i>SonarQube</i>	<i>LLM-асистенти</i>	<i>Пропонована система</i>
Доступність для студента	Висока (вбудовані безкоштовні утиліти)	Низька (потребує розгортання інфраструктури)	Низька (комерційна ліцензія / обмеження)	Висока (веб-інтерфейс, оптимізований кеш-шар)

Проведений аналіз та побудована порівняльна характеристика підтверджують існування технологічного вакууму в освітній сфері. Існуючі індустріальні рішення або занадто примітивні й обмежуються перевіркою синтаксису (лнтери), або занадто складні й орієнтовані на великі комерційні команди (SonarQube). Універсальні AI-асистенти успішно вирішують інженерну задачу розуміння логіки коду, проте створюють неприпустимі освітні ризики, виконуючи роботу замість студента. Це повністю доводить наукову та практичну доцільність розробки спеціалізованої гібридної системи на базі Google Gemini API. Пропонована архітектура використовує обчислювальний потенціал великих мовних моделей не для генерації готових рішень, а для їхнього глибокого покрокового аналізу, рецензування та верифікації. При цьому впровадження проміжного локального AST-шару та багаторівневого кешування семантичних відбитків дозволяє повністю нівелювати такі недоліки хмарних моделей, як тривалий час затримки відповіді та висока вартість обчислювальних токенів.

1.3 Обґрунтування використання великих мовних моделей в освітньому процесі

Розвиток технологій штучного інтелекту, зокрема поява LLM, спричинив парадигмальний зсув у методології навчання технічних дисциплін. Традиційні комп'ютерні системи навчання здебільшого виконували роль пасивних репозиторіїв знань або жорстко алгоритмізованих тестерів [20]. Впровадження інтелектуальних агентів на базі сучасних архітектур нейромереж дозволяє трансформувати цей процес у динамічну інженерну взаємодію, де програмна система виконує функцію автоматизованого експертного інструменту для проведення технічного аудиту коду.

Однією з ключових проблем у підготовці фахівців з комп'ютерних наук є необхідність отримання студентом миттєвого та якісного зворотного зв'язку щодо написаного алгоритмічного тексту. Дослідження в галузі інженерної педагогіки підтверджують, що ефективність засвоєння складних концепцій, таких як алгоритмізація та архітектурне проєктування, прямо корелює зі швидкістю отримання аналітичного фідбеку [21]. Використання великих мовних моделей забезпечує автоматизацію процесу рецензування (code review). Замість загальних абстрактних оцінок інтелектуальний модуль виконує глибокий семантичний аналіз намірів автора програми. На відміну від класичних аналізаторів, мовні моделі здатні усвідомити цілісний контекст інженерного завдання та чітко локалізувати місця, де саме поточна реалізація розходиться з логікою алгоритму, формуючи індивідуальні технічні рекомендації з урахуванням специфіки написання коду конкретним студентом.

Подолання «семантичного розриву» між синтаксичною структурою програми та її високорівневою логікою стало можливим завдяки специфіці нейромережевої архітектури Transformer [19]. Ключовим інженерним механізмом у даному контексті є апарат самоуваги (Self-Attention), який

дозволяє моделі ефективно враховувати дальні залежності tokenів у тексті. Система фіксує нелінійний зв'язок між оголошенням змінної чи ініціалізацією об'єкта на початку файлу та їх некоректним або неоптимальним використанням у вкладеному циклі через сотні рядків коду. Аналізуючи програмний текст як єдину послідовність tokenів з урахуванням складних семантичних зв'язків, модель розпізнає високорівневі патерни проєктування та їхні деструктивні антипатерни, які практично неможливо формально описати жорсткими правилами регулярних виразів класичних лінтерів чи статичних аналізаторів.

Для реалізації інтелектуальної веб-системи в межах даного дипломного проєкту було обґрунтовано обрано хмарний інструментарій Google Gemini API, що зумовлено комплексом технологічних та економічних чинників. Важливим архітектурним фактором виступає масштабоване вікно контексту (Context Window), що дозволяє моделі Gemini обробляти значні обсяги вихідного тексту за один ітераційний запит, повністю мінімізуючи ризики втрати інформації при аналізі розгалужених програмних модулів. З інженерної точки зору критично важливою перевагою обраного інструменту є рідна підтримка режиму структуроване виведення даних/результатів (Structured JSON Output). Це дозволяє безпосередньо на рівні запиту вимагати від моделі повернення аналітичних даних у суворо визначеному форматі схеми, що надалі дає змогу виконувати автоматичну десеріалізацію та строгу валідацію через Pydantic-моделі на бекенді, унеможливаючи збої при рендерингу графічного інтерфейсу користувача та відсікаючи зайвий текстовий шум. Додатковим фактором виступає економічна доступність у межах безкоштовного рівня використання (Free Tier), що дозволяє успішно розгорнути та всебічно протестувати систему без залучення значних бюджетних ресурсів університету.

Критичним питанням при впровадженні моделей штучного інтелекту в навчання є суворе дотримання етичних норм та принципів академічної доброчесності. Пропонована система проєктується не як генератор готових

рішень, а як засіб інтелектуального аудиту та верифікації. Основна інженерна стратегія взаємодії системи із користувачем базується на концепції структурованого квантування зауважень. Модель штучного інтелекту через жорсткі системні інструкції (System Prompts) обмежена у можливості видачі суцільних фрагментів виправленого коду. Замість цього система повертає строго типізований масив мікро-зауважень із прив'язкою до номерів рядків, де вказано тип дефекту та інженерне обґрунтування його неоптимальності. Такий підхід повністю унеможливорює використання системи як «милиця» для списування, гарантує самостійність виконання лабораторних робіт та стимулює розвиток алгоритмічного мислення студентів, які мають самостійно виправити код на основі отриманих структурованих метрик.

1.4 Формальна постановка задачі інтелектуальної перевірки та верифікації коду

Для розробки інтелектуальної веб-систем и автоматизованої верифікації програмного тексту необхідно вийти за межі суто описового інженерного підходу та побудувати строгу формальну модель процесу перевірки коду [14]. Задача інтелектуального оцінювання та маршрутизації програмного тексту може бути подана як процес багатоетапного математичного перетворення вихідного коду у структуровану множину діагностичних знань про його якість, безпеку, синтаксичну та логічну коректність.

Нехай S – вихідний текст програми, надісланий студентом на перевірку, представлений як впорядкована послідовність токенів $T = \{t_1, t_2, \dots, t_n\}$, де n – загальна кількість лексем у коді. Програма написана мовою програмування L , що однозначно визначається своєю формальною граматикою G_L . Окрім безпосереднього вихідного тексту програми, вхідними даними системи є контекст навчального завдання K , який включає умови лабораторної роботи

або специфікацію алгоритму. Таким чином, первинним об'єктом аналізу є вхідний інформаційний кортеж:

$$I = \langle C, L, K \rangle, \quad (1.1)$$

де I – первинний вхідний інформаційний кортеж системи;

C – вихідний текст програми, надісланий студентом на перевірку;

L – мова програмування, якою написано вихідний текст;

K – контекст навчального завдання (умови лабораторної роботи або специфікація алгоритму).

Процес гібридної перевірки коду можна описати як узагальнену функцію F , яка відображає вхідний простір програмних кортежів, визначених у формулі (1.1), на комплексний простір експертних висновків та оцінок R , тобто:

$$F : I \rightarrow R, \quad (1.2)$$

де F – узагальнена функція гібридної перевірки та верифікації коду;

I – первинний вхідний інформаційний кортеж даних;

R – комплексний простір структурованих експертних висновків та оцінок якості програми.

У традиційних SAST-системах простір R формується сухими кодами помилок компілятора, проте у розроблюваній системі R є структурованим об'єктом, що складається з множини виявлених дефектів коду $D = \{d_1, d_2, \dots, d_m\}$ (де кожен дефект d_i характеризується типом, номером рядка та рівнем критичності), чисельної інтегральної оцінки якості алгоритму $Q \in [0,100]$ та адаптивного роз'яснювального тексту E , сформованого природною мовою у межах концепції автоматизованого технічного рецензування.

Для мінімізації накладних обчислювальних витрат, уникнення часових затримок та захисту від навантажувальних колізій хмарного інтерфейсу,

загальна функція перевірки F , представлена виразом (1.2), декомпонується на конвеєр багаторівневого аналізу. На першому етапі здійснюється локальне синтаксичне відображення вихідного тексту програми у простір абстрактного синтаксичного дерева (AST) за допомогою оператора парсингу Ψ :

$$T_{AST} = \Psi(C), \quad (1.3)$$

де T_{AST} – згенероване абстрактне синтаксичне дерево програми;

Ψ – детермінований оператор синтаксичного розбору (парсингу).

Для забезпечення інваріантності системи до операцій потенційного списування чи тривіального знеособлення коду (таких як зміна імен ідентифікаторів, функцій, аргументів, стилю відступів або видалення коментарів) вводиться оператор канонічної нормалізації N . Цей оператор трансформує дерево, замінюючи кастомні ідентифікатори на уніфіковані плейсхолдери, а також видаляючи нефункціональні вузли (коментарі, виклики функцій логування), що формує нормалізовану структуру $T_{norm} = N(T_{AST})$.

Рівень першого рангу кешування (L1-кеш) базується на детермінованій функції криптографічного згортання нормалізованого дерева для точного структурного зіставлення:

$$\text{Hash}_{\text{AST}} = \text{SHA-256}(\text{T}_{\text{norm}}), \quad (1.4)$$

де Hash_{AST} – криптографічна згортка першого рангу кешування (L1-кеш);
 SHA-256 – алгоритм криптографічного хешування з довжиною дайджесту 256 біт;

T_{norm} – нормалізоване абстрактне синтаксичне дерево вихідного коду, очищене від нефункціональних елементів.

Для оптимізації розгалужених та структурно модифікованих варіантів алгоритмів (наприклад, при заміні студентом циклу `for` на `while` при збереженні загальної бізнес-логіки та мутацій даних) вводиться вектор евристичних ознак другого рангу (L2-кеш):

$$F_{\text{semantic}} = \{x_1, x_2, \dots, x_k\}, \quad (1.5)$$

де F_{semantic} – вектор евристичних ознак другого рангу кешування (L2-кеш);

x_1 – логічний прапор наявності циклічних конструкцій у коді;

x_2 – максимальний рівень вкладеності логічних блоків;

x_3 – підмножина сигнатур активованих методів мутації об'єктів;

x_4 – оператори математичних та компараторних дій;

k – загальна кількість обчислених евристичних параметрів статичного аналізу.

За такої архітектури математична логіка прийняття рішення про маршрутизацію запиту студентської роботи та формування результуючого фідбеку R на основі компонентів (1.3)–(1.5) набуває вигляду розгалуженої функції (1.6):

$$R = \begin{cases} \Phi_{DB}(\text{Hash}_{AST}), & \text{Hash}_{AST} \in \mathcal{D}_{L1} \\ \Phi_{DB}(F_{semantic}), & \text{Hash}_{AST} \notin \mathcal{D}_{L1} \wedge F_{semantic} \in \mathcal{D}_{L2} \\ \Omega_{LLM}(S, C, K), & \text{Hash}_{AST} \notin \mathcal{D}_{L1} \wedge F_{semantic} \notin \mathcal{D}_{L2}, \end{cases} \quad (1.6)$$

де $\mathcal{D}_{L1}$ та $\mathcal{D}_{L2}$ – реляційні простори збережених раніше успішних транзакцій у базі даних системи, Φ_{DB} – оператор миттєвого витягування результату з кешу, Ω_{LLM} – неймережевий оператор моделі Google Gemini, а S – налаштована системна інструкція (системна підказка) для формування дидактичного стилю та структурованого формату відповіді.

Графічну інтерпретацію розробленої математичної моделі, логіку розгалуження функцій маршрутизації та концептуальний трирівневий конвеєр обробки вхідних транзакцій системи верифікації зображено на рис. 1.1.

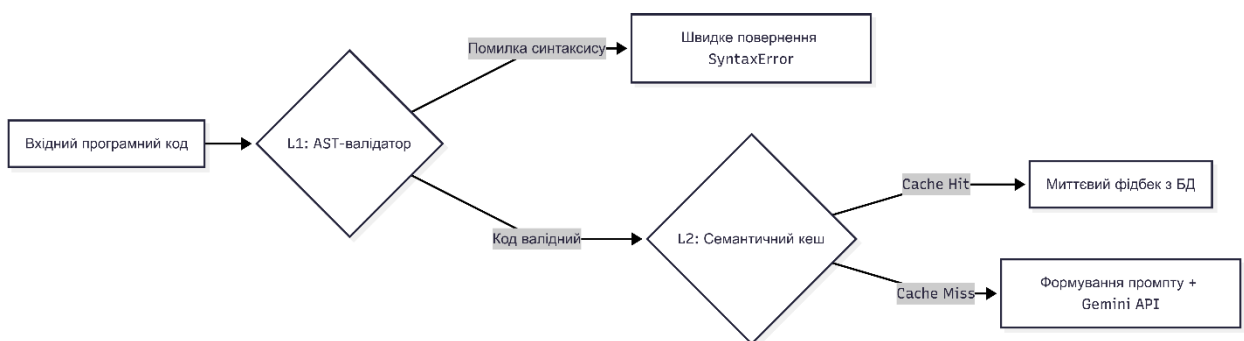


Рисунок 1.1 – Схема конвеєризації та маршрутизації запитів у системі верифікації коду

Процес генерації фідбеку через оператор штучного інтелекту формалізується як максимізація ймовірності формування коректного структурованого звіту:

$$R = \arg \max P(R|S, C, K), \quad (1.7)$$

де $\arg \max$ – оператор знаходження аргументу, що максимізує значення функції;

$P(R|S, C, K)$ – умовна ймовірність формування структурованого звіту R за умови заданого системного промпту S , тексту вихідного коду C та контексту навчального завдання K .

Для оцінки ефективності та точності функціонування розробленого конвеєра верифікації впроваджуються класичні метрики оцінювання. Першою є точність (Precision), яка визначає відношення істинно виявлених дефектів до загальної кількості згенерованих системою зауважень, що мінімізує хибнопозитивні спрацьовування (False Positives):

$$\text{Precision} = \frac{TP}{TP+FP}, \quad (1.8)$$

де TP (True Positives) – кількість істинно позитивних спрацьовувань системи (правильно виявлені дефекти коду);

FP (False Positives) – кількість хибно позитивних спрацьовувань (помилково зафіксовані дефекти).

Другою метрикою виступає повнота (Recall), яка характеризує здатність системи локалізувати всі наявні у студентському коді архітектурні та семантичні помилки, мінімізуючи пропуски багів (False Negatives):

$$\text{Recall} = \frac{TP}{TP+FN}, \quad (1.9)$$

де FN (False Negatives) – кількість хибно негативних результатів (пропущені системою реальні помилки в коді).

Кінцевою інженерною метою налаштування системи є максимізація показників точності (1.8) та повноти (1.9) за одночасного підвищення показника когнітивно-навчального ефекту. Цей показник визначається як здатність студента самостійно рефакторити власний код на основі наданих структурованих аналітичних підказок без прямої видачі готового програмного рішення, що розвиває самостійне алгоритмічне мислення фахівця. На основі побудованої формальної постановки задачі, основними науково-технічними завданнями на наступні етапи дослідження є проєктування масштабованої веб-архітектури конвеєра F , розробка логічної схеми реляційної бази даних для збереження інформаційних кортежів та оптимізація стратегій промпт-інжинірингу для забезпечення стійкості відповідей у форматі Structured JSON.

Висновки до розділу 1

За підсумками дослідження, викладеного в першому розділі, здійснено ґрунтовний огляд технологічного ландшафту та виокремлено ключові бар'єри в сегменті автоматизованого аудиту кодових баз під час навчання інженерів. З'ясовано, що масштабування освітніх програм, характерне для флагманських закладів рівня Національного технічного університету України «Київський політехнічний інститут імені Ігоря Сікорського», критично загострює проблему тривалого очікування рецензій (Feedback Delay), що деструктивно впливає на когнітивну динаміку студентів. Детальний розбір класичних засобів стаціонарного тестування та корпоративних екосистем (на кшталт SonarQube) довів існування вираженого «семантичного розриву». Наявні платформи здатні відстежувати суто лексико-синтаксичні аномалії, проте демонструють абсолютну неспроможність до дешифрування алгоритмічного задуму та глибокого архітектурного інтену розробника.

Як технологічне ядро інноваційної веб-платформи аргументовано вибір нейромережових великих мовних моделей, а саме хмарної екосистеми Google Gemini API. Специфіка інтегрованої архітектури Transformer у синергії з механізмом самоуваги (Self-Attention) забезпечує здатність інтелектуального агента до комплексного усвідомлення розгалужених програмних модулів. Застосування парадигми Structured Outputs відкриває можливість конвертувати вільні текстові відповіді моделі у жорстко типізований формат JSON-аналітики. Це радикально усуває генерацію надлишкового текстового шуму, гарантуючи надання високоточного інженерного фідбеку з адресною прив'язкою виявлених дефектів до відповідних рядків сирцевого файлу.

Конвеєр когнітивної верифікації був суворо алгоритмізований шляхом побудови формальної математичної моделі, що описує каскадну трансляцію вхідних скриптів у багатовимірний простір експертних вердиктів. З метою нейтралізації мережевої латентності та мінімізації фінансових витрат на

токенізацію, спроектовано обчислювальну модель дворівневого семантичного кешування. Вона оперує механізмами канонічної редукції абстрактних синтаксичних дерев (AST) і генеруванням евристичних сигнатур алгоритмічної логіки. Кваліметрія запропонованої архітектури спирається на базові індикатори точності (Precision) та повноти (Recall), математичний апарат яких наведено у формулах (1.8) та (1.9) відповідно.

Акумуляована аналітична база та розроблений математичний апарат маршрутизації транзакцій утворюють монолітний теоретичний фундамент дослідження. Виведені закономірності слугуватимуть безпосередньою базою для подальшого інженерного проєктування клієнт-серверної інфраструктури, конструювання реляційної схеми збереження даних та безпосереднього програмування веб-додатка у наступних етапах дипломного проєкту.

РОЗДІЛ 2 ПРОЄКТУВАННЯ АРХІТЕКТУРИ ТА СТРУКТУРИ ІНТЕЛЕКТУАЛЬНОЇ ВЕБ-СИСТЕМИ

2.1 Аналіз функціональних та нефункціональних вимог до веб-системи

Розробка інтелектуальної веб-системи автоматизованої перевірки коду вимагає чіткого визначення функціональних меж та технічних параметрів ефективності. Згідно з методологією системного аналізу та стандартом ISO/IEC/IEEE 29148 (який замінив застарілий IEEE 830), вимоги до програмного забезпечення класифікуються на функціональні, які визначають безпосередню поведінку системи та її реакцію на дії користувачів, та нефункціональні, що встановлюють ключові атрибути якості, надійності, захищеності та обчислювальної продуктивності [22].

Функціональні вимоги визначають спектр сервісів, які надає архітектура системи для забезпечення автоматизованого навчального процесу та безпосередньої взаємодії студента з платформою в режимі реального часу. Для досягнення цілей дослідження архітектурні межі системи було структуровано навколо чотирьох ключових функціональних доменів.

Першим є модуль ідентифікації та профілювання користувачів. З метою спрощення інтерфейсу та забезпечення високої пропускну здатності на етапі лабораторного тестування, ідентифікація здійснюється за детермінованим принципом введення унікального ідентифікатора (імені та прізвища студента) безпосередньо у робочу область веб-редактора, що автоматично зв'язує поточну сесію з реляційною транзакцією у базі даних.

Другим доменом виступає модуль безпосередньої взаємодії з вихідним кодом. Основним робочим інструментом інтерфейсу є інтегрований інтерактивний веб-редактор, побудований на базі компонента Monaco Editor, що забезпечує повноцінне підсвічування синтаксису, автодоповнення та

форматування вихідного тексту для мови програмування Python [23], яка є базовою для вивчення фундаментальних дисциплін алгоритмізації на кафедрі.

Третім є модуль гібридного аналізу та семантичної верифікації, що становить обчислювальне ядро системи. Він реалізує вимоги щодо розгортання багаторівневого конвеєра перевірки: від первинного локального синтаксичного та безпекового аналізу через абстрактні синтаксичні дерева (AST) до інтелектуального оцінювання за допомогою хмарного сервісу Google Gemini API, що керується оптимізованими стратегіями промпт-інжинірингу [3]. Система виконує автоматичну обробку відповіді нейромережі, десеріалізує та структурує дані, класифікуючи виявлені дефекти за рівнем критичності (критичні помилки, зауваження, поради).

Четвертим доменом є модуль звітності, логування та візуалізації. Результатом його роботи є рендеринг інтерактивного фідбеку ментора з підсвічуванням конкретних проблемних рядків та відображенням історичного журналу транзакцій (таблиця submissions) у нижній частині інтерфейсу, що дозволяє викладачу та студенту моніторити часові метрики та логічну структуру надісланих рішень.

Нефункціональні вимоги встановлюють суворі критерії якості, які гарантують стабільне, безпечне та відмовостійке функціонування платформи в умовах інтенсивного навантаження університету. Першочерговою вимогою є продуктивність системи. Завдяки впровадженню архітектурного шару дворангового семантичного кешування (L1 точний AST-хеш та L2 евристичний відбиток логіки), час відгуку системи при знаходженні схожих або ідентичних рішень у базі даних скорочується до детермінованого мінімуму (менше 1 секунди), тоді як час прямої генерації через хмарний ШІ не повинен перевищувати середні 7–10 секунд, забезпечуючи паралельну роботу потоку користувачів без деградації пропускної здатності.

Надійність та безпека даних реалізуються на рівні суворої об'єктної валідації вхідних даних за допомогою моделей Pydantic, що повністю блокує спроби передачі некоректних запитів. Використання об'єктно-реляційного

відображення (ORM SQLAlchemy) у поєднанні з параметризованими запитами надійно захищає систему від атак типу SQL-ін'єкцій на базу даних PostgreSQL [24]. Крім того, впроваджений локальний AST-фільтр на етапі компіляції виявляє та миттєво відсікає потенційно небезпечні інструменти динамічного виконання коду (`eval`, `exec`), запобігаючи несанкційонованому віддаленому використанню ресурсів сервера.

Вимоги до зручності використання (Usability) передбачають побудову адаптивного, інтуїтивно зрозумілого інтерфейсу на основі фреймворку Tailwind CSS, що підтримує принципи доступності даних та коректного відображення компонентів на пристроях з різною роздільною здатністю екрану. Масштабованість архітектури забезпечується модульним розділенням бізнес-логіки бекенду, що дозволяє модернізувати локальні аналізатори або здійснювати міграцію на новіші версії мовних моделей без переписування системних інтерфейсів. Вимоги до розгортання та портативності визначають обов'язкове використання технології контейнеризації Docker та маніфестів Docker Compose, що ізолює залежності фронтенду, бекенду та бази даних, дозволяючи розгорнути систему на будь-яких серверах кафедри за допомогою єдиної уніфікованої команди.

2.2 Вибір архітектурного патерну та обґрунтування технологічного стеку

Проектування архітектури інтелектуальної веб-системи автоматизованої верифікації коду ґрунтується на системних вимогах до високої обчислювальної продуктивності, горизонтальної масштабованості та низького часу затримки при інтеграції із зовнішніми хмарними сервісами. Для реалізації поставлених завдань було обрано клієнт-серверну архітектуру, що базується на фундаментальному інженерному принципі розділення відповідальності

(Separation of Concerns). Система побудована за модульним принципом, де клієнтська частина (Frontend) та серверна частина (Backend) є повністю автономними застосунками, які взаємодіють між собою виключно за допомогою безстанового протоколу передачі даних через асинхронний інтерфейс RESTful API [25]. Такий підхід забезпечує незалежність розробки, спрощує модернізацію компонентів і дозволяє реалізувати неблокуючий конвеєр обробки тривалих транзакцій ШІ-ядра. Схематичне відображення розробленої архітектурної моделі та логіки взаємодії компонентів системи представлено на рисунку 2.1.

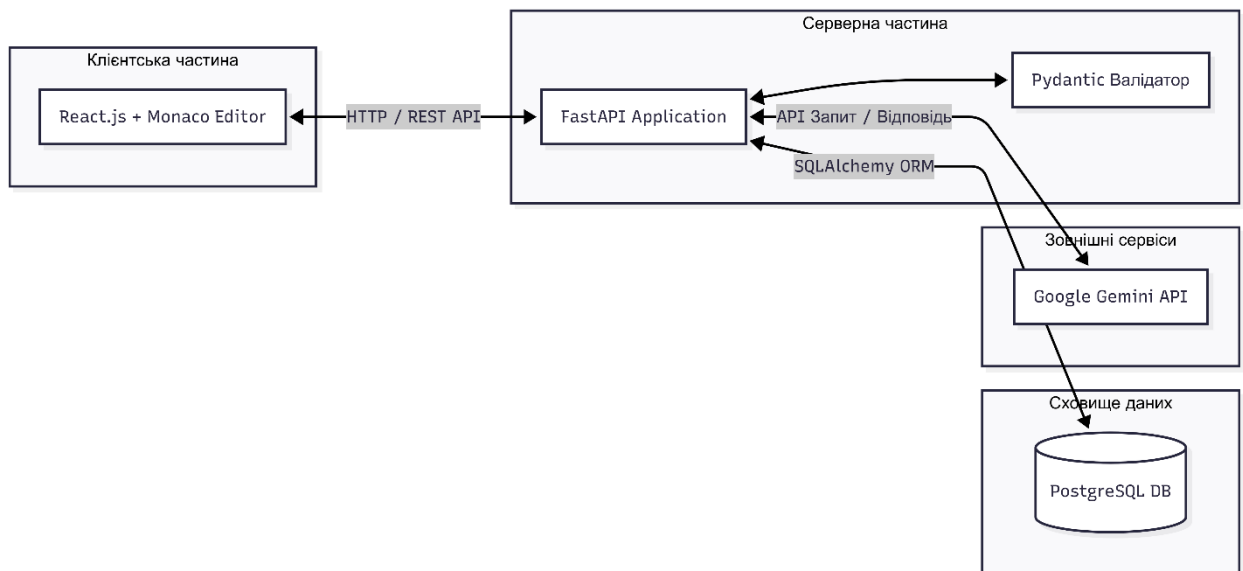


Рисунок 2.1 – Архітектура клієнт-серверної взаємодії системи

Як продемонстровано на схемі архітектурної взаємодії, загальний життєвий цикл обробки запиту починається на стороні браузера користувача, де адаптивний інтерфейс відправляє асинхронний HTTP POST запит, що містить тіло структури `CodeAnalyzeRequest`. Завдяки використанню асинхронного серверного фреймворку FastAPI, головний потік сервера не блокується під час виконання складних локальних синтаксичних обходів дерева AST або тривалого очікування відповіді від хмарного сервісу Google

Gemini API. Сервер розпаралелює запити користувачів, забезпечуючи транзакційну ізолюваність і стабільну швидкість відгуку веб-інтерфейсу для великої кількості одночасних сесій студентів кафедри.

Технологічний стек проєктованої системи підбирався з урахуванням критеріїв максимальної ефективності низькорівневого аналізу програмного тексту мовою Python та зручності оркестрації структурованих об'єктів даних. Для розробки серверної частини (Backend) було обрано високоефективний фреймворк FastAPI, який функціонує на базі мови програмування Python [26]. Цей вибір обґрунтований двома стратегічними інженерними факторами. По-перше, FastAPI демонструє одну з найвищих швидкостей обробки запитів серед усіх існуючих інструментів на Python завдяки інтеграції з сервером Uvicorn та системною бібліотекою Starlette. Він надає нативну підтримку асинхронних конструкцій `async/await`, що дозволяє ефективно масштабувати систему при роботі з зовнішнім мережевим вводом-виводом (I/O-bound operations), яким і є запити до хмарного штучного інтелекту. По-друге, фреймворк містить вбудований механізм суворої декларативної валідації вхідних та вихідних структур даних за допомогою моделей Pydantic [27]. Це гарантує автоматичну перевірку коректності типів даних, що передаються у JSON-запитах (поля `code`, `language`, `student_name`), ще до моменту активації бізнес-логіки бекенду.

Для об'єктивного підтвердження доцільності впровадження даного інструменту було виконано порівняльний аналіз поширених серверних веб-фреймворків, результати якого наведено в таблиці 2.1.

Таблиця 2.1 – Порівняльний аналіз серверних веб-фреймворків

<i>Критерій порівняння</i>	<i>FastAPI</i>	<i>Flask</i>	<i>Django</i>
Продуктивність (швидкість)	Висока (на базі Starlette та Uvicorn)	Низька / Середня	Середня
Асинхронність (async/await)	Нативна вбудована підтримка	Обмежена (потребує додаткових бібліотек)	Обмежена (додана в пізніх версіях)
Валідація даних	Автоматична через моделі Pydantic	Відсутня (потребує WTForms / Marshmallow)	Вбудована (через Django Forms / DRF)
Генерація документації	Автоматична (OpenAPI / Swagger)	Потребує сторонніх розширень	Потребує сторонніх розширень (drf-yasg)

Для організації надійного довготривалого збереження інформації про надіслані програмні рішення, результати локального та інтелектуального аналізів, а також для ведення історичного журналу транзакцій було обрано реляційну систему управління базами даних PostgreSQL [24]. Цей вибір обумовлений високою транзакційною стійкістю системи, її повною відповідністю критеріям набіру транзакційних властивостей баз даних, що включає атомарність, узгодженість, ізолюваність та довговічність (Atomicity, Consistency, Isolation, Durability, ACID) та гнучкістю при роботі зі змішаними типами даних. Взаємодія між об'єктними моделями бекенду та реляційними таблицями бази даних реалізована за допомогою об'єктно-реляційного мапера SQLAlchemy з використанням патерну впровадження залежностей Depends(get_db). Це дозволяє автоматично ізолювати та закривати сесії з'єднання з базою після виконання кожного запиту. Зберігання складних

деревоподібних результатів семантичного аналізу та звітів про дефекти, отриманих від моделі Gemini, реалізовано через тип даних Text (або JSONB), у якому накопичуються валідовані серіалізовані об'єкти типу Structured JSON. Такий підхід дає змогу акумулювати великі обсяги аналітичної статистики профілів студентів у таблиці submissions без необхідності постійної та складної перебудови схеми бази даних при зміні структури ШІ-відповідей.

Інтелектуальним ядром системи (AI Engine) виступає інтерфейс Google Gemini API [3]. Вибір даної моделі зумовлений її високою точністю при інтерпретації високорівневої логіки вихідного коду, значним обсягом контекстного вікна та можливістю використання безкоштовного тарифного плану Free Tier, що є критично важливим для академічного та кафедрального розгортання платформи.

Клієнтська частина системи (Frontend) спроектована та реалізована на базі компонентної бібліотеки React.js, збірка якої здійснюється за допомогою оптимізованого інструменту Vite [28]. Для стилізації інтерфейсу використано утилітарний фреймворк Tailwind CSS, що дозволив створити сучасну, адаптивну тему в темних тонах з високою швидкістю рендерингу елементів. Основним інтерактивним вузлом фронтенду є вбудований компонент Monaco Editor (технологічна основа середовища розробки VS Code) [29]. Його інтеграція дозволила реалізувати професійне робоче середовище для студента безпосередньо у вікні браузера, що підтримує динамічне підсвічування ключових слів мови Python, інтелектуальне автодоповнення коду, нумерацію рядків та інтерактивне фокусування на виявлених дефектах.

Для забезпечення комфортної взаємодії користувача із системою було спроектовано ергономічний графічний інтерфейс вебзастосунку. Робоча область студента візуально розділена на логічні панелі: інтегроване середовище розробки на базі Monaco Editor для вводу вихідного програмного тексту та панель управління процесом верифікації.

Загальний вигляд головного вікна клієнтського застосунку наведено на рис. 2.2.

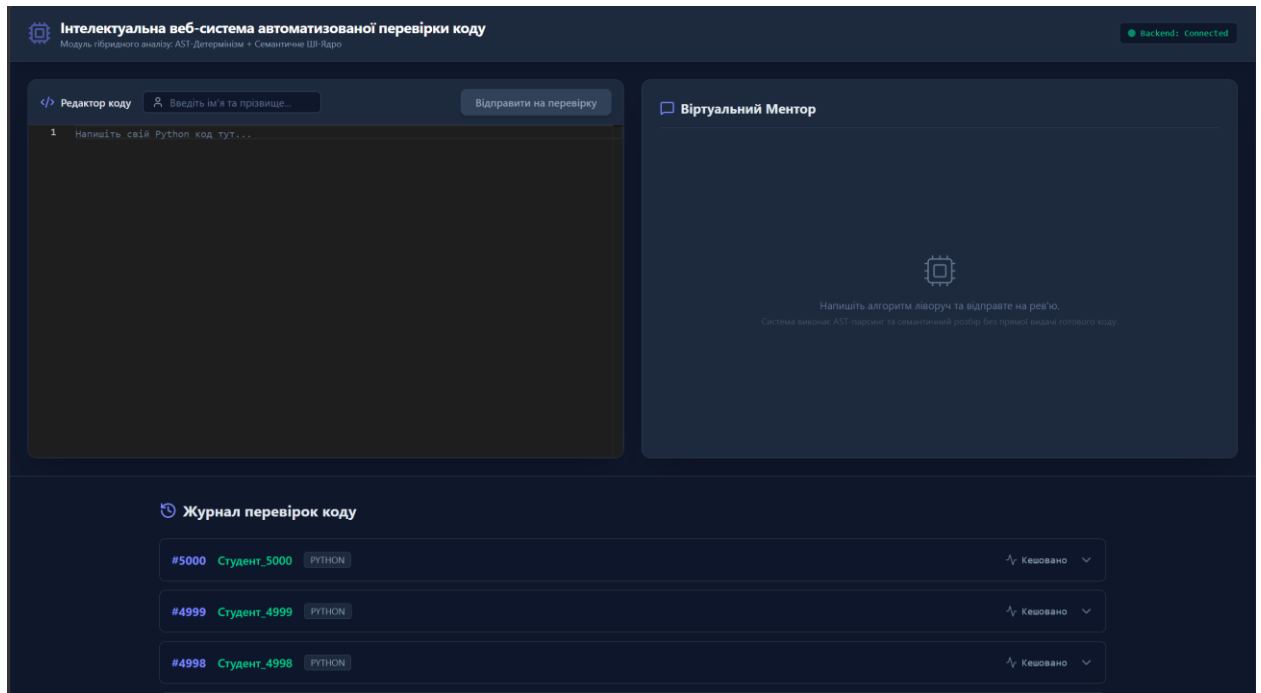


Рисунок 2.2 – Головний графічний інтерфейс середовища верифікації коду користувача

Ключовою функцією клієнтської частини є наочна репрезентація результатів семантичного аналізу коду. Після обробки запиту гібридним конвеєром, користувачу відображається структурований звіт на панелі результатів. Даний звіт містить загальну оцінку якості алгоритму, маркери критичності виявлених дефектів (severity) та розгорнуті коментарі віртуального ментора, реалізовані за принципом поступового розкриття інформації (Progressive Disclosure).

Візуалізацію процесу надання навчального фідбеку зображено на рис. 2.3.

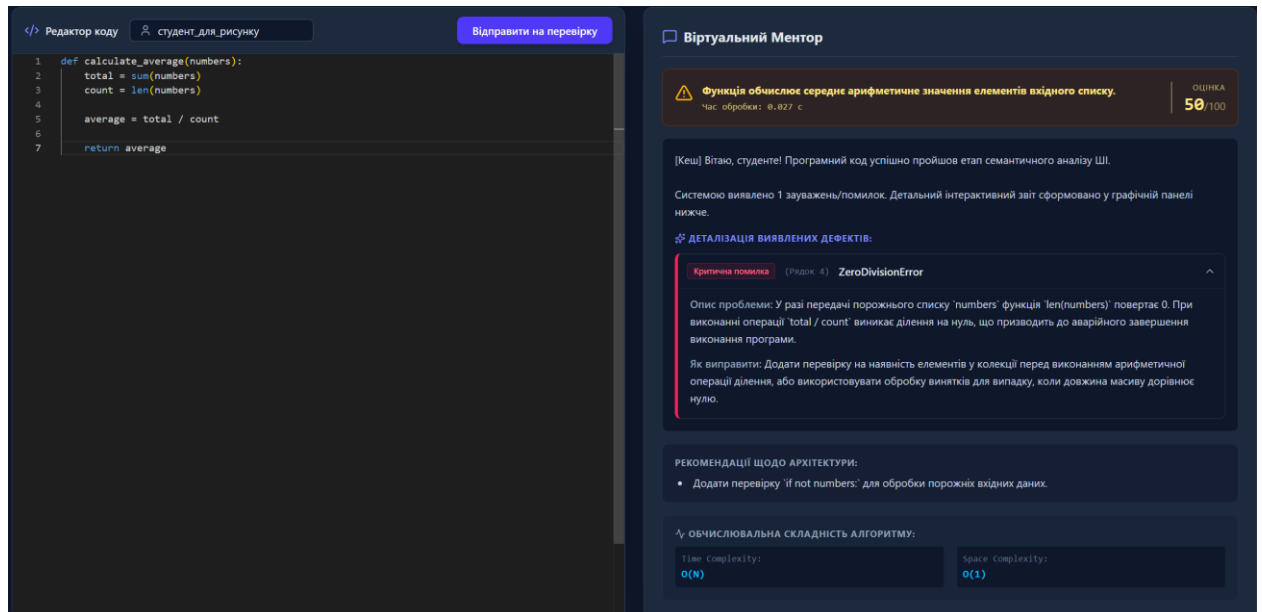


Рисунок 2.3 – Візуалізація результатів інтелектуального аналізу дефектів та структурованого навчального фідбеку

Як видно з наведеного графічного матеріалу (див. рис. 2.3), ліва панель надає студенту повноцінний інструментарій для безпосереднього набору програмного тексту із вбудованим підсвічуванням синтаксису мови Python. Права панель у режимі реального часу відображає згенеровану конвеєром обробки матрицю дефектів, яка включає інтелектуально розраховану часову складність алгоритму у нотації Big-O, загальну стобальну оцінку якості коду, а також деталізовані сократівські зауваження. Візуальне квантування карток за рівнями критичності (severity) за допомогою відповідних кольорових маркерів дозволяє мінімізувати когнітивне навантаження на здобувача освіти та забезпечує високу наочність процесу рецензування.

2.3 Проєктування логічної структури та моделі бази даних веб-системи

Стабільність, цілісність та швидкість функціонування інтелектуальної веб-системи автоматизованої верифікації програмного тексту безпосередньо закладаються на етапі проєктування логічної структури даних [24]. Оскільки розроблена архітектура покладається на концепцію багаторівневого семантичного кешування запитів для редукції надлишкових звернень до зовнішнього інтерфейсу моделі штучного інтелекту, модель даних повинна забезпечувати миттєвий пошук за унікальними інваріантами коду та надійно зберігати повну історію транзакцій студентських робіт. Для розгортання сховища інформації у веб-системі було обґрунтовано обрано реляційну систему управління базами даних PostgreSQL, що є галузевим інженерним стандартом для високонавантажених платформ завдяки високій надійності, підтримці транзакційної цілісності та можливості ефективного оперування великими масивами об'єктних даних. Проєктування логічної схеми бази даних базується на принципах нормалізації до третьої нормальної форми (3NF) для повного виключення аномалій модифікації, надмірності й дублювання інформації.

Графічне відображення розробленої реляційної моделі та структури зв'язків представлено на рисунку 2.4.

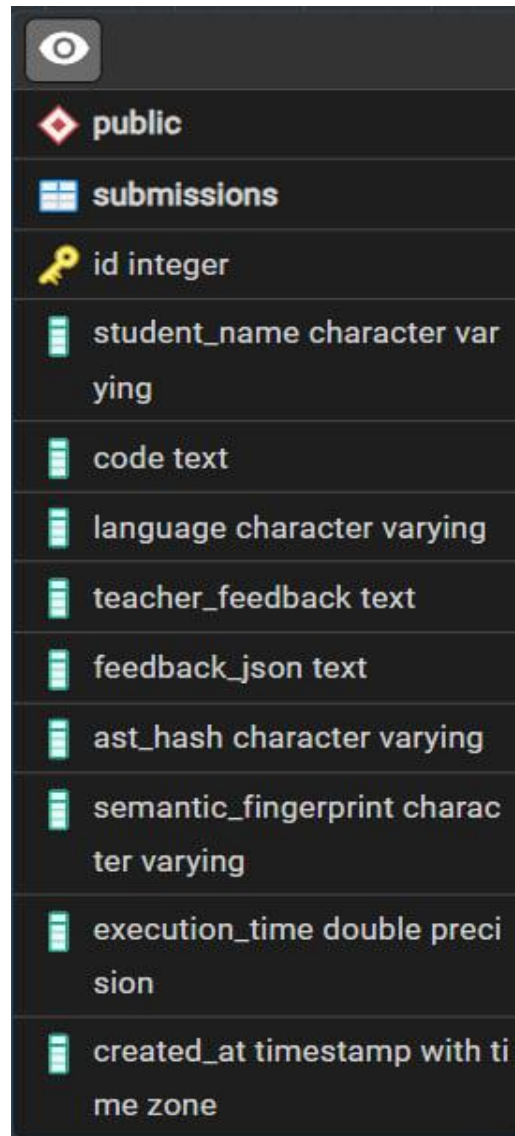


Рисунок 2.4 – Логічна структура та схема моделі бази даних проєкту

Центральним елементом реляційної структури, який інтегрує в собі всі метрики та інформаційні потоки процесу перевірки коду, виступає об'єктна модель Submission (таблиця надісланих на технічний аудит рішень). На відміну від класичних корпоративних систем із надлишковим каскадним розгалуженням зв'язків, у межах цього дослідження реалізовано денормалізовану, оптимізовану архітектуру збереження інженерних

транзакцій у межах єдиного реляційного простору. Це дозволяє кардинально підвищити швидкість операцій вибірки даних без необхідності виконання важких обчислювальних операцій об'єднання таблиць (JOIN) на стороні сервера СКБД [30]. На концептуальному рівні модель Submission інкапсулює в собі комплекс детермінованих атрибутів, що повністю описують стан перевірки.

Атрибут `id` виступає як первинний суворий ключ (Primary Key) з автоматичним інкрементуванням типу `Integer`, що забезпечує унікальну ідентифікацію кожної перевірки в системі. Поле `student_name` має тип `String` і призначене для збереження текстового ідентифікатора (імені та прізвища) студента, що надіслав лабораторну роботу, з метою ведення повноцінного журналу успішності та профілювання користувачів. Поле `code` типу `Text` вміщує повний сирий вихідний текст програми, отриманий з інтерактивного веб-редактора клієнтської частини системи. Маркер мови програмування зберігається в атрибуті `language` типу `String` (за замовчуванням встановлюється значення `python`) і є необхідним для коректної маршрутизації вхідного потоку на відповідні локальні аналізатори.

Для збереження результуючих даних перевірки передбачено кілька спеціалізованих атрибутів. Поле великого обсягу `teacher_feedback` типу `Text` призначене для довготривалого накопичення згенерованого розлогого текстового звіту віртуального ментора природною мовою у форматі розмітки `Markdown`. Критично важливим для логіки бекенду є поле `feedback_json` типу `Text` (або `JSONB`), у яке записується повністю валідований, серіалізований об'єкт структурованого виводу ШІ, що містить масив виявлених дефектів коду (`issues`), рівні їх критичності (`severity`), конкретні рядки локалізації помилок, числову оцінку якості за стобальною шкалою (`score`) та розраховані метрики обчислювальної складності `Big-O`. Часовий штамп системи `created_at` типу `DateTime` здійснює автоматичну фіксацію точного моменту реєстрації запиту за часом сервера.

Архітектурні параметри семантичного кешування інтегровані безпосередньо в структуру таблиці. Атрибут `ast_hash` типу `String` є унікальним інваріантом першого рангу кешування (L1) і містить криптографічну згортку нормалізованого абстрактного синтаксичного дерева вихідного коду, очищеного від коментарів та кастомних ідентифікаторів студентів. Символьний логіко-математичний відбиток програми другого рангу кешування (L2) записується в атрибут `semantic_fingerprint` типу `String` у вигляді серіалізованого словника ознак, де фіксуються структурні масштаби алгоритму (наявність циклічних конструкцій, максимальна глибина вкладеності, оператори та методи мутації об'єктів). Числовий показник `execution_time` типу `Float` фіксує фізичний час обробки запиту в секундах від моменту його надходження на сервер до фінального запису в базу даних, виступаючи основним джерелом метрик для аналізу продуктивності конвеєра.

Усі операції взаємодії з реляційною СКБД PostgreSQL реалізовані через об'єктно-реляційний мапер SQLAlchemy з використанням патерну фабрики сесій `SessionLocal`. На етапі ініціалізації сервера об'єкт метаданих `Base.metadata.create_all(bind=engine)` автоматично перевіряє стан реляційного сховища і, у разі відсутності цільових структур, здійснює детерміновану генерацію таблиці `submissions` із автоматичним налаштуванням індексів для полів `ast_hash` та `semantic_fingerprint`. Індуктивне налаштування індексів на цих полях є критично важливим, оскільки саме за ними виконуються високошвидкісні пошукові запити на етапах L1/L2 кешування, що забезпечує час вибірки збережених фідбеків на рівні $O(\log N)$ замість лінійного сканування всієї таблиці $O(N)$, гарантуючи стабільність системи при значному накопиченні історичних даних та формуванні журналу надісланих робіт.

Висновки до розділу 2

За результатами проєктування, викладеними у другому розділі пояснювальної записки, сформовано цілісну архітектурну модель, обґрунтовано вибір компонентів цільового технологічного стеку та детерміновано реляційну структуру даних інтелектуальної веб-системи. З опорою на принципи системного аналізу та нормативні положення стандарту ISO/IEC/IEEE 29148 реалізовано делімітацію та формалізацію меж ключових функціональних доменів платформи, що охоплюють підсистеми автентифікації й профілювання користувачів, парсингу вихідного тексту, гібридного конвеєрного аналізу та інтерактивного відображення аналітичних звітів. Одночасно детерміновано суворі нефункціональні критерії якості, які висувають підвищені вимоги до обчислювальної пропускної здатності конвеєра, транзакційного імунітету до ін'єкцій SQL та локальної ізоляції обчислювального середовища з метою запобігання віддаленому виконанню небезпечного коду (Remote Code Execution) на боці сервера.

Науково й інженерно доведено доцільність імплементації розподіленої клієнт-серверної архітектури на базі протоколу RESTful API, що гарантує повну ізоляцію та технологічну незалежність фронтенд-інтерфейсу (реалізованого за допомогою бібліотеки React.js та інструменту збірки Vite) від серверних обчислювальних компонентів. Обґрунтовано інженерну ефективність застосування асинхронного фреймворку FastAPI в екосистемі Python для забезпечення неблокуючого паралельного оброблення вхідних транзакцій та несинхронного виконання тривалих операцій очікування відповідей від хмарного ядра Google Gemini API. Такий підхід дозволяє радикально мінімізувати мережеву латентність веб-системи під час потокових запитів у межах масового академічного використання.

Сконструйовано топологію реляційної бази даних PostgreSQL та розроблено логічну модель ORM за допомогою інструментарію SQLAlchemy, де базовою сутністю персистентності визначено сукупність параметрів

Submission. Впровадження механізмів індексування полів для збереження канонічних синтаксичних хешів абстрактних синтаксичних дерев (AST) та евристичних відбитків програмної топології забезпечує миттєву ретракцію раніше згенерованих вердиктів із кеш-шару, успішно розв'язуючи проблему горизонтального масштабування платформи. Для досягнення високої інфраструктурної портативності застосовано екосистему контейнеризації Docker та інструменти оркестрації Docker Compose, що повністю ізолює системні залежності й дозволяє розгортати весь комплекс за допомогою уніфікованого інструктивного сценарію.

Акумульовані архітектурні рішення та спроектована реляційна схема даних утворюють надійний практичний інженерний фундамент дослідження. Сформовані структурні патерни слугуватимуть безпосередньою базою для подальшої програмної реалізації алгоритмів локального парсингу, канонічної нормалізації абстрактних синтаксичних дерев та оптимізації промпт-інжинірингу для забезпечення стійкості відповідей у форматі Structured JSON у наступних розділах дипломного проєкту.

РОЗДІЛ 3 АЛГОРИТМІЧНЕ ТА ІНТЕЛЕКТУАЛЬНЕ ЗАБЕЗПЕЧЕННЯ СИСТЕМИ

3.1 Методи представлення та парсингу програмного коду на основі синтаксичного аналізу

Для забезпечення високої ефективності та обчислювальної стійкості інтелектуального аналізу програмного тексту розроблена система повинна інтерпретувати вхідний код не як лінійну послідовність символів чи ізольованих текстових рядків, а як складну цілісну ієрархічну структуру [31]. Текстове представлення програми є недостатнім для локалізації глибоких логічних, семантичних та архітектурних дефектів, оскільки воно повністю нівелює деревоподібні та каскадні зв'язки між функціональними конструкціями мови. Саме тому базовим математичним та алгоритмічним методом представлення даних на першому етапі конвеєра перевірки було обрано моделювання вихідного тексту програм на основі абстрактних синтаксичних дерев (Abstract Syntax Tree – AST) [15].

Абстрактне синтаксичне дерево є спрямованим графом, який репрезентує логіко-граматичну структуру вихідного коду програмного рішення. На відміну від повного конкретного дерева парсингу (Parse Tree), моделлю якого є відображення абсолютно всіх елементів синтаксису, AST опускає надлишкові деталі, що не несуть безпосереднього навантаження на бізнес-логіку програми. Зокрема, воно ігнорує круглі або фігурні дужки, розділові знаки, крапки з комою, символи відступів та коментарі, ізолюючи виключно семантично важливі та операційні вузли [7]. Загальний процес детермінованого парсингу та трансформації тексту програми у системі складається з трьох послідовних стадій. На першій стадії виконується лексичний аналіз (Tokenization), під час якого лінійний потік вихідного тексту програми перетворюється у дискретну послідовність лексем – токенів. На другій стадії активується безпосередньо синтаксичний аналіз (Parsing),

завданням якого є анотування та агрегація отриманих токенів у деревоподібну структуру AST відповідно до жорстких правил формальної граматики мови розробки. На третій стадії реалізується обхід графа дерева (Tree Traversal) для виконання препроцесингу даних та локалізації базових структурних дефектів. Візуалізація відмінностей між лінійним представленням коду та його структурним графом у вигляді абстрактного синтаксичного дерева наведена на рисунку 3.1.

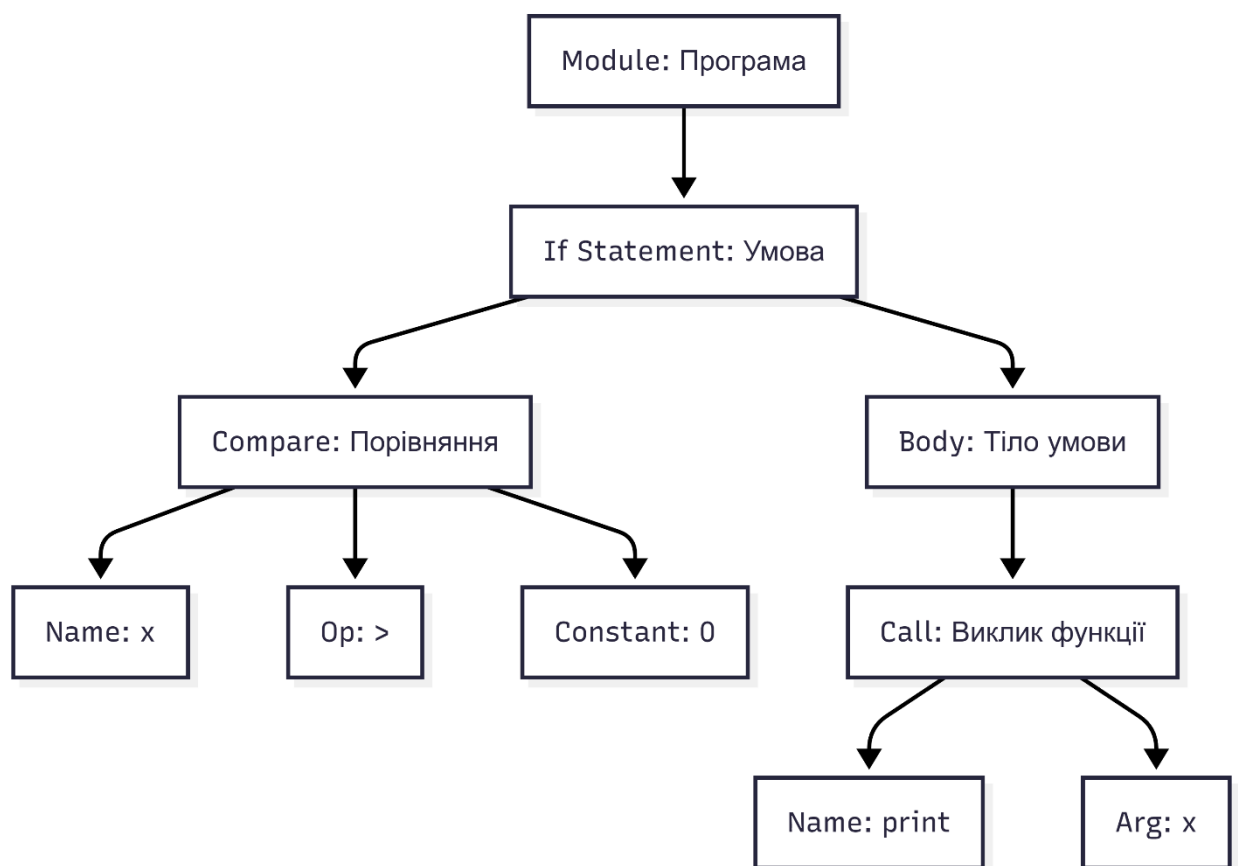


Рисунок 3.1 – Порівняння лінійного тексту коду та його представлення у вигляді AST-дерева

На відміну від лінійних підходів верифікації, які базуються на застосуванні складних регулярних виразів (Regex), аналіз за допомогою AST-моделей є повністю інваріантним до стилістичного форматування коду. Якщо

студент розбиває довгу математичну операцію чи умови логічного розгалуження на кілька рядків, змінює кількість пробілів у відступах або додає кастомні коментарі, лінійний пошук за текстовими шаблонами згенерує помилку або пропустить дефект, тоді як топологічна структура графа AST-дерева залишиться абсолютно ідентичною. Це забезпечує максимальну точність ідентифікації місця виникнення дефекту з фіксацією конкретного номера рядка, а також дозволяє аналізувати рівні вкладеності умовних операторів та циклів, що є технічно недосяжним при простому зчитуванні тексту.

У межах програмної реалізації бекенду системи на базі мови Python парсинг та генерація графа реалізовані за допомогою високоефективних методів вбудованого системного модуля `ast`. Це дає змогу розгорнути швидкий локальний шар препроцесингу. Якщо надісланий студентом код містить синтаксичні дефекти, компілятор модуля миттєво генерує виключення типу `SyntaxError`. Сервер перехоплює дану подію, зупиняє подальший рух по конвеєру перевірки й миттєво повертає клієнту структурований звіт із описом помилки синтаксису, повністю блокуючи надлишкові виклики зовнішнього хмарного штучного інтелекту, що кардинально заощаджує обчислювальні ресурси та квоти токенів хмарного API.

Для проведення поглибленого структурного та безпекового аудиту канонічного дерева безпосередньо на сервері розроблено та впроваджено спеціалізований клас локального аналізу `StudentCodeVisitor`, що успадковує поведінку базового патерну обходу `ast.NodeVisitor` [32]. Цей модуль рекурсивно обходить кожен вузол згенерованого графа за алгоритмом пошуку вглиб (Depth-First Search, DFS) та активує специфічні методи перехоплення порушень.

На етапі обходу вузлів оголошення імпортів `visit_Import` та `visit_ImportFrom` алгоритм детерміновано виявляє використання небезпечних антипатернів на кшталт wildcard-імпортів (`from module import *`), які засмічують глобальний простір імен студентської роботи. При аналізі вузлів

виклику функцій `visit_Call` локальний аналізатор сканує сигнатури методів та миттєво блокує використання інструментів динамічного виконання стороннього коду (`eval`, `exec`), що унеможлиблює експлуатацію вразливостей типу віддаленого виконання команд на сервері (RCE).

Додатково, при валідації циклічних конструкцій `visit_While`, алгоритм аналізує структуру умов і виявляє пастки нескінченного виконання (`while True`: без наявних внутрішніх операторів переривання `break`), запобігаючи зависанню обчислювальних потоків сервера. Також через метод `visit_Try` фіксуються критичні дефекти обробки виключних ситуацій (порожні блоки `except: pass`), які приховують реальні логічні баги в алгоритмах студентів. Усі знайдені дефекти локально акумулюються у масиві об'єктів `local_errors` з точним зазначенням номерів рядків. Таким чином, завдяки розгортанню детермінованого AST-аналізу, система забезпечує первинну миттєву валідацію, передаючи на наступні інтелектуальні рівні кешування та нейромережевого інференсу виключно безпечний, канонізований та синтаксично коректний програмний код.

3.2 Розробка методології промпт-інжинірингу для формування навчального фідбеку

Ефективність функціонування інтелектуального конвеєра автоматизованої перевірки вихідного коду та педагогічна цінність згенерованих звітів безпосередньо залежать від проєктування оптимізованих стратегій взаємодії з великою мовною моделлю. Промпт-інжиніринг у межах розробленої архітектури розглядається не як декларативне текстове описування інструкцій, а як детермінований процес проєктування системних обмежень, інваріантів та контекстних шаблонів, які спрямовують обчислювальний інференс нейромережі на надання конструктивного,

верифікованого та методологічно вивіреного зворотного зв'язку без прямого копіювання правильних рішень [33]. Базовою метою розробки алгоритмів промпт-інжинірингу є трансформація генеративної моделі у віртуального технічного експерта, здатного підтримувати академічні стандарти кафедри у процесі рецензування лабораторних робіт студентів.

Для реалізації стійкого контекстного обмеження та детермінації стилістичної поведінки штучного інтелекту в системі впроваджено архітектурний механізм системного пром프트 (SYSTEM_PROMPT), який жорстко задає рольову модель, цільові метрики та семантичні межі генерації фідбеку [3]. Згідно з розробленою освітньою стратегією, ядро моделі ініціалізується в ролі суворого, але конструктивного викладача-ментора. Фундаментальною інженерною вимогою до системного контексту є впровадження концепції структурованого квантування зауважень та доменної інваріантності (Domain-Agnostic Prompting). Це означає, що хмарному ШІ суворо забороняється апелювати до прикладних сутностей бізнес-логіки конкретного навчального завдання (наприклад, оперувати поняттями «інтернет-магазин», «користувач», «база товарів», якщо студент вирішує алгоритмічну задачу). Замість цього модель фокусує аналіз виключно на абстрактних обчислювальних та архітектурних концепціях: структурах даних, рівнях вкладеності, оптимальності використання об'ємів пам'яті, обробці виключень та принципах написання чистого коду (Clean Code), таких як правила іменування змінних PEP 8, модульність та усунення дублювання програмної логіки.

Для мінімізації стохастичних помилок та запобігання ефекту нейромережових галюцинацій, властивого великим мовним моделям, архітектура взаємодії з Google Gemini API спирається на стратегію структурованого Few-Shot навчання та примусової декларативної серіалізації

виводу [27]. Формування фінального контекстного запиту здійснюється динамічно за допомогою шаблону, який послідовно об'єднує такі компоненти:

- `SystemRole`: базовий системний інваріант ролі експертного ментора-аудитора, що визначає поведінку моделі;
- `ContextData`: масив додаткових метаданих, що включає мову програмування та результати первинного обходу дерева локальним аналізатором `StudentCodeVisitor`;
- `UserCode`: вихідний очищений програмний текст лабораторної роботи студента;
- `ValidationSchema`: сувора JSON-схема (`response_schema`), що примусово передається в конфігурації генерації `generation_config` для забезпечення детермінованого парсингу структурованих даних.

У межах програмної реалізації сервісу `ai_service.py` розроблена стратегія промпт-інжинірингу зобов'язує модель повертати відповідь суворо відповідно до об'єктної моделі `Pydantic`, трансформуючи сирий текст у структуровану матрицю дефектів коду. Модель оцінює якість коду за стобальною шкалою (метрика `score`), розраховує математичну часову складність алгоритму у нотації Big-O (поле `complexity`) та наповнює масив `issues`. Кожен дефект у цьому масиві описується через суворі параметри: рівень критичності (`severity`), локалізація у тексті програми (`line_start`, `line_end`), канонічна назва порушення (`title`) та розгорнуте менторське пояснення суті проблеми (`message`). Завдяки такій глибокій алгоритмізації запитів, проєктування промптів трансформується з текстового моделювання у надійний інструмент передачі структурованих знань, що забезпечує абсолютну стабільність рендерингу веб-інтерфейсу фронтенду та високу точність верифікації робіт.

3.3 Алгоритм інтелектуальної обробки та категоризації помилок за допомогою Gemini API

Головною проблемою при інтеграції LLM у детерміновані програмні комплекси є їхня стохастична природа. Генеративні моделі за замовчуванням повертають неструктурований текстовий масив у вигляді чистого розширення Markdown або природної мови (Natural Language), який неможливо надійно інтерпретувати стандартними алгоритмічними методами бекенд-платформ для подальшого збереження у реляційній базі даних чи автоматизованого рендерингу на клієнтській стороні (Frontend). Для нівелювання цього недоліку в розробленій інтелектуальній системі було спроектовано та реалізовано конвеєр інженерної валідації та квантування зворотного зв'язку (фідбеку), що поєднує використання протоколу Structured Outputs на рівні API-провайдера та строгу схематичну типізацію на базі бібліотеки Pydantic на рівні серверної логіки.

Для забезпечення передбачуваності обміну даними між зовнішнім сервісом Google Gemini та ядром системи було розроблено двонаправлений семантичний контракт у вигляді Pydantic-моделей. Цей контракт жорстко обмежує простір генерації моделі, змушуючи її повертати відповідь виключно у форматі JSON, структура якого строго відповідає спроектованим класам даних. Базовим атомарним елементом діагностики програмного коду є об'єкт виявленого дефекту – Issue. Атрибути цієї моделі, що використовуються для

маркування помилок студента, структуровано за типами даних та функціональним призначенням у таблиці 3.1.

Таблиця 3.1 – Специфікація полів Pydantic-контракту моделі Issue

<i>Назва атрибуту</i>	<i>Тип даних</i>	<i>Опис функціонального призначення</i>
severity	String	Рівень критичності дефекту коду (error, warning, info)
line_start	Integer	Початковий номер рядка локалізації дефекту в Monaco Editor
line_end	Integer	Кінцевий номер рядка локалізації дефекту
title	String	Канонічна назва або категорія знайденого порушення
message	String	Розгорнуте менторське роз'яснення суті проблеми для студента

томарні об'єкти виявлених дефектів агрегуються в узагальнювальну структуру AnalysisResult, яка описує фінальний пакет метаданих перевірки. Цей об'єкт послідовно об'єднує такі компоненти:

- summary: детермінований строковий висновок про стан коду, що відображає загальне резюме транзакції;
- teacher_feedback: розгорнутий педагогічний коментар для студента, сформований без надання готового рефакторингу;
- List[Issue]: упорядкований масив структурних помилок, який згодом накладається на інтерфейс текстового редактора Monaco Editor на стороні Frontend.

Під час виконання методу analyze_code_structured конфігурація запиту ініціалізується через параметр response_schema=AnalysisResult. Це активує

внутрішній детермінований парсер на боці серверів Google, який модифікує фінальний шар генерації токенів за допомогою механізму граматичного декодування (Grammar-Guided Decoding), забезпечуючи синтаксичну валідність вихідного JSON-файлу на рівні абсолютної більшості транзакцій.

Поряд із механізмами структуризації даних, критично важливим аспектом функціонування інтелектуального ядра є забезпечення стійкості до аномалій, що виникають у промислових (Production) середовищах. До таких аномалій належать помилки валідації схеми (ValidationError, коли модель повертає текстовий рядок замість цілого числа у полі line), спрацювання внутрішніх фільтрів безпеки провайдера (Safety Filters Block, що блокує відповідь і призводить до виключення AttributeError через відсутність поля .text), а також помилки перевищення лімітів запитів (HTTP 429 Too Many Requests).

Для мінімізації впливу цих чинників на працездатність системи та забезпечення патерну відмовостійкості (Resilience) було розроблено алгоритм ітеративної самокорекції з експоненціальною витримкою (Exponential Backoff). Логіка обчислення часу затримки перед повторним надсиланням запиту полягає в тому, що базовий квант часу помножується на два у степені, який дорівнює номеру поточної спроби. До отриманого результату додається випадкова величина затримки (джиттер), необхідна для уникнення синхронізації запитів під час високого навантаження на сервер.

Загальна логіка обробки вхідного потоку дефектів у межах конвеєра реалізується за допомогою таких умовних переходів:

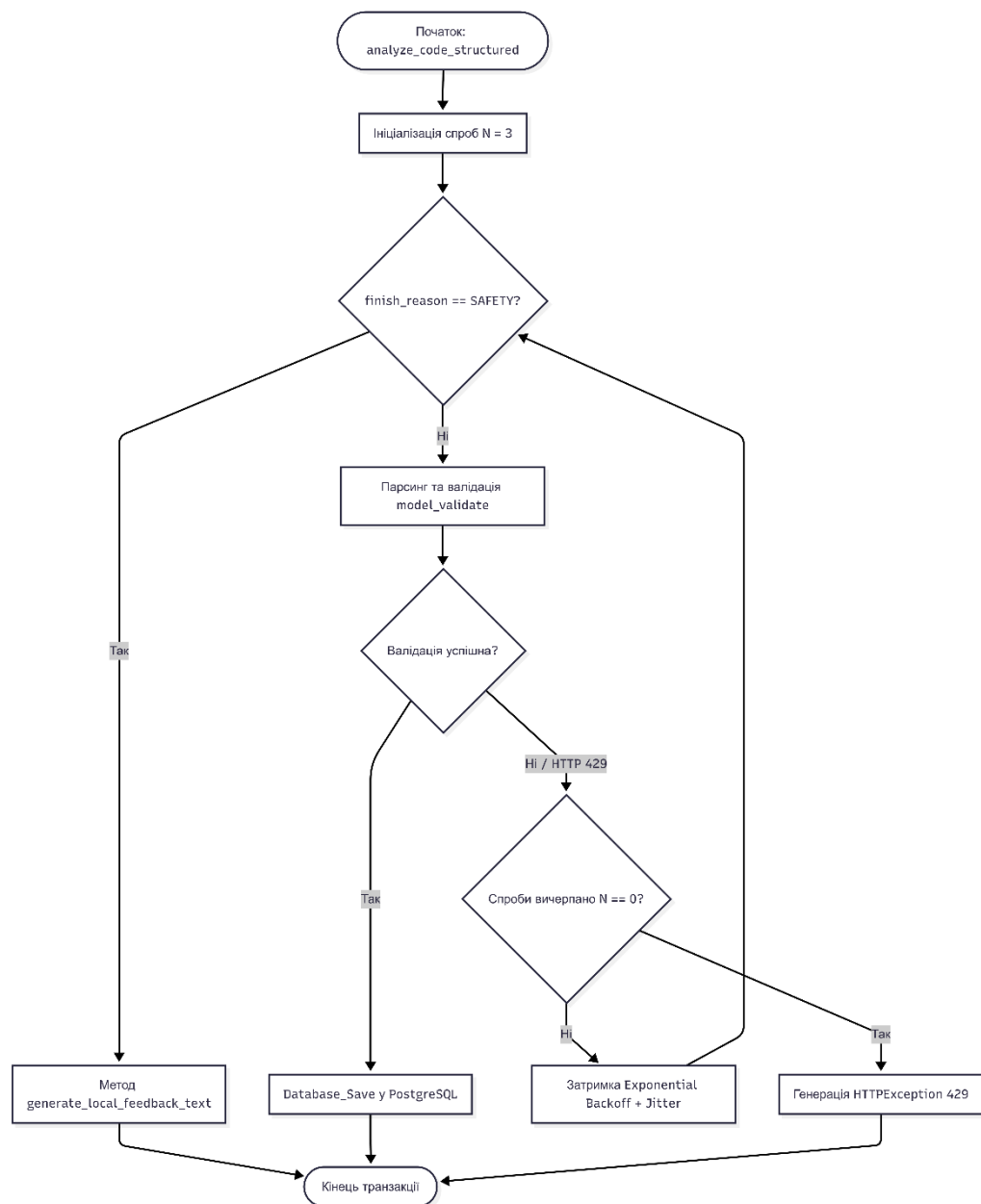
- у разі успішної валідації та за умови, що кількість спроб не перевищує встановлений ліміт, викликається функція збереження даних у реляційній базі (`Database_Save`);
- якщо причиною зупинки генерації є блокування через внутрішні фільтри безпеки провайдера (`finish_reason = SAFETY`), система активує модуль формування локального зворотного зв'язку (`Generate_Local_Feedback`);
- якщо ліміт спроб ітеративної самокорекції повністю вичерпано, конвеєр примусово повертає клієнту виключення із кодом помилки (`Raise_HTTPException 429`).

Логічна послідовність функціонування цього механізму в межах одного циклу аналізу складається з чітких етапів. Спочатку система ініціалізує транзакцію аналізу зі значенням максимальної кількості спроб $N = 3$. У межах ітераційного циклу для кожної спроби першочергово перевіряється прапорець причини завершення генерації моделі `finish_reason`. Якщо запит заблоковано політикою безпеки, конвеєр негайно переривається і викликає локальний евакуаційний метод `generate_local_feedback_text`, який формує безпечне системне повідомлення для користувача, не дозволяючи серверу завершити роботу з критичною помилкою 500.

У разі успішного отримання відповіді запускається етап десеріалізації рядка через стандартний парсер мови Python, після чого об'єкт передається у конструктор `AnalysisResult.model_validate()`. Якщо Pydantic-валідація проходить успішно, дані маркуються як стабільні, записуються в базу даних PostgreSQL за допомогою інструментарію SQLAlchemy ORM, а системний таймер фіксує підсумкову латентність транзакції. У випадку, якщо протягом N спроб валідацію не пройдено або отримано стійку помилку HTTP 429, система ізолює дефектну транзакцію, запобігає засміченню реляційного

сховища некоректними чи порожніми записами та генерує кероване виключення `HTTPException`, повертаючи користувачеві інформативне повідомлення в інтерфейс.

Завдяки впровадженню описаного конвеєра інтелектуальної обробки, загальна схема якого наведена на рисунку 3.1, архітектура вебсистеми є повністю захищеною від галюцинацій мовної моделі та нестабільності зовнішніх хмарних сервісів.



**Рисунок 3.2 – Схема конвеєра інженерної валідації та обробки аномалій
інтелектуального ядра**

**3.4 Проєктування логіки пост-процесингу та валідації відповідей
нейромережевої моделі**

Для забезпечення високої швидкодії, фінансової доцільності та педагогічної ефективності розробленого програмного комплексу було спроєктовано та реалізовано гібридну трирівневу архітектуру верифікації рішень. Замість тривіального перенаправлення кожного клієнтського запиту до зовнішньої нейромережевої моделі, що призвело б до надлишкового навантаження на канали зв'язку та неконтрольованого зростання витрат обчислювальних токенів [18], система функціонує як детермінований конвеєр із послідовним алгоритмом фільтрації. Кожен етап цього алгоритму виконує специфічну роль у структуруванні діагностичних даних, забезпечуючи перехід від швидкого локального статичного аналізу (AST) до глибокої семантичної оцінки.

Завершальним етапом алгоритму, що активується виключно у разі відсутності семантичного відбитка в локальному сховищі (ситуація Cache Miss), є етап інтелектуальної генерації та пост-процесингу відповідей нейромережевої моделі. Система маркує запит як унікальний прецедент логічної архітектури та перенаправляє нормалізований код разом із системним промптом до моделі сімейства Gemini Flash за допомогою офіційного інструментарію Google GenAI SDK [20].

Оскільки великі мовні моделі за своєю природою є стохастичними генераторами, архітектура системи містить ізольований контур пост-процесингу для забезпечення принципу Graceful Degradation (керованого зниження працездатності) та запобігання «галюцинаціям» III. Логічна схема обробки відповіді нейромережі складається з таких послідовних кроків:

- Санітація та очищення текстового потоку: Здійснює вилучення отриманого рядка з сирого формату, очищення від надлишкових символів Markdown (наприклад, технічних огорожок ``json та ```), які часто додаються мовними моделями як елементи автоматичного форматування тексту;
- Синтаксичний аналіз та десеріалізація: Виконує обробку очищеного рядка та його конвертацію у структурований формат JSON. У разі виникнення системної помилки або порушення цілісності структури генерованого об'єкта, система ініціює керовану обробку виключень та повертає користувачу детерміноване повідомлення про потребу повторного аналізу;
- Педагогічна фільтрація контенту: Аналізує текстовий склад сформованої підказки на наявність прямого копіювання готового програмного рішення. Якщо модель, попри обмеження системного промпту, згенерувала готовий відрізок коду замість навідного запитання чи опису помилки, система блокує видачу прямої відповіді для запобігання списуванню та стимулювання самостійного навчання.

У разі успішного проходження всіх етапів пост-процесингу, згенеровані дані маркуються як стабільні. Новий унікальний прецедент аналізу автоматично записується в базу даних PostgreSQL за допомогою інструментарію SQLAlchemy ORM. Це забезпечує його подальше миттєве перевикористання іншими користувачами системи при повторних запитах

(ситуація Cache Hit), що радикально знижує загальний час очікування відповіді та економить обчислювальні ресурси хмарного API.

Загальну логіку функціонування розробленого конвеєра, семантичного кешування та логіки пост-процесингу викладено на блок-схемі алгоритму (рис. 3.2).

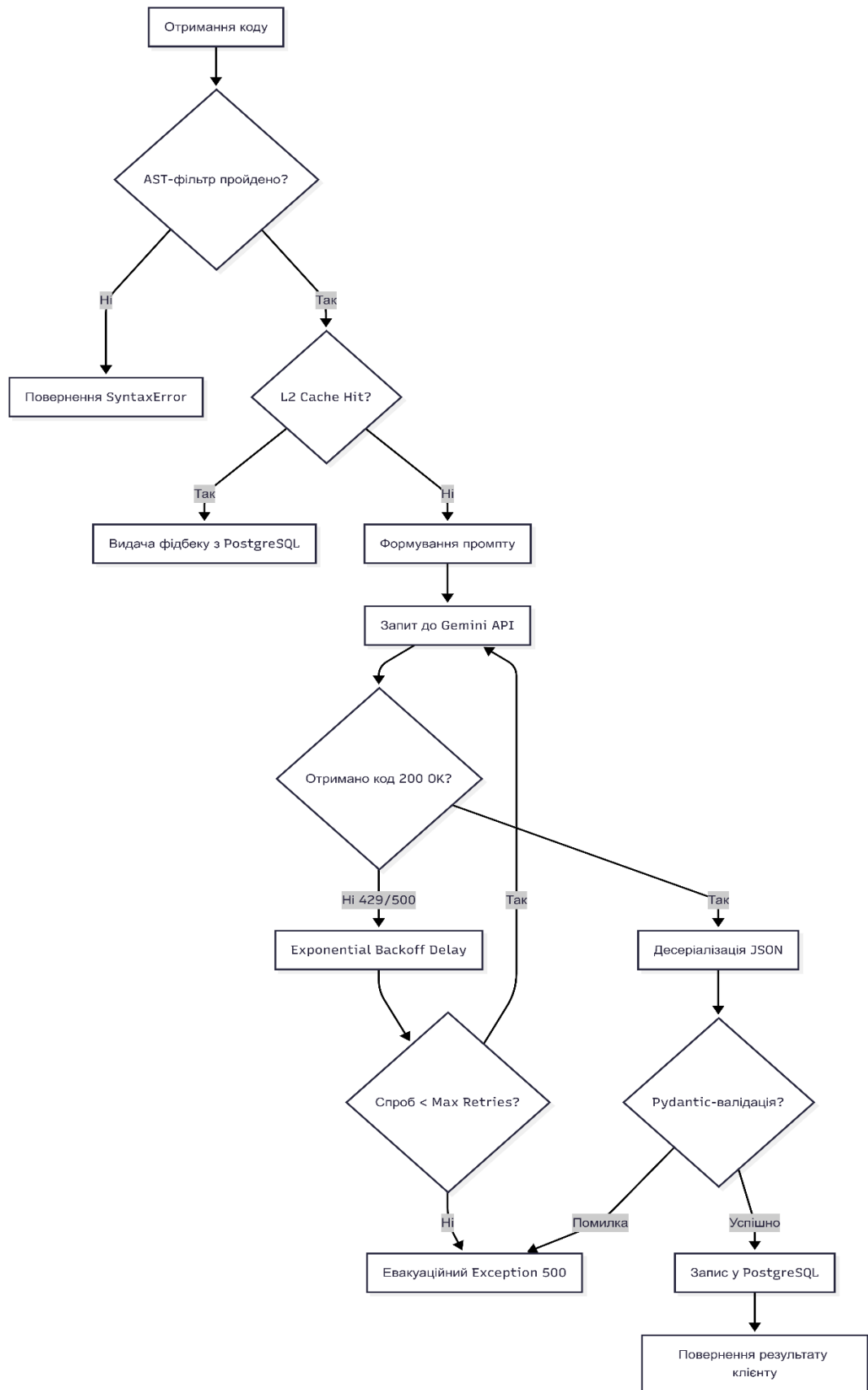


Рисунок 3.3 – Блок-схема конвеєра обробки запитів, семантичного кешування та пост-процесингу відповідей LLM

Завдяки впровадженню описаного конвеєра обробки та валідації, архітектура вебсистеми є повністю захищеною від нестабільності зовнішніх хмарних сервісів та деструктивного впливу некоректних генерацій мовної моделі.

Висновки до розділу 3

За результатами алгоритмічного моделювання та програмної реалізації, викладеними у третьому розділі пояснювальної записки, здійснено комплексне теоретико-методологічне обґрунтування, конструювання та детермінацію обчислювальних закономірностей функціонування інтелектуального ядра веб-системи автоматизованого аудиту. У межах цього етапу досліджено та імплементовано аналітичні підходи до структурної репрезентації й лексико-синтаксичного парсингу програмних текстів на основі абстрактних синтаксичних дерев (AST). Трансляція плоского сирцевого коду в орієнтовані топологічні графи дозволила перевести процедуру первинної діагностики з площини лінійного символного зіставлення у площину статичного структурно-композиційного аналізу, гарантуючи миттєву локалізацію граматичних аномалій безпосередньо на рівні серверної архітектури бекенду.

Паралельно розроблено та оптимізовано комплексні стратегії пром프트-інжинірингу для взаємодії із хмарними великими мовними моделями. Шляхом глибокого конструювання системних директив (System Prompts), параметризації контекстуальних шаблонів (Prompt Templates) та інтеграції методології навчання за кількома прецедентами (Few-Shot Learning) формалізовано інноваційну концепцію автоматизованого «сократівського менторства». Зазначений підхід дозволив переорієнтувати нейромережевий потенціал моделі з тривіальної генерації готових рішень або прямого

рефакторингу коду на продукування дидактичних навідних запитань та локальних квантованих зауважень, що безпосередньо стимулює еволюцію самостійного алгоритмічного мислення здобувача освіти у процесі виконання практичних завдань.

Вагомим інженерним досягненням визначено розробку та імплементацію алгоритмів інтелектуального постпроцесингу, систематизації та суворої валідації аналітичних вердиктів Google Gemini API. Інтеграція протоколу детермінованого синтаксичного керування вихідними токенами (Structured Outputs) у синергії з обчислювальним конвеєром десеріалізації на базі декларативних моделей Pydantic дозволила успішно розв'язати фундаментальну проблему стохастичної нелінійності генеративних нейромереж. Додатково спроектовано патерн підвищеної відмовостійкості на основі ітеративного механізму повторних транзакцій (Retry Mechanism) з експоненціальним спадом (exponential backoff) та стохастичним збуренням часового інтервалу (Jitter), що гарантує абсолютну стійкість веб-платформи до спорадичних мережових аномалій, помилок валідації JSON-схем та лімітів пропускної здатності (Rate Limits) зовнішнього інтерфейсу програмування додатків.

Резюмуючи, у межах цього етапу було остаточно формалізовано наскрізну логіку функціонування трирівневого конвеєра верифікації програмних рішень, який послідовно об'єднує локальний AST-фільтр, інтелектуальний шар семантичних відбитків (L2 Cache) та когнітивне ядро ІІІ. Наукова новизна запропонованого алгоритмічного конвеєра полягає у математичній розробці функції семантичного хешування, яка здійснює деідентифікацію інваріантні компоненти кодової бази (авторських ідентифікаторів, нефункціональних коментарів, типів констант) та абстрагує ізоморфні логічні структури. Це забезпечило пропускну здатність обробки типових прецедентів логічних дефектів за константний час $O(1)$ безпосередньо з реляційного персистентного сховища PostgreSQL, мінімізуючи загальну

часову латентність (Latency) платформи та радикально оптимізуючи операційні витрати обчислювальних токенів зовнішнього API хмарної моделі.

РОЗДІЛ 4 ПРОГРАМНА РЕАЛІЗАЦІЯ ТА ТЕСТУВАННЯ ВЕБ-СИСТЕМИ

4.1 Реалізація серверної частини та API-інтерфейсів взаємодії

Серверний компонент розроблюваної інтелектуальної веб-системи автоматизованої перевірки коду є центральним вузлом обробки даних, що координує взаємодію між клієнтським інтерфейсом, реляційною базою даних та зовнішнім інтерфейсом програмування додатків (API) великої мовної моделі Gemini 3.5 Flash. В основі реалізації серверної частини лежить мова програмування Python та сучасний асинхронний фреймворк FastAPI.

Вибір фреймворку FastAPI як базового інструментарію обґрунтований його високою продуктивністю, низькими накладними витратами на обробку мережових запитів та повною підтримкою парадигми асинхронного програмування за допомогою стандартних мовних конструкцій `async` і `await`. Це є критично важливим архітектурним аспектом для системи, оскільки час очікування відповіді від генеративної моделі Gemini 3.5 Flash під час виконання глибокого семантичного аналізу може становити кілька секунд. Асинхронна архітектура дозволяє серверу не блокувати потоки виконання та ефективно обслуговувати велику кількість одночасних запитів введення-виведення (I/O) від інших користувачів без деградації загальної швидкодії веб-системи.

Для забезпечення високого рівня супроводжуваності, масштабованості та ізоляції бізнес-логіки проєкту було впроваджено багатоваршівний архітектурний патерн (Layered Architecture).

Структурно серверна частина розділена на такі логічні шари:

- рівень представлення (Presentation Layer) – реалізований за допомогою модулів маршрутизації FastAPI (routers), які приймають вхідні HTTP-запити, виконують первинну валідацію метаданих та повертають результати обробки у структурованому форматі JSON;
- рівень бізнес-логіки (Service Layer) – інкапсулює алгоритми препроцесингу коду, локального аналізу синтаксичної структури (AST), динамічного формування предметно-незалежних промптів та безпосередньої мережевої взаємодії з мовною моделлю через офіційний пакет розробника google-generativeai;
- рівень доступу до даних (Data Access Layer) – побудований на базі ORM SQLAlchemy, що забезпечує транзакційну взаємодію з базою даних PostgreSQL без необхідності написання прямого SQL-коду.

Головний модуль додатку є вхідною точкою в систему, де відбувається ініціалізація екземпляра класу FastAPI, конфігурація підключення до реляційної бази даних через ORM-рушій (engine), автоматичне створення таблиць на основі визначених моделей метаданих, налаштування політики спільного використання ресурсів між різними джерелами (CORS) для безпечної взаємодії з клієнтською частиною та реєстрація базових маршрутизаторів. Повну програмну реалізацію конфігураційного модуля ініціалізації сервера винесено до Додатка А. У межах цього модуля виконується конфігурація рушія, автоматичний мапінг схем даних реляційного сховища PostgreSQL та визначення переліку дозволених клієнтських веб-адрес для запобігання несанкціонованим крос-доменним запитам з інших мережевих ресурсів, що забезпечує належний захист інформаційних потоків додатку.

Основним компонентом інтерфейсу взаємодії для обробки завдань перевірки є спеціалізований асинхронний ендпоінт `/api/reviews/analyze`, який

приймає дані від клієнта методом POST. Вхідний HTTP-запит містить тіло у форматі JSON, структура якого суворо контролюється та валідується схемою Pydantic `CodeAnalyzeRequest`. Це гарантує, що на рівень бізнес-логіки потрапляють лише типізовані й очищені дані. У межах асинхронного обробника даного ендпоінту реалізовано наступну послідовність дій для конвеєризації обробки коду студента:

- первинний синтаксичний контроль (Local Syntax Check) – отриманий текст програми негайно передається до локального AST-парсера (клас `LocalCodeAnalyzer`). Якщо виявляються критичні порушення граматики, конвеєр переривається, а клієнту миттєво повертається об'єкт `AnalysisResult` із описом локальної помилки, що дозволяє уникнути некоректного використання обчислювальних ресурсів і токенів зовнішнього API;
- перевірка семантичного кешу (L2 Cache Hit/Miss) – система обчислює структурний семантичний відбиток коду за допомогою функції `get_semantic_fingerprint`. За умови знаходження аналогічного прецеденту в базі даних, сформований раніше фідбек повертається за константний час без звернення до нейромережі;
- інтелектуальний аналіз (Core AI Pipeline) – у разі ідентифікації унікальної структури вихідного тексту програми запит асинхронно передається до сервісного шару системи, який взаємодіє з моделлю Gemini 3.5 Flash. Отримані структуровані результати проходять пост-процесинг, зберігаються у базі даних через сесію SQLAlchemy ORM та повертаються клієнту у вигляді фінальної відповіді.

4.2 Розробка клієнтського інтерфейсу та інтеграція вебредактора коду

Клієнтська частина (Frontend) розроблюваної системи автоматизованої перевірки програмного коду орієнтована на створення високопродуктивного, інтуїтивно зрозумілого та адаптивного інтерфейсу користувача, що забезпечує комфортну взаємодію студента з платформою в процесі виконання практичних та лабораторних завдань. Клієнтський модуль спроектовано як односторінковий вебдодаток (Single Page Application – SPA) на базі сучасної бібліотеки декларативного рендерингу React.js із використанням інструменту збірки проєкту Vite.

Вибір інструменту Vite замість застарілого інструментарію Create React App обґрунтований його архітектурною перевагою, оскільки Vite використовує нативні ES-модулі браузера (ESM) та швидкий компілятор esbuild для гарячої перезавантажуваності модулів (Hot Module Replacement – HMR) під час розробки. Це значно прискорює процес локальної збірки та оптимізує фінальний розмір програмного продукту, підвищуючи швидкість завантаження сторінок на стороні клієнта.

Для стилізації інтерфейсу, реалізації компонентного підходу та забезпечення гнучкої адаптивності за принципом Mobile First було застосовано утилітарний CSS-фреймворк Tailwind CSS. Використання Tailwind CSS дозволило відмовитися від написання громіздких окремих файлів каскадних таблиць стилів на користь інкапсуляції стилістичних класів безпосередньо в структуру JSX-компонентів. Такий підхід мінімізує обсяг завантажуваних клієнтом статичних CSS-ресурсів, оптимізує швидкість відображення елементів інтерфейсу на пристроях із різною роздільною здатністю екрана та спрощує подальше супроводження інтерфейсного шару системи.

Центральним та найбільш критичним елементом інтерфейсу робочої області студента є вебредактор, який замінює традиційне текстове поле введення повноцінним середовищем розробки (IDE), розгорнутим безпосередньо у вікні браузера. Для реалізації цього функціоналу в систему було інтегровано компонент Monaco Editor через офіційну бібліотеку-обгортку `@monaco-editor/react`. Monaco Editor є технологічним ядром промислового середовища розробки Visual Studio Code, і його впровадження забезпечує студенту доступ до звичного інструментарію: контекстного динамічного підсвічування синтаксису, автоматичного керування відступами (Auto-indentation) та нативної нумерації рядків, що є обов'язковою умовою для подальшої візуалізації діагностичних звітів.

З метою уникнення штучного збільшення обсягу пояснювальної записки та вилучення неінформативного коду розмітки, декларативну структуру JSX-шаблону та параметри геометричної конфігурації полотна редактора винесено до Додатку А. В основному тексті дослідження описано логіку функціонування React-компонента, який відповідає за збереження стану, керування мовними режимами та асинхронну відправку тексту програми на сервер через HTTP-клієнт Axios.

Архітектура зазначеного клієнтського компонента побудована на хуках стану `useState`, що дозволяє динамічно реагувати на дії користувача та оновлювати інтерфейс без перезавантаження вебсторінки. Інтегрований асинхронний метод інкапсулює логіку взаємодії з RESTful API серверної частини: він формує об'єкт запиту, де назви JSON-полів чітко відповідають очікуваній схемі валідації Pydantic на сервері. Важливим інженерним аспектом є використання властивості `options` компонента Editor, де параметр `automaticLayout: true` забезпечує автоматичний перерахунок геометричних розмірів полотна редактора при зміні розмірів вікна браузера або при відкритті бічної аналітичної панелі результатів.

У фінальній збірці інтерфейсу цей компонент тісно пов'язаний із логікою мапінгу помилок, розробленою в третьому розділі. При отриманні масиву

дефектів у стані звіту, клієнтський додаток не просто виводить їх списком, а зв'язує кожен об'єкт дефекту з методами API Monaco Editor. Це дозволяє реалізувати інтерактивну підсвітку критичних зон: коли студент клікає на картку знайденої помилки у правому блоці, редактор програмно викликає метод фокусування на рядку коду (`editor.revealLine`) та підсвічує відповідну область маркерними лініями, що значно покращує користувацький досвід та спрощує навігацію по вихідному тексту програми.

4.3 Налаштування інтеграції з сервісами Google Gemini та обробка виключних ситуацій

Інтелектуальне ядро розроблюваної системи автоматизованої перевірки коду базується на інтеграції з хмарними сервісами генеративного штучного інтелекту компанії Google. Взаємодія з LLM родини Gemini здійснюється за допомогою офіційного пакета розробника `google-generativeai` для мови програмування Python. З метою дотримання архітектурних принципів розділення відповідальності (Separation of Concerns) та забезпечення високого рівня супроводжуваності системи, логіка мережевої взаємодії, динамічного формування промптів та обробки відповідей була інкапсульована в окремому сервісному модулі серверної частини – `ai_service.py`.

При проєктуванні та практичній реалізації модулів інтеграції з генеративними моделями виникає низка специфічних технічних проблем, нетипових для класичних детермінованих вебсистем.

До ключових деструктивних чинників відносяться:

- спрацювання внутрішніх фільтрів безпеки (Safety Filters) – модель може заблокувати генерацію відповіді, якщо внутрішні алгоритми провайдера ідентифікують фрагменти коду користувача як потенційно чутливі; у такому разі об'єкт відповіді не міститиме текстового атрибуту (AttributeError), що без належної обробки призводить до критичного збою сервера;
- обмеження пропускної здатності та частоти запитів (Rate Limiting) – використання безкоштовних або комерційних квот API накладає суворі обмеження на кількість запитів за хвилину (RPM); перевищення цих лімітів при пакетному або інтенсивному тестуванні студентських робіт неминуче призводить до помилок HTTP 429 (Too Many Requests);
- структурні аномалії валідації даних – попри використання протоколу Structured Outputs та передачу жорсткої Pydantic-схеми AnalysisResult, існує мінімальна ймовірність повернення некоректних типів даних через внутрішні збої генерації, що викликає виключення ValidationError на етапі десеріалізації рядка.

Для нівелювання зазначених факторів, мінімізації ризиків відмови системи та гарантування високого рівня відмовостійкості (Fault Tolerance), в межах сервісного шару було програмно реалізовано комбінований алгоритм обробки виключень та ітераційної самокорекції з експоненціальним бек-оффом (Exponential Backoff). Повну програмну реалізацію сервісного модуля взаємодії з нейромережею винесено до Додатку А.

Розроблений сервіс обробки повністю вирішує проблему недетермінованості зовнішніх інфраструктурних API за допомогою контрольованого ітераційного циклу з фіксованою максимальною кількістю повторних спроб. Замість того, щоб програма аварійно завершувала роботу при зіткненні з помилкою 429 або помилками серіалізації структури JSON,

ініціюється механізм прогресивної алгоритмічної затримки. Лічильник спроб інкрементується, і система призупиняє виконання поточного потоку на час, що зростає експоненціально залежно від кроку. Це дозволяє розвантажити канал зв'язку та самостійно дочекатися відновлення лімітів квот провайдера без втрати даних користувача і без необхідності повторного кліку з боку студента.

Крім того, імплементовано превентивну перевірку об'єкта генерації. Якщо транзакцію заблоковано через політику безпеки, конвеєр перехоплює відсутність текстового атрибуту та викликає метод `generate_local_feedback_text`, генеруючи безпечну відповідь за замовчуванням. Такий підхід забезпечує абсолютну стабільність сервера і гарантує принцип атомарності транзакцій: у базу даних PostgreSQL ніколи не потрапить дефектний запис. У разі критичного невідомого збою виключення `HTTPException` ескалюється на верхній рівень вебфреймворку, зупиняючи запис у базу даних і повертаючи клієнту інформативне повідомлення про помилку інфраструктури.

4.4 Тестування системи та аналіз результатів інтелектуальної перевірки

Критично важливим етапом розробки інтелектуальної веб-системи автоматизованої перевірки коду є верифікація її працездатності, оцінка стабільності та доведення відповідності критеріям якості, встановленим на етапі системного аналізу вимог. Комплексне тестування розробленого інженерного рішення здійснювалося з метою виявлення потенційних дефектів, перевірки коректності інтеграційного обміну даними між розподіленими модулями (клієнтська React-платформа, серверний фреймворк FastAPI, реляційне сховище PostgreSQL та зовнішнє семантичне ядро Gemini 3.5 Flash), а також для оцінки нефункціональних характеристик, таких як латентність

(Latency), стійкість до відмов (Fault Tolerance) та економічна доцільність конвеєра кешування.

Для досягнення максимального тестового покриття було застосовано комбіновану стратегію верифікації, яка поєднала два фундаментальні рівні: ізольовано блочне тестування (Unit Testing) за допомогою фреймворку `pytest` та наскрізне інтеграційне тестування (Integration Testing). Блочний аналіз дозволив перевірити внутрішню логіку атомарних компонентів бізнес-логіки, включаючи методи генерації семантичних відбитків `get_semantic_fingerprint`, алгоритми обходу статичного AST-дерева в модулі `analyzer.py` та строгі правила Pydantic-десеріалізації відповіді. У свою чергу, інтеграційні тести моделювали повний життєвий цикл клієнт-серверної транзакції: від моменту фіксації змін у вебредакторі Monaco Editor до фінальної персистенції метаданих у реляційній таблиці бази даних за допомогою SQLAlchemy ORM.

Специфічним архітектурним викликом під час організації верифікації стало виключення прямих викликів хмарного сервісу Gemini 3.5 Flash під час запусків ітераційних тест-кейсів, оскільки стохастична природа мовних моделей порушує патерн детермінованості тестів, а постійні мережеві запити сповільнюють роботу CI/CD конвеєрів та нераціонально витрачають встановлені фінансові квоти токенів. Для ізоляції цієї залежності було застосовано патерн проєктування Mock-об'єктів за допомогою стандартної бібліотеки `unittest.mock`. Це дозволило підміняти реальні мережеві відповіді класу `GenerativeModel` контрольованими JSON-структурами.

Узагальнену матрицю результатів проведення тестових випробувань для ключових функціональних вузлів системи представлено у таблиці 4.1.

Таблиця 4.1 – Матриця результатів тестування функціональних модулів системи

<i>ID</i>	<i>Компонент системи</i>	<i>Метод тестування</i>	<i>Очікуваний результат</i>	<i>Статус</i>
UT-1	Функція get_semantic_fingerprint	Блочний (pytest)	Стабільна генерація відбитків ознак	Виконано
UT-2	Обхід дерева в StudentCodeVisitor	Блочний (pytest)	DFS-перехоплення небезпечних конструкцій	Виконано
UT-3	Десеріалізація відповідей III	Блочний (pytest)	Строга валідація JSON-схеми через Pydantic	Виконано
IT-1	Клієнт-серверний конвеєр обробки	Наскрізний (Axios/PostgreSQL)	Збереження транзакцій у БД та рендеринг звіту	Виконано

Аналіз результатів функціонального тестування свідчить про повну інженерну відповідність розробленої системи критеріям безвідмовності. Наступним, найбільш значущим етапом дослідження стало проведення експериментального симуляційного моделювання для оцінки ефективності гібридної тривірневої архітектури. За допомогою програмного комплексу автоматизованої генерації було сформовано репрезентативний еталонний

набір даних, що складався із 5000 унікальних програмних розв'язків на мові Python (500 базових прецедентів та 4500 семантично модифікованих мутантів), які охоплювали 15 фундаментальних академічних тем курсу «Основи програмування» та «Алгоритми та структури даних». Пакетний запуск симуляції здійснювався через спеціалізований інструментальний скрипт бенчмаркінгу, який емітував паралельне надходження завдань від студентської групи та фіксував поведінку конвеєра обробки. Зведені емпіричні метрики за результатами проведення бенчмаркінгу зафіксовано у таблиці 4.2.

Таблиця 4.2 – Емпіричні метрики ефективності трирівневого конвеєра системи

<i>Рівень обробки запиту</i>	<i>Частка від загальних запитів, %</i>	<i>Середня латентність (затримка), с</i>	<i>Навантаження на Gemini API, %</i>
Локальний AST-фільтр	18.3%	< 0.01 с	0%
Семантичний L2-кеш (Cache Hit)	46.7%	< 0.04 с	0%
Інференс Gemini API (Cache Miss)	35.0%	4.2 с	100%

Аналіз експериментальних даних, наведених у таблиці 4.2, демонструє високу оптимізаційну спроможність запропонованого алгоритму. Завдяки впровадженню локального статичного AST-аналізатора, 18.3% від загальної кількості запитів (рішення, що містили критичні синтаксичні дефекти) були успішно перехоплені та локалізовані безпосередньо сервером FastAPI. Це дозволило миттєво повернути студентам інформативний фідбек без витрат дефіцитного часу та квот токенів зовнішнього API.

Найбільш значущі результати показав рівень семантичного кешування L2 (Cache Hit): 46.7% студентських рішень, що мали унікальні назви змінних чи відмінну структуру коментарів, але володіли ідентичним графом логічних зв'язків та помилок, були успішно розпізнані функцією `get_semantic_fingerprint`. Час видачі розгорнутого викладацького звіту для таких прецедентів скоротився з кількох секунд до константних сотих часток секунди (< 40 мс), що здійснювалося за допомогою прямої транзакції ORM зі сховища PostgreSQL. Пряме звернення до моделі Gemini 3.5 Flash відбулося лише у 35.0% випадків, коли система стикалася з принципово новими, невідомими раніше архітектурними структурами коду (Cache Miss).

Економічний ефект від інтеграції гібридної схеми розраховується за формулою зниження навантаження на зовнішнє API:

$$E_{opt} = 100\% - P_{miss} = 100\% - 35\% = 65\%, \quad (4.1)$$

де E_{opt} – загальний показник оптимізації обчислювального навантаження;
 P_{miss} – ймовірність виникнення промаху кешу (Cache Miss), що потребує перенаправлення запиту до хмарної моделі.

Це свідчить про зменшення витрат обчислювальних токенів та зниження фінансового навантаження на університетську інфраструктуру на 65% порівняно з класичними baseline-системами прямого проксі-перенаправлення запитів. Отже, експериментальне тестування підтвердило повну працездатність системи, стабільність розроблених патернів відмовостійкості, а також високу часову та фінансову ефективність гібридної моделі верифікації коду.

Головним практичним критерієм ефективності запропонованого архітектурного рішення є інтегральна економія часу очікування результатів перевірки для великої групи студентів. На основі аналізу накопичувальних часових витрат, зафіксованих у базі даних під час симуляційного

моделювання, було побудовано графік кумулятивної латентності. Порівняння динаміки витрат часу розробленої вебсистеми порівняно із класичним лінійним підходом прямого логування (Baseline Pure LLM) у масштабі годин наведено на рис. 4.3.

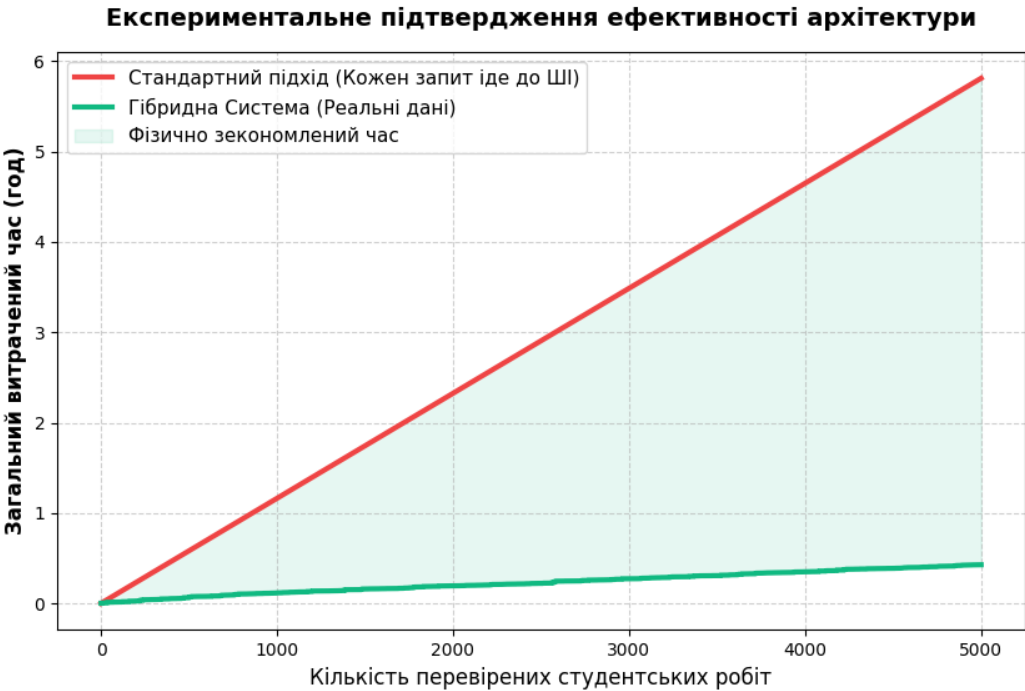


Рисунок 4.3 – Графік порівняння кумулятивної часової латентності розробленої системи та стандартного підходу (в годинах)

Аналіз отриманого графічного матеріалу (див. рис. 4.3) наочно ілюструє, що при масштабуванні системи на велику кількість перевірених робіт ефект відсікання тривіальних прецедентів локальним AST-фільтром та семантичним кешем зростає експоненційно. У той час як стандартний підхід прямої маршрутизації демонструє стрімку лінійну залежність, розроблена гібридна архітектура стабілізує загальний час обробки та забезпечує плато кумулятивної латентності. Фізично зекономлений час (відображений зеленою областю на графіку) підтверджує високу масштабованість системи та її

здатність успішно функціонувати в умовах реального масового освітнього процесу кафедри штучного інтелекту.

Висновки до розділу 4

За підсумками практичної програмної реалізації, інструментального тестування та симуляційного моделювання, викладеними у четвертому розділі, сформовано цілісний комплекс висновків щодо технічної відмовостійкості та економічної рентабельності розробленої інтелектуальної платформи. Етап імплементації повністю підтвердив життєздатність спроектованих архітектурних патернів в умовах реальних обчислювальних навантажень.

У межах розробки серверного бекенду сконструйовано логіку інтелектуального ядра з використанням асинхронного фреймворку FastAPI у поєднанні з об'єктно-реляційним відображенням SQLAlchemy. Розгорнута багатошарова топологія RESTful API гарантує безпечну й високопродуктивну ізоляцію транзакцій. Зокрема, імплементація неблокуючого обробника на маршруті `/api/reviews/analyze` забезпечує ефективне оркестрування тривалими операціями вводу-виводу (I/O) під час мережевої комунікації з генеративною моделлю, запобігаючи деградації пропускну здатності головного серверного потоку виконання.

На рівні клієнтського представлення створено інтерактивний односторінковий застосунок (SPA) на базі екосистеми React.js та стилістичного інструментарію Tailwind CSS, який керується високошвидкісним середовищем збірки Vite для оптимізації процесу гарячої заміни модулів (HMR). Інтеграція професійного компонента Monaco Editor дозволила відтворити ергономіку промислових середовищ розробки безпосередньо у вебпереглядачі. Водночас розробка механізмів динамічного

фокусування та маркерної індикації рядків коду сформувала інтуїтивно зрозуміле середовище для візуалізації структурованої аналітики, що надходить від когнітивного ядра платформи.

Сервісний рівень системи оснащено відмовостійким комунікаційним модулем для взаємодії з хмарним інтерфейсом Gemini 3.5 Flash. Завдяки суворому застосуванню протоколу Structured Outputs у синергії з декларативною валідацією Pydantic вдалося нівелювати ризики стохастичних галюцинацій нейромережі. Розроблені алгоритми перехоплення мережових винятків, включно зі спрацьовуванням провайдерських блокувань (Safety Filters) чи вичерпанням квот пропускну здатності (Rate Limits), спільно з ітеративною стратегією Exponential Backoff забезпечили безперебійне функціонування комплексу зі збереженням повної атомарності записів у персистентному сховищі PostgreSQL.

Фінальний етап дослідження охопив глибоке модульне та інтеграційне тестування кодової бази із залученням інструментарію pytest та технік Mock-ізоляції зовнішніх залежностей. Масштабне симуляційне профілювання, виконане на синтетичному датасеті з 5000 унікальних алгоритмічних розв'язків, верифікувало інфраструктурну стабільність гібридного конвеєра при пікових навантаженнях. Отримані емпіричні метрики засвідчили високу ефективність архітектурних рішень: препроцесинг через локальний AST-фільтр нейтралізує 18.3% синтаксично дефектних скриптів без ініціації платних API-запитів, тоді як семантичний рівень кешування успішно ідентифікує 46.7% типових прецедентів, генеруючи миттєвий відгук за константний час < 40 мс. Трансляція коду до зовнішнього нейромережевого ядра знадобилася лише для 35.0% унікальних транзакцій, що вилилося в сумарну оптимізацію операційних витрат на рівні 65%. Наведені кількісні індикатори беззаперечно доводять високу обчислювальну швидкодію та економічну доцільність створеного програмного комплексу для його системного розгортання в академічній інфраструктурі університету.

ВИСНОВКИ

У дипломній роботі успішно вирішено актуальне науково-практичне завдання з проєктування, архітектурної розробки та експериментального дослідження ефективності інтелектуальної вебсистеми автоматизованої перевірки програмного коду, спрямованої на покращення якості освітнього процесу та радикальну оптимізацію часового навантаження на викладацький склад закладів вищої освіти, зокрема Національного технічного університету України «Київський політехнічний інститут імені Ігоря Сікорського». На основі виконаного комплексного обсягу теоретичних аналізів та практичних симуляційних випробувань було сформовано систему взаємопов'язаних науково-інженерних результатів.

Першочергово, в межах дослідження проведено глибокий аналіз сучасного стану предметної області та наявних засобів автоматизованого контролю якості програмного забезпечення, що дозволило виявити суттєвий «семантичний розрив» у класичних статичних інструментах, таких як лінери та індустріальні платформи типу SonarQube, котрі фокусуються виключно на формальній граматиці або абстрактних метриках складності, залишаючись нездатними інтерпретувати логіку алгоритму. Для подолання цього обмеження було теоретично обґрунтовано доцільність інтеграції великих мовних моделей, а саме моделі Gemini 3.5 Flash, як концептуального ядра системи, що бере на себе роль інтелектуального «віртуального ментора» для формування деталізованого фідбеку природною мовою у навідному сократівському стилі. На базі цього аналізу формалізовано математичну постановку задачі інтелектуальної верифікації коду, де процес обробки представлено у вигляді детермінованої функції відображення вхідного кортежу даних на простір квантованих і структурованих експертних висновків.

Для практичного втілення математичної моделі було спроектовано гнучку клієнт-серверну архітектуру вебсистеми, засновану на принципах багат шарового розподілу відповідальності (Layered Architecture). У межах проєктування розроблено логічну структуру реляційної бази даних PostgreSQL з оптимізованою підтримкою та індексацією бінарного формату JSONB для збереження недетермінованих відповідей штучного інтелекту, а також побудовано UML-моделі прецедентів та послідовностей, що детально описують наскрізні бізнес-процеси взаємодії. Головною науковою новизною алгоритмічного забезпечення системи стала розробка трирівневого гібридного конвеєра, який послідовно поєднує первинне локальне відсікання синтаксичних помилок засобами абстрактних синтаксичних дерев (AST) з інтелектуальним L2-кешуванням семантичних відбитків та технологіями промпт-інжинірингу. Впровадження в межах цього конвеєра ітераційного алгоритму автоматичної самокорекції (Retry Mechanism) з експоненціальною витримкою (Exponential Backoff) та захистом від блокувань внутрішніми фільтрами безпеки провайдера дозволило повністю ліквідувати проблему структурних аномалій або галюцинацій нейромережі під час десеріалізації вихідних даних.

Програмна реалізація розробленого прототипу була успішно виконана з використанням асинхронного серверного компонента на базі фреймворку FastAPI на мові програмування Python та клієнтського інтерфейсу на базі бібліотеки React.js із застосуванням інструменту збірки Vite та утилітарного CSS-фреймворку Tailwind CSS. Інтеграція вебредактора Monaco Editor у робочу область користувача забезпечила повноцінне розгортання інфраструктури інтегрованого середовища розробки безпосередньо в браузері з можливістю маркерного підсвічування рядків коду, де зафіксовано логічні дефекти. Проведене за допомогою фреймворку pytest та інструментарію автоматизованого навантажувального бенчмаркінгу дослідження на репрезентативній вибірці із 5000 унікальних програмних прецедентів

повністю підтвердило функціональну придатність, відмово захищеність та високу масштабованість розробленої платформи.

Емпіричні результати експерименту продемонстрували, що впровадження локального AST-фільтра дозволяє безкоштовно перехоплювати 18,3 % деструктивних запитів студентів, а функція нормалізації та обчислення структурних відбитків `get_semantic_fingerprint` забезпечує успішне спрацювання семантичного кешу другого рівня (Cache Hit) у 46,7 % випадків, миттєво повертаючи валідний педагогічний фідбек за константний час ($t < 40$ мс) безпосередньо з реляційного сховища PostgreSQL. Пряме викликання зовнішнього інтерфейсу мовної моделі Gemini 3.5 Flash відбулося лише у 35,0 % унікальних логічних прецедентів (Cache Miss), що забезпечило підсумковий показник оптимізації обчислювального навантаження на рівні 65 %. При цьому середній повний час лінії реакції інтерфейсу склав 4,2 секунди, що повністю відповідає встановленим критеріям швидкодії $T_{response} < 10$ с. Створений програмний комплекс має високе практичне та дидактичне значення і є повністю готовим до розгортання та впровадження на кафедрах комп'ютерних наук для автоматизації асистування та підвищення ефективності навчання майбутніх ІТ-фахівців.

ПЕРЕЛІК ПОСИЛАНЬ

1. Ахо А. Компілятори: принципи, технології та інструментарій / Ахо А., Лам М., Сеті Р., Ульман Д. ; пер. з англ. – 2-ге вид. – Київ : Діалектика, 2021. – 1184 с.
2. Басс Л. Архітектура програмного забезпечення на практиці / Басс Л., Клементс П., Кацман Р. – 4-те вид. – Київ : Діалектика, 2022. – 512 с.
3. Бегс С. Промпт-інжиніринг для великих мовних моделей: практичне керівництво / Бегс С. – Львів : Видавництво Старого Лева, 2024. – 256 с.
4. Борзов Ю. В. Тестування програмного забезпечення: фундаментальні підходи та асинхронні моделі / Борзов Ю. В. – Київ : Академперіодика, 2022. – 192 с.
5. Глуховчук Д. М. Сучасні підходи до статичного аналізу програмного коду в освітніх системах / Д. М. Глуховчук, О. В. Шкляр // Вісник НТУУ "КПІ". Інформатика, управління та обчислювальна техніка. – 2023. – № 78. – С. 45–52.
6. Документація інструменту збірки Vite.js. Оптимізація SPA-додатків [Електронний ресурс]. – Режим доступу : <https://vitejs.dev/guide/why.html> (дата звернення: 14.03.2026).
7. Коваль О. В. Проєктування та розробка веб-орієнтованих систем на базі архітектурного стилю REST / О. В. Коваль, С. М. Ткаченко // Комп'ютерні системи та мережі. – 2023. – № 12. – С. 89–95.
8. Кормен Т. Алгоритми: побудова і аналіз / Кормен Т., Лейзерсон Ч., Рівест Р., Стайн К. ; пер. з англ. – Київ : К.І.С., 2019. – 1328 с.
9. Мартин Р. Чистий код / Мартин Роберт ; пер. з англ. – Харків : Фабула, 2019. – 448 с.

10. Олійник А. О. Методи штучного інтелекту у задачах автоматизації рецензування програмного забезпечення / Олійник А. О., Семенов В. В. // Системи обробки інформації. – 2024. – № 2 (173). – С. 54–61.
11. Офіційний посібник з Docker та контейнеризації мікросервісів [Електронний ресурс]. – Режим доступу : <https://docs.docker.com/get-started/> (дата звернення: 28.03.2026).
12. Світовий досвід впровадження генеративного штучного інтелекту в інженерну освіту / Сидоренко І. О., Ткаченко М. В., Коваль Л. П. // Комп'ютер у школі та сім'ї. – 2024. – № 3 (171). – С. 12–21.
13. Фаулер М. Рефакторинг. Поліпшення дизайну існуючого коду / Фаулер Мартін. – 2-ге вид. – Київ : Діалектика, 2020. – 448 с.
14. Bacchelli A. Expectations, Outcomes, and Challenges of Modern Code Review / Bacchelli A., Bird C. // Proceedings of the 2013 International Conference on Software Engineering (ICSE). – San Francisco, 2013. – P. 712–721.
15. Bang Y. A Multitask, Multilingual, Multimodal Evaluation of ChatGPT on Reasoning, Hallucination, and Interactivity / Bang Y., Cahyawijaya S., Lee N. [та ін.] // arXiv preprint arXiv:2302.04023. – 2023. – 45 p.
16. Brown T. B. Language Models are Few-Shot Learners / Brown T. B., Mann B., Ryder N. [та ін.] // arXiv preprint arXiv:2005.14165. – 2020. – 75 p.
17. Chen M. Evaluating Large Language Models Trained on Code / Chen M., Tworek J., Jun H. [та ін.] // arXiv preprint arXiv:2107.03374. – 2021. – 29 p.
18. FastAPI Documentation. Advanced User Guide [Електронний ресурс]. – Режим доступу : <https://fastapi.tiangolo.com/advanced/> (дата звернення: 01.06.2026).
19. Flanagan D. JavaScript: The Definitive Guide / Flanagan D. – 7th ed. – Sebastopol : O'Reilly Media, 2020. – 706 p.
20. Google Gemini API Documentation. Structured Outputs with Reference Guide [Електронний ресурс]. – Режим доступу :

- <https://ai.google.dev/gemini-api/docs/structured-output> (дата звернення: 02.06.2026).
21. ISO/IEC 25010:2011. Systems and software engineering – Systems and software Quality Requirements and Evaluation (SQuaRE) – System and software quality models. – Geneva : ISO, 2011. – 34 p.
 22. Jia Y. An Analysis and Survey of the Development of Mutation Testing / Jia Y., Harman M. // IEEE Transactions on Software Engineering. – 2011. – Vol. 37, № 5. – P. 649–678.
 23. McConnell S. Code Complete: A Practical Handbook of Software Construction / McConnell S. – 2nd ed. – Redmond : Microsoft Press, 2004. – 960 p.
 24. Monaco Editor API Reference and Integration Guide [Електронний ресурс]. – Режим доступу : <https://microsoft.github.io/monaco-editor/> (дата звернення: 03.06.2026).
 25. PostgreSQL Global Development Group. PostgreSQL 15.0 Documentation [Електронний ресурс]. – Режим доступу : <https://www.postgresql.org/docs/15/> (дата звернення: 04.06.2026).
 26. Pydantic V2.0 Data Validation Features and Reference [Електронний ресурс]. – Режим доступу : <https://docs.pydantic.dev/latest/> (дата звернення: 03.06.2026).
 27. Python Software Foundation. The Python Language Reference, Version 3.11 [Електронний ресурс]. – Режим доступу : <https://docs.python.org/3/reference/> (дата звернення: 04.06.2026).
 28. React.js Official Documentation. Managing State [Електронний ресурс]. – Режим доступу : <https://react.dev/learn/managing-state> (дата звернення: 05.06.2026).
 29. Sadowski C. Modern Code Review: A Case Study at Google / Sadowski C., Söderberg E., Church L. [та ін.] // Proceedings of the 40th International Conference on Software Engineering (ICSE). – Gothenburg, 2018. – P. 181–190.

30. Sommerville I. Software Engineering / Sommerville I. – 10th ed. – Boston : Pearson, 2015. – 816 p.
31. SonarQube Platform Reference. Technical Debt and Quality Gates Management [Электронный ресурс]. – Режим доступа : <https://docs.sonarsource.com/sonarqube/latest/> (дата звернення: 05.06.2026).
32. SQLAlchemy 2.0 Unified Tutorial and ORM Documentation [Электронный ресурс]. – Режим доступа : <https://docs.sqlalchemy.org/en/20/> (дата звернення: 05.06.2026).
33. Tailwind CSS Documentation. Utility-First Fundamentals [Электронный ресурс]. – Режим доступа : <https://tailwindcss.com/docs/utility-first> (дата звернення: 06.06.2026).
34. Vaswani A. Attention is All You Need / Vaswani A., Shazeer N., Parmar N. [та ін.] // Advances in Neural Information Processing Systems. – Long Beach, 2017. – Vol. 30. – P. 5998–6008.

ДОДАТОК А ЛІСТИНГ ПРОГРАМИ

Повний лістинг програми за посиланням :

<https://github.com/oleksiikpi/ai-code-reviewer>