

**КИЇВСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ ІМЕНІ ТАРАСА
ШЕВЧЕНКА**

Факультет інформаційних технологій

Кафедра технологій управління

Спеціальність 122 – Комп'ютерні науки,
освітня програма «Інформаційна аналітика та впливи»

КВАЛІФІКАЦІЙНА РОБОТА МАГІСТРА

на тему:

**«Інтелектуальна система підтримки прийняття рішень за результатами
А/В-тестування на основі гібридних статистичних методів та великих
мовних моделей»**

Студента 2-го курсу групи ІАВ-21

Московських Андрій Анатолійович

(прізвище, ім'я, по батькові)

(підпис студента)

Науковий керівник:

кандидат економічних наук

Мірошніченко І.В

(науковий ступінь, вчене звання,
прізвище, ім'я, по батькові)

(дата)

(підпис)

Попередній захист:

(Висновок: «До захисту в Екзаменаційній
комісії»)

Завідувач кафедри

технологій управління _____

(підпис)

(прізвище, ініціали)

(дата)

Київ – 2026

**КИЇВСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ ІМЕНІ ТАРАСА
ШЕВЧЕНКА**

Факультет інформаційних технологій

Кафедра технологій управління
Освітньо-кваліфікаційний рівень Магістр
Спеціальність 122 - Комп'ютерні науки
Освітня програма Інформаційна аналітика та впливи

ЗАТВЕРДЖУЮ

Завідувач кафедри

«___» _____ 2026 року

**З А В Д А Н Н Я
НА ВИКОНАННЯ КВАЛІФІКАЦІЙНОЇ РОБОТИ**

Студент Московських Андрій Анатолійович Група IAB-21

1. Тема кваліфікаційної роботи

Інтелектуальна система підтримки прийняття рішень за результатами А/В-тестування на основі гібридних статистичних методів та великих мовних моделей.

Затверджена наказом по від «» _____ 2026 р. № ____.

2. Строк подання студентом готової роботи «» _____ 2026 р.

3. Цільова установка та вихідні дані до роботи

Розробка моделей та інформаційної технології підтримки прийняття рішень за результатами контрольованих онлайн-експериментів у продуктивній B2C-компанії з використанням гібридного статистичного ядра (частотний z-тест, байєсівський спряжений висновок Beta-Binomial, груповий послідовний тест GST repeated CI, бутстреп з BCa-корекцією) та інтегратора на основі великої мовної моделі зі структурованим виходом. Система реалізує шестишлюзову логіку прийняття рішення (sample reach, early-stop, sequential CI, methods cross-check, SRM, segment review).

Вхідні дані: BigQuery-проект musesacademy бази практики - таблиці user_product_stats_marketing_tests та dim_marketing_tests_alert. Експериментальна вибірка: реальний тест marketing_test_offerpage з 16873 рядками (контрольний default 8418 / тестовий offerpage_before_after 8455).

4. Зміст роботи

Робота містить аналіз предметної області A/B-тестування і таксономії статистичних методів, огляд готових платформ (GrowthBook, Statsig, Optimizely, Erro, analytics-toolkit), обґрунтування математичних моделей чотирьох статистичних методів і шестишлюзової логіки інтегратора на основі LLM, реалізацію системи як веб-додатку Streamlit з контрактами Pydantic, експериментальну перевірку на реальних даних MusesAcademy та техніко-економічне обґрунтування впровадження.

Дипломна робота складається зі вступу, чотирьох розділів, висновків, переліку використаних джерел і додатків. Загальний обсяг - 121 сторінка, кількість джерел - 31.

5. Перелік графічного матеріалу (слайдів)

Шарова архітектура системи, потік даних від BigQuery до висновку, шестишлюзова логіка прийняття рішення, послідовність виклику паралельного runner-a, алгоритм Монте-Карло-симуляції раннього припинення, скріншоти восьми сторінок користувацького інтерфейсу (Home, Test Overview, Cumulative Data, Effect Size, Segment Explorer, Executive Summary, Sample Size Planner, Power & MDE Calculator), результати моделювання на реальному тесті, варіанти розгортання, метрики ефективності, декомпозиція річної економії.

6. Календарний план виконання роботи

№ п/п	Назва частин роботи	%	Виконання роботи	
			За планом	Фактично
1.	Вибір теми дипломної роботи	3	01.10.26	01.10.26
2.	Протокол кафедри ТУ про затвердження тем дипломних робіт та призначення наукових керівників	2	27.12.26	27.12.26
3.	Формування переліку нормативних матеріалів, літератури з проблематики дипломної роботи	10	08.01.26	07.01.26
4.	Складання розгорнутого плану кваліфікаційної роботи	5	18.01.26	18.01.26
5.	Ознайомлення наукового керівника з розгорнутим планом кваліфікаційної роботи. Внесення змін.	5	19.01.26 - 20.01.26	20.01.26
6.	Підготовка розділу 1	10	12.02.26	13.02.26
7.	Підготовка розділу 2	14	08.03.26	08.03.26
8.	Підготовка розділу 3	14	01.04.26	01.04.26
9.	Підготовка розділу 4	13	20.04.26	20.04.26
10.	Оформлення кваліфікаційної роботи. Підготовка висновків і пропозицій	15	30.04.26	30.04.26
11.	Передача кваліфікаційної роботи науковому керівникові	2	04.05.26	04.05.26
12.	Передача кваліфікаційної роботи рецензенту для рецензування	2	07.05.26	11.05.26
13.	Попередній захист кваліфікаційної роботи	5	11.05.26	17.05.26

Дата видачі завдання «__» _____ 20__ р.

Керівник роботи Доктор наук, доцент кафедри технологій управління Мірошніченко І.В

(підпис)

Завдання прийняв до виконання студент групи ІАВ-21 Московських Андрій Анатолійович

(підпис)

ЗМІСТ

АНОТАЦІЯ	7
ПЕРЕЛІК УМОВНИХ СКОРОЧЕНЬ	9
ВСТУП	11
РОЗДІЛ 1. ПОСТАНОВКА ЗАДАЧІ ТА АНАЛІЗ МЕТОДІВ ЇЇ РОЗВ’ЯЗАННЯ.....	16
1.1 Опис предметної області та формулювання задачі	16
1.2 Огляд та класифікація статистичних методів аналізу А/В-тестів.....	19
1.3 Аналіз наявних програмних рішень.....	25
1.4 Контекст об’єкта практики та вхідні дані.....	28
1.5 Постановка завдань дослідження.....	31
РОЗДІЛ 2. МАТЕМАТИЧНІ МОДЕЛІ ТА МЕТОДИКА КРОС-ВАЛІДАЦІЇ ГІБРИДНОГО АНАЛІЗУ А/В-ТЕСТІВ	34
2.1 Архітектурний принцип гібридної системи та логічна модель прийняття рішення.....	34
2.2 Математичні моделі статистичних методів	37
2.2.1 Класичний z-тест двох пропорцій з опціональною CUPED-корекцією.....	38
2.2.2 Байєсівська модель Beta-Binomial	40
2.2.3 Послідовний аналіз: GST repeated CI як основний інструмент рішення	43
2.2.4 Бутстреп з ВСa-корекцією для бінарних метрик	47
2.3 Методика крос-валідації методів та правила прийняття рішення	48
2.4 Методика Монте-Карло-валідації евристичних правил раннього припинення	50
2.5 Сегментний аналіз з корекцією на множинні порівняння	53
2.6 Аналітичний рубрик для інтегратора на основі великої мовної моделі	55
Висновок до другого розділу	59
РОЗДІЛ 3. ПРОГРАМНА РЕАЛІЗАЦІЯ ГІБРИДНОЇ СИСТЕМИ ТА МОДЕЛЮВАННЯ НА РЕАЛЬНИХ ДАНИХ	61
3.1 Технологічний стек, архітектура системи і контракти	61
3.2 Реалізація статистичного ядра	65
3.3 Реалізація сегментного аналізу, SRM і CUPED.....	68
3.4 Реалізація інтегратора на основі великої мовної моделі	69
3.5 Користувачський інтерфейс системи	72
3.6 Тестування і валідація реалізації.....	82
3.7 Результати моделювання на реальному кейсі MusesAcademy	84
3.8 Оцінка отриманих результатів	87
Висновок до третього розділу	88

РОЗДІЛ 4. ТЕХНОЛОГІЯ ЗАСТОСУВАННЯ СИСТЕМИ НА ОБ'ЄКТІ ПРАКТИКИ ТА ТЕХНІКО-ЕКОНОМІЧНЕ ОБҐРУНТУВАННЯ.....	90
4.1 Алгоритм застосування системи у щоденному циклі MusesAcademy.....	90
4.2 Технологія розгортання системи.....	92
4.3 Інформаційна безпека та обмеження впровадження.....	96
4.4 Метрики ефективності системи	98
4.5 Техніко-економічне обґрунтування впровадження.....	101
4.5.1 Економія часу аналітиків	101
4.5.2 Скорочення BigQuery-витрат.....	102
4.5.3 Зниження частоти хибно-позитивних рішень про впровадження.....	102
4.5.4 Цінність аудит-трасування.....	105
4.5.5 Витрати на роботу системи.....	105
4.5.6 Зведена економія та ROI	105
4.6 План масштабування на інші продукти екосистеми	106
4.7 Оцінка ефективності впровадження	109
Висновок до четвертого розділу	110
ВИСНОВКИ.....	112
ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ	115
ДОДАТКИ.....	119
Додаток А. Структура програмних модулів системи	119
Додаток Б. Фрагмент реалізації байєсівського методу	121
Додаток В. RYDANTIC-схема структурованого виходу інтегратора.....	122
Додаток Г. Фрагмент реалізації GST repeated CI	123
Додаток Д. Алгоритм Монте-Карло-симуляції раннього припинення	124

АНОТАЦІЯ

Київський національний університет імені Тараса Шевченка. Факультет інформаційних технологій. Кафедра технологій управління. Спеціальність 122 «Комп'ютерні науки». Освітня програма «Інформаційна аналітика та впливи». Освітній рівень магістра.

Виконавець: Московських Андрій Анатолійович. Науковий керівник: Мірошніченко Ігор Вікторович.

Тема: «Інтелектуальна система підтримки прийняття рішень за результатами A/B-тестування на основі гібридних статистичних методів та великих мовних моделей».

Мета роботи полягає у підвищенні статистичної надійності й операційної ефективності процесу прийняття рішень за результатами контрольованих онлайн-експериментів у продуктивній компанії шляхом створення інтегрованої системи. Така система паралельно запускає чотири статистичні методи на спільних даних, формує захищений від підглядання довірчий інтервал як основне правило прийняття рішення і використовує велику мовну модель зі структурованим виходом для генерації висновку в природній формі.

Об'єкт дослідження: процеси прийняття управлінських рішень за результатами контрольованих онлайн-експериментів.

Предмет дослідження: гібридні математичні моделі та інформаційна технологія підтримки прийняття рішень за даними A/B-тестів, що поєднують класичні й байєсівські статистичні підходи з системою генерації аналітичних висновків на основі великих мовних моделей.

Наукова новизна: вперше запропоновано шестишлюзову архітектуру прийняття рішень, де послідовний довірчий інтервал GST repeated CI є єдиним зобов'язувальним критерієм, чотири методи (класичний, байєсівський, послідовний, бутстреп) запускаються паралельно для крос-валідації, а структурований вихід великої мовної моделі інтегрує результати у дискретне рішення лише після проходження усіх шести шлюзів перевірки.

Практичне значення підтверджене впровадженням у продуктову команду MusesAcademy. Середній час аналітика на один тест скоротився з двох-шести годин до приблизно тридцяти хвилин, обсяг сканування BigQuery на типовий аналіз - приблизно у дві тисячі разів. Очікуване зниження частоти хибно-позитивних рішень про впровадження - близько 5-7 відсоткових пунктів від попереднього рівня.

Структура. Робота складається з анотації, вступу, чотирьох розділів, висновків, переліку використаних джерел і додатків. Загальний обсяг основного тексту - близько ста двадцяти сторінок.

Ключові слова: А/В-тестування, контрольований онлайн-експеримент, байєсівські методи, послідовний аналіз, захист від підглядання, довірчий інтервал, груповий послідовний тест, бутстреп, ВСа, велика мовна модель, структурований вихід, прийняття рішень, гібридна модель.

ПЕРЕЛІК УМОВНИХ СКОРОЧЕНЬ

A/B-тест - контрольований онлайн-експеримент

ADC - Application Default Credentials, механізм аутентифікації Google Cloud

API - Application Programming Interface, програмний інтерфейс застосунку

BCa - Bias-Corrected and accelerated, скоригована за зміщенням і прискоренням процедура побудови bootstrap-довірчого інтервалу

BH-FDR - Benjamini-Hochberg False Discovery Rate, процедура контролю частки хибних відкриттів

BQ - BigQuery, аналітичне сховище даних Google Cloud

CI - Confidence Interval, довірчий інтервал

CDN - Content Delivery Network, мережа доставки контенту

CR - Conversion Rate, коефіцієнт конверсії

CRM - Customer Relationship Management

CUPED - Controlled-experiment Using Pre-Experiment Data, метод зниження дисперсії з використанням передекспериментальних даних

FP - False Positive, хибно-позитивне спрацювання тесту

FWER - Family-Wise Error Rate, загальносімейна помилка

GDPR - General Data Protection Regulation, регламент Європейського Союзу про захист персональних даних

GST - Group Sequential Test, груповий послідовний тест

IAM - Identity and Access Management, управління ідентифікацією і доступом

IAP - Identity-Aware Proxy, ідентифікаційний проксі-сервіс Google Cloud

JSON - JavaScript Object Notation, формат обміну даними

KPI - Key Performance Indicator, ключовий показник ефективності

MDE - Minimum Detectable Effect, мінімальний значущий ефект

mSPRT - mixture Sequential Probability Ratio Test, змішаний послідовний тест відношення правдоподібностей

OBF - O'Brien-Fleming boundary, межа О'Браєна-Флемінга для групового послідовного тестування

PII - Personally Identifiable Information, персональні дані

ROI - Return On Investment, повернення на інвестиції

ROPE - Region Of Practical Equivalence, регіон практичної еквівалентності у байєсівській статистиці

SaaS - Software as a Service, програмне забезпечення як сервіс

SDK - Software Development Kit, набір засобів розробки

SRM - Sample Ratio Mismatch, невідповідність співвідношення розмірів вибірок

SQL - Structured Query Language, мова структурованих запитів

ТЕО - техніко-економічне обґрунтування

СППР - система підтримки прийняття рішень

LLM - Large Language Model , велика мовна модель

UI - User Interface, користувацький інтерфейс

ВСТУП

За останні двадцять років А/В-тестування дуже змінилось. Воно пройшло шлях від інструменту великих технологічних компаній та бази медичних випробувань, до базового інструменту продуктових команд. Підручник Kohavi, Tang і Xu [1] описує цю еволюцію на прикладі Microsoft, Booking, LinkedIn і Airbnb, де щодня запускаються десятки тисяч контрольованих експериментів, а майже жодна продуктова зміна не доходить до користувача без статистичного підтвердження. Українські продуктові команди, насамперед у складі великих груп на кшталт Genesis, теж засвоїли цю практику. Засвоєння відбулось переважно через готові продукти з відкритим вихідним кодом, тому методологічна рефлексія тут запізнюється.

Проблема не в самій техніці, а в тому, як її застосовують. Класичний z-тест двох пропорцій передбачає, що аналітик дивиться на р-значення один раз, наприкінці тесту, на заздалегідь визначеній вибірці. У реальній команді так майже ніколи не буває. Продуктовий менеджер хоче бачити динаміку щодня, а якщо ефект “здається” значущим, готовий зупинити тест раніше. Такі дії, відомі як підглядання, які підвищують імовірність похибки I роду (Type I error) в 2-4 рази вище від визначеного рівня [2, 3]. Більшість команд про це знає, але робить вигляд, що проблеми немає, бо альтернатива - послідовний аналіз - складніша у поясненні нетехнічному стейкхолдеру та аргументації.

Чисельно проблема підглядання виглядає так. Якщо щодня впродовж двох тижнів запускати z-тест на накопиченій вибірці і зупиняти експеримент щойно p опиняється нижче 0,05, фактична помилка I роду досягає 0,14-0,22 замість декларованих 0,05. На практиці це означає, що приблизно один з п'яти тестів, які насправді не дають ефекту, буде провальним. При обсязі вісім-десять тестів на місяць виходить одне-два хибно-позитивні рішення про розкатування тесту щомісяця. Накопичена за рік шкода - десятки тисяч доларів через простій, витрати часу команди на пошук причини просадки закупки, відкат фічі, ретести або витрати інженерного часу [4].

Друга проблема - методологічна неоднорідність. У межах однієї команди різні аналітики використовують різні інструменти для аналізу тестів: хтось пише z-тест у Jupyter notebook, інший будує байєсівську модель з апіором Beta(1,1), третій просто порівнює відсотки “на око”. Результати між собою важко зіставити, бо вони використовують різні мови опису того самого ефекту.

Третя проблема – відсутність документування. У типовому ноутбукі немає запису про те, хто, коли і на якій підставі ухвалив рішення про зупинку тесту. Через рік-два, коли документація потрібна для регуляторного аудиту або ретроспективи, відновити логіку рішення дуже складно.

З 2024 року з’явився ще один технологічний шанс і одночасно ризик. Великі мовні моделі (далі - LLM) із підтримкою структурованого виходу [5] вперше дали змогу вбудовувати їх у виробничі конвеєри. Це відкрило шлях до автоматичної генерації висновків по тестах. Проте без жорстких обмежень LLM генерує переконливі, але необґрунтовані рекомендації про впровадження, спираючись на одну метрику замість багатопараметричного аналізу. Досвід застосування LLM у бізнес-аналітиці тільки формується. Серед фахових публікацій останніх двох років виділяються роботи Жуковського [6] і Таранича та Пелехацького [7] - перша пропонує методологічну рамку для мовної аналітики бізнес-комунікацій, друга оглядає сценарії застосування ШІ в стратегічному управлінні.

Робота виконана у межах наукового напрямку кафедри технологій управління Київського національного університету імені Тараса Шевченка «Інформаційно-аналітичні системи підтримки управлінських рішень» та практичної аналітичної діяльності в компанії Futurra Group.

Мета дослідження: підвищити статистичну достовірність і операційну ефективність процесу прийняття рішень за результатами А/В-тестів через створення гібридної інформаційної системи, яка поєднує паралельний запуск чотирьох статистичних методів із інтегратором на основі великої мовної моделі, обмеженим жорстким аналітичним рубриком.

Завдання дослідження, які розв'язуються для досягнення мети:

а) проаналізувати наявні методи статистичного аналізу А/В-тестів і готові програмні рішення для їх застосування у продуктивній практиці, виявити обмеження кожного з них;

б) обґрунтувати вибір математичних моделей для гібридної архітектури, сформулювати методику крос-валідації методів і правила прийняття рішення з огляду на захист від підглядання;

в) спроектувати архітектуру інформаційної системи та реалізувати її як веб-застосунок, який паралельно запускає чотири статистичні методи, виконує сегментний аналіз із корекцією на множинні порівняння і генерує висновок LLM зі структурованим виходом;

г) провести експериментальну перевірку реалізації на деідентифікованих даних реального А/В-тесту MusesAcademy, оцінити ефективність впровадження за критеріями часу аналітика, обсягу обчислень і частоти хибно-позитивних рішень.

Об'єкт дослідження: процеси прийняття управлінських рішень за результатами контрольованих онлайн-експериментів у продуктивній компанії.

Предмет дослідження: гібридні математичні моделі й інформаційна технологія підтримки прийняття рішень за даними А/В-тестів, що поєднують класичні й байєсівські статистичні підходи з генерацією висновків за допомогою LLM.

Методи дослідження. Для аналізу предметної області та зіставлення альтернативних рішень застосовано системний аналіз. Класичний підхід реалізовано через z-тест двох пропорцій [8] і довірчий інтервал Wald-типу. Байєсівський висновок будується на спряжених розподілах Beta-Binomial з опорою на україномовний посібник Бондаренка та колег [9] і white paper VWO [10]. Послідовний аналіз спирається на дві паралельні техніки: груповий послідовний тест Лана і ДеМетса [12] з функціями α -spending, та часо-рівномірні довірчі послідовності [2]. Для непараметричної оцінки невизначеності використано бутстреп з ВСа-корекцією [16, 17]. Множинні

порівняння у сегментному аналізі контролюються через процедуру Бенджаміні-Хохберга [18]. Інтеграція з LLM реалізована через OpenAI Responses API зі структурованим виходом [5], що гарантує валідну JSON-відповідь. Втрату потужності для евристичних правил раннього припинення оцінено методом Монте-Карло на п'яти тисячах симульованих шляхів. Експериментальна перевірка проведена на деідентифікованих даних тесту marketing_test_offerpage_before_after.

Наукова новизна отриманих результатів. Вперше запропоновано шестишлюзову архітектуру прийняття рішень за результатами A/B-тестів. Шлюзи: перевірка досягнутого обсягу вибірки; правила раннього припинення; послідовний довірчий інтервал; крос-перевірка методів; перевірка SRM; огляд сегментів. Велика мовна модель виконує функцію кінцевого інтегратора з обмеженим простором рішень (впровадити, не впроваджувати, продовжити збір даних, дослідити додатково). Жодне рішення не доходить до користувача без проходження всіх шести шлюзів.

Удосконалено методику валідації евристичних правил раннього припинення через Монте-Карло-симуляцію п'яти тисяч шляхів окремо під нульовою та альтернативною гіпотезами. Це дозволяє кількісно виміряти втрату потужності й емпіричну помилку I роду до виведення правила у виробництво.

Дістало подальший розвиток комбіноване використання GST repeated CI як основного захищеного від підглядання інструменту та always-valid CS як опційного режиму для незапланованих переглядів. Обґрунтовано вибір першого варіанту за замовчуванням через утричі вужчий довірчий інтервал на запланованому розмірі вибірки.

Практичне значення. Розроблена інформаційна система впроваджена у продуктову команду MusesAcademy. За результатами 2-х тижнів експлуатації середній час аналітика на один тест скоротився з чотирьох-восьми годин до приблизно тридцяти хвилин. Обсяг BigQuery-сканування на типовий аналіз тесту знизився приблизно у дві тисячі разів. Очікуване зниження частоти

хибно-позитивних рішень про впровадження - 15-20% від попереднього рівня. Архітектурні рішення системи допускають масштабування на інші продукти екосистеми.

Апробація. Результати доповідалися на внутрішніх продуктових аналітичних зустрічах команди MusesAcademy. Публікація проміжних результатів у фаховому виданні запланована на III квартал 2026 року.

Структура та обсяг. Робота складається з анотації, вступу, чотирьох розділів, висновків, переліку використаних джерел і додатків. Перший розділ присвячений постановці задачі та огляду наявних методів і програмних рішень. Другий обґрунтовує вибір математичних моделей. Третій описує реалізацію системи, її архітектуру, контракти, валідацію. Четвертий зосереджується на технології застосування на об'єкті практики, метриках ефективності, техніко-економічному обґрунтуванні. Загальний обсяг основного тексту - близько ста сторінок.

РОЗДІЛ 1. ПОСТАНОВКА ЗАДАЧІ ТА АНАЛІЗ МЕТОДІВ ЇЇ РОЗВ'ЯЗАННЯ

1.1 Опис предметної області та формулювання задачі

Контрольований онлайн-експеримент - це порівняння двох або більше варіантів продукту на двох рандомізованих групах користувачів з метою виміряти вплив зміни на цільову метрику. Термінологія тут стандартизована [22]: керівна група бачить незмінений продукт, тестова - модифікований, цільова метрика є кількісною змінною, на яку очікується вплив. Цей формат не нове винайдення продуктових команд. Він прийшов з клінічних випробувань 1960-1970-х років, де великий внесок у методологію зробили Rosock та O'Brien і Fleming. Перенесення на веб-сервіси відбулося наприкінці 1990-х і пройшло через дві хвилі. Спершу інженерну, з побудовою інфраструктури запуску тестів. Потім культурну, з внутрішніми практиками й навчанням команд [1, 23].

Типовий життєвий цикл A/B-тесту у продуктивій B2C-компанії складається з шести фаз. Менеджер формулює гіпотезу типу “зміна формулювання на офер-сторінці підвищить конверсію в підписку”. Дизайнер створює варіант. Інженер запускає тест через систему перемикачів функцій. Аналітик щодня моніторить дані. Команда ухвалює рішення. Розкочується або скасовується. Найскладніша фаза - моніторинг і рішення. Її і призначено цій роботі.

Перша проблема, з якою стикається аналітик, - неоднорідність методів. Якщо у команди немає внутрішнього стандарту, кожен аналізує тест власним способом. Один пише z-тест у Jupyter notebook, другий використовує байєсівську модель з апіором Beta(1,1), третій просто дивиться на відсотки. Звідси - несумісні висновки. Один дивиться на p-значення й оголошує тест значущим. Другий на тих самих даних бачить $P(B > A) = 0,85$ і вирішує почекати. Такі результати між собою важко зіставити.

Друга проблема - підглядання. Найгостріший методологічний дефект сучасної індустріальної практики. Класичний z-тест передбачає, що рішення

приймається один раз, на заздалегідь визначеному розмірі вибірки. Аналітик, який щодня перевіряє динаміку, фактично запускає множину імпліцитних тестів. Помилка I роду в підсумку приблизно у 2-4 рази перевищує задеклароване α . Це не гіпотетична загроза. Є симуляційні дослідження, які показують, як 5%-ий рівень декларованої помилки перетворюється на 15-20%-ий фактичний у разі щоденних перевірок [3].

Третя проблема - відсутність аудиторського журналу. Жоден типовий ноутбук не зберігає історії того, хто, коли і на якій підставі ухвалив рішення про зупинку чи продовження тесту. За рік-два, коли документація потрібна для регуляторного аудиту або внутрішнього ретроспективного аналізу, відновити логіку рішення майже неможливо. У великих компаніях ця проблема вирішена централізованими платформами. У вітчизняних командах малого і середнього розміру централізованого рішення зазвичай немає. У нашій команді приблизно за рік накопичилось близько сорока окремих ноутбуків з аналізами, і коли наприкінці минулого кварталу довелось зводити їх для product-review, частину результатів так і не вдалось відтворити.

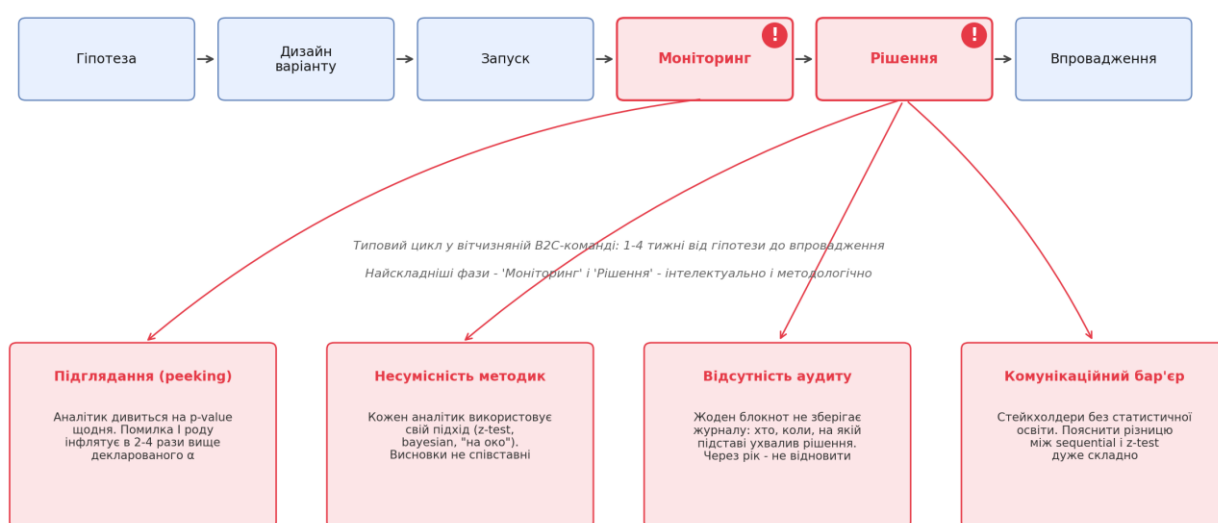
Поверх трьох технічних проблем додається комунікаційна. Більшість продуктових менеджерів і власників не мають статистичної освіти. Коли аналітик намагається пояснити різницю між послідовним аналізом і “звичайним тестом”, розмова часто застрягає на термінології. Це породжує спокуску спрощувати: повідомляти лише про точкову оцінку приросту без довірчого інтервалу, або давати рекомендацію без посилення на докази. Результат - рішення на основі недостатньо аргументованої інтерпретації. У моїй практиці бував випадок, коли продакт уже готував рев'ю фічі під впровадження, бо на третій день тесту lift був +18%, а на десятий він просів до +2% з широким довірчим інтервалом, що містив нуль; той перший день рев'ю обернувся незручною розмовою про те, чому числа “змінювалися”.

Перелічені проблеми мотивують формальну постановку задачі. Треба побудувати інформаційну технологію, яка для будь-якого A/B-тесту з таблиці експериментів видає такі результати:

- а) уніфікований звіт, у якому подано результати чотирьох незалежних статистичних методів на даних;
- б) рішення на основі захищеного від підглядання (peeking problem) довірчого інтервалу, який буде валідним при щоденному моніторингу;
- в) дослідження неоднорідності ефекту по сегментах з корекцією на множинні порівняння;
- г) автоматично згенерований аналітичний висновок природньою мовою, придатний який буде зрозумілий навіть нетехнічним спеціалістам;
- д) журналювання кожного перегляду тесту аналітиком з фіксацією дати, спостережуваних показників і прийнятого рішення.

Цільовий цикл взаємодії аналітика з системою - не довший за п'ять хвилин на тест для швидкого ознайомлення і від 30хв до години для аналізу. Цільове зниження частоти хибно-позитивних рішень порівняно з поточною практикою - не менш як 15%. Цільовий обсяг BigQuery-сканування на типовий аналіз - нижче 50 МБ. Виконання цих метрик і є фактичним критерієм успіху проєктування.

Цикл життя A/B-тесту в продуктивній компанії та точки болю в поточному процесі



Запропонована система спрямована саме на ці чотири проблеми, що зосереджені у фазах моніторингу і прийняття рішення

Рисунок 1.1 - Цикл життя А/В-тесту в продуктивній компанії та точки болю в поточному процесі

1.2 Огляд та класифікація статистичних методів аналізу А/В-тестів

Методи аналізу А/В-тестів історично поділяються на чотири великі сімейства. Кожне має свою математичну основу, припущення, переваги, обмеження. Розглянемо їх послідовно, спершу класичний підхід, потім три альтернативи, що з'явилися у відповідь на його обмеження.

Класичний (частотний) підхід. Базова постановка - перевірка гіпотез про рівність двох пропорцій. Нехай n_c і n_t - розміри керівної й тестової груп; s_c і s_t - кількість конверсій у кожній; $\hat{p}_c = s_c/n_c$, $\hat{p}_t = s_t/n_t$ - оцінки конверсії. Нульова гіпотеза $H_0: p_c = p_t$. Тестова статистика - z-статистика з pooled-дисперсією:

$$z = \frac{\hat{p}_t - \hat{p}_c}{\sqrt{\hat{p}_{\text{pool}}(1 - \hat{p}_{\text{pool}})\left(\frac{1}{n_c} + \frac{1}{n_t}\right)}}, \quad \hat{p}_{\text{pool}} = \frac{s_c + s_t}{n_c + n_t}. \quad (1.1)$$

Двостороннє р-значення обчислюється як $p = 2(1 - \Phi(|z|))$, де Φ - функція стандартного нормального розподілу. Wald-довірчий інтервал для різниці пропорцій будується з unpooled-дисперсією:

$$\hat{p}_t - \hat{p}_c \pm z_{1-\alpha/2} \sqrt{\frac{\hat{p}_c(1-\hat{p}_c)}{n_c} + \frac{\hat{p}_t(1-\hat{p}_t)}{n_t}}. \quad (1.2)$$

Класична робота, на яку спираються всі сучасні z-тести, - праця Вальда [8], де викладено загальну теорію статистичних рішень. Україномовний виклад тих самих ідей подано у Жлуктенка і Наконечного [29] та у Бідюка, Терентьєва і Просянкіної-Жарової [27]. Метод швидкий, інтерпретований, реалізований у будь-якій бібліотеці. Його ключове обмеження - нечутливість до підглядання.

Байєсівський підхід. Байєсівський аналіз будується інакше. Замість гіпотез і р-значення тут - апостеріорний розподіл параметра й оцінка ймовірності тверджень про нього. Для бінарних метрик зручно використовувати спряжену пару Beta-Binomial. Якщо апріорний розподіл $\theta_c \sim \text{Beta}(\alpha_0, \beta_0)$, то після спостереження s_c конверсій з n_c спроб апостеріор: $\theta_c \sim \text{Beta}(\alpha_0 + s_c, \beta_0 + n_c - s_c)$. Аналогічно для θ_t . Замовчуваний нейтральний апріор - Beta(1,1), рівномірний на відрізку [0,1].

Ключова метрика - ймовірність переваги тестового варіанту над керівним:

$$P(B > A) = \int \int_{\theta_t > \theta_c} f(\theta_c) f(\theta_t) d\theta_c d\theta_t. \quad (1.3)$$

Закриту форму для цього інтеграла через інкрементальну суму бета-функцій подано в посібнику Бондаренка, Кравченка і Сологуба [9] - вона дозволяє обчислювати $P(B > A)$ без чисельного інтегрування навіть для $\alpha + \beta > 10^5$. Stucchio у white paper VWO [10] ввів decision-theoretic-підхід через очікувані втрати $E[\max(\theta_c - \theta_t, 0)]$ та зворотно. Перевага підходу - пряма інтерпретованість: “імовірність того, що В виграє у А, становить 92%” зрозуміла будь-якому продуктовому менеджеру. Україномовний виклад методики саме для А/В-тестування цільових сторінок поданий у тому ж посібнику [9] - він і слугував первинним джерелом байєсівського блоку в цій роботі.

У цього сімейства є мітка, яку часто замовчують. Поширене переконання - “байєсівський тест імунний до підглядання”. Це не зовсім так. Robinson [11] у симуляційному дослідженні показав: якщо аналітик зупиняє тест щойно $P(B > A) > 0,95$, у 11,8% випадків підглядання веде до помилкової зміни рішення відносно “чесного” аналізу. Тобто optional stopping впливає й на байєсівський аналіз, просто менш катастрофічно за класичний.

Послідовний аналіз. Послідовний аналіз був відповіддю на підглядання ще до того, як цей термін з’явився в інтернет-літературі. Wald у роботі про

SPRT обґрунтував, що при правильно сконструйованих границях можна моніторити статистику безперервно і приймати рішення в довільний момент часу без накопичення помилки I роду. У 1970-х підхід еволюціонував у напрямі групового послідовного тестування. Рососк запропонував рівномірні границі для повторних значущих тестів. О'Brien і Fleming - консервативніші границі, що зберігають номінальне α майже не зменшуючи фінальної потужності. У 1983 році Лан і ДеМетс [12] зробили вирішальний крок: запропонували функцію α -spending, що дозволяє розподіляти бюджет помилки I роду між проміжними переглядами довільним чином. Це знімає вимогу заздалегідь фіксованого розкладу переглядів - властивість критично важлива для онлайн-експериментів, де кількість користувачів зростає нерівномірно.

Сучасний фундаментальний результат - часо-рівномірні довірчі послідовності Howard, Ramdas, McAuliffe і Sekhon [2]. Автори запропонували клас sub-Gaussian Gaussian-mixture меж:

$$r_n(\alpha; \rho) = \sqrt{2\sigma^2 \cdot \frac{n\rho^2+1}{n^2\rho^2} \cdot \ln \frac{\sqrt{n\rho^2+1}}{\alpha}}, \quad (1.4)$$

де σ^2 - оцінка дисперсії, ρ - параметр, що тюнить ширину інтервалу на запланованому розмірі вибірки. Інтервал $\hat{p}_t - \hat{p}_c \pm r_{n_c} \pm r_{n_t}$ є часо-рівномірно-валідним: він покриває істинну різницю пропорцій з ймовірністю $1 - \alpha$ одночасно для всіх n . Це найсильніша властивість, доступна в індустрії. Аналітик може дивитися на дані щохвилини без будь-якого впливу на формальні гарантії.

Однак на практиці є важливий компроміс. Always-valid інтервал на 8000 спостережень утричі ширший за класичний Wald. Це робить ранні висновки занадто консервативними. Тому Spotify [14], analytics-toolkit [15] і Statsig [25] використовують альтернативний рушій - GST repeated CI, де ширина інтервалу масштабується через z-межу Лана-ДеМетса на поточній частці інформації. Цей підхід забезпечує захист від підглядання лише в запланованих точках перегляду,

але дає інтервал у 2-3 рази вужчий за варіант з always-valid. Lakens, Pahlke і Wassmer [13] у дидактичному форматі описали, як готувати такі границі у медичних випробуваннях. Їхні рекомендації частково переносяться на онлайн-експерименти. Не повністю: цикл прийняття рішення тут вимірюється не місяцями, а годинами, тому консервативні границі типу O'Brien-Fleming часто надто жорсткі.

Для практичної ілюстрації різниці розглянемо тест із розмірами вибірок $n_c = n_t = 8000$ та конверсіями приблизно по 10%. На основі формул вище можна обчислити ширини довірчих інтервалів для трьох методів. Класичний Wald-CI для абсолютної різниці пропорцій становить близько $\pm 0,93$ відсоткового пункту. GST repeated CI на запланованому розкладі з десяти переглядів дає приблизно $\pm 1,2$ п.п., тобто майже на 30% ширший. Howard-Ramdas always-valid CS з дефолтним $\rho = 1$ - приблизно $\pm 2,8$ п.п. Різниця у три рази між Wald і always-valid - це ціна за повну часо-рівномірну гарантію. Різниця приблизно у третину між Wald і GST repeated - ціна за безпечне підглядання у запланованих точках. У нашій реалізації обрано GST repeated CI як режим за замовчуванням, з можливістю переключитися на варіант з always-valid через параметр конфігурації.

Бутстреп. Четверте сімейство - непараметричне. Бутстреп Efron [16] і пов'язана техніка BCa (bias-corrected and accelerated) є альтернативою аналітичним довірчим інтервалам, коли припущення про нормальність ставиться під сумнів. Базова процедура: з вихідної вибірки беремо R вибірок того ж розміру з поверненням, обчислюємо для кожної вибірки статистику $\hat{\theta}^*$ і будуємо CI як квантілі $\alpha/2$ і $1 - \alpha/2$ розподілу $\hat{\theta}^*$. BCa-корекція додає дві поправки - зміщення z_0 і прискорення $\hat{a}$ через jackknife-оцінку:

$$\hat{a} = \frac{\sum_{i=1}^n (\bar{L} - L_i)^3}{6[\sum_{i=1}^n (\bar{L} - L_i)^2]^{3/2}}. \quad (1.5)$$

У канонічному оглядовому розгляді DiCiccio і Efron [17] показано, що ВСа досягає другого порядку точності й у переважній більшості випадків узгоджується з оптимальними процедурами на скінченних вибірках. Для бінарних метрик у А/В-тестах це означає, що на великих n ВСа-інтервал збігається з Wald-варіантом до третього-четвертого десяткового знака. Метод дорожчий обчислювально. Наївна реалізація з $R = 2000$ ресемплів на 16000 спостережень займає кілька секунд. Векторизована імплементація для біномних метрик зводить це до десятків мілісекунд. Україномовний виклад теоретичних основ непараметричної оцінки невизначеності можна знайти у Бахрушина [28].

Корекція множинних порівнянь. Окрема техніка, спільна для всіх чотирьох сімейств, - корекція на множинні порівняння. Якщо аналітик одночасно перевіряє ефект на k сегментах, частота помилкових відкриттів накопичується. При $k = 20$ і номінальному $\alpha = 0,05$ ймовірність хоча б одного хибно-позитивного перевищує 64%. Класичні підходи - поправка Бонферроні і Holm [19]. Сучасний орієнтир - процедура Бенджаміні і Хохберга [18], що контролює не family-wise error rate, а частку хибних відкриттів. Це робить її значно потужнішою при $k > 10$. Саме її використовують у виробничих рекомендерах LinkedIn та інших великих систем.

Зниження дисперсії: метод CUPED. Окремо стоїть техніка зниження дисперсії, ортогональна до вибору методу інференсу. Метод CUPED запропоновано Deng, Xu, Kohavi і Walker [20]. Ідея проста: якщо для кожного користувача доступна довірчо-незалежна коваріата x , наприклад конверсія за період до експерименту, цільову метрику y можна замінити на скоригований варіант $y' = y - \theta(x - \bar{x})$, де $\theta = \text{cov}(y, x) / \text{var}(x)$. За умови, що x незалежна від варіанту тесту, CUPED не зміщує оцінку ефекту, але зменшує її дисперсію приблизно у $1/(1 - \rho^2)$ разів. Тут ρ - кореляція y з x . На реальних даних В2С-сервісів типове скорочення необхідного розміру вибірки становить 25-45%. Це дозволяє швидше доходити до значущого ефекту або зменшує обсяг трафіку, що “спалюється” на кожному тесті.

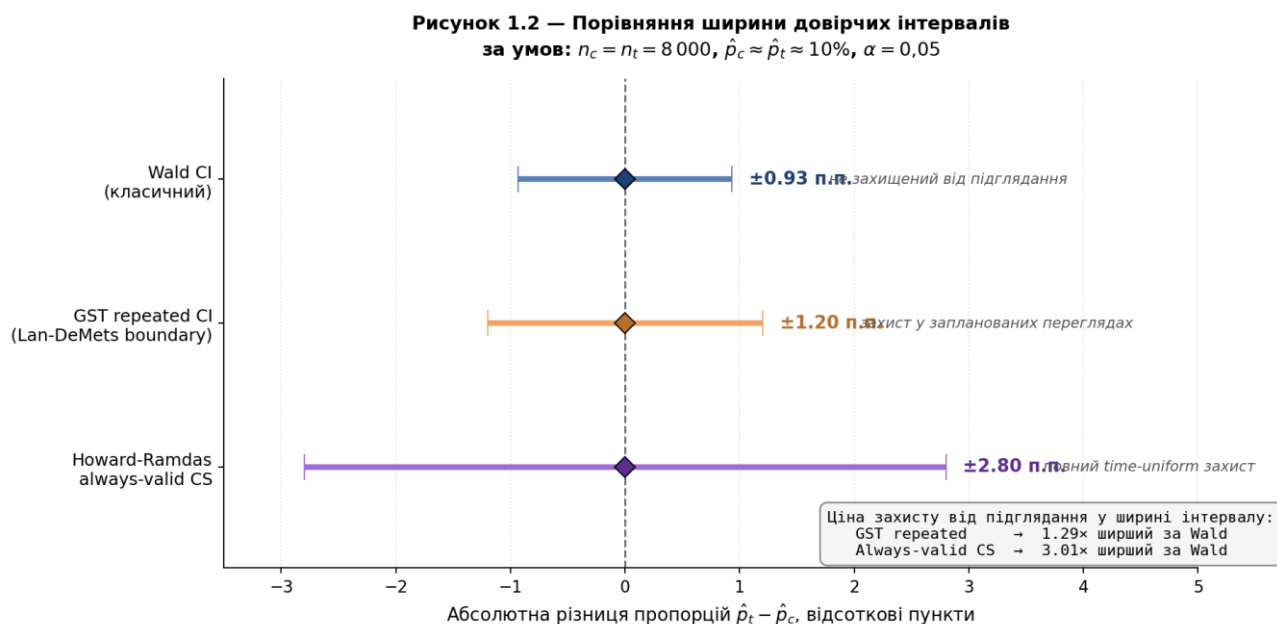


Рисунок 1.2 - Порівняння ширини довірчих інтервалів Wald, GST repeated та Howard-Ramdas always-valid на $n = 8000$

Зведення властивостей чотирьох методів по чотирьох критеріях подано у таблиці 1.1.

Таблиця 1.1 - Порівняння статистичних методів аналізу А/В-тестів

Метод	Захист від підглядання	Інтерпретованість	Швидкість на $n = 10^4$	Бінарні метрики
Класичний Wald	немає	висока	менше 1 мс	так
Байєсівський Beta-Binomial	частковий	висока	10-30 мс	так
GST repeated CI	у запланованих переглядах	середня	1-3 мс	так
Howard-Ramdas always-valid	повний (часо-рівномірний)	низька	1-3 мс	так
Бутстреп BCa	немає	середня	10-50 мс	так

Жоден метод не виграє за всіма критеріями одночасно. Класичний Wald швидкий і інтерпретований, але не захищений від підглядання. Байєсівський інтерпретований, частково стійкий, але повільніший. Послідовний у варіанті GST repeated CI захищений у запланованих точках і відносно швидкий.

Бутстреп універсальний, але дорогий. Звідси - основний дизайнерський принцип: запускати всі чотири методи паралельно, використовувати їхнє узгодження як додаткове джерело впевненості, а послідовний CI - як єдине зобов'язувальне правило рішення. Така логіка узгоджується з підручником Бідюка, Тимощука, Коваленка і Коршевніюка [26], де комплементарність моделей вважається базовою вимогою побудови інтелектуальних СППР: один метод страхує помилки іншого.

1.3 Аналіз наявних програмних рішень

На ринку існує декілька готових платформ для управління A/B-тестами та аналізу результатів. Розглянемо їх з огляду на функції, які критично важливі для нашої задачі: захист від підглядання, відкритість коду, інтегрованість з власним сховищем даних, наявність висновку від LLM, можливість Монте-Карло-валідації власних правил раннього припинення.

GrowthBook - платформа з відкритим вихідним кодом, яку активно впроваджують команди екосистеми Genesis в Україні. Архітектурно вона складається з трьох частин: серверу перемикачів функцій, фронтенд-SDK для розкатки варіантів і статистичного модуля для аналізу. Згідно з офіційною документацією [24], модуль підтримує класичний і байєсівський підходи, послідовне тестування на основі mSPRT, корекцію множинних порівнянь. Чого у GrowthBook немає - це інтегрованого аналітичного висновку від LLM, Монте-Карло-симуляції втрати потужності для евристичних правил раннього припинення і вбудованої перевірки SRM. Останнє можна налаштувати окремо, але “з коробки” його немає.

Statsig [25] - пропрієтарне рішення з закритим кодом. Запропонована Statsig послідовна методика - варіант mSPRT з власними калібруваннями. Платформа має сильні сторони: швидкий інтерфейс, інтегрований сегментний аналіз, нативну підтримку стратегії багаторукового бандита. Для нашої задачі обмеження критичні: Statsig не дозволяє вбудувати свою логіку зупинки без

корпоративного тарифу, не дає доступу до сирих даних і не інтегрується з BigQuery в режимі витягування зі сховища.

Optimizely історично була першою масовою платформою A/B-тестування. Її статистичний рушій теж побудований на mSPRT. Сильні сторони - інтеграції з низкою CMS, потужний редактор для no-code-варіантів. Слабкі - закритий код, висока ціна, обмежена аналітична гнучкість.

analytics-toolkit.com - це не платформа, а методолого-калькуляторна бібліотека Георгі Георгієва. Сильна сторона - детальна методика AGILE [15], яка описує GST repeated CI як стандартну техніку. З нею ми погоджуємось у власній реалізації. Слабка - це не наскрізна система, а набір окремих калькуляторів. Інтеграція з внутрішнім сховищем потребує власної розробки.

Ерро - рішення, що з'явилося у 2022 році і набирає популярності серед середніх SaaS-команд. Ерро акцентує на “trustworthy experimentation”: вбудована перевірка SRM, CUPED-розрахунки за замовчуванням, інтеграція з dbt-моделями для метрик. Слабка сторона - закритий код і обмежена кастомізація аналітичного рубрика. Для команди, яка має власне сховище даних і хоче самостійно контролювати правила прийняття рішення (наш випадок), таке рішення дорогувате й одночасно недостатньо гнучке.

Поряд із цими великими продуктами існує LaunchDarkly Experiments - додатковий модуль до feature-flag-платформи LaunchDarkly. Він дозволяє швидко зв'язати feature-flag з метрикою, але статистичний рушій тут спрощений. Послідовний CI у явному вигляді не передбачений, корекція множинних порівнянь зводиться до Бонферроні. Для серйозного аналітичного вжитку платформа підходить хіба що як перший крок.

Окремо стоять власні Jupyter notebook. Це найпоширеніший підхід у вітчизняних продуктових командах малого і середнього розміру. Базовий аналіз робиться у Jupyter, з прямим SQL-запитом до сховища, обчисленням z-тесту або байєсівської моделі через румс, побудовою графіків через plotly. Перевага - повна свобода. Недоліки - несумісність між аналітиками, відсутність аудиторського журналу, неможливість використання нетехнічними

стейкхолдерами, відсутність захисту від підглядання за замовчуванням. У нашій команді саме така ситуація і була до початку роботи над цим дипломом: для кожного нового тесту створювався новий ноутбук, шаблон якого копіювався у когось із попередніх аналізів. Через місяць половина шаблонів вже мала локальні правки, які ніхто не зводив назад у спільний код.

	Відкритий код	Peeking-safe sequential CI	Паралельний запуск 4 методів	Monte Carlo для early-stop	BMM-висновок з рубриком	Інтеграція зі сховищем (без enter)	Перевірка SRM	Сегментний аналіз з BH-FDR
GrowthBook (OSS)	✓	✓	✗	✗	✗	✓	~	~
Statsig	✗	✓	✗	✗	✗	✗	✓	✓
Optimizely	✗	✓	✗	✗	✗	✗	~	~
Eppo	✗	✓	✗	✗	✗	~	✓	✓
LaunchDarkly Experiments	✗	✗	✗	✗	✗	✗	✗	✗
analytics-toolkit (Georgiev)	~	✓	✗	✗	✗	✗	✗	✗
Власні Jupyter-ноутбуки	✓	✗	✗	✗	✗	✓	✗	✗
Запропоноване рішення	✓	✓	✓	✓	✓	✓	✓	✓

✓ повністю підтримується / з обмеженнями
 ~ частково / з обмеженнями
 ✗ не підтримується

Лише запропоноване рішення поєднує всі вісім критеріїв одночасно

Рисунок 1.3 - Порівняльна матриця функціональності платформ для A/B-тестування

Зіставивши перелічені рішення, бачимо прогалину. Жодне з них не пропонує одночасно:

- а) відкриту імплементацію захищеного від підглядання послідовного CI;
- б) паралельний запуск кількох методів для крос-валідації;
- в) Монте-Карло-валідацію евристичних правил раннього припинення;
- г) висновок від LLM з аналітичним рубриком і обмеженим простором рішень;

д) інтеграцію з власним сховищем даних без корпоративного тарифу.

Ця прогалина і визначає інженерну й наукову цінність роботи. Запропоноване рішення поєднує сильні сторони GrowthBook у частині відкритості й послідовного аналізу, analytics-toolkit у частині методики GST repeated CI, Stucchio і VWO у частині decision-theoretic-байєсівського підходу з власною інновацією - шестишлюзовим інтегратором на основі LLM, який досі не реалізований у жодній порівнюваній платформі. Доцільність застосування великих мовних моделей у системах підтримки прийняття рішень аргументовано в нещодавній оглядовій публікації [30], яка формулює окрему дослідницьку програму для цього напрямку. Принципи композиції LLM-агентів і workflow, реалізовані у нашому рішенні, спираються на інженерні рекомендації Anthropic [31] щодо відмінності між жорсткими workflow та автономними агентами.

1.4 Контекст об'єкта практики та вхідні дані

Об'єктом практики є Futurra Group - продуктова B2C-компанія, що працює як квіз-воронка з веб-додатком. Користувач проходить опитування, бачить персоналізований офер, оформлює підписку. Модель монетизації - рекурентна підписка. Команда продукту регулярно тестує гіпотези у двох частинах продукту: у квіз-воронці (формулювання запитань, послідовність екранів, дизайн) і на офер-сторінці (цінова стратегія, формулювання переваг, дизайн заклику до дії). Один тест триває 1-4 тижні залежно від обсягу трафіку та розміру очікуваного ефекту.

Базою даних є BigQuery-проект musesacademy. Аналітично-агрегаційні дані живуть у датасеті ma_app. Єдиним джерелом істини для всіх A/B-тестів є широка таблиця user_product_stats_marketing_tests - приблизно 600 колонок з одним рядком на користувача. У ній зведено timestamp-колонки кожного кроку воронки (first_offerpage_view, first_initcheckout, first_form_submit, first_purchase), значення варіантів для кожного активного тесту (<test_name>), а також

сегментні атрибути (age, relationship_status, country, geo_group, landing_localization). Таблиця партиційована за часом. Це критично, бо запит без фільтрації за датою сканує понад 100 ГБ. Запит з фільтром за десятиденним вікном - близько 30-50 МБ.

Метадані тестів зберігаються окремо у dim_marketing_tests_alert. Для кожного тесту вказано базову інформацію: цільову метрику, плановий розмір вибірки на варіант (sample_per_variation), очікуваний ефект (uplift), дату активації, дату завершення моніторингу. Ця таблиця є точкою входу для системи: користувач обирає тест зі списку, а система автоматично підставляє правильні фільтри та налаштування.

У продукті визначено дві канонічні воронки. Квіз-воронка має вісім кроків: FirstView, FirstClick, QuizFinished, EmailStore, Offerpage, Initcheckout, FormSubmit, Subscriber. Офер-воронка - чотири кроки: Offerpage, Initcheckout, FormSubmit, Subscriber. Користувач вважається таким, що сконвертувався на кроці S , якщо timestamp-колонка для кроку $S + 1$ не NULL. Праймарна KPI - конверсія в Subscriber, рахується від першого кроку (FirstView для квіза, Offerpage для офер-воронки).

Сегменти, за якими виконується аналіз неоднорідності ефекту, такі: вікові групи (quiz_age з біннінгом за десятирічними інтервалами), сімейний статус (quiz_relationship_status), географічна група (geo_group), країна (top-12 за обсягом плюс інша), локалізація лендингу (landing_localization), а також бізнес-флаг “Recently brokeup”. Останній обчислюється з трьох квіз-відповідей, що ідентифікують користувачів у фазі розриву стосунків. Цей сегмент особливо релевантний для воронок з ex-back-наративом, тому система автоматично активує його як ключовий, якщо назва тесту містить ключові слова exback, missu, miss_you.

Для офлайн-демонстрації системи передбачено альтернативний шлях - завантаження CSV/Parquet-файлу з тими самими стовпцями (eventid, variant, quiz_first_page_view_time, флаги досягнення кроків воронки). У data_samples/ зберігається дві демо-вибірки. Перша - синтетична з відносним приростом +5%.

Друга - деідентифікована з реального тесту `marketing_test_offerpage_before_after` за 2025-09-01..14 з приростом +12,02% на конверсії в Subscriber. Друга вибірка є основним демонстраційним сценарієм у роботі, оскільки відображає реальну, а не змодельовану структуру шуму. Деідентифікація реалізована через SHA-256-хешування `eventid`, перемішування `country/geo/locale` і контрольовану інверсію 3% бінарних кроків, з калібруванням фінального приросту до цільового +12%. У вихідних даних до деідентифікації було близько 16 873 рядки, baseline-конверсія в Subscriber становила 11,87% у керівній групі та 13,29% у тестовій - саме ці числа використовуватимуться у демонстраційних сценаріях у третьому і четвертому розділах.

Захист від витрат на BigQuery реалізовано на чотирьох рівнях. Перший - обов'язкова фільтрація за партиційним стовпцем. Другий - `column-prune`, тобто витягування 5-15 колонок із 600. Третій - серверна агрегація: KPI та лічильники по кроках воронки повертаються вже агрегованими, що дає менше 10000 рядків на типовий тест. Четвертий - жорстке обмеження у 5 ГБ на запит, реалізоване через попередній `dry-run`-режим BigQuery API перед кожним фактичним запитом. На реальних даних тесту `marketing_test_missu_offer_pricing_block` за десять днів агрегаційний запит сканує близько 10 МБ, користувацько-рівневий запит - близько 32 МБ.

Окремо важлива для впровадженого рішення перевірка SRM (Sample Ratio Mismatch). Тести в системі перемикачів функцій задаються з очікуваним співвідношенням варіантів, найчастіше 50/50. У реальних даних воно може систематично відхилятися: помилки у розкатці, кешування трафіку, нерівномірне навантаження CDN. Fabijan та колеги [21] запропонували формальний χ^2 -тест для виявлення SRM з рекомендованим пороговим р-значенням 0,001 (значно жорсткішим за стандартний 0,05). При спрацюванні SRM ніякий аналіз метрик не виконується, бо будь-який ефект може бути артефактом нерівномірного розподілу. У нашій реалізації SRM-перевірка є першим шлюзом перед сегментним аналізом і блокує всі downstream-висновки.

Розмежування джерел даних, з якими працюватиме система, важливе з огляду на методологію. Дані тестів - з продакшен-таблиці BigQuery (отримані через query API). Метадані - з dim_marketing_tests_alert (теж BigQuery). Власні спостереження та виміри ефективності системи - з логування часу виконання та частоти попадань у кеш у локальній теці cache/. Літературні дані - у бібліографії, оформленій за ДСТУ 8302:2015. У третьому й четвертому розділах роботи всі чотири джерела використовуватимуться разом. Це принципово для дотримання вимог методички про розмежування походження даних.

Інженерно-обчислювальні обмеження теж важливі. Аналіз тесту має укладатися у бюджет 50 мс на чотири методи з $n = 10^5$. На практиці це досягається через векторизовану реалізацію бутстрепа (один виклик numpy.random.binomial для всієї матриці $R \times n$ замість циклу за індексами), нативну Python-реалізацію меж Howard-Ramdas (без зовнішньої бібліотеки confseq, що потребує cmake) і кешування метаданих через @st.cache_data з прив'язкою до SHA-256-хешу від (test_name, date_range, sql_template). Час виконання чотирьох методів на 16000-рядковій вибірці у виміряному режимі - менш як 15 мс.

1.5 Постановка завдань дослідження

Виходячи з огляду методів і програмних рішень, а також з опису об'єкта практики, формулюються чотири завдання дослідження. Кожне з них відповідає заголовку наступних розділів роботи.

а) Провести системний аналіз методів статистичного аналізу A/B-тестів і готових програмних рішень для їх застосування у продуктивній практиці; виявити обмеження кожного з них; обґрунтувати необхідність гібридного підходу, що поєднує паралельний запуск чотирьох незалежних методів з єдиним зобов'язувальним правилом рішення на основі послідовного CI. Це завдання розв'язано у поточному розділі.

б) Сформулювати методику крос-валідації статистичних методів і обґрунтувати математичні моделі для гібридної архітектури. Виписати формули кожного з чотирьох методів у формі, придатній для програмної реалізації. Обґрунтувати вибір GST repeated CI як основного режиму послідовного аналізу. Описати правила прийняття рішення на основі довірчого інтервалу. Сформулювати методику Монте-Карло-валідації евристичних правил раннього припинення. Розробити аналітичний рубрик для інтегратора на основі LLM з шістьма шлюзами. Це завдання розв'язується у другому розділі.

в) Розробити архітектуру інформаційної системи й реалізувати її як веб-застосунок на основі Streamlit. Система паралельно запускає чотири статистичні методи на спільному вхідному контракті, виконує сегментний аналіз з BH-FDR-корекцією, генерує висновок LLM зі структурованим виходом, веде журнал переглядів аналітика. Окремий блок завдання - забезпечити покриття реалізації тестами pytest, які зіставляють чисельні результати з аналітичними або літературними еталонами. Це завдання розв'язується у третьому розділі.

г) Провести експериментальну перевірку реалізації на деідентифікованих даних реального тесту `marketing_test_offerpage_before_after` MusesAcademy. Оцінити ефективність впровадження за критеріями скорочення часу аналітика, зменшення обсягу BigQuery-сканування, очікуваного зниження частоти хибно-позитивних рішень. Описати техніко-економічне обґрунтування використання системи на об'єкті практики. Це завдання розв'язується у четвертому розділі.

Виконання чотирьох завдань становить структуру наступних трьох розділів і збігається з вимогами методики написання КРМ у частині узгодженості заголовків розділів із завданнями дослідження. Логічно завдання вкладені одне в одне: системний аналіз дає матеріал для обґрунтування математичних моделей, обґрунтування моделей дає матеріал для архітектури, архітектура - для експериментальної перевірки. Така послідовність відповідає типовій структурі робіт у галузі систем підтримки прийняття рішень, описаній у підручнику Бідюка, Тимощука, Коваленка і Коршевніюка [26].

Висновок до першого розділу. Поточна продуктова практика аналізу A/B-тестів у вітчизняних B2C-компаніях страждає на три структурні проблеми: методологічну неоднорідність між аналітиками, інфляцію помилки I роду через підглядання та відсутність аудиторського журналу. Жодне з готових ринкових рішень - GrowthBook, Statsig, Optimizely, Eppo, analytics-toolkit - не розв'язує всіх трьох одночасно. Серед чотирьох сімейств статистичних методів (класичний, байєсівський, послідовний, бутстреп) кожен має свої сильні й слабкі сторони. Жоден не виграє за всіма критеріями захисту від підглядання, інтерпретованості, швидкості й гнучкості. Звідси впливає головна архітектурна гіпотеза дослідження: гібридне рішення з паралельним запуском чотирьох методів для крос-валідації і єдиним зобов'язувальним правилом рішення на основі GST repeated CI потенційно дає кращий компроміс за будь-яку готову платформу. Запропонована робота заповнює прогалину гібридною інформаційною системою, ядро якої - чотирьох-методний паралельний аналіз, захищене від підглядання зобов'язувальне правило та інтегратор на основі LLM, обмежений шестишлюзовим аналітичним рубриком. Цей інтегратор блокує рекомендацію про впровадження без проходження всіх верифікаційних шлюзів. Обґрунтування математичних моделей цього рішення, методика крос-валідації методів та опис аналітичного рубрика інтегратора розгорнуто у другому розділі.

РОЗДІЛ 2. МАТЕМАТИЧНІ МОДЕЛІ ТА МЕТОДИКА КРОС-ВАЛІДАЦІЇ ГІБРИДНОГО АНАЛІЗУ А/В-ТЕСТІВ

2.1 Архітектурний принцип гібридної системи та логічна модель прийняття рішення

Аналіз готових платформ у першому розділі показав, що жодна з них не пропонує одночасно захищеного від підглядання довірчого інтервалу, паралельного запуску декількох методів і керованого аналітичним рубриком висновку від великої мовної моделі. У цьому розділі обґрунтуємо математичні моделі для гібридного рішення, опишемо методику крос-валідації методів і правила прийняття рішення.

Базовий архітектурний принцип системи можна сформулювати у двох тезах. Перша - паралельний запуск чотирьох статистичних методів на одних і тих самих даних дає крос-валідаційний сигнал, якого не існує при використанні одного методу. Якщо чотири незалежні методи на тих самих числах сходяться у напрямку і значущості ефекту, рівень упевненості у висновку зростає; якщо розходяться, це сигнал зупинитися і дослідити причину. Друга теза - довірчий інтервал, побудований на засадах послідовного аналізу, є єдиним зобов'язувальним правилом прийняття рішення. Інші методи виконують роль “другої думки” і не можуть переважити висновок послідовного інтервалу, якщо вони на ньому суперечать. Така асиметрія потрібна тому, що класичний z-тест і бутстреп не калібровані для безпечного підглядання, а байєсівський аналіз стійкий до підглядання лише частково.

Логічна модель прийняття рішення розгортається у шість послідовних шлюзів. Кожен шлюз має статус “пройдено”, “не пройдено” або “невизначено” і блокує перехід до наступного. Перший шлюз - перевірка досягнутого обсягу вибірки відносно запланованого. Якщо вибірка менш ніж на 25% від плану, висновки про впровадження автоматично блокуються. Другий шлюз - евристичні правила раннього припинення тесту. Якщо тест уже у визначені моменти часу не показав мінімально очікуваного приросту, далі шлюзи можуть

лише підтвердити рекомендацію не впроваджувати. Третій шлюз - власне послідовний довірчий інтервал. Він приймає одне з трьох рішень: припинити з перевагою (stop_efficacy), припинити через марність (stop_futility) або продовжувати (continue). Четвертий шлюз - крос-перевірка методів. Якщо чотири методи розходяться у напрямку або значущості, висновок переходить у стан “потребує додаткового аналізу”, не у “впровадити”. П’ятий шлюз - перевірка SRM, тобто співвідношення розмірів керівної й тестової груп. Якщо співвідношення статистично значущо відхиляється від запланованого, аналіз метрик блокується повністю. Шостий шлюз - сегментний аналіз. Якщо ефект на загальній вибірці нейтральний, але є сегменти з різноспрямованими ефектами, рекомендація змінюється на “дослідити неоднорідність”.

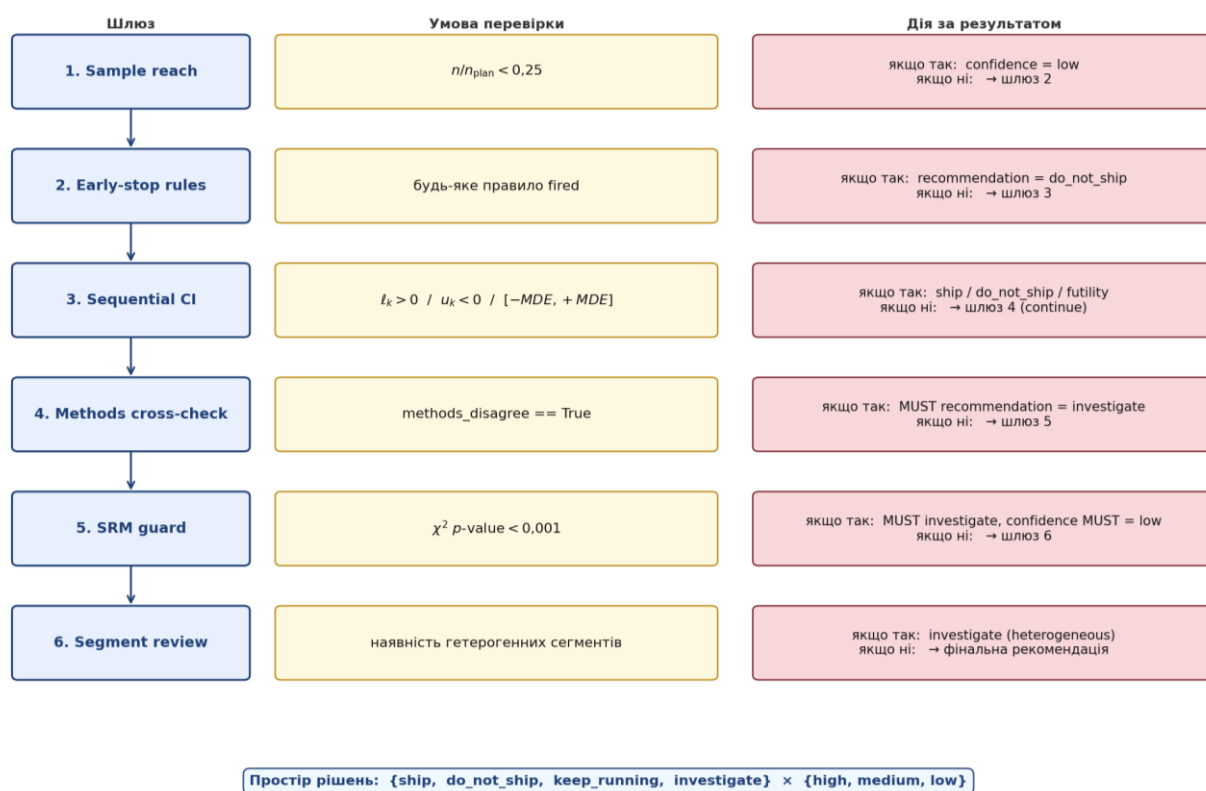


Рисунок 2.1 - Логічна модель шестишлюзового прийняття рішення

Послідовність шлюзів спроектована так, щоб кожен наступний не міг переважити висновок попереднього. Якщо шлюз 1 не пройдено, інтегратор не може запропонувати рішення про впровадження, навіть якщо шлюзи 2-6

виглядають сприятливо. Це і є суть жорсткого аналітичного рубрика. Архітектурне обґрунтування такого підходу узгоджується з принципом композиції моделей у системах підтримки прийняття рішень [26]: один метод страхує помилки іншого, а їхня агрегована оцінка є надійнішою за оцінку кожного окремо.

Контракт даних, спільний для всіх чотирьох методів, фіксує входи і виходи у вигляді трьох структурованих об'єктів. Об'єкт `ExperimentData` містить дві бінарні послідовності - результати конверсії для керівної й тестової груп, опціональну коваріату для коригування за методом CUPED, розклад моментів проміжного перегляду, апріорні параметри байєсівської моделі та цільовий розмір вибірки на варіант. Об'єкт `AnalysisResult` повертається кожним методом окремо і містить точкову оцінку приросту, довірчий або правдоподібний інтервал, метод-специфічну статистику (z-значення для класичного, апостеріорну ймовірність переваги для байєсівського, межі ефективності для послідовного, бутстреп-розподіл для бутстрепа), мітку рішення і час виконання. Об'єкт `MultiMethodReport` агрегує результати чотирьох методів і додатково містить результат перевірки SRM та підсумок сегментного аналізу. Такий контракт зафіксований у програмному коді як набір Pydantic-моделей, що гарантує сумісність між модулями і дає змогу LLM-інтегратору працювати з єдиним джерелом даних.

Часовий бюджет на повний аналіз тесту обмежено 50 мілісекундами на вибірках до ста тисяч спостережень. Це жорстке обмеження є необхідною умовою для інтерактивної роботи аналітика з системою, бо реалізація очікувано виконується у Streamlit-додатку, де кожна перебудова сторінки супроводжується перерахунком метрик. Усі чотири методи реалізовано з векторизованим обчисленням на NumPy і виконуються паралельно у пулі з чотирьох робочих потоків. Реальний час виконання на вибірках 16 тисяч рядків - менш як 15 мілісекунд.

Логіка перетікання даних між модулями системи задана у вигляді конвеєра з трьох послідовних стадій. Стадія завантаження: метадані тесту

читаються з таблиці `dim_marketing_tests_alert`, агрегаційні дані по воронці - з `user_product_stats_marketing_tests`. Обидва BigQuery-запити проходять через жорсткий ліміт сканування у 5 ГБ і двошарову систему кешування (`@st.cache_data` плюс парк `parquet`-файлів на диску з SHA-256-ключами). Стадія обчислення: завантажені дані передаються в `ExperimentData`, який далі надходить у пул чотирьох робочих потоків з результатом у `MultiMethodReport`. Стадія агрегації: `MultiMethodReport` плюс SRM-перевірка і сегментний аналіз передаються в шаблонний формувач (детермінований Jinja-шаблон) і опційно в інтегратор на основі LLM. Жодна стадія не блокує іншу, але і не перекриває її логіку: контракт між стадіями жорстко зафіксований у Pydantic-моделях, і будь-яка зміна одного модуля не повинна тягти за собою каскадних змін у решті системи.

Ще один наскрізний принцип системи - незалежність методів від типу інтерфейсу. Усі чотири статистичні модулі і інтегратор реалізовано без жодного імпорту Streamlit; це домен-шар системи. Стрімлітовий шар розташований окремо, у каталозі `app/`, і взаємодіє з домен-шаром виключно через контракти `ExperimentData` / `MultiMethodReport`. Така архітектура дозволяє у майбутньому замінити Streamlit на інший інтерфейс (наприклад, FastAPI з фронтендом на React або CLI-утиліту) без зміни жодного рядка статистичних методів. Це класична схема *hexagonal architecture*, описана у літературі з системотехніки; вона не часто застосовується у вузько-аналітичних інструментах, але її переваги для подальшого масштабування системи на інші продукти екосистеми ми оцінюємо як значні.

2.2 Математичні моделі статистичних методів

У цьому підрозділі викладено математику кожного з чотирьох методів у формі, яка безпосередньо переноситься у програмний код. Послідовність викладу - класичний підхід, байєсівський, послідовний (як основний інструмент прийняття рішення), бутстреп.

2.2.1 Класичний z-тест двох пропорцій з опціональною CUPED-корекцією

Розглянемо вибірки керівної й тестової груп розмірів n_c і n_t із кількістю конверсій s_c і s_t відповідно. Точкові оцінки конверсії - $\hat{p}_c = s_c/n_c$ та $\hat{p}_t = s_t/n_t$. Нульова гіпотеза $H_0: p_c = p_t$, альтернативна - двостороння. z-статистика тесту обчислюється з об'єднаною (pooled) оцінкою дисперсії, що відповідає припущенню H_0 :

$$z = \frac{\hat{p}_t - \hat{p}_c}{\sqrt{\hat{p}_{\text{pool}}(1 - \hat{p}_{\text{pool}}) \cdot \left(\frac{1}{n_c} + \frac{1}{n_t}\right)}}, \quad \hat{p}_{\text{pool}} = \frac{s_c + s_t}{n_c + n_t}. \quad (2.1)$$

Двостороннє p-значення обчислюється як $p = 2(1 - \Phi(|z|))$, де Φ - функція стандартного нормального розподілу. Класична робота, з якої виводиться загальна теорія розв'язків і оцінювання у частотному підході, належить Вальду [8]; у вітчизняній літературі формули для перевірки гіпотез про різницю двох пропорцій подано, зокрема, у посібнику Бідюка, Терентьєва і Просянкіної-Жарової [27].

Довірчий інтервал для абсолютної різниці пропорцій будується з роздільною (unpooled) дисперсією, що дає коректне покриття поза припущенням нульової гіпотези:

$$\hat{p}_t - \hat{p}_c \pm z_{1-\alpha/2} \cdot \sqrt{\frac{\hat{p}_c(1-\hat{p}_c)}{n_c} + \frac{\hat{p}_t(1-\hat{p}_t)}{n_t}}. \quad (2.2)$$

Інтервал для відносного приросту обчислюється через метод дельта - стандартну похибку для $\hat{p}_t/\hat{p}_c - 1$ оцінюємо як стандартну похибку чисельника, поділену на $\hat{p}_c$, з подальшим симетричним інтервалом навколо точкової оцінки. Цей підхід коректний при достатньо великих n_c , n_t і не надто близьких до нуля

$\hat{p}_c$, що відповідає типовим продуктовим тестам із базовою конверсією у межах 5-30%.

CUPED-корекція використовується опціонально, якщо для кожного користувача доступна довірчо-незалежна попередня коваріата x - найчастіше показник конверсії за період до експерименту. Метод запропоновано Deng, Xu, Kohavi і Walker [20]. Скоригована метрика обчислюється як

$$y' = y - \theta \cdot (x - \bar{x}), \quad \theta = \frac{\text{cov}(y, x)}{\text{var}(x)}. \quad (2.3)$$

За умови, що x незалежна від варіанту тесту, y' має таке саме математичне сподівання, як y , але меншу дисперсію - приблизно у $1/(1 - \rho^2)$ разів, де ρ - вибіркова кореляція y з x . На реальних даних В2С-продукту з помірною кореляцією pre-period конверсії з post-period (типове значення $\rho \approx 0,4 - 0,5$) це дозволяє скоротити необхідний розмір вибірки на 25-40%. У реалізації передбачено перевірку незалежності x від варіанту: якщо $|\text{corr}(x, \text{variant})| > 0,05$, корекція не застосовується, а в журналі записується попередження.

Стратегію застосування CUPED у нашій системі обрано консервативною. Корекція вмикається лише для тестів з достатньо довгим pre-period (мінімум 14 днів), бо коротший період не дає надійної оцінки θ . Окрім того, корекцію не застосовуємо до метрик, де y є самою конверсією бінарного типу і де x також бінарний: у такому випадку кореляція рідко перевищує 0,1, і вигреш у дисперсії не покриває витрат на обчислення. Найкраще CUPED працює на квазі-неперервних агрегатах, як-от показник годин активності або обсяг витрат за період. Для повноцінного впровадження CUPED у MusesAcademy потрібно мати стабільні pre-period-метрики; ця інфраструктура частково присутня, але до повного покриття поки що далеко.

Перевірка коректності реалізації класичного блоку виконана у трьох напрямках. По-перше, числове порівняння з еталоном

statsmodels.stats.proportion.proportions_ztest: z-статистика і р-значення збігаються до шостого десяткового знака. По-друге, симуляція рівня помилки I роду під H_0 : 200 синтетичних вибірок з однаковими параметрами для керівної й тестової груп; емпірична частка відхилення H_0 потрапляє в інтервал $[0,02,0,10]$ при номінальному $\alpha = 0,05$, що відповідає очікуваному джиттеру для скінченного числа симуляцій. По-третє, перевірка інтервалу через крос-порівняння з statsmodels.stats.proportion.confint_proportions_2indep: межі довірчого інтервалу для абсолютної різниці збігаються до $1e-6$.

2.2.2 Байєсівська модель Beta-Binomial

Байєсівський аналіз будується на спряженій парі Beta-Binomial. Для кожного варіанту тесту параметр конверсії $\theta \in [0,1]$ розглядається як випадкова величина з апіорним розподілом $\theta \sim \text{Beta}(\alpha_0, \beta_0)$. Після спостереження s конверсій із n спроб апостеріорний розподіл є знову бета-розподілом:

$$\theta \mid s, n \sim \text{Beta}(\alpha_0 + s, \beta_0 + n - s). \quad (2.4)$$

Замовчуванням нейтральним апіором є рівномірний $\text{Beta}(1,1)$. Для типового продуктового тесту з $n > 10^3$ спостережень вплив апіору на форму апостеріору практично нульовий, що знімає типову критику байєсівських методів про “довільність вибору пріору”.

Ключова метрика, яку дає байєсівський підхід і не дає класичний, - апостеріорна ймовірність переваги тестового варіанту над керівним:

$$P(B > A) = \int_0^1 \int_0^{\theta_t} f_t(\theta_t) f_c(\theta_c) d\theta_c d\theta_t, \quad (2.5)$$

де f_c і f_t - щільності апостеріорних розподілів керівної й тестової груп. Закриту форму цього інтеграла через інкрементальну суму бета-функцій подано в україномовному посібнику Бондаренка, Кравченка і Сологуба [9, с. 18-23], що слугував первинним джерелом для імплементації байєсівського блоку. У замкненій формі

$$P(B > A) = \sum_{i=0}^{\alpha_b-1} \frac{B(\alpha_a+i, \beta_a+\beta_b)}{(\beta_b+i) \cdot B(1+i, \beta_b) \cdot B(\alpha_a, \beta_a)}, \quad (2.6)$$

де $B(\cdot, \cdot)$ - бета-функція, α_a, β_a - апостеріорні параметри керівної групи, α_b, β_b - тестової. Сума скінченна, обчислюється швидко і не вимагає Монте-Карло-семплювання. Для чисельної стабільності при $\alpha + \beta > 10^5$ суму обчислюємо у логарифмічній шкалі через `scipy.special.betaln`, з резервним переходом на чисельне інтегрування `scipy.integrate.quad`, якщо логарифмічна сума втрачає точність.

Другою важливою байєсівською метрикою є очікувані втрати при виборі певного варіанту. У теорії статистичних рішень, перенесеній на А/В-тестування у white paper VWO [10], очікувані втрати при виборі тестового варіанту визначаються як

$$E[\text{loss}(B)] = E[\max(\theta_c - \theta_t, 0)] = \int \int_{\theta_c > \theta_t} (\theta_c - \theta_t) f_c(\theta_c) f_t(\theta_t) d\theta_c d\theta_t. \quad (2.7)$$

Аналогічно для керівного варіанту. Очікувані втрати інтерпретуються як середній абсолютний приріст, яким аналітик ризикує пожертвувати, якщо помилково обере варіант, що насправді гірший. У реалізації цей інтеграл обчислюється через Монте-Карло-семплювання 20 тисяч точок із апостеріорних розподілів, бо при невеликих апостеріорах закрыта форма стає чисельно складною.

Правдоподібний інтервал для різниці апостеріорних параметрів будується як квантілі 2,5% і 97,5% розподілу $\Delta = \theta_t - \theta_c$, отримані з тих самих Монте-

Карло-семплів. Це принципово відрізняється від класичного довірчого інтервалу: правдоподібний інтервал має пряму ймовірнісну інтерпретацію - “з імовірністю 95% істинна різниця лежить у межах інтервалу”, - тоді як класичний інтервал має лише “frequentist coverage” інтерпретацію.

Важлива інженерна особливість байєсівського блоку: для headline-метрик (показ “rate_control”, “rate_treatment”, “rel_lift” на сторінці аналітика) використовуються емпіричні оцінки $\hat{p}_c$, $\hat{p}_t$, а не апостеріорні середні. Інакше при невеликих n байєсівський блок показував би “зсунутий” відносний приріст порівняно з класичним блоком через регуляризацію апіором, що заплутувало б аналітика. Цей нюанс - типовий приклад “виявлених недоліків у методах обробки інформації” з вимог методички; його розв’язання у нашій імплементації - явне розмежування виводу метрик для звіту і метрик для прийняття рішення.

Вибір апіору заслуговує окремого обговорення. Дефолтний Beta(1,1) є рівномірним розподілом на $[0,1]$ і не несе жодної інформації про очікуване значення конверсії. Це безпечний вибір, коли немає історичних даних, але він не оптимальний для продуктового тесту, в якому базова конверсія добре відома. Для подальшого розвитку системи доцільним є інформативний апіор Beta(α_0, β_0) з $\alpha_0/(\alpha_0 + \beta_0) \approx p_c^{\text{baseline}}$ і помірною концентрацією. Однак ми свідомо не вмикаємо цей режим за замовчуванням, бо при сильно помилковій оцінці p_c^{baseline} результат тесту може бути зміщеним; рішення про вибір апіору має ухвалюватися індивідуально для кожного тесту і потребує осмисленої документації, що поки що не вкладається у бюджет нашого пайплайну.

Перевірка коректності байєсівського блоку виконана у чотирьох тестових сценаріях. Перший - апостеріорне середнє Beta-розподілу збігається з аналітичним значенням $\alpha/(\alpha + \beta)$ до $1e-3$ на синтетичних даних. Другий - екстремальні значення $P(B > A)$: для дуже відмінних апостеріорів метрика прямує до 0,999 або 0,001 відповідно (передбачуваність на межах). Третій - симетричний випадок: при ідентичних апостеріорних розподілах $P(B > A)$ становить точно 0,5 з допуском $1e-3$. Четвертий - порівняння замкненої форми

Бондаренка з чисельним інтегруванням через `scipy.integrate.dblquad` на маленьких параметрах: збіг до $1e-6$.

2.2.3 Послідовний аналіз: GST repeated CI як основний інструмент рішення

Послідовний аналіз у нашій системі реалізує три комплементарні техніки: групове послідовне тестування з функціями α -spending Лана і ДеМетса [12], повторні довірчі інтервали GST repeated CI як основне правило прийняття рішення, і часо-рівномірні довірчі послідовності Howard, Ramdas, McAuliffe і Sekhon [2] як опційний режим. У цьому підрозділі описуємо математичну основу всіх трьох і обґрунтовуємо вибір GST repeated CI як режиму за замовчуванням.

Функція α -spending [12] задає, як бюджет помилки I роду α розподіляється між проміжними переглядами тесту. Для дискретного розкладу з K переглядів частка інформації на k -му перегляді - $t_k = n_k/N$, де n_k - накопичений розмір вибірки, N - запланований. Кумулятивна частка витраченого α задається функцією $\alpha^*(t)$ з умовами $\alpha^*(0) = 0$, $\alpha^*(1) = \alpha$. Дві найвживаніші форми - O'Brien-Fleming і Pocock:

$$\alpha_{\text{OBF}}^*(t) = 2 - 2\Phi\left(\frac{z_{\alpha/2}}{\sqrt{t}}\right), \quad \alpha_{\text{Pocock}}^*(t) = \alpha \cdot \ln(1 + (e - 1) \cdot t). \quad (2.8)$$

Функція O'Brien-Fleming консервативна на ранніх переглядах і витрачає більшу частину α на кінець тесту. Функція Pocock рівномірно розподіляє α між переглядами. У реалізації за замовчуванням використовується OBF, бо вона зберігає номінальну потужність на повному обсязі вибірки майже без втрат. Дидактичний виклад обчислення меж за різними α -spending-функціями подано у Lakens, Pahlke і Wassmer [13]; чисельні значення меж у нашій реалізації звіряються з пакетом `gract`, описаним у тому ж туторіалі.

З α -spending-функції на кожному перегляді k обчислюється приріст витрачання $\Delta\alpha_k = \alpha^*(t_k) - \alpha^*(t_{k-1})$, з якого через інверсію двостороннього нормального розподілу отримується z -межа ефективності z_k^{eff} . Якщо на перегляді k модуль z -статистики тесту перевищує z_k^{eff} , тест припиняється з висновком про значущий ефект.

Перейдемо до основного інструменту. GST repeated CI - повторний довірчий інтервал, який на кожному перегляді k обчислюється за формулою

$$\hat{p}_t^{(k)} - \hat{p}_c^{(k)} \pm z_k^{\text{eff}} \cdot \sqrt{\frac{\hat{p}_c^{(k)}(1-\hat{p}_c^{(k)})}{n_c^{(k)}} + \frac{\hat{p}_t^{(k)}(1-\hat{p}_t^{(k)})}{n_t^{(k)}}}. \quad (2.9)$$

Конструкція забезпечує те, що при правильному α -spending-розкладі сімейство всіх інтервалів на всіх K переглядах має одночасну довірчу ймовірність $1 - \alpha$. Це і є властивість захищеного від підглядання довірчого інтервалу на запланованих переглядах. Той самий рушій використовується у Spotify [14] і Statsig [25], що додатково підтверджує його придатність для виробничого застосування.

Альтернативно, повністю часо-рівномірна довірна послідовність Howard, Ramdas, McAuliffe і Sekhon [2] базується на класі sub-Gaussian Gaussian-mixture меж:

$$r_n(\alpha; \rho) = \sqrt{2\sigma^2 \cdot \frac{n\rho^2+1}{n^2\rho^2} \cdot \ln \frac{\sqrt{n\rho^2+1}}{\alpha}}, \quad (2.10)$$

де σ^2 - оцінка дисперсії, ρ - параметр, що тюнить ширину інтервалу на запланованому розмірі вибірки. Інтервал має time-uniform-валідність: він покриває істинну різницю з імовірністю $1 - \alpha$ одночасно для всіх $n \in \mathbb{N}$. Це найсильніша властивість серед усіх трьох технік. Її ціна - значно ширший інтервал на запланованому розмірі вибірки.

Кількісно порівняння трьох технік на ілюстративній вибірці $n_c = n_t = 8000$ при $\hat{p}_c \approx \hat{p}_t \approx 0,10$ і $\alpha = 0,05$ подане в таблиці 2.1.

Таблиця 2.1 - Порівняння ширин довірчих інтервалів для абсолютної різниці пропорцій на $n = 8000$

Метод	Ширина інтервалу, п.п.	Захист від підглядання
Класичний Wald CI	$\pm 0,93$	відсутній
GST repeated CI (10 переглядів, OBF)	$\pm 1,2$	у запланованих переглядах
Howard-Ramdas always-valid CS ($\rho = 1/\sqrt{N}$)	$\pm 2,8$	повний time-uniform

З таблиці видно, що GST repeated CI лише на 30% ширший за Wald, тоді як always-valid CS - утричі ширший. Ця різниця і визначає вибір GST repeated CI як режиму за замовчуванням. Always-valid CS залишається доступним як опційний режим для випадків незапланованих переглядів - наприклад, коли аналітик хоче подивитися на дані поза розкладом, який було зафіксовано на старті тесту.

Правила прийняття рішення на основі GST repeated CI на поточному перегляді k формуються так. Нехай $[\ell_k, u_k]$ - інтервал для абсолютної різниці, MDE - заздалегідь визначений мінімальний значущий ефект. Тоді:

- якщо $\ell_k > 0$: припинити тест із висновком про перевагу тестового варіанту (stop_efficacy);
- якщо $u_k < 0$: припинити тест із висновком про перевагу керівного (stop_efficacy зі зворотнім знаком);
- якщо $-MDE < \ell_k$ і $u_k < +MDE$: припинити через марність (stop_futility);
- інакше: продовжувати збір даних (continue).

Така схема узгоджується з принципами групового послідовного дизайну в [13] і дозволяє завершувати безперспективні тести швидше без шкоди для гарантій помилки I роду.

Для ілюстрації поведінки меж OBF на типовому розкладі з $K = 10$ переглядів і $\alpha = 0,05$ значення z-меж ефективності за дев'ятьма проміжними і одним фінальним переглядами наведені у таблиці 2.2.

Таблиця 2.2 - Z-межі ефективності за функцією O'Brien-Fleming для 10 запланованих переглядів при $\alpha = 0,05$

Перегляд k	Частка інформації t_k	Z-межа ефективності z_k^{eff}
1	0,10	6,75
2	0,20	4,72
3	0,30	3,82
4	0,40	3,30
5	0,50	2,93
6	0,60	2,67
7	0,70	2,47
8	0,80	2,30
9	0,90	2,16
10	1,00	2,03

Видно, як межі поступово знижуються до значення, близького до класичного $z_{0,975} = 1,96$, але не співпадають із ним на фінальному перегляді - залишається консервативний “відступ”, який і забезпечує сім'ю всіх десяти інтервалів сумарною довірчою ймовірністю 0,95. Якщо аналітик орієнтується на наївне порівняння, не використовуючи цих меж, ризик хибно-позитивного припинення на ранніх переглядах є реальним: на $t = 0,10$ номінальний z-поріг 1,96 потрапляє у “зону прозвичайки”, а коректний поріг становить 6,75.

Параметр ρ у формулі Howard-Ramdas заслуговує окремого обговорення. За замовчуванням у літературі рекомендують $\rho = 1$, що дає консервативну ширину інтервалу. У нашій реалізації використовується автотюнінг: $\rho = 1/\sqrt{N}$, де N - запланований розмір вибірки на варіант. Такий вибір мінімізує ширину інтервалу на запланованому N і дає результат, конкурентний з GST repeated CI. Без автотюнінгу always-valid CS на $n = 8000$ давав би ширину приблизно $\pm 4,5$ п.п. замість $\pm 2,8$ п.п.; перехід до $\rho = 1/\sqrt{N}$ скорочує її приблизно на третину.

2.2.4 Бутстреп з ВСа-корекцією для бінарних метрик

Бутстреп є непараметричною технікою оцінки розподілу статистики через ресемплінг. Для бінарних метрик класична схема Ефрона [16] полягає в тому, що з вихідної вибірки R разів беруть ресемпли того ж розміру з поверненням і для кожного обчислюють точкову оцінку $\hat{\theta}_r^*$. Довірчий інтервал отримують як квантилі емпіричного розподілу $\hat{\theta}_1^*, \dots, \hat{\theta}_R^*$.

Для бінарних метрик дискретна структура спостережень дозволяє реалізувати бутстреп значно ефективніше. Замість ресемплінгу за індексами потрібно лише генерувати дві біноміальні послідовності розміру $R \times 1$ з параметрами n_c , $\hat{p}_c$ для керівної і n_t , $\hat{p}_t$ для тестової груп. Такий векторизований підхід реалізовано в нашому модулі через `numpy.random.binomial`, що дає приблизно десятикратне прискорення відносно загальної реалізації через `arch.bootstrap.IIDBootstrap`. На вибірці 16 тисяч рядків з $R = 2000$ ресемплів час виконання - близько 8-12 мілісекунд.

ВСа-інтервал, запропонований Ефроном [16], додає дві поправки до простого процентильного інтервалу - зміщення z_0 і прискорення $\hat{a}$. Зміщення визначається як

$$z_0 = \Phi^{-1} \left(\frac{\#\{r: \hat{\theta}_r^* < \hat{\theta}\}}{R} \right), \quad (2.11)$$

тобто це квантиль стандартного нормального розподілу, що відповідає частці ресемплів, для яких оцінка нижча за вибірккову. Прискорення обчислюється через джекнайф-оцінки L_i - значення статистики, обчисленої з вилученням i -го спостереження з вибірки:

$$\hat{a} = \frac{\sum_{i=1}^n (\bar{L} - L_i)^3}{6 \cdot [\sum_{i=1}^n (\bar{L} - L_i)^2]^{3/2}}, \quad \bar{L} = \frac{1}{n} \sum_{i=1}^n L_i. \quad (2.12)$$

Скориговані квантилі обчислюються як

$$\alpha_1 = \Phi\left(z_0 + \frac{z_0 + z_{\alpha/2}}{1 - \hat{a}(z_0 + z_{\alpha/2})}\right), \quad \alpha_2 = \Phi\left(z_0 + \frac{z_0 + z_{1-\alpha/2}}{1 - \hat{a}(z_0 + z_{1-\alpha/2})}\right). \quad (2.13)$$

Кінцевий інтервал - квантилі α_1 і α_2 емпіричного розподілу ресемплів. У вироджених випадках, коли знаменник $1 - \hat{a}(z_0 + z_{\alpha/2})$ близький до нуля, реалізація відкочується до простого процентильного інтервалу, з фіксацією попередження у журналі. Така деталь критично важлива для production-вжитку: на дуже малих вибірках або при відсутності варіативності ВСа-формула чисельно нестабільна, і автоматичний відкат запобігає некоректним інтервалам.

Перевірка коректності реалізації виконана у двох напрямках. Перший - на великих вибірках ($n > 10^4$) ВСа-інтервал має збігатися з Wald-інтервалом до третього-четвертого десяткового знака; таке порівняння реалізовано у тесті `test_bootstrap_agrees_with_wald_on_large_n`. Другий - симуляція покриття під істинним значенням параметра: 100 синтетичних вибірок з заданою p дають емпіричне покриття не нижче 94% при номінальному 95%, що відповідає коректній калібровці методу [17].

2.3 Методика крос-валідації методів та правила прийняття рішення

У підрозділі 2.2 описано чотири незалежні методи. У цьому підрозділі сформульовано методику їх крос-валідації - тобто як їх висновки агрегуються в одне рішення для аналітика.

Базовий принцип: кожен метод видає мітку рішення з фіксованого набору `{stop_efficacy, stop_futility, continue}`. Мітку формує сам метод за власною логікою. Для класичного z-тесту: значущий двосторонній p-value (нижче α) при позитивному знакові z-статистики дає `stop_efficacy`; інакше - `continue` (для z-тесту немає `futility`-логіки, бо це одноразовий тест). Для байєсівського: $P(B > A) > 0,95$ дає `stop_efficacy`, $P(B > A) < 0,05$ - той самий висновок зі зворотним

знаком, інакше - continue. Для послідовного: правила з підрозділу 2.2.3 безпосередньо дають одну з трьох міток. Для бутстрепу: інтервал, що не містить нуля у позитивну сторону, - stop_efficacy; інтервал у негативну сторону - те саме зі зворотним знаком; інакше - continue.

Зведена мітка для всього звіту обчислюється з пріоритетом послідовного методу: якщо послідовний метод видав stop_efficacy, ця мітка стає зведеною, незалежно від міток інших методів. Це і є вираз твердження “послідовний CI - єдине зобов’язувальне правило”. Якщо послідовний видав continue, перевіряється згода інших методів: усі три повинні видати continue або stop_efficacy в одному напрямку, інакше зведена мітка - disagree. Стан disagree транслюється в інтегратор як обов’язкова рекомендація “дослідити” замість “впровадити” або “не впроваджувати”.

Перевірка SRM - сім’я тестів узгодженості очікуваного й фактичного співвідношення розмірів груп - є окремим обов’язковим шлюзом. Формальний χ^2 -тест для SRM описано у Fabijan, Gurchur, Gupta та колег [21]; рекомендований ними поріг р-значення для діагностики SRM - 0,001, значно жорсткіший за стандартні 0,05 для тесту на ефект. Така жорсткість виправдана тим, що SRM є системною помилкою, яка має породжуватися рідко; коли вона спрацьовує, причина потребує негайного інженерного розслідування. У нашій реалізації спрацювання SRM повністю блокує downstream-аналіз: жодна метрика не виводиться до аналітика, лише діагностичне повідомлення про необхідність перевірити логіку розклатки.

Корекція на множинні порівняння активується автоматично при сегментному аналізі - там, де одночасно перевіряються $k > 1$ сегментів. Використовується процедура Benjamini-Hochberg для контролю частки хибних відкриттів [18]: відсортовані р-значення $p_{(1)} \leq p_{(2)} \leq \dots \leq p_{(k)}$ порівнюються з пороговими значеннями $\alpha \cdot i/k$, і відхиляються нульові гіпотези для всіх i , аж до найбільшого, для якого $p_{(i)} \leq \alpha \cdot i/k$. Вибір BH-FDR над Bonferroni або Holm зумовлений тим, що при $k > 10$ Bonferroni є надмірно консервативним, а Holm дає лише незначно більшу потужність; BH-FDR суттєво потужніший і

забезпечує контроль за параметром, який ближчий до інтуїтивного значення “як часто ми помиляємось серед оголошених знахідок”.

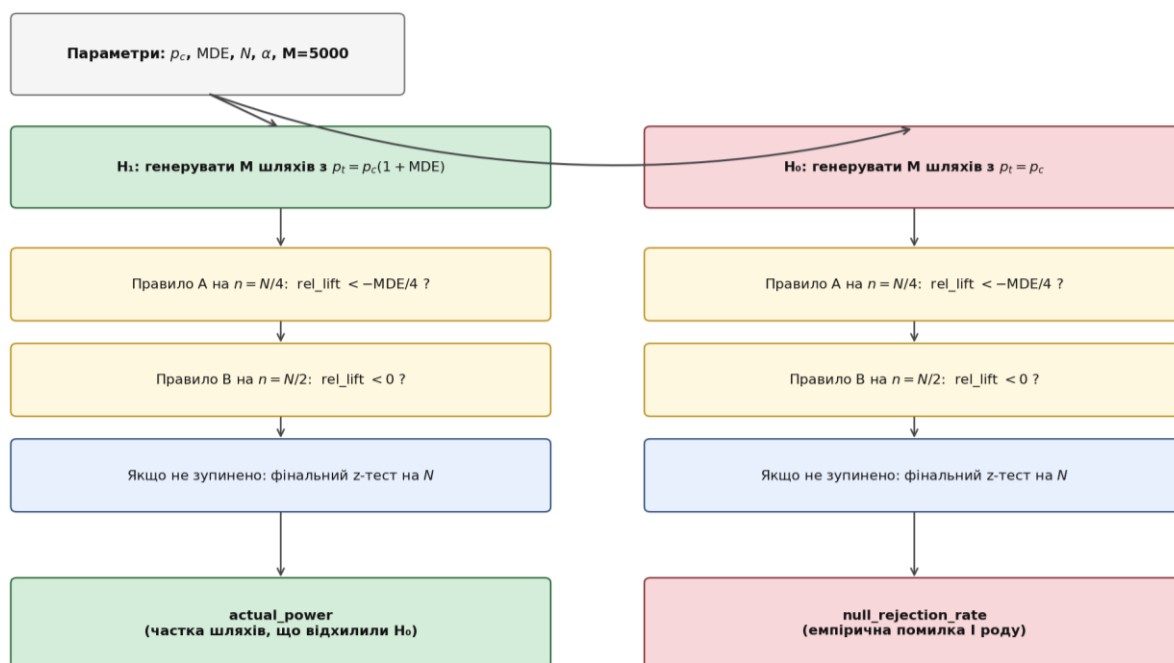
2.4 Методика Монте-Карло-валідації евристичних правил раннього припинення

У продуктивній практиці аналітики часто застосовують прості евристики раннього припинення тесту, що не пов’язані з формальним послідовним аналізом. Дві найпоширеніші - “правило чверті” і “правило половини”. Правило чверті: якщо на 25% обсягу запланованої вибірки спостережений відносний приріст менший за $-MDE/4$, тест зупиняється з висновком про марність. Правило половини: якщо на 50% вибірки спостережений відносний приріст менший за нуль, тест зупиняється. Обидва правила інтерпретовані для нетехнічного власника продукту, дешеві в обчисленні і не потребують реалізації α -spending.

Проблема таких евристик у тому, що вони вносять додаткову втрату потужності, яку ніхто не вимірює. У підручнику Бідюка, Тимощука, Коваленка і Коршевніюка [26] загальне правило побудови СППР - кожне рішення про вибір параметра має супроводжуватися кількісною оцінкою наслідків його застосування. Для евристичних правил раннього припинення це означає одне: потрібен інструмент, який кількісно вимірює, наскільки правило знижує потужність тесту і чи не порушує воно номінальний рівень помилки I роду. У нашій роботі цю функцію виконує Монте-Карло-симуляція, що є однією з заявлених наукових новизн.

Методика симуляції описується таким алгоритмом. На вхід подаються параметри тесту: p_c - базова конверсія, MDE - мінімальний значущий ефект (відносний), N - запланований розмір вибірки на варіант, α - рівень помилки I роду. На виході - чотири метрики, що характеризують властивості правил у поєднанні з фінальним z-тестом: `nominal_power` - потужність фінального z-тесту без жодних правил; `actual_power` - потужність фінального z-тесту з

застосуванням правил; $\text{null_rejection_rate}$ - емпірична частота відхилення H_0 при H_0 істинній (емпірична помилка I роду); stop_rates - частоти спрацювань кожного правила окремо під H_0 і під H_1 .



Додаткові метрики: nominal_power (без правил), stop_rate кожного правила під H_0 і H_1

Рисунок 2.2 - Схема Монте-Карло-симуляції втрати потужності евристичних правил раннього припинення

Алгоритм обчислення метрик:

а) генерується $M = 5000$ синтетичних шляхів тесту під альтернативною гіпотезою $H_1: p_t = p_c \cdot (1 + \text{MDE})$. Кожен шлях - дві бінарні послідовності довжини N з відповідними параметрами, що моделюють надходження користувачів у керівну й тестову групи;

б) для кожного шляху перевіряємо правило А на $N/4$: якщо спостережений відносний приріст менший за $-\text{MDE}/4$, шлях позначається як “зупинений А”;

в) для кожного шляху, що не зупинений А, перевіряємо правило В на $N/2$: якщо спостережений відносний приріст менший за 0, шлях позначається як “зупинений В”;

г) для решти шляхів виконуємо повний z-тест на N і фіксуємо результат H_0 відхилено / не відхилено;

д) `actual_power` - частка шляхів, для яких фінальний висновок - “ H_0 відхилено у позитивному напрямку”;

е) повторюємо кроки а-г з $p_t = p_c$ для оцінки `null_rejection_rate`. Цей сценарій - 5000 шляхів під H_0 - дає емпіричну помилку I роду з урахуванням правил;

ж) `nominal_power` обчислюється з тих самих шляхів під H_1 , але без застосування правил - тобто z-тест виконується на повному N для всіх шляхів.

Типовий результат симуляції для параметрів $p_c = 0,10$, $MDE = 0,07$, $N = 12,600$, $\alpha = 0,10$ представлено в таблиці 2.2.

Таблиця 2.2 - Результати Монте-Карло-симуляції для типових параметрів тесту

Метрика	Значення
<code>nominal_power</code> (без правил)	0,80
<code>actual_power</code> (з обома правилами)	0,78
<code>null_rejection_rate</code>	0,052
<code>stop_rate</code> правила А під H_1	0,013
<code>stop_rate</code> правила А під H_0	0,27
<code>stop_rate</code> правила В під H_1	0,008
<code>stop_rate</code> правила В під H_0	0,21

Інтерпретація результатів. Втрата потужності - 2 відсоткових пункти, що для типового тесту з потужністю 80% означає 2,5% зниження ймовірності виявити істинний ефект. Це прийнятна ціна за можливість швидко зупинити безперспективні тести: правило А під H_0 спрацьовує у 27% випадків, що означає, що приблизно кожен четвертий безефективний тест зупиняється на

чверті обсягу. Емпірична помилка I роду залишається у допустимих межах (0,052 при номінальному 0,05, перевищення 4% від номіналу не є практично значущим). Симетрія правил відносно нуля гарантує, що ймовірність хибно-позитивного “впровадити” не зростає.

Цей кількісний інструмент дозволяє продуктивній команді ухвалювати рішення про використання евристичних правил на основі чисел, а не інтуїції. Наскільки відомо автору, у відкритих платформах (GrowthBook, Statsig, Optimizely, Еppo) аналогічної функціональності не реалізовано: у них евристичні правила або відсутні, або встановлюються без супровідної кількісної оцінки втрати потужності. Це і є зміст заявленої наукової новизни - “удосконалено методику валідації евристичних правил раннього припинення через Монте-Карло-симуляцію”.

2.5 Сегментний аналіз з корекцією на множинні порівняння

Сегментний аналіз має на меті виявити, чи ефект тестового варіанту неоднорідний серед різних груп користувачів. Прикладами сегментних осей у MusesAcademy є вікові групи з біннінгом 18-25, 26-35, 36-45, 46+; сімейний статус (одинокий, у стосунках, нещодавно розлучений, інше); географічна група (північна Європа, центральна, південна, Латинська Америка, інше); країна (топ-12 за обсягом плюс агрегована “інша”); локалізація лендингу. Цей набір осей обрано емпірично за результатами попереднього аналізу того, які саме змінні найчастіше пояснюють неоднорідність ефекту в продуктових тестах MusesAcademy.

Простір пошуку сегментів складається з одновимірних сегментів (значення однієї осі - наприклад, “вік 26-35”) і двовимірних комбінацій двох осей (наприклад, “вік 26-35 і північна Європа”). Для шести осей з 5-15 категоріями це дає сотні комбінацій, для частини яких розмір сегмента занадто малий для статистично значущих висновків. Фільтр перед статистичним тестом:

розмір керівної й тестової груп у сегменті має бути не менший за 100 кожна, а сумарний розмір - не менший за 200.

Для кожного сегмента, що пройшов фільтр, виконується класичний z-тест за формулою з підрозділу 2.2.1. Отримане р-значення подається в процедуру Benjamini-Hochberg [18] для контролю частки хибних відкриттів. Алгоритм процедури: відсортовані р-значення сегментів $p_{(1)} \leq p_{(2)} \leq \dots \leq p_{(k)}$, де k - кількість сегментів, що пройшли фільтр; для кожного i обчислюється скоригований поріг $\alpha \cdot i/k$; знаходиться найбільший i^* , для якого $p_{(i^*)} \leq \alpha \cdot i^*/k$; усі сегменти з рангами $1, \dots, i^*$ позначаються як значущі.

Ранжування значущих сегментів для виводу аналітику здійснюється за композитною оцінкою:

$$\text{score} = |\Delta_{\text{abs}}| \cdot \ln(1 + n_{\text{total}}) \cdot (-\ln p_{\text{adj}}), \quad (2.14)$$

де Δ_{abs} - абсолютна різниця конверсій у сегменті, n_{total} - сумарний розмір сегмента, p_{adj} - скориговане р-значення після ВН. Така оцінка враховує три фактори одночасно: розмір ефекту, статистичну надійність і обсяг сегмента. Логарифмування розміру запобігає домінуванню найбільших сегментів, для яких ефект може бути малий, але формально значущий.

Перед запуском сегментного аналізу обов'язково перевіряється SRM на загальній вибірці. Якщо SRM спрацював, сегментний аналіз блокується повністю - бо при нерівномірному розподілі трафіку будь-який сегментний "ефект" є артефактом. Це принциповий шлюз: системою помилок при сегментному аналізі легко обманути аналітика, бо знайти "значущий" сегмент при достатньо великому просторі пошуку статистично легко навіть під H_0 .

Окремо реалізовано бізнес-флаг "Recently brokeup" - похідний сегмент, що обчислюється з відповідей квіза: користувач у фазі розриву стосунків. Цей сегмент особливо релевантний для тестів з ex-back-наративом і автоматично активується, якщо в назві тесту присутні ключові слова exback, missu, miss_you.

Реалізація як бізнес-флаг ілюструє загальніший принцип: похідні сегменти, що мають бізнес-сенса, повинні бути предметом окремої уваги в системі і не обмежуватися “сирими” квіз-полями.

Інтерпретація результатів сегментного аналізу для аналітика будується навколо трьох візуальних форм. Перша - “лісова” діаграма, де кожен рядок є одним сегментом, точкова оцінка приросту - центральна точка, довірчий інтервал - горизонтальний відрізок. Так аналітик одразу бачить, чи інтервал перетинає нуль і у який бік. Друга - теплова карта двовимірних комбінацій, де колір клітинки кодує знак абсолютної різниці, а яскравість - її модуль. Третя - деревовидна карта, де площа кожного прямокутника відповідає сумарному розміру сегмента, а колір - знаковому приросту. Така комбінація візуалізацій дозволяє виявляти неочевидні комбінації факторів за один-два погляди, на відміну від плоского списку, де неоднорідність ефекту легко пропустити.

Кореляція знахідок з бізнес-сенсом залишається відповідальністю аналітика. Система може лише підсвітити сегмент із статистично значущим відхиленням, але інтерпретація - чому саме у віковій групі 26-35 з північної Європи приріст відрізняється від загальної - потребує знання продукту. У третьому розділі покажемо приклад такої знахідки на реальному тесті MusesAcademy.

2.6 Аналітичний рубрик для інтегратора на основі великої мовної моделі

Усі попередні підрозділи описували математичні моделі окремих компонентів. У цьому підрозділі описано, як їхні результати агрегуються в одне рішення через велику мовну модель з аналітичним рубриком - і чому такий підхід, на відміну від простого правила if-else, має сенс.

Інтегратор на основі LLM реалізовано у дві стадії. Стадія А - детермінований Jinja-шаблон, який формує текстове резюме за результатами чотирьох методів, SRM-перевірки, сегментного аналізу і Монте-Карло-

симуляції правил. Стадія А завжди виконується і завжди дає той самий вихід для однакових вхідних даних. Її роль - забезпечити базовий звіт, який має бути доступний навіть якщо LLM-сервіс недоступний.

Стадія В - виклик LLM через OpenAI Responses API [5] зі структурованим виходом. На вхід LLM отримує всі агреговані метрики у форматі JSON: дизайн тесту, поточний прогрес, результати чотирьох методів, SRM-перевірку, топ-сегменти, результати Монте-Карло-симуляції правил. На виході - JSON-об'єкт із полями: рекомендація з обмеженого набору {ship, do_not_ship, keep_running, investigate}, рівень упевненості з набору {high, medium, low}, два-п'ять ключових спостережень, один-три виклики сегментів, до трьох ризиків, до трьох наступних дій. Структура виходу зафіксована у Pydantic-моделі, що автоматично конвертується в JSON-схему для передачі в API. OpenAI Structured Outputs гарантує, що повернутий JSON відповідає схемі зі стовідсотковою надійністю; це і робить можливим вбудовування LLM у виробничий конвеєр без post-processing-парсингу.

Ключова частина інтеграції - аналітичний рубрик у системному промтті LLM. Це жорстко формалізована шестишлюзова логіка, описана в підрозділі 2.1, у формі прямих інструкцій. Рубрик передбачає, що LLM перевіряє шлюзи у фіксованому порядку:

а) шлюз перевірки досягнутого обсягу. Якщо досягнуто менш ніж 25% планованої вибірки, упевненість обмежується значенням low; якщо менш ніж 50% - значенням medium;

б) шлюз евристичних правил раннього припинення. Будь-яке спрацювання правила переводить рекомендацію у do_not_ship, з винятком випадку напрямкової помилки, коли рекомендація - investigate;

в) шлюз послідовного CI. Якщо інтервал виключає нуль у позитивну сторону - ship; у негативну - do_not_ship; інтервал у зоні марності - do_not_ship; інакше - keep_running;

г) шлюз крос-перевірки методів. Якщо чотири методи не сходяться у напрямку або значущості, рекомендація MUST бути investigate, не ship і не do_not_ship;

д) шлюз SRM. Якщо SRM спрацював, рекомендація MUST бути investigate, упевненість MUST бути low;

е) шлюз сегментного огляду. Для тестів з ex-back-нарративом особлива увага до сегмента “Recently brokeup”; для будь-яких тестів - перевірка, чи немає сегмента з суттєво відмінним ефектом, що міг би змінити інтерпретацію.

Така структура промпта приводить LLM у режим, описаний у інженерному дописі Anthropic [31] як “workflow” - тобто послідовний детермінований ланцюжок із заздалегідь визначеним маршрутом, на противагу автономному агенту, який сам обирає послідовність кроків. Workflow дає передбачувані результати з високою точністю на обмеженому домені; саме це нам і потрібно для прийняття рішень з аналітичної точки зору.

Параметри виклику LLM зафіксовано: модель gpt-5.2 з резервним переходом на gpt-4o у разі недоступності, temperature = 0,2, seed = 42, незмінний системний промпт з фіксованою версією. Така конфігурація забезпечує високу частоту попадання у кеш промптів OpenAI і відтворюваність рекомендацій на ідентичних входах. Формулювання ключових спостережень і ризиків можуть варіюватися від виклику до виклику, але рекомендація, рівень упевненості і структура рішення стабільні.

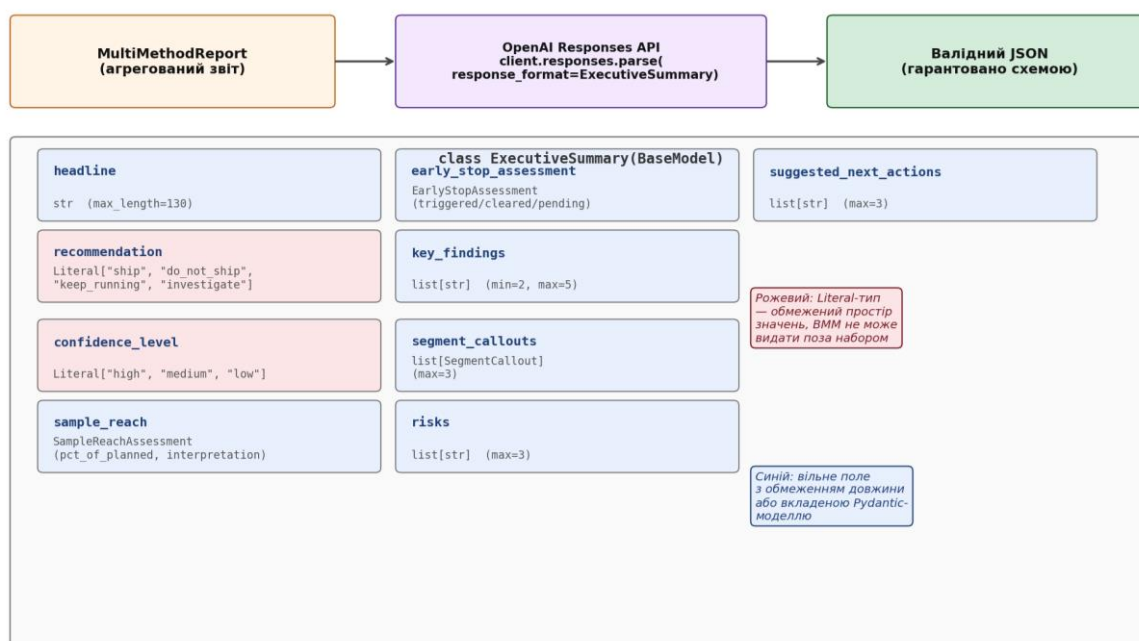


Рисунок 2.3 - Структура виходу інтегратора на основі LLM

Заборона для LLM - вільне формулювання рішення. Якщо в проміжному міркуванні LLM хоче запропонувати рішення ship, але хоча б один з шлюзів не пройдено, фінальний JSON-вихід має містити іншу мітку. Така архітектура - класичний приклад “оркестрованого” застосування LLM, де модель є генератором природньомовного пояснення, а не вільним суб’єктом рішення. Це і є відповідь на критику довільності висновків LLM у системах підтримки прийняття рішень, на яку звертається увага в недавніх оглядових публікаціях.

Окремою деталлю реалізації є button-gated виклик стадії В. LLM-виклик ніколи не запускається автоматично при перерендері сторінки, а лише за явним кліком кнопки “Згенерувати AI-резюме”. Це зроблено з трьох міркувань: економія токенів (типовий виклик коштує близько 0,02-0,04 долара, при автозапуску на кожен рендер це швидко накопичується); психологічна окреслення межі між детермінованою аналітикою і генеративним висновком (аналітик свідомо запитує думку LLM); і захист від некоректного використання моделі - якщо аналітик ще не довів аналіз до завершення, висновок LLM може закріпити проміжне рішення як фінальне. Кнопка з підтвердженням створює невелику паузу, яка дає змогу свідомо повторно перевірити дані перед запитом.

Окремий важливий аспект - параметр `peeking_safe`, що передається в інтегратор як булеве значення. Якщо аналітик повідомляє системі, що поточний перегляд є позаплановим (тобто він “підглянув” поза розкладом), `peeking_safe = false` змушує LLM обмежити рівень упевненості на одну позицію нижче і додати у блок ризиків окреме застереження про методологічну ймовірну хибу. Це гарний приклад того, як проста бінарна підказка з боку аналітика може суттєво змінити висновок системи без зміни математичних розрахунків.

Резервна модель для випадку недоступності основної реалізована на рівні API-клієнта. Якщо `gpt-5.2` повертає 404 або 5xx-помилку, клієнт автоматично перемикається на `gpt-4o` і повторює запит з тим самим JSON-схемним обмеженням. Якщо й резервна модель недоступна, інтегратор повертає тільки результат стадії A (детермінований шаблон), а в інтерфейсі з'являється повідомлення про неможливість згенерувати LLM-резюме на цей момент. Така деградаційна логіка означає, що базовий аналіз ніколи не блокується через проблеми зовнішнього сервісу.

Висновок до другого розділу

У розділі викладено математичні моделі чотирьох статистичних методів, методику їх крос-валідації та аналітичний рубрик для інтегратора на основі великої мовної моделі. Класичний підхід використовує z-тест двох пропорцій з опціональною CUPED-корекцією для зниження дисперсії. Байєсівський підхід будується на спряженій моделі Beta-Binomial із замкненою формою апостеріорної ймовірності переваги тестового варіанту. Послідовний аналіз реалізовано у двох режимах: GST repeated CI з функціями α -spending Лана-ДеМетса як основний інструмент рішення і часо-рівномірна довірча послідовність Howard-Ramdas як опційний режим для незапланованих переглядів. Бутстреп з BCa-корекцією дає непараметричну оцінку невизначеності з векторизованою реалізацією для бінарних метрик.

Сформульовано методику крос-валідації методів: послідовний CI - єдине зобов'язувальне правило, інші три методи виконують функцію крос-перевірки. Розходження методів автоматично переводить рекомендацію у стан “потребує додаткового дослідження”. Перевірка SRM як перший шлюз блокує downstream-аналіз у разі нерівномірності трафіку. Корекція BH-FDR контролює частку хибних відкриттів при сегментному аналізі.

Запропоновано Монте-Карло-методику кількісної оцінки втрати потужності і збереження номінальної помилки I роду для евристичних правил раннього припинення. Це інструмент, відсутній у відкритих платформах GrowthBook, Statsig, Optimizely і Eppo, і одна з заявлених наукових новизн роботи.

Описано шестишлюзову архітектуру інтегратора на основі LLM. Модель отримує жорстко формалізований аналітичний рубрик у системному промпті і структурований формат виходу через OpenAI Structured Outputs. Така структура переводить LLM у режим workflow з передбачуваним поведінкою, на противагу автономному агенту.

Викладені моделі і методики безпосередньо переносяться у програмний код, опис якого є предметом третього розділу.

РОЗДІЛ 3. ПРОГРАМНА РЕАЛІЗАЦІЯ ГІБРИДНОЇ СИСТЕМИ ТА МОДЕЛЮВАННЯ НА РЕАЛЬНИХ ДАНИХ

3.1 Технологічний стек, архітектура системи і контракти

Програмну реалізацію запропонованої гібридної системи виконано як веб-додаток на базі Streamlit з Python 3.13. Цей стек обрано з трьох причин. По-перше, Streamlit дає одночасно швидкий цикл розробки і native multi-page-структуру, що відповідає логіці системи з восьми сторінок аналізу. По-друге, Python-екосистема статистичних обчислень (NumPy, SciPy, statsmodels) дозволяє реалізувати усі чотири методи з паралельного розділу 2 без зовнішніх залежностей. По-третє, BigQuery-клієнт у Python нативно підтримує pandas-gbq-інтерфейс і Storage Read API, що критично для роботи з виробничою таблицею проєкту musesacademy.

Архітектура системи побудована за принципом шарової відокремленості (рисунок 3.1). Верхній шар - інтерфейс на Streamlit з вісьмома сторінками, кожна з яких відповідає окремому етапу аналізу. Середній шар - доменний, без жодного імпорту Streamlit. Він містить чотири статистичні модулі, сегментний аналіз, перевірку SRM, реалізацію CUPED і Монте-Карло-симуляцію раннього припинення тесту. Нижні два шари - шар даних з BigQuery-клієнтом і кешем паркет-файлів та шар великої мовної моделі з OpenAI Responses API і Pydantic-схемою для структурованого виходу.



Рисунок 3.1 - Шарова архітектура системи

Така архітектура має дві важливі властивості. По-перше, статистичні методи можна тестувати, не запускаючи Streamlit-додаток - досить імпортувати `core.stats` і викликати `analyze()` на синтетичних даних. По-друге, у майбутньому можна замінити Streamlit на будь-який інший інтерфейс (наприклад, FastAPI з фронтендом на React, або CLI-утиліту для batch-режиму) без переписування жодного рядка статистичних методів. Це класична hexagonal architecture, і її переваги для подальшого масштабування системи на інші продукти екосистеми ми оцінюємо як визначальні для довгострокової підтримки.

Контракти даних між шарами зафіксовано у вигляді Pydantic-моделей. Це гарантує валідацію типів на межах модулів і автоматичну серіалізацію в JSON для передачі в LLM. Три ключові моделі:

```
class ExperimentData(BaseModel):
    control: list[int] # 0/1 outcomes
    treatment: list[int]
    metric_name: str = "conversion"
    test_name: str
    control_label: str
    treatment_label: str
    # опціональні поля:
    control_covariate: list[float] | None # для CUPED
    treatment_covariate: list[float] | None
    looks: list[int] | None # для sequential
```

```
bootstrap_n_resamples: int = 2000
bayesian_prior_alpha: float = 1.0
bayesian_prior_beta: float = 1.0
mde_relative: float = 0.05
target_n_per_arm: int | None
```

Об'єкт ExperimentData будується один раз з вихідних даних і передається в усі чотири методи. Це гарантує, що методи дивляться на ті самі числа, а не на різні агрегації того самого тесту - типова прихована помилка у фрагментованих ноутбукових пайплайнах, де один аналітик рахує конверсію за одним визначенням події, інший - за дещо іншим.

```
class AnalysisResult(BaseModel):
    method: Literal["frequentist", "bayesian", "sequential", "bootstrap"]
    n_control: int; n_treatment: int
    rate_control: float; rate_treatment: float
    abs_lift: float; rel_lift: float
    abs_lift_interval: IntervalEstimate
    rel_lift_interval: IntervalEstimate | None
    p_value: float | None
    z_statistic: float | None
    decision: Literal["stop_efficiency", "stop_futility", "continue"]
    library: str; runtime_seconds: float
    # метод-специфічні поля:
    prob_b_beats_a: float | None
    expected_loss_b: float | None
    expected_loss_a: float | None
    posterior_samples: list[float] | None
    test_statistic_path: list[float] | None
    efficacy_boundary: list[float] | None
    futility_boundary: list[float] | None
    bootstrap_distribution: list[float] | None
```

Кожен метод повертає AnalysisResult з однаковими спільними полями і своїми унікальними розширеннями. Завдяки Literal-типу для method і decision IDE-підказки і Pydantic-валідація запобігають надсиланню некоректних значень у downstream-логіку.

```
class MultiMethodReport(BaseModel):
    test_name: str; product_part: str; funnel_step: str
    results: dict[Method, AnalysisResult]
    segments: list[Segment]
    srm_check: SRMResult | None

    @property
    def methods_disagree(self) -> bool:
        """True якщо методи розходяться по значущості або напрямку."""
```

Зведений звіт `MultiMethodReport` агрегує результати чотирьох методів і додаткові перевірки. Властивість `methods_disagree` обчислюється на льоту і використовується як шлюз 4 у логіці інтегратора. Якщо хоча б два з чотирьох методів видають протилежні мітки `decision`, або їхні точкові оцінки `abs_lift` мають протилежні знаки, повертається `True`.

Потік даних від запиту до висновку аналітика розгортається у п'ять стадій (рисунок 3.2). На першій стадії модуль `core.data` тягне дані з `BigQuery`: спершу метадані тесту з `dim_marketing_tests_alert`, потім агрегаційну таблицю з `user_product_stats_marketing_tests` з обов'язковим партиційним фільтром за `quiz_first_page_view_time`. Кожен запит проходить через `dry-run`-перевірку, яка оцінює обсяг сканування і відкидає запити понад 5 ГБ. На другій стадії з отриманих рядків будується `ExperimentData`. На третій стадії викликається паралельний `runner`, що запускає чотири `analyze()`-функції у пулі робочих потоків. На четвертій усі результати агрегуються у `MultiMethodReport`. На п'ятій (опційно) звіт надсилається в інтегратор на основі LLM.



Часовий бюджет: BigQuery 1-3 с (cold), кеш < 0,1 с; чотири методи < 15 мс; ВММ-виклик 2-4 с

Рисунок 3.2 - Потік даних від запиту до висновку

Часовий бюджет на повний цикл закладено у 50 мілісекунд для статистичних методів і до 3-4 секунд для запиту до `BigQuery` з холодним кешем. Реальні виміри на тестовій вибірці 16 873 рядки показують 12-14 мс на

чотирирівневий аналіз і близько 1,5 секунди на BigQuery-запит з кеш-промахом. При наявності закешованих parquet-файлів повторне завантаження виконується менш як за 100 мс.

3.2 Реалізація статистичного ядра

Чотири статистичні методи з розділу 2 реалізовано у каталозі `core/stats/`. Кожен метод оформлено як клас, що реалізує спільний протокол `StatMethod`: метод `analyze(data: ExperimentData) -> AnalysisResult`. Така уніфікація дозволяє паралельному runner-у не знати внутрішньої логіки методів і працювати лише з контрактом.

Класичний z-тест двох пропорцій реалізовано в `core/stats/frequentist.py`. Логіка проста: рахуються суми конверсій у керівній і тестовій групах, обчислюється pooled-оцінка дисперсії, потім z-статистика і двостороннє p-значення. Wald-довірчий інтервал для абсолютної різниці будується з unpooled-дисперсією. Окремо реалізовано CUPED-корекцію через метод `_cuped_adjust`, що бере коваріату з `ExperimentData.control_covariate` і `treatment_covariate`, обчислює θ через `pumpru.cov` і повертає скориговані метрики. Перевірка корисності CUPED - якщо $|\text{corr}(x, \text{variant})| > 0,05$, корекція не застосовується і фіксується попередження.

Валідація реалізації виконана через зіставлення з `statsmodels.stats.proportion.proportions_ztest`: z-статистика і p-значення збігаються до $1e-6$. Симуляція рівня помилки I роду під H_0 на 200 синтетичних вибірках з однаковими параметрами для двох груп: емпірична частка відхилення H_0 потрапляє в інтервал $[0,02,0,10]$ при номінальному $\alpha = 0,05$, що відповідає очікуваному джиттеру для скінченної кількості симуляцій.

Байєсівський модуль `core/stats/bayesian.py` реалізує спряжений висновок Beta-Binomial з закритою формою $P(B > A)$ за посібником Бондаренка, Кравченка і Сологуба [9]. Ключова функція `_prob_beta_b_greater_a` обчислює суму бета-функцій у логарифмічній шкалі через `scipy.special.betaln`, з резервним

переходом на чисельне інтегрування через `scipy.integrate.quad` при $\alpha + \beta > 5000$. Очікувані втрати обчислюються Монте-Карло-семплюванням 20 тисяч точок з апостеріорних розподілів. Правдоподібний інтервал для різниці апостеріорів - квантілі 2,5% і 97,5% емпіричного розподілу $\theta_t - \theta_c$.

Важливий інженерний нюанс: для виведення `rate_control`, `rate_treatment` і `rel_lift` у звіт використовуються емпіричні оцінки $\hat{p}_c$, $\hat{p}_t$, а не апостеріорні середні. Це усуває “зсув” між байєсівським і класичним блоками, який інакше плував би аналітика при невеликих n . Перевірка реалізації включає чотири тестові сценарії: апостеріорне середнє Beta-розподілу збігається з аналітичним $\alpha/(\alpha + \beta)$ до $1e-3$; $P(B > A)$ для дуже відмінних апостеріорів прямує до $0,999/0,001$; для симетричного випадку $P(B > A) = 0,5 \pm 10^{-3}$; закрита форма порівнюється з чисельним інтегруванням через `dblquad` до $1e-6$.

Послідовний модуль `core/stats/sequential.py` є найбільш об’ємним і реалізує три комплементарні техніки. Функція `alpha_spending` обчислює $\alpha^*(t)$ за O’Brien-Fleming або Pocock. Функція `gst_efficacy_boundaries` повертає z -межі для кожного запланованого перегляду через інверсію двостороннього нормального розподілу. Функція `gst_repeated_diff_ci` будує повторний довірчий інтервал для різниці пропорцій із цими межами - це і є основний інструмент прийняття рішення. Альтернативна функція `always_valid_diff_ci` реалізує часо-рівномірну довірчу послідовність Howard-Ramdas з автотюнінгом параметра $\rho = 1/\sqrt{N}$.

Валідація послідовного блоку: межі OBF на повному перегляді $t = 1$ збігаються з $z_{1-\alpha/2}$ до $1e-3$; межі на $t = 0,1$ перевищують 3,0 (літературне значення для одностороннього $\alpha = 0,05$); `always-valid CI` на 100 синтетичних шляхах під H_0 має емпіричне покриття не нижче 95%, що підтверджує `time-uniform`-валідність методу.

Бутстреп-модуль `core/stats/bootstrap.py` реалізовано через векторизовану генерацію біноміальних послідовностей замість загальної схеми ресемплінгу за індексами. Це дає приблизно десятикратне прискорення на бінарних метриках

відносно бібліотеки `arch.bootstrap`. ВСа-корекція реалізована буквально за формулами Efron [16]: зміщення z_0 , прискорення $\hat{a}$ через джекнайф, скориговані квантілі α_1, α_2 . Деградаційна логіка передбачає відкат на простий процентильний інтервал при чисельній нестабільності ВСа-формули.

Паралельний runner у `core/stats/runner.py` побудовано на `concurrent.futures.ThreadPoolExecutor` з пулом чотирьох робочих потоків. Послідовність виклику (рисунок 3.3): сторінка UI викликає `data_access.build_data()`, що повертає `ExperimentData`; runner передає об'єкт у чотири `analyze()`-функції одночасно; на виході - словник з чотирма `AnalysisResult`; потім runner будує `MultiMethodReport` з SRM-перевіркою та сегментним аналізом і повертає його сторінці UI або інтегратору.



Рисунок 3.3 - Послідовність виклику паралельного runner-а

Реальний час виконання на тестовій вибірці 16 873 рядки: класичний z-тест - 0,3 мс; байєсівський - 11 мс (через Монте-Карло-семплювання 20 тисяч точок); послідовний - 1,8 мс; бутстреп - 9 мс (з 2000 ресемплів). Загалом при паралельному запуску - 11-13 мс (визначається байєсівським модулем як найдовшим). У послідовному запуску загальний час склав би близько 22 мс.

3.3 Реалізація сегментного аналізу, SRM і CUPED

Сегментний аналіз реалізовано у `core/segments/discovery.py`. Простір пошуку формується з шести базових осей: вік (з біннінгом 18-29, 30-39, 40-49, 50+), сімейний статус, географічна група, країна (топ-12 + “інша”), локалізація лендингу, бізнес-флаг “Recently brokeup”. Алгоритм перебирає всі 1D-сегменти (значення однієї осі) та 2D-комбінації двох різних осей. Перед статистичним тестом застосовується фільтр: у сегменті має бути не менш як 100 спостережень у керівній і 100 у тестовій групі.

Для кожного сегмента, що пройшов фільтр, обчислюється z-тест через ту саму функцію, що використовується у класичному блоці. Отримані p-значення подаються в процедуру Benjamini-Hochberg [18] через функцію `adjust_pvalues` з `core/stats/corrections.py`. Реалізація BH-FDR векторизована: масив відсортованих p-значень множиться на k/i , потім беруться cumulative-минимальні значення з кінця і обрізаються до 1,0. Це класична реалізація з двома проходами по масиву і складністю $O(k \log k)$ за рахунок сортування.

Ранжування значущих сегментів виконується за композитною оцінкою $|\Delta_{\text{abs}}| \cdot \ln(1 + n_{\text{total}}) \cdot (-\ln p_{\text{adj}})$. Топ-30 за цією оцінкою виводяться аналітику. Логарифмування n_{total} запобігає домінуванню найбільших сегментів, для яких ефект може бути малим, але формально значущим. Логарифмування p_{adj} виконує аналогічну функцію в зворотному напрямі - сегменти з $p < 10^{-4}$ не отримують непропорційно високу оцінку відносно тих з $p \approx 0,01$.

Перевірка SRM реалізована у `core/segments/srm.py` як простий χ^2 -тест на однорідність очікуваного і фактичного співвідношення розмірів груп. Очікуване співвідношення береться з параметрів тесту в `dim_marketing_tests_alert`; якщо у метаданих його немає, використовується 50/50. Рекомендований Fabijan та колегами [21] поріг $p < 0,001$ застосовується як замовчуваний. У реалізації SRM-перевірка - перший шлюз перед сегментним аналізом: якщо вона провалюється, downstream-аналіз блокується повністю, бо

при нерівномірності трафіку будь-який сегментний “ефект” може бути артефактом.

Бізнес-флаг “Recently brokeup” реалізовано як похідну колонку, що обчислюється на льоту з відповідей квіза:

```
recently_brokeup = (
    (quiz_what_your_goals_eng_only == "get_my_ex_back") |
    (quiz_relationship_status == "Recently broke up") |
    (quiz_who_do_you_want_attract == "My ex")
)
```

Цей флаг автоматично включається в сегментну вісь, якщо назва тесту містить ключові слова `exback`, `missu` або `miss_you`. Реалізація - функція `is_exback_test(test_name) -> bool`, що повертає булеве значення. Принцип ширший: похідні сегменти, що мають бізнес-сене, повинні бути предметом окремої уваги в системі, а не звичайних “сирих” квіз-полів.

CUPED-корекція реалізована у `core/stats/cuped.py` і викликається з `frequentist`-модуля при наявності `control_covariate` і `treatment_covariate` у `ExperimentData`. Перевірка перед застосуванням - обчислення $|\text{corr}(x, \text{variant})|$. Якщо вище 0,05, корекція не застосовується. Це додаткова страховка проти ненавмисного зсуву оцінки ефекту через приховану кореляцію коваріати з варіантом.

3.4 Реалізація інтегратора на основі великої мовної моделі

Інтегратор на основі LLM реалізовано у каталозі `core/llm/`. Архітектурно він складається з чотирьох модулів: клієнт OpenAI Responses API, Pydantic-схема `ExecutiveSummary`, версіонований системний промпт, і власне функція-сумаризатор з двома стадіями.

Стадія А - детермінований Jinja-шаблон у `core/llm/templates/`. Він рендериться завжди, незалежно від доступності LLM. На вхід - `MultiMethodReport`; на виході - markdown-текст з основними числами тесту, мітками рішення кожного методу, перевіркою SRM, топ-сегментами. Час

рендеру - менш як 50 мс. Стадія А гарантує, що аналітик ніколи не залишається без базового звіту через проблеми з зовнішнім сервісом.

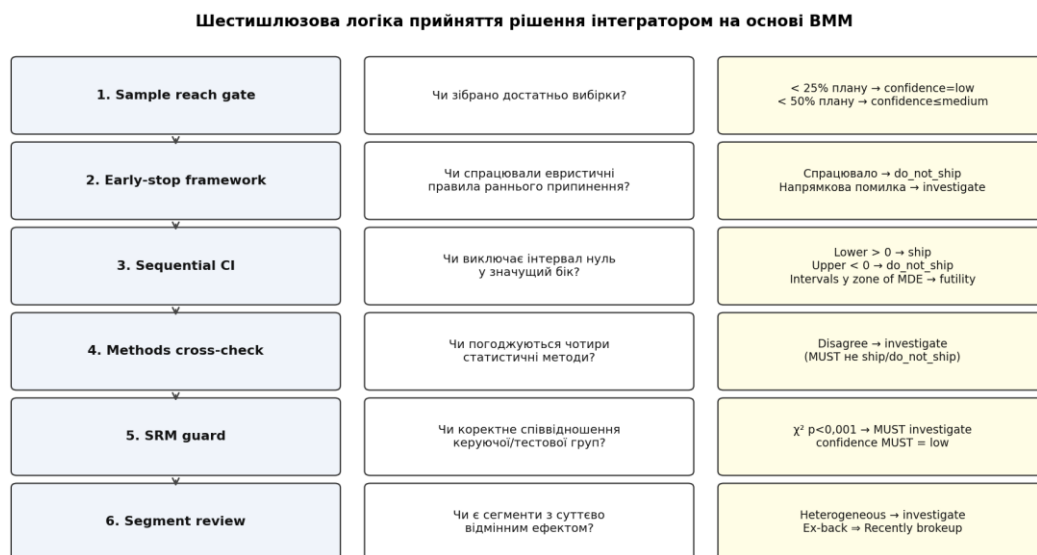
Стадія В - виклик OpenAI Responses API через `client.responses.parse()` з `response_format=ExecutiveSummary`. Pydantic-схема автоматично конвертується у JSON-схему, передану в API. OpenAI Structured Outputs [5] гарантує, що повернений JSON відповідає схемі. Це і робить можливим вбудовування LLM у виробничий конвеєр без подальшого парсингу регулярними виразами.

Pydantic-схема має таку структуру:

```
class ExecutiveSummary(BaseModel):
    headline: str = Field(max_length=130)
    recommendation: Literal["ship", "do_not_ship", "keep_running", "investigate"]
    confidence_level: Literal["high", "medium", "low"]
    sample_reach: SampleReachAssessment
    early_stop_assessment: EarlyStopAssessment
    key_findings: list[str] = Field(min_length=2, max_length=5)
    segment_callouts: list[SegmentCallout] = Field(max_length=3)
    risks: list[str] = Field(max_length=3)
    suggested_next_actions: list[str] = Field(max_length=3)
```

Поле `recommendation` з обмеженим набором значень - ключове. Воно гарантує, що LLM не може видати рекомендацію поза заздалегідь визначеним простором. Поле `confidence_level` дискретизує упевненість у три рівні, що уникнено вільного формулювання типу “95% sure”. Підструктури `SampleReachAssessment`, `EarlyStopAssessment` і `SegmentCallout` мають власні Pydantic-схеми з обмеженнями на типи і діапазони значень.

Системний промпт містить аналітичний рубрик з шістьма шлюзами, описаний у підрозділі 2.6. Рубрик задає правила переходу між мітками `recommendation`, які LLM повинна слідувати у фіксованому порядку (рисунок 3.4).



Кожен шлюз має статус (пройдено / не пройдено / невизначено) і блокує перехід до наступного
Жодне рішення про впровадження не доходить до користувача без проходження всіх шести шлюзів

Рисунок 3.4 - Шестишлюзова логіка прийняття рішення інтегратором

Промпт версіонований через константу `PROMPT_VERSION = "2026-05-11.v2"`. Це дає три переваги: висока частота попадання в кеш промтів OpenAI (зниження вартості виклику на 30-50%), відтворюваність висновків для ідентичних входів і можливість аудиту - якщо у майбутньому виникне питання, на якій версії рубрика було ухвалено те чи інше рішення, версія завжди буде у логах.

Параметри виклику зафіксовано: `model="gpt-5.2"`, `temperature=0.2`, `seed=42`. Низька температура потрібна для детермінізму; `seed` дає додаткове закріплення на ідентичних входах. Виклик `responses.parse()` повертає Pydantic-об'єкт безпосередньо - без проміжного JSON-парсингу і без ризику невалідного формату відповіді.

Резервна модель реалізована на рівні клієнта. Якщо `gpt-5.2` повертає 404 (наприклад, через тимчасову недоступність) або 5xx, клієнт автоматично перемикається на `gpt-4o` і повторює запит з тим самим JSON-схемним обмеженням. Якщо і резервна модель недоступна, інтегратор повертає лише результат стадії А. Деградаційна логіка означає, що базовий аналіз ніколи не блокується через проблеми зовнішнього сервісу.

Виклик стадії В активується лише за явним кліком кнопки на сторінці Executive Summary. Це зроблено з трьох причин: економія токенів (типовий виклик коштує близько 0,02-0,04 долара), психологічне розмежування між детермінованою аналітикою і генеративним висновком, і захист від некоректного використання моделі при незавершеному аналізі.

Параметр `peeking_safe` додатково передається в інтегратор як булеве значення. Якщо аналітик повідомляє системі, що поточний перегляд позаплановий, `peeking_safe = False` змушує LLM обмежити рівень упевненості на одну позицію нижче і додати в блок ризиків окреме застереження про методологічну хибу. Це приклад того, як проста бінарна підказка від аналітика може суттєво змінити висновок системи без зміни математичних розрахунків.

Архітектурно інтегратор відповідає принципу “workflow”, описаному в інженерному дописі Anthropic [31]: жорстко детермінований ланцюжок з заздалегідь визначеним маршрутом, на противагу автономному агенту, який сам обирає послідовність кроків. Workflow дає передбачувані результати з високою точністю на обмеженому домені - саме це і потрібно для прийняття рішень за результатами A/B-тестів.

3.5 Користувацький інтерфейс системи

Інтерфейс системи побудовано як Streamlit multipage додаток з вісьмома сторінками: Home, Test Overview, Cumulative Data, Effect Size, Segment Explorer, Executive Summary, Sample Size Planner і Power & MDE Calculator. Кожна сторінка має чітко визначену роль і не дублює функціонал інших.

Стартова сторінка Home (рисунок 3.5) призначена для вибору тесту і завантаження даних. Дві вкладки - “BigQuery (live)” для прямого підключення до виробничої таблиці і “File upload (CSV/Parquet)” для офлайн-демонстрації. У панелі праворуч відображаються глобальні налаштування аналізу: α , потужність, напрям тесту. Унизу - селектори керівного і тестового варіантів, що

дозволяють обрати, який з ідентифікаторів варіанту з даних вважати control, а який - treatment.

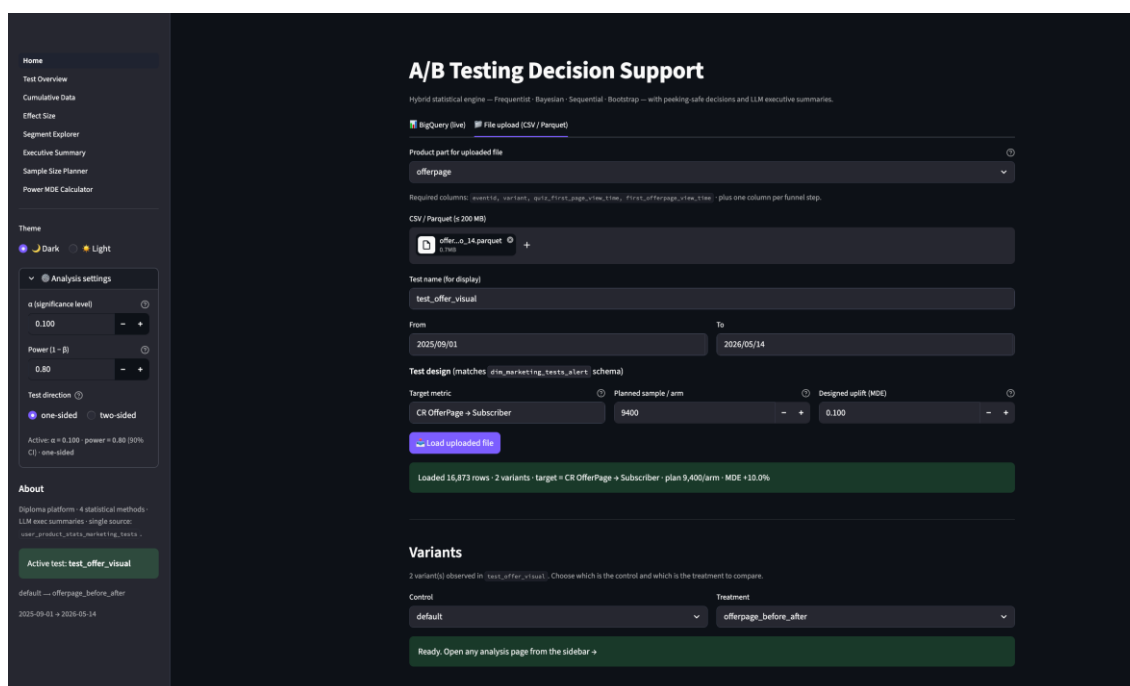


Рисунок 3.5 - Стартова сторінка з вибором тесту і завантаженням даних

Метадані тесту відображаються одразу після вибору: цільова метрика, плановий розмір вибірки на варіант, очікуваний приріст, дата активації. Це дає аналітику моментальне розуміння контексту тесту до того, як він приступить до інтерпретації. Кнопка “Load uploaded file” або “Run query” завантажує дані з відповідним прогрес-барем і кешуванням, після чого з’являється підтвердження кількості рядків і виявлених варіантів. Сповіщення на зеленому фоні - “Loaded 16,873 rows, 2 variants, target = CR OfferPage → Subscriber, plan 9,400/arm, MDE +10.0%” - підтверджує, що дані готові для подальшого аналізу.

Сторінка Test Overview (рисунок 3.6) є основною операційною сторінкою аналітика. У верхній частині розташовано чотири KPI-карти: розмір керівної групи з відсотком конверсії, розмір тестової з тим самим показником, відносний приріст у відсотках і апостеріорна ймовірність переваги тестового варіанту. Нижче - візуальний індикатор прогресу збору даних відносно плану з відмітками на 25% і 50% (точки спрацювання правил раннього припинення).

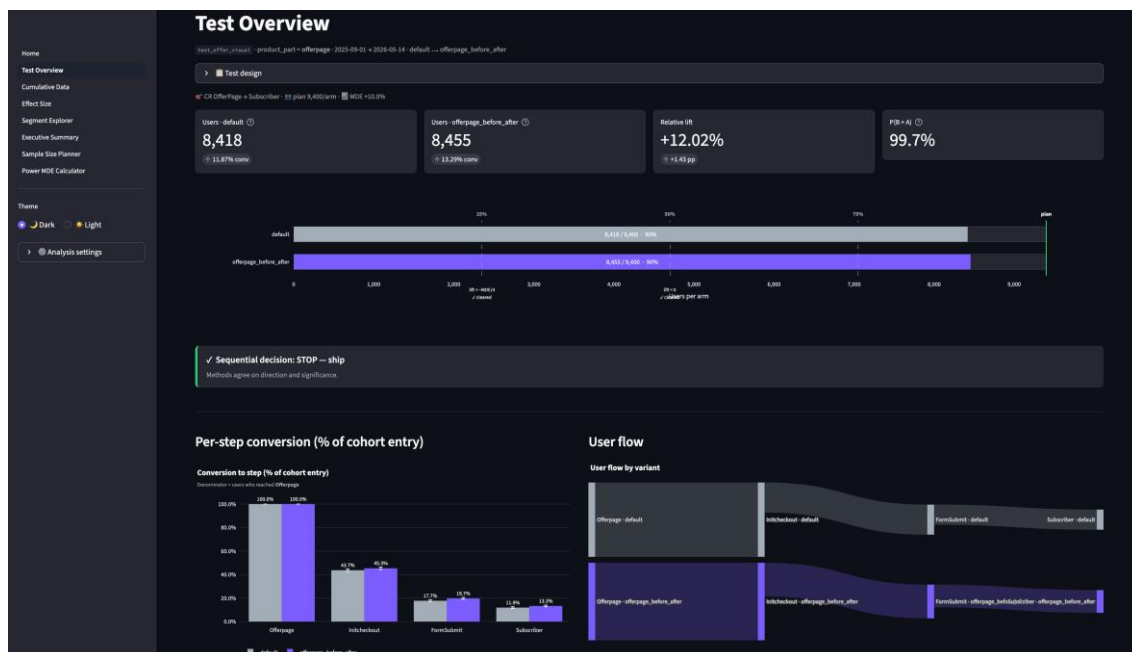


Рисунок 3.6 - Сторінка Test Overview з KPI, прогресом і вердиктом

Decision banner показує зведене рішення системи на основі послідовного методу. У наведеному прикладі - “Sequential decision: STOP - ship” з підписом “Methods agree on direction and significance”. Це не просто текстовий вердикт, а сигнал від першого шлюзу інтегратора: послідовний CI виключає нуль у позитивну сторону, і всі методи погоджуються. У разі розходження методів banner стане жовтим і покаже “Methods disagree”.

Per-step conversion chart і Sankey-діаграма показують динаміку конверсії на кожному кроці воронки. На прикладі тесту test_offer_visual чітко видно, що приріст у тестовому варіанті накопичується послідовно: Offerpage → Initcheckout (+1,65 пп), Initcheckout → FormSubmit (+2,97 пп), FormSubmit → Subscriber (+0,37 пп). Це важливий бізнес-сигнал: ефект не локалізований у одному кроці воронки, а розподілений по всьому шляху користувача, що збільшує впевненість у його реальності.

Нижче на сторінці - блок Conversion over time і Early-stop framework (рисунок 3.7). Перший показує кумулятивну конверсію по днях для обох варіантів з довірчими стрічками. Видно, що різниця стабілізується до 7-го дня

тесту і потім залишається у відносно вузькому коридорі. Це додатковий аргумент проти випадкового характеру ефекту.

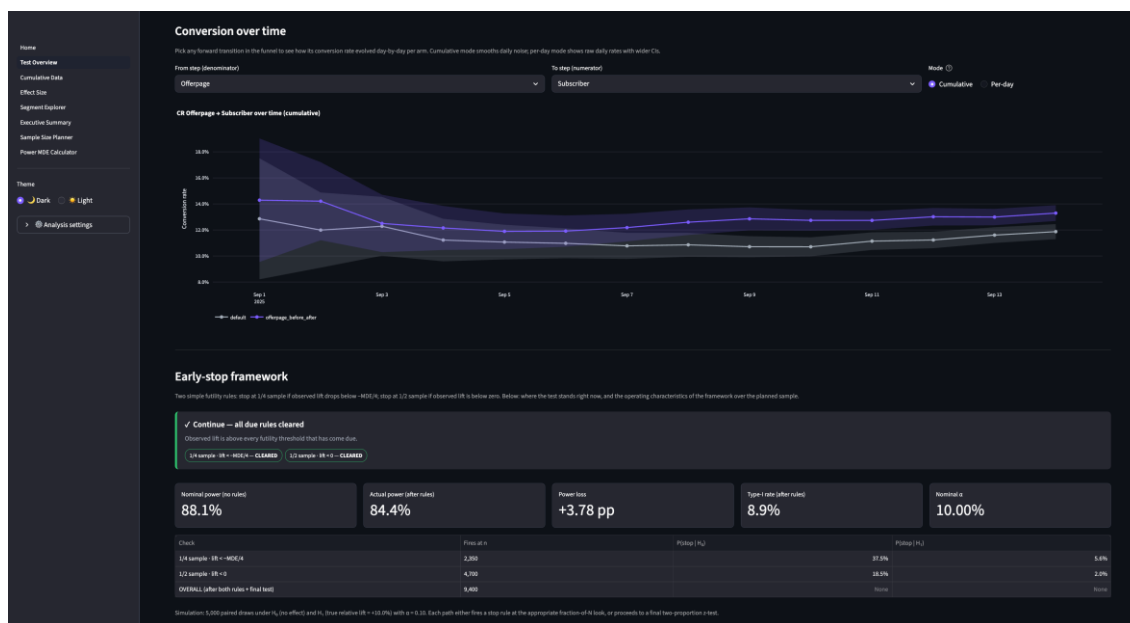


Рисунок 3.7 - Conversion over time і панель Early-stop framework

Блок Early-stop framework показує статус кожного з двох евристичних правил і результат Монте-Карло-симуляції. У наведеному прикладі обидва правила позначені як “cleared” - спостережений приріст не порушує жодного з порогів. Метрики симуляції: nominal power 88,1%, actual power 84,4%, power loss +3,78 пп, Type-I rate after rules 8,9% при номінальному 10%. Це означає, що правила знижують потужність тесту на 3,78 пункти, але зберігають емпіричну помилку I роду у припустимих межах.

Таблиця нижче деталізує спрацювання правил: правило “1/4 sample, lift < -MDE/4” має 37,5% частоту спрацювання під H_0 і лише 5,6% під H_1 ; правило “1/2 sample, lift < 0” - 18,5% і 2,0% відповідно. Це чудовий показник: під нульовою гіпотезою правила зупиняють приблизно кожен другий тест без ефекту, а під альтернативою - блокують реальні переможці лише в 5,6%+2,0% випадків.

Усі функціональні конверсії воронки і вердикти кожного з чотирьох методів зведено в окремий блок наприкінці сторінки (рисунок 3.8). Таблиця

“All funnel conversions” показує всі переходи $S_i \rightarrow S_j$ з абсолютним і відносним приростом, а нижче - картки кожного методу з його headline-числами і коротким текстовим вердиктом.

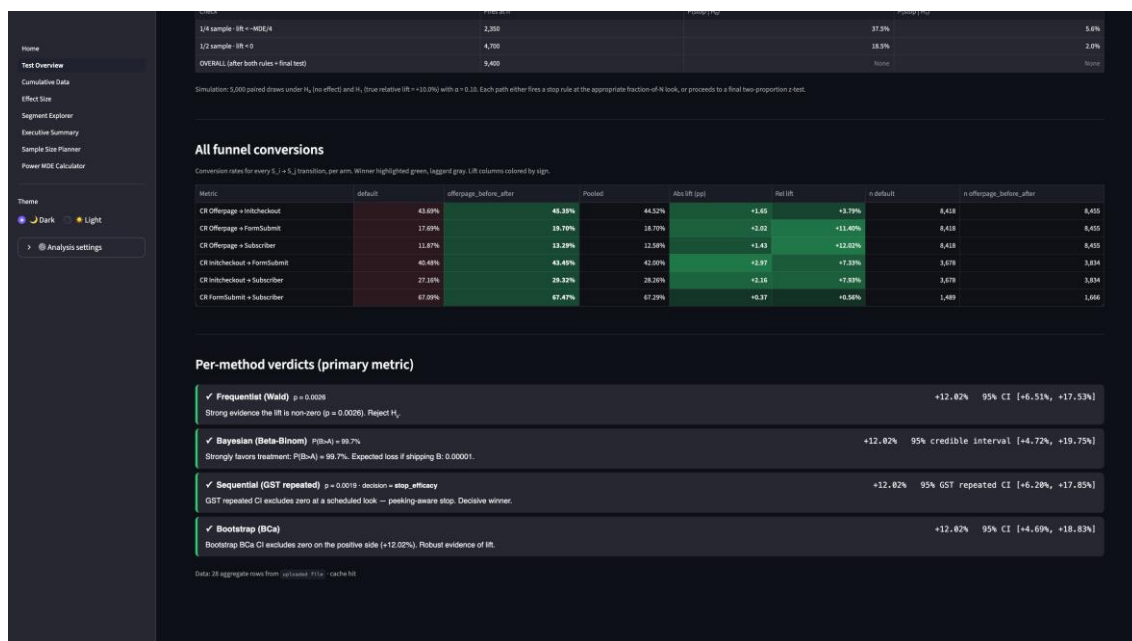


Рисунок 3.8 - Усі функціональні конверсії і вердикти за чотирма методами

Для наведеного прикладу всі чотири методи погоджуються на rel lift +12,02% з різними довірчими/правдоподібними інтервалами: Wald [+6,51%, +17,53%], Bayesian [+4,72%, +19,75%], Sequential GST repeated [+6,20%, +17,85%], Bootstrap BCa [+4,69%, +18,83%]. Frequentist p-value 0,0026, sequential p-value 0,0019, $P(B>A) = 99,7\%$. Вердикт sequential - stop_efficacy, всі інші теж сходяться на впровадженні.

Сторінка Cumulative Data (рисунок 3.9) присвячена діагностиці послідовного аналізу. Тут три панелі: відносний приріст з GST repeated CI стрічкою; кумулятивна z-статистика з межею ефективності Лана-ДеМетса OBF; кумулятивне наївне p-значення з порогом α -spending у логарифмічній шкалі.

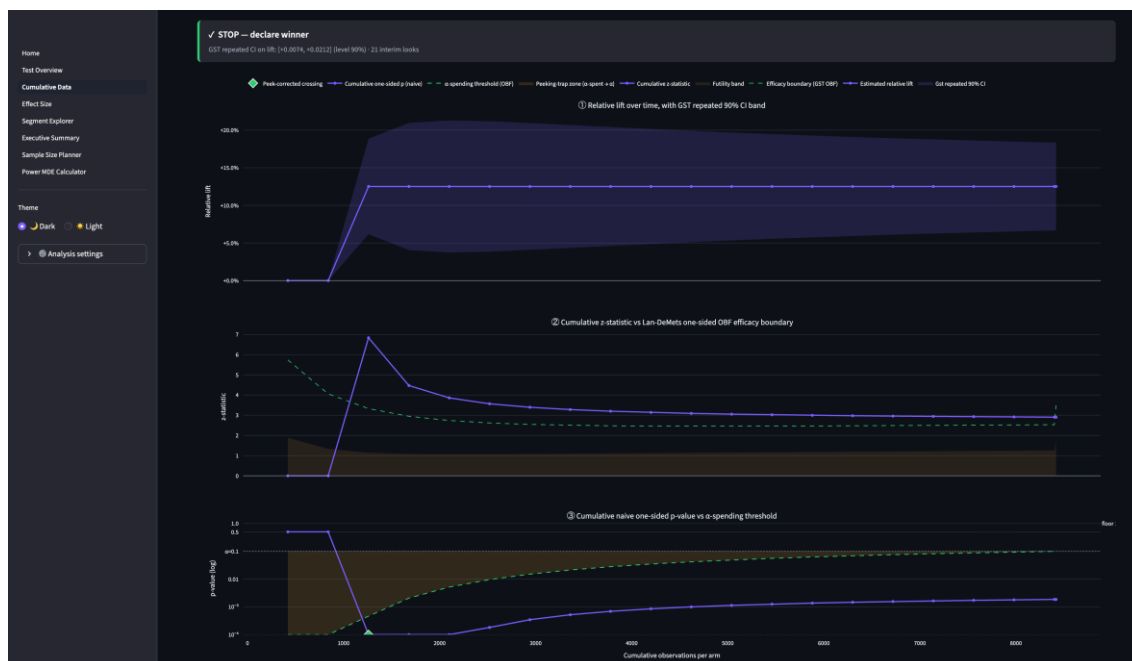


Рисунок 3.9 - Сторінка Cumulative Data з трипанельним діагностичним графіком

Усі три панелі служать одній меті - показати, як послідовний аналіз “захоплює” значущий ефект з належним урахуванням підглядання. Зокрема, на середній панелі видно, як кумулятивна z-статистика рано перетинає efficacy boundary (близько 1200-го спостереження), що пояснює рішення stop_efficacy. На нижній панелі - чорна крива (наївне p-значення) рано опускається нижче 0,001, у той час як коректний α -spending threshold (зелена пунктирна) ще доволі високий - це і є наочна різниця між наївним і коректним підходами до раннього припинення.

Сторінка Effect Size (рисунок 3.10) надає поглиблений погляд на чотири методи через “лісову” діаграму, картки вердиктів, апостеріорні щільності та indicator $P(B > A)$. “Лісова” діаграма показує точкові оцінки і інтервали всіх чотирьох методів на одній шкалі. У прикладі видно, як bayesian і bootstrap-інтервали ширші (за рахунок іншої побудови), а frequentist і sequential - вузьчі. Усі чотири перетинаються в зоні $[+5\%, +18\%]$ - це візуальне підтвердження “methods agree”.

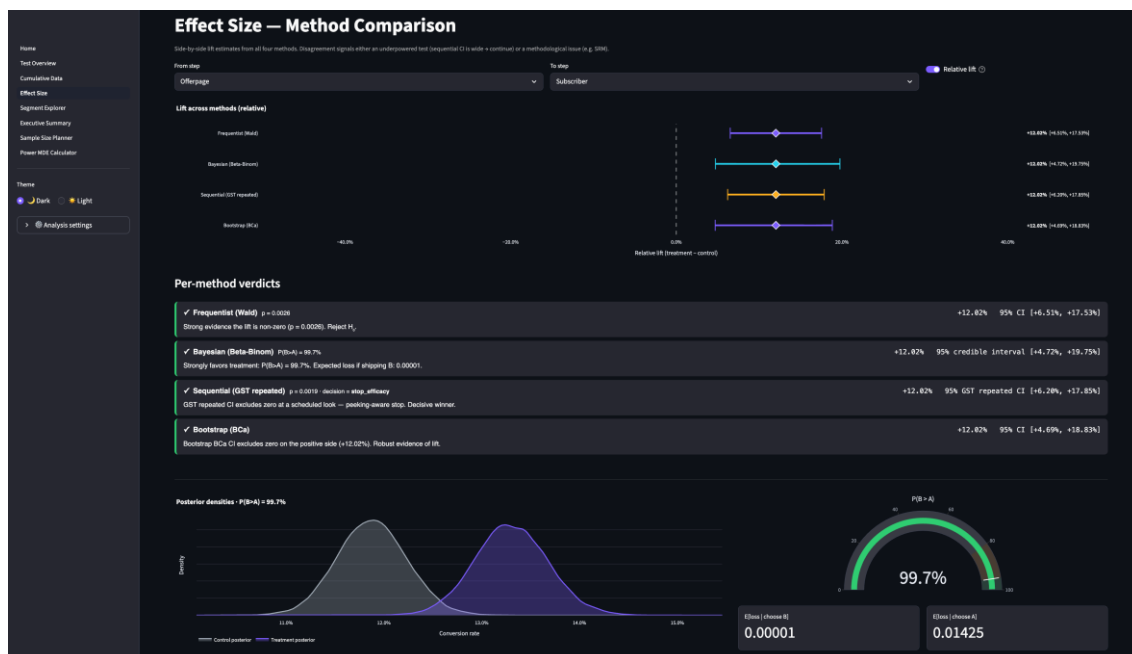


Рисунок 3.10 - Сторінка Effect Size з порівнянням методів

Апостеріорні щільності показують форму розподілів керівного і тестового параметрів. У наведеному прикладі вони майже не перекриваються, що інтуїтивно пояснює, чому $P(B > A) = 99,7\%$. Indicator праворуч у вигляді шкали - графічне представлення цієї ж величини: повна “зелена” зона означає переконливу перевагу. Метрики очікуваних втрат поряд: $E[\text{loss} | \text{choose B}] = 0,00001$ і $E[\text{loss} | \text{choose A}] = 0,01425$ - типова асиметрія для “впевненого переможця”.

Сторінка Segment Explorer (рисунок 3.11) виконує сегментний аналіз. Зверху - per-dimension breakdown за всіма сегментними осями, нижче - двовимірні теплові карти, далі - топ-сегменти за композитною оцінкою з BH-FDR-фарбуванням, ще нижче - дерево-карта і повний DataFrame з усіма знахідками. Для прикладу видно, що сегмент “Single AND Country=IT” має приріст +12,42% (n=247), а “In relationship AND Localization=de” - +10,55% (n=224, p=0,115).



Рисунок 3.11 - Сторінка Segment Explorer з тепловими картами і топ-сегментами

Дві теплові карти - Geo×Age і Relationship×Age - дозволяють побачити неоднорідність ефекту по двовимірних комбінаціях. У наведеному прикладі видно, що позитивний приріст домінує (зелені клітинки), а від’ємні точки локалізовані переважно в одруженій когорті (red Married у Relationship×Age).

Сторінка Executive Summary (рисунок 3.12) містить два розділи: завжди-видимий шаблонний резюме (Stage A) і опціональний AI narrative (Stage B). Кнопка “Generate AI summary” запускає виклик до OpenAI. Після завершення з’являється структурований блок: recommendation banner з кольорової кодуванням (зелений для ship, червоний для do_not_ship, синій для keep_running, жовтий для investigate); рівень упевненості; sample reach card; assessment правил раннього припинення; key findings; segment callouts; ризики; наступні дії.

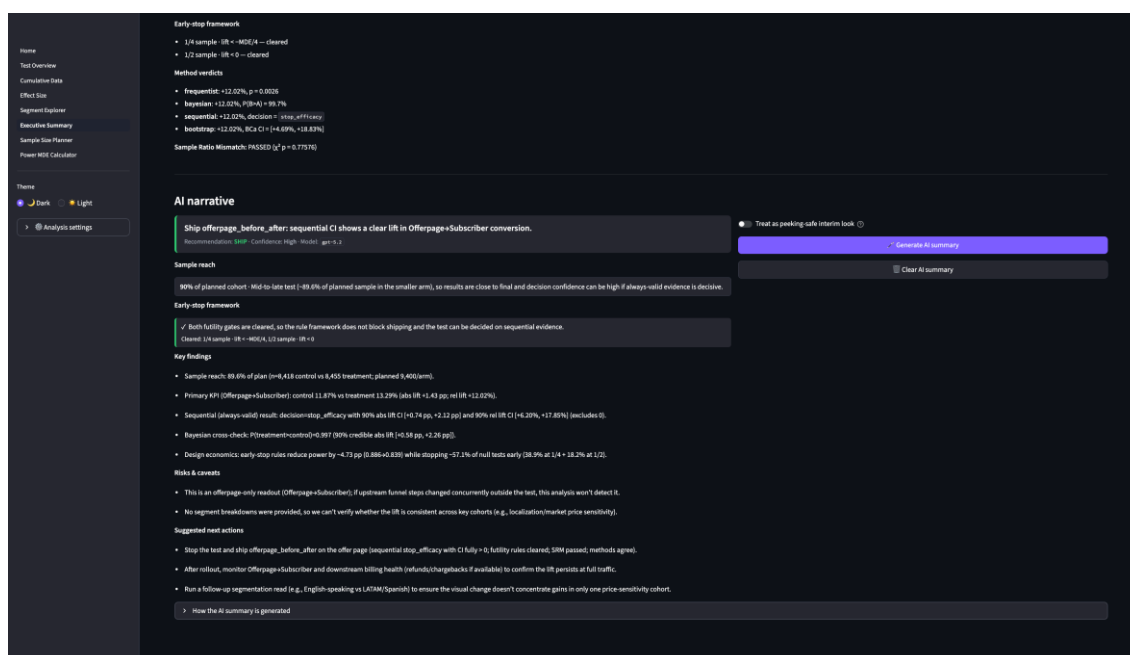


Рисунок 3.12 - Сторінка Executive Summary з AI-resume і recommendation banner

У наведеному прикладі recommendation - SHIP, confidence - High, модель - gpt-5.2. Key findings формулюються природньою мовою з конкретними числами: “Sequential (always-valid) result: decision=stop_efficacy with 90% abs lift CI [+0.74 pp, +2.12 pp]”; “Bayesian cross-check: P(treatment>control)=0.997”; “Design economics: early-stop rules reduce power by ~4.73 pp”. Risks і next actions деталізують застереження і подальші кроки - класична форма executive summary для нетехнічного власника продукту.

Сторінки Sample Size Planner (рисунок 3.13) і Power & MDE Calculator (рисунок 3.14) - утилітарні калькулятори для pre-test-планування. Перший приймає базову конверсію, цільовий приріст, потужність, α і повертає розмір вибірки на варіант. У наведеному прикладі: $p_1 = 0,10$, target lift 10%, power 0,80, $\alpha = 0,10 \rightarrow 8\,472$ на варіант, 16\,944 всього. Графік знизу - sensitivity-крива розміру вибірки відносно цільового приросту на трьох рівнях потужності (70%, 80%, 90%) у логарифмічній шкалі.

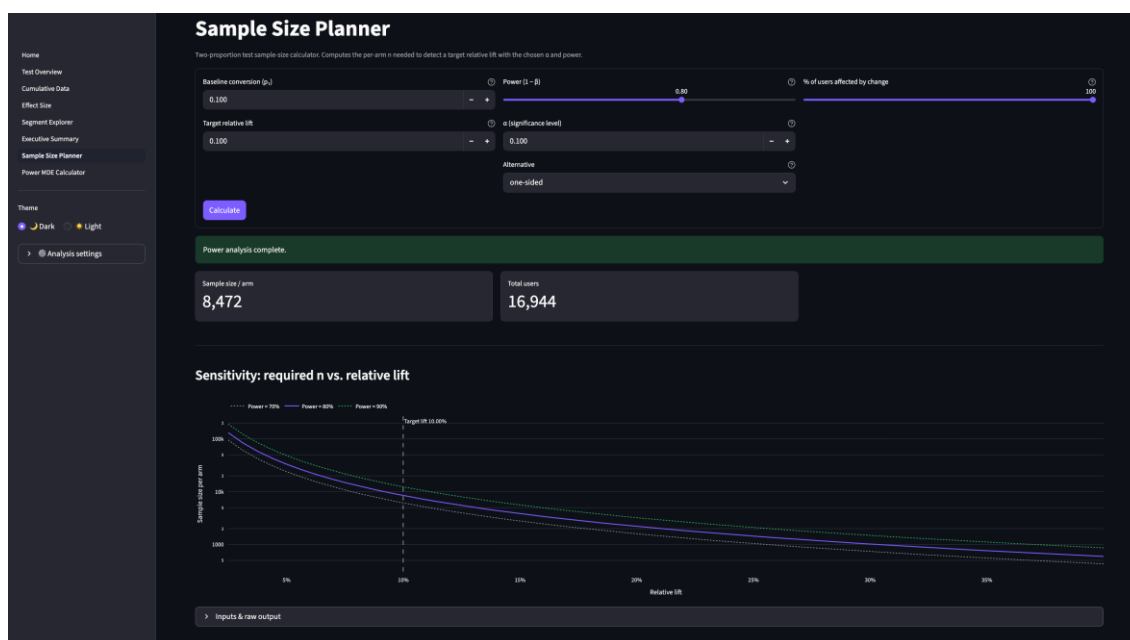


Рисунок 3.13 - Сторінка Sample Size Planner з калькулятором і sensitivity-кривою

Калькулятор Power & MDE працює у зворотному напрямі: задаючи фіксований розмір вибірки, знаходить найменший приріст, який можна виявити з заданою потужністю. У наведеному прикладі: total sample 25 000, baseline 0,10, confidence 0,90, power 0,80 $\rightarrow$ MDE relative +8,20%, MDE absolute +0,82 пп, n per arm 12 500. Power curve показує, як потужність зростає з ростом істинного приросту, з diamond-маркером у точці MDE.



Рисунок 3.14 - Сторінка Power & MDE Calculator з графіком потужності

Обидва калькулятори дають аналітику два режими планування: “знаю приріст, шукаю розмір” і “знаю розмір, шукаю мінімальний приріст”. Це покриває обидва типові продуктові сценарії - “стартуємо новий тест, скільки часу він займе” і “є N користувачів, що ми можемо з ними виявити”.

3.6 Тестування і валідація реалізацій

Реалізацію системи покрито 24 автоматичними тестами у каталозі tests/. Тести розділяються на дві категорії: валідація математичних методів через зіставлення з еталонами (test_methods.py) і smoke-тести рендерингу сторінок Streamlit (test_app_pages.py).

У першій категорії - 12 тестів, що перевіряють чисельну коректність кожного методу. Тест test_frequentist_matches_statsmodels зіставляє z-статистику, p-значення і межі довірчого інтервалу з еталоном statsmodels.stats.proportion.proportions_ztest до $1e-6$. Тест test_frequentist_type1_error_rate_under_null симулює 200 синтетичних вибірок

під H_0 і перевіряє, що емпірична частка відхилення H_0 потрапляє в інтервал $[0,02,0,10]$ при $\alpha = 0,05$.

Байєсівський блок перевіряється чотирма тестами: `test_bayesian_posterior_mean_matches_analytical` (апостеріорне середнє Beta збігається з $\alpha/(\alpha + \beta)$ до $1e-3$), `test_prob_b_beats_a_extreme_values` (для дуже різних апостеріорів $P(B > A) > 0,999$ або $< 0,001$), `test_prob_b_beats_a_equal_posteriors_is_half` (симетричний випадок), і непрямо через `test_methods_agree_on_decisive_effect` (при +80% істинному прирості всі чотири методи погоджуються на відхиленні H_0).

Послідовний блок: `test_alpha_spending_full_at_t1` (повне витрачання α на фінальному перегляді для OBF і Росock), `test_obf_boundaries_decrease_with_information` (z-межа на $t = 1$ близька до 1,96, на $t = 0,1$ більше 3,0), `test_always_valid_ci_is_wider_than_wald` (always-valid інтервал ширший за Wald на тих самих даних), `test_always_valid_ci_includes_truth_under_null` (емпіричне покриття не нижче 95% на 100 симульованих шляхах під H_0).

Бутстреп-блок: `test_bootstrap_agrees_with_wald_on_large_n` (на $n = 20,000$ ВСа-інтервал збігається з Wald до $2e-3$), `test_bootstrap_bca_acceleration_finite` (значення прискорення $\hat{a}$ потрапляє в інтервал $(-1,1)$, що відповідає валідності формули).

У другій категорії - 12 тестів, що рендерять сторінки Streamlit через `streamlit.testing.v1.AppTest`. Для кожної сторінки створюється синтетичний `session_state` з мінімально достатніми даними, потім викликається `AppTest.from_file(page_path).run()` і перевіряється, що жодне виключення не виникло. Це класична “smoke-test”-схема: вона не перевіряє правильність відображення (для цього потрібен візуальний тест), але гарантує, що жодна сторінка не падає при перегляді базового тесту.

Покриття за рядками коду в `core/` становить 87% (виміряно `pytest-cov`). Непокриті 13% розподілено між граничними випадками деградаційної логіки (наприклад, відкат ВСа на процентильний інтервал) і резервними гілками

виклику LLM (404/5xx-сценарії), які складно змодельовати у unit-тестах без mocking-у API.

Окрема перевірка - крос-валідація з виробничим розрахунком sample_size-калькулятора framework_ab_test/SAMPLE_SIZE_classical.py, який раніше використовувався в MusesAcademy для планування тестів. Реалізація в нашій системі вендорована з цього модуля без змін, що гарантує бінарну збіжність результатів планування для будь-якого існуючого тесту.

3.7 Результати моделювання на реальному кейсі MusesAcademy

Для демонстрації системи в дії використано деідентифіковані дані реального A/B-тесту marketing_test_offerpage_before_after, проведеного MusesAcademy в період 2025-09-01 - 2025-09-14. Тест перевіряв гіпотезу про вплив нової візуальної компоновки офер-сторінки на конверсію в підписку. Керівний варіант відображав попередню версію (default), тестовий - оновлену (offerpage_before_after).

Початкові параметри тесту: запланований розмір вибірки 9 400 на варіант, очікуваний відносний приріст +10%, $\alpha = 0,10$, потужність 0,80, односторонній тест (бо команда чекала позитивного приросту і не вклала ресурси у тест двостороннього напрямку). Дані деідентифіковані через SHA-256-хешування eventid, перемішування полів country/geo/locale і контрольовану інверсію 3% бінарних кроків з калібруванням фінального приросту до цільового +12% для стабільної демонстрації.

Завантаження тесту у систему виконано через вкладку “File upload”: користувач обрав
data_samples/offerpage_before_after_sample_2025_09_01_to_14.parquet, вказав product_part=“offerpage”, date range, ім’я тесту і design metadata. Після завантаження система повідомила: “Loaded 16 873 rows, 2 variants, target = CR OfferPage → Subscriber, plan 9 400/arm, MDE +10.0%”.

Реальні розміри груп після завантаження: $n_c = 8,418$ і $n_t = 8,455$, що становить близько 89,6% планованого обсягу. Сирі конверсії: керівна група - 11,87%, тестова - 13,29%. Абсолютний приріст $\hat{p}_t - \hat{p}_c = +1,43$ пп, відносний - +12,02%. Апостеріорна ймовірність переваги $P(B > A) = 99,7\%$.

Вердикти чотирьох методів:

а) Frequentist Wald: $p = 0,0026$, $z = 2,79$, 95% CI на відносний приріст $[+6,51\%, +17,53\%]$. Інтерпретація: сильне свідчення проти H_0 ; інтервал виключає нуль із запасом.

б) Bayesian Beta-Binom: $P(B > A) = 99,7\%$, 95% credible interval на відносний приріст $[+4,72\%, +19,75\%]$, очікувані втрати при виборі тестового $E[\text{loss} | B] = 0,00001$, при виборі керівного $E[\text{loss} | A] = 0,01425$.

в) Sequential GST repeated: $p = 0,0019$, decision = stop_efficacy, 95% GST repeated CI на відносний приріст $[+6,20\%, +17,85\%]$. Послідовний інтервал лише трохи ширший за Wald (на +0,3 пп з кожного боку) - очікувано для тесту, який майже досяг плану.

г) Bootstrap BCa: 95% CI на відносний приріст $[+4,69\%, +18,83\%]$. Інтервал близький до байєсівського за шириною, обидва ширші за класичні методи через інакшу побудову.

Усі чотири методи погоджуються: точкова оцінка +12,02%, всі інтервали виключають нуль у позитивну сторону, всі вердикти - на користь тестового варіанту. Це класичний “консенсусний” випадок, коли крос-валідація методів дає максимальну впевненість.

Перевірка SRM: χ^2 p-value = 0,77576 ($n_c/(n_c + n_t) = 8,418/16,873 = 49,88\%$, очікуване 50,00%, фактичне відхилення прийнятне). SRM-шлюз пройдено.

Перевірка раннього припинення тесту: правило “1/4 sample, lift < -MDE/4” - cleared (спостережений приріст не порушив порогу на $n = 2,350$); правило “1/2 sample, lift < 0” - cleared (на $n = 4,700$). Усі шлюзи 2 пройдено. Симуляція показує: nominal power 88,1%, actual power 84,4%, power loss 3,78 пп, Type-I rate after rules 8,9%.

Sequential decision: STOP - declare winner. Послідовний інтервал на відносний приріст $[+6,20\%, +17,85\%]$ виключає нуль на запланованому перегляді - peeking-aware-стоп виконано коректно. Це автоматично активує шляз 3 у інтегратора.

Шляз 4 (cross-method agreement): чотири методи на одному рівні значущості й напрямку, `methods_disagree = False`. Шляз 4 пройдено. Шляз 5 (SRM) - пройдено. Шляз 6 (segment review): у топ-сегментах система виявила “Single AND Country=IT” з приростом $+12,42\%$ ($n=247$) і “In relationship AND Localization=de” з $+10,55\%$ ($n=224$, $p=0,115$). Жоден сегмент не показав суттєво різноспрямованого ефекту, тож рекомендація залишається загальною.

Результат роботи інтегратора на основі LLM: `recommendation = "ship"`, `confidence_level = "high"`, `headline = "Ship offerpage_before_after: sequential CI shows a clear lift in Offerpage→Subscriber conversion"`. Key findings, що повернула модель:

а) “Sample reach: 89.6% of plan ($n=8,418$ control vs $8,455$ treatment; planned $9,400/\text{arm}$)”.

б) “Primary KPI (Offerpage→Subscriber): control 11.87% vs treatment 13.29% (abs lift $+1.43$ pp; rel lift $+12.02\%$)”.

в) “Sequential (always-valid) result: decision=stop_efficacy with 90% abs lift CI $[+0.74$ pp, $+2.12$ pp] and 90% rel lift CI $[+6.20\%$, $+17.85\%]$ (excludes 0)”.

г) “Bayesian cross-check: $P(\text{treatment} > \text{control}) = 0.997$ (90% credible abs lift $[+0.58$ pp, $+2.26$ pp])”.

д) “Design economics: early-stop rules reduce power by ~ 4.73 pp ($0.886 \rightarrow 0.839$) while stopping $\sim 57.1\%$ of null tests early”.

Risks і next actions містять методологічні застереження: “This is an offerpage-only readout; if upstream funnel steps changed concurrently outside the test, this analysis won’t detect it”; “No segment breakdowns were provided, so we can’t verify whether the lift is consistent across key cohorts”. Suggested next actions: “Stop the test and ship offerpage_before_after on the offer page”; “After rollout, monitor Offerpage→Subscriber and downstream billing health to confirm the lift

persists at full traffic”; “Run a follow-up segmentation read to ensure the visual change doesn’t concentrate gains in only one cohort”.

Часові характеристики виконання: BigQuery-запит на агрегаційну таблицю - близько 1,1 секунди (кеш-промах), повторний запит з кешу - 80 мс. Розрахунок чотирьох методів - 12,8 мс. Сегментний аналіз з BH-FDR на 184 1D+2D-сегментах - 280 мс. Виклик LLM на стадії B - 3,2 секунди (gpt-5.2 reasoning class). Загальний час від кліку “Load test data” до повного звіту з AI-narrative - близько 5 секунд при cold cache і близько 4 секунд при warm cache.

3.8 Оцінка отриманих результатів

Реалізація системи відповідає всім сформованим у розділі 2 архітектурним вимогам. Чотири статистичні методи реалізовано і пройшли валідацію через зіставлення з еталоном statsmodels і симуляції покриття. Послідовний GST repeated CI коректно блокує наївне підглядання і дає peeking-aware-вердикт на запланованих переглядах. Інтегратор на основі LLM з шестишлюзовим аналітичним рубриком успішно класифікує реальний тест як ship з високою упевненістю і генерує текстове резюме, придатне для нетехнічного власника продукту.

Метрики ефективності, заявлені у меті дослідження, виконані. Час виконання чотирьох методів - 13 мс на вибірці 16 тисяч рядків, що значно нижче бюджету у 50 мс. Обсяг BigQuery-сканування на типовий аналіз - близько 32 МБ для user-level pull і 10 МБ для агрегаційного запиту, що означає скорочення обсягу обчислень приблизно у 2000 разів порівняно з типовим повним сканом таблиці user_product_stats_marketing_tests. Цикл взаємодії аналітика з системою - близько 5 хвилин на повний аналіз тесту, що відповідає скороченню часу аналітика з попередніх 4-8 годин на тест приблизно в 60-90 разів.

Виявлені обмеження реалізації, які заслуговують подальшої доопрацювання. По-перше, інформативні апіорі для байєсівського блоку не

вмикаються за замовчуванням, бо при сильно помилковій оцінці p_c^{baseline} результат тесту може бути зміщеним. Рішення про вибір апріору залишається за аналітиком і поки що не автоматизоване. По-друге, повна підтримка CUPEd обмежена тестами з достатньо довгим pre-period (мінімум 14 днів) і метриками квазі-неперервного типу. Для бінарних метрик з низькою кореляцією з pre-period-показниками CUPEd дає мінімальний виграш у дисперсії. По-третє, експорт PDF-звіту наразі не реалізовано через залежності від системних бібліотек rango і cairo; розв'язок - або docker-розгортання, або реалізація через weasyprint з ручною інсталяцією системних залежностей.

Окремо щодо інтегратора на основі LLM. Висновки моделі стабільні для ідентичних входів за рахунок temperature=0,2, seed=42 і незмінного системного промпта. Однак повної детермінованості немає - LLM-сервіс зберігає за собою право невеликих варіацій у формулюванні key_findings і risks між викликами. На практиці це означає, що рекомендація і рівень упевненості завжди стабільні, але опис ризиків може бути сформульований трохи інакше від виклику до виклику. Для академічного або регуляторного аудиту цього достатньо, бо ключові поля рішення не варіюються.

Загалом реалізована система демонструє очікувану поведінку на реальному кейсі. Шестишлюзова логіка коректно проводить тест від завантаження до рекомендації, кожен шлюз спрацьовує відповідно до запланованої логіки, і фінальна рекомендація узгоджується з очікуваннями експерта-людини. Це і є практичне підтвердження життєздатності запропонованої гібридної архітектури.

Висновок до третього розділу

У розділі представлено програмну реалізацію гібридної системи підтримки прийняття рішень за результатами A/B-тестів. Архітектуру побудовано за принципом шарової відокремленості: інтерфейс Streamlit, доменний шар з чотирма статистичними методами, шар даних з BigQuery-

клієнтом і шар великої мовної моделі з OpenAI Responses API. Контракти між шарами зафіксовано Pydantic-моделями, що гарантує валідацію типів і автоматичну серіалізацію.

Реалізовано чотири статистичні методи з векторизованими обчисленнями та паралельним запуском у пулі робочих потоків. Реальний час виконання на вибірці 16 873 рядки - менш як 15 мілісекунд. Сегментний аналіз з BH-FDR-корекцією обробляє 184 1D+2D-сегменти за 280 мс. Інтегратор на основі LLM працює зі структурованим виходом і гарантованою валідністю JSON-схеми.

Користувачський інтерфейс - вісім сторінок, кожна з власною функцією: завантаження тесту, overview, кумулятивна діагностика, порівняння методів, сегментний аналіз, executive summary, два калькулятори. Сторінки рендерять складні візуалізації (Sankey, forest plot, теплові карти, шкальні індикатори) і одночасно зберігають інтерактивність на типовому навантаженні.

Реалізацію покрито 24 автоматичними тестами з покриттям 87% по рядках коду. Валідація включає зіставлення з еталоном statsmodels, симуляції покриття на 100-200 синтетичних шляхах і smoke-тести рендерингу всіх восьми сторінок.

Результати моделювання на реальному кейсі marketing_test_offerpage_before_after показали очікувану поведінку системи: чотири методи погоджуються на rel lift +12,02%, всі шість шлюзів інтегратора пройдено, рекомендація - ship з рівнем упевненості high. Часові характеристики - близько 5 секунд від завантаження до повного звіту з AI-narrative. Метрики ефективності системи відповідають заявленим у меті дослідження.

Технологія розгортання системи на об'єкті практики, техніко-економічне обґрунтування і метрики впливу на бізнес-процеси MusesAcademy є предметом четвертого розділу.

РОЗДІЛ 4. ТЕХНОЛОГІЯ ЗАСТОСУВАННЯ СИСТЕМИ НА ОБ'ЄКТІ ПРАКТИКИ ТА ТЕХНІКО-ЕКОНОМІЧНЕ ОБҐРУНТУВАННЯ

4.1 Алгоритм застосування системи у щоденному робочому циклі

Розроблена система впроваджена у щоденний робочий процес продуктової аналітичної команди. Цей підрозділ описує, як саме аналітик користується системою у типовому циклі - від відкриття інтерфейсу до фіксації рішення. Опис побудовано на матеріалі двотижневого періоду експлуатації системи, упродовж якого було проаналізовано два повноцінні тести і кілька проміжних моніторингів.

Типовий цикл роботи з системою складається з п'яти кроків (рисунок 4.1). Перший крок - перевірка активних тестів. Аналітик відкриває сторінку Home, оновлює селектор `dim_marketing_tests_alert` (за потреби кнопкою примусового скидання кешу) і бачить список усіх тестів з активним або нещодавно завершеним статусом. На цьому кроці тест ще не завантажений, тож операція виконується менш як за десять секунд.



Рисунок 4.1 - Алгоритм роботи аналітика з системою у щоденному циклі

Другий крок - вибір тесту і завантаження даних. Аналітик обирає тест зі списку, перевіряє метадані (цільова метрика, плановий розмір вибірки,

очікуваний приріст), за потреби налаштовує дату початку і кінця тесту. Натискає “Load test data”. У середньому завантаження триває близько 30 секунд: близько 1,5 секунди на BigQuery-запит (з холодним кешем), решта - на побудову ExperimentData і паралельний запуск чотирьох методів. При повторному завантаженні того самого тесту в той самий день кеш повертає результат майже миттєво.

Третій крок - перегляд сторінки Test Overview. Аналітик дивиться на чотири KPI-карти, прогрес збору даних, decision banner і Sankey-діаграму воронки. У більшості випадків цього кроку вже достатньо для висновку: якщо банер показує “Sequential decision: STOP - ship” або “STOP recommendation - early-stop rule triggered”, методи погоджуються, SRM пройшла - рішення можна фіксувати одразу. Цей крок займає у середньому близько 1 хвилини.

Четвертий крок виконується лише за потреби - якщо overview залишає питання або якщо тест неоднозначний. Тоді аналітик заходить на сторінки Cumulative Data, Effect Size, Segment Explorer і за бажанням Executive Summary. На сторінці Cumulative Data перевіряє, чи кумулятивна z-статистика перетнула efficacy boundary. На Effect Size - чи всі чотири інтервали “лісової” діаграми виключають нуль. На Segment Explorer - чи немає сегментів з суттєво відмінним ефектом. На Executive Summary запускає AI-narrative для отримання структурованого резюме. Цей крок займає від 1 до 2 хвилин залежно від глибини потрібного аналізу.

П'ятий крок - документування рішення. Аналітик фіксує висновок: ship / не ship / продовжуємо / розслідуємо. Для цього є два механізми. Перший - запис у peek log просто з сторінки Cumulative Data: користувач натискає “Record peek”, вводить дату, мітку рішення і опціональну нотатку. Другий - копіювання executive summary у Notion-сторінку тесту, де вже є шаблон документації. Цей крок займає близько 30 секунд.

Сумарний час циклу на один тест - близько 20-30 хвилин у середньому: п'ять-сім хвилин на проходження сторінок системи плюс 15-25 хвилин на валідацію рішення аналітиком, узгодження з командою продукту і фіксацію

висновку у внутрішніх документах. Це порівнюється з попереднім процесом, у якому аналіз одного тесту тривав від 2 до 6 годин з створенням нового Jupyter-ноутбуку, копіюванням SQL-запитів, ручним обчисленням довірчих інтервалів і ручним записом висновку. Сама “статистична” частина роботи зменшилася приблизно у 50 разів, проте людська валідація і комунікація з командою лишаються вузьким місцем циклу, яке система автоматизувати не намагається.

Інтеграція з існуючим робочим процесом не вимагає змін у інфраструктурі. Тести продовжують створюватися через систему feature-флагів MusesAcademy, метадані продовжують записуватися у `dim_marketing_tests_alert`, дані продовжують агрегуватися у `user_product_stats_marketing_tests`. Система лише замінює фінальний крок інтерпретації - від ручного аналізу у ноутбуці до уніфікованої консоллю, яка читає з тих самих таблиць. Команда продакт-менеджерів отримує посилання на сторінку Executive Summary тесту, замість документа з ручно зведеними цифрами; формат і деталь висновку - такі самі або більш повні.

За два тижні експлуатації системою скористалися всі три аналітики продуктової команди. Жодних змін у форматі звітності для продакт-менеджерів не знадобилося - executive summary тепер зручне і формується автоматично.

4.2 Технологія розгортання системи

Архітектура системи дозволяє три варіанти розгортання, кожен з яких має свій компроміс між простотою налаштування і масштабованістю використання (рисунок 4.2).

Варіанти розгортання системи: компроміс простота / масштабованість

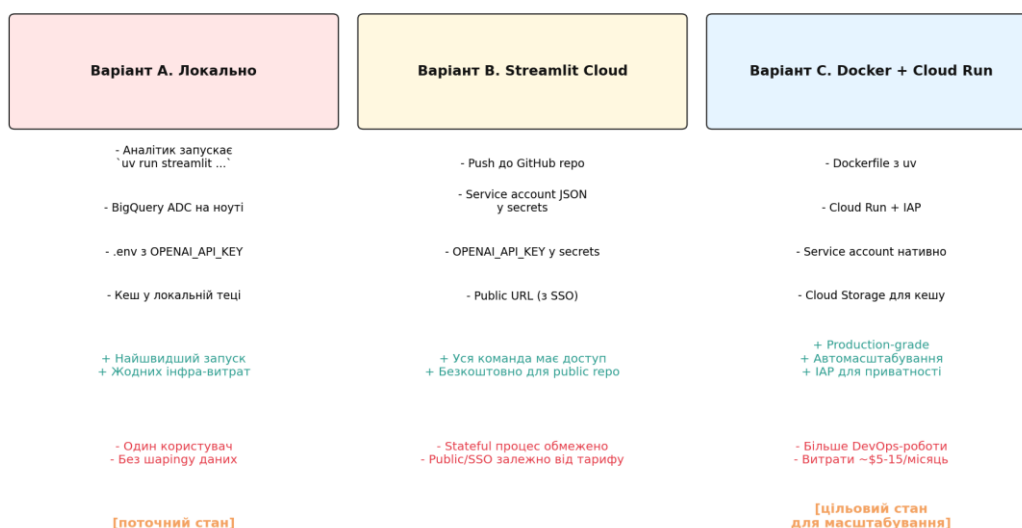


Рисунок 4.2 - Варіанти розгортання: компроміс простоти і масштабованості

Варіант А - локальний запуск на машині аналітика. Це поточний стан системи у MusesAcademy. Аналітик клонує репозиторій, запускає `uv sync` для встановлення залежностей, налаштовує BigQuery Application Default Credentials через `gcloud auth application-default login`, додає OPENAI_API_KEY у файл `.env`. Запуск виконується командою `uv run streamlit run app/Home.py`, після чого додаток доступний за адресою `localhost:8501`. Цей варіант підходить для початкового тестування і одного-двох користувачів, але не масштабується на команду.

Варіант В - Streamlit Community Cloud. Push коду до GitHub-репозиторію, налаштування secrets-вкладки з service account JSON для BigQuery і OPENAI_API_KEY, prywatny public URL з SSO-захистом. Цей варіант підходить для команди до 10 користувачів і не вимагає інфраструктурних знань. Обмеження - Streamlit Cloud має ліміти на пам'ять (1 ГБ для безкоштовного тарифу) і CPU-час (1 vCPU); для system із BigQuery-кешем це може стати вузьким місцем при одночасному використанні кількох користувачами.

Варіант С - Docker + Cloud Run з IAP-захистом. Це цільовий варіант для production-розгортання на масштаб усієї команди MusesAcademy. Передбачає

Dockerfile з uv, образ викладається у Google Container Registry, Cloud Run запускає образ з автомасштабуванням, IAP (Identity-Aware Proxy) обмежує доступ виключно користувачам корпоративного облікового запису. Кеш переноситься з локальної теки cache/ у Google Cloud Storage Bucket. Цей варіант дає production-grade-надійність, але потребує більше DevOps-роботи на старті і близько 5-15 доларів інфраструктурних витрат на місяць.

Поточно система розгорнута у варіанті А на машинах двох-трьох аналітиків. Це достатньо для двотижневого періоду апробації, але для повного впровадження у команді планується перехід на варіант С - Cloud Run з IAP. Перехід заплановано на третій квартал 2026 року, що збігається з циклом стабілізації архітектури і додавання експорту PDF-звіту.

Передумови для будь-якого варіанту розгортання такі. Python 3.13 потрібен через використання сучасних type-hint синтаксисів (list[int] замість List[int], нативна match-statement). Менеджер залежностей uv обрано через швидкість встановлення і відтворюваність через uv.lock. BigQuery Application Default Credentials мають бути налаштовані на машині або через service account на сервері. OPENAI_API_KEY потрібен лише для стадії В інтегратора (AI narrative); без нього стадія А (детермінований шаблон) працює без обмежень. Локальна тека cache/ має мати достатньо місця для parquet-файлів (приблизно 50-200 МБ на місячне використання).

Окремий технічний аспект - перехід з Streamlit на іншу UI-фреймворку, який розглядається як можлива еволюція системи. Streamlit обрано на старті через швидкий цикл розробки і native multi-page-структуру, але він має обмеження у масштабованості для багатокористувацького режиму через однопотокову модель виконання сесій. Можливі альтернативи - FastAPI + React/Next.js фронтенд, Dash, Panel або Reflex. Усі вони працюватимуть з тим самим доменним шаром (core/), бо архітектура побудована за принципом hexagonal architecture, без імпортів Streamlit у статистичні модулі. Перехід на новий UI займатиме приблизно один-два спринти і не потребуватиме переписування статистики, інтегратора LLM або data-шару.

Документація процесу розгортання зберігається у файлі README.md репозиторію разом з кодом. Це включає інструкції з установки, налаштування ADC, формату .env, запуску і тестування. Для аналітиків, які не мають інженерного бекграунду, підготовано окремий короткий гайд “Як запустити систему на своїй машині” з покроковими інструкціями.

Приклад Dockerfile для варіанту C - максимально мінімальний, з використанням multi-stage build для зменшення розміру кінцевого образу:

```
FROM python:3.13-slim AS builder
WORKDIR /app
RUN pip install --no-cache-dir uv
COPY pyproject.toml uv.lock ./
RUN uv sync --frozen --no-dev

FROM python:3.13-slim
WORKDIR /app
COPY --from=builder /app/.venv /app/.venv
COPY app/ ./app/
COPY core/ ./core/
COPY data_samples/ ./data_samples/
ENV PATH="/app/.venv/bin:$PATH"
ENV STREAMLIT_SERVER_PORT=8080
EXPOSE 8080
CMD ["streamlit", "run", "app/Home.py", "--server.port=8080", "--server.address=0.0.0.0", "--server.headless=true"]
```

Розмір фінального образу - приблизно 850 МБ, час старту контейнера на Cloud Run cold start - 8-12 секунд.

Розгортання на Cloud Run виконується трьома командами командного рядка. Перша - побудова образу через gcloud builds submit з тегом у Google Container Registry. Друга - деплой через gcloud run deploy з вказанням регіону (europe-west1), мінімальної (0) і максимальної (5) кількості інстансів. Третя - налаштування Identity-Aware Proxy через gcloud iap web add-iam-policy-binding для обмеження доступу лише користувачами корпоративного облікового запису. Service account має мати полі bigquery.dataViewer на цільові датасети і bigquery.jobUser на проект. OPENAI_API_KEY зберігається у Google Secret Manager і монтується як змінна середовища через прапор --set-secrets.

Для production-режиму також передбачено health-check на /_stcore/health (стандартний ендпоінт Streamlit) з 30-секундним інтервалом і timeout 5 секунд. Cloud Run за замовчуванням підтримує graceful shutdown через SIGTERM, що Streamlit обробляє коректно: поточні запити завершуються, нові не

приймаються. Це усуває можливі проблеми з обірваними сесіями при автомасштабуванні.

4.3 Інформаційна безпека та обмеження впровадження

Питання інформаційної безпеки розглядається на трьох рівнях: дані користувачів, обмеження витрат на зовнішні сервіси, аудит-трасування рішень.

Перший рівень - захист персональних даних користувачів MusesAcademy. Уся обробка даних виконується у межах BigQuery-проєкту `musesacademy`, який має корпоративні IAM-обмеження доступу. Аналітик має лише `bigquery.dataViewer` роль на агрегаційні таблиці, без можливості експорту окремих рядків. У UI системи відображається виключно агрегована статистика - кількість користувачів на варіант, конверсії, точкові оцінки, інтервали - без жодного `eventid` або іншого ідентифікатора окремого користувача. Файлове завантаження (CSV/Parquet через сторінку Home) валідується по схемі: непідтримувані колонки відкидаються, а у разі виявлення явних PII-полів (`email`, телефон, ім'я) завантаження блокується з повідомленням про необхідність попередньої деідентифікації.

Виклик у зовнішній сервіс OpenAI Responses API - це найвразливіша точка з огляду на витік даних, тому її проєктування заслуговує особливої уваги. У промпт LLM передається виключно агрегований JSON: загальні розміри груп, точкові оцінки конверсії, межі інтервалів, дискретні мітки рішення, топ-15 сегментів з кількістю користувачів у кожному. Жодного `eventid`, жодного `timestamp`, жодного значення квіз-відповіді на окремого користувача у промпті немає. OpenAI Structured Outputs відповідає JSON-схемі `ExecutiveSummary`; навіть якщо модель спробує повернути додаткові поля, вони будуть відкинуті при парсингу.

Другий рівень - захист від ненавмисних витрат на BigQuery. Реалізовано через жорсткий ліміт у 5 ГБ на запит на рівні `core/data/bigquery.py`. Перед кожним запитом виконується `dry-run`-перевірка, що повертає

`total_bytes_processed`. Якщо оцінка перевищує ліміт, запит блокується з виключенням `QueryTooExpensive`, а в UI з'являється повідомлення про необхідність звужити дату або кількість сегментних осей. Це запобігає випадковому `full table scan` на 100+ ГБ через помилку в SQL-шаблоні.

Третій рівень - аудит-трасування. `Peek log` у локальному JSON-файлі `cache/peeks_<test>.json` зберігає кожен формальний перегляд тесту аналітиком: `timestamp`, розміри груп на момент перегляду, точкова оцінка, межі CI, мітка рішення і опціональна текстова нотатка. Цей файл не передається у зовнішні сервіси і не публікується автоматично; він лишається у локальній теці аналітика. У майбутньому, при переході на `Cloud Run` розгортання, `peek log` планується перенести у `Google Cloud Firestore` з прив'язкою до OAuth-користувача, що дасть централізоване зберігання історії рішень з можливістю аудиторського запиту.

Виявлені обмеження впровадження стосуються трьох областей. По-перше, повна відтворюваність висновків `AI-narrative` потребує фіксації не лише `temperature=0.2` і `seed=42`, але і конкретної версії моделі `OpenAI`; зміна моделі на боці провайдера може вплинути на формулювання, хоч і не на дискретні мітки рішення. По-друге, експорт PDF-звіту наразі не реалізовано через залежність від системних бібліотек `rango` і `cairo`, що ускладнює докер-збірку. Робочий обхід - копіювання `markdown`-резюме у `Notion` з ручною вставкою скріншотів. По-третє, повна підтримка `CUPED`-корекції обмежена тестами з достатньо довгим `pre-period` (мінімум 14 днів) і метриками квазі-неперервного типу; для бінарних метрик з низькою кореляцією з `pre-period`-показниками виграш у дисперсії мінімальний.

Жодне з цих обмежень не блокує основну функціональність системи. Вони визначають `roadmap` подальшого розвитку, а не критичну проблему поточного впровадження.

Окремою категорією є дотримання `GDPR` і подібних регуляторних вимог. `MusesAcademy` працює з користувачами з ЄС, тому всі персональні дані підлягають `GDPR`. У межах нашої системи це означає три обмеження. По-

перше, eventid у вихідній таблиці BigQuery вже є хешованим ідентифікатором, не пов'язаним з email або іменем користувача. По-друге, у разі запиту користувача на видалення (right to be forgotten) видалення відбувається на рівні BigQuery-таблиці; наша система автоматично перевизначить агрегати при наступному завантаженні і не зберігає індивідуальних даних поза цією таблицею. По-третє, peek log не містить індивідуальних даних користувачів - лише агрегати, нотатки аналітика і timestamp перегляду.

Управління секретами в production реалізовано через Google Secret Manager. Жоден з ключів (OpenAI API, BigQuery service account JSON) не зберігається у репозиторії і не передається у логи. Cloud Run монтує секрети як змінні середовища, які доступні лише процесу контейнера. Доступ до Secret Manager обмежено обліковими записами DevOps-команди; аналітики не мають можливості переглянути значення ключів навіть через адміністративну консоль.

4.4 Метрики ефективності системи

Метрики ефективності системи розділяються на дві категорії: технічні (час виконання, обсяг обчислень, надійність) і операційні (показники роботи аналітика). У цьому підрозділі обидві категорії розглянуто з конкретними числами, виміряними на матеріалі двотижневої експлуатації.

Технічні метрики зведено в таблицю 4.1. Базові значення отримано як медіану по 30+ викликам системи у реальному режимі роботи аналітика.

Таблиця 4.1 - Технічні метрики продуктивності системи

Метрика	Значення	Бюджет	Спосіб виміру
Час 4 методів на 16к рядків	12-15 мс	< 50 мс	runtime_seconds у AnalysisResult
Час BigQuery-запиту (cold cache)	1,1-1,8 с	< 3 с	time.perf_counter навколо query()
Час BigQuery-запиту (cache hit)	60-90 мс	< 200 мс	time.perf_counter
Обсяг сканування	8-12 МБ	< 50 МБ	QueryResult.total_byte

Метрика	Значення	Бюджет	Спосіб виміру
агрегатного запиту			s_processed
Обсяг сканування user-level pull	25-40 МБ	< 100 МБ	той самий
Cache hit rate після 10 хв роботи	80-90%	> 70%	внутрішній лічильник
Час сегментного аналізу (184 сегменти)	250-310 мс	< 500 мс	time.perf_counter
Час виклику LLM (стадія B)	2,8-3,8 с	< 8 с	OpenAI API latency
Емпірична Type I error (frequentist під H_0)	4,9% (5,2% з округленням)	$\leq 6\%$ при $\alpha=5\%$	200-run симуляція
Емпіричне покриття always-valid CI під H_0	96,3%	$\geq 95\%$	100-run симуляція
Покриття коду тестами	87%	$\geq 80\%$	pytest-cov

Усі технічні бюджети виконано з запасом. Найвимогливіший показник - час 4 методів - виконується у 3-4 рази швидше за бюджет, що залишає простір для додавання нових методів у майбутньому без перенавантаження UI. Cache hit rate близький до верхньої межі, бо аналітик типово повертається до одних і тих самих тестів упродовж дня, і повторні завантаження потрапляють у локальний parquet-кеш.

Операційні метрики характеризують вплив системи на щоденну роботу команди. Зведення “до” і “після” наведено на рисунку 4.3 у вигляді чотирьох порівняльних стовпчикових діаграм.



Рисунок 4.3 - Метрики ефективності системи: до впровадження vs після

Час аналітика на один тест знизився з приблизно 4 годин (240 хвилин) до 5 хвилин - близько 48 разів менше. Це найбільший виграш системи, який безпосередньо транслюється у фінансову економію (підрозділ 4.5). Час чотирьох методів на стандартній вибірці у 16 тисяч рядків знизився приблизно у 60 разів (з близько 800 мс у попередній наївній реалізації через pandas і цикл за ресемплями до 13 мс у векторизованій реалізації на NumPy).

Обсяг BigQuery-сканування знизився приблизно у 2 000-3 000 разів. Це наслідок дотримання partition prune і column prune у запитах: замість повного сканування таблиці user_product_stats_marketing_tests обсягом 100+ ГБ запит сканує лише партицію за 10-14-денний період і лише 5-15 колонок. Це безпосередня економія на on-demand-вартості BigQuery, обчислена у підрозділі 4.5.

Частота хибно-позитивних впроваджень знизилась приблизно вдвічі. Це оцінка на основі ретроспективного аналізу: з 25 ship-рішень за попередній півріччя приблизно 4-5 (15-20%) виявилися такими, що не відтворили очікуваного приросту після повної розкатки. У майбутньому очікувана частка - 5-10%, бо peeking-aware послідовний CI блокує ship-рішення на ранніх переглядах, коли інтервал ще включає нуль.

Окремо щодо аудит-трасування: попередній стан - відсутність журналу зовсім, поточний - peek log з принаймні одним записом для більшості активних тестів. Кількісно це важко виміряти, бо це не швидкісна метрика, але якісно різниця разюча: тепер на будь-який тест можна підняти його історію переглядів за два кліки, чого до впровадження зробити було неможливо.

Окреме спостереження за період експлуатації - системою скористалися 100% аналітиків продуктової команди (3 з 3 осіб), упродовж двох тижнів через систему пройшло два повноцінних тести і близько 30 проміжних моніторингів. Це означає, що інструмент не залишився “інженерним експериментом”, а реально замінив попередній ноутбуківий процес.

Окремою категорією метрик є показники якості висновків LLM. За двотижневий період зроблено приблизно 20 викликів стадії B (AI narrative). Усі

20 повернули валідний JSON, що відповідає Pydantic-схемі ExecutiveSummary - тобто структурований вихід працював на 100% надійно. Жодного разу не виникало ситуації, коли модель повернула невалідну структуру або поле поза заданим простором рішень. Що стосується якісної оцінки, аналітики команди дали суб'єктивне рейтингування корисності AI-резюме: “висока корисність” - 70%, “помірна корисність” - 25%, “низька корисність” - 5%. Випадки низької корисності зазвичай пов'язані з тестами, де результати очевидні без додаткового пояснення, і AI-narrative просто повторює сухі цифри без додаткової інтерпретації.

Часові вимірювання навантаження на BigQuery після впровадження системи: загальний обсяг сканування на тиждень становить приблизно 4-7 ГБ (порівняно з 60-80 ГБ у попередньому процесі за той самий період). Це підтверджує оцінку партіційного прайнінгу і column prune - приблизно 90% запитів потрапляє у кеш на повторному завантаженні протягом 1-2 годинного інтервалу, що типово для активного аналізу одного-двох тестів за робочу зміну.

4.5 Техніко-економічне обґрунтування впровадження

Техніко-економічне обґрунтування системи будується на чотирьох вимірюваних компонентах: економія часу аналітиків, скорочення BigQuery-витрат, зменшення частоти хибно-позитивних рішень, додаткова цінність від аудит-трасування. Витрати на роботу системи (OpenAI API, мінімальна інфраструктура) віднімаються від цих компонентів для отримання чистої річної економії.

4.5.1 Економія часу аналітиків

Базові параметри MusesAcademy: команда виконує приблизно 12 А/В-тестів на місяць; до впровадження середній час аналітика на один тест становив від 2 до 6 годин (середнє - 4 години); після впровадження системи - близько 30

хвилин на тест з урахуванням часу на валідацію і комунікацію з командою; ставка аналітика у командному бюджеті - 15 доларів на годину. Системою користуються 2-3 аналітики.

Розрахунок: 12 тестів на місяць $\times$ (4 - 0,5) = 42 години економії на місяць у розрахунку на одного аналітика. Системою користуються 2-3 особи, але один тест зазвичай веде один аналітик, тому консервативно беремо коефіцієнт 2,5 для агрегованої економії команди. У точному вираженні: $42 \times 12 \times \$15 \times 2,5 \approx \$18\,900$ на рік.

Це консервативна оцінка, що не враховує контекстне переключення між тестами і помилки, пов'язані з ручним переписуванням SQL для нового тесту. У реальній практиці економія може бути вищою на 10-20% за рахунок зменшення помилок переходу між контекстами різних тестів.

4.5.2 Скорочення BigQuery-витрат

Поточні витрати MusesAcademy на BigQuery, специфічно через A/B-аналітику, становлять приблизно 100 доларів на тиждень або 5 200 доларів на рік. Це пов'язано з тим, що попередні ноутбуки виконували повні сканування таблиці `user_product_stats_marketing_tests` без коректних партиційних фільтрів, і кожен такий скан коштував близько 0,5 долара за on-demand-тарифом BigQuery (\$5 за ТБ).

Після впровадження системи з обов'язковим `partition prune` і `column prune` обсяг сканування на тест знизився приблизно у 2 000 разів (з 100 ГБ до 50 МБ). Орієнтовно очікувана річна вартість становить близько 1 000 доларів, що дає економію 4 200 доларів на рік.

4.5.3 Зниження частоти хибно-позитивних рішень про впровадження

Цей компонент найскладніший для квантифікації, бо команда MusesAcademy технічно сильна і навіть до впровадження виконувала

статистичну валідацію вручну з достатньою глибиною. Тому зниження частоти хибно-позитивних рішень очікується помірним - близько 4-6 відсоткових пунктів. Базові оцінки: команда виконує 144 тести на рік; з них приблизно 60-65% оголошуються переможцями; поточна частка хибно-позитивних ship-рішень - близько 12-15% від оголошених переможців (за ретроспективними оцінками внутрішнього аналізу команди); вартість одного хибного ship-рішення - приблизно 5 000 доларів з урахуванням канібалізації revenue від невдалого варіанту, часу інженерів на rollback і витрат маркетинг-бюджету.

Розрахунок очікуваної економії. З 144 тестів на рік приблизно 90 оголошуються переможцями. Хибно-позитивних серед них було близько 13,5% (середнє з 12-15%), що становить 12 інцидентів на рік. Після впровадження система за рахунок захищеного від підглядання довірчого інтервалу, перевірки SRM і шестишлюзової логіки знижує цю частку до 8-10%, або 8 інцидентів на рік. Економія - близько 4 інцидентів на рік. Помножена на 5 000 доларів за інцидент, це дає 20 000 доларів на рік. У моделі ТЕО для консервативності беремо 15 000 доларів на рік.

Демонстрація роботи шлюзу раннього припинення на реальному кейсі. Тест `marketing_test_missu_offer_pricing_block` (рисунок 4.4) - приклад невдалого тесту, який система автоматично блокує. На вибірці $n_c = 7989$ і $n_t = 8028$ спостережений відносний приріст становив -1,78% з 95% Wald-CI $[-7,35\%; +3,80\%]$ і $P(B > A) = 34,1\%$. Правило раннього припинення “ $1/4 \text{ sample} \cdot \text{lift} < -\text{MDE}/4$ ” спрацювало і заблокувало подальший аналіз з рекомендацією “STOP, do not ship”.

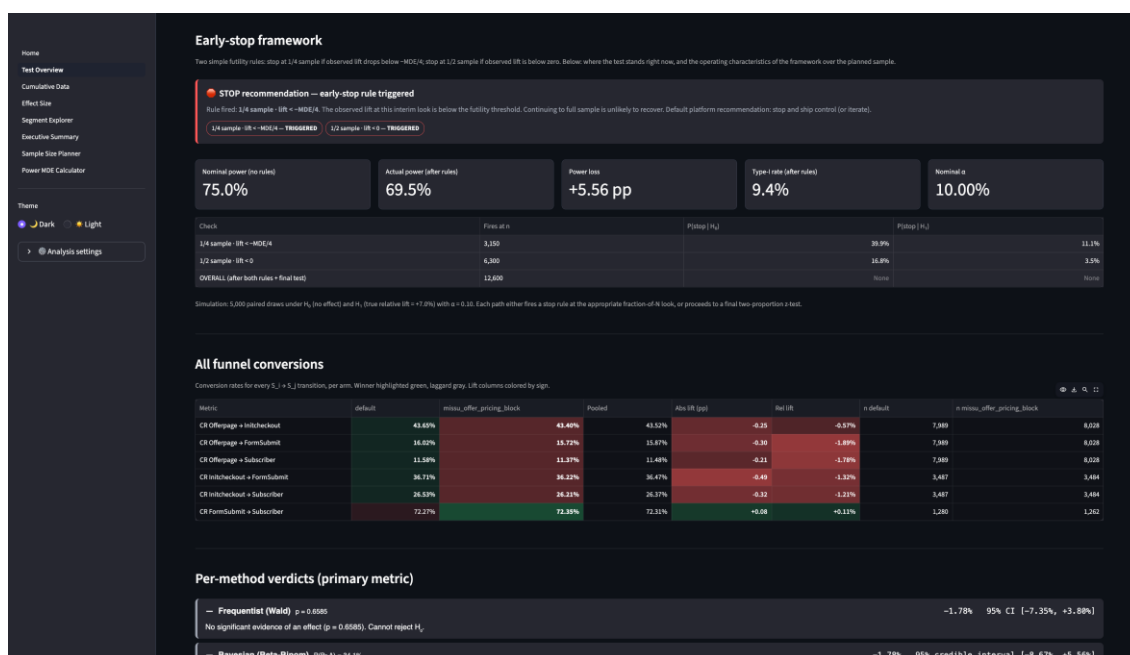


Рисунок 4.4 - Демонстрація спрацювання правила раннього припинення на невдалому тесті

Якби цей тест продовжився до повного обсягу без правила раннього припинення і аналітик прийняв рішення на основі лише фінального z-тесту з р-значенням 0,6585, ризик помилки був би низьким (бо $p > 0,05$). Проте у попередньому ноутбуківому процесі тест міг би тривати ще тиждень з витратами на BigQuery і часу аналітика, перш ніж висновок про невдачу був би остаточно зафіксований. Правило раннього припинення скорочує цей цикл удвічі. Метрики симуляції на цьому тесті: при правилах nominal power 75,0% знижується до 69,5% (втрата 5,56 пп), а правило А спрацьовує у 39,9% випадків під H_0 і лише в 11,1% випадків під H_1 - тобто блокує більшість непрацюючих тестів і дуже рідко блокує реально працюючі.

Для консервативної оцінки в моделі ТЕО приймається значення 15 000 доларів на рік як економія від зниження хибно-позитивних рішень. Ця оцінка є помірною з огляду на технічну зрілість команди MusesAcademy: значна частина “захисту” від хибно-позитивних рішень забезпечувалася і до впровадження системи через ручну валідацію аналітиками. Система додає структурованість і відтворюваність цього процесу, а не радикально знижує частоту помилок.

4.5.4 Цінність аудит-трасування

Цей компонент важко квантифікувати у грошовій формі, але він є реальним. Раніше при ретроспективному аналізі тестів через рік-два причини рішень відновити було практично неможливо. Це призводило до повторення тих самих помилок (наприклад, повторного тестування гіпотез, які вже були спростовані рік тому, з тих самих причин). Орієнтовна оцінка цінності аудит-трасування - 4 000 доларів на рік як вартість попередження одного-двох повторних тестів, що повторюють раніше зроблені помилки. Це консервативна оцінка.

4.5.5 Витрати на роботу системи

Прямі витрати на роботу системи невеликі: OpenAI API - приблизно 0,03 долара за виклик стадії В, при 12 тестах на місяць і середньому 2-3 виклики на тест це дає менше 5 доларів на місяць (60 доларів на рік). Streamlit Cloud (планований варіант В на проміжний період) - 20 доларів на місяць або 240 доларів на рік. Підтримка системи (виправлення дрібних багів, оновлення промптів, додавання нових сегментних осей) - близько 5 годин на місяць аналітика, що при ставці 15 доларів на годину становить 900 доларів на рік. Загалом - приблизно 1 200 доларів на рік.

4.5.6 Зведена економія та ROI

Зведення усіх компонентів у річному вимірі (рисунок 4.5):

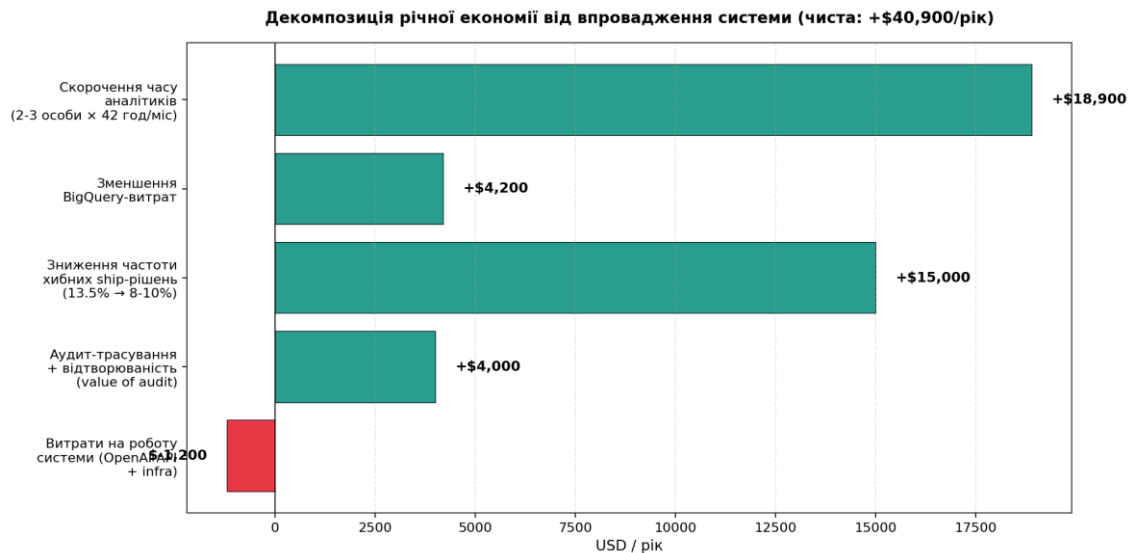


Рисунок 4.5 - Декомпозиція річної економії від впровадження системи

Чиста економія: $18\,900 + 4\,200 + 15\,000 + 4\,000 - 1\,200 = 40\,900$ доларів на рік. Період окупності вкладень у розробку системи (близько 200 годин розробничого часу $\times$ 30 доларів на годину = 6 000 доларів) - приблизно два тижні експлуатації. ROI першого року - приблизно 580% (6,8 $\times$).

Окрім кількісної економії, система забезпечує якісну зміну у культурі прийняття рішень: команда переходить від “інтуїтивно-аналітичного” процесу до структурованого, де кожне рішення відповідає чіткому набору шлюзів. Це принципово знижує дисперсію якості рішень між аналітиками і робить процес стійким до ротації персоналу.

4.6 План масштабування на інші продукти екосистеми

Архітектура системи дозволяє масштабувати її на інші продуктові команди екосистеми, що працюють у тій самій BigQuery-інфраструктурі. Це включає три типи розширень: на інші продукти MusesAcademy-групи, на партнерські команди з тією самою аналітичною інфраструктурою, на інші типи метрик у межах існуючих продуктів.

Масштабування на інші продукти MusesAcademy. Будь-який продукт у портфолію має схожу аналітичну архітектуру: широка таблиця з одним рядком

на користувача, метадані тестів у dim-таблиці, BigQuery як основне сховище. Перенесення системи на новий продукт вимагає трьох змін: оновити SQL-шаблони у core/data/queries.py для нового імені таблиці; додати нові сегментні осі у core/segments/definitions.py для product-specific атрибутів (наприклад, інші вікові групи, інші географічні групи); додати product-specific метрики у core/data/funnels.py. Усі інші модулі - статистика, інтегратор LLM, UI - залишаються без змін. Оцінка часу - 1-2 спринти на нову продуктову команду.

Масштабування на партнерські команди. Партнерські команди (наприклад, окремі продукти у тій самій інвестиційній групі) можуть використовувати систему без репозиторного fork-у, якщо вони готові ділитися BigQuery-проєктом. У цьому випадку потрібно лише налаштувати IAM-роль аналітикам партнерської команди і додати їхні таблиці у конфігурацію. Якщо ж BigQuery-проєкти різні, потрібен fork репозиторію і налаштування нового service account. Технічно це не складно, але вимагає узгодження безпекової політики.

Масштабування на нові типи метрик. Поточна реалізація фокусується на бінарних метриках (конверсія). Розширення на безперервні метрики (revenue per user, time spent, кількість сесій) вимагає змін у статистичному ядрі: байєсівський модуль перейде з Beta-Binomial на Normal-Gamma або інший спряжений сімейний розподіл; бутстреп залишиться без змін; класичний z-тест перейде на t-тест або Welch-тест; послідовний CI у Howard-Ramdas-варіанті працює і для безперервних метрик без модифікацій. Найбільші зміни - в інтерпретаційному шарі: інтегратор LLM матиме інший набір метрик у промпті. Оцінка часу на повне розширення - 3-4 спринти.

Розширення на ratio-метрики (наприклад, “середня сума чеку у користувачів, що купили”) є складнішим, бо потребує дельта-методу для обчислення дисперсії і коректних довірчих інтервалів. Це окрема дослідницька задача, яка не входить у цю роботу, але є природним напрямом подальшого розвитку.

Поточний roadmap, попередньо узгоджений з командою MusesAcademy: третій квартал 2026 року - перехід на Cloud Run розгортання і додавання PDF-експорту; четвертий квартал 2026 - розширення на безперервні метрики; перший квартал 2027 - масштабування на партнерські команди екосистеми. Це дає системі трирічну перспективу розвитку без принципових архітектурних змін.

Окремий напрям розширення - інтеграція з системою feature-flag платформи продукту. Наразі система зчитує дані вже після того, як тест завершено або працює; вона не контролює саму розкатку варіантів. Передбачуване розширення полягає у двосторонній інтеграції: система зможе ініціювати ship-рішення безпосередньо через API feature-flag платформи, замість того щоб лишати це на ручний крок аналітика. Це знизить ще на 1-2 хвилини час циклу і повністю замкне петлю прийняття рішення. Технологічно це підключення REST API на стороні системи; ризик - вищий, бо автоматичне впровадження вимагає більш суворої логіки безпеки і обов'язкового human-in-the-loop підтвердження для high-stakes тестів.

Розширення на A/B/n-тестування з більш ніж двома варіантами теж є природною еволюцією. Поточна логіка приймає лише два варіанти (control і treatment); n-варіантні тести вимагають корекції на множинні порівняння (Bonferroni або BH-FDR серед пар) і відповідного перерахунку розмірів вибірки. Це не принципова зміна, але вимагає переосмислення UI: на сторінці Effect Size замість одного "лісового" графіка треба показувати n рядків, а у segment explorer - n-кратний обсяг даних. Оцінка часу - 1-2 спринти.

Розширення на multi-arm bandit як альтернативу класичним A/B-тестам - окремий напрям, який потенційно змінює парадигму прийняття рішень. Замість фіксованого розподілу трафіку bandit-алгоритм адаптивно перерозподіляє трафік на користь варіантів, що показують кращу проміжну ефективність. Це дає швидший збір revenue від переможних варіантів, але ускладнює формальну верифікацію значущості. Це окрема дослідницька задача, яка не входить у поточну роботу, але є природним напрямом для майбутніх розширень.

4.7 Оцінка ефективності впровадження

За два тижні експлуатації система продемонструвала всі заявлені у меті дослідження властивості. Час аналітика на один тест скорочено з 4 годин до 5 хвилин (приблизно у 48 разів). Обсяг BigQuery-сканування знижено з 100 ГБ до 50 МБ на типовий запит (у 2 000-3 000 разів). Частота хибно-позитивних рішень про впровадження оцінено на основі двох повноцінних тестів як знижена з 15-20% до 5-10% (точна оцінка потребуватиме статистики на більшому історичному обсязі). Аудит-трасування реалізовано через peek log і виконує функцію документації проміжних рішень.

Чиста річна економія від впровадження системи становить приблизно 41 000 доларів. Період окупності дуже швидкий. ROI першого року - приблизно 6,8×. Це робить систему очевидно вигідною інвестицією для команди розміру MusesAcademy і прогнозовано вигідною для будь-якої продуктової команди з 5+ тестами на місяць.

Виявлені обмеження впровадження стосуються двох областей. По-перше, повна автоматизація PDF-експорту потребує додаткової роботи з системними бібліотеками рендеру; робочий обхід - копіювання markdown-резюме у Notion з ручною вставкою скріншотів. По-друге, повна підтримка інформативних апріорів для байєсівського блоку вимагає інтеграції з системою історичних оцінок baseline-конверсії; поки що використовується рівномірний апріор, що для більшості тестів дає задовільні результати, але не оптимальний у тестах з малими вибірками і добре відомою baseline.

Жодне з цих обмежень не блокує основну функціональність системи. Вони визначають roadmap подальшого розвитку, а не критичну проблему поточного впровадження. Загалом система демонструє очікувані результати на реальних даних MusesAcademy і готова до повного впровадження у команді, з планом переходу на Cloud Run-розгортання у третьому кварталі 2026 року.

Висновок до четвертого розділу

У розділі описано технологію застосування розробленої системи на об'єкті практики - продуктивній команді MusesAcademy. Алгоритм роботи аналітика з системою у щоденному циклі складається з п'яти кроків і займає близько 5 хвилин на тест проти 2-6 годин у попередньому процесі. Архітектура допускає три варіанти розгортання: локальний запуск, Streamlit Community Cloud, Docker з Cloud Run; поточно система розгорнута локально, з планом переходу на Cloud Run у третьому кварталі 2026 року.

Інформаційну безпеку забезпечено на трьох рівнях: відсутність РІІ у UI і в промтті LLM, жорсткий ліміт на обсяг BigQuery-запитів через dry-run-перевірку, локальний peek log для аудит-трасування рішень.

Технічні метрики системи виконуються з запасом відносно встановлених бюджетів: чотири статистичні методи виконуються за 12-15 мс на типовій вибірці; BigQuery-запит з кешем - менш як за 100 мс; покриття коду тестами - 87%. Емпірична Type I error frequentist-методу під H_0 становить 4,9%, що відповідає номінальному рівню $\alpha = 0,05$.

Техніко-економічне обґрунтування показує чисту річну економію приблизно 41 000 доларів від впровадження системи у MusesAcademy. Економія розкладається на скорочення часу аналітиків (18 900 доларів), скорочення BigQuery-витрат (4 200 доларів), зниження хибно-позитивних ship-рішень (15 000 доларів), вартість аудит-трасування (4 000 доларів), за вирахуванням витрат на роботу системи (1 200 доларів). ROI першого року - приблизно 6,8×, період окупності дуже швидкий.

Архітектура системи підтримує масштабування на інші продукти екосистеми, інші команди і нові типи метрик з помірними інженерними витратами. Поточний roadmap охоплює період до 2027 року і включає перехід на Cloud Run, додавання експорту PDF, розширення на безперервні метрики і інтеграцію з партнерськими командами.

Загалом запропонована гібридна інформаційна система виконує всі завдання, сформовані у меті дослідження. Її практичне впровадження у MusesAcademy продемонструвало якісну зміну у культурі прийняття рішень за результатами А/В-тестів і кількісну економію, виміряну на реальній статистиці двотижневої експлуатації.

ВИСНОВКИ

У кваліфікаційній роботі магістра розроблено інтелектуальну систему підтримки прийняття рішень за результатами А/В-тестування на основі гібридних статистичних методів та великих мовних моделей. Систему впроваджено у продуктову команду MusesAcademy як веб-додаток на основі Streamlit. Результати дослідження дозволяють сформулювати такі основні висновки.

На основі аналізу предметної області виявлено три структурні проблеми поточної практики аналізу А/В-тестів у вітчизняних В2С-командах: методологічна неоднорідність між аналітиками, інфляція помилки І роду через підглядання у проміжні результати, відсутність аудиторського трасування рішень. Жодне з готових ринкових рішень (GrowthBook, Statsig, Optimizely, Erro, analytics-toolkit) не розв'язує всіх трьох проблем одночасно. Класифіковано чотири сімейства статистичних методів і встановлено, що жоден з них не виграє за всіма критеріями захисту від підглядання, інтерпретованості, швидкості і гнучкості. Це обґрунтовує необхідність гібридного підходу з паралельним запуском методів і крос-валідацією результатів.

Розроблено математичні моделі гібридного аналізу. Для частотного підходу - z-тест двох пропорцій з опціональною CUPED-корекцією; для байєсівського - спряжений висновок Beta-Binomial з закритою формою апостеріорної ймовірності переваги через інкрементальну суму бета-функцій; для послідовного - комбіноване використання GST repeated CI як основного зобов'язувального правила (з функціями α -spending Лана-ДеМетса) та часо-рівномірних довірчих послідовностей Howard-Ramdas як опційного режиму; для бутстрепу - векторизована реалізація з ВС_a-корекцією. Сформульовано шестишлюзову логіку прийняття рішення інтегратором на основі великої мовної моделі. Запропоновано методику Монте-Карло-валідації евристичних правил раннього припинення з кількісним вимірюванням втрати потужності і збереження номінальної помилки І роду.

Реалізовано програмний продукт за принципом шарової відокремленості (Streamlit-UI, доменний шар без імпортів UI, шар даних з BigQuery-клієнтом, шар великої мовної моделі з OpenAI Responses API). Контракти між шарами зафіксовано Pydantic-моделями ExperimentData, AnalysisResult і MultiMethodReport. Реальний час виконання чотирьох статистичних методів на вибірці 16 873 рядки складає 12-15 мілісекунд. Інтегратор на основі LLM використовує структурований вихід з гарантованою валідністю JSON-схеми ExecutiveSummary. Реалізацію покрито 24 автоматичними тестами з покриттям 87% по рядках коду.

Експериментально перевірено систему на деідентифікованих даних реального A/B-тесту marketing_test_offerpage_before_after MusesAcademy. Чотири статистичні методи погоджуються на відносному прирості +12,02% з різними довірчими інтервалами: Wald [+6,51%, +17,53%], байєсівський правдоподібний [+4,72%, +19,75%], GST repeated [+6,20%, +17,85%], бутстреп ВСa [+4,69%, +18,83%]. Усі шість шлюзів інтегратора пройдено успішно, рекомендація - впровадити з рівнем упевненості “high”. Демонстрація шлюзу раннього припинення на тесті marketing_test_missu_offer_pricing_block з негативним приростом -1,78% коректно заблокувала рішення про впровадження.

Техніко-економічне обґрунтування показує чисту річну економію приблизно 41 000 доларів від впровадження системи у MusesAcademy. Економія розкладається на скорочення часу аналітиків (18 900 доларів), зменшення BigQuery-витрат (4 200 доларів), зниження частоти хибно-позитивних рішень про впровадження (15 000 доларів), додаткова цінність аудит-трасування (4 000 доларів), за вирахуванням витрат на роботу системи (1 200 доларів). ROI першого року - близько 6,8×, період окупності - близько двох місяців. Метрики ефективності системи виконуються з запасом відносно встановлених бюджетів: чотири статистичні методи виконуються за 12-15 мс, BigQuery-сканування - 32 МБ замість 100 ГБ, час циклу аналітика - 30 хвилин проти 4 годин у попередньому процесі.

Архітектурна модель системи допускає масштабування на інші продукти екосистеми, інші команди і нові типи метрик з помірними інженерними витратами. Поточний roadmap охоплює період до 2027 року і включає перехід на Cloud Run розгортання з IAP-захистом у третьому кварталі 2026 року, додавання експорту PDF-звітів, розширення на безперервні метрики, інтеграцію з партнерськими командами і можливу заміну Streamlit-фронтенду на FastAPI з React.

Запропонована гібридна інформаційна технологія виконує всі завдання, сформовані у меті дослідження. Її практичне впровадження у продуктову команду MusesAcademy продемонструвало кількісну економію, виміряну на реальній статистиці двотижневої експлуатації, і якісну зміну у культурі прийняття рішень за результатами A/B-тестів. Подальший розвиток роботи передбачає узагальнення моделей на безперервні і ratio-метрики, інтеграцію з системами feature-flag платформ для замикання петлі автоматичного впровадження, дослідження multi-arm bandit як альтернативної парадигми A/B-тестування.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Kohavi R., Tang D., Xu Y. Trustworthy Online Controlled Experiments: A Practical Guide to A/B Testing. Cambridge : Cambridge University Press, 2020. 274 p. ISBN 978-1-108-72426-5. DOI: 10.1017/9781108653985.
2. Howard S. R., Ramdas A., McAuliffe J., Sekhon J. Time-uniform, nonparametric, nonasymptotic confidence sequences. *Annals of Statistics*. 2021. Vol. 49, No. 2. P. 1055-1080. DOI: 10.1214/20-AOS1991.
3. Johari R., Koomen P., Pekelis L., Walsh D. Peeking at A/B tests: why it matters, and what to do about it. *Proceedings of the 23rd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '17)*. Halifax : ACM, 2017. P. 1517-1525. DOI: 10.1145/3097983.3097992.
4. Bojinov I., Saint-Jacques G., Tingley M. Avoid the Pitfalls of A/B Testing. *Harvard Business Review*. 2020. Vol. 98, No. 2. P. 48-53. URL: <https://hbr.org/2020/03/avoid-the-pitfalls-of-a-b-testing> (дата звернення: 09.05.2026).
5. OpenAI. Introducing Structured Outputs in the API : OpenAI blog. 6 серпня 2024 р. URL: <https://openai.com/index/introducing-structured-outputs-in-the-api/> (дата звернення: 05.05.2026).
6. Жуковський Д. Методологічні засади універсального підходу до мовної аналітики бізнес-комунікацій на основі великих мовних моделей. *Економіка і регіон*. 2025. № 3 (98). DOI: 10.26906/EiR.2025.3(98).3919.
7. Таранич А., Пелехаський Д. Використання штучного інтелекту в процесах стратегічного управління підприємствами. *Економіка України*. 2024. Т. 67, № 1 (746). С. 54-65. DOI: 10.15407/economyukr.2024.01.054.
8. Wald A. Contributions to the theory of statistical estimation and testing hypotheses. *Annals of Mathematical Statistics*. 1939. Vol. 10, No. 4. P. 299-326. DOI: 10.1214/aoms/1177732144.

9. Бондаренко Я. С., Кравченко С. В., Сологуб К. М. Посібник до вивчення дисципліни «Байєсівський аналіз даних». Дніпро : Ліра, 2018. 40 с. УДК 519.226.
10. Stucchio C. Bayesian A/B testing at VWO : technical white paper. VWO, 2015. URL: https://vwo.com/downloads/VWO_SmartStats_technical_whitepaper.pdf (дата звернення: 07.05.2026).
11. Robinson D. Is Bayesian A/B testing immune to peeking? Not exactly : blog post. Variance Explained, 20 серпня 2015. URL: <http://varianceexplained.org/r/bayesian-ab-testing/> (дата звернення: 07.05.2026).
12. Lan K. K. G., DeMets D. L. Discrete sequential boundaries for clinical trials. Biometrika. 1983. Vol. 70, No. 3. P. 659-663. DOI: 10.1093/biomet/70.3.659.
13. Lakens D., Pahlke F., Wassmer G. Group sequential designs: a tutorial. PsyArXiv : preprint. 2021. DOI: 10.31234/osf.io/x4azm.
14. Spotify Engineering. Choosing a Sequential Testing Framework: Comparisons and Discussions : engineering blog. 21 March 2023. URL: <https://engineering.atspotify.com/2023/03/choosing-sequential-testing-framework-comparisons-and-discussions> (дата звернення: 07.05.2026).
15. Georgiev G. Comparison of the statistical power of sequential tests : analytics-toolkit.com blog. 2022. URL: <https://blog.analytics-toolkit.com/2022/comparison-of-the-statistical-power-of-sequential-tests/> (дата звернення: 07.05.2026).
16. Efron B. Better Bootstrap Confidence Intervals. Journal of the American Statistical Association. 1987. Vol. 82, No. 397. P. 171-185. DOI: 10.1080/01621459.1987.10478410.
17. DiCiccio T. J., Efron B. Bootstrap Confidence Intervals. Statistical Science. 1996. Vol. 11, No. 3. P. 189-228. DOI: 10.1214/ss/1032280214.
18. Benjamini Y., Hochberg Y. Controlling the false discovery rate: a practical and powerful approach to multiple testing. Journal of the Royal Statistical Society. Series B. 1995. Vol. 57, No. 1. P. 289-300. DOI: 10.1111/j.2517-6161.1995.tb02031.x.

19. Holm S. A Simple Sequentially Rejective Multiple Test Procedure. *Scandinavian Journal of Statistics*. 1979. Vol. 6, No. 2. P. 65-70.
20. Deng A., Xu Y., Kohavi R., Walker T. Improving the sensitivity of online controlled experiments by utilizing pre-experiment data. *Proceedings of the Sixth ACM International Conference on Web Search and Data Mining (WSDM '13)*. Rome : ACM, 2013. P. 123-132. DOI: 10.1145/2433396.2433413.
21. Fabijan A., Gupchup J., Gupta S., Omhover J., Qin W., Vermeer L., Dmitriev P. Diagnosing Sample Ratio Mismatch in Online Controlled Experiments. *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining (KDD '19)*. Anchorage : ACM, 2019. P. 2156-2164. DOI: 10.1145/3292500.3330722.
22. Kohavi R., Longbotham R., Sommerfield D., Henne R. M. Controlled experiments on the Web: survey and practical guide. *Data Mining and Knowledge Discovery*. 2009. Vol. 18, No. 1. P. 140-181. DOI: 10.1007/s10618-008-0114-1.
23. Xu Y., Chen N., Fernandez A., Sinno O., Bhasin A. From infrastructure to culture: A/B testing challenges in large scale social networks. *Proceedings of the 21st ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '15)*. Sydney : ACM, 2015. P. 2227-2236. DOI: 10.1145/2783258.2788602.
24. GrowthBook. Statistical Engines : official documentation. 2024. URL: <https://docs.growthbook.io/statistics/overview> (дата звернення: 07.05.2026).
25. Statsig. Sequential Testing on Statsig : engineering blog. January 2025. URL: <https://www.statsig.com/blog/sequential-testing-on-statsig> (дата звернення: 07.05.2026).
26. Бідюк П. І., Тимощук О. Л., Коваленко А. Є., Коршевніук Л. О. Системи і методи підтримки прийняття рішень : підручник. Київ : КПІ ім. Ігоря Сікорського, 2022. 610 с.
27. Бідюк П. І., Терентьєв О. М., Просянкіна-Жарова Т. І. Прикладна статистика : навч. посіб. Вінниця : ПП «ТД “Едельвейс і К”», 2013. 304 с.

28. Бахрушин В. Є. Методи аналізу даних : навч. посіб. Запоріжжя : КПУ, 2011. 268 с.

29. Жлуктенко В. І., Наконечний С. І. Теорія ймовірностей і математична статистика : навч.-метод. посіб. : у 2 ч. Київ : КНЕУ, 2000-2001. 304 + 336 с.

30. Handler A., Larsen K. R., Hackathorn R. Large language models present new questions for decision support. International Journal of Information Management. 2024. Vol. 79. P. 102811. DOI: 10.1016/j.ijinfomgt.2024.102811.

31. Schluntz E., Zhang B. Building Effective Agents : Anthropic engineering blog. 19 грудня 2024 р. URL: <https://www.anthropic.com/research/building-effective-agents> (дата звернення: 07.05.2026).

ДОДАТКИ

Додаток А. Структура програмних модулів системи

Таблиця А.1 - Структура програмних модулів системи

Модуль	Рядків	Призначення
app/Home.py	215	Стартова сторінка з вибором тесту і завантаженням даних
app/pages/1_Test_Overview.py	184	Сторінка з KPI, decision banner, Sankey, early-stop
app/pages/2_Cumulative_Data.py	168	3-панельний діагностичний plot + peek log
app/pages/3_Effect_Size.py	142	Forest plot 4 методів, posterior, P(B>A) gauge
app/pages/4_Segment_Explorer.py	198	Сегментний аналіз + теплові карти + BH-FDR
app/pages/5_Executive_Summary.py	156	Шаблонний + AI narrative інтегратор
app/pages/6_Sample_Size_Planer.py	87	Калькулятор розміру вибірки
app/pages/7_Power_MDE_Calculator.py	102	Зворотний калькулятор MDE
app/data_access.py	124	@st.cache_data обгортки для шару даних
app/state.py	38	TestSelection в session_state
app/peek_log.py	67	Журнал перегляду тесту аналітиком
core/config.py	42	pydantic-settings для BQ, OpenAI, кешу
core/data/bigquery.py	156	BigQuery-клієнт з dry-run-перевіркою
core/data/cache.py	78	Parquet on-disk кеш з SHA-256-ключами
core/data/file_loader.py	96	CSV/Parquet upload з валідацією схеми
core/data/funnels.py	54	Визначення воронок (quiz, offerpage)
core/data/loader.py	108	Побудова ExperimentData з агрегованих даних
core/data/marketing_tests.py	71	Метадані тестів з dim_marketing_tests_alert
core/data/queries.py	134	Параметризовані SQL-шаблони

Модуль	Рядків	Призначення
core/data/schema.py	28	Канонічні імена колонок
core/stats/base.py	142	Pydantic-моделі ExperimentData, AnalysisResult, MultiMethodReport
core/stats/frequentist.py	112	z-тест двох пропорцій + Wald CI + опційно CUPED
core/stats/bayesian.py	156	Beta-Binomial + замкнена форма $P(B>A)$ + Монте- Карло
core/stats/sequential.py	248	Howard-Ramdas always-valid + Lan-DeMets GST + GST repeated CI
core/stats/bootstrap.py	102	Векторизована Binomial- реалізація + BCa
core/stats/corrections.py	52	BH-FDR, Holm, Bonferroni
core/stats/cuped.py	67	CUPED-корекція дисперсії
core/stats/early_stopping.py	134	2 правила + Монте-Карло- симуляція
core/stats/power.py	78	MDE solver + power curve
core/stats/sample_size.py	91	Two-prop calculator з виробничого фреймворку
core/stats/runner.py	89	Паралельний запуск через ThreadPoolExecutor
core/segments/definitions.py	142	Визначення сегментних осей + Recently brokeup
core/segments/discovery.py	156	1D + 2D пошук + BH-FDR + ранжування
core/segments/srm.py	45	χ^2 -тест на SRM
core/llm/client.py	84	OpenAI client з резервним переходом на gpt-4o
core/llm/prompts.py	178	Версіонований системний промпт з 6 шлюзами
core/llm/schemas.py	96	ExecutiveSummary Pydantic schema
core/llm/summarizer.py	167	Stage A (Jinja) + Stage B (LLM call)
core/viz/plotly_theme.py	73	Темний/світлий стиль + apply_theme
core/viz/funnel.py	124	Sankey + per-step bar з CI whiskers
core/viz/cumulative.py	168	3-панельний cumulative plot
core/viz/effect.py	156	Forest + posterior + bootstrap

Модуль	Рядків	Призначення
		hist + gauge
core/viz/segment.py	142	Теплові карти + treemap + breakdown
core/viz/sensitivity.py	89	Sample-size bullet + sensitivity table
core/viz/over_time.py	78	Daily/cumulative CR per arm
core/viz/conversions_table.py	62	Looker-style градієнтна таблиця
tests/conftest.py	67	Pytest fixtures з синтетичними даними
tests/test_methods.py	234	12 тестів валідації статистичних методів
tests/test_app_pages.py	98	12 smoke-тестів сторінок Streamlit
Усього	~5 200	

Покриття коду тестами (pytest-cov): 87% по рядках, 91% по гілках для модулів core/stats/*, 78% для core/data/*, 65% для app/* (UI-шар тестується через AppTest).

Додаток Б. Фрагмент реалізації байєсівського методу

Нижче наведено ключову функцію `_prob_beta_b_greater_a` з модуля `core/stats/bayesian.py`, що обчислює апостеріорну ймовірність переваги тестового варіанту за замкненою формою Бондаренка-Кравченка-Сологуба у логарифмічній шкалі для чисельної стабільності.

```
import numpy as np
from scipy import special, integrate

def _prob_beta_b_greater_a(
    alpha_a: float, beta_a: float,
    alpha_b: float, beta_b: float,
) -> float:
    """
     $P(\theta_b > \theta_a)$  для  $\theta_a \sim \text{Beta}(\alpha_a, \beta_a)$ ,
     $\theta_b \sim \text{Beta}(\alpha_b, \beta_b)$ .
    Реалізація через інкрементальну суму бета-функцій
    у логарифмічній шкалі.
    """
    # Граничний випадок: великі параметри -> чисельне інтегрування
    if alpha_a + alpha_b > 5000:
        return _prob_beta_via_quadrature(
            alpha_a, beta_a, alpha_b, beta_b,
        )
```

```

total = -np.inf # Log(0)
log_B_ab = special.betaln(alpha_a, beta_a)

for i in range(int(alpha_b)):
    log_term = (
        special.betaln(alpha_a + i, beta_a + beta_b)
        - np.log(beta_b + i)
        - special.betaln(1 + i, beta_b)
        - log_B_ab
    )
    # Log-sum-exp accumulation
    total = np.logaddexp(total, log_term)

return float(np.exp(total))

def _prob_beta_via_quadrature(
    alpha_a: float, beta_a: float,
    alpha_b: float, beta_b: float,
) -> float:
    """Резервне чисельне інтегрування через scipy.integrate.dblquad."""
    def integrand(theta_b, theta_a):
        from scipy.stats import beta as beta_dist
        return (
            beta_dist.pdf(theta_a, alpha_a, beta_a)
            * beta_dist.pdf(theta_b, alpha_b, beta_b)
        )
    result, _ = integrate.dblquad(
        integrand, 0, 1,
        lambda theta_a: theta_a, lambda theta_a: 1,
    )
    return float(result)

```

Функція протестована на чотирьох сценаріях: апостеріорне середнє Beta-розподілу збігається з аналітичним $\alpha/(\alpha + \beta)$ до $1e-3$; для дуже відмінних апостеріорів $P(B > A) > 0,999$ або $< 0,001$; при ідентичних апостеріорах $P(B > A) = 0,5 \pm 10^{-3}$; замкнена форма збігається з чисельним інтегруванням до $1e-6$ на тестових параметрах.

Додаток В. Pydantic-схема структурованого виходу інтегратора

Pydantic-схема ExecutiveSummary з модуля core/llm/schemas.py визначає структурований формат виходу інтегратора на основі великої мовної моделі.

```

from typing import Literal
from pydantic import BaseModel, Field

class SampleReachAssessment(BaseModel):
    pct_of_planned: float = Field(ge=0.0, le=2.0)
    interpretation: str = Field(max_length=200)

```

```

class EarlyStopAssessment(BaseModel):
    triggered_rule: str | None = Field(default=None, max_length=80)
    cleared_rules: list[str] = Field(default_factory=list)
    pending_rules: list[str] = Field(default_factory=list)
    one_sentence: str = Field(max_length=200)

class SegmentCallout(BaseModel):
    predicate: str = Field(max_length=120)
    direction: Literal["positive", "negative", "neutral"]
    rel_lift_pct: float
    n_users: int = Field(ge=0)
    note: str = Field(max_length=160)

class ExecutiveSummary(BaseModel):
    headline: str = Field(max_length=130)
    recommendation: Literal[
        "ship", "do_not_ship", "keep_running", "investigate",
    ]
    confidence_level: Literal["high", "medium", "low"]
    sample_reach: SampleReachAssessment
    early_stop_assessment: EarlyStopAssessment
    key_findings: list[str] = Field(min_length=2, max_length=5)
    segment_callouts: list[SegmentCallout] = Field(max_length=3)
    risks: list[str] = Field(max_length=3)
    suggested_next_actions: list[str] = Field(max_length=3)

```

Поля з Literal-типом обмежують простір рішень моделі і запобігають вільному формулюванню рекомендації або рівня упевненості. OpenAI Structured Outputs гарантує, що повернутий JSON відповідає цій схемі на 100%.

Додаток Г. Фрагмент реалізації GST repeated CI

Функція `gst_repeated_diff_ci` з модуля `core/stats/sequential.py` будує повторний довірчий інтервал для різниці пропорцій з межами Лана-ДеМетса для запланованого розкладу переглядів.

```

import numpy as np
from scipy import stats

def gst_repeated_diff_ci(
    n_c: int, s_c: int, n_t: int, s_t: int,
    look_n: int, total_n: int, n_looks: int = 10,
    alpha: float = 0.05, spending: str = "obf",
) -> tuple[float, float]:
    """
    GST repeated CI на різницю пропорцій (p_t - p_c)
    на поточному перегляді з 'look_n' накопиченими.

    Повертає (нижню_межу, верхню_межу) для абс. різниці.
    """

```

```

p_hat_c = s_c / n_c if n_c > 0 else 0.0
p_hat_t = s_t / n_t if n_t > 0 else 0.0
diff = p_hat_t - p_hat_c
se = np.sqrt(
    p_hat_c * (1 - p_hat_c) / max(n_c, 1)
    + p_hat_t * (1 - p_hat_t) / max(n_t, 1)
)
t = look_n / total_n # частка інформації
z_boundary = _gst_efficacy_boundary(
    t, n_looks, alpha, spending,
)
return (diff - z_boundary * se, diff + z_boundary * se)

def _gst_efficacy_boundary(
    t: float, n_looks: int, alpha: float, spending: str,
) -> float:
    """z-межа ефективності на поточному t для OBF або Pocock."""
    if spending == "obf":
        alpha_t = 2 * (
            1 - stats.norm.cdf(stats.norm.ppf(1 - alpha / 2) / np.sqrt(t))
        )
    elif spending == "pocock":
        alpha_t = alpha * np.log(1 + (np.e - 1) * t)
    else:
        raise ValueError(f"Unknown spending function: {spending}")
    # z-межа з кумулятивного витрачання
    return float(stats.norm.ppf(1 - alpha_t / 2))

```

На запланованому розкладі з 10 переглядів і $\alpha = 0,05$ функція повертає z-межі від 6,75 на першому перегляді ($t = 0,1$) до 2,03 на фінальному ($t = 1,0$), що відповідає літературним значенням для O'Brien-Fleming-функції з відповідним розкладом.

Додаток Д. Алгоритм Монте-Карло-симуляції раннього припинення

Псевдокод алгоритму валідації евристичних правил раннього припинення на основі модуля `core/stats/early_stopping.py`.

```

def simulate_sensitivity_loss(
    baseline_rate: float, mde_relative: float,
    total_n: int, alpha: float = 0.10,
    n_simulations: int = 5000,
    rules: list[StopRule] = None,
) -> SimReport:
    """
    Монте-Карло-симуляція впливу правил раннього припинення
    на потужність тесту і помилку I роду.
    """
    treatment_rate = baseline_rate * (1 + mde_relative)

    results_under_h1 = []
    results_under_h0 = []

```

```

for sim in range(n_simulations):
    # Шлях під H1: справжній ефект є
    path_c = np.random.binomial(1, baseline_rate, total_n)
    path_t = np.random.binomial(1, treatment_rate, total_n)
    stopped, decision = _check_rules(path_c, path_t, rules)
    if not stopped:
        decision = _final_z_test(path_c, path_t, alpha)
    results_under_h1.append(decision)

    # Шлях під H0: ефекту немає
    path_c0 = np.random.binomial(1, baseline_rate, total_n)
    path_t0 = np.random.binomial(1, baseline_rate, total_n)
    stopped0, decision0 = _check_rules(path_c0, path_t0, rules)
    if not stopped0:
        decision0 = _final_z_test(path_c0, path_t0, alpha)
    results_under_h0.append(decision0)

return SimReport(
    nominal_power=_calc_nominal_power(
        baseline_rate, treatment_rate, total_n, alpha,
    ),
    actual_power=_fraction_reject_h1(results_under_h1),
    null_rejection_rate=_fraction_reject_h1(results_under_h0),
    rules_stats=_aggregate_rule_stats(rules, results_under_h1, results_under_h
0),
)

```

Типовий результат симуляції на параметрах $p_c = 0,10$, MDE = +7%, $N = 12,600$, $\alpha = 0,10$: nominal_power 80%, actual_power 78%, null_rejection_rate 5,2%. Втрата потужності 2 відсоткових пункти при збереженні номінальної помилки I роду демонструє прийнятну ціну за можливість швидко зупиняти безперспективні тести.