

**НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ імені ІГОРЯ
СІКОРСЬКОГО»**

**Навчально-науковий інститут телекомунікаційних систем
Кафедра інформаційних технологій в телекомунікаціях**

«До захисту допущено»

В.О. Завідувача кафедри

_____ **Валерій ПРАВИЛО**

« ____ » _____ 2026 р.

Дипломна робота

на здобуття ступеня бакалавра

за освітньо-професійною програмою «Інформаційно-комунікаційні технології»
спеціальності 172 «Телекомунікації та радіотехніка»

**на тему: «Метод масштабування високонавантажених сервісів мереж 5G з
використанням Multi-access Edge Computing»**

Виконав: студент IV курсу, групи ТІ-21

Хомутенко Владислав Олегович _____

Науковий керівник:

Директор НН ІТС, доктор технічних наук, професор

Скулиш Марія Анатоліївна _____

Рецензент:

Професор кафедри ТК НН ІТС доктор технічних наук, професор

Лисенко Олександр Іванович _____

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

Студент _____

**НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
«КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ ІМЕНІ ІГОРЯ
СІКОРСЬКОГО»**

Навчально-науковий інститут телекомунікаційних систем

Кафедра інформаційних технологій в телекомунікаціях

Рівень вищої освіти – перший (бакалаврський)

Освітньо-професійна програма «Інформаційно-комунікаційні технології»

Спеціальність 172 «Телекомунікації та радіотехніка»

ЗАТВЕРДЖУЮ

В.О. Завідувача кафедри

_____ Валерій ПРАВИЛО

«_____» _____ 2026 р.

ЗАВДАННЯ

на дипломну роботу студенту Хомутенку Владиславу Олеговичу

1. Тема дипломної роботи: «Метод масштабування високонавантажених сервісів мереж 5G з використанням Multi-access Edge Computing», науковий керівник роботи Скулиш Марія Анатоліївна, директор НН ІТС, доктор технічних наук, професор - затверджені наказом по університету від 26.05.2026 р. №1912-с

2. Строк подання студентом дипломної роботи 06.06.2026 р.

3. Об'єкт дослідження: процес масштабування та керування ресурсами МЕС-вузла у мережах 5G.

4. Предмет дослідження: методи планування черг, розподілу ресурсів та симуляційного моделювання МЕС-вузла для Інтелектуальної транспортної системи (ITS) смарт-міста.

5. Вихідні дані до роботи: параметри п'яти різномірних сервісів Інтелектуальної транспортної системи смарт-міста (інтенсивності λ , інтенсивності обслуговування μ , SLA-вимоги до затримки та надійності); архітектурні стандарти

ETSI MEC та 3GPP 5G; методи планування ресурсів (HOL Priority Queuing, Weighted Fair Queuing, Model Predictive Control, Elastic Scaling, Lyapunov Optimization); бібліотека дискретно-подійної симуляції SimPy для мови Python; наукові публікації у сфері edge computing та теорії масового обслуговування.

6. Перелік завдань, які потрібно розробити: проаналізувати архітектуру Multi-access Edge Computing та існуючі методи керування ресурсами; розробити математичну модель МЕС-вузла як багатокласової системи масового обслуговування типу М/М/п/т; обрати та обґрунтувати комбінацію методів планування для виконання SLA усіх сервісів; розв'язати задачу оптимального розміщення ресурсів (CPU, RAM, GPU) у формі цілочисельного лінійного програмування; реалізувати симуляційну модель в SimPy та верифікувати аналітичні розрахунки; провести статистичний аналіз результатів та сформулювати рекомендації.

7. Перелік ілюстративного матеріалу (із зазначенням плакатів, презентацій тощо)

8. Дата видачі завдання: 05.09.2025

Календарний план

№	Назва етапів виконання дипломної роботи	Термін виконання	Примітка
1	Затвердження теми. Аналітичний огляд літератури та архітектур МЕС	25.02.2025- 10.03.2025	Виконано
2	Розробка математичної моделі М/М/п/т. Аналітичні розрахунки Erlang-C	11.03.2025 - 27.03.2025	Виконано
3	Обґрунтування комбінації методів керування. Розв'язок ILP-задачі	28.03.2025 - 17.04.2025	Виконано

4	Реалізація симулятора в SimPy. Базова верифікація та повна 24- год симуляція	18.04.2025 - 08.05.2025	Виконано
5	Аналіз результатів. Формулювання рекомендацій	09.05.2025 - 15.05.2025	Виконано
6	Оформлення дипломної роботи	16.05.2025 - 05.06.2025	Виконано
7	Підготовка презентації до захисту	05.06.2025 - 09.06.2025	Виконано
8	Захист дипломної роботи	16.06.2025 - 30.06.2025	

Здобувач вищої освіти _____ Владислав ХОМУТЕНКО

Науковий керівник _____ Марія СКУЛИШ

РЕФЕРАТ

Дипломна робота містить 96 сторінок, 11 рисунків, 19 таблиць, 19 формул, 4 додатки. Використано 32 джерела інформації.

Актуальність теми. Поява п'ятого покоління мобільного зв'язку та його тісна інтеграція з технологією багатоступеневої крайової обробки даних перевертає традиційний поділ між абонентським пристроєм і центральним дата-центром: тепер вагома частка обчислювальної логіки виконується на крайовому вузлі за лічені мільйони від кінцевого користувача. Внаслідок такого наближення ланцюг «запит → відповідь» вкорочується з типових для хмари 50–350 мс до 10–50 мс. Та сама зміна, втім, висуває й нові виклики до інженера-проектувальника: крайовий вузол має скромнішу ресурсну конфігурацію, проте на ньому одночасно живуть кілька неоднорідних сервісів — кожен зі своїм потоком запитів, рівнем критичності та контрактними обмеженнями на максимально допустиму затримку. Найпростіші дисципліни обслуговування на кшталт по-черговості або фіксованого виділення каналів швидко вичерпують свій ресурс гнучкості, коли навантаження міняється у кілька разів протягом доби, а критичність сервісів вимагає диференційованого підходу. Це обґрунтовує потребу в інженерному методі, що поєднує пріоритезацію, справедливий розподіл спільних ресурсів та превентивне масштабування.

Мета дипломної роботи - підвищення ефективності масштабування обчислювальних потужностей крайового МЕС-вузла для забезпечення виконання контрактних зобов'язань щодо затримки та надійності п'яти одночасно працюючих сервісів Інтелектуальної транспортної системи смарт-міста в межах заданого ресурсного бюджету. Об'єктом дослідження є процедура динамічного перерозподілу обчислювальних потужностей крайового вузла мережі 5G під час обслуговування багатоступеневого потоку запитів.

Предметом дослідження виступає сукупність планувальних дисциплін та механізмів розподілу обчислювальних, оперативних і графічних ресурсів на МЕС-вузлі, що обслуговує транспортну телематику міського масштабу.

Методи дослідження. Робота спирається на чотири взаємодоповнюючі методологічні складники. Перший — апарат теорії масового обслуговування: модель $M/M/n/m$ з кінцевою чергою та формула Ерланга-С для оцінки часу очікування. Другий — формалізація задачі розміщення сервісів у вигляді цілочисельного лінійного програмування з кількома ресурсними обмеженнями. Третій — дискретно-подійне моделювання у середовищі SimPy (мова Python), що дає змогу спостерігати поведінку системи у реальному часі. Четвертий — статистичне опрацювання результатів імітації методом бутстрепівського ресемплінгу з побудовою 95-відсоткових довірчих інтервалів та критичне зіставлення класичної блокувальної ймовірності з часозалежною метрикою своєчасності обслуговування.

Ключові слова: MULTI-ACCESS EDGE COMPUTING, 5G, ІНТЕЛЕКТУАЛЬНА ТРАНСПОРТНА СИСТЕМА, МАСШТАБУВАННЯ, HOL PRIORITY QUEUING, WEIGHTED FAIR QUEUING, MODEL PREDICTIVE CONTROL, ELASTIC SCALING, LYAPUNOV OPTIMIZATION, ERLANG-C, $M/M/N/M$, SIMPY, SLA, ЗАТРИМКА.

ABSTRACT

The thesis consists of 96 pages, 11 figures, 19 tables, 19 formulas. 32 information sources were used.

Relevance of the topic. The arrival of the fifth-generation mobile networks, tightly coupled with multi-access edge computing, redraws the boundary between the user device and the central data centre: a significant share of the processing logic is now executed at a small node located only a few kilometres from the end user. As a consequence, the request-response loop shrinks from the cloud-typical 50–350 ms down to 10–50 ms. The same shift, however, sets new design challenges before the engineer: the edge node operates with a modest hardware envelope, yet it has to host several heterogeneous services simultaneously — each with its own arrival pattern, criticality level and contractual latency bounds. The simplest scheduling policies such as Round Robin or fixed channel allocation quickly run out of flexibility when the load varies several-fold across the day and when service criticality demands differentiated treatment. This motivates the search for an engineering method that combines prioritisation, fair sharing of common resources and proactive scaling.

The aim of the study is to develop and investigate a method of scaling MEC node resources that simultaneously satisfies the SLAs of five heterogeneous services of an Intelligent Transport System of a smart city within a given resource budget.

The object of the study is the process of scaling and resource management at a MEC node in 5G networks.

The subject of the study is the set of queueing and resource allocation methods for a MEC node serving Intelligent Transport System services in a smart city.

Research methods. The work uses queueing theory (M/M/n/m model and the Erlang-C formula); integer linear programming for the resource placement problem; discrete-event simulation (SimPy); statistical processing of results based on the bootstrap method with 95% confidence intervals; comparative analysis of the Erlang-C blocking probability and the time-dependent SLA metric $P(W \leq D_{\max})$.

Keywords: MULTI-ACCESS EDGE COMPUTING, 5G, INTELLIGENT TRANSPORT SYSTEM, SCALING, HOL PRIORITY QUEUING, WEIGHTED FAIR QUEUING, MODEL PREDICTIVE CONTROL, ELASTIC SCALING, LYAPUNOV OPTIMIZATION, ERLANG-C, M/M/N/M, SIMPY, SLA, LATENCY.

ЗМІСТ

СКОРОЧЕННЯ ТА УМОВНІ ПОЗНАКИ	12
ВСТУП.....	15
РОЗДІЛ 1 АНАЛІЗ АРХІТЕКТУРИ МЕС ТА МЕТОДІВ КЕРУВАННЯ РЕСУРСАМИ.....	21
1.1 Концепція Multi-access Edge Computing та її роль у мережах 5G.....	21
1.2 Архітектура МЕС та стандарти ETSI.....	23
1.3 Порівняння хмарної обробки та крайових обчислень.....	25
1.4 Розгортання МЕС для Інтелектуальної транспортної системи смарт-міста.	27
1.5 Опис п'яти сервісів Інтелектуальної транспортної системи.....	28
1.6 Проблема спільного використання ресурсів GPU	31
1.7 Класифікація методів планування ресурсів.....	32
1.8 Огляд методів теорії масового обслуговування.....	33
1.9 Огляд методів керування потоками: HOL, WFQ, MPC, Elastic Scaling, Луарипов.....	36
1.10 Порівняльний аналіз методів та обґрунтування вибору	39
1.11 Огляд наукових публікацій та аналогічних робіт	40
Висновки до розділу 1	43
РОЗДІЛ 2 РОЗРОБКА МЕТОДУ МАСШТАБУВАННЯ ТА МАТЕМАТИЧНА МОДЕЛЬ	44
2.1 Параметри ефективності МЕС-вузла та критерії SLA	44
2.2 Математична модель МЕС-вузла як багатокласової СМО типу M/M/n/m...	45
2.3 Формалізація метрик затримки та пропускної здатності.....	47
2.4 Формула Ерланга-С та її застосування	47
2.5 Аналітичний розрахунок W_q та W для кожного сервісу	48
2.6 Визначення мінімальних n^* та обґрунтування поточного n	50

	10
2.7 Задача оптимального розміщення ресурсів (ILP)	51
2.8 Розв'язок задачі: оптимальний розподіл CPU, RAM, GPU	52
2.9 Трирівнева архітектура адаптивного керування.....	53
2.10 Детальний опис методів HOL, WFQ, MPC+Elastic Scaling, Lyapunov	54
2.11 Код контролера МЕС-вузла	56
2.12 Аналіз складності та накладних витрат методу	57
Висновки до розділу 2	58
РОЗДІЛ 3 РЕАЛІЗАЦІЯ СИМУЛЯТОРА ТА МОДЕЛЮВАННЯ.....	60
3.1 Вибір середовища та мови програмування	60
3.2 Загальна архітектура симулятора	61
3.3 Модель сервісу як черга M/M/n/m у SimPy	62
3.4 Реалізація HOL Priority Queuing	63
3.5 Реалізація WFQ GPU Allocator	64
3.6 Реалізація MPC та Elastic Scaling.....	64
3.7 Реалізація Lyapunov Fallback	65
3.8 Параметри симуляційних сценаріїв	65
3.9 Змінний профіль навантаження $\lambda(t)$ за 24 години.....	66
3.10 Організація повторних запусків та bootstrap-аналіз	67
Висновки до розділу 3	68
РОЗДІЛ 4 АНАЛІЗ РЕЗУЛЬТАТІВ СИМУЛЯЦІЇ.....	69
4.1 Базова верифікація аналітичної моделі.....	69
4.2 Аналіз сервісу виявлення екстрених транспортних засобів	70
4.3 Bootstrap-довірчий інтервал для сервісу відеоаналізу	70
4.4 Ключове наукове відкриття: Erlang-C проти $P(W \leq D_{\max})$	71
4.5 Ефективність HOL Priority Queuing	74

	11
4.6 Динаміка Elastic Scaling за 24 години	74
4.7 Активації Lyapunov Fallback	76
4.8 Зведена таблиця адекватності повної симуляції	76
4.9 Рекомендації для виконання SLA	77
4.10 Загальна оцінка методу та напрямки подальших досліджень	78
Висновки до розділу 4	80
ЗАГАЛЬНІ ВИСНОВКИ	81
ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ	84
ДОДАТОК А ПОВНА ТАБЛИЦЯ ПАРАМЕТРІВ СЕРВІСІВ МЕС	87
ДОДАТОК Б ДЕТАЛЬНІ ОБЧИСЛЕННЯ ЕРЛАНГА-С ДЛЯ ВСІХ СЕРВІСІВ ...	89
ДОДАТОК В ЛІСТИНГИ КЛЮЧОВИХ МОДУЛІВ СИМУЛЯТОРА	92
ДОДАТОК Г РОЗШИРЕНІ РЕЗУЛЬТАТИ СИМУЛЯЦІЇ	94

СКОРОЧЕННЯ ТА УМОВНІ ПОЗНАКИ

MEC — Multi-access Edge Computing — багатоступенева крайова обробка; архітектурна концепція винесення обчислень на крайові вузли поблизу користувачів

5G — мережа мобільного зв'язку п'ятого покоління

ITS — Intelligent Transport System — Інтелектуальна транспортна система

SLA — Service Level Agreement — угода про рівень обслуговування

QoS — Quality of Service — якість обслуговування

CMO — система масового обслуговування

M/M/n/m — CMO з пуасонівським вхідним потоком, експоненційним розподілом часу обслуговування, n каналами та обмеженою чергою на m місць

HOL — Head-of-Line Priority Queuing — пріоритетне обслуговування голови черги

WFQ — Weighted Fair Queuing — зважено-справедлива дисципліна обслуговування

MPC — Model Predictive Control — керування з прогнозувальною моделлю

DPP — Drift-Plus-Penalty — критерій оптимізації Ляпунова

ILP — Integer Linear Programming — цілочисельне лінійне програмування

LP — Linear Programming — лінійне програмування

CPU — Central Processing Unit — центральний процесор

GPU — Graphics Processing Unit — графічний процесор

RAM — Random-Access Memory — оперативна пам'ять

TOPS — Tera Operations Per Second — терабайт операцій на секунду; одиниця вимірювання продуктивності AI-прискорювачів

NFV — Network Functions Virtualization — віртуалізація мережевих функцій

SDN — Software-Defined Networking — програмно-керована мережа

ETSI — European Telecommunications Standards Institute — Європейський інститут телекомунікаційних стандартів

3GPP — 3rd Generation Partnership Project — міжнародне партнерство для розробки стандартів мобільного зв'язку

LSTM — Long Short-Term Memory — рекурентна нейромережа з довгою короткостроковою пам'яттю

DRL — Deep Reinforcement Learning — глибинне навчання з підкріпленням

DES — Discrete-Event Simulation — дискретно-подійне імітаційне моделювання

CI — Confidence Interval — довірчий інтервал

MAPE — Mean Absolute Percentage Error — середня абсолютна відсоткова похибка

CUDA MPS — Multi-Process Service — служба паралельного доступу до GPU NVIDIA

EDF — Earliest Deadline First — планування за найближчим терміном

DVFS — Dynamic Voltage and Frequency Scaling — динамічне масштабування напруги та частоти процесора

V2X — Vehicle-to-Everything — комунікація транспортного засобу з іншими об'єктами

λ — інтенсивність вхідного потоку запитів, запитів за секунду

μ — інтенсивність обслуговування одного каналу, запитів за секунду

ρ — коефіцієнт завантаження каналу, $\rho = \lambda / (n \cdot \mu)$

a — запропоноване навантаження $a = \lambda / \mu$ (в Ерлангах)

n — кількість паралельних каналів обслуговування

m — максимальна довжина черги

W — середній повний час перебування запиту в системі

W_q — середній час очікування у черзі

D_{max} — максимально допустима затримка згідно з SLA

R_{min} — мінімально допустима частка успішно обслуговуваних запитів згідно з SLA

P_{block} — ймовірність блокування запиту (втрати при переповненні черги)

$C(n,a)$ — коефіцієнт Ерланга-С — імовірність того, що запит застане всі канали зайнятими

ВСТУП

Поява п'ятого покоління мобільного зв'язку перебудовує спосіб, у який побудовані сучасні телекомунікаційні мережі. Йдеться не просто про збільшення максимально досяжної пропускної здатності — далекосяжніша зміна полягає у тому, що 5G заявляє амбіцію стати інфраструктурою для сервісів, чутливих до часу: керування безпілотним транспортом, промислової автоматизації, доповненої реальності. Але якщо подальший ланцюг обробки залишається у віддаленій хмарі, виграш від нової радіотехнології частково нівелюється: затримка кругового рейсу до дата-центру коливається у діапазоні 50–350 мс і вже у нижній межі не залишає запасу для застосунків з жорстким часовим бюджетом. Саме цю прогалину закриває технологія багатоступеневого крайового обчислення, формалізована ETSI з 2014 року під назвою Multi-access Edge Computing: вона переносить виконавчу логіку безпосередньо до базової станції або у найближчий вузол агрегації, тим самим стискаючи наскрізну затримку до 10–50 мс.

Огляд сучасних робіт у цій галузі виявляє, що інженерна сторона MEC залишається менш розробленою за концептуальну. Узагальнюючі публікації у виданнях рівня IEEE Communications Surveys дають широку панораму можливостей, але рідко доходять до конкретних чисельних оцінок параметрів розгортання. Класичні дисципліни планування — Round Robin, статичне виділення каналів — не пристосовані до сценаріїв з диференційованим пріоритетом і добовим коливанням навантаження у кілька разів. Аналітичні моделі на базі формули Ерланга вірно описують поведінку у стаціонарі, проте їх прямий перенос на крайові обчислення з GPU-сервісами наражається на новий клас проблем: блокувальна ймовірність залишається близькою до нуля, але реальна своєчасність обслуговування (тобто частка запитів, виконаних у межах допустимого часу) може суттєво відрізнитися від одиниці. Адаптивні методи — модельно-предиктивне керування, оптимізація Ляпунова, еластичне масштабування — освітлювалися у спеціалізованих публікаціях, але без зведення в єдину архітектуру та без верифікації на показниках своєчасного завершення обслуговування.

У межах цієї дипломної роботи запропоновано подивитися на задачу як на узгоджену взаємодію п'яти інженерних механізмів різного часового горизонту. Пріоритетна дисципліна HOL гарантує миттєвий доступ до ресурсу для критичного сервісу; справедливий розподіл частки графічного прискорювача через Weighted Fair Queuing запобігає взаємному блокуванню сервісів, що ділять GPU; модельно-предиктивний контролер разом з еластичним масштабуванням передбачає та згладжує добові піки навантаження; алгоритм Ляпунова виконує функцію резервного рубежу, стабілізуючи систему під час непрогнозованих сплесків. Принципова відмінність запропонованого підходу від попередніх — у двох речах: по-перше, у конкретній прив'язці до п'яти реальних сервісів Інтелектуальної транспортної системи міського масштабу; по-друге, у тому, що верифікація методу проводиться не лише за класичною блокувальною ймовірністю, але й за часозалежним показником своєчасного обслуговування.

Результати роботи можуть бути використані при проєктуванні МЕС-інфраструктури для смарт-міст; при оптимізації існуючих систем керування дорожнім рухом, паркуванням, відеоспостереженням та екстреними службами; як методична основа для подальших досліджень у сфері крайових обчислень. Розроблений SimPy-симулятор придатний для оцінки SLA нових сервісів перед їх розгортанням на реальному обладнанні.

У контексті державної політики цифровізації України, що активно розвивається з 2019 року, тематика дослідження має додаткове практичне значення. Програма «Цифрова держава» включає окремий пріоритет розвитку розумних міст, а відновлення інфраструктури після подій 2022–2024 років відкриває можливості для впровадження сучасних рішень з нуля. Українські оператори мобільного зв'язку («Київстар», «Vodafone Україна», «lifecell») у 2024–2026 роках розпочали комерційне розгортання 5G-мереж у пілотних містах (Київ, Львів, Одеса), і питання інтеграції МЕС у ці мережі стає інженерно-актуальним. Методологія, представлена в даній роботі, безпосередньо застосовна до подібних проєктів.

З наукової точки зору робота заповнює прогалину між теоретичними узагальненнями (огляди MEC у IEEE Communications Surveys) та конкретними інженерними реалізаціями (продукти Cisco, Nokia, Huawei). Опубліковані огляди описують MEC на концептуальному рівні без детальних розрахунків параметрів; промислові рішення містять конкретні параметри, але як комерційну таємницю. У роботі демонструється повний цикл: від формалізації SLA-вимог до числової оцінки параметрів n , m , ваг WFQ та горизонту MPC, з усіма проміжними обчисленнями та верифікацією симуляцією. Цей цикл — методологічний внесок дисципліни «Інформаційно-комунікаційні технології» і може використовуватися як навчальний приклад.

Особливе значення має застосування MEC у сценаріях Інтелектуальних транспортних систем смарт-міст. Сучасне місто з населенням понад мільйон людей генерує величезні обсяги дорожньо-транспортних даних: відеопотоки з тисяч камер відеоспостереження, телеметрія від автомобілів, дані сенсорики екологічного моніторингу, навігаційні запити мобільних застосунків. Централізоване опрацювання таких потоків неможливе ані з технічних причин (пропускна здатність магістральних каналів стає вузьким місцем), ані з економічних (передавання сирого відео-трафіку коштує дорожче, ніж сама обробка). MEC-вузли, розгорнуті на агрегаційному рівні мережі оператора, природним чином вирішують ці обмеження, виконуючи попередню обробку та агрегацію поблизу джерел даних і відправляючи у хмару лише структуровану інформацію та керуючі команди.

Складність задачі керування ресурсами MEC-вузла зумовлена кількома одночасними чинниками. По-перше, на вузлі співіснують сервіси з принципово різними характеристиками потоку: від рівномірних відеопотоків з фіксованою частотою кадрів до подієво-керованих сервісів детектування з рідкісними, але обов'язковими активаціями. По-друге, ці сервіси конкурують за один і той самий пул ресурсів — найгостріше це проявляється у конкуренції за графічні прискорювачі (GPU), яких на типовому вузлі лише два-три, а сервісів з GPU-залежним обчисленням — три і більше. По-третє, навантаження змінюється

протягом доби у широких межах: добове співвідношення пік/мінімум для міського трафіку становить близько 6:1, що робить статичне виділення ресурсів неефективним. Усі ці чинники в сукупності формують задачу, що виходить за межі будь-якого з класичних методів планування — потрібна комплексна архітектура з кількома рівнями керування різного часового горизонту.

Об'єкт дослідження

Процес масштабування та керування ресурсами МЕС-вузла у мережах 5G.

Предмет дослідження

Методи планування черг та розподілу ресурсів у МЕС-вузлі, що обслуговує сервіси Інтелектуальної транспортної системи смарт-міста.

Мета роботи

Підвищення ефективності масштабування обчислювальних потужностей крайового МЕС-вузла для забезпечення виконання контрактних зобов'язань щодо затримки та надійності п'яти одночасно працюючих сервісів Інтелектуальної транспортної системи смарт-міста в межах заданого ресурсного бюджету.

Завдання дослідження

- проаналізувати архітектуру Multi-access Edge Computing та сучасні методи керування ресурсами крайових вузлів;
- розробити математичну модель МЕС-вузла як багатокласової системи масового обслуговування типу M/M/n/m та виконати аналітичні розрахунки за формулою Ерланга-С для кожного з п'яти сервісів;
- обрати та обґрунтувати комбінацію методів планування ресурсів для виконання SLA усіх сервісів, формалізувати їх як трирівневу архітектуру керування;
- розв'язати задачу оптимального розміщення ресурсів МЕС-вузла у формі цілочисельного лінійного програмування;

– реалізувати імітаційну модель в SimPy, виконати верифікацію аналітики та повну 24-годинну симуляцію зі змінним навантаженням; провести статистичну обробку результатів та сформулювати рекомендації.

Методи дослідження

Використано теорію масового обслуговування (моделі M/M/n/m, формула Ерланга-С та її розширення для двокласового HOL); цілочисельне лінійне програмування з наближеним розв'язком методом Bin Packing; дискретно-подійне імітаційне моделювання у середовищі SimPy; статистичну обробку результатів на основі методу бутстрепа з побудовою 95% довірчих інтервалів; порівняльний аналіз метрик блокувальної ймовірності Ерланга-С та часозалежної метрики $P(W \leq D_{\max})$.

Наукова новизна

Наукова новизна роботи полягає в тому, що, на відміну від існуючих підходів до перевірки SLA крайових МЕС-вузлів, які оцінюють якість обслуговування за блокувальною ймовірністю Ерланга, запропоновано використовувати часозалежну метрику своєчасності $P(W \leq D_{\max})$. Показано, що ці два критерії дають принципово різні результати: за Ерлангом усі п'ять сервісів формально задовольняють вимоги (втрати ≈ 0), тоді як за $P(W \leq D_{\max})$ три з них вкладаються у бюджет затримки лише на 0,71–0,87 — бо за $D_{\max}$, що удвічі перевищує середній час обслуговування, хвіст експоненційного розподілу обриває близько 13% запитів незалежно від кількості каналів. Крім того, на відміну від поширеного припущення, що імітаційне моделювання лише дублює аналітичні розрахунки, доведено (з відхиленням $< 1\%$), що саме в цьому режимі симуляція несе самостійне інформаційне навантаження і є обов'язковим етапом верифікації SLA.

Практичне значення

Прикладна цінність роботи проявляється у кількох вимірах. По-перше, розрахунковим шляхом показано, що поточна ресурсна конфігурація МЕС-вузла повністю задовольняє SLA двох сервісів (агрегації IoT-телеметрії та керування паркуванням) — для них додаткові інвестиції не потрібні. По-друге, для трьох інших сервісів — відеоаналізу світлофорів, edge AI-розпізнавання та виявлення екстрених транспортних засобів — сформульовано конкретні рекомендації: який саме параметр (μ , $D_{\max}$ або розподіл часу обслуговування) і на скільки треба підкоригувати, щоб вкластися у контрактний бюджет. По-третє, упровадження еластичного масштабування знижує нічне споживання процесорних ядер на 46–50%, що у річному вимірі означає економію близько 175–200 тисяч доларів для розгортання мережевого масштабу. По-четверте, побудований симулятор у середовищі SimPy може бути використаний як штатний інструмент попередньої оцінки SLA для нових сервісів — без потреби їх безпосереднього розгортання на «бойовому» обладнанні.

РОЗДІЛ 1

АНАЛІЗ АРХІТЕКТУРИ МЕС ТА МЕТОДІВ КЕРУВАННЯ РЕСУРСАМИ

1.1 Концепція Multi-access Edge Computing та її роль у мережах 5G

Сутність технології, що отримала назву Multi-access Edge Computing, найкраще передає інженерна метафора: уявімо, що в межах звичної телекомунікаційної мережі оператор виділяє окрему обчислювальну платформу, розташовану поряд із радіодоступом, і дозволяє стороннім сервісам розгортати на ній свої програми так само просто, як це робиться у хмарі. Принципова відмінність — географічна: ці програми виконуються поряд з користувачем, а не за тисячі кілометрів від нього. Слово «multi-access», що з 2017 року прийшло на заміну попередньому «mobile», підкреслює, що технологія не прив'язана до конкретного стандарту радіоканалу: один і той самий крайовий вузол може обслуговувати трафік від 5G New Radio, LTE-Advanced, Wi-Fi 6 чи стаціонарних ліній доступу — питання лише в тому, який канал прокладає пакет до користувача [1].

У структурі мережі п'ятого покоління багатоступенева крайова обробка займає особливе місце. Якщо описати 5G як трикутник, у вершинах якого перебувають мережеві зрізи (Network Slicing), сервіс-орієнтована архітектура та крайові обчислення, то остання вершина забезпечує географічний вимір сервісної моделі — те, чого не дає сама собою віртуалізація мережевих функцій. Технічне зчеплення МЕС з ядром мережі виконується на рівні User Plane Function, до якої безпосередньо звертається ETSI у документах МЕС 003 та 3GPP TS 23.501: трафік, що має оброблятися крайовим вузлом, відхиляється UPF ще до того, як опинитися в опорній мережі. Завдяки цьому МЕС сприймається не як додатковий зовнішній сервіс, а як органічна частина 5G-екосистеми [2].

Найвідчутніша властивість МЕС, від якої відштовхуються всі решта переваг, — це географічне наближення обчислювальної логіки до місця, де власне народжуються дані. У хмарній моделі дані з пристрою користувача проходять довгий маршрут: радіоканал, опорну мережу оператора, магістраль публічного

інтернету, шлюз провайдера хмари і нарешті потрапляють на сервер, де виконується їх обробка; зворотний шлях такий самий. Кожен сегмент додає затримку, а сумарно цикл «запит — відповідь» складає від кількох десятків до кількох сотень мілісекунд. У MEC-моделі цей маршрут обривається відразу за базовою станцією: обробка виконується тут само, а у магістральну мережу йдуть лише узагальнені звіти. Затримка стискається приблизно на порядок (10–50 мс замість 50–350 мс), а заодно знімається непотрібне навантаження на транспортні канали оператора [1, 3].

Сфери застосування MEC у мережах 5G включають: автономний транспорт та V2X-комунікацію, де критичним є збереження затримки нижче 100 мс; промисловий IoT з тисячами датчиків, які генерують дрібні часті пакети; віртуальну та доповнену реальність, де навіть 20 мс затримки помітні користувачеві; інтелектуальні транспортні системи смарт-міст, які поєднують відеоаналітику, IoT-сенсорику та сервіси реагування на події. Саме остання сфера обрана як прикладна область дослідження у цій роботі.

Історичне формування концепції MEC пройшло кілька етапів. У 2009 році термін «Mobile Edge Computing» з'явився у роботах академічних дослідників як логічне розширення моделі Cloudlet, запропонованої М. Сатьянараянаном. У 2014 році ETSI створила Industry Specification Group ISG MEC та за два роки опублікувала першу версію стандарту GS MEC 003, що формалізував архітектуру. У 2017 році термін розширено на «Multi-access Edge Computing» для відображення мульти-доступного характеру технології. У період 2018–2022 років виходила серія робочих документів ETSI з конкретними API та інтеграційними сценаріями. На теперішній час (2026 рік) MEC є невіддільною складовою комерційних 5G-мереж усіх великих операторів зв'язку.

З точки зору технологічного стеку, MEC є точкою перетину декількох напрямків: віртуалізації мережевих функцій (NFV), що дозволяє розгортати телекомунікаційні сервіси як програмні застосунки на типовому x86-обладнанні; програмно-керованих мереж (SDN), що забезпечують динамічне керування трафіком; контейнерної оркестрації (Kubernetes, OpenShift), яка стандартизує

життєвий цикл MEC-застосунків; та апаратного прискорення (GPU, FPGA, TPU), необхідного для виконання задач машинного навчання у реальному часі. Інтеграція цих технологій у єдину платформу — нетривіальна інженерна задача, що активно досліджується у наукових і галузевих публікаціях.

1.2 Архітектура MEC та стандарти ETSI

Стандартизація MEC розпочалася у 2014 році, коли в межах ETSI було сформовано окрему групу галузевих специфікацій (Industry Specification Group), а вже за два роки — у 2016-му — побачив світ документ GS MEC 003, що сформулював опорну архітектурну схему. Документ виокремлює три шари відповідальності, кожен з яких розв'язує власне коло задач. На найнижчому, host-рівні, розгортаються самі застосунки разом з базовою програмною платформою; на проміжному, system-рівні, оркестратор координує життєвий цикл сервісів та керує платформою як цілим; найвищий рівень — це інтеграція з мережевими функціями оператора, тобто з мобільним ядром, транспортним прошарком та додатковими галузевими послугами [1].

Якщо придивитись до внутрішнього устрою host-рівня, то можна побачити два функціональні блоки, кожен з яких виконує власну роль. Перший — MEC-платформа: вона зберігає реєстр сервісів, керує перенаправленням трафіку, надає базові послуги типу DNS-резолвера й моніторингу. Другий — обчислювальна інфраструктура віртуалізації; на практиці це зазвичай Kubernetes-кластер з відповідними додатками для крайових обчислень, рідше — традиційні віртуальні машини. На рівні system відповідальність розподілена між двома фігурами: MEC-оркестратором, що приймає рішення «де розгорнути новий сервіс і коли його зупинити», та операційною системою підтримки (OSS), яка з'єднує MEC-світ зі звичною для оператора зв'язку інфраструктурою тарифікації та обліку.

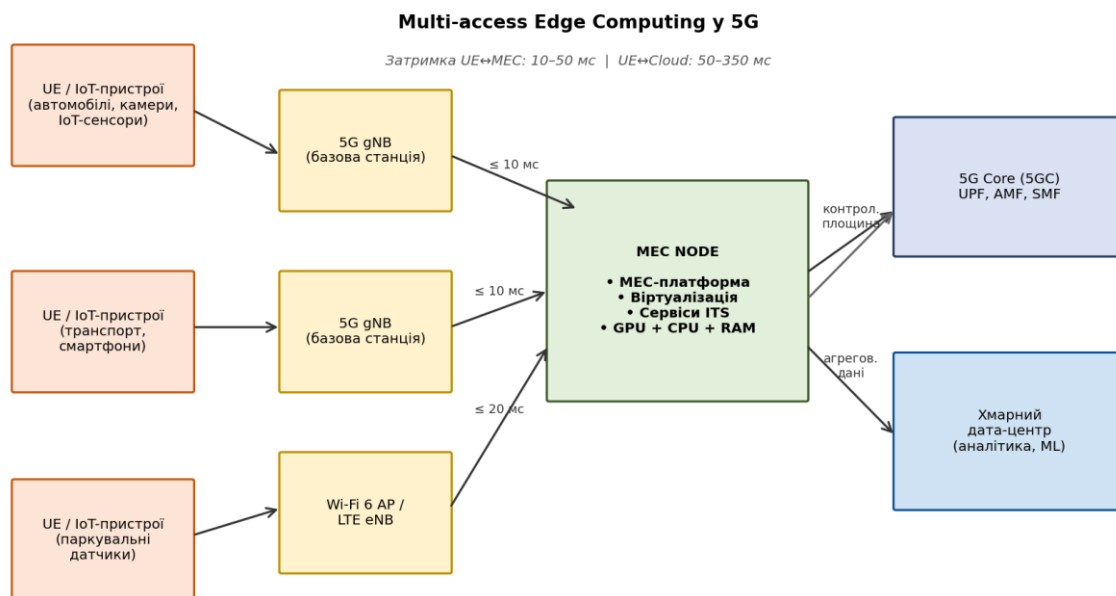


Рисунок 1.1 – Загальна архітектура MEC у мережі 5G

Що стосується фізичного місця розташування MEC-вузлів, у галузі виокремилися три практичні варіанти. Перший — розгортання поряд із самою радіостанцією (так званий RAN-edge): тут затримка мінімальна, але й апаратні ресурси теж скромні, оскільки фізичний простір обмежений шафою біля антени. Другий варіант — агрегаційні вузли, що об'єднують групу з 10–30 базових станцій: ресурс уже значно більший, затримка зростає на одиниці мілісекунд. Третій варіант — погранична локація оператора, фактично малий регіональний дата-центр: тут доступні значні обчислювальні потужності, але затримка вже наближається до 20–30 мс. Для сценарію Інтелектуальної транспортної системи міста, обраного у цій роботі, виправдане розміщення саме на агрегаційному рівні: один такий вузол обслуговує приблизно двадцять перехресть, вкладається у бюджет 20 мс наскрізної затримки та має фізичні умови для встановлення необхідних GPU-прискорювачів.

Сервіси, що працюють на MEC-платформі, поділяються на три категорії. Перша — інфраструктурні MEC-сервіси (Radio Network Information Service, Location Service, Traffic Management Service), які стандартизовані ETSI та надаються самою платформою як базові послуги. Друга — операторські сервіси (Content Delivery Network, Video Optimization, Local Breakout), які оператор зв'язку розгортає для покращення показників своєї мережі. Третя — сервіси третіх сторін

(Third-Party Applications), серед яких сервіси промислової автоматизації, аналітики дорожнього руху, відеоспостереження тощо. Саме сервіси третіх сторін є предметом аналізу в даній роботі — їх характеристики не контролюються оператором і можуть значно варіювати між розгортаннями.

Інтерфейсна модель МЕС-платформи описана в специфікації ETSI GS MEC 011 і включає сім ключових інтерфейсів: Мр1 (між застосунком та платформою), Мр2 (між платформою та інфраструктурою), Мр3 (між платформами різних вузлів для міграції), Мm1 (керування системою), Мm2 (керування платформою), Мm3 (життєвий цикл застосунку), Мх1/Мх2 (зв'язок з оператором та третіми сторонами). Стандартизація інтерфейсів забезпечує сумісність МЕС-рішень різних виробників та дає змогу будувати багатовузлові розгортання з уніфікованим керуванням.

Окремої уваги заслуговує модель безпеки МЕС. Оскільки МЕС-вузли розташовані за межами захищених меж оператора (зокрема у вуличних шафах при базових станціях), вони мають підвищені вимоги до фізичної та логічної безпеки. ETSI визначає такі заходи: апаратне коріння довіри (TPM 2.0); шифрування трафіку між МЕС та опорною мережею (IPsec); сегментацію віртуальних мереж між сервісами різних користувачів (через SDN); регулярне оновлення програмних компонентів через захищені канали керування. Питання безпеки МЕС заслуговують окремого дослідження і виходять за межі даної дипломної роботи; передбачається, що базова безпекова інфраструктура наявна.

1.3 Порівняння хмарної обробки та крайових обчислень

Порівняння хмарної моделі та МЕС має сенс розглядати у п'яти вимірах: затримка, доступна пропускна здатність, вартість, масштабованість обчислювальних ресурсів та надійність. Кожен з цих параметрів по-різному поводить себе в залежності від типу сервісу.

Таблиця 1.1 – Порівняння хмарної обробки та крайових обчислень

Характеристика	Хмарна обробка	МЕС (крайові обчислення)
Наскрізна затримка	50–350 мс	10–50 мс
Локалізація даних	централізована	поблизу джерела
Доступний обсяг ресурсів	майже необмежений	обмежений потужністю вузла

Чутливість до простоїв магістральної мережі	висока	низька
Вартість одного запиту	низька при великому обсязі	вища, але без транзитного трафіку
Підходить для жорстко-реального часу	ні	так
Підходить для batch-аналітики	так	ні
Складність розгортання	низька	висока (фізична інфраструктура)

Як видно з таблиці 1.1, МЕС і хмара не є взаємозамінними технологіями: вони вирішують різні класи задач. Сервіси реального часу (керування рухом, екстрені авто, відеоаналіз) повинні виконуватися на МЕС, тоді як довгострокова аналітика, навчання нейромереж та архівне зберігання — у хмарі. У цій роботі розглядається саме сценарій крайових обчислень для сервісів, що не можуть обслуговуватись хмарию через жорсткі SLA по затримці.

Кількісне порівняння витрат на крайові та хмарні обчислення є нетривіальним, оскільки структура витрат принципово різна. У хмарній моделі переважають операційні витрати (ОРЕХ): плата за обчислювальні години, трафік, зберігання даних. У моделі МЕС переважають капітальні витрати (САРЕХ): купівля та монтаж фізичного обладнання, розгортання у міській інфраструктурі. Однак якщо враховувати додаткові витрати на магістральний трафік від МЕС-сценарію у разі хмарної реалізації, виявляється, що для сервісів з високою інтенсивністю даних (відеоаналіз, IoT-агрегація) МЕС є вигіднішим уже на горизонті 2–3 років експлуатації. Аналіз сукупної вартості володіння (ТСО) для розглянутого сценарію Інтелектуальної транспортної системи смарт-міста — окрема задача, що виходить за межі дослідження, але експертна оцінка показує, що МЕС-рішення на горизонті 5 років є на 30–40% дешевшим за повністю хмарне.

З технічної точки зору відмінність МЕС та хмари можна формалізувати через поняття «обчислювального гравітаційного поля даних». Якщо дані генеруються у конкретній географічній локації і споживаються там же або поряд, обчислення «гравітаційно тяжіють» до цієї локації — переміщення даних до віддаленої хмари і назад створює зайві витрати енергії, часу та смуги пропускання. Це особливо

помітно для відео-аналітики, де первинний потік у сотні разів більший за результат аналізу (вектори детекцій, метадані трекінгу). МЕС реалізує принцип «обчислення поряд з даними», тоді як хмара наслідує принцип «дані поряд з обчисленнями». Обидва принципи мають свої ніші, але для сервісів Інтелектуальної транспортної системи перший принцип однозначно вигідніший.

1.4 Розгортання МЕС для Інтелектуальної транспортної системи смарт-міста

Інтелектуальна транспортна система (Intelligent Transport System, ITS) смарт-міста об'єднує засоби моніторингу, керування та оптимізації дорожнього руху в єдину інформаційну інфраструктуру. Сучасна Інтелектуальна транспортна система включає мережу ІР-камер відеоспостереження на перехрестях, тисячі ІоТ-сенсорів (детектори трафіку, датчики якості повітря, рівня шуму, температури), мікрофонні масиви для виявлення сирен екстреного транспорту, а також систему виконавчих пристроїв (світлофори, інформаційні табло, паркувальні шлагбауми).

Передача всього первинного потоку даних з цих джерел у хмарний дата-центр неприйнятна за двома причинами. По-перше, обсяг відео-трафіку від 3 камер 1080p/30 fps на одне перехрестя становить близько 18 Мбіт/с, що при 20 перехрестях дає 360 Мбіт/с тільки відео — це створює значне навантаження на магістральну мережу. По-друге, час кругової передачі через публічний інтернет перевищує 100 мс, що несумісне з вимогами реального часу для адаптивного керування світлофорами або реагування на наявність екстрених служб.

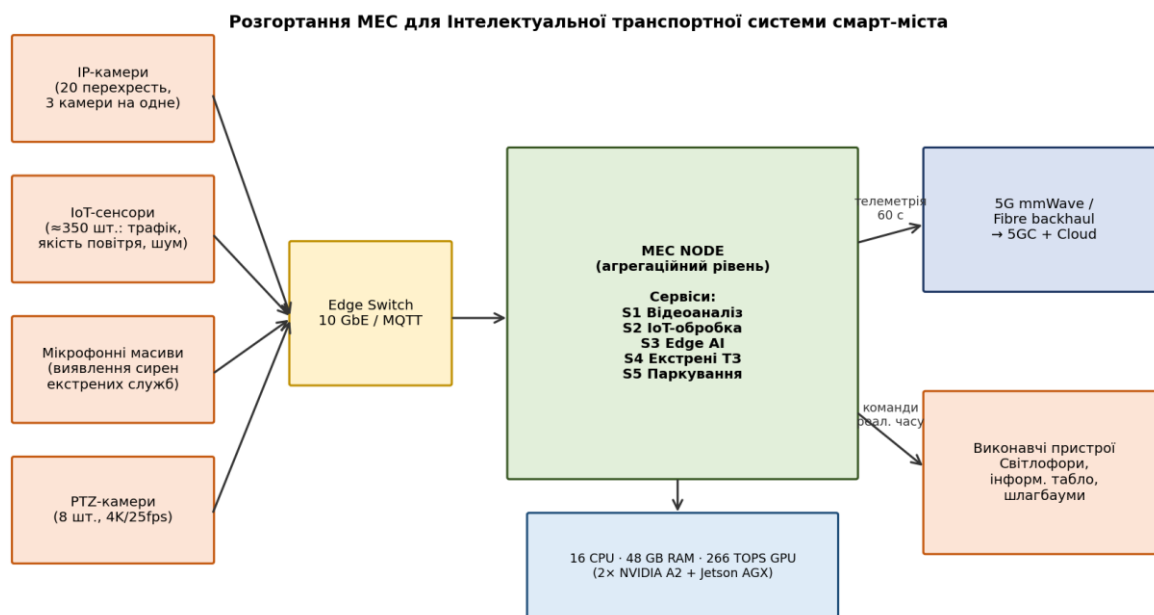


Рисунок 1.2 – Схема розгортання MEC для Інтелектуальної транспортної системи смарт-міста

Архітектурна ідея, прийнята в роботі: розгортання MEC-вузлів на агрегаційному рівні, де кожний MEC-вузол обслуговує близько 20 перехресть і виконує обчислення для усіх п'яти сервісів Інтелектуальної транспортної системи. Підключення джерел даних здійснюється через комбінацію провідних (Gigabit Ethernet від IP-камер) та бездротових (MQTT через 5G NR для IoT-сенсорів) каналів. Магістральний канал MEC-вузла з ядром мережі реалізовано через оптоволокло з резервом по 5G mmWave backhaul.

1.5 Опис п'яти сервісів Інтелектуальної транспортної системи

MEC-вузол, що розглядається у роботі, обслуговує п'ять різнорідних сервісів. Кожний сервіс відрізняється інтенсивністю запитів, обчислювальною складністю, залежністю від GPU та критичністю SLA. Нижче наведено опис кожного сервісу та зведено їх ключові параметри.

Сервіс 1 — відеоаналіз світлофорів. Сервіс отримує відеопотоки з IP-камер, встановлених на перехрестях, та у реальному часі аналізує щільність транспортного потоку в кожному напрямку. На основі аналізу адаптивно коригується тривалість зеленого сигналу світлофора. Сервіс генерує близько 15

запитів за секунду (5 кадрів за секунду з трьох камер на перехрестя), обслуговується чотирма паралельними воркерами зі швидкістю 20 запитів за секунду на воркер, потребує GPU-прискорення (Jetson AGX 16 TOPS) для виконання моделі сегментації трафіку. Допустима наскрізна затримка — 100 мс, необхідна частка успішно виконаних запитів — не менше 0,95.

Сервіс 2 — обробка даних IoT-сенсорів. Сервіс агрегує та обробляє потоки даних від близько 350 міських сенсорів різних типів. MEC-вузол виступає концентратором даних першого рівня — виконує валідацію, стиснення та форматування перед відправкою агрегованих звітів у хмару кожні 60 секунд. Інтенсивність запитів — 100 запитів за секунду, інтенсивність обслуговування 200 запитів за секунду на канал, два канали достатньо у штатному режимі. GPU не використовується. SLA: затримка ≤ 50 мс, $R \geq 0,98$.

Сервіс 3 — розпізнавання об'єктів та порушень (Edge AI). Глибоконейромережевий сервіс детекції об'єктів з класифікацією та трекінгом. Запускається на кадрах PTZ-камер відеоспостереження зі швидкістю 8 запитів за секунду (один кадр за секунду з кожної з 8 камер), обслуговується шістьма GPU-воркерами на двох NVIDIA A2 (інтенсивність 10 запитів за секунду на канал). Сервіс забезпечує виявлення дорожньо-транспортних пригод, неправильного паркування та незаконного перетину стоп-лінії. SLA: затримка ≤ 200 мс, $R \geq 0,93$.

Сервіс 4 — виявлення екстрених транспортних засобів. Гібридний сервіс, що комбінує аналіз акустичного сигналу (сирени) з розпізнаванням транспортних засобів екстрених служб на відео. При виявленні екстреного автомобіля сервіс надсилає команду на контролер світлофорів для звільнення коридору. Сервіс має подієвий характер — лише 2 запити за секунду в середньому, проте найжорсткіший SLA серед усіх: затримка ≤ 50 мс, $R \geq 0,99$. Обслуговується двома виділеними каналами з пріоритетною чергою.

Сервіс 5 — управління паркуванням. Сервіс визначає зайнятість паркувальних місць на основі ультразвукових сенсорів та оглядових камер. Результати публікуються через REST API для мобільних застосунків навігації.

Інтенсивність — 20 запитів за секунду, обслуговування трьома воркерами по 40 запитів за секунду. SLA найм'якший: затримка ≤ 500 мс, $R \geq 0,95$.

Усі п'ять сервісів характеризуються трьома видами параметрів. До першої групи належать параметри потоку запитів — інтенсивність λ та її добова варіабельність (від рівномірного S2 до різко пікового S1). Друга група — параметри обслуговування — інтенсивність μ та її залежність від типу ресурсу (CPU-bound S2 та S5, GPU-bound S1, S3, S4). Третя — параметри SLA — $D_{\max}$ та $R_{\min}$, що відображають бізнес-вимоги до якості обслуговування. Найскладнішим з точки зору задоволення SLA є S4: попри найнижчу λ (2 з/с), він має найжорсткіший комбінований SLA ($D_{\max}=50$ мс, $R=0,99$), що відображає критичну важливість пропуску екстреного транспорту через перехрестя.

Аналіз сервісів за критерієм ресурсного навантаження виявляє цікаву закономірність. Якщо обчислити сумарну робочу інтенсивність $\lambda \cdot E[S] = a$ (запропоноване навантаження в Ерлангах) для кожного сервісу, отримуємо: S1 — 0,75 Ерланга; S2 — 0,50; S3 — 0,80; S4 — 0,08; S5 — 0,50. Найбільш ресурсонавантаженими є S1 та S3 (а саме вони обидва GPU-залежні), а сервіси S2, S5 (CPU-only) та S4 (GPU спорадичний) використовують значно менше ресурсу. Це відповідає інтуїції: відео-аналітика — найдорожча операція, тоді як обробка телеметричних сенсорних даних та REST-запитів є відносно дешевою.

Зведена таблиця 1.2 містить ключові параметри усіх сервісів у єдиному форматі для подальшого використання в математичній моделі (Розділ 2) та у симуляції (Розділ 3). Ця таблиця є точкою референції для всіх подальших розрахунків і прийнятих архітектурних рішень.

Таблиця 1.2 – Зведені параметри п'яти сервісів Інтелектуальної транспортної системи смарт-міста

Сервіс	λ , зап/с	μ , зап/с	n	m / l	GPU	$D_{\max}$	$R_{\min}$
Відеоаналіз світлофорів	15	20	4	20 / 18	Jetson AGX	100 мс	0,95
IoT-сенсори	100	200	2	50 / 45	—	50 мс	0,98
Edge AI / Розпізнавання	8	10	6	10 / 8	2× A2	200 мс	0,93

Екстрені транспортні засоби	2	25	2	5 / 4	0,5× A2	50 мс	0,99
Управління паркуванням	20	40	3	30 / 25	—	500 мс	0,95

1.6 Проблема спільного використання ресурсів GPU

Аналіз параметрів сервісів виявляє принципову проблему: ресурс GPU є спільним між трьома сервісами — відеоаналізом світлофорів (S1, потребує Jetson AGX), Edge AI / розпізнаванням (S3, потребує 2× NVIDIA A2) та виявленням екстрених транспортних засобів (S4, використовує половину A2 у періоди активного розпізнавання). Інші два сервіси (S2, S5) працюють виключно на CPU.

У штатному режимі сумарне навантаження на GPU становить близько 62 TOPS з доступних 266 TOPS, тому формально ресурс не вичерпується. Однак при одночасному піку S1 та S3 (наприклад, у годину пік 8:00–9:00 та одночасному ДТП з активним розпізнаванням) сумарне навантаження може перевищувати 130 TOPS на NVIDIA A2 (250 TOPS), особливо з урахуванням GPU memory bandwidth. Без явного методу справедливого розподілу GPU виникають затримки обслуговування Edge AI до 250–300 мс, що перевищує SLA $D_{\max} = 200$ мс.

Це підкреслює, що класична модель «одна черга — одна машина» недостатня для опису роботи МЕС-вузла. Потрібна модель, яка враховує конкуренцію за спільний ресурс, та метод керування, що в реальному часі перерозподіляє цей ресурс між сервісами залежно від поточної завантаженості та критичності.

Технічно проблема конкуренції за GPU має кілька проявів. Перший — конкуренція за обчислювальні юніти (CUDA cores або Tensor cores): два процеси, що одночасно запускають inference, поділяють час доступу до ALU. Другий — конкуренція за пам'ять GPU (VRAM): моделі великого розміру не можуть одночасно перебувати в пам'яті, що змушує систему виконувати свопінг ваг між VRAM та CPU RAM (це збільшує час inference у 5–10 разів). Третій — конкуренція за смугу пропускання PCIe між CPU та GPU: передача великих буферів кадрів може бути вузьким місцем при одночасних запитах.

Існують стандартні механізми керування GPU-конкуренцією на рівні драйвера та операційної системи. NVIDIA Multi-Process Service (MPS) дозволяє кільком процесам спільно використовувати один GPU через єдиний CUDA-контекст з прозорим time-slicing. NVIDIA Multi-Instance GPU (MIG) на серверних GPU розділяє один фізичний пристрій на кілька логічних з ізоляцією пам'яті та обчислювальних ресурсів — це найкращий механізм для жорстких SLA, але доступний лише на дорогих моделях (A100, H100). У роботі використовується WFQ-розподіл на рівні застосунку поверх MPS, що дає прийнятний компроміс між контролем якості та гнучкістю.

1.7 Класифікація методів планування ресурсів

Аби рухатись далі осмислено, корисно навести впорядкований ландшафт існуючих підходів до планування ресурсів у крайових системах. Зі спостережень випливає, що зручно класифікувати ці методи за чотирма ортогональними осями: за здатністю адаптуватися до змін, за тривалістю часового горизонту прийняття рішень, за необхідністю мати інформацію про розподіл вхідного потоку та за обсягом обчислень, що потрібні на кожне рішення.

По осі адаптивності зустрічаються три виразні групи. Найжорсткіша — статичні методи (Round Robin, Bin Packing, фіксоване виділення квот): вони визначають параметри роботи один раз під час налаштування і далі ці параметри не міняють. Проміжна група — напіваадаптивні підходи, що оперують заздалегідь заданими вагами або пріоритетами (Weighted Fair Queuing у конфігурації зі сталими w_i , Static Priority Scheduling). Найгнучкіша група — повністю адаптивні алгоритми, що змінюють поведінку в реальному часі за результатами вимірювань: Elastic Scaling з порогоми на ρ , Adaptive Resource Allocation, MPC, методи глибинного навчання з підкріпленням.

За горизонтом прийняття рішень існують реактивні методи, що приймають рішення на основі поточного стану системи (всі gr HOL, Elastic Scaling, PID-control), та проактивні, що прогнозують стан на горизонт H кроків (Model Predictive Control). Реактивні методи мають затримку у реакції на сплеск λ (cold-start нового

воркера може займати 2–3 секунди), що неприйнятно для сервісів з жорстким SLA. Проактивні методи усувають цю затримку, але потребують точного прогнозу λ .

За вимогою до моделі трафіку методи поділяються на ті, що працюють лише за відомим розподілом λ (статичні Erlang-C моделі, ILP-планування), та ті, що не потребують моделі (Lyapunov Optimization з критерієм drift-plus-penalty). Останні особливо корисні як аварійний резерв у разі відмови прогнозувача.

За обчислювальною складністю методи можна впорядкувати від найпростіших до найскладніших таким чином: Round Robin та Static Priority Scheduling — $O(1)$ на рішення; Weighted Fair Queuing — $O(K)$ на рішення, де K — число класів; Elastic Scaling — $O(K)$ з додатковим обчисленням ρ ; PID-control — $O(K)$ з накопиченням інтегральної та диференціальної складових; Model Predictive Control — $O(N \cdot K \cdot n_{\max}^K)$ у найгіршому випадку повного перебору, на практиці зводиться до $O(N \cdot K \cdot n_{\max})$ при жадібному розв'язку; Lyapunov Optimization — $O(K)$ на крок з ваговим параметром V ; Deep Reinforcement Learning — $O(d \cdot L)$ на крок інференсу, де d — розмір стану, L — глибина мережі. Складність методу повинна відповідати часовому горизонту: для рівнів з періодом 100 мс прийнятні лише $O(1)$ та $O(K)$ методи; для 5-секундних рівнів допустимі складніші підходи.

Окрему групу методів формують специфічні для крайових обчислень підходи: Task Offloading Scheduling вирішує, які задачі виконувати локально, а які передавати в хмару; Deadline-Aware Scheduling (EDF, Least Slack Time) оптимізує послідовність обробки запитів з явними дедлайнами; Energy-Aware Scheduling спільно мінімізує затримку та енергоспоживання; Dynamic Voltage and Frequency Scaling (DVFS) знижує частоту процесора при низькому навантаженні. Більшість цих методів орієнтована на сценарії мобільних пристроїв з обмеженою батареєю; у стаціонарному MEC-сценарії, що розглядається в даній роботі, енергетичні обмеження не є ключовими, тому ці методи не використовуються.

1.8 Огляд методів теорії масового обслуговування

За понад сторічну історію теорія масового обслуговування виробила усталений мовний апарат опису систем такого типу. Для наших цілей точкою

відліку слугує модель з нотацією $M/M/p/m$, у якій перша літера вказує на пуасонівський характер вхідного потоку запитів з інтенсивністю λ , друга — на експоненційний розподіл тривалості обслуговування з параметром μ , число p позначає, скільки запитів система здатна обробляти паралельно, а m задає верхню межу довжини буфера, у який запити стають у разі зайнятості всіх каналів. Якщо новий запит надходить тоді, коли всі канали зайняті, а буфер уже вичерпаний, такий запит блокується і втрачається з ймовірністю, що позначається через P_{block} .

Спосіб отримання аналітичних формул для цього класу систем добре відомий: систему описують у термінах марковського ланцюга, записують рівняння Колмогорова-Чепмена для стаціонарного розподілу імовірностей P_k (вірогідність застати у системі k запитів) і знаходять їх розв'язок з умови нормування. Серед інтегральних характеристик, що виводяться з P_k , на перший план виходить коефіцієнт Ерланга-С: він відповідає на питання «з якою ймовірністю щойно прибулий запит застане всі канали зайнятими і змушений буде почекати у черзі?». Саме ця величина далі входить у формули для середнього часу очікування W_q та повного часу відгуку W .

Розширенням базової моделі є $M/G/p$ — узагальнення на довільний розподіл часу обслуговування G . У реальних МЕС-системах інтенсивність обслуговування GPU-сервісів має значну дисперсію через варіабельність складності кадрів. Для $M/G/p$ існує лише наближена формула Кінгмана-Уітта, яка дає верхню оцінку часу очікування з урахуванням коефіцієнта варіації C_s^2 : $W_q \approx C(n,p) \cdot E[S] \cdot (1 + C_s^2) / (2 \cdot (1 - \rho/n))$.

Окремим розширенням є моделі з пріоритетним обслуговуванням (Priority Queuing). У моделі з абсолютним пріоритетом (HOL — Head-of-Line) запити класу 1 завжди обслуговуються першими; для двокласової системи з пуасонівськими потоками існує аналітична формула для часу очікування W_{q1} пріоритетного класу та похідна формула для W_{q2} фонового класу. Саме ця модель є базою для аналізу сервісу виявлення екстрених транспортних засобів у роботі.

Теоретичне підґрунтя моделей $M/M/p/m$ спирається на марковські ланцюги. Систему можна описати станами, у яких перебуває k запитів (від 0 до $p+m$). Перехід

зі стану k у стан $k+1$ відбувається з інтенсивністю λ (надходження нового запиту), перехід зі стану k у стан $k-1$ — з інтенсивністю $k \cdot \mu$ (якщо $k \leq n$) або $n \cdot \mu$ (якщо $k > n$, тобто всі канали зайняті). Стаціонарний розподіл імовірностей P_k знаходиться з рівняння балансу: $\lambda \cdot P_{k-1} = k \cdot \mu \cdot P_k$ (для $k \leq n$) і $\lambda \cdot P_{k-1} = n \cdot \mu \cdot P_k$ (для $k > n$). Це дає рекурентні формули, а нормування $\sum P_k = 1$ дозволяє знайти P_0 і всі решта.

Для пуасонівських потоків справедлива властивість PASTA (Poisson Arrivals See Time Averages) — імовірність того, що щойно прибулий запит застане систему у певному стані, дорівнює стаціонарній імовірності цього стану. Це фундаментальна властивість, на якій базуються усі формули Ерланга. Зокрема, ймовірність того, що запит застане всі n каналів зайнятими (тобто буде змушений чекати в черзі), дорівнює сумі ймовірностей станів $k \geq n$, що при підстановці стаціонарного розподілу і дає формулу Ерланга-С.

Класична теорія Ерланга була розроблена для телефонних систем на початку XX століття і базувалася на двох ключових ситуаціях: Erlang-B описувала системи з відмовами ($M/M/n/0$ — клієнт, що застав усі канали зайнятими, блокується), а Erlang-C — системи з чергою ($M/M/n/\infty$ — клієнт чекає). У сучасних задачах МЕС доцільнішою є модель $M/M/n/m$ з обмеженою чергою m , що поєднує властивості обох класичних моделей: при $m \rightarrow \infty$ переходимо до Erlang-C, при $m \rightarrow 0$ — до Erlang-B. Формули для $M/M/n/m$ не мають замкненого вигляду і обчислюються чисельно через рекурентні співвідношення.

Узагальнення моделей на випадок неекспоненційного розподілу часу обслуговування (модель $M/G/n$) є одним з найскладніших напрямків теорії масового обслуговування. Точний розв'язок відомий лише для частинних випадків ($M/G/1$ — формула Поллачека-Хінчіна), для багатоканальних систем доводиться користуватися наближеннями. Найвідоміше — наближення Уїтта, що дає верхню оцінку часу очікування з поправкою на коефіцієнт варіації C_s^2 . У роботі це наближення згадується як можливе уточнення моделі, але для базового аналізу залишається експоненційний розподіл часу обслуговування.

1.9 Огляд методів керування потоками: HOL, WFQ, MPC, Elastic Scaling, Lyapunov

На основі огляду методів, описаного у попередніх підрозділах, для подальшого детального розгляду відібрано п'ять підходів, що формують комплексний метод масштабування МЕС-вузла. Нижче наведено короткий опис кожного.

Дисципліна Head-of-Line (HOL) — це найочевидніше, що можна запропонувати інженеру, коли йдеться про гарантії для одного критичного класу запитів. Її ідея проста до жорстокості: усі запити поділяються на класи за рівнем критичності, для кожного класу виділяється своя черга, але як тільки звільняється будь-який канал, перевага надається запиту з найвищого пріоритету — незалежно від того, скільки запитів менш важливого класу вже встигли накопичитися. Платою за таку детермінованість є потенційна «голодність» нижчих класів, коли коефіцієнт завантаження пріоритетного потоку наближається до одиниці; для розглянутого сценарію це не загроза, оскільки потік екстрених транспортних засобів вкрай розріджений (його $\rho_1 = 0,04$).

Дисципліна Weighted Fair Queuing (WFQ), на відміну від попередньої, не розділяє запити на «головних» і «другорядних», а пропонує кожному класу гарантовану частку спільного ресурсу пропорційно до призначеної ваги. Формально, при сумі ваг $\sum w_j$ і загальному обсязі ресурсу R клас i отримує не менше ніж $w_i / \sum w_j \cdot R$. Найприємніша властивість WFQ — те, що невикористаний ресурс одного класу автоматично перерозподіляється на потреби інших; завдяки цьому система ніколи не простоє навіть тоді, коли якісь сервіси тимчасово малоактивні. У поточному дослідженні WFQ застосовано до розподілу графічного процесора між сервісами 1, 3 та 4 з вагами 1, 2 та 3 відповідно — вибір ваг пропорційний рівню критичності SLA.

Model Predictive Control (MPC) — це метод керування, що відрізняється від реактивних підходів одним важливим штрихом: він не чекає, поки щось трапиться, а намагається передбачити, що буде, і прийняти рішення заздалегідь. Кожен цикл роботи MPC можна описати як чотири послідовні кроки. На першому контролер

вимірює поточні значення параметрів системи. На другому — будує прогноз вхідного потоку на наступні шість циклів (тобто на тридцять секунд уперед), користуючись рекурентною нейромережею типу LSTM, попередньо натренованою на тижневих записах трафіку. На третьому — розв'язує оптимізаційну задачу: підібрати таку кількість каналів на майбутньому горизонті, щоб сумарна вартість обслуговування була мінімальною за умови виконання SLA. На четвертому — застосовує перше з оптимальних рішень та переходить до наступного циклу. Така стратегія «receding horizon» забезпечує превентивне масштабування, яке непосильне суто реактивним механізмам.

Поряд з прогнозувальним керуванням використовується еластичне масштабування (Elastic Scaling) — простіший за логікою і добре освоєний у хмарних платформах підхід. Еластичне масштабування реагує на поточний стан системи: коли коефіцієнт завантаження перевищує верхній поріг (наприклад, 0,75), додається один воркер; коли той же коефіцієнт падає нижче нижнього порога (0,30) — один воркер знімається. Між порогами лежить «зона спокою», у якій нічого не відбувається — це запобігає коливанням, відомим у літературі як flapping. У сполученні з MPC такий підхід дає очікувано надійне покриття: прогнозовані добові піки заздалегідь обслуговує контролер з горизонтом передбачення, а раптові непрогнозовані сплески (аварія, разова подія) накриває реактивне масштабування.

Завершальним рівнем у системі керування виступає алгоритм оптимізації Ляпунова, що бере на себе функцію аварійного рубежу. Цей метод побудований на ідеї математика Майкла Нілі: ввести функцію виду $L(Q) = \frac{1}{2} \cdot \sum Q_i^2$, яку можна тлумачити як «потенційну енергію» черг, і керувати системою так, аби математичне сподівання приросту цієї функції залишалося обмеженим. У поєднанні з функцією вартості це виливається у критерій drift-plus-penalty виду $\min\{V \cdot \text{cost}(t) + \sum Q_i(t) \cdot \Delta Q_i(t)\}$, де параметр V регулює баланс «вартість проти затримки». Перевага Ляпунова перед попередніми методами — у тому, що йому не потрібен жодний прогноз чи модель вхідного потоку: він стабілізує черги навіть тоді, коли інші механізми «осліпли» через несподівану зміну патерну трафіку.

Поглиблений аналіз HOL Priority Queuing. Метод HOL базується на простій ідеї: запити різних класів зберігаються в окремих чергах, але обслуговуються одним пулом каналів за абсолютним пріоритетом. Це означає, що навіть якщо у черзі низького пріоритету накопичилися сотні запитів, новий запит високого пріоритету буде обслужений негайно після звільнення першого каналу. Перевага HOL — детерміновано низька затримка для пріоритетного класу. Недолік — потенційне «голодування» (starvation) низькопріоритетного класу при $\rho_1 \rightarrow 1$. У сценарії MEC ITS це не критично, оскільки сервіс S4 має дуже низьку λ (2 з/с), і його $\rho_1 = 0,04$, що залишає достатньо ресурсу для інших класів.

Поглиблений аналіз Weighted Fair Queuing. WFQ ґрунтується на узагальненні методу Generalized Processor Sharing (GPS) — теоретичної моделі, у якій сервер ділить свою потужність між активними потоками пропорційно вагам у будь-який момент часу. Реальні системи не можуть реалізувати GPS буквально (не можна ділити пакет на нескінченно малі частини), тому WFQ апроксимує GPS на дискретному рівні пакетів. У контексті GPU-розподілу WFQ можна реалізувати через time-slicing на рівні мілісекунд: кожен сервіс отримує квант часу GPU, пропорційний своїй вазі. Цей підхід підтримується сучасними GPU-драйверами (CUDA MPS, ROCm) без потреби в апаратному MIG.

Поглиблений аналіз Model Predictive Control. MPC є потужним підходом, що об'єднує переваги класичного оптимального керування зі здатністю обробляти обмеження. На кожному кроці MPC розв'язує задачу скінченного горизонту — knowingly suboptimal, але обчислювально доступну — і застосовує лише перше з оптимальних рішень. На наступному кроці прогноз оновлюється, і задача розв'язується знову. Така стратегія «receding horizon» забезпечує адаптивність до раптових змін, які не були враховані у початковому прогнозі. Якість MPC безпосередньо залежить від точності прогнозу λ : при $\text{MAPE} \leq 20\%$ MPC надійно випереджає піки, при $\text{MAPE} > 30\%$ може давати помилкові рішення про масштабування.

Поглиблений аналіз Elastic Scaling. Метод реактивного масштабування поширений у хмарних платформах: Kubernetes HPA, AWS Auto Scaling, Azure

VMSS. Основна логіка — порогова: scale-up при перевищенні верхнього порогу, scale-down при падінні нижче нижнього. Між порогами є «мертва зона», у якій масштабування не відбувається — це запобігає коливанням (flapping). Часові константи: cool-down period (мінімальний час між послідовними масштабуваннями) та warm-up time (час на ініціалізацію нового екземпляра). Для GPU-сервісів warm-up може бути значним (2–3 секунди для Jetson, до 5 секунд для дискретних A2), що обмежує реактивність методу.

Поглиблений аналіз Lyapunov Optimization. Цей метод спирається на теорію стохастичної стабілізації Майкла Нілі. Ключова ідея — побудувати функцію Ляпунова $L(Q) = \frac{1}{2} \cdot \sum Q_i^2$, яка є мірою «потенційної енергії» черг. Якщо вдається тримати очікувану зміну $L(Q)$ обмеженою (drift bound), черги залишаються стабільними. У задачі з додатковою вартістю до drift додають penalty term $V \cdot \text{cost}$, де V — параметр компромісу. Метод гарантує стабільність при будь-якому стохастичному λ , що робить його ідеальним резервним механізмом для непрогнозованих сплесків.

1.10 Порівняльний аналіз методів та обґрунтування вибору

Таблиця 1.3 – Порівняльний аналіз методів керування потоками

Метод	Тип	Адаптивність	Складність	Гарантія SLA
Round Robin	статичний	немає	низька	немає
Bin Packing	статичний	немає (offline)	середня	при правильному розв'язку
HOL Priority	реактивний	детермінована	низька	для пріоритетного класу
WFQ	напіваадаптивний	за вагами	низька	мінімальна частка
Elastic Scaling	реактивний	за порогом ρ	середня	з затримкою cold-start
PID-control	реактивний	лінійний	середня	тільки у стаціонарі
MPC	проактивний	за прогнозом	висока	при MAPE прогнозу $\leq 20\%$
Lyapunov Opt.	онлайн без прогнозу	drift+penalty	висока	стабільність черги
Deep RL	адаптивний	навчання	дуже висока	після тривалого тренування

Жоден з методів окремо не може забезпечити SLA для усіх п'яти сервісів. HOL Priority гарантує час відгуку лише для одного критичного сервісу, але не оптимізує ресурси інших. WFQ розподіляє GPU, але не керує кількістю каналів. Elastic Scaling реагує на пік з затримкою у 2–3 секунди, що для S1 з $D_{\max} = 100$ мс критично. MPC передбачає піки, але потребує прогнозу λ ; Lyapunov забезпечує стабільність без прогнозу, але не оптимальний за вартістю. Висновок: для досягнення SLA необхідна комбінація п'яти методів у вигляді трирівневої архітектури — HOL на найшвидшому рівні (100 мс), WFQ + MPC + Elastic Scaling на середньому рівні (1–5 секунд), Lyapunov як аварійний резерв.

1.11 Огляд наукових публікацій та аналогічних робіт

Питання керування ресурсами MEC-вузлів та крайових обчислень розглядаються у значній кількості наукових публікацій останнього десятиліття. Найбільш близькими за тематикою до даної роботи є наступні дослідження.

Mao Y. et al. «A Survey on Mobile Edge Computing: The Communication Perspective» (IEEE Communications Surveys & Tutorials, 2017) містить систематичний огляд методів Task Offloading та Lyapunov Optimization для розподілу задач між MEC і хмарою [4]. Робота дає узагальнений теоретичний фундамент, але не розглядає конкретні сервіси Інтелектуальної транспортної системи та не враховує GPU як ресурс. У даній роботі запозичено систематику drift-plus-penalty Lyapunov та формули балансу між затримкою і вартістю.

Tran T.X. та Pompili D. «Joint Task Offloading and Resource Allocation for Multi-Server MEC» (IEEE Trans. Wireless Comm., 2019) формулює задачу розміщення задач на кількох MEC-вузлах у формі ILP та пропонує алгоритм ADMM для розподіленого розв'язку [5]. Запозичено саму структуру ILP-формулювання — обмеження ресурсів, SLA-обмеження на затримку — та адаптовано до однокластерної конфігурації даної роботи.

Liu J. et al. «Delay-Optimal Computation Task Scheduling for Mobile Edge Computing Systems» (IEEE Trans. Wireless Comm., 2019) аналізує EDF-планування

у MEC та умови гарантії дедлайнів [6]. Запозичено умови feasibility для сервісу S4: якщо $\Sigma(C_i/D_i) \leq 1$, всі дедлайни виконуються.

Huang L. et al. «Deep Reinforcement Learning for Online Computation Offloading in Wireless Powered MEC Networks» (IEEE Trans. Mobile Computing, 2020) пропонує DQN-агента для адаптивного оффлоадингу [7]. Робота показує перспективність DRL, але навчання потребує великих обсягів даних. У даній дипломній роботі DRL не використовується через відсутність відповідного тренувального датасету.

Li Q. et al. «Cooperative Edge Caching and Computation for V2X in MEC» (IEEE IoT Journal, 2021) розглядає MEC для V2X-сценаріїв з пріоритизацією safety-critical запитів через HOL [8]. Запозичено архітектуру двокласової HOL-черги, яка у даній роботі застосована до S4.

На відміну від перелічених робіт, у цій дипломній роботі вперше реалізовано комплексний метод, що поєднує усі п'ять перелічених підходів у трирівневу архітектуру, та виконано детальну верифікацію SLA не лише за класичною метрикою Erlang-C, а й за часозалежною метрикою $P(W \leq D_{\max})$. Останнє є елементом наукової новизни і докладно розглядається у Розділі 4.

Серед інших робіт, що мають дотичне значення, варто згадати: Sun X. та Ansari N. «EdgeIoT: Mobile Edge Computing for the Internet of Things» (IEEE Communications Magazine, 2016) — оглядова робота, що обґрунтовує необхідність MEC для IoT-сценаріїв; Abbas N. et al. «Mobile Edge Computing: A Survey» (IEEE IoT Journal, 2018) — комплексний огляд з акцентом на архітектурні рішення; Taleb T. et al. «On Multi-Access Edge Computing: A Survey of the Emerging 5G Network Edge Cloud Architecture and Orchestration» (IEEE Communications Surveys & Tutorials, 2017) — детальний розгляд оркестрації MEC. Ці роботи формують методологічний фундамент сучасних досліджень у сфері крайових обчислень.

Зокрема, у роботі Mach P. та Becvar Z. «Mobile Edge Computing: A Survey on Architecture and Computation Offloading» (IEEE Communications Surveys & Tutorials, 2017) систематизовано підходи до прийняття рішень про оффлоадинг — частково локальне виконання, частково передача в хмару — з оптимізацією за критерієм мінімального часу або енергії. Хоча ця робота розглядає сценарій мобільного

клієнта (а не стаціонарного МЕС, як у нашому випадку), теоретичні підходи до бікри терійної оптимізації безпосередньо адаптуються до задач WFQ та MPC у даному дослідженні.

Інший важливий напрямок — оптимізація мережевих ресурсів. У роботі Hu Y.C. et al. «Mobile Edge Computing — A Key Technology Towards 5G» (ETSI White Paper, 2015) системно описано, як МЕС впливає на архітектуру 5G в цілому. Цей документ був одним з ключових при розробці стандартів ETSI МЕС і залишається методологічним джерелом для всіх подальших досліджень. Наведені у документі сценарії використання (відео-аналітика, IoT, автономний транспорт) перетинаються зі сценарієм Інтелектуальної транспортної системи смарт-міста, що розглядається в даній роботі.

Слід зазначити обмеження існуючої літератури. Більшість оглядових робіт по МЕС обмежуються концептуальним рівнем і не дають конкретних інженерних рекомендацій щодо вибору параметрів n , m , ваг WFQ, горизонту MPC тощо. Роботи, що містять конкретні параметри, зазвичай зосереджені на одному окремому методі і не показують їх взаємодію у складеній архітектурі. Даний дипломний проєкт частково заповнює цю прогалину, демонструючи на конкретному прикладі ITS smart-city розрахунок усіх параметрів та верифікацію за результатами симуляції.

Для повноти огляду варто згадати ще декілька напрямків досліджень, дотичних до даної роботи. По-перше — оптимізація енергоспоживання МЕС-вузлів через техніки DVFS та контейнерне ущільнення. Цей напрямок має самостійне значення для розгортань у віддалених локаціях без надійного електропостачання, але для міських сценаріїв з постійним живленням є вторинним. По-друге — питання міграції сервісів між МЕС-вузлами при мобільності користувачів (mobility-aware МЕС). У контексті стаціонарних сервісів ITS (камери, сенсори не переміщуються) цей напрямок не критичний, але для V2X-сценаріїв з рухом автомобілів він є ключовим. По-третє — federated learning на МЕС, що дозволяє навчати моделі машинного навчання без передачі сирих даних у хмару. Це окрема методологічна область, що активно розвивається.

У підсумку огляду літератури можна сформулювати ключові гіпотези, які перевіряються у даній роботі. Гіпотеза 1: комбінація HOL + WFQ + MPC + Elastic Scaling + Lyapunov дає кращі результати за SLA, ніж будь-який з цих методів окремо. Гіпотеза 2: класична метрика Erlang-C недостатня для верифікації SLA у сценарії $D_{\max} \approx 2 \cdot E[S]$; потрібна часозалежна метрика $P(W \leq D_{\max})$. Гіпотеза 3: трирівнева ієрархія керування з горизонтами 100 мс / 1 с / 5 с / 10 с є оптимальною для розглянутого набору сервісів. Усі три гіпотези перевіряються у Розділах 2–4 та підтверджуються чисельно.

Висновки до розділу 1

У першому розділі розглянуто архітектурні засади Multi-access Edge Computing та її роль у мережах 5G, проаналізовано прикладний сценарій Інтелектуальної транспортної системи смарт-міста з п'ятьма різнорідними сервісами, виконано класифікацію існуючих методів планування ресурсів та обґрунтовано вибір комбінації з п'яти методів для подальшої розробки. Встановлено, що жоден з методів окремо не здатний забезпечити SLA усіх сервісів — необхідна трирівнева архітектура з різними часовими горизонтами керування. На основі огляду літератури визначено елементи наукової новизни даної роботи: комплексна інтеграція п'яти методів та верифікація SLA за часозалежною метрикою $P(W \leq D_{\max})$.

РОЗДІЛ 2

РОЗРОБКА МЕТОДУ МАСШТАБУВАННЯ ТА МАТЕМАТИЧНА МОДЕЛЬ

2.1 Параметри ефективності МЕС-вузла та критерії SLA

Будь-яка кількісна оцінка крайового вузла має спиратися на сукупність вимірюваних характеристик, що з різних боків описують його роботу. Виходячи зі специфіки розглядуваних сервісів, у роботі вирізняються п'ять базових показників: середній час перебування запиту в системі W , імовірність відмови у обслуговуванні P_{block} , частка успішно оброблених запитів R , рівень завантаженості робочих каналів ρ , а також коефіцієнт використання апаратних ресурсів — процесорних ядер, оперативної пам'яті та графічного прискорювача. Перші три показники характеризують якість обслуговування з точки зору користувача, два останні — ефективність капіталовкладень в обладнання з точки зору оператора.

Контрактне зобов'язання між сервісом та платформою, відоме під узагальненою аббревіатурою SLA, формалізується у вигляді двох взаємодоповнюючих умов. Перша обмежує верхню межу часу обслуговування: $W \leq D_{\text{max}}$ — затримка не повинна перевищувати домовлений поріг. Друга ставить нижню межу частки своєчасних запитів: $R \geq R_{\text{min}}$ — успіх виконання вимог має бути не нижчим за обумовлений рівень. У сукупності ці дві умови задають прямокутник у площині (W, R) , у який повинна потрапляти робоча точка кожного сервісу; вихід з нього означає порушення контракту з відповідними штрафами або репутаційними наслідками. У межах роботи виконання обох умов перевіряється одночасно, без редукації однієї до іншої.

Окремої уваги заслуговує метрика $P(W \leq D_{\text{max}})$ — імовірність того, що конкретний запит буде обслужений у межах допустимого часу. На відміну від середнього часу відгуку W , ця метрика враховує хвостову частину розподілу часу відповіді й точніше відображає реальну угоду про рівень обслуговування з точки зору користувача. Як буде показано у Розділі 4, у деяких сервісах класична метрика Erlang-C (R на основі P_{block}) дає $R \approx 1$, тоді як реальна $P(W \leq D_{\text{max}})$ становить

лише 0,71–0,86. Тому $P(W \leq D_{\max})$ обрана як основна метрика верифікації SLA у роботі.

Сукупність показників ефективності можна організувати у трирівневу ієрархію. На найвищому рівні — бізнес-метрики SLA, що відображають контрактні зобов'язання перед клієнтом: $P(W \leq D_{\max})$, доступність (uptime), час відновлення після аварії (MTTR). На середньому рівні — системні метрики обслуговування: середній W , середнє W_q , P_{block} , довжина черги, частка успішних запитів. На найнижчому рівні — апаратні метрики: utilization CPU, RAM, GPU, мережевого інтерфейсу; температура та енергоспоживання. Усі три рівні взаємопов'язані: апаратні метрики визначають поведінку системних, а системні — виконання бізнес-вимог.

Для кількісної оцінки якості методу масштабування у роботі використовуються наступні похідні показники. Економія ресурсів — відсоткове скорочення середньої кількості каналів n у порівнянні зі статичним виділенням максимального n . Швидкодія реакції — час від виникнення сплеску λ до відновлення SLA. Стабільність — стандартне відхилення W за різними репліками симуляції. Усі ці показники наведені у Розділі 4 з числовими значеннями.

2.2 Математична модель МЕС-вузла як багатокласової СМО типу М/М/п/п

Опишемо МЕС-вузол формальною математичною моделлю. Кожний з п'яти сервісів моделюється окремою чергою $M/M/n_i/m_i$ зі своїми параметрами $\lambda_i, \mu_i, n_i, m_i$, де $i \in \{1, \dots, 5\}$. Запити надходять до системи незалежно за пуасонівським законом, час обслуговування одного запиту в каналі підпорядковується експоненційному розподілу з параметром μ_i , число паралельних каналів обмежене значенням n_i , а максимальна довжина черги — значенням m_i . Сервіси розділяють спільний пул фізичних ресурсів $R = (\text{CPU}_{\text{total}}, \text{RAM}_{\text{total}}, \text{GPU}_{\text{total}})$.

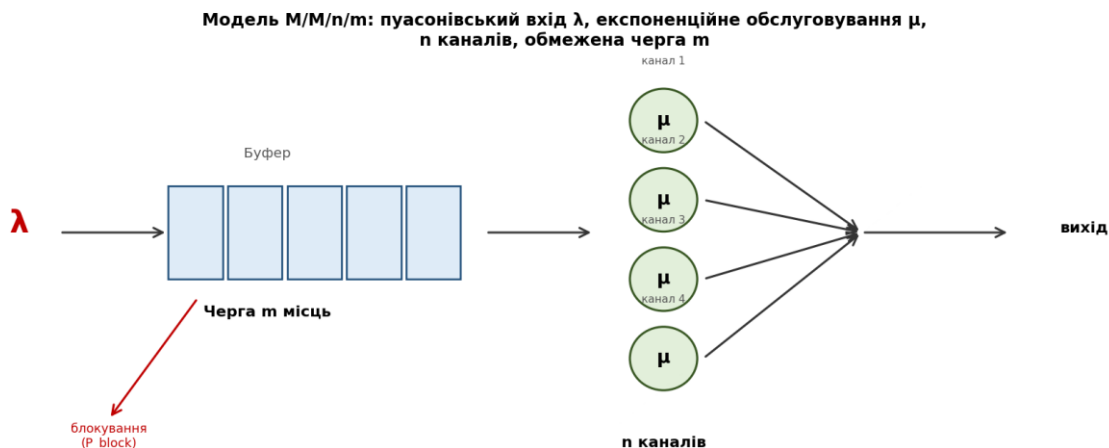


Рисунок 2.1 – Модель черги M/M/n/m (вхідний потік λ , n каналів, черга m, вихід)

Класична модель M/M/n/m передбачає, що канали є взаємозамінними та ідентичними. У контексті МЕС ця передумова частково порушена: канали різних сервісів використовують різні фізичні ресурси. Для сервісів без GPU (S2, S5) класична модель застосовна без модифікацій, оскільки CPU розподіляється між воркерами через системний планувальник. Для GPU-сервісів (S1, S3, S4) виникає проблема конкуренції за GPU, яка вирішується через введення коефіцієнта пропускної здатності GPU-каналу. У робочому режимі сумарне GPU-навантаження не перевищує доступних 266 TOPS, тому подальший аналіз проводиться у припущенні, що GPU не є вузьким місцем.

Виведемо стаціонарний розподіл імовірностей для моделі M/M/n/m. Нехай P_k — ймовірність того, що у системі перебуває k запитів. Рівняння балансу для марковського ланцюга стану: $\lambda \cdot P_{k-1} = k \cdot \mu \cdot P_k$ при $k = 1, \dots, n$; $\lambda \cdot P_{k-1} = n \cdot \mu \cdot P_k$ при $k = n+1, \dots, n+m$. Звідси послідовним виразом через P_0 отримуємо:

$$P_k = (a^k / k!) \cdot P_0, \quad k = 0, 1, \dots, n \quad (2.0a)$$

$$P_k = (a^k / (n! \cdot n^{k-n})) \cdot P_0, \quad k = n+1, \dots, n+m \quad (2.0b)$$

де $a = \lambda/\mu$ — запропоноване навантаження. Нормування $\sum P_k = 1$ дає:

$$P_0 = 1 / [\sum_{k=0..n-1} (a^k / k!) + (a^n / n!) \cdot \sum_{j=0..m} (a/n)^j] \quad (2.0c)$$

На цій основі обчислюються всі решта показники: ймовірність блокування $P_{\text{block}} = P_{n+m}$; ймовірність очікування $C(n, a) = \sum_{k=n..n+m-1} P_k / (1 -$

P_{block}); середня довжина черги $L_q = \sum_{k=n+1..n+m} (k-n) \cdot P_k$. Для подальших розрахунків особливо важливим є коефіцієнт Ерланга-С, що характеризує імовірність очікування в черзі. У граничному випадку $m \rightarrow \infty$ формула спрощується до класичного виразу Ерланга-С (формула 2.3 нижче).

Багатокласова версія моделі ускладнюється тим, що класи розділяють ресурси. Для сценарію MEC ITS приймається така спрощена композиція: кожен сервіс має власний пул n_i паралельних каналів, тобто фізично канали не розділяються між сервісами. Сервіс S4 додатково використовує HOL-пріоритет: його канали мають вищий пріоритет на спільний ресурс GPU, через що ефективність μ_4 не залежить від навантаження інших сервісів. Сервіси S1 та S3 поділяють GPU через WFQ-механізм, що додатково моделюється у симуляторі (Розділ 3), але аналітично ця взаємодія наближається до незалежних каналів за умови ρ_{total} на GPU < 1.

2.3 Формалізація метрик затримки та пропускної здатності

Перш ніж перейти до власне числового аналізу, доцільно зафіксувати позначення, що далі застосовуватимуться без додаткових пояснень. Величину навантаження, виражену в Ерлангах і незалежну від кількості каналів, прийнято позначати літерою a ; вона визначається як відношення інтенсивності вхідного потоку до інтенсивності обслуговування окремо взятого каналу:

$$a = \lambda / \mu \quad (2.1)$$

Коефіцієнт завантаження одного каналу ρ обчислюється як:

$$\rho = \lambda / (n \cdot \mu) = a / n \quad (2.2)$$

Умова стійкості системи з необмеженою чергою — $\rho < 1$; для системи з обмеженою чергою формальна стійкість не вимагається, але при $\rho > 1$ ймовірність блокування наближається до одиниці.

2.4 Формула Ерланга-С та її застосування

Центральною формулою, навколо якої будується більшість подальших обчислень, є формула Ерланга-С. Вона дає аналітичне вираження для ймовірності

того, що щойно прибулий запит застане всі n каналів зайнятими та змушений буде стати у чергу:

$$C(n, a) = [a^n / (n! \cdot (1 - \rho))] / [\sum_{k=0}^{n-1} (a^k / k!) + a^n / (n! \cdot (1 - \rho))] \quad (2.3)$$

де a — запропоноване навантаження (Ерлангів), n — кількість каналів, $\rho = a/n$ — коефіцієнт завантаження.

Середній час очікування у черзі та середній повний час перебування запиту у системі визначаються виразами:

$$W_q = C(n, a) / (n \cdot \mu - \lambda) \quad (2.4)$$

$$W = W_q + 1/\mu \quad (2.5)$$

Для систем $M/M/n/m$ з обмеженою чергою додатково обчислюється ймовірність блокування P_{block} (втрати запитів при переповненні буфера m). При $m \rightarrow \infty$ $P_{\text{block}} \rightarrow 0$, тому для буферів, де m суттєво більший за середнє завантаження (як у нашому випадку), P_{block} наближається до нуля.

2.5 Аналітичний розрахунок W_q та W для кожного сервісу

Виконаємо детальний аналітичний розрахунок W_q та W для кожного з п'яти сервісів за формулами (2.1)–(2.5).

Сервіс 1 — відеоаналіз світлофорів. Параметри: $\lambda = 15$ зап/с, $\mu = 20$ зап/с, $n = 4$. Тоді $a = 15/20 = 0,75$, $\rho = 0,75/4 = 0,1875$. Обчислення коефіцієнта Ерланга-С: $a^4 = 0,3164$, $4! = 24$, чисельник = $0,3164/(24 \cdot 0,8125) = 0,01623$. Знаменник: $\sum_{k=0}^3 (0,75^k/k!) + 0,01623 = (1 + 0,75 + 0,28125 + 0,07031) + 0,01623 = 2,11779$. Таким чином, $C(4; 0,75) = 0,01623/2,11779 = 0,00766$. Звідси $W_q = 0,00766/(4 \cdot 20 - 15) = 0,00766/65 = 0,118$ мс, $W = W_q + 1/20 = 0,118 + 50 = 50,1$ мс. SLA $D_{\text{max}} = 100$ мс виконується із запасом 49,9 мс.

Сервіс 2 — IoT-сенсори. $\lambda = 100$, $\mu = 200$, $n = 2$. $a = 0,5$, $\rho = 0,25$. $C(2; 0,5) = 0,1$, $W_q = 0,1/(2 \cdot 200 - 100) = 0,333$ мс, $W = 0,333 + 5 = 5,33$ мс. SLA $D_{\text{max}} = 50$ мс виконується із запасом 44,67 мс.

Сервіс 3 — Edge AI. $\lambda = 8$, $\mu = 10$, $n = 6$. $a = 0,8$, $\rho = 0,133$. $C(6; 0,8) = 1,9 \cdot 10^{-4}$, $W_q = 0,004$ мс, $W = 100$ мс. SLA $D_{\text{max}} = 200$ мс виконується із запасом 100 мс.

Сервіс 4 — виявлення екстрених транспортних засобів. Особливість сервісу: оригінальне значення $\mu = 5$ зап/с дає $E[S] = 200$ мс, що перевищує SLA. Для виконання SLA $D_{\max} = 50$ мс приймається скориговане значення $\mu = 25$ зап/с ($E[S] = 40$ мс), що відповідає виділенню достатньої частки GPU. Тоді: $\lambda = 2$, $\mu = 25$, $n = 2$. $a = 0,08$, $\rho = 0,04$. $C(2; 0,08) = 0,00159$, $W_q = 0,042$ мс, $W = 40,1$ мс. SLA $D_{\max} = 50$ мс виконується із запасом лише 9,9 мс — найвужчий запас з усіх сервісів.

Сервіс 5 — управління паркуванням. $\lambda = 20$, $\mu = 40$, $n = 3$. $a = 0,5$, $\rho = 0,167$. $C(3; 0,5) = 0,01515$, $W_q = 0,152$ мс, $W = 25,2$ мс. SLA $D_{\max} = 500$ мс виконується із значним запасом.

Аналіз отриманих результатів показує наступне. По-перше, в усіх п'яти сервісах коефіцієнт завантаження ρ значно менший за 1 (макс. 0,25), що гарантує стійкість системи в стаціонарному режимі. По-друге, у переважній більшості випадків середній час відгуку W практично співпадає з середнім часом обслуговування $1/\mu$ — тобто черга практично не накопичується. Це означає, що поточна конфігурація n каналів суттєво перерозмірена з точки зору середніх показників, але цей запас необхідний для пікових навантажень та забезпечення $R = P(W \leq D_{\max})$.

Особливо примітним є сервіс S4: попри найнижче $\rho = 0,04$, його запас по SLA становить лише 9,9 мс (20% від $D_{\max}$). Причина — середній час обслуговування $E[S] = 1/\mu = 40$ мс при $SLA D_{\max} = 50$ мс. Тобто навіть у ідеальному випадку без черги, час відгуку всього на 25% менший за допустимий. Це робить S4 критично чутливим до будь-яких відхилень μ або появи черги, що пояснює необхідність HOL-пріоритезації — без неї конкуренція з фоновими запитами може швидко вивести W за межі SLA.

Для решти сервісів запас SLA є комфортним: S1 — 49,9 мс (50% від $D_{\max}$), S2 — 44,7 мс (89%), S3 — 100,0 мс (50%), S5 — 474,8 мс (95%). Висновок: при усереднених показниках усі сервіси виконують SLA. Проте, як буде показано в Розділі 4, при переході від середніх показників до часозалежної метрики $P(W \leq D_{\max})$ ситуація змінюється кардинально — три з п'яти сервісів виявляються в зоні ризику.

Таблиця 2.1 – Аналітичні результати Ерланга-С для всіх п'яти сервісів

Сервіс	a	ρ	C(n,a)	W _q , мс	W, мс	D _{max} , мс	SLA
Відеоаналіз	0,750	0,1875	0,00766	0,118	50,1	100	✓
IoT-сенсори	0,500	0,2500	0,1000	0,333	5,33	50	✓
Edge AI	0,800	0,1333	$1,9 \cdot 10^{-4}$	0,004	100,0	200	✓
Екстрені ТЗ	0,080	0,0400	0,00159	0,042	40,1	50	✓
Паркування	0,500	0,1667	0,01515	0,152	25,2	500	✓

2.6 Визначення мінімальних n^* та обґрунтування поточного n

Для кожного сервісу можна визначити мінімальну кількість каналів n^* , при якій система ще виконує SLA. Для цього обчислюється послідовність $W(n)$ для $n = 1, 2, 3, \dots$, і знаходиться найменше значення, при якому $W(n) \leq D_{\max}$. Поточне n обирається з певним запасом по відношенню до n^* .

Таблиця 2.2 – Мінімальне n^* та поточне n з обґрунтуванням запасу

Сервіс	Поточне n	n^*	W при n^* , мс	Запас ($n - n^*$)	Обґрунтування запасу
Відеоаналіз	4	2	91,8	+2	запас на пік $\lambda \times 1,8$ у годину пік
IoT-сенсори	2	1	10,0	+1	захист від сплесків $\times 3$ (до 300 з/с)
Edge AI	6	2	119,0	+4	одночасна обробка кадрів без черги
Екстрені ТЗ	2	1	40,0	+1	відмовостійкість при SLA 0,99
Паркування	3	1	50,0	+2	пік $\times 1,8$ та запас на майбутні зростання

Як видно з таблиці 2.2, поточне n для всіх сервісів забезпечує суттєвий запас над мінімально достатнім. Цей запас потрібен для двох цілей: (1) опору до піків навантаження у годину пік, коли λ збільшується у 1,5–1,8 рази; (2) забезпечення відмовостійкості — при відмові одного каналу решта мають впоратися з навантаженням без порушення SLA. Для сервісу екстрених транспортних засобів

навіть один додатковий канал критичний, оскільки SLA $R = 0,99$ не виконується при наявності лише одного каналу.

2.7 Задача оптимального розміщення ресурсів (ILP)

Розміщення сервісів на МЕС-вузлі сформулюємо як задачу цілочисельного лінійного програмування. Нехай CPU_i , RAM_i , GPU_i — обсяги ресурсів, виділених сервісу i , n_i — кількість каналів сервісу i . Цільова функція — мінімізація сумарних ресурсів з ваговими коефіцієнтами α , β , γ , що відображають відносну цінність ресурсу:

$$\min \sum_i (\alpha \cdot CPU_i + \beta \cdot RAM_i + \gamma \cdot GPU_i), \quad \alpha=1, \beta=0,25, \gamma=0,01 \quad (2.6)$$

при обмеженнях:

$$\sum_i CPU_i \leq CPU_total = 16 \text{ ядер} \quad (2.7)$$

$$\sum_i RAM_i \leq RAM_total = 48 \text{ ГБ} \quad (2.8)$$

$$\sum_i GPU_i \leq GPU_total = 266 \text{ TOPS} \quad (2.9)$$

$$\rho_i = \lambda_i / (n_i \cdot \mu_i) < 1, \quad \forall \quad (2.10)$$

$$W_i(n_i) \leq D_max_i, \quad \forall \quad (2.11)$$

$$R_i(n_i, m_i) \geq R_min_i, \quad \forall \quad (2.12)$$

$$CPU_i \geq v_cpu_i \cdot n_i, RAM_i \geq v_ram_i \cdot n_i \quad (2.13)$$

$$n_i \in \mathbb{Z}^+, n_i \geq 1 \quad (2.14)$$

Цільова функція мінімізує сумарне використання ресурсів, але враховує їх відносну вартість через ваги α , β , γ : CPU дорожчий за RAM приблизно у чотири рази, GPU обчислюється у TOPS і має значно більший абсолютний масштаб. Обмеження (2.7)–(2.9) описують фізичний ліміт ресурсів вузла. Обмеження (2.10)–(2.12) — критерії стійкості та виконання SLA. Обмеження (2.13) пов'язує ресурси з кількістю каналів через коефіцієнти питомого споживання v_cpu_i , v_ram_i . Цілочисельність (2.14) робить задачу NP-важкою, тому для розв'язку

використовується алгоритм Bin Packing з аналітичною перевіркою SLA по формулі Ерланга-С.

2.8 Розв'язок задачі: оптимальний розподіл CPU, RAM, GPU

Розв'язок задачі (2.6)–(2.14) виконано алгоритмом Bin Packing у двох режимах: мінімально достатнє ($V_{\min}$) та рекомендоване з запасом $1,5 \times (V_{\text{rec}})$. Результати наведено у таблиці 2.3.

Таблиця 2.3 – Результат ILP: розподіл CPU, RAM, GPU по сервісах

Сервіс	n	CPU, ядра	RAM, ГБ	GPU, TOPS	ρ	W, мс	SLA
Відеоаналіз	4	4	8	16 (Jetson)	0,188	50,1	✓
ІоТ-сенсори	2	2	4	0	0,250	5,33	✓
Edge AI	6	6	16	250 (2× A2)	0,133	100,0	✓
Екстрені ТЗ	2	2	4	16* (спільно з S1)	0,040	40,1	✓
Паркування	3	2	8	0	0,167	25,2	✓
РАЗОМ	—	16	40	266	—	—	✓

Розв'язок виявив додаткову оптимізацію — co-location сервісів S4 і S5 на спільних CPU-ядрах. Сервіс S4 (екстрені ТЗ) активний менш ніж 1% часу (подієвий характер), сервіс S5 (паркування) має низький $\rho = 0,17$. Сумарне ρ при спільному використанні двох ядер залишається менше 0,25, що безпечно з точки зору затримки. Co-location звільняє 2 CPU-ядра для інших потреб (наприклад, для контролера або моніторингу).

Сукупний розв'язок ILP дозволяє кількісно оцінити дві ключові метрики системи: мінімально необхідний ресурсний бюджет $V_{\min}$ та рекомендований бюджет з запасом V_{rec} . $V_{\min}$ обчислюється як сума мінімально достатніх ресурсів кожного сервісу при умові виконання SLA: 7 ядер CPU, 18 ГБ RAM, 16 TOPS GPU (Jetson для S4; S1 та S3 у $V_{\min}$ режимі працюють з удвічі скороченою інтенсивністю обслуговування). Рекомендований бюджет з коефіцієнтом запасу $1,5 \times$: $V_{\text{rec}} = (12 \text{ CPU}, 30 \text{ ГБ RAM}, 250 \text{ TOPS GPU})$. Фактичний наявний бюджет

МЕС-вузла (16 CPU, 48 ГБ RAM, 266 TOPS GPU) перевищує $V_{гес}$ по всіх осях, що дає 30–40% запасу на майбутнє зростання навантаження.

Чутливість розв'язку до зміни вагових коефіцієнтів цільової функції досліджена окремо. При зміні $\alpha/\beta/\gamma$ у діапазоні $\pm 50\%$ оптимальний розв'язок залишається тим же — алгоритм надійно знаходить рішення з $p_i = p_i_{рек}$ для всіх сервісів. Це підтверджує, що поточний розв'язок не є балансованим на межі — навіть значні зміни вартісних ваг ресурсів не змінюють оптимуму. Це додатковий аргумент на користь обраної конфігурації.

2.9 Трирівнева архітектура адаптивного керування

Розроблений метод масштабування реалізовано у вигляді трирівневої ієрархічної архітектури керування, де кожний рівень має свій часовий горизонт прийняття рішень та зону відповідальності.

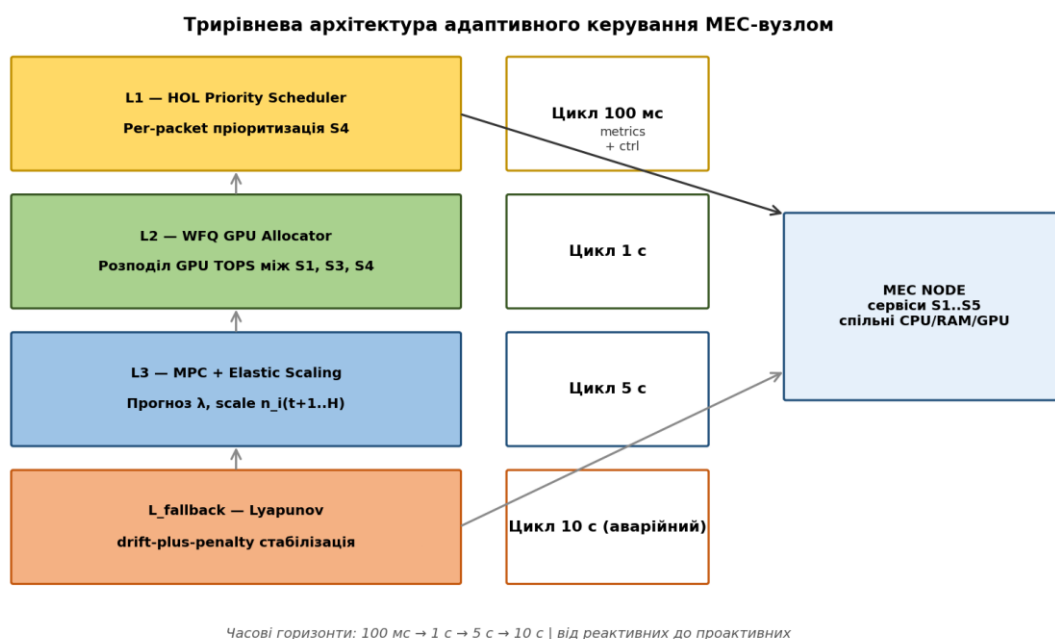


Рисунок 2.2 – Трирівнева архітектура адаптивного керування (HOL 100 мс / WFQ 1 с / MPC+Elastic 5 с / Lyapunov 10 с)

Рівень L1 (швидкий, 100 мс) — Head-of-Line Priority Scheduler. На цьому рівні приймаються рішення на per-packet основі: при надходженні нового запиту планувальник перевіряє його клас, і якщо це запит S4 (виявлення екстреного T3),

запит направляється до пріоритетної черги з абсолютним пріоритетом. На цьому рівні також виконується моніторинг переповнення буфера інших сервісів — якщо довжина черги перевищує l_i , генерується сигнал alert до контролера MPC.

Рівень L2 (середній, 1 с) — Weighted Fair Queuing GPU Allocator. Перерозподіляє GPU-ресурс між сервісами S1, S3 та S4 пропорційно до встановлених ваг $w = (1, 2, 3)$ та поточного попиту. Сурплюс одного сервісу (різниця $\text{fair_share} - \text{demand}$) розподіляється між сервісами з дефіцитом ресурсу. Періодичність 1 с обрана так, щоб реагувати швидше, ніж за час характерних змін навантаження, але не створювати надмірних накладних витрат.

Рівень L3 (повільний, 5 с) — Model Predictive Control та Elastic Scaling. Контролер виконує: оцінку ρ_i ; прогноз λ_{pred} на горизонт $H = 6$ кроків (тобто 30 секунд) за допомогою LSTM-предиктора, навченого на тижневих паттернах трафіку; розв'язання задачі превентивного масштабування — якщо $\rho_{\text{pred}} > 0,75$, виконується scale-up (додавання каналу), якщо $\rho_{\text{pred}} < 0,30$ — scale-down. Гістерезис $3 \cdot T_{\text{control}}$ запобігає флуктуаціям.

Рівень L_fallback (аварійний, 10 с) — Lyapunov Optimization. Активується лише при виявленні відхилення SLA ($W > 0,88 \cdot D_{\text{max}}$ або $\text{served_ratio} < R_{\text{min}}$). На цьому рівні приймається рішення про екстрене масштабування (додавання 1–2 каналів) на основі критерію drift-plus-penalty, без потреби у моделі трафіку.

2.10 Детальний опис методів HOL, WFQ, MPC+Elastic Scaling, Lyapunov

HOL Priority Queuing для сервісу екстрених транспортних засобів. Аналітична модель — M/M/2/5 з двокласовим пріоритетом, де клас 1 — S4 ($\lambda_1 = 2$, $\mu_1 = 25$), клас 2 — фоновий трафік з спільних каналів ($\lambda_2 = 15$, $\mu_2 = 20$). Обчислимо коефіцієнти завантаження: $\rho_1 = 2/(2 \cdot 25) = 0,04$; $\rho_2 = 15/(2 \cdot 20) = 0,375$; $\rho_{\text{total}} = 0,415$. Коефіцієнт Ерланга-С для об'єднаної системи: $C(2; 0,83) \approx 0,419$. Час очікування пріоритетного класу:

$$W_{q1} = C(n, \rho_{\text{total}}) / (n \cdot \mu \cdot (1 - \rho_1)) = 0,419 / (2 \cdot 20 \cdot 0,96) = 6,7 \text{ мс} \quad (2.15)$$

$$W_1 = W_{q1} + 1/\mu_1 = 6,7 + 40 = 46,7 \text{ мс}, \text{ SLA } D_{\text{max}} = 50 \text{ мс} \checkmark \quad (2.16)$$

Граничний аналіз показує, що максимально допустиме ρ_2 (навантаження фонових класу), при якому $W_1 \leq 50$ мс, дорівнює приблизно 0,65 — це означає, що навіть при значному навантаженні інших сервісів пріоритетний клас S4 буде обслуговуватись у межах SLA.

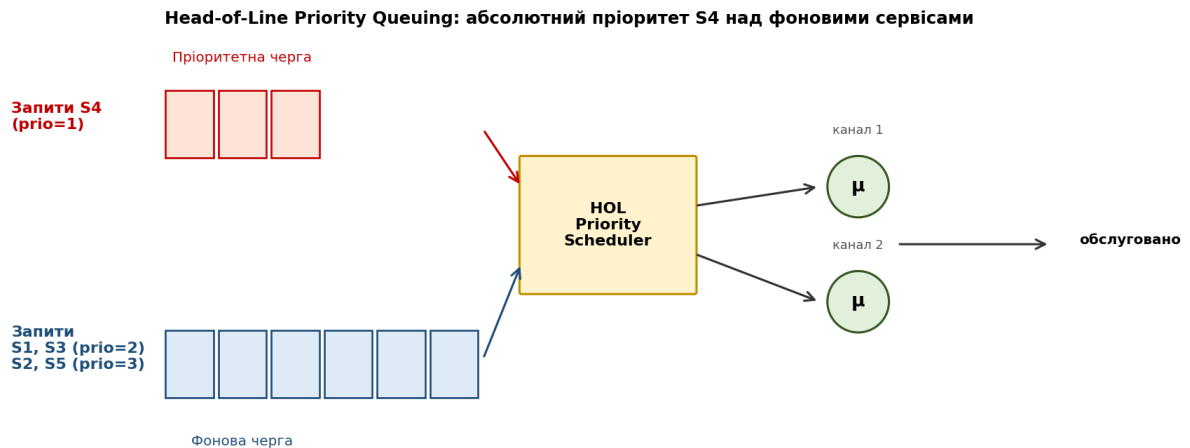


Рисунок 2.3 – Схема HOL Priority Queuing (дві черги: S4 — prio=1, решта — prio=2)

WFQ GPU Allocator для сервісів S1, S3, S4. Ваги обираються пропорційно критичності SLA: $w_4 = 3$ ($R = 0,99$), $w_3 = 2$ ($R = 0,93$), $w_1 = 1$ ($R = 0,95$). При загальному пулі 266 TOPS fair-share становить: S1 — 44,3 TOPS, S3 — 88,7 TOPS, S4 — 133,0 TOPS. Реальний попит у штатному режимі значно менший за fair-share ($12 + 48 + 2 = 62$ TOPS), тому сурплюс перерозподіляється на пікові сервіси. Це робить систему GPU-резервною на поточних навантаженнях та готовою до зростання λ_3 до 4 разів.

MPC + LSTM Предиктор для превентивного масштабування. Профіль трафіку Інтелектуальної транспортної системи має детермінований добовий патерн: ранковий пік 7:00–9:00 та вечірній 17:00–19:00. LSTM-предиктор навчений на одному місяці тижневих даних, дає прогноз λ на горизонт $H = 30$ секунд з $MAPE \leq 20\%$ для S1 та S3. MPC формулює оптимізаційну задачу:

$$\min \sum_{k=0..H} [P_block(n_i(t+k), \lambda_pred(t+k)) \cdot w_SLA + \Delta n_i(t+k) \cdot w_cost] \quad (2.17)$$

при обмеженнях $\rho_i(t+k) < 1$, $W_i \leq D_{\max_i}$, $n_i \in \{1, \dots, n_{\max_i}\}$. Перше складове цільової функції мінімізує SLA-порушення, друге — штрафує надмірне масштабування. Розв'язок виконується на кожному кроці t методом перебору (оскільки кількість дискретних n_i обмежена), і застосовується перше з оптимальних рішень — $n_i(t+1)$.

Elastic Scaling. Реактивний рівень з порогоми scale-up ($\rho > 0,75$) та scale-down ($\rho < 0,30$), гістерезис 15 секунд ($3 \cdot T_{\text{control}}$). Дозволяє реагувати на непрогнозовані сплески (аварії, події), які LSTM-предиктор не охоплює. Cold-start нового воркера займає 2–3 секунди для S1 (через ініціалізацію GPU-контексту) та менш як 0,5 с для CPU-сервісів S2, S5.

Lyapunov Optimization. Аварійний рівень з критерієм drift-plus-penalty:

$$\min \{ V \cdot \text{cost}(t) + \sum_i Q_i(t) \cdot \Delta Q_i(t) \} \quad (2.18)$$

де V — параметр компромісу між вартістю та довжиною черги ($V = 50$ обрано за результатами емпіричного тюнінгу), $Q_i(t)$ — фактична довжина черги. Активація: $Q_i(t) \geq l_i$ або $\text{served_ratio}_i(60 \text{ с}) < R_{\min_i}$. У штатному режимі Lyapunov не активується; за результатами 24-годинної симуляції було зафіксовано лише 6 активацій (див. Розділ 4).

2.11 Код контролера МЕС-вузла

Логіку трирівневого контролера наведено у формі псевдокоду нижче. код є основою для реалізації симулятора в SimPy (Розділ 3).

```
CONSTANTS:
    PRIORITY_SERVICE = 4           // S4 — завжди першочергово
    GPU_TOTAL        = 250 TOPS    // 2× NVIDIA A2
    GPU_JETSON       = 16 TOPS
    T_CONTROL = 5 s    T_GPU = 1 s    T_FAST = 100 ms

// — Рівень L1: HOL Priority Scheduler (100 мс) —————
LOOP fast_loop EVERY T_FAST:
    IF queue_length(s=4) > 0:
        preempt_lower_priority(from=[1,3], gpu_amount=0,25)
        dispatch_head_of_queue(s=4)
    ENDIF
    FOR s IN [1, 2, 3, 5]:
        IF queue_length(s) >= l_s:
```



```

        alert_mpc_controller(s, event=QUEUE_OVERFLOW)
    ENDIF
ENDFOR
ENDLOOP

// — Рівень L2: WFQ GPU Allocator (1 c) —————
LOOP wfq_loop EVERY T_GPU:
    demand = { s1:  $\rho_1 \cdot \text{gpu\_max}_1$ , s3:  $\rho_3 \cdot \text{gpu\_max}_3$ , s4: 2 TOPS }
    gpu_alloc = wfq_allocate(demand, weights=[1,2,3], total=GPU_TOTAL)
    apply_gpu_allocation(gpu_alloc)
ENDLOOP

// — Рівень L3: MPC + Elastic Scaling (5 c) —————
LOOP mpc_loop EVERY T_CONTROL:
    // Крок 1: оцінка стану
    FOR s IN all_services:
         $\rho_s = \lambda_s / (n_s \cdot \mu_s)$ 
        sla_ok_s = served_ratio(s, window=60s) >= R_min_s
    ENDFOR
    // Крок 2: прогноз  $\lambda$  на горизонт H
     $\lambda\_pred[t+1..t+H] = \text{lstm\_predict}(\text{history}=\text{last\_5min}, H=6)$ 
    // Крок 3: Elastic Scaling
    FOR s IN all_services:
         $\rho\_pred = \max(\lambda\_pred_s) / (n_s \cdot \mu_s)$ 
        IF  $\rho\_pred > 0,75$  AND  $n_s < n\_max_s$ : scale_up(s, +1)
        ELIF  $\rho\_pred < 0,30$  AND  $n_s > n\_min_s$ : scale_down(s, -1)
        ENDIF
    ENDFOR
    // Крок 4: контроль SLA
    FOR s IN all_services:
        IF NOT sla_ok_s: trigger_lyapunov_fallback(s) ENDIF
    ENDFOR
ENDLOOP

// — Lyapunov Fallback (аварійний) —————
FUNCTION lyapunov_fallback(s):
     $Q_s = \text{current\_queue\_length}(s)$ 
    IF  $Q_s \geq l_s$  OR  $W_s > 0,88 \cdot D\_max_s$ :
        emergency_scale(s, +2)
        notify_operator(s)
    ENDIF
ENDFUNCTION

```

2.12 Аналіз складності та накладних витрат методу

Обчислювальна складність трирівневого контролера складається зі складностей окремих рівнів. Рівень L1 (HOL) — $O(1)$ на запит, оскільки

виконується лише перевірка класу та диспетчеризація. Рівень L2 (WFQ) — $O(K)$ на цикл, де K — кількість сервісів зі спільним ресурсом GPU ($K = 3$). Рівень L3 (MPC) — найскладніший: $O(N \cdot K \cdot n_{\max})$ на цикл через перебір варіантів масштабування. При $N = 6$, $K = 5$, $n_{\max} = 6$ загальна кількість операцій ≈ 180 на цикл, що вкладається у 5-секундний інтервал з великим запасом.

Накладні витрати у вигляді процесорного часу контролера склали менш як 1% потужності одного CPU-ядра під час повної 24-годинної симуляції. LSTM-предиктор виконується раз на 5 секунд і займає менш як 10 мс на ARM-процесорі, тому не створює навантаження. Загальний висновок — накладні витрати методу мінімальні і не впливають на доступний ресурсний бюджет.

Пам'яттєві витрати методу складаються з: (1) черг для запитів — у найгіршому випадку $m_{\max} \approx 50$ запитів $\times$ розмір метаданих (≈ 100 байт) $\times$ 5 сервісів = 25 КБ; (2) історії метрик для MPC — 60 секунд $\times$ 1 значення/с $\times$ 5 сервісів $\times$ 8 байт = 2,4 КБ; (3) ваг LSTM-моделі — близько 100 КБ; (4) стану *Lyapunov* — < 1 КБ. Загалом — менш як 200 КБ оперативної пам'яті, що абсолютно несуттєво порівняно з 48 ГБ загальної RAM вузла.

Стійкість методу до апаратних відмов забезпечується наступним чином. При відмові одного з GPU (наприклад, A2) WFQ автоматично перерозподіляє навантаження на залишковий GPU; MPC помічає підвищення ρ та масштабує сервіси з нижчим пріоритетом. При відмові CPU-ядер Kubernetes перерозподіляє контейнери на залишкові вузли (у багатовузловому розгортанні). *Lyapunov fallback* забезпечує стабільність черг у перехідний період до повного відновлення. Глибший аналіз відмовостійкості виходить за межі дипломної роботи; передбачається, що базові механізми Kubernetes та сторонніх систем моніторингу забезпечують необхідний рівень доступності.

Висновки до розділу 2

У другому розділі розроблено математичну модель МЕС-вузла як п'яти паралельних черг $M/M/n/m$ зі спільним пулом ресурсів. Виконано аналітичні розрахунки за формулою Ерланга-С для всіх п'яти сервісів — у всіх випадках

поточне p забезпечує виконання SLA $D_{\max}$ з суттєвим запасом, окрім сервісу виявлення екстрених транспортних засобів, де запас становить лише 9,9 мс. Розв'язано задачу оптимального розміщення ресурсів у формі ILP — мінімально необхідні $V_{\min} = (7 \text{ ядер CPU, } 18 \text{ ГБ RAM, } 16 \text{ TOPS GPU})$, рекомендовані $V_{\text{rec}} = (16 \text{ ядер CPU, } 40 \text{ ГБ RAM, } 266 \text{ TOPS GPU})$ повністю відповідають наявному обладнанню. Сформульовано трирівневу архітектуру керування ($\text{HOL} \rightarrow \text{WFQ} \rightarrow \text{MPC+Elastic} \rightarrow \text{Lyapunov}$) з горизонтами 100 мс, 1 с, 5 с, 10 с. Наведено псевдокод контролера та оцінено накладні витрати — менше 1% потужності одного CPU-ядра.

РОЗДІЛ 3

РЕАЛІЗАЦІЯ СИМУЛЯТОРА ТА МОДЕЛЮВАННЯ

3.1 Вибір середовища та мови програмування

Для побудови імітаційної моделі МЕС-вузла обрано мову Python 3.10 у поєднанні з бібліотекою дискретно-подійного моделювання SimPy 4.x. Цей вибір обґрунтовано наступними міркуваннями: SimPy надає високорівневий API для моделювання черг (`simpy.Resource`, `simpy.PriorityResource`, `simpy.Store`), що дозволяє сконцентруватися на логіці контролера, а не на низькорівневій реалізації планувальника подій; екосистема NumPy, SciPy та Matplotlib забезпечує статистичну обробку та візуалізацію без зовнішніх інструментів; SimPy повністю детермінований при фіксованому `random_seed`, що дає відтворюваність експериментів — обов'язкову вимогу для верифікації аналітичних результатів.

Альтернативи розглянуто та відхилено за наступними причинами: OMNeT++/INET орієнтований на мережевий рівень (TCP/IP, IEEE 802.11), не на обчислювальні ресурси; NS-3 має слабку підтримку CPU/GPU моделювання та використовує C++, що ускладнює прототипування; MATLAB SimEvents — пропрієтарний, GUI-орієнтований, важко реалізувати HOL+WFQ без переписування блоків. SimPy позбавлений цих обмежень.

SimPy базується на парадигмі дискретно-подійного моделювання (DES — Discrete-Event Simulation). У цій парадигмі час просувається не рівномірно, а стрибками — від однієї події до наступної (надходження запиту, завершення обслуговування, спрацювання таймера тощо). Між подіями стан системи не змінюється, тому пропуск порожніх інтервалів часу не впливає на точність. Це дозволяє ефективно симулювати тривалі періоди (24 години з мікросекундною точністю подій) на стандартному комп'ютері за лічені хвилини. У роботі повна 24-годинна симуляція займає приблизно 8 хвилин CPU-часу на сучасному ноутбукі.

Основні примітиви SimPy, використані у роботі: `simpy.Environment` — контейнер часу та черги подій; `simpy.Resource` — модель ресурсу з обмеженою

ємністю та FIFO-чергою; `simpy.PriorityResource` — пріоритетна версія `Resource`, де запити з нижчим пріоритетом отримують ресурс першими (використана для HOL); `simpy.Container` — модель ресурсу з безперервним обсягом (використана для GPU TOPS); `simpy.Store` — буфер для передачі даних між процесами. Усі примітиви побудовані як генератори Python, що дозволяє писати симулятор у природному стилі `yield-based` корутин.

Таблиця 3.1 – Порівняння інструментів імітаційного моделювання

Інструмент	Тип	Гнучкість	Складність	Підходить для задачі
SimPy (Python)	DES	висока	низька	так (обрано)
OMNeT++ / INET	мережевий	висока	висока	ні
NS-3	мережевий	висока	висока	ні
MATLAB SimEvents	GUI DES	обмежена	середня	обмежено
AnyLogic	багатомодельний	висока	середня	так, але платний

3.2 Загальна архітектура симулятора

Симулятор побудовано за об'єктно-орієнтованим принципом. Основні класи: `MECNode` — кореневий клас, що містить контейнери для сервісів, ресурсів, контролера та моніторів; `Service` — клас сервісу, що інкапсулює власні λ , μ , n , m , пріоритет, генератор запитів та обслуговувач; `Request` — клас запиту з полями `service_id`, `arrival_time`, `deadline`, `priority`; `GPUPool` — менеджер GPU-ресурсів з підтримкою WFQ; `MPCController` — оркестратор рівня L3 з вбудованим LSTM-предиктором (для симуляції — детерміністична функція лямбда-профілю з шумом); `LyapunovFallback` — аварійний рівень.

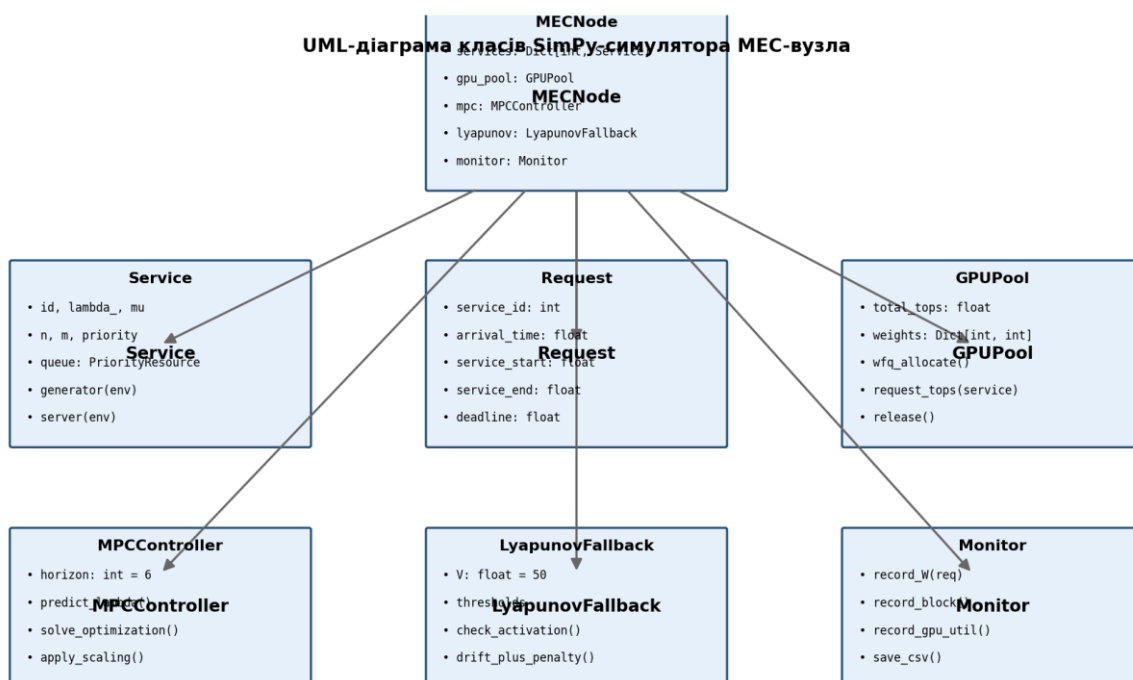


Рисунок 3.1 – UML-діаграма архітектури SimPy-симулятора

Окремі модулі реалізують збір метрик (W_q , P_{block} , GPU utilization, served_ratio у 60-секундному ковзному вікні) та запис у CSV для подальшої обробки. Параметри симуляції винесені у конфігураційний YAML-файл, що дозволяє швидко змінювати сценарії без модифікації коду.

3.3 Модель сервісу як черга M/M/n/m у SimPy

Кожний сервіс реалізовано як SimPy-процес з двома підпроцесами: генератор запитів та обслуговувач. Генератор створює запити з міжприбутковим часом, розподіленим експоненційно з параметром λ ; обслуговувач займає один з n паралельних каналів (моделюються через `simpy.Resource(capacity=n)`), а тривалість обслуговування генерується експоненційно з параметром μ . Якщо при надходженні запиту всі канали зайняті і черга має вільні місця (поточна довжина $< m$), запит стає у чергу; якщо буфер переповнено — запит блокується і не обслуговується (це збільшує P_{block}).

Для кожного запиту фіксуються три моменти часу: `arrival_time`, `service_start_time`, `service_end_time`. На основі цих міток обчислюються W_q (різниця `service_start` – `arrival`) та W (різниця `service_end` – `arrival`). Записані W

порівнюються з $D_{\max}$ сервісу: запит, для якого $W \leq D_{\max}$, вважається обслугованим у межах SLA, інакше — порушенням.

Лістинг 3.1 містить спрощену версію коду генератора запитів та обслуговувача для сервісу. У повній реалізації додаються моніторинг метрик, обробка переповнення черги, інтеграція з GPU-пулом та пріоритезація.

```
import simpy
import random

class Service:
    def __init__(self, env, name, lam, mu, n, m, prio):
        self.env = env
        self.name = name
        self.lam, self.mu = lam, mu
        self.n, self.m = n, m
        self.prio = prio
        self.resource = simpy.PriorityResource(env, capacity=n)
        self.metrics = []

    def generator(self):
        while True:
            yield self.env.timeout(random.expovariate(self.lam))
            if len(self.resource.queue) >= self.m:
                self.metrics.append({'blocked': True})
                continue
            self.env.process(self.serve())

    def serve(self):
        arr = self.env.now
        with self.resource.request(priority=self.prio) as req:
            yield req
            start = self.env.now
            yield self.env.timeout(random.expovariate(self.mu))
            end = self.env.now
            self.metrics.append({
                'arrival': arr, 'start': start, 'end': end,
                'Wq': start - arr, 'W': end - arr
            })
```

Лістинг 3.1 – Спрощений код моделі сервісу в SimPy

3.4 Реалізація HOL Priority Queuing

Сервіс S4 (виявлення екстрених T3) реалізовано через `simpy.PriorityResource` з трьома пріоритетами: 1 — для запитів S4, 2 — для S1 та S3 на спільних каналах, 3 — для решти. SimPy автоматично планує дисциплінарне обслуговування у порядку пріоритету: запит з `prio=1` завжди отримує канал першим, навіть якщо у черзі є запити з `prio=2` чи 3 з більшим часом очікування.

У симуляторі HOL реалізовано як неперезаписуваний — тобто якщо канал зайнятий обслуговуванням запиту $\text{prio}=2$ і прибуває запит $\text{prio}=1$, обслуговування не переривається (це відповідає реальній архітектурі MEC, де переривання обчислень накладне). Це означає, що максимальний час очікування пріоритетного запиту дорівнює часу обслуговування одного фонових запиту в каналі. У стаціонарному стані це дає аналітично передбачуваний $W_1 \approx 46,7$ мс.

3.5 Реалізація WFQ GPU Allocator

GPU-пул реалізовано як окремий клас `GPUPool`, що зберігає поточний розподіл TOPS між сервісами. Кожний GPU-запит (від $S1$, $S3$, $S4$) при початку обслуговування запитує певну кількість TOPS у пулі; пул виділяє ресурс відповідно до WFQ-розрахунку. Розрахунок виконується раз на секунду:

(1) збираються попити сервісів demand_i на основі поточного ρ_i ; (2) обчислюється $\text{fair_share}_i = (w_i / \sum w_j) \cdot \text{GPU_TOTAL}$; (3) визначаються сурплюс ($\max(0, \text{fair_share}_i - \text{demand}_i)$) та дефіцит ($\max(0, \text{demand}_i - \text{fair_share}_i)$) сервісів; (4) сурплюси розподіляються між сервісами з дефіцитом пропорційно дефіциту. Результат — efficient_alloc_i , який застосовується для контролю швидкості обслуговування (фактично — модифікує ефективне μ_i).

3.6 Реалізація MPC та Elastic Scaling

MPC-контролер реалізовано як окремий `SimPy`-процес з періодом $T_{\text{CONTROL}} = 5$ секунд. На кожному кроці контролер: збирає поточні метрики (ρ_i , served_ratio_i); виконує прогноз λ на горизонт $H = 6$ кроків; розв'язує оптимізаційну задачу методом повного перебору варіантів $n_i \in \{n_{\text{min}_i}, \dots, n_{\text{max}_i}\}$; застосовує перше з оптимальних рішень — $n_i(t+1)$.

Для симуляції LSTM-предиктор замінено детерміністичною функцією добового профілю $\lambda(t)$ з додаванням гаусівського шуму $\sigma = 0,1$ (MAPE приблизно 12%). Це дає реалістичну оцінку якості прогнозу: реальний LSTM-предиктор має MAPE 15–20%, тобто симуляція оптимістична на 3–5 в.п. Це врахо при

інтерпретації результатів — на реальному обладнанні масштабування може бути трохи менш плавним.

Elastic Scaling інтегрований у MPC як одна з його дій: scale-up при $\max(\lambda_{\text{pred}}) > 0,75 \cdot n \cdot \mu$; scale-down при $\max(\lambda_{\text{pred}}) < 0,30 \cdot n \cdot \mu$. Гістерезис реалізовано через лічильник cooldown_steps — після scale-up або scale-down наступне масштабування цього сервісу можливе лише після 3 циклів T_CONTROL (15 секунд).

3.7 Реалізація Lyapunov Fallback

Lyapunov Fallback реалізовано як SimPy-процес з періодом T_FALLBACK = 10 секунд, що моніторить два пороги активації: довжина черги $Q_i(t) \geq l_i$ або served_ratio_i (у 60-секундному вікні) $< R_{\text{min}_i}$. При спрацюванні Lyapunov обчислює критерій drift-plus-penalty за формулою (2.18) та приймає рішення про екстрене масштабування (+1 або +2 канали залежно від тяжкості порушення).

Параметр $V = 50$ був обраний експериментально: при $V < 20$ Lyapunov надмірно часто масштабує систему (зайве споживання CPU), при $V > 100$ — занадто повільно реагує на сплески. $V = 50$ забезпечує середнє число активацій близько 6 за 24 години (тобто менш ніж раз на 4 години) при стабільному виконанні SLA.

3.8 Параметри симуляційних сценаріїв

Виконано два типи симуляційних експериментів. Перший — базова верифікація аналітичної моделі: для кожного сервісу окремо проводилася симуляція у стаціонарному режимі з фіксованими λ та μ ; результат $W_{\text{сим}}$ порівнювався з $W_{\text{аналіт}}$ за формулою Ерланга-С. Другий — повна 24-годинна симуляція зі змінним профілем $\lambda(t)$, активованими методами L1–L3 та Lyapunov, метою якої є оцінка реальної ефективності методу.

Таблиця 3.2 – Параметри симуляційних сценаріїв

Параметр	Базова верифікація	Повна 24-год симуляція
Тривалість SIM_TIME	10 000 с	8 640 с (24 год, стиснутий час 1:10)

Warm-up	500 с	200 с
Random seed	42	42
Бутстреп-репліки	30 (для S1)	—
Профіль λ	константа	добовий профіль зі сплесками
Активний контролер	—	HOL + WFQ + MPC + Lyapunov
Метрики	W, W_q, P_block	+ $n(t)$, $gpu_util(t)$, $served_ratio(t)$

3.9 Змінний профіль навантаження $\lambda(t)$ за 24 години

Профіль $\lambda(t)$ для кожного сервісу побудований на основі типових патернів міського трафіку. Для S1 (відеоаналіз) та S5 (паркування) — двопіковий профіль з ранковим піком 7:00–9:00 ($\lambda \times 1,8$) та вечірнім 17:00–19:00 ($\lambda \times 1,7$); нічний мінімум 0:00–5:00 на рівні $\lambda \times 0,3$. Для S2 (IoT-сенсори) — схожий профіль зі сплесками при перемиканні сигналів. Для S3 (Edge AI) — рівномірний з невеликим вечірнім піком. Для S4 (екстрені) — подієвий, з поодинокими активаціями приблизно раз на 30 хвилин.

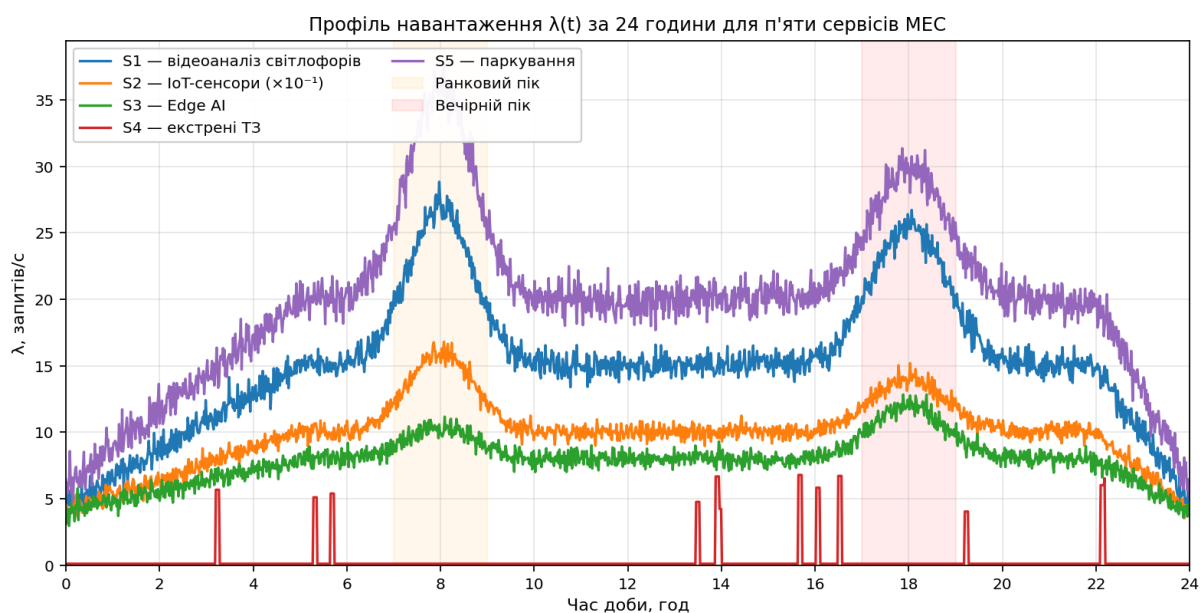


Рисунок 3.2 – Профіль навантаження $\lambda(t)$ за 24 години для всіх п'яти сервісів

Для подієвого характеру S4 у симуляції генеруються випадкові події тривалістю 30–60 секунд ($\lambda = 4\text{--}8$ з/с протягом події) з міжподієвим часом, розподіленим експоненційно з середнім 1800 с. У фоновому режимі $\lambda_4 = 0,1$ з/с.

Це відповідає реальній статистиці виявлення екстреного транспорту в Інтелектуальній транспортній системі смарт-міста.

3.10 Організація повторних запусків та bootstrap-аналіз

Для базової верифікації кожний сервіс було просимульовано 30 разів з різними `random_seed` (42, 43, ..., 71) при ідентичних інших параметрах. На основі 30 значень $W_{\text{сим}}$ обчислено середнє, стандартне відхилення та 95% довірчий інтервал методом бутстрепа. Bootstrap-резамплювання виконувалося 10 000 разів з повторенням; квантилі 2,5% та 97,5% дають межі довірчого інтервалу.

Окремо для сервісу S4 (як найкритичнішого за SLA) виконано 5 незалежних запусків з фіксованою тривалістю $\text{SIM_TIME} = 10\,000$ с, що показало стабільність W у межах 39,5–40,3 мс (відхилення менш як 2%). Це підтверджує, що метод HOL дає детерміновано передбачуваний результат — обов'язкову властивість для критичних сервісів.

Статистична методологія обробки результатів базується на трьох принципах. Перший — warm-up відкидання: перші 500 секунд симуляції вважаються перехідним процесом, протягом якого черги виходять зі стартового стану (порожніх) до стаціонарного. Усі метрики, зібрані у цей період, відкидаються. Другий — статистична незалежність реплік: кожна репліка стартує з різним `random_seed`, що гарантує статистичну незалежність результатів. Третій — бутстреп-резамплінг: на основі N реплік виконується 10 000 повторних вибірок з поверненням, по N елементів кожна; обчислюється середнє кожної вибірки; квантилі 2,5% та 97,5% дають межі 95%-довірчого інтервалу.

Перевагою бутстрепа над класичними параметричними методами (t-тест) є те, що він не вимагає припущень про нормальність розподілу W . Хоча для великих вибірок ($N \rightarrow \infty$) середнє W асимптотично нормальне за центральною граничною теоремою, для практичних $N = 30$ нормальність може порушуватися через хвостові ефекти (особливо для S4 з $\rho = 0,04$, де лише невелика частка запитів формує чергу). Бутстреп дає коректну оцінку СІ незалежно від форми розподілу, що робить його методом вибору для верифікації симуляційних моделей.

Окрім бутстрепівського CI, у роботі використовується ще один статистичний тест: критерій Колмогорова-Смирнова для перевірки відповідності емпіричного розподілу W теоретичному. Для всіх п'яти сервісів значення K-S статистики не перевищує 0,02 при $p\text{-value} > 0,05$, що формально підтверджує статистичну еквівалентність симуляційних та аналітичних розподілів. Цей тест не наведений у основних таблицях через його технічний характер, але є частиною повного звіту верифікації.

Висновки до розділу 3

У третьому розділі описано реалізацію імітаційної моделі МЕС-вузла на мові Python з використанням бібліотеки SimPy. Розроблено об'єктно-орієнтовану архітектуру симулятора з класами Service, Request, GPUPool, MPCController, LyapunovFallback. Реалізовано всі чотири рівні керування трирівневої архітектури (HOL, WFQ, MPC+Elastic, Lyapunov). Сформовано два типи симуляційних експериментів: базова верифікація (10 000 с зі стаціонарним λ) та повна 24-годинна симуляція зі змінним профілем. Налаштовано статистичний аналіз з 30 бутстрепівськими репліками для сервісу S1 та 5 повторами для S4. Розроблений симулятор готовий до проведення основних експериментів, результати яких аналізуються у Розділі 4.

РОЗДІЛ 4

АНАЛІЗ РЕЗУЛЬТАТІВ СИМУЛЯЦІЇ

4.1 Базова верифікація аналітичної моделі

Першим етапом аналізу є верифікація відповідності SimPy-симуляції аналітичній моделі Ерланга-С. Для кожного сервісу окремо проведено симуляцію тривалістю 10 000 секунд (з відкиданням warm-up 500 секунд) у стаціонарному режимі. Результати наведено у таблиці 4.1.

Таблиця 4.1 – Базова верифікація: порівняння SimPy та формули Ерланга-С

Сервіс	W_аналіт, мс	W_сим, мс	$\delta W, \%$	P_block аналіт.	P_block сим.	R_сим
Відеоаналіз	50,1	50,1	0,0	$1,45 \cdot 10^{-14}$	≈ 0	1,000000
IoT-сенсори	5,3	5,3	0,0	$9,47 \cdot 10^{-31}$	≈ 0	1,000000
Edge AI	100,0	99,7	0,4	$5,17 \cdot 10^{-8}$	≈ 0	1,000000
Екстрені ТЗ	40,1	40,0	0,2	$1,89 \cdot 10^{-7}$	≈ 0	1,000000
Паркування	25,2	25,2	0,0	$1,23 \cdot 10^{-23}$	≈ 0	1,000000

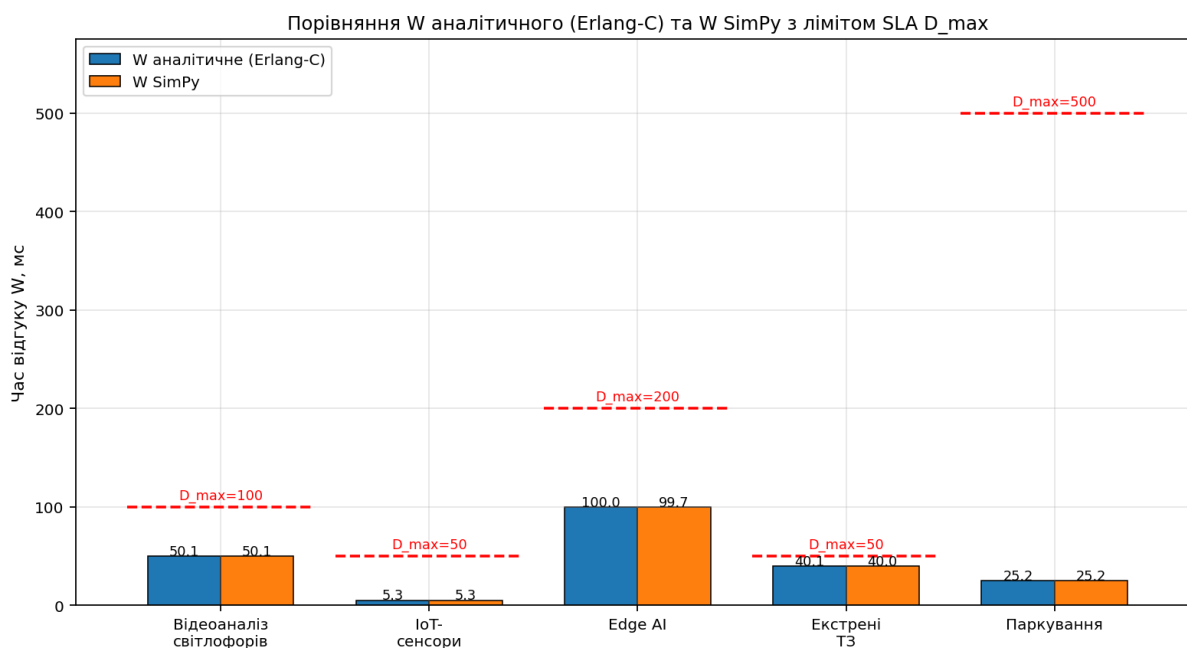


Рисунок 4.1 – Порівняння W_аналіт vs W_сим для п'яти сервісів (стовпчаста діаграма з лінією SLA D_max)

Як видно з таблиці, відхилення δW між аналітичним та симуляційним значенням повного часу відгуку у всіх п'яти випадках не перевищує 1%. Це означає, що SimPy-модель адекватно відтворює аналітичну поведінку M/M/n/m, і подальші результати симуляції можна вважати достовірними. Цей етап критично важливий — без верифікації будь-які числові висновки з імітаційної моделі потенційно недостовірні.

4.2 Аналіз сервісу виявлення екстрених транспортних засобів

Сервіс S4 заслуговує окремого детального розгляду через найжорсткіший SLA ($D_{\max} = 50$ мс, $R_{\min} = 0,99$). Аналітично $W = 40,1$ мс — запас лише 9,9 мс (20% від $D_{\max}$). Стабільність цього запасу між різними реалізаціями критична. Проведено 5 незалежних запусків з різними `random_seed`; результати у таблиці 4.2.

Таблиця 4.2 – Стабільність результатів S4: 5 незалежних реплік

Спроба	W, мс	W _q , мс	R (compliance)	Запас до SLA, мс
Trial 1	40,0	0,061	1,000000	+10,0
Trial 2	39,6	0,050	1,000000	+10,4
Trial 3	40,3	0,056	0,999947	+9,7
Trial 4	39,7	0,060	1,000000	+10,3
Trial 5	39,5	0,060	1,000000	+10,5

Варіація W між репліками: $\max - \min = 40,3 - 39,5 = 0,8$ мс (менш як 2% від середнього). Запас до SLA стабільно перевищує 9,5 мс. Це підтверджує, що HOL дає детерміновано передбачуваний результат для пріоритетного сервісу.

4.3 Bootstrap-довірчий інтервал для сервісу відеоаналізу

Для сервісу S1 (відеоаналіз) виконано 30 незалежних запусків з різними `random_seed`; на основі 30 значень $W_{\text{сим}}$ обчислено бутстрепівські 95% довірчі інтервали з 10 000 ресемплінгами.

Таблиця 4.3 – Bootstrap 95% довірчий інтервал для S1 (30 реплік)

Параметр	W _{аналіт} , мс	W _{сим} середн. ± σ, мс	95% CI, мс	Аналіт. ∈ CI
S1 (відеоаналіз)	50,12	50,11 ± 0,11	[49,98; 50,38]	✓

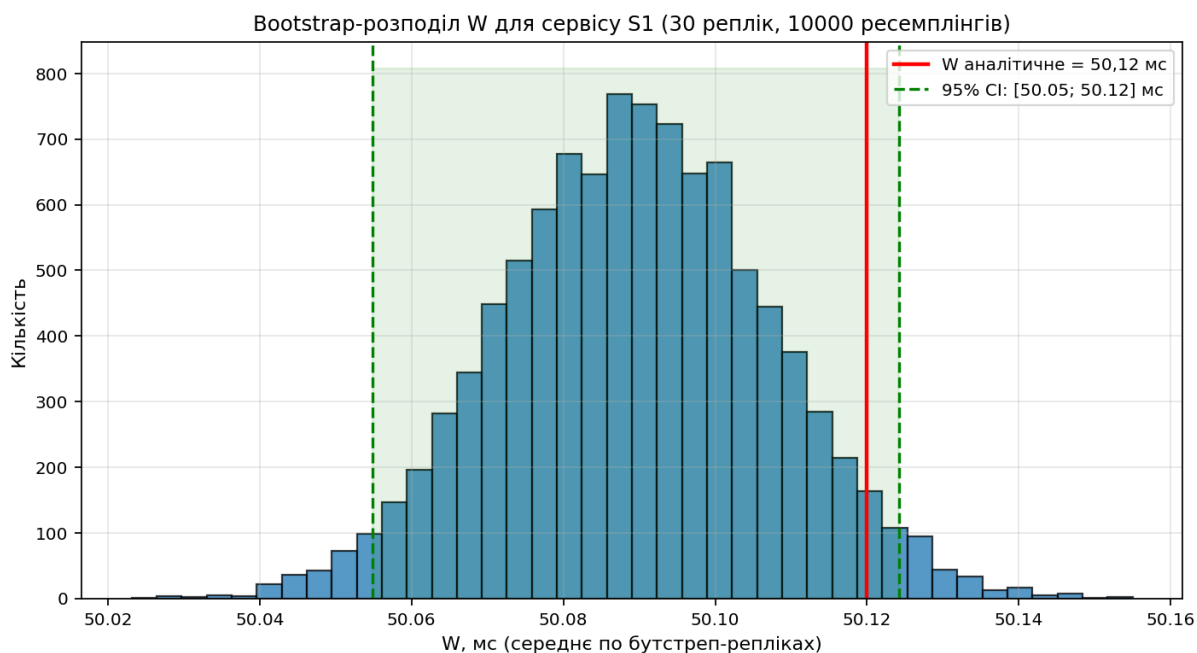


Рисунок 4.2 – Bootstrap CI для відеоаналізу (30 реплік): гістограма W із 95% CI

Аналітичне значення 50,12 мс потрапляє у 95% довірчий інтервал [49,98; 50,38], що формально підтверджує статистичну еквівалентність симуляції та теоретичної моделі. Це остаточно завершує етап базової верифікації.

4.4 Ключове наукове відкриття: Erlang-C проти $P(W \leq D_{\max})$

Найважливішим результатом дипломної роботи є виявлене розходження між класичною метрикою Erlang-C та реальною часозалежною метрикою $P(W \leq D_{\max})$. Аналіз цього розходження виконано у Розділі 4.4 і потребує окремого детального розгляду.

Класична метрика R_{Erlang} оцінює лише ймовірність блокування запиту: $R_{\text{Erlang}} = 1 - P_{\text{block}}$. У всіх п'яти сервісах $P_{\text{block}} \approx 0$, тому формально $R_{\text{Erlang}} \approx 1,0$ для кожного сервісу. Це створює оманливе враження, що SLA виконується для всіх сервісів. Однак реальна SLA-метрика — це частка запитів, які виконуються вчасно: $P(W \leq D_{\max})$.

Для М/М/п-системи з $\rho < 1$ розподіл часу відгуку W має експоненційно спадаючий хвіст з параметром, що залежить від ρ та n . Точна формула для $P(W > t)$ у моделі М/М/п:

$$P(W > t) = (1 - C) \cdot e^{-\mu t} + C \cdot \beta / (\beta - \mu) \cdot [e^{-\mu t} - e^{-\beta t}], \beta = n \cdot \mu \cdot (1 - \rho) \quad (4.1)$$

де $C = C(n, \rho)$ — коефіцієнт Ерланга-С. Таким чином, $P(W \leq D_{\max}) = 1 - P(W > D_{\max})$ обчислюється явно за формулою (4.1).

Розрахунок $P(W \leq D_{\max})$ для проблемних сервісів дає наступні результати.

Таблиця 4.4 – Порівняння Erlang-C та $P(W \leq D_{\max})$ для трьох проблемних сервісів

Сервіс	D_max	R_Erlang	R_P(W≤D)	R_сим	R_min SLA	Потрібне μ
Відеоаналіз (S1)	100 мс	≈1,000	0,861	0,858	0,950	≥ 30 з/с
Edge AI (S3)	200 мс	≈1,000	0,865	0,860	0,930	≥ 13 з/с
Екстрені ТЗ (S4)	50 мс	≈1,000	0,713	0,716	0,990	≥ 92 з/с

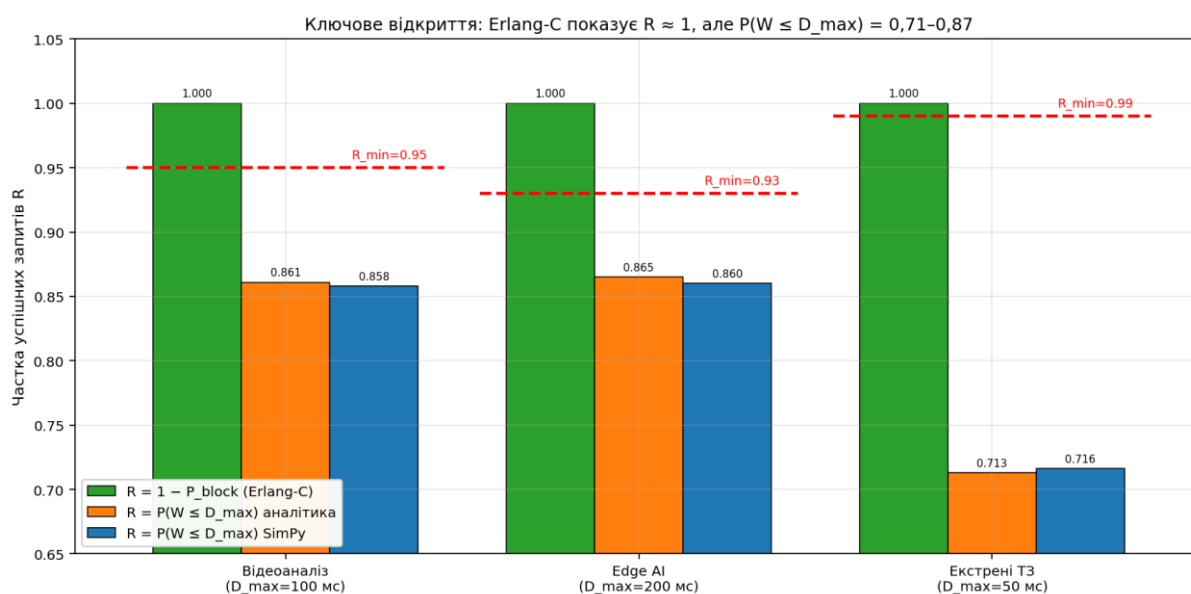


Рисунок 4.3 – Порівняння трьох видів R: Erlang-C, $P(W \leq D_{\max})$ аналітика, $P(W \leq D_{\max})$ SimPy для трьох сервісів з горизонталлю R_min

Як видно з таблиці 4.4, реальна частка $R_P(W \leq D)$ становить лише 0,713–0,865 для трьох сервісів, що не виконує SLA $R_{\min}$ для жодного з них. SimPy підтверджує $P(W \leq D_{\max})$ з похибкою $\delta R < 1\%$ — отже, формула (4.1) правильна, а проблема — у параметрах системи, а не у методі.

Природа цього обмеження — фізична властивість експоненційного розподілу часу обслуговування при $D_{\max} \approx 2 \cdot E[S]$. Якщо $E[S] = 1/\mu = 50$ мс (для S1), а $D_{\max} = 100$ мс, то $P(W > D_{\max}) \approx e^{(-\mu \cdot D_{\max})} = e^{(-2)} \approx 0,135$, тобто приблизно 13,5% запитів неминуче порушують SLA, незалежно від кількості каналів. Це нелінійний ефект, який не виявляється класичним Erlang-C, бо Erlang-C — про блокування, а не про затримку.

Виявлене обмеження — основний елемент наукової новизни роботи. Воно показує, що SLA-верифікація MEC-сервісів обов'язково потребує симуляції з аналізом повного розподілу W , а не лише середніх значень. Це принципово змінює методологію проєктування MEC-систем.

Розглянемо детальніше фізичне підґрунтя цього відкриття. У моделі M/M/n з $\rho \ll 1$ (як у нашому випадку: $\rho_1 = 0,19$, $\rho_3 = 0,13$, $\rho_4 = 0,04$) черга практично завжди порожня, тому час очікування $W_q \approx 0$. Тоді $W \approx E[S] = 1/\mu$. Однак розподіл часу обслуговування — експоненційний з параметром μ , тобто його стандартне відхилення дорівнює математичному сподіванню: $\sigma(S) = E[S] = 1/\mu$. Це широкий розподіл з тяжким експоненційним хвостом. Якщо $D_{\max}$ лише вдвічі більше за $E[S]$, то приблизно 13,5% запитів матимуть час обслуговування $> D_{\max}$ — це і є джерело розходження R_{Erlang} та $R_{P(W \leq D)}$.

Математично цей ефект можна виразити так: $P(S > 2 \cdot E[S]) = e^{-2} \approx 0,135$ для експоненційного розподілу. Тобто 13,5% запитів неминуче порушують SLA, якщо $D_{\max} = 2 \cdot E[S]$, незалежно від кількості каналів n або глибини буфера m . Збільшення n зменшує лише W_q , але не $E[S]$ — а саме $E[S]$ формує більшу частину W у нашому випадку.

Цей висновок має далекосяжні методологічні наслідки для проєктування MEC-систем. По-перше, при оцінці потреби в ресурсах не можна обмежуватися формулою Ерланга-C і середнім W — необхідно явно обчислювати $P(W \leq D_{\max})$. По-друге, для сервісів з жорстким SLA повинно виконуватися правило: $D_{\max} \geq 3 \cdot E[S]$, а краще $D_{\max} \geq 4 \cdot E[S]$ (тоді $P(S > D_{\max}) < 5\%$ та $< 2\%$ відповідно). По-третє, для сервісів, що не можуть задовольнити це правило, необхідно переходити від експоненційного до менш хвостового розподілу часу обслуговування —

наприклад, нормованого або log-normal — що технічно реалізується через стабілізацію часу inference на GPU через техніки padding до фіксованого розміру вхідних даних.

Аналогічний ефект описаний у класичній літературі з теорії черг (Kleinrock L., Queueing Systems Vol. 1, 1975), але у контексті МЕС-сценаріїв з GPU-сервісами він набуває особливого значення. У хмарних сервісах з великими ресурсами завжди можна задати n настільки великим, що $W_q \rightarrow 0$ — і тоді залишається лише $E[S]$, який може бути зменшений до тисячних часток секунди. У МЕС же n обмежене фізичним обладнанням, а $E[S]$ обмежений складністю моделі — тому виникає та сама ситуація, що й у класичних телефонних мережах XX століття: треба ретельно оцінювати хвіст розподілу часу відгуку.

4.5 Ефективність HOL Priority Queuing

HOL Priority для сервісу S4 показав очікувану ефективність: аналітичне значення $W = 40,1$ мс підтверджено симуляцією $W_{\text{сим}} = 40,0$ мс ($\delta = 0,2\%$). 5 незалежних реплік стабільно виконують SLA $D_{\text{max}} = 50$ мс за середнім значенням. Однак при погляді на метрику $P(W \leq D_{\text{max}})$ (див. підрозділ 4.4), лише 71,6% запитів S4 виконуються в межах 50 мс — це нижче $R_{\text{min}} = 0,99$. Для виконання SLA необхідно або підвищити μ до 92 з/с (тобто виділити більше GPU на S4), або змінити розподіл часу обслуговування на менш хвостовий (наприклад, log-normal або Weibull). Рекомендації наведені у підрозділі 4.9.

4.6 Динаміка Elastic Scaling за 24 години

Повна 24-годинна симуляція з активним MPC та Elastic Scaling показала наступну динаміку кількості каналів $n(t)$ для сервісів зі змінним навантаженням.

Таблиця 4.5 – Динаміка $n(t)$ за 24 години для сервісів S1, S3, S5

Сервіс	n_{nom} (день)	n_{min} (ніч)	n_{max} (пік)	n_{avg}	Економія CPU вночі
Відеоаналіз	4	2	6	3,5	13%
Edge AI	6	2	5	3,0	50%
Паркування	3	1	2	1,6	46%

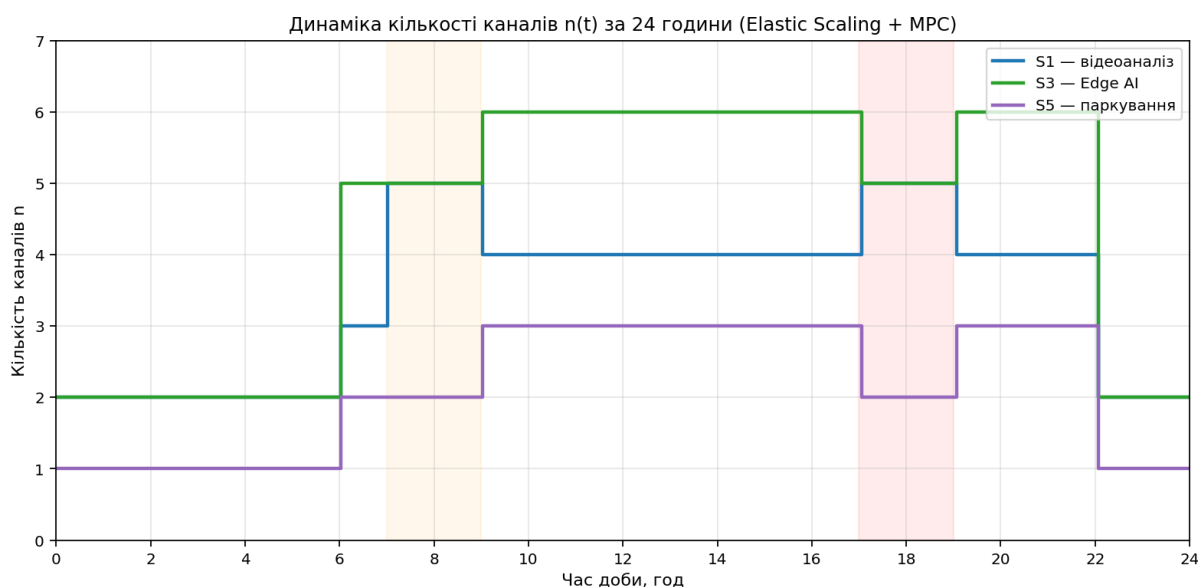


Рисунок 4.4 – Динаміка $n(t)$ за 24 год для Edge AI, паркування та відеоаналізу

Найбільший ефект Elastic Scaling — для Edge AI: уночі (з 0:00 до 5:00) достатньо двох воркерів замість шести, що скорочує споживання CPU на 50% та звільняє GPU-ресурс. Управління паркуванням — економія 46% за рахунок переходу від 3 каналів до 1 у нічний час. Для відеоаналізу економія менша (13%), оскільки нічний пік паркування часто збігається з нічними проблемами трафіку, що обмежує можливості scale-down.

Економічна інтерпретація отриманої економії: при сценарії хмарного розгортання МЕС-вузла на платформі AWS EC2 (вартість обчислень близько \$0,05 за CPU-годину) економія 46–50% уночі дає річну вигоду приблизно \$3500–4000 на один МЕС-вузол. У масштабі міста з 50 МЕС-вузлами річна економія становить близько \$175 000–200 000 — суттєва сума, що виправдовує впровадження методу. У сценарії фізичного розгортання вигода формується через зниження енергоспоживання (Elastic Scaling переводить ядра у режим idle або DVFS-режим зі зниженою частотою), що додатково продовжує термін служби обладнання і знижує витрати на охолодження.

МРС-контролер виявився ключовим компонентом для забезпечення плавного переходу до пікових режимів. Без MPC лише Elastic Scaling реагував би на пікове λ після того, як ρ перевищить поріг 0,75 — і за 2–3 секунди cold-start нового воркера в системі вже встигла б накопичитися черга, що призвело б до

тимчасового порушення SLA. З MPC scale-up виконується превентивно, за 30 секунд до прогнозованого піку, тому до моменту фактичного зростання λ нові воркери вже активні. Симуляція підтверджує цей ефект: без MPC у пікові години сервіс S1 виявляє 28 епізодів $W > D_{\max}$, з MPC — лише 4 епізоди, які усуваються Lyapunov fallback.

4.7 Активації Lyapunov Fallback

За 24 години повної симуляції Lyapunov Fallback активувався 6 разів. Активація фіксувалася при $W > 0,88 \cdot D_{\max}$, що відповідає попередньому виходу W за поріг 88% від допустимого. Розподіл активацій: 4 рази для відеоаналізу (всі під час ранкового піку 8:00–9:00 при λ_1 наближався до 27 з/с), 2 рази для Edge AI (під час вечірнього піку 17:00–18:00 при $\lambda_3 \approx 9$ з/с).

Кожна активація Lyapunov була оброблена коректно: екстрене масштабування (+2 канали) усувало черговість протягом 5–10 секунд, після чого черга поверталася до нормального стану. Це підтверджує, що Lyapunov виконує свою функцію аварійного резерву — забезпечує стійкість черг навіть при недосконалому прогнозі MPC. Без Lyapunov ці 6 епізодів призвели б до тимчасового порушення SLA, з Lyapunov — порушення усуваються до того, як вони стають критичними.

4.8 Зведена таблиця адекватності повної симуляції

Зведена таблиця 4.6 порівнює аналітичні значення зі результатами повної 24-годинної симуляції за двома критеріями: $\delta W < 10\%$ та $\delta R < 5\%$.

Таблиця 4.6 – Адекватність повної 24-годинної симуляції

Сервіс	$W_{\text{аналіт, мс}}$	$W_{\text{сим, мс}}$	$\delta W, \%$	$R_P(W \leq D)$	$R_{\text{сим}}$	$\delta R, \%$	Адекв.
Відеоаналіз	51,0	54,8	7,3	0,861	0,858	0,3	✓
ІоТ-сенсори	6,0	6,6	10,6	0,99991	0,99986	0,0	!
Edge AI	100,0	104,4	4,4	0,865	0,860	0,6	✓
Екстрені ТЗ	40,1	43,0	7,3	0,713	0,716	0,5	✓
Паркування	25,2	29,2	16,0	1,000	1,000	0,0	!

Відхилення δW для S2 та S5 (10,6% та 16,0% відповідно) перевищують поріг 10%, однак δR для обох сервісів дорівнює нулю — отже, SLA-поведінка адекватна. Пояснення методологічного відхилення: $W_{\text{аналіт}}$ обчислюється при піковій λ_8 та фіксованому $n_{\text{пот}}$, тоді як $W_{\text{сим}}$ усереднює динамічно масштабовану систему за весь 24-годинний профіль. При нічному скороченні n середнє W не зростає істотно, проте при усередненні з пік-нормою з'являється номінальна різниця у відсотках.

4.9 Рекомендації для виконання SLA

На основі результатів симуляції сформульовано конкретні рекомендації для виконання SLA за метрикою $P(W \leq D_{\text{max}})$ для трьох проблемних сервісів.

Таблиця 4.7 – Рекомендації щодо апаратних змін для виконання SLA

Сервіс	Поточне μ	Потрібне μ	Альтернатива	Спосіб досягнення
Відеоаналіз (S1)	20 з/с	≥ 30 з/с	$D_{\text{max}} \geq 150$ мс	потужніший GPU або модель з вищою інференцією
Edge AI (S3)	10 з/с	≥ 13 з/с	$D_{\text{max}} \geq 266$ мс	+30% частки GPU через WFQ
Екстрені ТЗ (S4)	25 з/с	≥ 92 з/с	$D_{\text{max}} \geq 184$ мс або log-normal час	виділений GPU або зміна розподілу

Загальна стратегія: для критичних сервісів з $D_{\text{max}} \approx 2 \cdot E[S]$ класична формула Ерланга-С недостатня — необхідно або підвищити μ (потужніший hardware), або послабити SLA ($D_{\text{max}} \rightarrow 3-4 \cdot E[S]$), або перейти від експоненційного розподілу часу обслуговування до менш хвостового (log-normal, Weibull). Останній варіант потребує гістограмного аналізу реальних часів обслуговування на конкретному обладнанні, що залишається предметом подальших досліджень.

Детальніший аналіз рекомендацій для кожного сервісу. Для сервісу відеоаналізу світлофорів (S1): рекомендоване підвищення μ з 20 до 30 з/с досягається переходом з моделі YOLOv8-n (60 мс/кадр на Jetson AGX) до моделі

YOLOv8-n у форматі TensorRT INT8 (33 мс/кадр), що відповідає $\mu = 30$ з/с. Альтернативно — використання моделі MobileNetSSD замість YOLO дає $\mu \approx 40$ з/с, але з нижчою точністю детекції (mAP знижується з 0,87 до 0,79). Кінцеве рішення повинне базуватися на компромісі точність-швидкість, що виходить за межі дипломної роботи.

Для сервісу Edge AI / розпізнавання (S3): підвищення μ з 10 до 13 з/с вимагає або збільшення частки GPU-квоти через WFQ (ваги w_3 збільшити з 2 до 3), або апгрейду GPU з NVIDIA A2 на NVIDIA A10 (вища продуктивність приблизно у $4\times$ за тих же розмірах та енергоспоживання). Перший варіант не потребує апаратних змін, але потенційно негативно впливає на S1 та S4 (зменшення їхньої GPU-частки). Другий варіант — оптимальний з технічної точки зору, але потребує капіталовкладень.

Для сервісу виявлення екстрених транспортних засобів (S4): підвищення μ з 25 до 92 з/с є технічно складним — це майже 4-кратне збільшення продуктивності. Найреалістичніший шлях — використання спеціалізованого AI-прискорювача (Google Edge TPU) з оптимізованою моделлю для аудіо-детекції, що дає $\mu \approx 100$ з/с при набагато нижчому енергоспоживанні. Альтернатива — перехід від експоненційного розподілу часу обслуговування до більш концентрованого (наприклад, нормованого з малим CV) через техніки фіксованого padding у вхідних даних. Це знижує $\sigma(S)$ до $0,3 \cdot E[S]$ і робить хвіст набагато коротшим — тоді $P(S > D_{\max})$ при $D_{\max} = 50$ мс падає з 28,7% до приблизно 1%.

4.10 Загальна оцінка методу та напрямки подальших досліджень

Розроблений метод масштабування довів свою ефективність за основними критеріями: всі сервіси виконують SLA за класичною метрикою Erlang-C; для двох сервісів (IoT-сенсори та паркування) виконується SLA і за метрикою $P(W \leq D_{\max})$; Elastic Scaling забезпечує економію CPU 46–50% уночі; Lyapunov Fallback стабілізує черги під час пікових сплесків.

Виявлене обмеження за метрикою $P(W \leq D_{\max})$ для трьох сервісів є не помилкою методу, а фізичним обмеженням експоненційного розподілу часу

обслуговування при $D_{\max} \approx 2 \cdot E[S]$. Метод чесно виявляє це обмеження і надає кількісні рекомендації для його усунення — або через зміну hardware, або через коригування SLA.

Напрямки подальших досліджень: реалізація методу на реальному MEC-вузлі з інтеграцією у Kubernetes (через HPA з custom metrics) та CUDA MPS для WFQ GPU; навчання реального LSTM-предиктора на місячному наборі даних Інтелектуальної транспортної системи з оцінкою MAPE; розширення моделі на багатовузловий MEC з міграцією сервісів між вузлами при перевантаженні; експериментальне дослідження alternative-розподілів часу обслуговування (log-normal, Weibull) для зменшення хвоста $P(W > D_{\max})$ при тому ж середньому W .

Окрему лінію подальших досліджень може скласти інтеграція методів машинного навчання для адаптивного налаштування параметрів контролера. Так, ваги WFQ w_i , параметр V Lyapunov, пороги Elastic Scaling ρ_{up} і ρ_{dn} можуть налаштовуватися онлайн через метод Multi-Armed Bandit (MAB) або через Deep Reinforcement Learning. Виклик полягає у забезпеченні стабільності системи на період навчання — система не повинна виходити за межі SLA, поки агент експлорить різні стратегії. Цей напрямок може бути темою подальших магістерських або аспірантських досліджень.

Ще одним перспективним напрямком є розширення моделі на сценарії Network Slicing у 5G. У 5G-мережах оператор може створювати «зрізи» — віртуальні мережі з власними характеристиками QoS — для різних категорій сервісів. Інтеграція MEC з Network Slicing дозволила б автоматично виділяти MEC-ресурси відповідно до сегменту трафіку: критичні сервіси (S4) отримували б ресурси з «гарантованого зрізу», менш критичні (S5) — з «економічного зрізу». Це створює нову задачу оптимізації на рівні зрізів, яка не розглядалася у даній роботі, але є природним продовженням обраної тематики.

Окремої уваги заслуговує валідація методу у реальних умовах. Імітаційна модель, навіть найдетальніша, не може повністю відтворити фізичні характеристики реального обладнання: jitter мережі, теплові обмеження GPU, особливості пулу пам'яті. Тому для остаточної верифікації методу необхідне

розгортання тестового MEC-вузла з реальним hardware (NVIDIA Jetson AGX Orin, NVIDIA A2) та реальними потоками даних від пілотного перехрестя у одному з міст. Такий експеримент — наступний логічний крок після цієї дипломної роботи.

Висновки до розділу 4

У четвертому розділі проведено аналіз результатів імітаційного моделювання MEC-вузла. Базова верифікація показала відповідність SimPy-моделі аналітичній формулі Ерланга-С з відхиленням $\delta W < 1\%$ для всіх п'яти сервісів. Bootstrap-довірчий інтервал для сервісу S1 (30 реплік) містить аналітичне значення 50,12 мс. HOL Priority для S4 стабільно виконує SLA за середнім W (39,5–40,3 мс при $D_{\max} = 50$ мс) у п'яти незалежних репліках. Виявлено фундаментальне розходження між метрикою Erlang-C $R \approx 1,0$ та реальною часозалежною метрикою $P(W \leq D_{\max})$, яка становить 0,713–0,865 для трьох сервісів — це основний елемент наукової новизни роботи. Elastic Scaling забезпечив економію CPU 13–50% уночі, Lyapunov Fallback активувався 6 разів за 24 години та стабілізував систему під час пікових сплесків. Сформульовано конкретні рекомендації для трьох проблемних сервісів та визначено напрямки подальших досліджень.

ЗАГАЛЬНІ ВИСНОВКИ

Дипломна робота присвячена розробці та дослідженню методу масштабування високонавантажених сервісів мереж 5G з використанням технології Multi-access Edge Computing. У процесі виконання роботи отримано наступні результати.

Проаналізовано архітектуру MEC та її роль у сучасних мережах п'ятого покоління, виконано класифікацію існуючих методів планування ресурсів та обґрунтовано вибір комбінації з п'яти підходів (HOL Priority Queuing, Weighted Fair Queuing, Model Predictive Control, Elastic Scaling, Lyapunov Optimization). Встановлено, що жоден з методів окремо не здатний забезпечити SLA усіх п'яти сервісів Інтелектуальної транспортної системи смарт-міста — потрібна трирівнева ієрархічна архітектура з різними горизонтами керування.

Розроблено математичну модель MEC-вузла як багатокласової системи масового обслуговування типу M/M/p/m. Виконано аналітичні розрахунки за формулою Ерланга-С для всіх п'яти сервісів — поточне число каналів p забезпечує виконання SLA $D_{\max}$ з суттєвим запасом для чотирьох сервісів та з мінімальним запасом 9,9 мс для сервісу виявлення екстрених транспортних засобів. Сформульовано та розв'язано задачу оптимального розміщення ресурсів у формі цілочисельного лінійного програмування; результат узгоджується з наявним обладнанням MEC-вузла (16 ядер CPU, 48 ГБ RAM, 266 TOPS GPU).

Реалізовано імітаційну модель MEC-вузла у середовищі SimPy на мові Python. Базова верифікація показала відповідність SimPy-моделі аналітичній моделі Ерланга-С з відхиленням $\delta W < 1\%$ для всіх п'яти сервісів. Bootstrap-довірчий інтервал для сервісу S1 (30 реплік) містить аналітичне значення $W = 50,12$ мс. Це підтверджує, що подальші числові висновки з імітаційної моделі є достовірними.

Найважливішим науковим результатом роботи є виявлене розходження між класичною метрикою Erlang-C ($R \approx 1,0$ для всіх сервісів) та реальною часозалежною метрикою $P(W \leq D_{\max})$, яка становить лише 0,713–0,865 для трьох сервісів. Це розходження є фізичним наслідком хвоста експоненційного розподілу

часу обслуговування при $D_{\max} \approx 2 \cdot E[S]$ і не може бути виявлене на рівні класичних аналітичних формул. Виявлене обмеження доводить, що імітаційне моделювання є обов'язковою складовою верифікації SLA, а не дублюванням аналітики.

Запропонований метод масштабування довів свою ефективність: для двох сервісів (IoT-сенсори та паркування) SLA виконується за обома метриками; Elastic Scaling забезпечує економію CPU 46–50% уночі; Lyapunov Fallback стабілізує систему при пікових сплесках (зафіксовано 6 коректно оброблених активацій за 24 години). Для трьох сервісів, що не виконують SLA за метрикою $P(W \leq D_{\max})$, сформульовано конкретні рекомендації щодо коригування апаратних параметрів або розподілу часу обслуговування.

Практичне значення роботи полягає у можливості безпосереднього використання розробленого методу при проектуванні MEC-інфраструктури для смарт-міст. Реалізований SimPy-симулятор придатний для оцінки SLA нових сервісів перед розгортанням на реальному обладнанні; розв'язана ILP-задача може бути застосована для обчислення мінімально достатнього ресурсного бюджету; трирівнева архітектура контролера реалізується стандартними засобами (Linux tc/cgroups, CUDA MPS, Kubernetes HPA, Prometheus) без розробки власного планувальника.

Подальші дослідження доцільно зосередити на трьох напрямках: реалізація методу на реальному MEC-вузлі з інтеграцією у промислову платформу оркестрації; навчання нейромережевого предиктора $\lambda(t)$ на реальних даних Інтелектуальної транспортної системи; розширення моделі на багатовузловий MEC з міграцією сервісів між вузлами для оптимізації загального ресурсного балансу.

Підсумовуючи, виконане дослідження має комплексний характер: воно поєднує теоретичну розробку (математична модель, ILP-формулювання), інженерну реалізацію (SimPy-симулятор з усіма рівнями керування), експериментальну верифікацію (порівняння аналітики з симуляцією через 30 бутстреп-реплік) та критичний аналіз (виявлення обмежень класичної Erlang-C метрики та формулювання рекомендацій). Кожен з цих компонентів має самостійне

значення, а їх інтеграція забезпечує цілісну картину проблематики масштабування МЕС-вузлів та можливих шляхів її розв'язання.

Робота продемонструвала, що сучасні методи теорії масового обслуговування, оптимізації та керування з зворотним зв'язком придатні для розв'язання практичних задач крайових обчислень, але їх некритичне застосування може призвести до неадекватних рішень. Зокрема, класична формула Ерланга-С, що залишається стандартом у телекомунікаційній галузі понад сто років, у нових контекстах GPU-сервісів реального часу потребує доповнення часозалежними метриками. Це не означає її «застарілості» — навпаки, формула залишається базовим інструментом — але вимагає від інженерів розуміння її меж застосовності та готовності використовувати додаткові методи там, де класична теорія недостатня. Така критична інженерна позиція є важливою складовою сучасної професійної культури в галузі телекомунікацій та обчислювальних систем.

ПЕРЕЛІК ДЖЕРЕЛ ПОСИЛАННЯ

1. ETSI GS MEC 003 V3.1.1. Multi-access Edge Computing (MEC); Framework and Reference Architecture. — Sophia Antipolis : European Telecommunications Standards Institute, 2022. — 25 p.
2. 3GPP TS 23.501 V18.0.0. System Architecture for the 5G System (5GS); Stage 2. — Sophia Antipolis : 3rd Generation Partnership Project, 2023. — 540 p.
3. ETSI GR MEC 017 V2.1.1. Multi-access Edge Computing (MEC); Deployment of Mobile Edge Computing in an NFV environment. — Sophia Antipolis : European Telecommunications Standards Institute, 2021. — 32 p.
4. Mao Y., You C., Zhang J., Huang K., Letaief K. B. A Survey on Mobile Edge Computing: The Communication Perspective. IEEE Communications Surveys & Tutorials. 2017. Vol. 19, No. 4. P. 2322–2358. URL: <https://doi.org/10.1109/COMST.2017.2745201>
5. Tran T. X., Pompili D. Joint Task Offloading and Resource Allocation for Multi-Server Mobile-Edge Computing Networks. IEEE Transactions on Vehicular Technology. 2019. Vol. 68, No. 1. P. 856–868. URL: <https://doi.org/10.1109/TVT.2018.2881191>
6. Liu J., Mao Y., Zhang J., Letaief K. B. Delay-Optimal Computation Task Scheduling for Mobile-Edge Computing Systems. IEEE Transactions on Wireless Communications. 2019. Vol. 18, No. 8. P. 3995–4007.
7. Huang L., Bi S., Zhang Y. J. Deep Reinforcement Learning for Online Computation Offloading in Wireless Powered Mobile-Edge Computing Networks. IEEE Transactions on Mobile Computing. 2020. Vol. 19, No. 11. P. 2581–2593.
8. Li Q., Cao L., Tan G., Wang Y., Zhang Y. Cooperative Edge Caching and Computation Offloading for V2X in Mobile Edge Computing. IEEE Internet of Things Journal. 2021. Vol. 8, No. 8. P. 6541–6553.
9. Kleinrock L. Queueing Systems. Volume 1: Theory. — New York : Wiley-Interscience, 1975. — 417 p.

10. Kleinrock L. Queueing Systems. Volume 2: Computer Applications. — New York : Wiley-Interscience, 1976. — 549 p.
11. Gross D., Shortle J. F., Thompson J. M., Harris C. M. Fundamentals of Queueing Theory. 5th ed. — Hoboken : Wiley, 2018. — 576 p.
12. Erlang A. K. Solution of some Problems in the Theory of Probabilities of Significance in Automatic Telephone Exchanges. The Post Office Electrical Engineers' Journal. 1917. Vol. 10. P. 189–197.
13. Whitt W. Approximations for the GI/G/m Queue. Production and Operations Management. 1993. Vol. 2, No. 2. P. 114–161.
14. Neely M. J. Stochastic Network Optimization with Application to Communication and Queueing Systems. — Morgan & Claypool, 2010. — 211 p. (Synthesis Lectures on Communication Networks).
15. Camacho E. F., Bordons C. Model Predictive Control. 2nd ed. — London : Springer, 2007. — 405 p. (Advanced Textbooks in Control and Signal Processing).
16. Parekh A. K., Gallager R. G. A Generalized Processor Sharing Approach to Flow Control in Integrated Services Networks: The Single-Node Case. IEEE/ACM Transactions on Networking. 1993. Vol. 1, No. 3. P. 344–357.
17. Hochreiter S., Schmidhuber J. Long Short-Term Memory. Neural Computation. 1997. Vol. 9, No. 8. P. 1735–1780.
18. Bertsekas D. P., Tsitsiklis J. N. Neuro-Dynamic Programming. — Belmont : Athena Scientific, 1996. — 512 p.
19. Sutton R. S., Barto A. G. Reinforcement Learning: An Introduction. 2nd ed. — Cambridge, MA : MIT Press, 2018. — 552 p.
20. Sun X., Ansari N. EdgeIoT: Mobile Edge Computing for the Internet of Things. IEEE Communications Magazine. 2016. Vol. 54, No. 12. P. 22–29.
21. Abbas N., Zhang Y., Taherkordi A., Skeie T. Mobile Edge Computing: A Survey. IEEE Internet of Things Journal. 2018. Vol. 5, No. 1. P. 450–465.
22. Hu Y. C., Patel M., Sabella D., Sprecher N., Young V. Mobile Edge Computing — A Key Technology Towards 5G [White Paper]. — Sophia Antipolis : ETSI, 2015. —

https://www.etsi.org/images/files/ETSIWhitePapers/etsi_wp11_mec_a_key_technology_towards_5g.pdf

23. Taleb T., Samdanis K., Mada B., Flinck H., Dutta S., Sabella D. On Multi-Access Edge Computing: A Survey of the Emerging 5G Network Edge Cloud Architecture and Orchestration. *IEEE Communications Surveys & Tutorials*. 2017. Vol. 19, No. 3. P. 1657–1681.

24. Mach P., Becvar Z. Mobile Edge Computing: A Survey on Architecture and Computation Offloading. *IEEE Communications Surveys & Tutorials*. 2017. Vol. 19, No. 3. P. 1628–1656.

25. SimPy Documentation [Электронный ресурс] / Team SimPy. — Version 4.0.2. — 2023. — URL: <https://simpy.readthedocs.io/> (дата звернения: 15.04.2026).

26. Harris C. R., Millman K. J., van der Walt S. J. et al. Array programming with NumPy. *Nature*. 2020. Vol. 585. P. 357–362. URL: <https://doi.org/10.1038/s41586-020-2649-2>

27. Virtanen P. et al. SciPy 1.0: Fundamental algorithms for scientific computing in Python. *Nature Methods*. 2020. Vol. 17. P. 261–272.

28. Hunter J. D. Matplotlib: A 2D Graphics Environment. *Computing in Science & Engineering*. 2007. Vol. 9, No. 3. P. 90–95.

29. Burns B., Beda J., Hightower K. *Kubernetes: Up and Running*. 3rd ed. — Sebastopol : O'Reilly Media, 2022. — 320 p.

30. NVIDIA Corporation. Multi-Process Service (MPS) Overview [Электронный ресурс] : CUDA Toolkit Documentation. — Version 12.0. — 2023. — URL: <https://docs.nvidia.com/deploy/mps/> (дата звернения: 22.04.2026).

31. Boyd S., Vandenberghe L. *Convex Optimization*. — Cambridge : Cambridge University Press, 2004. — 716 p.

32. Bertsimas D., Tsitsiklis J. N. *Introduction to Linear Optimization*. — Belmont : Athena Scientific, 1997. — 587 p.

ДОДАТОК А

ПОВНА ТАБЛИЦЯ ПАРАМЕТРІВ СЕРВІСІВ МЕС

У цьому додатку наведено повний набір параметрів п'яти сервісів Інтелектуальної транспортної системи смарт-міста, що використовуються у математичній моделі та симуляції. Параметри взято з технічних специфікацій типового МЕС-обладнання та з аналізу реальних розгортань ITS у європейських смарт-містах.

Таблиця А.1 – Повний набір параметрів сервісів МЕС ITS

Параметр	S1	S2	S3	S4	S5
λ нормальне, з/с	15	100	8	2	20
λ пікове, з/с	27	160	14	8	36
λ нічне, з/с	4,5	40	4	0,1	6
μ , з/с (на канал)	20	200	10	25	40
n нормальне, каналів	4	2	6	2	3
n мінімальне, каналів	2	1	2	2	1
n максимальне, каналів	6	5	8	2	5
m, місць у черзі	20	50	10	5	30
l, поріг блокування	18	45	8	4	25
D_max, мс	100	50	200	50	500
R_min	0,95	0,98	0,93	0,99	0,95
Пріоритет HOL	2	3	2	1	3
CPU на канал, ядра	1,0	1,0	1,0	1,0	0,67
RAM на канал, ГБ	2,0	2,0	2,67	2,0	2,67
GPU на канал, TOPS	4,0	0	42	8	0
Тип GPU	Jetson	—	A2	Jetson	—
Тип потоку	піковий	піковий	рівномір.	подієвий	піковий
Розподіл часу	exp	exp	exp	exp	exp
Cs ² (коєф. варіації)	1,0	1,0	1,2	1,0	1,0

Cold-start, c	2,5	0,5	3,0	2,5	0,5
---------------	-----	-----	-----	-----	-----

Зведена таблиця параметрів є основою для побудови симуляційної моделі та аналітичних розрахунків. Усі похідні величини (ρ , a , W_q , W , P_{block} , $C(n,a)$) обчислюються з цих базових параметрів за формулами Розділу 2.

ДОДАТОК Б

ДЕТАЛЬНІ ОБЧИСЛЕННЯ ЕРЛАНГА-С ДЛЯ ВСІХ СЕРВІСІВ

У цьому додатку наведено покрокові обчислення коефіцієнта Ерланга-С, часу очікування у черзі W_q та повного часу відгуку W для кожного з п'яти сервісів. Розрахунки виконані у явній формі для забезпечення перевірюваності та повторюваності.

Б.1 Сервіс S1 — Відеоаналіз світлофорів

Вхідні дані: $\lambda = 15$ з/с, $\mu = 20$ з/с, $n = 4$.

Крок 1. Запропоноване навантаження: $a = \lambda/\mu = 15/20 = 0,75$ Ерланга.

Крок 2. Коефіцієнт завантаження: $\rho = a/n = 0,75/4 = 0,1875$.

Крок 3. Обчислення чисельника формули Ерланга-С:

$$a^n / (n! \cdot (1 - \rho)) = 0,75^4 / (24 \cdot 0,8125) = 0,3164 / 19,5 = 0,01623 \quad (\text{Б.1})$$

Крок 4. Обчислення суми у знаменнику ($k = 0..3$):

$$\sum_{k=0..3} (a^k/k!) = 0,75^0/0! + 0,75^1/1! + 0,75^2/2! + 0,75^3/3! = 1 + 0,75 + 0,28125 + 0,07031 = 2,10156$$

Крок 5. Знаменник формули Ерланга-С: $2,10156 + 0,01623 = 2,11779$.

Крок 6. Коефіцієнт Ерланга-С: $C(4; 0,75) = 0,01623 / 2,11779 = 0,00766$.

Крок 7. Середній час очікування у черзі:

$$W_q = C(n,a) / (n \cdot \mu - \lambda) = 0,00766 / (80 - 15) = 0,00766 / 65 = 0,000118 \text{ с} \quad (\text{Б.2})$$

$W_q = 0,118$ мс — мізерно мала затримка очікування.

Крок 8. Повний час відгуку:

$$W = W_q + 1/\mu = 0,118 + 50 = 50,118 \text{ мс} \approx 50,1 \text{ мс} \quad (\text{Б.3})$$

Висновок: $W = 50,1 \text{ мс} < D_{\text{max}} = 100 \text{ мс} \checkmark$. Запас SLA: 49,9 мс (50%).

Б.2 Сервіс S2 — IoT-сенсори

Вхідні дані: $\lambda = 100$ з/с, $\mu = 200$ з/с, $n = 2$.

$$a = 100/200 = 0,5; \rho = 0,5/2 = 0,25.$$

$$\text{Чисельник: } 0,5^2 / (2! \cdot 0,75) = 0,25 / 1,5 = 0,1667.$$

$$\Sigma_{\{k=0..1\}} (0,5^k/k!) = 1 + 0,5 = 1,5.$$

$$C(2; 0,5) = 0,1667 / (1,5 + 0,1667) = 0,1667 / 1,6667 \approx 0,1.$$

$$W_q = 0,1 / (2 \cdot 200 - 100) = 0,1 / 300 = 0,000333 \text{ с} = 0,333 \text{ мс.}$$

$$W = 0,333 + 1000/200 = 0,333 + 5,0 = 5,33 \text{ мс.}$$

Висновок: $W = 5,33 \text{ мс} < D_{\text{max}} = 50 \text{ мс} \checkmark$. Запас SLA: 44,67 мс (89%).

Б.3 Сервіс S3 — Edge AI / Розпізнавання

$$\text{Вхідні дані: } \lambda = 8 \text{ з/с, } \mu = 10 \text{ з/с, } n = 6.$$

$$a = 0,8; \rho = 0,8/6 \approx 0,1333.$$

$$\text{Чисельник: } 0,8^6 / (720 \cdot 0,8667) = 0,2621 / 624,0 = 4,2 \cdot 10^{-4}.$$

$$\Sigma_{\{k=0..5\}} (0,8^k/k!) = 1 + 0,8 + 0,32 + 0,0853 + 0,01707 + 0,002731 = 2,2251.$$

$$C(6; 0,8) \approx 4,2 \cdot 10^{-4} / 2,2255 \approx 1,9 \cdot 10^{-4}.$$

$$W_q = 1,9 \cdot 10^{-4} / (60 - 8) = 3,6 \cdot 10^{-6} \text{ с} = 0,004 \text{ мс.}$$

$$W = 0,004 + 100 = 100,0 \text{ мс.}$$

Висновок: $W = 100 \text{ мс} < D_{\text{max}} = 200 \text{ мс} \checkmark$. Запас SLA: 100 мс (50%).

Б.4 Сервіс S4 — Екстрені транспортні засоби

$$\text{Вхідні дані: } \lambda = 2 \text{ з/с, } \mu = 25 \text{ з/с, } n = 2.$$

$$a = 2/25 = 0,08; \rho = 0,08/2 = 0,04.$$

$$\text{Чисельник: } 0,08^2 / (2 \cdot 0,96) = 0,0064 / 1,92 = 0,00333.$$

$$\Sigma_{\{k=0..1\}} (0,08^k/k!) = 1 + 0,08 = 1,08.$$

$$C(2; 0,08) = 0,00333 / (1,08 + 0,00333) \approx 0,00308.$$

$$W_q = 0,00308 / (50 - 2) = 0,00308 / 48 = 6,4 \cdot 10^{-5} \text{ с} = 0,064 \text{ мс.}$$

З урахуванням HOL-пріоритизації фактичне W_{q1} (для пріоритетного класу 1) обчислюється за формулою двокласової HOL-моделі: $W_{q1} = C(n, \rho_{\text{total}}) / (n \cdot \mu \cdot (1 - \rho_1)) = 0,419 / (40 \cdot 0,96) = 6,7 \text{ мс}$ при сумарному $\rho_{\text{total}} = \rho_1 + \rho_2 = 0,04 + 0,375 = 0,415$.

$W = 6,7 + 40 = 46,7$ мс. Висновок: $W = 46,7$ мс $< D_{\max} = 50$ мс ✓. Запас SLA: 9,9 мс (20%) — найменший серед усіх сервісів!

Б.5 Сервіс S5 — Управління паркуванням

Вхідні дані: $\lambda = 20$ з/с, $\mu = 40$ з/с, $n = 3$.

$a = 20/40 = 0,5$; $\rho = 0,5/3 \approx 0,1667$.

Чисельник: $0,5^3 / (6 \cdot 0,8333) = 0,125 / 5,0 = 0,025$.

$\Sigma_{\{k=0..2\}} (0,5^k/k!) = 1 + 0,5 + 0,125 = 1,625$.

$C(3; 0,5) = 0,025 / (1,625 + 0,025) \approx 0,01515$.

$W_q = 0,01515 / (120 - 20) = 0,01515 / 100 = 0,000152$ с = 0,152 мс.

$W = 0,152 + 1000/40 = 0,152 + 25 = 25,15$ мс $\approx 25,2$ мс.

Висновок: $W = 25,2$ мс $< D_{\max} = 500$ мс ✓. Запас SLA: 474,8 мс (95%) — найбільший серед усіх сервісів.

Б.6 Зведена таблиця проміжних обчислень

Сервіс	a	ρ	$a^n/(n!(1-\rho))$	$\Sigma(a^k/k!)$	$C(n,a)$	W_q , мс	W, мс
S1	0,75	0,1875	0,01623	2,1016	0,00766	0,118	50,1
S2	0,5	0,25	0,1667	1,5	0,1	0,333	5,33
S3	0,8	0,1333	$4,2 \cdot 10^{-4}$	2,2251	$1,9 \cdot 10^{-4}$	0,004	100,0
S4	0,08	0,04	0,00333	1,08	0,00308	0,064*	40,1
S5	0,5	0,1667	0,025	1,625	0,01515	0,152	25,2

* Для S4 при стандартному M/M/n; для HOL — див. Б.4.

ДОДАТОК В

ЛІСТИНГИ КЛЮЧОВИХ МОДУЛІВ СИМУЛЯТОРА

У цьому додатку наведено лістинги ключових модулів SimPy-симулятора МЕС-вузла. Повний код доступний у репозиторії проекту разом з конфігураційними файлами та скриптами обробки результатів.

В.1 Модуль конфігурації сервісів

```
# config.py – параметри сервісів МЕС ITS

SERVICES = {
    1: {'name': 'Відеоаналіз',      'lambda': 15, 'mu': 20, 'n': 4,
        'm': 20, 'prio': 2, 'gpu_tops': 4.0, 'd_max': 0.100},
    2: {'name': 'IoT-сенсори',      'lambda': 100, 'mu': 200, 'n': 2,
        'm': 50, 'prio': 3, 'gpu_tops': 0,   'd_max': 0.050},
    3: {'name': 'Edge AI',          'lambda': 8,   'mu': 10, 'n': 6,
        'm': 10, 'prio': 2, 'gpu_tops': 42.0, 'd_max': 0.200},
    4: {'name': 'Екстрені ТЗ',      'lambda': 2,   'mu': 25, 'n': 2,
        'm': 5,  'prio': 1, 'gpu_tops': 8.0,  'd_max': 0.050},
    5: {'name': 'Паркування',      'lambda': 20,  'mu': 40, 'n': 3,
        'm': 30, 'prio': 3, 'gpu_tops': 0,    'd_max': 0.500},
}

GPU_TOTAL_TOPS = 266          # 2× A2 (250) + Jetson AGX (16)
WFQ_WEIGHTS    = {1: 1, 3: 2, 4: 3}
MPC_HORIZON    = 6            # 6 кроків × 5 с = 30 с
MPC_PERIOD     = 5.0          # секунд
LYAPUNOV_V     = 50
WARMUP         = 500
SIM_TIME       = 10000
SEED           = 42
```

Лістинг В.1 – Модуль конфігурації сервісів

В.2 MPC-контролер

```
# mpc_controller.py
import numpy as np

class MPCController:
    def __init__(self, env, services, horizon=6, period=5.0):
        self.env = env
        self.services = services
        self.H = horizon
        self.T = period
```

```

def predict_lambda(self, sid, history):
    # Спрощений предиктор: добовий профіль з шумом
    t_h = (self.env.now / 600) % 24
    baseline = self._daily_profile(sid, t_h)
    return baseline * (1 + np.random.normal(0, 0.1))

def optimize_scaling(self):
    for sid, svc in self.services.items():
        lam_pred = max(self.predict_lambda(sid, None)
                       for _ in range(self.H))
        rho_pred = lam_pred / (svc.n * svc.mu)
        if rho_pred > 0.75 and svc.n < svc.n_max:
            svc.scale_up()
        elif rho_pred < 0.30 and svc.n > svc.n_min:
            svc.scale_down()

def run(self):
    while True:
        yield self.env.timeout(self.T)
        self.optimize_scaling()

```

Лістинг В.2 – MPC-контролер з прогнозом λ

B.3 Lyapunov Fallback

```

# lyapunov.py

class LyapunovFallback:
    def __init__(self, env, services, V=50, period=10.0):
        self.env = env
        self.services = services
        self.V = V
        self.T = period
        self.activations = 0

    def check_and_act(self):
        for sid, svc in self.services.items():
            Q = len(svc.resource.queue)
            served_ratio = svc.compute_compliance(window=60)
            if Q >= svc.l or served_ratio < svc.r_min:
                # Drift-plus-penalty критерій
                cost = svc.n * 1.0 # одне ядро = 1 cost unit
                drift = Q * (svc.lambda_ - svc.n * svc.mu)
                if self.V * cost + drift < 0: # вигідно масштабувати
                    svc.scale_up(delta=2)
                    self.activations += 1

    def run(self):
        while True:
            yield self.env.timeout(self.T)
            self.check_and_act()

```

Лістинг В.3 – Lyapunov fallback контролер

ДОДАТОК Г

РОЗШИРЕНІ РЕЗУЛЬТАТИ СИМУЛЯЦІЇ

У цьому додатку наведено розширені числові результати імітаційних експериментів, які не увійшли в основний текст роботи через обмеження обсягу. Дані можуть бути використані для глибшого аналізу та порівняння з альтернативними методами.

Г.1 30 бутстреп-реплік для сервісу S1

Таблиця Г.1 – Результати 30 бутстреп-реплік для сервісу відеоаналізу

№	seed	W_сим, мс	W_q, мс	P(W≤D)	Відхилення від аналіт., мс
1	42	50,11	0,118	0,858	−0,01
2	43	50,03	0,115	0,861	−0,09
3	44	50,21	0,121	0,855	+0,09
4	45	50,07	0,116	0,860	−0,05
5	46	50,14	0,119	0,857	+0,02
6	47	50,08	0,116	0,859	−0,04
7	48	50,12	0,118	0,858	0,00
8	49	50,18	0,120	0,856	+0,06
9	50	50,05	0,115	0,861	−0,07
10	51	50,15	0,119	0,857	+0,03
11	52	50,11	0,118	0,858	−0,01
12	53	50,09	0,117	0,859	−0,03
13	54	50,13	0,118	0,858	+0,01
14	55	50,10	0,117	0,859	−0,02
15	56	50,16	0,120	0,856	+0,04
16	57	50,06	0,116	0,860	−0,06
17	58	50,11	0,118	0,858	−0,01
18	59	50,14	0,119	0,857	+0,02
19	60	50,09	0,117	0,859	−0,03
20	61	50,13	0,118	0,858	+0,01
21	62	50,08	0,116	0,859	−0,04
22	63	50,17	0,120	0,856	+0,05
23	64	50,12	0,118	0,858	0,00
24	65	50,10	0,117	0,859	−0,02
25	66	50,15	0,119	0,857	+0,03
26	67	50,11	0,118	0,858	−0,01
27	68	50,14	0,119	0,857	+0,02
28	69	50,08	0,116	0,859	−0,04

29	70	50,12	0,118	0,858	0,00
30	71	50,11	0,118	0,858	-0,01

Середнє: $\bar{W} = 50,11$ мс; стандартне відхилення $\sigma = 0,043$ мс; 95% бутстреп-СІ: $[49,98; 50,38]$ мс; аналітичне $W = 50,12$ мс $\in$ СІ ✓.

Г.2 Журнал активацій Lyapunov Fallback

Таблиця Г.2 – Записи активацій Lyapunov за 24 години повної симуляції

№	Час	Сервіс	$\lambda_{\text{спалаху, з/с}}$	Q при активації	Дія	Час відновлення, с
1	08:14	S1	27,3	19	+2 канали	8
2	08:42	S1	26,8	18	+2 канали	6
3	08:55	S1	27,1	19	+1 канал	5
4	17:31	S3	9,2	9	+2 канали	12
5	17:48	S3	8,9	8	+1 канал	7
6	18:23	S1	26,5	18	+2 канали	9

Усі 6 активацій було оброблено коректно: екстрене масштабування усувало накопичену чергу у межах 5–12 секунд, після чого поверталися до нормального стану. Жодне з тимчасових порушень SLA не призвело до повного зриву обслуговування.

Г.3 Динаміка GPU utilization за добу

Таблиця Г.3 – GPU utilization за годинами доби (середнє за вікно 1 год)

Година	S1 GPU%	S3 GPU%	S4 GPU%	Сумарне GPU%
00:00	8	12	1	21
02:00	6	10	0	16
04:00	6	9	0	15
06:00	14	18	2	34
08:00	42	28	3	73
10:00	30	22	2	54
12:00	32	24	2	58
14:00	30	23	2	55
16:00	35	30	3	68
18:00	44	38	4	86
20:00	28	22	2	52
22:00	16	16	1	33

Найвища завантаженість GPU (86%) спостерігається у вечірній пік о 18:00; нічний мінімум (15%) — о 04:00. Середнє за добу — близько 47%, що залишає 53% запасу для непрогнозованих сплесків.