

***БАКАЛАВРСЬКА
РОБОТА***

БР.ІІ –22.00.00.000 ІІЗ

Група ІІ-22-1

Польнюк Юрій

2026

Івано-Франківський національний технічний університет нафти і газу

Факультет інформаційних технологій

Кафедра інженерії програмного забезпечення

Польнюк Юрій Миколайович

(прізвище, ім'я, по батькові)

УДК 004.942

(індекс)

БАКАЛАВРСЬКА РОБОТА

Використання контейнеризації для масштабування додатків

(назва роботи)

Інженерія програмного забезпечення

(назва освітньої програми)

121 - Інженерія програмного забезпечення

(шифр і назва спеціальності)

Робота містить результати власних досліджень, використання ідей, результатів і текстів інших авторів мають посилання на відповідне джерело:

Здобувач освітнього ступеня Ю. М. Польнюк
(підпис, ініціали та прізвище здобувача)

Науковий керівник Бандура Вікторія Валеріївна, к.т.н., доцент
(підпис, прізвище, ім'я, по батькові, науковий ступінь, вчене звання керівника)

Допущено до захисту

Завідувач кафедри

доц. Бандура В.В.
(посада) (підпис) (дата) (ініціали та прізвище)

Івано-Франківський національний технічний університет нафти і газу

Інститут, факультет інформаційних технологій

Кафедра інженерії програмного забезпечення

Ступінь вищої освіти бакалавр

Спеціальність 121 – Інженерія програмного забезпечення

ЗАТВЕРДЖУЮ:

Зав. кафедрою

ІПЗ

доцент.

В.В. Бандура

“ ” 2026 р.

ЗАВДАННЯ НА БАКАЛАВРСЬКУ РОБОТУ СТУДЕНТОВІ

Польнюк Юрій Миколайович

(прізвище, ім'я, по-батькові)

1. Тема проекту (роботи) "Використання контейнеризації для масштабування додатків"

керівник проекту (роботи) Бандура Вікторія Валеріївна, к.т.н., доцент

затверджені наказом вищого навчального закладу від “ 21 ” квітня 2026 р. № 179/7

2. Строк подання студентом проекту (роботи) 10 червня 2026 р.

3. Вихідні дані до проекту (роботи) Результати і матеріали отримані під час проходження переддипломної практики

4. Зміст розрахунково - пояснювальної записки(перелік питань, які потрібно розробити)

1. Аналіз вимог до програмного забезпечення

2. Аналіз вимог і постановка задачі

3. Проектування системи

4. Реалізація проекту

5. Тестування та впровадження

5. Перелік графічного матеріалу (з точним зазначенням обов'язкових креслень)

1 Програма для віртуалізації контейнерів на базі Docker (рис. 1.2, ст. 14)

2 Архітектура Docker (рис. 2.1, ст. 21)

3 Контейнери Linux або контейнери LXC (рис. 2.4, ст. 35)

4 Загальний огляд серверної архітектури eShare (рис.3.1, ст. 46)

5 Діаграма класів інтерфейсу IMessageBus та його реалізацій KafkaMessageBus і MessageBus (рис. 3.5, ст. 55)

6. Консультанти розділів проекту (роботи)

Розділ	Консультант	Підпис, дата	
		Завдання видав	Завдання прийняв

2. Дата видачі завдання 2026 р.

Керівник Бандура В.В.
(підпис)

Завдання прийняв до виконання Польнюк Ю.М.
(підпис)

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів дипломного проекту (роботи)	Строк виконання етапів проекту	Примітка
1	Визначення та обґрунтування теми роботи	17.01.2026	виконано
2	Огляд існуючих концепцій, рішень та сервісів в даній області	21.02.2026	виконано
3	Побудова моделі або алгоритму власного рішення	01.03.2026	виконано
4	Документування реалізації власного оригінального рішення вибраними засобами	21.04.2026	виконано
5	Оформлення пояснювальної записки бакалаврської роботи	02.05.2026	виконано

Студент Польнюк Ю.М
(підпис)

Керівник роботи Бандура В.В
(підпис)

АНОТАЦІЯ

Бакалаврська робота містить 68 сторінок, 10 рисунків, 1 таблиця, список використаних джерел із 40 найменувань.

Метою роботи є дослідження підходів до використання контейнеризації для масштабування додатків та аналіз сучасних інструментів і технологій, що забезпечують ефективне керування та масштабування програмних систем.

Об'єкт дослідження: процеси контейнеризації додатків у сучасних програмних системах.

Предмет дослідження: методи, моделі та інструменти контейнеризації (Docker, Docker Compose, контейнерні архітектури) для забезпечення масштабованості, гнучкості та ефективності розгортання додатків.

Результати дослідження: проведено аналіз принципів контейнеризації, визначено її переваги над традиційною віртуалізацією, досліджено інструменти створення та керування контейнерами.

У першому розділі розглянуто теоретичні основи контейнеризації, її роль у сучасній розробці програмного забезпечення, а також порівняння контейнерів і віртуальних машин та визначено сфери їх застосування.

У другому розділі досліджено основні інструменти контейнеризації, зокрема Docker, розглянуто практичні сценарії використання контейнерів.

У третьому розділі описано процес розробки та впровадження контейнеризованого рішення, використання Docker Compose для керування сервісами.

Висновок: використання контейнеризації дозволяє значно підвищити масштабованість, гнучкість та ефективність розгортання програмних систем, спростити процеси управління інфраструктурою та забезпечити стабільність роботи додатків у різних середовищах.

КЛЮЧОВІ СЛОВА: КОНТЕЙНЕРИЗАЦІЯ, DOCKER, DOCKER COMPOSE, МІКРОСЕРВІСИ, ВІРТУАЛІЗАЦІЯ, DEVOPS, МАСШТАБОВАНІСТЬ, ХМАРНІ ОБЧИСЛЕННЯ.

ANNOTATION

The bachelor's thesis contains of 68 pages, 10 figures, 1 tables, and a reference list containing 40 sources.

The aim of the work is to study approaches to the use of containerization for application scaling and to analyze modern tools and technologies that ensure efficient deployment, management, and scalability of software systems.

Research object: processes of application containerization in modern software systems.

Research subject: methods, models, and tools of containerization (Docker, Docker Compose, container architectures) for ensuring scalability, flexibility, and efficient application deployment.

Research results: the principles of containerization were analyzed, its advantages over traditional virtualization were identified, container management tools were studied.

The first section describes the theoretical foundations of containerization, its role in modern software development, and provides a comparison between containers and virtual machines along with their application areas.

The second section analyzes containerization tools, particularly Docker, examines practical use cases, and explores architectural approaches to building containerized systems.

The third section presents the development and implementation of a containerized solution, the use of Docker Compose for service orchestration.

Conclusion: the use of containerization significantly improves scalability, flexibility, and efficiency of software deployment, simplifies infrastructure management, and ensures stable application performance across different environments.

KEY WORDS: CONTAINERIZATION, DOCKER, DOCKER COMPOSE, MICROSERVICES, VIRTUALIZATION, DEVOPS, SCALABILITY, CLOUD COMPUTING.

ЗМІСТ

ВСТУП	7
РОЗДІЛ 1. ТЕОРЕТИЧНІ ОСНОВИ КОНТЕЙНЕРИЗАЦІЇ ТА МАСШТАБУВАННЯ ДОДАТКІВ	9
1.1 Сутність контейнеризації та її роль у сучасній розробці програмного забезпечення	9
1.2 Порівняння контейнеризації та віртуалізації	15
1.3 Застосування контейнеризації у програмній інженерії та масштабуванні додатків	18
1.4 Висновки по розділу	19
РОЗДІЛ 2. МЕТОДИ ТА ІНСТРУМЕНТИ КОНТЕЙНЕРИЗАЦІЇ	20
2.1 Використання Docker як базового інструменту контейнеризації	20
2.2 Практичні сценарії використання контейнерів для розгортання додатків...	25
2.3 Контейнерні архітектури та підходи до їх організації	32
2.4 Висновки по розділу	40
РОЗДІЛ 3. РОЗРОБКА, РОЗГОРТАННЯ ТА ОЦІНКА ЕФЕКТИВНОСТІ КОНТЕЙНЕРИЗОВАНИХ РІШЕНЬ	41
3.1 Керування контейнерами за допомогою Docker Compose	41
3.2 Проєктування та впровадження рішення для контейнеризації додатку	48
3.3 Оцінювання ефективності та масштабованості контейнеризованого рішення	57
3.4 Висновки по розділу	66
ВИСНОВКИ	67
СПИСОК ПОСИЛАНЬ НА ДЖЕРЕЛА	69

					БР.ІП – 22.00.00.000 ПЗ			
Змн.	Арк.	№ докум.	Підпис	Дата	Використання контейнеризації для масштабування додатків Пояснювальна записка	Літ.	Арк.	Акрушів
Розроб.		Польнюк Ю. М.						
Перевір.		Бандура В. В.					7	
Реценз.		Піх В.Я.				ІФНТУНГ ІП-22-1		
Н. Контр.		Піх М.М.						
Затверд.		Бандура В. В.						

ВСТУП

У сучасному світі цифрові технології відіграють ключову роль у розвитку бізнесу, науки та інженерії, а програмні системи стають дедалі складнішими, розподіленими та ресурсоємними. Зростання обсягів даних, підвищення вимог до швидкодії, доступності та надійності програмного забезпечення обумовлюють необхідність використання ефективних підходів до розгортання та масштабування додатків. У цих умовах контейнеризація виступає одним із провідних інструментів, що дозволяє забезпечити гнучкість, портативність і стабільність роботи програмних систем у різних середовищах.

Актуальність теми зумовлена стрімким розвитком хмарних технологій, мікросервісної архітектури та DevOps-підходів, які потребують швидкого та надійного розгортання додатків. Традиційні підходи до віртуалізації часто не забезпечують достатньої ефективності та швидкості масштабування, що робить контейнеризацію більш привабливим рішенням. Використання контейнерів дозволяє ізолювати додатки, оптимізувати використання ресурсів і значно спростити процеси розробки, тестування та розгортання.

Метою роботи є дослідження та розробка підходів до використання контейнеризації для масштабування додатків, а також аналіз інструментів і технологій, що забезпечують ефективне управління контейнеризованими середовищами.

Завданнями дослідження є: аналіз теоретичних основ контейнеризації та її ролі в сучасній програмній інженерії; порівняння контейнерів і віртуальних машин; дослідження інструментів контейнеризації, зокрема Docker; вивчення підходів до побудови контейнерних архітектур; розробка та реалізація контейнеризованого рішення; оцінка ефективності та масштабованості запропонованого підходу.

Об'єктом дослідження є процеси розробки та розгортання програмних систем із використанням контейнеризації.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						7
Змн.	Арк.	№ докум.	Підпис	Дата		

Предметом дослідження є методи, моделі та інструменти контейнеризації, що застосовуються для забезпечення масштабованості, ефективності та гнучкості програмних додатків.

Методи дослідження включають аналіз наукових і технічних джерел, порівняльний аналіз технологій контейнеризації та віртуалізації, моделювання архітектури системи, експериментальний підхід до реалізації контейнеризованого рішення, а також оцінювання його ефективності за визначеними критеріями.

У роботі розглянуто теоретичні аспекти контейнеризації, її відмінності від традиційної віртуалізації та основні сфери застосування. Досліджено сучасні інструменти контейнеризації, зокрема Docker, а також підходи до побудови контейнерних архітектур. У практичній частині реалізовано контейнеризоване рішення з використанням Docker Compose, що дозволяє ефективно керувати декількома сервісами, а також проведено оцінку його продуктивності та масштабованості.

Очікується, що результати роботи сприятимуть підвищенню ефективності розгортання програмних систем, спрощенню процесів управління інфраструктурою та забезпеченню стабільної роботи додатків у різних середовищах.

Бакалаврська робота містить 68 сторінок, 10 рисунків, 21 таблиця, 3 розділи, список використаних джерел із 40 найменувань.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						8
Змн.	Арк.	№ докум.	Підпис	Дата		

РОЗДІЛ 1. ТЕОРЕТИЧНІ ОСНОВИ КОНТЕЙНЕРИЗАЦІЇ ТА МАСШТАБУВАННЯ ДОДАТКІВ

1.1 Сутність контейнеризації та її роль у сучасній розробці програмного забезпечення

Контейнеризація – це метод легкої віртуалізації програм, який сприяє широкому впровадженню хмарних обчислень. Як описано, оркестрація та розгортання контейнерів окремо та в групах було значною проблемою. Було проведено численні дослідження знань про контейнери в хмарі для виявлення тенденцій, прогалин у знаннях та майбутніх напрямків. Дослідження надають порівняльну картину стану досліджень шляхом ідентифікації, класифікації та синтезу доступних даних. У цьому дослідженні було використано систематичне картографування (SMS), щоб допомогти визначити та структурувати нові теми дослідження.

Як процеси розробки, так і розгортання отримують вигоду від використання контейнерів. Наприклад, хмарна архітектура розвивається в напрямку підходів DevOps, що дозволяє безперервний розвиток та конвеєр розповсюдження на основі контейнерів та оркестрації, що включає результати хмарної архітектури.

Згідно з результатами, контейнери можуть забезпечувати безперервний прогрес у хмарі, коли вони використовуються з хмарними платформами для розробки та розподілу, проте вони вимагають складного або координованого забезпечення. Таким чином, методи оркестрації на основі контейнерів стають інструментом для оркестрації обчислень у хмарних, групованих системах [1]. Результати цього дослідження показують, що такі методи реалізуються для досягнення рівноваги між необхідністю розумного контролю якості, наприклад, оптимального використання ресурсів, та продуктивністю, що є питанням вартості в хмарі (через код оцінки корисності).

Ця робота має на меті допомогти дослідникам, які працюють у сфері

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						9
Змн.	Арк.	№ докум.	Підпис	Дата		

програмної інженерії, розподілених структур та хмарних обчислень. Систематична презентація дослідження створює основу знань, на якій можна будувати теорії та пояснення, оцінювати наслідки дослідження та визначати майбутні напрямки. Крім того, вона допомагає практикам, які займаються цією діяльністю, використовуючи існуючі інструменти та технології.

Оскільки віртуалізація дозволяє віртуалізувати комп'ютерні, сховищні та мережеві компоненти, вона є основою хмарних обчислень. Завдяки використанню спеціалізованого програмного рівня, відомого як гіпервізор, віртуалізація дозволяє створювати кілька окремих середовищ виконання, відомих як віртуальні машини (VM), які можна легко переміщувати з однієї хмари в іншу. KVM, Xen, WM-Ware, Virtual Box та Virtual PC – це деякі програмні рішення для віртуалізації гіпервізора. Віртуалізація контейнерів Linux нещодавно з'явилася як легка альтернатива HVV (LCV). Цей метод розділяє фактичні ресурси комп'ютера на кілька окремих екземплярів у просторі користувача. Docker, Kubernetes, OpenVZ та Virtuozzo – лише кілька прикладів програмних рішень LCV[2].

Основна відмінність між HVV та LCV полягає в тому, що HVV абстрагує операційні системи відвідувачів, тоді як LCV об'єднує одне ядро ОС між контейнерами. Коротко кажучи, HVV встановлює концептуальний рівень апаратного забезпечення комп'ютера, тоді як LCV використовує системні виклики. Контейнери та віртуальні машини еквівалентні з точки зору користувача. Крім того, оскільки LCV споживає менше апаратних ресурсів, ніж HVV, він дозволяє розміщувати більше контейнерів на одному фізичному хості, ніж віртуальні машини.

Контейнери доступні у двох різних конфігураціях.

Контейнер застосунку: Окремий контейнер спеціально створений для виконання застосунку одного типу;

Системний контейнер: Простір користувача міститься всередині іншого контейнера. LCV забезпечує більшу адаптивність з точки зору проектування, оптимізації та адміністрування послуг.

У наступному розділі порівнюються HVV та LCV.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						10
Змн.	Арк.	№ докум.	Підпис	Дата		

- Час запуску. На етапі запуску LCV швидший, ніж HVV.
- Динамічне керування середовищем виконання. LCV дозволяє хост-ОС запускати або зупиняти контейнеризовану програму, тоді як HVV зазвичай потребує від хост-ОС встановлення віртуального мережевого/послідовного з'єднання.
- Швидкість. Тільки експерименти, проведені за певних умов, здатні визначити затримку та накладні витрати на пропускну здатність.
- Ізоляція. Порівняно з LCV, HVV має високий рівень ізоляції.
- Споживання флеш-пам'яті. LCV дозволяє роботу ядра ОС та інших частин середовища користувача, тоді як HVV не дозволяє спільний доступ через розділені образи віртуальних машин.
- Динамічне призначення або розділення ресурсів. Обидва рішення здатні реалізувати цю функціональність.
- Прямий доступ до комп'ютерного обладнання. На відміну від HVV, LCV потребує спеціалізованих драйверів у ядрі ОС для доступу до периферійних пристроїв комп'ютера. На основі цього дослідження ми робимо висновок, що LCV перемагає HVV.

Постачальники хмарних послуг Інтернету речей можуть використовувати моделі як HVV, так і LCV для надання своїх послуг. Пристрої Інтернету речей, які працюють на спеціалізованому програмному забезпеченні, підключені до хмарної бази, яка може керувати різнорідними ідентифікаційними даними з кількох датчиків. Через обмежені обчислювальні можливості пристрій Інтернету речей зазвичай виконує лише примітивні програми для виявлення та керування. З іншого боку, хмарний центр обробки даних виконує складні обчислювальні завдання, такі як зберігання та обробка даних датчиків. Кілька проектів розглядали застосунки на основі LCV для пристроїв Інтернету речей, щоб скористатися перевагами як HVV, так і LCV. На рисунку 1.1 зображено хмару Інтернету речей, яка використовує знання як HVV, так і LCV. Структура

Інтернету речей відповідає за наступне:

- Створення екземпляра;

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						11
Змн.	Арк.	№ докум.	Підпис	Дата		

- Гарантування узгодженості, якості обслуговування, безпеки та масштабованості;
- Контроль та оптимізація IoTaaS.

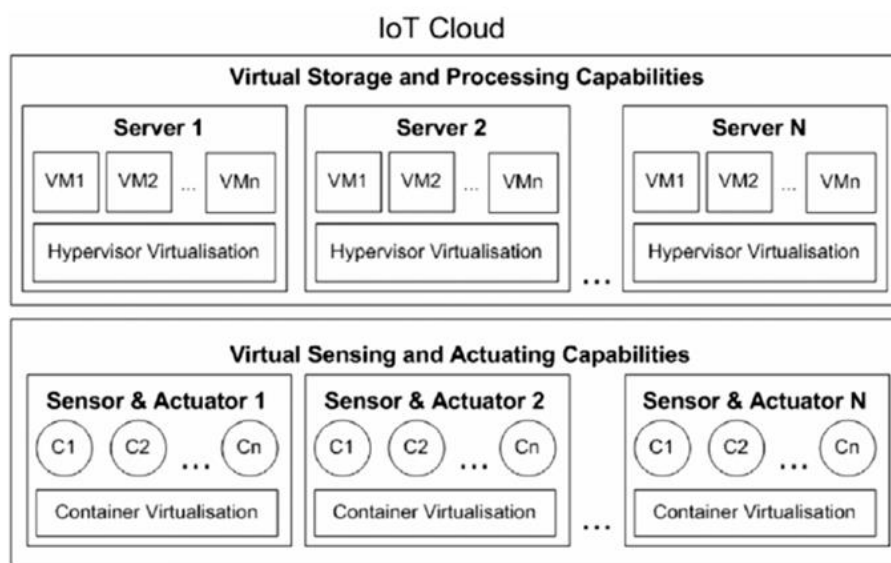


Рисунок 1.1 - Демонстрація хмарної платформи «Інтернету речей», що використовує досвід у сфері віртуалізації як HVV, так і LCV

Крім того, оскільки пристрої Інтернету речей включають можливості LCV (LCV), IoTaaS може бути розгорнутий у багатьох розподілених контейнерах. Інформація, зібрана з пристроїв Інтернету речей, стандартизована та зберігається на хмарному рівні, активуючи виконавчі механізми, пов'язані з пристроями Інтернету речей. HVV (високоякісний вибіркового обсягу) та LCV (низькоякісний вибіркового обсягу) створюють плівку абстракції, яка приховує всі фізичні навички. Зазвичай, працівники хмарних технологій Інтернету речей пропонують обидва шаблони, але в пристроях Інтернету речей використовується лише LCV. Щоб динамічно пропонувати послуги своїм клієнтам, оператори хмар Інтернету речей встановлювати різноманітні контейнери та віртуальні машини (ВМ) у свою інфраструктуру [3]. Це дозволяє працівнику реорганізовувати, вдосконалювати та переміщувати свої комп'ютерні ресурси. Оператор хмари Інтернету речей може виконувати будь-які запити на розподіл послуг, зроблені його клієнтами,

використовуючи цю інфраструктуру. Хмари Інтернету речей можуть виконувати різні дії завдяки функціональності LCV для одноплатних комп'ютерів (SBC).

- Розгортання розподіленого IoTaaS. Розподілений IoTaaS можна налаштувати шляхом об'єднання контейнерів, встановлених на різних IoT-пристроях.

- Перепозиціонування та оптимізація IoTaaS. Хмара IoT може переміщувати контейнер з одного IoT-гаджета до іншого, щоб впровадити методи балансування навантаження;

- Консолідація IoTaaS. Використання IoTaaS з урахуванням місцезнаходження, побудованого шляхом об'єднання контейнерів, розгорнутих на різних пристроях IoT, дозволяє запровадити методи консолідації програмного забезпечення на основі логіки програми. Це відкриває нові потоки доходів для працівників IoT Cloud у сферах економічно ефективного переміщення та оптимізації активів, енергозбереження та постачання ресурсів на вимогу.

Концепція контейнерів та ізоляції процесів є відносно старою та існує вже кілька десятиліть, з ранніми прикладами, такими як системний виклик *chroot* у версії 7 Unix у 1979 році. Операція *chroot* дозволяла змінювати видимий кореневий каталог процесів та їхніх дочірніх процесів [8]. Отже, процеси зі зміненими кореневими каталогами не могли отримати доступ до будь-яких ресурсів поза межами призначеного їм кореневого каталогу, тим самим створюючи певною мірою ізольоване середовище. Однак сучасна форма контейнерів з'явилася лише в останні два десятиліття. Широке впровадження контейнеризації можна пояснити запуском *Docker* у 2013 році, який з того часу став фактичним стандартом для розгортання та управління контейнерами.

Основною рушійною силою, що стимулювала розвиток контейнерних технологій, була потреба в більш ефективному та гнучкому методі розгортання та управління додатками у великомасштабних середовищах. З цією метою було розроблено різні методи та технології контейнеризації, кожна з яких базується на попередніх інноваціях [4]. Наприклад, у 2000 році FreeBSD представила *Jails*, інструмент, який розширив традиційну концепцію *chroot*-ізоляції процесів, що

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						13
Змн.	Арк.	№ докум.	Підпис	Дата		

передбачає зміну кореневого каталогу набору процесів, щоб обмежити їхній доступ до файлів та ресурсів поза їхнім кореневим каталогом. Jails ще більше розширили це, обмеживши також доступ для системних користувачів та мережевої підсистеми, яка раніше була спільною для процесів хоста.

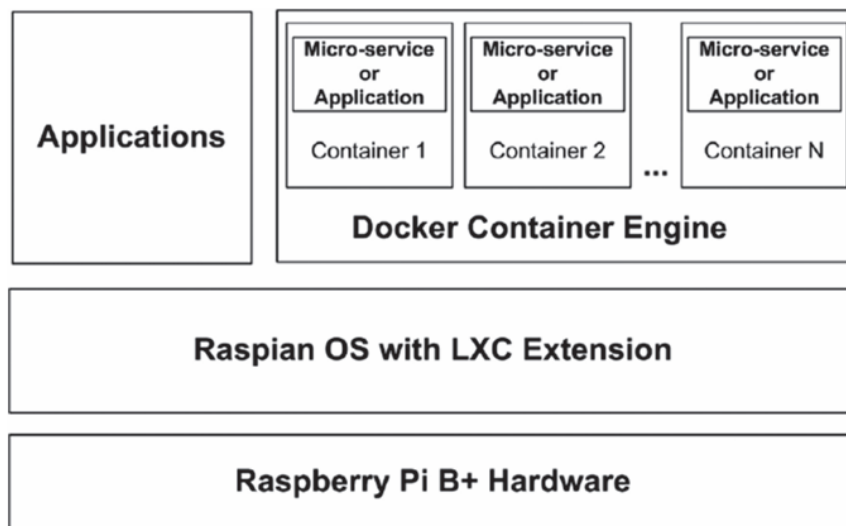


Рисунок 1.2 Програма для віртуалізації контейнерів на базі Docker

У Linux також у 2001 році була розроблена подібна концепція з *Linux VServer*, яка розділила простір користувача на окремі *віртуальні приватні сервери (VPS)*, кожен з яких функціонував як окремий сервер з точки зору внутрішніх процесів. [10] У 2006 році Google запустила *cgroups* (контрольні групи), які дозволили пріоритезувати, обмежувати та враховувати ресурси, що використовуються групою процесів, гарантуючи, що вони не перевищують призначений обсяг системних ресурсів, таких як використання процесора та оперативної пам'яті. Cgroups були об'єднані з ядром Linux у 2008 році. [10] [11]

У 2008 році також було випущено проект *Linux Containers (LXC)*, який спирався на функціональність контрольних груп та просторів імен Linux. LXC забезпечував середовища що дуже нагадували стандартні окремі інсталяції Linux без потреби в окремому ядрі. [10] [12]

У 2013 році було випущено *Docker*, який швидко здобув популярність завдяки простоті використання та ефективності. Поточну популярність

контейнерних технологій можна значною мірою пояснити успіхом Docker. Docker базувався на всіх попередніх поступових удосконаленнях, які були запроваджені старими методами, і зробив його доступнішим для розробників. Наприклад, на початкових етапах свого випуску Docker використовував LXC як середовище виконання за замовчуванням. Однак, після випуску версії 0.9 через рік, Docker замінив його на *libcontainer*, власну реалізацію середовища виконання. [13] Через свою популярність та статус фактичної контейнерної технології, ця робота з цього моменту буде зосереджена лише на використанні Docker як обраної контейнерної технології.

1.2 Порівняння контейнеризації та віртуалізації

Контейнери та віртуальні машини – це схожі підходи до віртуалізації. Вони створюють рівень абстракції поверх комп'ютерного обладнання та інфраструктури, який дозволяє представляти окремі системні ресурси, такі як процесор, пам'ять та мережа, та розділяти їх на кілька віртуальних ресурсів. Основна відмінність між ними полягає у способі досягнення цього. Віртуальні машини віртуалізують всю машину, починаючи з апаратних рівнів, причому кожна віртуальна машина працює під керуванням унікальної гостьової операційної системи. Контейнери, з іншого боку, віртуалізують лише рівні над операційною системою, причому кожен контейнер використовує ядро операційної системи хост-машини. Ця відмінність візуалізована на рисунку 1.3 .

Форма використання контейнерів віртуалізації називається *віртуалізацією на рівні ОС*. У віртуалізації на рівні ОС кожен контейнер використовує ядро операційної системи хоста, а часто й інші ресурси, такі як бібліотеки та бінарні файли, але вони працюють у власній окремій файловій системі та представленні простору процесів [5].

З іншого боку, віртуальні машини використовують віртуалізацію на апаратному рівні для створення ізольованих середовищ, які можуть запускати кілька операційних систем на одній фізичній машині. Кожна віртуальна машина

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						15
Змн.	Арк.	№ докум.	Підпис	Дата		

має власну гостьову операційну систему та віртуальне обладнання.

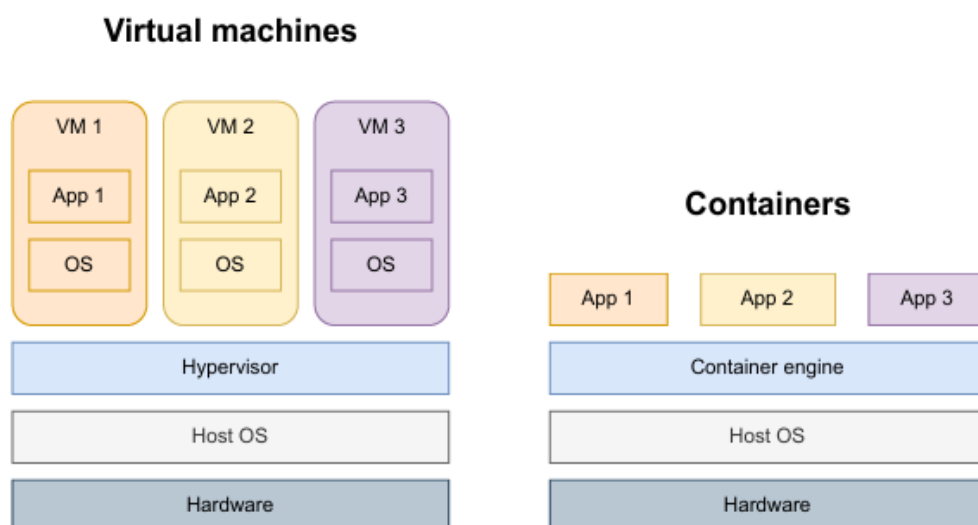


Рисунок 1.3 - Базові рівні типових віртуальних машин у порівнянні з контейнерами.

В останні роки контейнерні технології здобули широку популярність як рішення численних проблем у розробці та розгортанні програмного забезпечення. Основною мотивацією широкого впровадження контейнерів є ізоляція, модульність та відтворюваність, які вони пропонують.

Контейнери інкапсулюють усі залежності та конфігурації, необхідні для запуску програми, в автономні одиниці, забезпечуючи послідовну та передбачувану поведінку в різних середовищах [7]. Це усуває занепокоєння щодо базової інфраструктури з точки зору розробника та дозволяє йому легше маневрувати між різними етапами конвеєра розробки. Наприклад, переміщення програми з ноутбука розробника в тестове середовище та з тестового середовища в робоче середовище відбувається плавніше та призводить до послідовної поведінки. Як результат, організації можуть застосовувати більш гнучкі підходи до розгортання програм та надання послуг, оскільки контейнери можуть запускатися в різних середовищах, таких як фізичні та віртуальні машини, хмарні сервіси, локальні сервери або ноутбук розробника, за умови, що на хост-системі встановлено та налаштовано відповідне середовище виконання контейнера.

Слабо ізольовані середовища контейнерів також забезпечують підвищену безпеку. Оскільки контейнери працюють в окремих обмежених середовищах, ризик поширення атак, спрямованих на певний контейнер, по всій системі зменшується [6]. Крім того, контейнери є мінімальними за своєю природою, оскільки вони містять лише необхідні компоненти для запуску програми, що призводить до мінімізації поверхні атаки.

Масштабованість є ще однією важливою мотивацією для використання контейнерів. Завдяки своїй портативності та модульності, контейнери можна швидко та легко додавати та видаляти, щоб враховувати зміни в попиті, що дозволяє програмам масштабуватися вгору та вниз за потреби. Для довідки, контейнери Docker можна запускати за лічені секунди [15]. Разом з методами оркестрації контейнерів організації можуть оптимізувати використання контейнерів, використовуючи лише стільки, скільки необхідно для забезпечення оптимального досвіду користувача та клієнта в певний момент часу.

Ще однією перевагою контейнерів є їхня легкість порівняно з традиційними методами віртуалізації, такими як віртуальні машини. Контейнери займають значно менший простір і мають менші накладні витрати, пов'язані з їх використанням. Основна причина цього полягає в тому, що окремі контейнери не потребують окремих персональних операційних систем, як віртуальні машини. Натомість контейнери використовують ту саму операційну систему хост-системи та містять лише необхідні компоненти для запуску програми. Як видно на рисунку 1.3, типова віртуальна машина є, по суті, копією операційної системи, що працює поверх гіпервізора, який, у свою чергу, працює поверх фізичного обладнання. Це означає, що програма працює поверх усіх цих рівнів усередині віртуальної машини, що створює певні проблеми для продуктивності та ефективності. Тоді як контейнери працюють безпосередньо поверх контейнерного механізму та хост-операційної системи.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						17
Змн.	Арк.	№ докум.	Підпис	Дата		

1.3 Застосування контейнеризації у програмній інженерії та масштабуванні додатків

Для застосунків, які потребують високого рівня портативності, таких як веб-застосунки, контейнеризація є ідеальним варіантом, оскільки вона дозволяє легко упаковувати та розгортати застосунок за допомогою образів контейнерів. Контейнери об'єднують застосунок та його залежності в незалежні одиниці, які можна легко та надійно розгортати в різних середовищах, таких як середовища розробки, тестування та виробничі середовища [8].

Впровадження *архітектури мікросервісів* – це ще одна сфера, де - контейнеризація є вигідною. Архітектура мікросервісів передбачає розбиття різних завдань та компонентів складного програмного застосунку на менші, незалежно розгортані компоненти, якими можна легко керувати за допомогою контейнерів. Модульна природа архітектури мікросервісів добре підходить для контейнеризації, оскільки дозволяє розгортати та масштабувати кожен компонент незалежно, без впливаючи на решту програми.

Контейнери також є корисним інструментом для управління процесами *безперервної інтеграції/безперервного розгортання* (CI/CD). Традиційною перешкодою в розгортанні програмних застосунків є відмінності між середовищами розробки та виробництва. Розробники створюють та постачають застосунки, які працюють на їхніх локальних машинах, лише для того, щоб виявити, що застосунок не працює належним чином у тестовому або виробничому середовищі [15]. Це часто пов'язано з відмінностями в базовій інфраструктурі та залежностями між середовищами. Автоматизовані конвеєри CI/CD, що використовують контейнери, можуть допомогти зменшити цю проблему, гарантуючи, що застосунок упакований та розгорнутий узгоджено в різних середовищах. Це дозволяє розробникам зосередитися на логіці та функціональності застосунку, а не на базовій інфраструктурі.

Хоча контейнеризація вигідна для багатьох програм, вона може бути не найкращим вибором для всіх випадків використання [9]. Програми, які

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						18
Змн.	Арк.	№ докум.	Підпис	Дата		

потребують низькорівневого доступу до системи, такі як компоненти рівня ядра, або ті, що потребують доступу на рівні апаратного забезпечення, такі як драйвери пристроїв, можуть не підходити для контейнеризації. Це пояснюється тим, що контейнери розроблені як апаратно-агностичні, що означає, що вони не мають прямого доступу до апаратних ресурсів. Тому програми, які потребують прямого доступу до апаратного забезпечення, можуть не підходити для контейнеризації.

1.4 Висновки по розділу

У першому розділі було розглянуто теоретичні основи контейнеризації та масштабування додатків. Проаналізовано сутність контейнеризації, її роль у сучасній розробці програмного забезпечення та вплив на ефективність розгортання програмних систем. Проведено порівняння контейнеризації та віртуалізації, визначено їх переваги, недоліки та особливості використання. Також досліджено застосування контейнерних технологій у програмній інженерії та процесах масштабування додатків. Отримані результати дозволяють сформулювати теоретичну основу для подальшого дослідження методів і практичної реалізації контейнеризованих рішень.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						19
Змн.	Арк.	№ докум.	Підпис	Дата		

РОЗДІЛ 2. МЕТОДИ ТА ІНСТРУМЕНТИ КОНТЕЙНЕРИЗАЦІЇ

2.1 Використання Docker як базового інструменту контейнеризації

Docker + це програмна платформа, яка реалізує концепцію програмних контейнерів. Вперше її представив на PyCon (*конференції Python*) у 2013 році Соломон Хайкс, інженер-програміст dotCloud, компанії-платформи як послуги (*PaaS*), яка надавала середовище для створення та розгортання веб-додатків. Спочатку,

Docker було створено як внутрішній інструмент для dotCloud, щоб допомогти ефективніше керувати додатками та розгортати їх. [16] Однак проект швидко здобув популярність у спільноті розробників і зрештою був випущений як проект з відкритим кодом.

Як обговорювалося в розділі, лише з появою Docker у 2013 році контейнерні технології стали широко поширеними в індустрії програмного забезпечення. Docker базується на попередніх методах контейнеризації, головним чином надаючи зручний для розробників інтерфейс для створення, керування та спільного використання контейнерів. Ця простота використання та зручність для розробників є одними з основних причин популярності Docker.

Docker - це клієнт-серверна програма. Клієнт Docker взаємодіє з сервером {*Docker daemon* } через REST API, через сокети UNIX або мережевий інтерфейс [12]. Як клієнтські, так і серверні компоненти {*daemon* } можуть запускатися в одній системі або підключатися через віддалене з'єднання. *Docker CLI* та *Docker Compose* - це клієнтські програми, з якими користувачі Docker взаємодіють, щоб надавати команди та інструкції демону Docker. Демон Docker, у свою чергу, діє як сервер, який керує так званими *об'єктами Docker*, такими як образи Docker, контейнери, мережі та томи. Загальна структура Docker проілюстрована в документації Docker {див. рис. 2.1) .

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						20
Змн.	Арк.	№ докум.	Підпис	Дата		

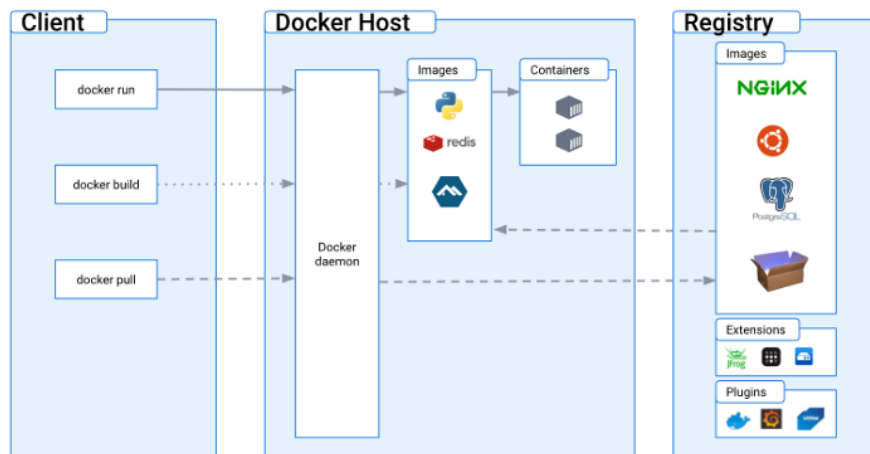


Рисунок 2.1 - Архітектура Docker.

Об'єкти Docker - це різні компоненти, що складають середовище Docker. Вони створюються та управляються Docker Engine, основним компонентом Docker, який дозволяє створювати, запускати та керувати контейнерами. Об'єкти Docker інкапсулюють різні аспекти екосистеми Docker та надають спосіб взаємодії з ними та маніпулювання ними.

Образи Docker є основною частиною запуску контейнерів Docker. Вони є незмінними знімками інструкцій для створення контейнерів і використовуються щоразу, коли потрібно створити контейнер. Образи є дуже портативними, і ними можна ділитися між користувачами, які можуть захотіти запускати контейнери на їх основі або створювати поверх них [13]. Часто буває, що образи не створюються з нуля, а навпаки, вони створюються поверх існуючих образів з додатковою налаштуванням. Наприклад, можна створити образ на основі загальнодоступного образу операційної системи Ubuntu та додатково встановити поверх нього веб-сервер разом із необхідними деталями конфігурації.

Dockerfile, строго кажучи, не є об'єктом Docker, але він є важливою частиною створення образу Docker, оскільки служить скриптом для визначення вмісту та конфігурації образу. Це звичайний текстовий файл, який містить низку інструкцій, що визначають, як налаштувати середовище контейнера. Інструкцією може бути, наприклад, визначення змінної середовища або bash-скрипт для виконання. Синтаксис, який використовується в Dockerfiles для визначення

інструкцій, досить простий і легкий для читання, і він підтримує різні параметри конфігурації, що дозволяють точно налаштувати вміст образу та середовище виконання. Лістинг 2.2 на сторінці 25 містить приклад простого Dockerfile, який можна використовувати для створення образу застосунку *Node.js*.

Кожна інструкція в Dockerfile представляє шар у образі, а результируючий образ є стеком усіх шарів, визначених у Dockerfile. Щоразу, коли Dockerfile змінюється та використовується для створення образу, перебудовуються лише ті шари, що змінилися. Згідно з документацією Docker, це одна з причин, чому образи Docker такі легкі та швидкі порівняно з іншими технологіями віртуалізації [17].

Деякі з найпоширеніших інструкцій Dockerfile включають:

- FROM : Визначає базове зображення, на якому буде базуватися нове зображення. Кожен Dockerfile має починатися з інструкції FROM . Щоб створити зображення з нуля без базового зображення, можна скористатися зарезервованим у Docker явно порожнім *скретчем зображень*.

- RUN : Використовується для виконання команди всередині контейнера. Зазвичай використовується для

встановлення програмних пакетів або налаштування середовища.

- КОПІЮВАННЯ : Копіює файли з хост-комп'ютера до контейнера. Наприклад, файли конфігурації, скрипти або вихідний код.

- EXPOSE : Надає доступ до порту з контейнера для хост-машини.

- CMD: Визначає команду за замовчуванням, яка буде виконана під час запуску контейнера. Для Dockerfile дозволено використовувати лише одну інструкцію CMD .

Контейнери Docker - це екземпляри Docker-образів під час виконання, що забезпечують ізольовані середовища для запуску програм та їхніх залежностей, визначених образом [14]. Контейнери створюються з образів, які діють як незмінні шаблони, що містять усі необхідні компоненти та конфігурації для запуску програми.

Контейнери за замовчуванням не зберігають жодних даних про свій стан,

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						22
Змн.	Арк.	№ докум.	Підпис	Дата		

натомість їх можна налаштувати для зберігання даних у томах, які є окремими від контейнера та можуть бути спільними для інших контейнерів. Це гарантує збереження даних навіть після знищення та повторного створення контейнера [15]. За замовчуванням контейнери також забезпечують високий ступінь ізоляції від хост-машини та інших контейнерів, який можна налаштувати за допомогою додаткової конфігурації.

Реєстри Docker не є життєво важливими з точки зору роботи контейнерів Docker, проте вони необхідні під час створення образів Docker, оскільки вони діють як механізми зберігання та розповсюдження образів Docker, дозволяючи користувачам обмінюватися образами контейнерів та отримувати до них доступ.

Як згадувалося раніше, щоб зменшити навантаження розробника, образи зазвичай створюються поверх уже існуючих образів. Ці образи витягуються з публічних або приватних реєстрів. Наприклад, *Docker Hub*- це публічний реєстр, який Docker використовує для пошуку образів за замовчуванням.

Публічні реєстри використовуються для публічного обміну та завантаження зображень, тоді як приватні реєстри в основному використовуються в організаціях для безпечного зберігання та розповсюдження зображень у межах їхньої інфраструктури. [17]

Томи Docker - це функція Docker, яка забезпечує спосіб зберігання та обміну даними між контейнерами Docker та хост-системою. Вони дозволяють відокремити дані від самого контейнера, забезпечуючи їм збереження після перезапуску, оновлення та навіть видалення контейнера.

Коли контейнер створюється, він ізолює його від хост-системи та будь-яких інших контейнерів. За замовчуванням Docker надає файлову систему всередині контейнера, але ця файлова система є тимчасовою, тобто будь-які зміни, внесені всередині контейнера, втрачаються, коли контейнер зупиняється або видаляється.

Томи Docker надають механізм для створення постійної області сховища, яка може використовуватися одним або кількома контейнерами. Замість зберігання даних у файловій системі контейнера, ви можете змонтувати том до

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						23
Змн.	Арк.	№ докум.	Підпис	Дата		

певного каталогу всередині контейнера. Це дозволяє вам читати з тому та записувати на нього, а будь-які зміни, внесені до тому, зберігаються, навіть якщо контейнер зупинено або видалено. Один і той самий том можна підключити до кількох різних контейнерів, що дозволяє створювати конвеєри даних та спільне сховище між контейнерами. [18]

Мережі Docker є невід'ємним елементом мережевих можливостей Docker. Вони забезпечують спосіб об'єднання контейнерів Docker та забезпечення зв'язку між ними, як в межах одного хоста Docker, так і між кількома хостами Docker.

Після встановлення Docker автоматично генерує три окремі мережі: *bridge*, *host* та *none* . Поточні активні мережі Docker можна перевірити за допомогою команди `docker network ls` :

Лістинг 2.1 - Перевірка за допомогою команди `docker network ls`

```
$ docker network ls
NETWORK ID          NAME                DRIVER              SCOPE
1dd11129834b        bridge             bridge              local
9a94cb0c01ef        host               host                local
65e7bc68b100        none               null                local
```

за замовчуванням призначається мережі мосту . Контейнери, підключені до мережі мосту, можуть взаємодіяти один з одним, використовуючи IP-адреси в межах одного хоста Docker. Docker призначає IP-адреси контейнерам у мережі **мосту** , і контейнери можуть використовувати ці адреси для взаємодії один з одним.

На відміну від цього, мережа **хоста** – це мережа хост-машини, тобто контейнер, підключений до неї, буде безпосередньо використовувати мережеву інфраструктуру хост-машини [16]. Отже, це впливає на ізоляцію контейнера, повністю усуваючи мережеву ізоляцію між контейнером і хостом. З іншого боку, параметр «**без мережі** » служить для вимкнення мережевих можливостей контейнера.

Окрім попередньо визначених мереж, Docker дозволяє створювати власні

мережі. Ці налаштовані мережі пропонують покращену ізоляцію та сегментацію контейнерів, що дозволяє створювати окремі мережі для різних груп контейнерів [17]. Коли створюється власна мережа, Docker встановлює віртуальну мережу, до якої можуть приєднуватися контейнери. Контейнери, пов'язані з однією й тією ж власною мережею, можуть спілкуватися один з одним, використовуючи свої відповідні IP-адреси.

2.2 Практичні сценарії використання контейнерів для розгортання додатків

Припустимо, ви розробили простий застосунок Node.js (див. лістинг 2.2), який має запускатися в контейнері. Першим кроком для його виконання, після встановлення необхідних інструментів командного рядка Docker, є визначення Dockerfile, який визначає кроки, необхідні для створення образу. Відповідний Dockerfile можна побачити в лістингу 2.2. Після визначення Dockerfile створюється образ, який потім використовується для створення запущеного екземпляра контейнера.

Хоча можна почати з нуля без базового образу, зазвичай зручніше використовувати базовий образ для подальшої розробки, наприклад, офіційний образ Ubuntu. Однак кожен Dockerfile повинен починатися з інструкції **FROM**, і у випадку лістингу образ `node:18` вказується як базовий образ для подальшої розробки. Аргумент інструкції `node` посилається на загальнодоступний образ середовища виконання Node.js, а тег `18` – на його версію. Якщо версію пропустити, інструкція `FROM` за замовчуванням використовуватиме останню версію образу, що те саме, що й при вказівці `FROM node:latest` [18].

Після інструкції `FROM` робочим каталогом програми визначається каталог `/app` за допомогою інструкції `WORKDIR`. Тепер кожна наступна інструкція починатиметься з каталогу `/app`.

Лістинг 2.2 - Мінімальний веб-додаток на JavaScript, що видає числа Фібоначчі на основі індексу запиту.

```

1 // index.js
2
3 const app = require("express")();
4 const { readFile, writeFile, mkdir } = require("fs/promises");
5 const { dirname } = require("path");
6
7 const PORT = process.env.PORT || 8080;
8 const INDEX_PATH = "./data/index.txt";
9
10 app.get("/fibonacci", async (req, res) => {
11   const index = await readIndex();
12   const fib = fibonacci(index);
13   await writeIndex(index + 1);
14   res.json({ fib, index });
15 });
16
17 app.listen(PORT,
18   () => console.log(`App listening on ${PORT}...`));
19
20 function fibonacci(n) {
21   if (n === 0) return 0;
22   if (n === 1) return 1;
23   return fibonacci(n - 1) + fibonacci(n - 2);
24 }
25
26 async function readIndex() {
27   try {
28     const data = await readFile(INDEX_PATH, "utf8");
29     return parseInt(data);
30   } catch {
31     return 0;
32   }
33 }
34
35 async function writeIndex(index) {
36   await mkdir(dirname(INDEX_PATH), { recursive: true });
37   await writeFile(INDEX_PATH, index.toString(), { flag: "w+" });
38 }

```

Далі файли `package.json` та `package-lock.json` копіюються з хост-комп'ютера до робочого каталогу програми. Копіювання файлів здійснюється за допомогою інструкції `COPY`, де перший аргумент – це шлях до файлу на хост-комп'ютері, а другий – шлях до робочого каталогу контейнера.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						26
Змн.	Арк.	№ докум.	Підпис	Дата		

Лістинг 2.3 - Приклад файлу Dockerfile для створення образу Docker

```
1 # Dockerfile
2
3 FROM node:18
4
5 WORKDIR /app
6
7 COPY package*.json ./
8
9 RUN npm install
10
11 COPY index.js ./
12
13 ENV PORT=8080
14
15 EXPOSE 8080
16
17 CMD ["node", "index.js"]
```

Файли `package.json` та `package-lock.json` – це файли конфігурації, які визначають залежності програми від інших зовнішніх модулів NPM (*Node Package Manager*), серед іншого. Наприклад, як видно з рядка 3 лістингу 2.1, програма використовує фреймворк *Express*, який є зовнішньою залежністю від стороннього розробника [19].

Тепер, коли інформація про залежності знаходиться в контейнері, їх можна встановити за допомогою команди `npm install`, яку можна виконати всередині контейнера за допомогою інструкції `RUN`.

Після встановлення залежностей вихідний код програми буде скопійовано до робочого каталогу за допомогою інструкції `COPY`. Причина, чому гарною практикою є копіювання вихідного коду лише після встановлення залежностей, полягає в тому, як Docker обробляє та кешує шари Dockerfile. Як обговорювалося в розділі 3.5.3, кожна інструкція в Dockerfile відповідає шару в результуючому зображенні, і щоразу, коли зображення будується з Dockerfile, Docker намагатиметься кешувати шари, якщо нічого не змінилося. Оскільки можна очікувати, що вихідний код буде змінюватися частіше, ніж файли `package.json`, шар `RUN npm install` буде зібрано лише тоді, коли змінюється файл `package.json`. Тобто модулі NPM не потрібно встановлювати щоразу, коли змінюється вихідний код програми. У складніших проектах з більшою кількістю залежностей це

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						27
Змн.	Арк.	№ докум.	Підпис	Дата		

потенційно заощадить багато часу під час створення образу.

Як видно з рядка 7 лістингу 2.1 , застосунок посилається на змінну середовища під назвою `PORT` . Щоб застосунок мав доступ до цієї змінної середовища під час виконання, вона визначається в `Dockerfile` за допомогою інструкції `ENV` [20]. Зауважте, що якщо змінна середовища має залишатися секретною та не передаватися до системи контролю версій разом із `Dockerfile`, її також можна передати контейнеру під час його створення за допомогою опції `-e` . Наприклад, `docker run -e PORT=8080`.

Відкриття портів

Далі, інструкція `EXPOSE` використовується для інформування `Docker` про те, що контейнер прослуховуватиме вказаний порт під час виконання. Інструкція фактично не публікує порт; радше вона служить механізмом документування, який вказує, які порти слід відкрити під час запуску контейнера. Щоб фактично опублікувати порт і дозволити доступ до програми з хост-машини, правила переадресації портів необхідно вказати під час запуску контейнера.

Нарешті, інструкція `CMD` визначає, як контейнер має запускати програму після її запуску. У кожному `Dockerfile` дозволено використовувати лише одну інструкцію `CMD` [20].

Визначивши набір інструкцій для збірки образу `Docker` у вигляді `Dockerfile`, образ можна зібрати за допомогою команди `docker build -t fib-app`. Опція `-t` використовується для назви образу, а аргумент `.` вказує шлях до контексту збірки `.` У цьому випадку шлях посилається на поточний каталог, в якому знаходиться `Dockerfile` разом з іншими необхідними файлами.

Після успішного створення образу, його можна поширювати через реєстр контейнерів та використовувати як базовий образ для інших образів або для створення контейнерів [21].

Створення контейнера за допомогою образу можна виконати за допомогою команди `docker run` . Однак, оскільки програма має бути доступною ззовні контейнера, переадресацію портів з контейнера на локальну машину необхідно реалізувати за допомогою опції `-p` . Команда `docker run --name fib-server -p`

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						28
Змн.	Арк.	№ докум.	Підпис	Дата		

7000:8080 fib-app

Команду можна використовувати для створення контейнера за допомогою образу демонстраційної програми. Опція `--name` надає контейнеру назву для ідентифікації, тоді як опція `-p` зіставляє порт 7000 локальної машини з портом 8080 контейнера. Завдяки переадресації портів, до програми можна отримати доступ зовні контейнера, наприклад, через веб-браузер. На рисунку 2.2 показано результат запиту до `localhost:7000/fibonacci` на хост-машині за допомогою веб-браузера.

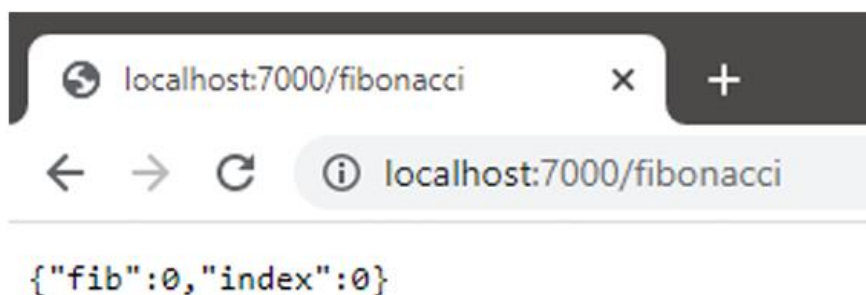


Рисунок 2.2 - JSON-дані, повернуті демонстраційним додатком, що працює в контейнері

Додаток Fibonacci використовує текстовий файл як грубу та незграбну базу даних для зберігання інформації про поточне значення лічильника індексу Фібоначчі. Однак у робочих середовищах зазвичай потрібна фактична надійна реалізація бази даних. Для цього Docker Hub містить готові образи для багатьох популярних систем керування базами даних (СКБД), таких як Oracle Database, MySQL, PostgreSQL та MariaDB [22].

Щоб розмежувати питання бізнес-логіки застосунку та збереження даних, розумним підходом було б виділити окремий контейнер для розміщення СУБД. Dockerfile для створення образу контейнера для розміщення бази даних можна побачити в лістингу 2.4. Він використовує останню версію системи керування реляційними базами даних MariaDB (RDBMS) як базовий образ. Dockerfile також визначає необхідні змінні середовища та створює таблицю, яка відповідає

потребам програми Фібоначчі, разом із рядком, що містить значення за замовчуванням.

Після додавання нового контейнера до програми, він повинен мати можливість взаємодіяти та взаємодіяти з контейнером, на якому запущено програму Фібоначчі, що можна досягти за допомогою мережевих можливостей Docker. Як описано за замовчуванням Docker виділяє контейнери в мережу- *міст* , якщо не вказано інше. Однак для демонстраційних цілей буде використана окрема мережа з назвою *fibnet*. Її створення можна виконати за допомогою команди `docker network create fibnet` [23]. Групування пов'язаних контейнерів, які потребують зв'язку в межах окремих мереж, є гарною практикою з точки зору ізоляції, оскільки загалом контейнери не здатні взаємодіяти за межами своїх відповідних мереж.

```
1 FROM mariadb:latest
2
3 ENV MYSQL_ROOT_PASSWORD=fib-admin
4 ENV MYSQL_DATABASE=fib-db
5
6 RUN echo "USE fib-db;" > /docker-entrypoint-initdb.d/init.sql
7
8 RUN echo "CREATE TABLE IF NOT EXISTS fib_index (" \
9   >> /docker-entrypoint-initdb.d/init.sql
10 RUN echo "id INTEGER PRIMARY KEY," \
11   >> /docker-entrypoint-initdb.d/init.sql
12 RUN echo "current_index INTEGER NOT NULL DEFAULT 0);" \
13   >> /docker-entrypoint-initdb.d/init.sql
14
15 RUN echo "INSERT IGNORE INTO fib_index VALUES (1, 0);" \
16   >> /docker-entrypoint-initdb.d/init.sql
17
18 EXPOSE 3306
```

Лістинг 2.4 - Файл Dockerfile для ініціалізації бази даних MariaDB із заздалегідь заповненою таблицею під назвою `fib_index`

Врахування міграції зі сховища даних на основі текстових файлів до бази даних можна побачити в лістингу 2.5. Примітно, що в рядку 9 лістингу адреса контейнера бази даних визначається шляхом посилання на ім'я контейнера. Під час виконання ім'я контейнера буде перетворено на відповідну IP-адресу. Аналогічно DNS (*системі доменних імен*) в Інтернеті, це спрощує керування кількома контейнерами в мережі, оскільки не потрібно запам'ятовувати конкретні

IP-адреси - достатньо імен контейнерів.

```
1 // index.js
2
3 const app = require("express")();
4 const mysql = require("mysql");
5
6 const PORT = process.env.PORT || 8080;
7
8 const dbConnection = mysql.createConnection({
9   host: "fib-db",
10  user: process.env.DB_USER,
11  password: process.env.DB_PWD,
12  database: process.env.DB_NAME,
13 });
14 dbConnection.connect((error) =>
15   console.log(error ? error : "Database connected")
16 );
17
18 app.get("/fibonacci", async (req, res) => {
19   const index = await readIndex();
20   const fib = fibonacci(index);
21   await writeIndex(index + 1);
22   res.json({ fib, index });
23 });
24
25 app.listen(PORT,
26   () => console.log(`App listening on ${PORT}...`));
27
28 function fibonacci(n) {
29   if (n === 0) return 0;
30   if (n === 1) return 1;
31   return fibonacci(n - 1) + fibonacci(n - 2);
32 }
33
34 function readIndex() {
35   return new Promise((resolve) => {
36     dbConnection.query(
37       "SELECT current_index FROM fib_index WHERE id = 1",
38       (error, results) => {
39         if (error) return 0;
40         resolve(results[0].current_index);
41       }
42     );
43   });
44 }
45
46 function writeIndex(index) {
47   return new Promise((resolve, reject) => {
48     dbConnection.query(
49       "UPDATE fib_index SET current_index = ? WHERE id = 1",
50       [index],
51       (error, results) =>
52         (error ? reject(error) : resolve(results))
53     );
54   });
55 }
```

Лістинг 2.5 - Мінімальний веб-додаток на JavaScript, який видає числа Фібоначчі на основі індексу, що зберігається у зовнішній базі даних.

Щоб призначити контейнер мережі, відмінній від мережі *мосту* за замовчуванням, під час створення екземпляра контейнера необхідно вказати опцію *-network*. Наприклад, щоб створити контейнер, на якому розміщено застосунок Фібоначчі, та призначити його мережі *fibnet*, можна використовувати таку команду: `docker run --name fib-server - network fibnet fib-app`. Згодом будь-які наступні контейнери в тій самій мережі *fibnet* зможуть зв'язуватися з ним через HTTP-запити, посилаючись на ім'я контейнера [24]. Наприклад, після створення

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						31
Змн.	Арк.	№ докум.	Підпис	Дата		

образу на основі Dockerfile, наведеного в лістингу 2.4 , та призначення контейнера, створеного з цього образу, мережі *fibnet* , застосунок Fibonассі може зв'язуватися з ним, просто посилаючись на його ім'я та порт, і навпаки.

У представленому застосунку Фібоначчі повернене число Фібоначчі обчислюється на основі лічильника, який збільшується на одиницю під час кожного запиту, зробленого до кінцевої точки /fibonассі . Однак, внутрішній стан контейнерів є тимчасовим. Будь-який файл або дані, які контейнер може записати до своєї файлової системи, не переживуть перезапуски контейнера, якщо їх не запишуть на том Docker. Таким чином, щоб зберегти базу даних застосунку, а отже, і значення лічильника, після життєвого циклу контейнера, необхідно створити том і підключити його до контейнера [25].

Том з назвою *fib-vol* можна створити за допомогою команди `docker volume create fib-vol` . Том потім можна приєднати до контейнера після його створення з параметром `-v` . Наприклад, команда `docker run --name fib-db --network fibnet -v fib-vol:/var/lib/mysql fib-db` створює контейнер бази даних, визначений у лістингу 2.4 , з томом *fib-vol*, змонтованим до каталогу `/var/lib/mysql` контейнера, який MariaDB використовує за замовчуванням для зберігання даних бази даних.

2.3 Контейнерні архітектури та підходи до їх організації

Віртуалізація використовується в хмарних обчисленнях для забезпечення еластичності великомасштабних спільних ресурсів. На рівні інфраструктури віртуальні машини (ВМ) зазвичай служать основою. На противагу цьому, контейнеризація дозволяє спрощувати віртуалізацію, налаштовуючи контейнери як відповідність додатків.

з окремих образів (часто отриманих з джерела образів), які використовують менше властивостей та часу. Крім того, вони забезпечують більш сумісну упаковку програм, що необхідно для хмарних переміщуваних та сумісних пакетних програм [26]. Контейнеризація створена для створення, тестування та розгортання програм у розподіленій мережі комп'ютерів та зв'язування цих

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						32
Змн.	Арк.	№ докум.	Підпис	Дата		

контейнерів. Таким чином, контейнери вирішують проблеми, пов'язані з хмарним PaaS. Враховуючи загальну важливість хмари, критично важливо мати консолідовану картину поточних операцій.

Для виконання програм потрібні упаковані, самостійні, готові до розповсюдження запитів контейнерні сховища та, за необхідності, проміжне програмне забезпечення та комерційна логіка (у вигляді бінарних файлів та бібліотек). Контейнерні апарати, такі як Docker, базуються на контейнерних машинах, які служать пересувними контейнерами для програм. Як наслідок, у багаторівневих системах необхідно обробляти залежності контейнерів. У багаторівневому плані план оркестрації може визначати компоненти, їхні потреби та термін їхньої служби. Після цього хмара PaaS може виконувати процеси через посередники (такі як контейнерний пристрій). Як результат, хмари PaaS можуть забезпечити розміщення контейнерних програм [27]. Їх скоординована розробка, розгортання та безперервне адміністрування в цьому контексті називають оркестрацією.

Численні контейнерні системи побудовані на фреймворку Linux LXC. Поточні версії Linux, що є частиною контейнерної схеми Linux LXC, мають основні функції, такі як простори імен та групи, що дозволяють процеси, які мають бути ізольовані в спільній операційній системі [35]. Docker зараз є максимально поширеним поясненням контейнера та використовувався для демонстрації контейнеризації. Копія Docker складається з багатошарових файлових систем, аналогічних інструментам LXC стеку віртуалізації Linux; Docker використовує комбіноване монтування, щоб додати файлову структуру з можливістю запису поверх файлової структури лише для читання. Це дозволяє співіснувати кілька файлових систем лише для читання. Ця функція дозволяє створювати нові зображення шляхом нашарування існуючих. Редагувати можна лише верхній шар контейнера [28].

Контейнеризація дозволяє перехід від автономних контейнерних програм до кластерів контейнерних хостів, здатних запускати контейнерні програми на кластерних хостах. Властива контейнерам сумісність сприяє останньому.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						33
Змн.	Арк.	№ докум.	Підпис	Дата		

Кластери пов'язаних контейнерних хостів формуються з окремих контейнерних хостів. Кожен кластер складається з численних вузлів (хостів). Засоби програми - це розумні кластери контейнерів, що належать до однакового образу. Масштабування програми на багатьох вузлах хоста здійснюється за допомогою служб програм. Вони використовуються для забезпечення методів збереження в програмах, які їх потребують. Ці вони можуть бути змонтовані для зберігання в контейнерах. Використання посилань забезпечує підключення та зв'язок двох або більше контейнерів. Для розгортання та управління цими контейнерними кластерами необхідне забезпечення оркестрації для міжконтейнерних операторів, підключень та складання забезпечення.

Оркестрація контейнерів передбачає більше, ніж просто запуск, зупинку та переміщення програм (тобто контейнерів) між серверами. Процес створення та безперервної підтримки груп комп'ютерних програм на основі контейнерів, які можуть бути географічно розкидані, відомий як оркестрація. Оркестрація контейнерів дозволяє клієнтам вказувати, як контейнери в Хмарі будуть координуватися під час встановлення запиту на багато контейнерів [29]. Оркестрація контейнерів – це процес розгортання та підтримки контейнерів як єдиного цілого. Він гарантує, що контейнери завжди доступні, масштабовані та пов'язані. Побудова контейнерів у хмарі – це просто форма оркестрації в розподіленому хмарному середовищі. Хмарну архітектуру можна уявити як багатошарову та розсіяну структуру з основними шарами інфраструктури, платформи та програмних додатків, розподіленими по кількох хмарних середовищах.

Контейнерні технології мають потенціал для допомоги. Таким чином, контейнерні технології є критично важливими для майбутнього адміністрування програм, особливо в контексті Cloud PaaS [30]. Контейнерні архітектури, які нещодавно набули популярності, можуть бути реалізовані в цій хмарній платформі. З огляду на цю зміну в архітектурному стилі, вторинні дослідження можуть допомогти практикам у виборі найкращої технології. (див. рис.2.4)

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						34
Змн.	Арк.	№ докум.	Підпис	Дата		

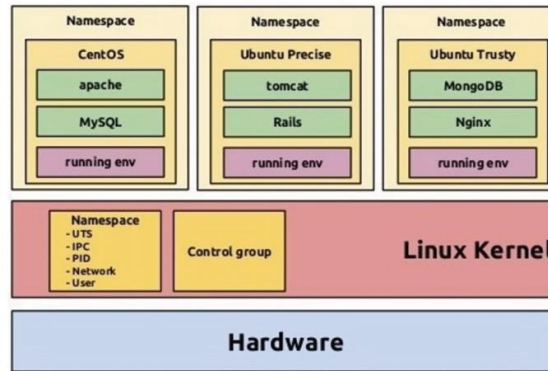


Рисунок 2.4 - Контейнери Linux або контейнери LXC

Віртуалізація ресурсів зазвичай передбачає встановлення програмного рівня поверх операційної системи хоста для керування численними ресурсами. Ці віртуальні машини (VM) можна розглядати як їхнє середовище виконання. У процесі віртуалізації використовується безліч методів. Віртуалізація на основі гіпервізора є одним із поширених підходів. KVM та VMware - дві відомі системи віртуалізації, засновані на гіпервізорах. Монітор віртуальної машини має бути розміщений поверх базової фізичної системи, щоб скористатися перевагами цієї технології. Крім того, кожна віртуальна машина підтримує (окремі) гостьові операційні системи. Ця стратегія віртуалізації є можливою для однієї операційної системи хоста для підтримки багатьох операційних систем відвідувачів.

Інший підхід представлений віртуалізацією на основі контейнерів. У цьому підході апаратні ресурси розділені шляхом реалізації багатьох (безпечних) ізоляцій. Гостьові процеси миттєво отримують абстракції за допомогою технологій на основі контейнерів, оскільки вони функціонують безпосередньо на рівні операційної системи (ОС) через рівень віртуалізації. Як правило, в системах на основі контейнерів одне ядро ОС об'єднане між віртуальними екземплярами. Користувачі сприймають контейнери як автономні операційні системи, здатні самостійно працювати з апаратним та програмним забезпеченням [31].

Функція ізоляції контейнерів обробляється просторами імен ядра. Це особлива властивість ядра Linux, яка дозволяє розробникам досягати необхідних етапів абстракції. Хоча контейнери не взаємодіють зовні з рівнем простору імен, існує розділення між хост-ОС та запрошеними процесами, і окремий контейнер

має свою операційну систему. Бідерман (EW Biederman, 2006) стверджує, що простори імен розділяють файлові системи, ідентифікатори процесів, мережі та міжпроцесну комунікацію. Однак рішення для віртуалізації на основі контейнерів накладають обмеження на використання ресурсів відповідно до груп процесів.

Якщо бути точнішим, у віртуалізації на основі контейнерів, групи с відповідають за розподіл пріоритетів використання процесора, пам'яті та вводу-виводу. Однак деякі технології, що використовують контейнери, послідовно - керують своїми ресурсами за допомогою груп с. Використовуючи рішення на основі контейнерів, можна динамічно розгортати та використовувати мікросервіси в умовах пакетного хостингу. Шаблони мікросервісів не є новою концепцією у світі архітектури програмного забезпечення. Зараз вони загальновизнані як економічно ефективний метод розробки додатків. До розробки мікросервісів стандартним підходом до розробки сервісів було переважно створення монолітних систем. З практичної точки зору, це вимагає створення єдиної платформи, здатної керувати всім. Багато з цих проблем можна вирішити в хмарних умовах за допомогою методів сценаріїв, що підтримують інфраструктуру як послугу (IaaS), платформу як послугу (PaaS) та програмне забезпечення як послугу (SaaS). Однак такі пояснення стикаються з труднощами, коли задіяно кілька спеціальних планів та груп, кожна з яких має свої різні вимоги до програмного забезпечення та інформації, як у випадку, коли образи віртуальних апаратів, специфічні для спільноти, недоступні для хмари. Легкі хмарні рішення є перевагами в таких умовах. Мікросервіси є прикладом такої архітектури.

Архітектура мікросервісів базується на ідеї «розділи та володарюй». Мікросервіси принципово замінюють єдину велику базу коду меншими базовими кодами, якими керують невеликі, гнучкі команди. Ці бази коду взаємодіють через єдиний API. Переваги цього підходу полягають у тому, що кожна група може працювати незалежно та бути захищеною/відключеною одна від одної - процес, який називається циклами винятків. Однак у певних випадках, наприклад, коли сервіси різних організацій взаємозалежні, ці цикли винятків можуть бути

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						36
Змн.	Арк.	№ докум.	Підпис	Дата		

пов'язані. Наразі існує безліч постачальників мікрохостингу на основі контейнерів. Docker, CoreOS та LXC є найвідомішими з них. Docker спрощує процес інкапсуляції програми в контейнер разом з будь-якими необхідними аксесуарами для роботи.

Це досягається за допомогою цілеспрямованого набору інструментів та комбінованих знань інтерфейсу прикладного програмування, що включають структури рівня ядра, такі як контейнери Linux, кластери керування та файлова структура копіювання під час запису. Збираючи файлову структуру контейнера, Docker спирається на розширену багатошарову файлову систему об'єднання (AMULFS) (AuFS) [32]. AuFS підтримує точне накладання однієї або кількох доступних файлових систем. Це дозволяє Docker використовувати образи, необхідні для основи контейнера. Наприклад, користувач може використовувати один образ Ubuntu як основу для багатьох інших контейнерів. Docker використовує одну копію Ubuntu, використовуючи розширену багатошарову схему об'єднання файлів. Це значно зменшує обсяг використовуваного сховища та обсяг оперативної пам'яті, що використовується контейнерами, завдяки їх швидкому запуску. AuFS також пропонує ще одну перевагу: він підтримує образи. Кожне нове видання позначається тегом diff, що є функцією впорядкування файлів, яка вказує на різницю між двома записами. Це дозволяє, наприклад, зменшити кількість файлів зображень.

Ще однією перевагою AuFS є те, що вона дозволяє відстежувати зміни, внесені до версій образів, подібно до того, як працюють схеми керування версіями коду розширення програмного забезпечення. CoreOS - це абсолютно нова версія Linux, створена для роботи з безліччю програмних структур. Цей метод дозволяє використовувати спрощене ядро Linux з метою максимально мінімізувати витрати [33]. Крім того, CoreOS надає можливості для забезпечення резервування та захисту від системних збоїв. CoreOS анонсувала оригінальне середовище виконання контейнерів Rocket (rkt), яке реалізує стандарт контейнерів відправлення. Основною метою цієї технології є створення налаштованої моделі контейнера.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						37
Змн.	Арк.	№ докум.	Підпис	Дата		

Крім того, цей метод підтримує образи машин Amazon. Середовище виконання контейнера Rocket є альтернативою Docker, включаючи покращену безпеку та інші вимоги для серверного виробництва. Середовище виконання контейнера Rocket відповідає стандарту Submission Container, який визначає нову колекцію презентацій контейнерів, що роблять їх легко транспортованими або переміщуваними. Docker виконує кожен процес через демон, який не пропонує такого ж рівня гарантії безпеки, як Rocket. Щоб вирішити цю проблему, деякі запропонували повністю переробити Docker.

LXC інтегровано у стандартне ядро Linux, і проект дозволяє керувати образами контейнерів та операційних систем за допомогою інструментарію. Контейнери можна розглядати як легкі операційні системи, що працюють разом з основною операційною системою. Оскільки контейнери не емулюють апаратну плівку, вони можуть працювати майже на власних швидкостях без додаткових накладних витрат на продуктивність [34]. У звичайному режимі роботи програми та веб-завантаження створюються та встановлюються в певній конфігурації на серверах без підзарядки з певних причин. Наприклад, можлива одночасна установка та налаштування PHP, MySQL, Nginx та Drupal. Однак наразі програми встановлюються на певний комп'ютер і їх не можна легко перемістити. Віртуальний комп'ютер можна встановити та перемістити, і хоча це створює ілюзію портативності, в результаті страждає продуктивність. Контейнер LXC забезпечує майже продуктивність без підзарядки та гнучкість для легкого переміщення між системами шляхом інкапсуляції порівнянних стеків у контейнери. Підвищена швидкість та гнучкість контейнера LXC створюють видимість виділеного сервера. Можливі клони, резервні копії та знімки контейнерів [35]. LXC спрощує керування контейнерами та додає нові рівні свободи для виконання та розгортання програм. Flockport дозволяє швидко розгортати веб-стеки та програми, упаковані в контейнери LXC, на будь-якому комп'ютері на базі Linux. Гнучкість та пропускну здатність безпосередньо підтримуються контейнерами LXC, тоді як Flockport - це механізм для розподілу та використання контейнерів LXC.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						38
Змн.	Арк.	№ докум.	Підпис	Дата		

Численні науковці зосередилися на зменшенні розриву в продуктивності та оптимізації між віртуалізацією та невіртуалізованими методами. Як правило, ці дослідження використовують унікальний підхід та набір інструментів. Крім того, різняться склад методів, що вивчалися та порівнювалися. Наприклад, нещодавні дослідницькі статті аналізують та порівнюють розрив між власними системами та їхніми передбачуваними аналогами, а також відмінності у представленні методів віртуалізації на основі контейнерів та гіпервізорів. Ми визнаємо, що це галузь, що швидко розвивається, і як наслідок, деякі зі старіших статей використовують застарілі інструменти. Крім того, сучасна технологія візуалізації не вивчалася. У нещодавній статті оцінили чотири різні системи віртуалізації на основі гіпервізора. Автори не знайшли гіпервізора, який би працював набагато краще. Як наслідок, вони рекомендували клієнтам використовувати різні програмні та апаратні етапи в хмарних сервісах для задоволення своїх потреб [36].

Порівняно презентацію VMware, Microsoft Hyper -V та Citrix Xen у різних ситуаціях. Автори використовували модифікований метод екземплярів SQL для проведення оцінки. Вони створили мільйони товарів, споживачів та замовлень за допомогою цього методу.

Хоча докази щодо накладних витрат віртуалізації були суперечливими встановлено, що плівка віртуалізації практично для кожного гіпервізора суттєво вплинула на представлення віртуалізованих систем, особливо для високопродуктивних обчислювальних сфер. У сучасних статтях порівнюються та протиставляються гіпервізори з контейнерними методами. За даними, контейнери набувають популярності для забезпечення можливостей PaaS. Estrada et al. провели порівняльний аналіз технологій віртуалізації на основі KVM та Xen, зокрема KVM, Xen та LXC [37]. Їхній основний підхід полягав у порівнянні та протиставленні продуктивності кожної з перелічених ними технологій. Їхні тести та бенчмарки були мотивовані метою сприяння розробці серійних додатків.

Heroku, компанія-розробник платформи як послуги, стала піонером у використанні контейнерів для професійної та повторюваної організації додатків. Замість того, щоб розглядати контейнер як сервер, згенерований комп'ютером,

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						39
Змн.	Арк.	№ докум.	Підпис	Дата		

Heroku визначає його як процедуру з покращеними характеристиками розділення. Як результат, розгортання додатків на основі контейнерів забезпечує рішення з низьким навантаженням та майже такою ж ізоляцією, як і віртуальні машини. Як і звичайні процеси, воно також включає можливості спільного використання ресурсів. Інфраструктура Google значною мірою спиралася на такі контейнери. Крім того, Docker надав стандартизований формат образу для контейнерів додатків та інструменти адміністрування для них.

2.4 Висновки по розділу

У другому розділі було досліджено методи та інструменти контейнеризації. Розглянуто використання Docker як базового інструменту для створення, ізоляції та розгортання контейнерів. Проаналізовано практичні сценарії використання контейнерів для розгортання додатків у різних середовищах, а також особливості побудови контейнерних архітектур. Окрему увагу приділено підходам до організації контейнеризованих систем, що забезпечують гнучкість, масштабованість і спрощення процесів супроводу програмного забезпечення. У результаті дослідження визначено ефективні підходи до використання контейнеризації у сучасних ІТ-проєктах.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						40
Змн.	Арк.	№ докум.	Підпис	Дата		

РОЗДІЛ 3. РОЗРОБКА, РОЗГОРТАННЯ ТА ОЦІНКА ЕФЕКТИВНОСТІ КОНТЕЙНЕРИЗОВАНИХ РІШЕНЬ

3.1 Керування контейнерами за допомогою Docker Compose

Зі зростанням кількості контейнерів у системі керування ними окремо стає дедалі складнішим. Як видно з попереднього прикладу використання двоконтейнерної системи, кожен контейнер вимагає окремих конфігурацій для різних аспектів, таких як порти, томи та мережі. Таке розповсюдження конфігурацій може швидко стати складним для обробки, що призводить до потенційних помилок та труднощів із підтримкою узгодженості між кількома контейнерами. Для вирішення цих проблем та оптимізації управління контейнерними додатками *Docker Compose* пропонує рішення.

Docker Compose - це інструмент, який спрощує розгортання та керування багатоконтейнерними застосунками. Він дозволяє розробникам визначати та оркеструвати кон-

фігурація кількох контейнерів в одному конфігураційному файлі, відомому як файл *Compose*. Цей файл служить централізованим місцем для визначення бажаного стану програми, включаючи залежності контейнерів, мережеві підключення, томи та змінні середовища, серед іншого.

За допомогою *Docker Compose* визначення контейнерів та пов'язані з ними конфігурації можуть бути інкапсульовані в одному файлі *Compose*. Це усуває необхідність налаштовувати кожен контейнер окремо вручну, зменшуючи складність та потенційну можливість помилок.

Однією з ключових можливостей *Docker Compose* є його здатність визначати та керувати зв'язками між контейнерами. Залежності між контейнерами, такі як залежність одного контейнера від доступності іншого, можна виразити у файлі *Compose*. Потім *Docker Compose* гарантує, що контейнери запускаються в правильному порядку, піклуючись про зв'язок та синхронізацію між контейнерами.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						41
Змн.	Арк.	№ докум.	Підпис	Дата		

Крім того, Docker Compose спрощує керування мережами, томами та змінними середовища в різних контейнерах. Він забезпечує єдиний інтерфейс для визначення та налаштування цих ресурсів, дозволяючи розробникам централізовано вказувати мережеві підключення, монтувати томи та встановлювати змінні середовища для всіх контейнерів. Це не лише зменшує зусилля, необхідні для управління цими аспектами, але й покращує узгодженість та відтворюваність середовища програми.

Наприклад, замість того, щоб запускати контейнери окремо та вручну вказувати кожен з їхніх конфігурацій та параметрів, таких як змінні середовища, правила переадресації портів, томи та мережі, ми можемо написати один файл створення (див. Лістинг 3.1) , використовуючи синтаксис YAML, який обробляє вищезгадані змінні. Файли Docker-compose складаються з взаємопов'язаних елементів, які разом визначають структуру та поведінку багатоконтейнерного застосунку. Примітно, що вони складаються з сервісів, томів та мереж. У випадку лістингу 3.1 визначено два основні сервіси: **fib-db** та **fib-server** , кожен з яких представляє окремий контейнеризований компонент. Сервіс **fib-db** містить контейнер, на якому запущено контейнер MariaDB, тоді як **fib-server** відповідає за контейнер, на якому запущено основну програму. Шляхи до відповідних каталогів Dockerfile (контексти збірки) визначаються у властивості build сервісу. Крім того, необхідні змінні середовища (властивість environment), томи (властивість volumes), порт для правил переадресації (властивість ports) та мережі (властивість networks) також визначаються для кожного сервісу окремо. Однак, щоб мати змогу вказати томи та мережі, їх необхідно визначити відповідно в розділах томів та мереж [38].

У сфері сучасної розробки програмного забезпечення контейнеризація стала популярним підходом до розгортання та доставки програмного забезпечення. Переваги та недоліки контейнеризації були обговорені в попередньому розділі, і цей розділ буде базуватися на цьому обговоренні, досліджуючи застосування контейнеризації в контексті реальної програмної системи у формі тематичного дослідження. Тематичне дослідження зосереджено

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						42
Змн.	Арк.	№ докум.	Підпис	Дата		

на eShare, веб-програмному продукті для управління інформацією, який був більш детально описаний у розділі.

```
1 # docker-compose.yml
2 version: "3.8"
3 services:
4   fib-db:
5     build:
6       context: ./fib-db
7     environment:
8       - MYSQL_ROOT_PASSWORD=fib-admin
9       - MYSQL_DATABASE=fib-db
10    volumes:
11      - fib-vol:/var/lib/mysql fib-db
12    ports:
13      - 3306:3306
14    networks:
15      - fib-net
16  fib-server:
17    build:
18      context: ./fib-server
19    ports:
20      - 7000:8080
21    networks:
22      - fib-net
23  volumes:
24    fib-vol:
25  networks:
26    fib-net:
```

Лістинг 3.1 - Файл Docker-compose для налаштування багатоконтейнерного додатка.

Цей розділ розпочнеться зі вступу до тематичного дослідження, в якому будуть окреслені цілі та мотивація дослідження. Далі буде надано технічний огляд архітектури eShare та потенційних проблем і вузьких місць, які можуть виникнути в процесі контейнеризації. Далі в розділі буде представлено рішення для контейнеризації eShare, включаючи план та реалізацію цього рішення. Розділ завершиться зауваженнями щодо реалізації та пропозиціями для майбутньої роботи.

Основна мета цього дослідження є двоякою: по-перше, провести та оцінити доцільність часткової контейнеризації eShare з метою отримання вигоди від переваг контейнеризації, визначених у попередньому розділі, а по-друге, підготувати основу для майбутнього переходу до хмарного розгортання та моделі програмного забезпечення як послуги (SaaS), яку Cadmatic досліджує та реалізує як потенційну майбутню бізнес-модель для eShare. Остання є довгостроковою

					БР.ІІІ – 22.00.00.000 ПЗ	Арк.
						43
Змн.	Арк.	№ докум.	Підпис	Дата		

метою, яка не входить до обсягу цього дослідження, але тим не менш є важливим фактором у контексті даного тематичного дослідження.

Центральним для цих цілей є оцінка контейнерного рішення в контексті eShare, зокрема оцінка його продуктивності та визначення переваг і обмежень, властивих процесу контейнеризації.

Наразі eShare періодично стикається з операційними вузькими місцями, головним чином через певні обчислювально ресурсоємні серверні завдання. Ці завдання, такі як *обробка хмари точок, введення посилань на документи, перетворення документів та публікація моделей*, страждають від низької продуктивності, що іноді негативно впливає на взаємодію з користувачем. Цю проблему, звичайно, можна пом'якшити, масштабуючи інфраструктуру вертикально (збільшуючи потужність існуючої системи), але це не є розумним рішенням з точки зору економічної ефективності, особливо в хмарних середовищах, куди eShare планує перейти в майбутньому, оскільки eShare не вимагає високого рівня апаратних ресурсів більшу частину часу. Але коли це іноді вимагає, поточна інфраструктура не дуже добре підходить для обробки цих піків попиту.

Основна мета цього дослідження є двоякою: по-перше, провести та оцінити доцільність часткової контейнеризації eShare з метою отримання вигоди від переваг контейнеризації, визначених у попередньому розділі, а по-друге, підготувати основу для майбутнього переходу до хмарного розгортання та моделі програмного забезпечення як послуги (SaaS), яку Cadmatic досліджує та реалізує як потенційну майбутню бізнес-модель для eShare [39]. Остання є довгостроковою метою, яка не входить до обсягу цього дослідження, але тим не менш є важливим фактором у контексті даного тематичного дослідження.

Центральним для цих цілей є оцінка контейнерного рішення в контексті eShare, зокрема оцінка його продуктивності та визначення переваг і обмежень, властивих процесу контейнеризації.

Наразі eShare періодично стикається з операційними вузькими місцями, головним чином через певні обчислювально ресурсоємні серверні завдання. Ці

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						44
Змн.	Арк.	№ докум.	Підпис	Дата		

завдання, такі як *обробка хмари точок*, *введення посилань на документи*, *перетворення документів* та *публікація моделей*, страждають від низької продуктивності, що іноді негативно впливає на взаємодію з користувачем. Цю проблему, звичайно, можна пом'якшити, масштабуючи інфраструктуру вертикально (збільшуючи потужність існуючої системи), але це не є розумним рішенням з точки зору економічної ефективності, особливо в хмарних середовищах, куди eShare планує перейти в майбутньому, оскільки eShare не вимагає високого рівня апаратних ресурсів більшу частину часу. Але коли це іноді вимагає, поточна інфраструктура не дуже добре підходить для обробки цих піків попиту.

Сервер eShare - це монолітний застосунок, написаний на C# та працює на платформі .NET. Сервер відповідає за обробку бізнес-логіки, обробку даних та зв'язок із зовнішніми системами. Сервер також відповідає за обслуговування клієнтських застосунків даними, які вони запитують. Для доступу до зовнішніх систем eShare надає набір налаштовуваних *адаптерів джерел даних* для взаємодії з різними типами джерел даних, такими як бази даних, REST API та файли електронних таблиць. Сервер також виконує періодичні фонові завдання, такі як індексування документів та обробка хмар точок, які часто є ресурсоемними з точки зору обчислень і можуть бути вузьким місцем для системи. Сервер розміщено на IIS (*Інтернет-інформаційні служби*) та використовує Microsoft SQL Server як систему керування базами даних. Загальний огляд серверної архітектури eShare показано на рисунку 3.1.

Оскільки більшість важких завдань виконує сервер, зусилля з контейнеризації будуть зосереджені на компонентах та сервісах, що працюють на сервері.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						45
Змн.	Арк.	№ докум.	Підпис	Дата		

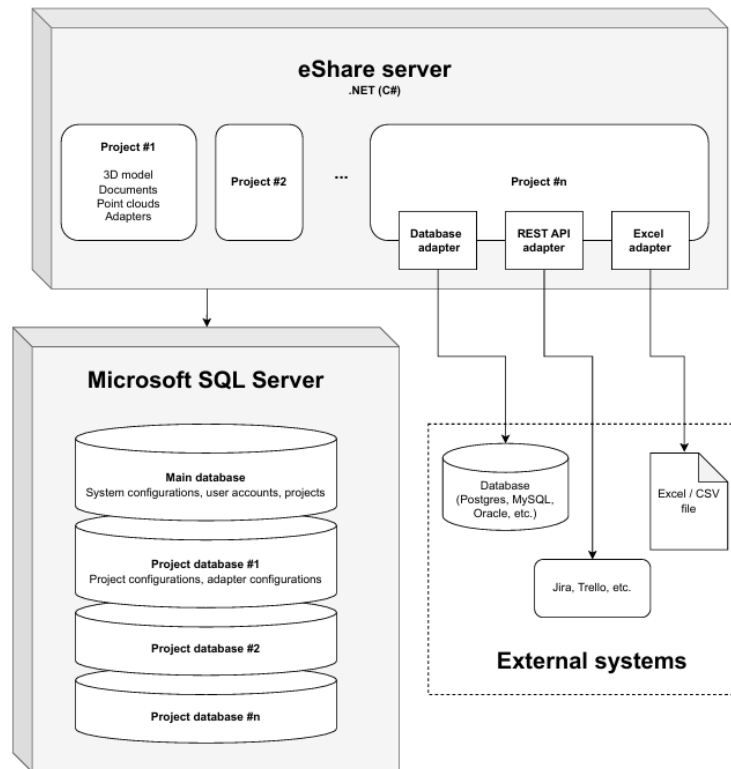


Рисунок 3.1 - Загальний огляд серверної архітектури eShare

Клієнтський застосунок eShare - це веб-орієнтований односторінковий застосунок, написаний на TypeScript та Angular. Клієнтський застосунок відповідає за візуалізацію даних та надання користувацького інтерфейсу для взаємодії з даними. 3D-візуалізація моделей та хмар точок обробляється за допомогою власної технології візуалізації Cadmatic, написаної на C++ та скомпільованої у Web Assembly за допомогою Emscripten. Зв'язок між клієнтом та сервером здійснюється за допомогою REST API, що надається eShare сервер.

Для розгортання eShare потрібне середовище Windows, оскільки воно розміщене на IIS (Інтернет-інформаційні служби) та використовує Microsoft SQL Server як систему керування базами даних. Наразі eShare найчастіше розгортається локально, для кожного клієнта окремо. Хмарне розгортання також можливе за допомогою віртуальної машини на базі Windows, але поки що це не дуже поширене серед клієнтів eShare.

Перехід на контейнери викликає деякі питання щодо потенційних труднощів, які можуть виникнути під час процесу контейнеризації. Зокрема,

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						46
Змн.	Арк.	№ докум.	Підпис	Дата		

оскільки Docker залежить від ядра Linux, той факт, що сервер eShare є додатком на базі Windows, може створювати певні труднощі. Однак платформа .NET, на якій побудовано eShare, наразі здебільшого є кросплатформною завдяки появі .NET Core. Це означає, що кількість залежностей Windows має бути мінімальною або навіть нульовою, що має полегшити перехід на контейнери.

Ще однією можливою проблемою є обробка великих файлів та їх передача до контейнерів і з них. Наприклад, розмір деяких файлів хмари точок може перевищувати 100 гігабайт. Якщо контейнер, до якого призначено файл хмари точок, працює на тому ж хості, що й сервер eShare, проблема полягає лише в призначенні відповідного тома для монтування до контейнера. Однак, якщо контейнер знаходиться на зовсім іншому хості, проблему може бути складніше вирішити належним чином.

Окрім проблем, характерних для існуючої архітектури eShare, аналіз літератури в розділі 4 також виділив деякі потенційні проблеми, які можуть виникнути в процесі контейнеризації. Зокрема, складність керування та оркестрації контейнерів, а також наслідки контейнеризації для безпеки – це два аспекти, які необхідно враховувати. Однак, це тематичне дослідження буде зосереджено на контейнеризації лише частини сервера eShare, і таким чином проблеми, пов'язані зі складністю керування та оркестрації контейнерів, мають бути досить обмеженими.

Як було виявлено в попередньому розділі, контейнеризація разом з архітектурою мікросервісів пропонує перспективне рішення для проблем і вузьких місць. Наприклад, підвищена продуктивність і масштабованість були визначені як найчастіше згадувані переваги контейнеризації в аналізі літератури, і це саме ті проблеми, з якими eShare зараз стикається щодо певних серверних завдань і служб. Інкапсулюючи ці серверні завдання, служби та обов'язки в контейнери в мікросервісі, програма може досягти більш точно налаштованих можливостей масштабування та спростити перехід між локальним та хмарним розгортанням [40]. Використання контейнерів та архітектури мікросервісів не лише пропонує потенційне рішення для безпосередніх операційних вузьких місць,

але й закладає основу для більш економічно ефективної, зручної в обслуговуванні та масштабованої системної архітектури в майбутньому.

3.2 Проєктування та впровадження рішення для контейнеризації додатку

У багатьох випадках, описаних у аналізі літератури, зазначені переваги контейнеризації щодо продуктивності та масштабованості були досягнуті завдяки використанню архітектури мікросервісів. Таким чином, після ретельної оцінки результатів аналізу літератури та обговорення з деякими розробниками, які працюють з eShare, було вирішено, що найкращим підходом до вирішення проблем продуктивності за допомогою контейнеризації буде створення незалежно масштабованих мікросервісів, які інкапсулюють обчислювально ресурсоємні завдання.

Як відправну точку, ціллю контейнеризації було обрано серверні завдання, пов'язані з читанням, записом та конвертацією документів з одного формату в інший. Ці завдання наразі реалізовані як незалежні автономні виконувані файли та відносно відокремлені від решти системи, тому їх вилучення з eShare та контейнеризація мають бути досить простими.

Наразі ці завдання обробки документів запускаються щоразу, коли користувач відкриває документ для перегляду в eShare. Залежно від апаратного забезпечення сервера та складності документа, виконання цих завдань може тривати кілька секунд, що часто суттєво негативно впливає на взаємодію з користувачем. Вирішення цієї проблеми за допомогою завдань обробки документів значно покращить взаємодію з користувачем eShare.

Початковий план полягає в тому, щоб створити мікросервіс, який інкапсулює ці завдання в контейнеризований застосунок. Мікросервіс потім буде розгорнуто як незалежний сервіс, який відповідатиме за обробку документів. Щоразу, коли запитується обробка документа, сервер eShare надсилатиме повідомлення мікросервісу обробки документів, який обробить документ і в

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						48
Змн.	Арк.	№ докум.	Підпис	Дата		

результаті поверне відповідні дані на сервер eShare.

План також полягає в тому, що в робочому середовищі мікросервіс обробки документів буде розгорнутий на окремій потужнішій інфраструктурі, що зменшить навантаження, яке має обробляти один екземпляр сервера eShare. Такий підхід також дозволить мікросервісу обробки документів обслуговувати кілька екземплярів eShare одночасно, що підвищить економічну ефективність системи в цілому, оскільки час простою потужнішого та дорожчого обладнання мікросервісу обробки документів буде мінімізовано.

Після обговорення деталей потенційного рішення для контейнеризації з деякими розробниками eShare було вирішено, що для зв'язку між сервером eShare та мікросервісом обробки документів буде використовуватися механізм зв'язку на основі подій. Зв'язок на основі подій – це шаблон, за якого відправник повідомлення (події) не повинен знати одержувача повідомлення. Натомість відправник публікує повідомлення на шині повідомлень, а одержувач підписується на шину повідомлень для отримання повідомлення. Цей шаблон добре підходить для запропонованого рішення для контейнеризації, оскільки він зменшує зв'язок та дозволяє чітко розділити сервер eShare та мікросервіс обробки документів (і будь-які майбутні мікросервіси), що ідеально підходить для розподіленого середовища.

Щоразу, коли запитується обробка документа, сервер eShare може опублікувати повідомлення на призначеній шині повідомлень, а мікросервіс обробки документів, у свою чергу, може підписатися на цю шину для отримання повідомлення та обробки документа. Це особливо корисно, коли мікросервіс обробки документів масштабується горизонтально до кількох екземплярів, оскільки повідомлення можуть розподілятися між екземплярами брокером повідомлень, що зменшує складність керування кількома мікросервісами.

Для цієї мети обрана технологія шини повідомлень Apache Kafka, яка є розподіленою платформою потокової передачі подій, що добре підходить для такого типу комунікації на основі подій. Kafka є популярним вибором для такого типу комунікації, оскільки вона має високу відмовостійкість і розроблена для

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						49
Змн.	Арк.	№ докум.	Підпис	Дата		

високої пропускної здатності.

Архітектуру запропонованого рішення для контейнеризації візуалізовано на рисунку 3.2 .

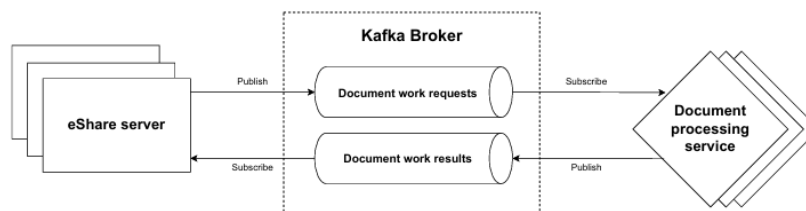


Рисунок 3.2 - Загальний огляд запропонованого рішення з контейнеризації для обробки документів у системі eShare.

Запропоноване рішення базується на брокері повідомлень Kafka, який використовується для зв'язку на основі подій між сервером eShare, мікросервісом обробки документів та будь-якими іншими потенційними майбутніми мікросервісами. Екземпляри серверів eShare використовують мікросервіси, публікуючи повідомлення в певних темах Kafka, які потім споживаються (підписуються) відповідними мікросервісами. І навпаки, мікросервіси публікують повідомлення назад в окремі теми Kafka, які споживаються сервером eShare для отримання результатів обробки. Потік запиту на обробку документа, починаючи з клієнтської програми eShare і закінчуючи передачею документа клієнту, візуалізовано у вигляді діаграми послідовності на рисунку 3.3.

Запити на обробку документів та результати обробки документів - це теми Kafka, до яких сервер eShare та мікросервіс обробки документів публікують повідомлення та підписуються на них.

Першим кроком у процесі контейнеризації було впровадження - мікросервісу обробки документів, який інкапсулює завдання обробки документів у контейнеризований додаток. Мікросервіс обробки документів був реалізований як додаток .NET, що є тією ж платформою, на якій побудовано сервер eShare. Цей вибір був зроблений через знайомство розробників з платформою та легкість інтеграції з існуючою кодовою базою eShare.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						50
Змн.	Арк.	№ докум.	Підпис	Дата		

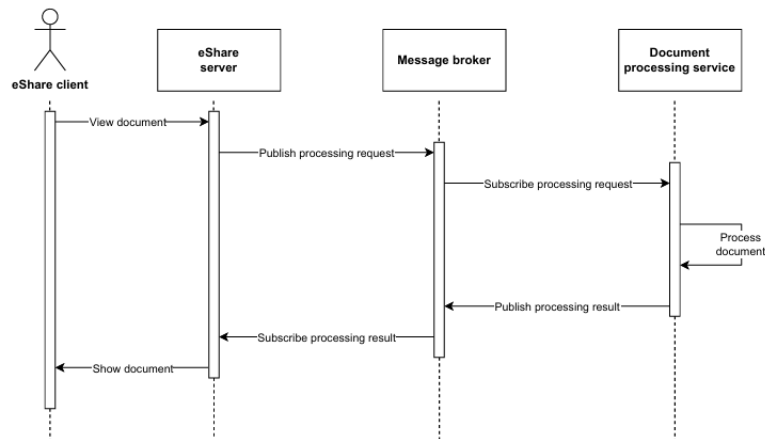


Рисунок 3.3 - Діаграма послідовності, що ілюструє хід обробки запиту на обробку документа в запропонованому рішенні з контейнеризації.

Значною розробкою стало створення абстрактного базового класу, призначеного для оптимізації розробки мікросервісів на основі Kafka. Ця абстракція спрощує розробку та розгортання мікросервісів, забезпечуючи загальну структуру та набір функціональних можливостей, необхідних для роботи мікросервісів в екосистемі eShare. Мікросервіс обробки документів потім був реалізований шляхом розширення цього абстрактного базового класу та забезпечення необхідної функціональності для обробки документів. Це проілюстровано на діаграмі класів на рисунку 3.4 .

Такий підхід абстрагує деталі комунікації Kafka від розробника мікросервісу, дозволяючи йому зосередитися на фактичній логіці обробки документів. Абстрактний базовий клас обробляє комунікацію Kafka, включаючи серіалізацію та десеріалізацію повідомлень, і надає простий інтерфейс для розробника мікросервісів для реалізації фактичної логіки обробки. Це також зменшує ризик прив'язки до постачальника, оскільки базову реалізацію шини повідомлень можна легко замінити на іншу, якщо це необхідно.

Наступним кроком було модифікація та рефакторинг можливостей обробки документів сервера eShare для підтримки нового подієвого підходу. Це включало реалізацію подієвої обробки запитів на обробку документів, коли сервер eShare публікує повідомлення на шину повідомлень щоразу, коли запитується обробка документа, та підписується на іншу шину повідомлень для отримання

результатів обробки. Ця система розроблена агностичною до фактичної обробки документів, що дозволяє полегшити зв'язок та інтеграцію мікросервісів у розподіленому середовищі.

Ця конфігурація, незалежність від реалізації обробки документів, також - дозволяє eShare працювати у двох режимах на основі параметра конфігурації. Якщо з якоїсь причини мікросервіс обробки документів не потрібен, eShare можна налаштувати для локальної обробки документів, як це було до процесу контейнеризації. Це досягається шляхом створення двох різних реалізацій шини повідомлень в eShare, які відповідають одному і тому ж інтерфейсу, але працюють по-різному:

- Шина повідомлень, яка взаємодіє з брокером Kafka, що використовується для обробки мікросервісів.
- Шина повідомлень, реалізована за допомогою коду C# та функціонує локально, подібно до процесу контейнеризації.

Відповідна реалізація шини повідомлень визначається під час виконання на основі конфігурації сервера. Таким чином, використання мікросервісу обробки документів може бути увімкнено адміністратором цього конкретного екземпляра сервера eShare, а не нав'язано йому.

Однією з головних проблем, що виникли під час впровадження, була відсутність підтримки PDF-бібліотеки eShare в Linux. PDF-бібліотека, що використовується eShare, підтримується лише у Windows та macOS, а це означало, що мікросервіс обробки документів мав запускатися в контейнері Windows. Це не ідеально для хмарного середовища, оскільки вимагає, щоб хост-ОС була Windows, що обмежує гнучкість рішення. Однак для цілей цієї роботи та тематичного дослідження використання контейнерів Windows було визнано прийнятним.

Оптимальний шлях розвитку передбачає перехід на контейнери Linux та впровадження кросплатформної бібліотеки PDF, що дозволить більш хмарне розгортання. Використання контейнерів Windows розглядається як тимчасове рішення для демонстрації доцільності процесу контейнеризації та підтвердження концепції для майбутньої розробки.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						52
Змн.	Арк.	№ докум.	Підпис	Дата		

Побудова мікросервісу обробки документів включала створення універсального абстрактного класу **KafkaWorkerService<TRequest, TResponse>**, який слугує основою для мікросервісів на базі Kafka в екосистемі eShare. Цей абстрактний клас обробляє зв'язок Kafka та забезпечує основу для реалізації фактичної логіки обробки. Розробник мікросервісу повинен розширити цей клас та реалізувати метод **HandleMessage**, який отримує вхідне повідомлення типу **TRequest** як вхідні дані та повертає відповідь типу **TResponse** як вихідні. Вхідні та вихідні дані передаються як повідомлення JSON через Kafka, що дозволяє легко серіалізувати та десеріалізувати дані. Абстрактний клас **KafkaWorkerService** по суті є розміщеним на ASP.NET

сервіс [35], який працює у фоновому режимі та прослуховує повідомлення на тему Kafka. Коли повідомлення отримано, викликається метод **HandleMessage**, і мікросервіс обробляє повідомлення та повертає відповідь.

Мікросервіс обробки документів було реалізовано шляхом створення класу **DocumentWorkerService**, який розширює клас **KafkaWorkerService** та забезпечує реалізацію методу **HandleMessage**. Діаграма класів на рисунку 5.4 ілюструє зв'язок між абстрактним класом **KafkaWorkerService**, класом **DocumentWorkerService** та будь-якими потенційними майбутніми мікросервісами, які можуть бути реалізовані за допомогою того ж фреймворку.

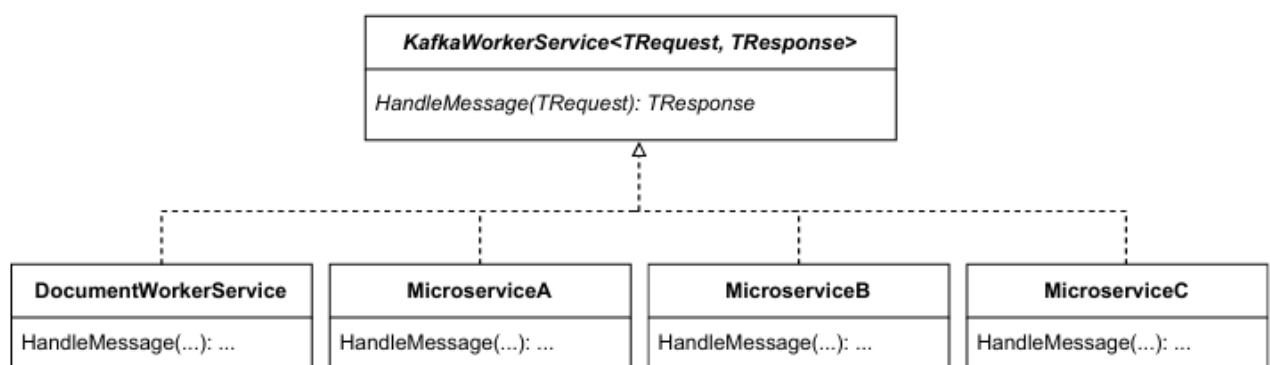


Рисунок 3.4 - Діаграма класів, що ілюструє взаємозв'язки між абстрактним класом **KafkaWorkerService** та реалізаціями мікросервісів

Kafka не дуже добре підходить для передачі великих двійкових файлів,

таких як PDF-документи, через обмеження розміру повідомлень та вплив на продуктивність. Таким чином, посилання на шлях до файлу документа передається через Kafka, а сам документ передається за допомогою спільної файлової системи, а не через конвеєр Kafka. У майбутньому для передачі та спільного використання фактичних файлів документів, можливо, можна було б використовувати окремий сервіс зберігання об'єктів, такий як *Amazon S3* або *Azure Blob Storage*, для створення більш надійного та масштабованого рішення. Однак, реалізація такого рішення була визнана поза межами цього тематичного дослідження, але є важливим фактором для подальшої розробки.

Щоб адаптуватися до нового мікросервісу обробки документів, сервер eShare було модифіковано для підтримки нової моделі комунікації на основі подій. У цій новій конструкції сервер eShare публікує повідомлення на шину повідомлень щоразу, коли запитується обробка документа, та підписується на іншу шину повідомлень для отримання результатів обробки. Наразі в eShare існує дві реалізації шини повідомлень: одна для локальної обробки документів, а інша для обробки мікросервісів. Відповідна реалізація визначається під час виконання на основі конфігурації сервера.

Як показано на діаграмі класів на рисунку 3.5 , обидві реалізації шини повідомлень базуються на одному й тому ж інтерфейсі *IMessageBus* , який забезпечує спільний інтерфейс для публікації та підписки на повідомлення. Локальна реалізація шини повідомлень, *MessageBus* , - це проста шина повідомлень в пам'яті, яка використовується для локальної обробки документів, тоді як реалізація шини повідомлень *Kafka*, *KafkaMessageBus* , використовується для обробки мікросервісів. Реалізація шини повідомлень *Kafka* використовує клієнтську бібліотеку *Confluent Kafka .NET* для зв'язку з брокером *Kafka* та внутрішню бібліотеку *JSON* для серіалізації та десеріалізації повідомлень.

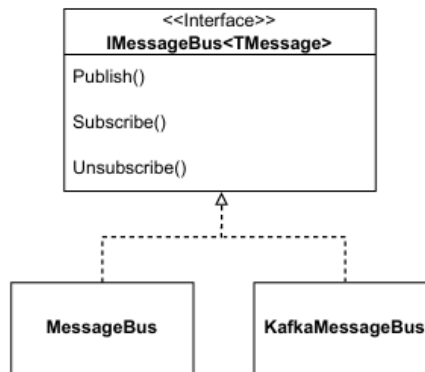


Рисунок 3.5 - Діаграма класів інтерфейсу IMessageBus та його реалізацій KafkaMessageBus і MessageBus.

Контейнеризація мікросервісу обробки документів передбачала створення Dockerfile, який визначає конфігурацію образу контейнера та залежності. Dockerfile для мікросервісу обробки документів базується на офіційних образах mcr.microsoft.com/dotnet/aspnet:8.0-windowsservercore-ltsc2019 та mcr.microsoft.com/dotnet/sdk:8.0-windowsservercore-ltsc2019, які надають базові образи для запуску та створення .NET-застосунків у Windows. Образ .NET SDK використовується для створення застосунку, тоді як образ .NET ASP.NET використовується для запуску застосунку. Використаний Dockerfile показано в лістингу 3.2 .

Оскільки отримане зображення є контейнером Windows, варто зазначити, що хост-ОС також має бути Windows, а Docker має бути налаштований на використання контейнерів Windows для запуску контейнера. Це обмеження поточної реалізації, оскільки воно обмежує можливості розгортання мікросервісу. Однак, як згадувалося раніше, це розглядається як тимчасове рішення, і план полягає в переході на контейнери Linux у майбутньому.


```

1 FROM mcr.microsoft.com/dotnet/aspnet:8.0-windowsservercore-
  ltsc2019 AS base
2 WORKDIR /app
3
4 FROM mcr.microsoft.com/dotnet/sdk:8.0-windowsservercore-ltsc2019
  AS build
5 WORKDIR /src
6
7 # Copy PdfWorkerService project
8 COPY ["/DocumentWorkerService", "/DocumentWorkerService"]
9
10 # Copy dependencies
11 COPY ["/KafkaWorkerService", "/KafkaWorkerService"]
12 COPY ["/Dto", "/Dto"]
13 COPY ["/Json", "/Json"]
14 COPY ["/VersionInfo.cs", "/VersionInfo.cs"]
15
16 # Restore and build
17 WORKDIR "/src/DocumentWorkerService"
18
19 RUN dotnet restore "/DocumentWorkerService.csproj"
20 RUN dotnet build "/DocumentWorkerService.csproj" -c Release -o /
  app/build /p:IsDockerBuild=true
21 # Publish
22 FROM build AS publish
23 RUN dotnet publish "/DocumentWorkerService.csproj" -c Release -o
  /app/publish /p:UseAppHost=false /p:IsDockerBuild=true
24
25 FROM base AS final
26 WORKDIR /app
27 COPY --from=publish /app/publish .
28 ENTRYPOINT ["dotnet", "DocumentWorkerService.dll"]

```

Лістинг 3.2 - Файл Dockerfile для мікросервісу обробки документів.

Впровадження мікросервісу обробки документів у контейнерах було - здебільшого успішним та функціональним. Мікросервіс може обробляти документи в контейнерному середовищі та взаємодіяти з сервером eShare через брокер повідомлень Kafka. Однак, перш ніж рішення можна буде вважати готовим до використання, необхідно вдосконалити кілька аспектів:

Аутентифікація та авторизація Kafka: Поточна реалізація спирається на те, що екземпляри сервера eShare безпосередньо підключаються до брокера Kafka за допомогою єдиного спільного ключа API. Перш ніж переходити до робочого середовища, важливо використовувати більш надійний та детальний механізм автентифікації та авторизації для Kafka. Слід забезпечити, щоб лише уповноважені сторони могли публікувати та підписуватися на відповідні теми Kafka, а також щоб конфіденційні дані були захищені. Цього можна було б досягти шляхом більш детального надання ключів API або шляхом маршрутизації

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						56
Змн.	Арк.	№ докум.	Підпис	Дата		

трафіку Kafka через безпечний шлюз замість безпосереднього підключення до брокера.

Міжплатформна сумісність : Поточна реалізація обмежена Windows.

контейнери dows через залежність від Pdfium.Net SDK для обробки PDF. Перехід на кросплатформну бібліотеку для маніпулювання PDF дозволить використовувати контейнери Linux, які більше підходять для хмарних середовищ.

Автоматичне масштабування та оркестрація: Поточна реалізація не підтримує автоматичне масштабування мікросервісу обробки документів. Реалізація автоматичного масштабування мікросервісу на основі навантаження була б корисною для ефективної обробки різних робочих навантажень та зниження витрат, коли мікросервіс використовується недостатньо. Такий функціонал можна було б реалізувати, наприклад, за допомогою Kubernetes [36] або Azure Container Apps [37].

Служба зберігання об'єктів: наразі файли документів передаються до контейнера через спільну файлову систему, яку може бути не найпростішою в налаштуванні та обслуговуванні. Використання замість неї служби зберігання об'єктів, такої як Amazon S3 або Azure Blob Storage, можливо, може стати предметом подальших досліджень та розробок.

3.3 Оцінювання ефективності та масштабованості контейнеризованого рішення

Найбільш поширені переваги та недоліки, пов'язані з контейнеризацією, використовуються як критерії для оцінки запропонованого рішення контейнеризації для eShare. Переваги та недоліки такі:

- Накладні витрати : сума накладних витрат, що виникають внаслідок нового запропонованого рішення ція оцінюється на основі часу, необхідного для обробки документів у новій архітектурі порівняно з немодифікованою архітектурою.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						57
Змн.	Арк.	№ докум.	Підпис	Дата		

- Масштабованість та продуктивність: Масштабованість та продуктивність запропонованого рішення оцінюється на основі часу, необхідного для відкриття документа в клієнтській програмі eShare при збільшенні кількості користувачів .

- Складність : Складність рішення оцінюється на основі мого досвіду та суб'єктивної думки щодо складності впровадження та підтримки рішення.

Оцінювальні тести проводилися на ноутбучі для розробників від компанії Cadmatic під керуванням Windows 10. Ноутбук оснащений процесором Intel Core i7-9850H, 32 ГБ оперативної пам'яті та твердотільним накопичувачем ємністю 1 ТБ. У тестах використовувалася налагоджувальна збірка серверної програми eShare (версія 2024T1), тоді як мікросервіс обробки документів був зібраний як релізна збірка, яка працювала в контейнері Docker. Брокер повідомлень Kafka, який використовувався в тестах, працював у хмарі за допомогою сервісу Confluent Cloud.

Важливо зазначити, що результати тестів не можна узагальнити для всіх систем та середовищ. Результати, найімовірніше, відрізнятимуться залежно від системи та середовища. Наприклад, проведення тестів на кластері потужних серверів у хмарі, ймовірно, дасть інші результати порівняно з проведенням тестів на ноутбучі розробника. Також важливо зазначити, що тестове середовище не було контрольованим. Інші програми та фонові процеси, що працюють на тестовій машині, могли вплинути на доступні ресурси та, отже, на результати тестів. Тести не є вичерпними, а радше дають загальне уявлення про вплив рішення для контейнеризації на eShare.

Для оцінки накладних витрат запропонованого рішення для контейнеризації було виміряно час, необхідний для обробки документів в eShare одним користувачем, за допомогою трьох різних конфігурацій:

- Конфігурація 1 : Поточна незмінена архітектура eShare.
- Конфігурація 2: Нова запропонована архітектура локальної шини повідомлень без використання мікросервісу обробки документів.
- Конфігурація 3: Нова запропонована архітектура віддаленої шини

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						58
Змн.	Арк.	№ докум.	Підпис	Дата		

повідомлень Kafka з одним екземпляром мікросервісу обробки документів.

Мета конфігурації 2 полягає в оцінці накладних витрат, що вносяться виключно новою подієво-керованою архітектурою, без додаткових накладних витрат, що вносяться мікросервісом обробки документів, тоді як мета конфігурації 3 полягає в оцінці накладних витрат, що вносяться новою подієво-керованою архітектурою разом із мікросервісом обробки документів. Логічно припущення полягає в тому, що накладні витрати, що вносяться конфігурацією 3, повинні бути вищими за накладні витрати, що вносяться конфігурацією 2, оскільки запуск мікросервісу обробки документів у контейнері повинен призвести до додаткових накладних витрат.

Вимірювання проводилися шляхом обробки набору з 9 різних документів (D1-D9) у серверній програмі eShare та вимірювання часу, необхідного для обробки кожного документа. Вимірювання повторювалися 20 разів для кожного документа за допомогою автоматизованого скрипта. Медіанний час, витрачений на обробку кожного документа, стандартне відхилення вимірювань та накладні витрати, внесені новими конфігураціями порівняно з поточною незміненою конфігурацією, представлені в таблиці 3.1. Крім того, зведена інформація про вимірювання представлена у вигляді коробкових діаграм на рисунку 3.1.

Накладні витрати, що виникають внаслідок нової архітектури шини повідомлень, керованої подіями (Конфіг 2), та нової архітектури шини повідомлень, керованої подіями, разом з обробкою мікросервісів (Конфіг 3), були розраховані шляхом порівняння медіанного часу, необхідного для обробки документа, з поточною незміненою архітектурою (Конфіг 1).

Таблиця, що відображає медіанний час, витрачений на обробку документа з різними конфігураціями, стандартне відхилення вимірювань та накладні витрати, внесені новими конфігураціями, порівняно з поточною незміненою конфігурацією. Накладні витрати повідомляються як відсоткове збільшення медіанного часу порівняно з незміненою конфігурацією (Конфігурація 1)

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						59
Змн.	Арк.	№ докум.	Підпис	Дата		

Статистичні дані

Документ	Розмір (КБ)	Config 1		Config 2			Config 3			Документ
		Медіа на (мс)	S D	Медіа на (мс)	SD	Overhead	Медіа на (мс)	SD	Overhead	
D1	4	1345	96	1374	192	2%	1901	173	41%	D1
D2	7	1063	95	1087	68	2%	1722	109	61%	D2
D3	20	1132	74	1159	98	2%	1818	133	60%	D3
D4	43	1133	67	1180	133	4%	1756	77	54%	D4
D5	49	1112	68	1126	90	1%	1737	133	56%	D5
D6	1544	1923	418	2079	104	8%	2540	109	32%	D6
D7	1652	1735	213	2951	411	70%	3224	234	85%	D7
D8	2062	7915	289	14941	1122	88%	11075	1133	39%	D8

Результати, наведені в Таблиці 3.1 та на Рисунку 3.3, показують, що архітектура шини повідомлень, керована подіями (Конфігурація 2), вносить приблизно 0-4% накладних витрат порівняно з поточною немодифікованою архітектурою eShare з невеликими документами (D1-D5). Однак, з більшими документами (D6-D9) накладні витрати, що вносяться архітектурою, керованою подіями, становлять 70-80%. Це значне збільшення порівняно з накладними витратами, що вносяться архітектурою, керованою подіями, з меншими документами, і причина цього не одразу зрозуміла.

З іншого боку, накладні витрати Config 3 становлять близько 30-60% порівняно з поточною немодифікованою архітектурою eShare. Це, здебільшого, очікувано. Однак, накладні витрати, внесені Config 3, менші для D8 та D9 порівняно з Config 2, що цікаво.

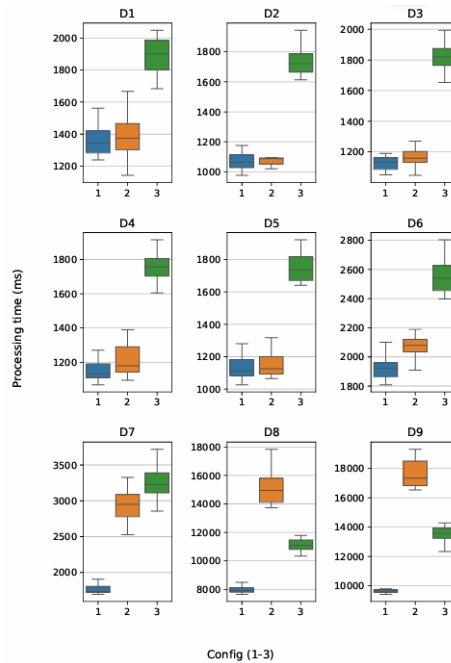


Рисунок 3.3 - Коробчасті діаграми, що візуалізують вимірювання накладних витрат різних конфігурацій під час обробки різних документів (D1-D9). Викиди на діаграмах опущені для наочності.

Через обмеження в часі тести масштабованості та продуктивності не були реалізовані. Натомість цей розділ зосереджений лише на описі того, як би проводилися тести та оцінювання, якби час дозволяв.

Подібно до тестів накладних витрат, тести масштабованості та продуктивності включатимуть вимірювання часу, необхідного для обробки документів в eShare з різними конфігураціями. Різниця полягатиме в тому, що будуть введені одночасні запити на обробку документів для імітації більшої кількості користувачів. Для виконання тестів будуть використані такі дві конфігурації :

- Конфігурація 1 : Поточна незмінена архітектура eShare.
- Конфігурація 2 : EShare з кількома екземплярами запропонованого контейнеризованого мікросервісу обробки документів. В ідеалі, мікросервіс обробки документів буде розгорнутий у хмарі на потужнішій інфраструктурі, ніж та, що використовується в конфігурації 1, для імітації реалістичного виробничого середовища.

Цього разу, замість вимірювання накладних витрат, що виникають

внаслідок нової конфігурації, тести масштабованості та продуктивності будуть зосереджені на вимірюванні часу, необхідного для обробки документів, коли кількість користувачів збільшується. Припущення та очікуваний результат полягатимуть у тому, що Конфігурація 1 буде обмежена через гірше апаратне забезпечення та відсутність можливостей паралельної обробки порівняно з Конфігурацією 2.

Для подальшої оцінки масштабованості та продуктивності запропонованого рішення, кількість одночасних запитів на обробку документів та екземплярів мікросервісів можна варіювати, щоб отримати більше уявлення про масштабованість та продуктивність контейнерного рішення.

Складність заданої архітектури та конфігурації програмного забезпечення може бути важко об'єктивно виміряти. У цьому випадку складність запропонованого рішення для контейнеризації оцінювалася на основі того, як я сприймав складність рішення та наскільки складним було побудувати та впровадити запропоноване рішення для контейнеризації, яке було представлено в розділі .

З точки зору архітектури системи, джерела складності контейнерного рішення включають необхідність керування кількома окремими компонентами та службами, такими як брокер Kafka, мікросервіс обробки документів та сервери eShare. Це також створює складність з точки зору безпеки, моніторингу та ведення журналу. Збільшена кількість компонентів та служб у системі не тільки збільшує поверхню атаки, але й ускладнює керування політиками безпеки між кількома компонентами. Моніторинг та ведення журналу також стають складнішими з кількома компонентами, що вимагає надійного рішення для ведення журналу та моніторингу, яке може агрегувати журнали з кількох джерел та забезпечувати цілісне уявлення про систему. З точки зору розробки та обслуговування, контейнерне рішення може бути більш складним та трудомістким для розробки та обслуговування.

Мікросервіс може бути складним і вимагати додаткових зусиль порівняно з розробкою та тестуванням монолітного застосунку.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						62
Змн.	Арк.	№ докум.	Підпис	Дата		

З іншого боку, контейнерне рішення також спрощує певні аспекти системи. Наприклад, завдяки ізольованій природі контейнерів і мікросервісів, кожен мікросервіс можна розробляти, розгортати та масштабувати незалежно. Це спрощує життєві цикли окремих сервісів і полегшує розробнику розуміння, роботу та зосередження на одному сервісі. Керування помилками та несправностями також спрощується завдяки ізольованій природі мікросервісів. Збій у мікросервісі з меншою ймовірністю вплине на інші частини системи, що спрощує процеси відновлення після несправностей.

Тести оцінки накладних витрат показали, що архітектура шини повідомлень, керована подіями (конфігурація 2), здається, створює приблизно 0-4% накладних витрат порівняно з поточною немодифікованою архітектурою (конфігурація 1) eShare для невеликих документів, тоді як для більших документів (D7, D8, D9) накладні витрати становили близько 70-80%. Таке радикальне збільшення накладних витрат для більших документів є неочікуваним, і причина цього не одразу зрозуміла. Можливо, є недоліки у реалізації рішення або що налаштування тестування має інші недоліки.

Оскільки конфігурації 1 та 2 відрізняються лише способом відправлення запитів на обробку документів, накладні витрати, що вносяться подієвою архітектурою, повинні бути подібними незалежно від розміру документа. У конфігурації 1, щоразу, коли обробляється документ, виконуваний файл, відповідальний за обробку документа, викликається безпосередньо з серверної програми eShare. У конфігурації 2 документ обробляється шляхом надсилання запиту на обробку документа до шини повідомлень у пам'яті, до якої підключено слухача, який потім викликає той самий виконуваний файл, що викликається безпосередньо в конфігурації 1. І, як згадувалося раніше, двійкові файли документів не надсилаються через шину повідомлень, а лише метадані документа. Натомість для передачі файлів використовується спільна файлова система. Це свідчить про те, що причиною збільшення накладних витрат для великих документів буде недолік у способі, яким реалізація шини повідомлень у пам'яті або слухач, який обробляє запити на обробку документів, обробляє вхідні та

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						63
Змн.	Арк.	№ докум.	Підпис	Дата		

вихідні дані виконуваного файлу для обробки документів. Для визначення точної причини збільшення накладних витрат для великих документів необхідні подальші дослідження та більш детальний підхід до тестування та бенчмаркінгу. Однак, через обмеження в часі, це, на жаль, неможливо.

З іншого боку, конфігурація 3 тестів накладних витрат показала, що віддалена шина повідомлень Kafka разом із мікросервісом обробки документів створює приблизно 30-60% накладних витрат порівняно з конфігурацією 1. Це очікувано та здається розумним, оскільки брокер Kafka, який використовувався в тестах, працює в хмарі, а мікросервіс обробки документів працює в контейнері Docker, що створює додаткові накладні витрати порівняно з обробкою документів безпосередньо в серверній програмі eShare. Однак накладні витрати з D8 та D9 менші порівняно з конфігурацією 2, що цікаво. Це, ймовірно, пов'язано з необґрунтовано високими накладними витратами, що створюються конфігурацією 2 для більших документів, і щоб з'ясувати причину цього, потрібні подальші дослідження.

Що стосується складності запропонованого контейнерного рішення, очевидно, що нова архітектура вносить певну складність у систему, водночас спрощуючи певні її аспекти. Збільшення складності головним чином пов'язане з необхідністю керування кількома компонентами та сервісами, такими як брокер Kafka, мікросервіс обробки документів та сервери eShare. Збільшена кількість компонентів та сервісів у системі ускладнює захист, моніторинг та реєстрацію системи в цілому. Однак ізольований характер контейнерів та мікросервісів також спрощує певні аспекти системи. Кожен мікросервіс є дуже незалежним та відокремленим від інших частин системи, що спрощує його розробку, розгортання та масштабування.

Загалом, запропоноване рішення для контейнеризації eShare, здається, створює деякі накладні витрати порівняно з поточною немодифікованою архітектурою eShare. Окрім аномалій із накладними витратами, що виникають через подієво-керовану архітектуру з більшими документами, накладні витрати, що виникають через нову архітектуру, здаються розумними та очікуваними.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						64
Змн.	Арк.	№ докум.	Підпис	Дата		

Збільшена складність нової архітектури є компромісом на користь переваг, які вона приносить з точки зору масштабованості, продуктивності та зручності обслуговування. Збільшена складність має бути чимось, чим можна керувати та терпіти за допомогою належних інструментів, процесів та практик, враховуючи переваги нової архітектури. Через обмеження в часі тести масштабованості та продуктивності не були проведені, і тому масштабованість та продуктивність запропонованого рішення залишаються неперевіреними. Для оцінки масштабованості та продуктивності запропонованого рішення для контейнеризації eShare необхідні подальші дослідження та тестування.

Щодо RQ1, переваги використання контейнерних технологій із застосунком, таким як eShare, в основному пов'язані з масштабованістю, продуктивністю та зручністю обслуговування. Аспекти масштабованості та продуктивності залишаються неперевіреними, але запропоноване рішення для контейнеризації теоретично має покращити масштабованість та продуктивність eShare. З іншого боку, зручність обслуговування системи стане зрозумілою лише з часом, у міру розробки та підтримки системи. Однак теоретично ізольований характер контейнерів та мікросервісів має спростити розробку та обслуговування системи. Недоліки використання контейнерних технологій із застосунком, таким як eShare, в основному пов'язані зі збільшенням складності системи та накладними витратами, що виникають через додатковий рівень віртуалізації контейнерів Docker.

З точки зору RQ2, функціональність у такому застосунку, як eShare, яка найбільше підходить для контейнеризації, - це функціональність, яка є відокремленою та легко ізолюється від решти системи. Наприклад, функціональність обробки документів у

eShare був гарним кандидатом для контейнеризації, оскільки це окрема та дуже незалежна частина системи.

Метою цієї роботи було дослідження та аналіз контейнеризації та контейнерних технологій у контексті веб-застосунків. Використання контейнерних технологій як рішення для різних проблем у розробці та розгортанні

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						65
Змн.	Арк.	№ докум.	Підпис	Дата		

програмного забезпечення було популярною темою останніх років. Контейнери інкапсулюють застосунок, його залежності та конфігурацію в автономні одиниці, які можна легко розгортати та керувати, а також забезпечують стабільну та передбачувану роботу застосунку в різних середовищах. Як результат, запуск застосунків у контейнерах абстрагує базову інфраструктуру та дозволяє розробникам зосередитися на написанні коду та створенні застосунків, а не на управлінні інфраструктурою

3.4 Висновки по розділу

У третьому розділі було розглянуто процес розробки, розгортання та оцінювання ефективності контейнеризованих рішень. Досліджено можливості керування контейнерами за допомогою Docker Compose та особливості організації взаємодії між сервісами. Проведено проєктування й реалізацію рішення для контейнеризації додатку, а також оцінено його ефективність і масштабованість. Аналіз отриманих результатів показав, що використання контейнеризації дозволяє підвищити стабільність роботи системи, спростити процес розгортання та забезпечити ефективне масштабування програмних додатків. Отримані результати підтверджують доцільність застосування контейнерних технологій у сучасній розробці програмного забезпечення.

ВИСНОВКИ

У ході виконання дипломної роботи на тему «Використання контейнеризації для масштабування додатків» було проведено комплексне

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						66
Змн.	Арк.	№ докум.	Підпис	Дата		

дослідження сучасних підходів до побудови масштабованих програмних систем із застосуванням контейнерних технологій. Робота охоплює аналіз теоретичних основ контейнеризації, дослідження інструментів і методів розробки, а також практичну реалізацію та оцінку ефективності запропонованого рішення. Це дозволило сформулювати цілісне бачення процесу впровадження контейнеризації як ключового інструменту забезпечення гнучкості, продуктивності та масштабованості сучасних додатків.

На початковому етапі було розглянуто базові принципи контейнеризації, її відмінності від традиційної віртуалізації, а також визначено основні сценарії використання у програмній інженерії. Особливу увагу приділено перевагам контейнерів, таким як легковаговість, швидкість розгортання, ізолюваність середовищ виконання та зручність масштабування. Це дало змогу обґрунтувати доцільність використання контейнеризації для побудови сучасних розподілених систем.

У другому розділі було досліджено інструменти та підходи до реалізації контейнеризованих рішень. Зокрема, розглянуто можливості платформи Docker як основного інструменту створення та управління контейнерами, а також проаналізовано приклади її практичного використання. Важливим аспектом стало вивчення контейнерних архітектур, включаючи мікросервісний підхід, що забезпечує незалежність компонентів системи та спрощує їх масштабування й супровід.

У процесі розробки було реалізовано практичне рішення із застосуванням Docker та Docker Compose для організації взаємодії між кількома контейнерами. Було спроектовано структуру контейнеризованого додатку, визначено залежності між компонентами та реалізовано механізми автоматичного розгортання середовища. Це дозволило забезпечити швидке налаштування системи та її відтворюваність у різних середовищах.

Значну увагу приділено оцінці ефективності запропонованого рішення. Проведений аналіз показав, що використання контейнеризації суттєво підвищує швидкість розгортання додатків, спрощує процес масштабування та знижує

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						67
Змн.	Арк.	№ докум.	Підпис	Дата		

витрати на підтримку інфраструктури. Крім того, контейнеризація забезпечує кращу ізоляцію компонентів, що позитивно впливає на стабільність та безпеку системи.

Загалом, результати роботи підтверджують доцільність використання контейнерних технологій для побудови масштабованих програмних систем. Основними перевагами запропонованого підходу є гнучкість архітектури, можливість швидкого масштабування, ефективне використання ресурсів та спрощення процесів розробки й розгортання. Водночас слід враховувати певні обмеження, зокрема складність управління великою кількістю контейнерів та необхідність використання додаткових інструментів оркестрації.

У перспективі подальших досліджень доцільно розглянути інтеграцію систем оркестрації контейнерів (наприклад, Kubernetes), автоматизацію процесів CI/CD, а також впровадження механізмів моніторингу та балансування навантаження. Це дозволить підвищити ефективність використання контейнеризації та забезпечити ще вищий рівень масштабованості й надійності програмних систем.

					БР.ІІ – 22.00.00.000 ПЗ	Арк.
						68
Змн.	Арк.	№ докум.	Підпис	Дата		

СПИСОК ПОСИЛАНЬ НА ДЖЕРЕЛА

1. Docker Documentation. (2025). docs.docker.com.
<https://docs.docker.com/>
2. Kubernetes Documentation. (2025). kubernetes.io.
<https://kubernetes.io/docs/>
3. Microsoft Azure Kubernetes Service (AKS). (2025). learn.microsoft.com.
<https://learn.microsoft.com/en-us/azure/aks/>
4. Google Kubernetes Engine Documentation. (2025). cloud.google.com.
<https://cloud.google.com/kubernetes-engine/docs>
5. Amazon Elastic Kubernetes Service (EKS). (2025). aws.amazon.com.
<https://docs.aws.amazon.com/eks/>
6. Docker Compose Documentation. (2025). docs.docker.com.
<https://docs.docker.com/compose/>
7. Red Hat OpenShift Documentation. (2025). redhat.com.
<https://docs.openshift.com/>
8. Newman, S. (2021). Building Microservices (2nd ed.). O'Reilly Media.
<https://www.oreilly.com/>
9. Burns, B., Grant, B., Oppenheimer, D., Brewer, E., & Wilkes, J. (2016). Borg, Omega, and Kubernetes. ACM Queue, 14(1), 70–93.
<https://doi.org/10.1145/2898442>
10. Merkel, D. (2014). Docker: Lightweight Linux Containers. Linux Journal, 2014(239).
11. Pahl, C. (2015). Containerization and the PaaS Cloud. IEEE Cloud Computing, 2(3), 24–31. <https://doi.org/10.1109/MCC.2015.51>
12. Turnbull, J. (2014). The Docker Book. <https://dockerbook.com/>
13. Hightower, K., Burns, B., & Beda, J. (2017). Kubernetes: Up and Running. O'Reilly Media.
14. Poulton, N. (2023). Docker Deep Dive. Leanpub.
<https://leanpub.com/dockerdeeplive>

					БР.ІІІ – 22.00.00.000 ІІЗ	Арк.
						69
Змн.	Арк.	№ докум.	Підпис	Дата		

15. IBM Cloud Kubernetes Service Docs. (2025). cloud.ibm.com.
<https://cloud.ibm.com/docs/containers>
16. HashiCorp Nomad Documentation. (2025). developer.hashicorp.com.
<https://developer.hashicorp.com/nomad/docs>
17. CNCF Cloud Native Landscape. (2025). landscape.cncf.io.
<https://landscape.cncf.io/>
18. NGINX Documentation (Container Deployment). (2025). nginx.org.
<https://nginx.org/en/docs/>
19. Traefik Proxy Documentation. (2025). doc.traefik.io.
<https://doc.traefik.io/traefik/>
20. Prometheus Monitoring Documentation. (2025). prometheus.io.
<https://prometheus.io/docs/>
21. Grafana Documentation. (2025). grafana.com. <https://grafana.com/docs/>
22. Istio Service Mesh Documentation. (2025). istio.io.
<https://istio.io/latest/docs/>
23. Linkerd Documentation. (2025). linkerd.io.
<https://linkerd.io/2.14/overview/>
24. Open Container Initiative (OCI). (2025). opencontainers.org.
<https://opencontainers.org/>
25. Kubernetes Patterns. (2020). Bilgin Ibryam, Roland Huß. O'Reilly Media.
26. Villamizar, M., et al. (2015). Evaluating Docker Containers. IEEE Cloud Computing. <https://doi.org/10.1109/MCC.2015.51>
27. Zhang, Q., Chen, M. (2018). Container-Based Cloud Architecture. IEEE Access. <https://doi.org/10.1109/ACCESS.2018.2810597>
28. Chen, L. (2018). Microservices Architecture. IEEE Software, 35(3), 96–100.
29. Dragoni, N., et al. (2017). Microservices: Yesterday, Today, and Tomorrow. Springer.
30. Richardson, C. (2018). Microservices Patterns. Manning Publications.
<https://www.manning.com/books/microservices-patterns>

					БР.ІІІ – 22.00.00.000 ІІЗ	Арк.
						70
Змн.	Арк.	№ докум.	Підпис	Дата		

31. Google Cloud Architecture Framework. (2025). cloud.google.com.
<https://cloud.google.com/architecture/framework>
32. AWS Well-Architected Framework. (2025). aws.amazon.com.
<https://aws.amazon.com/architecture/well-architected/>
33. Azure Architecture Center. (2025). learn.microsoft.com.
<https://learn.microsoft.com/en-us/azure/architecture/>
34. Kubernetes Security Best Practices. (2024). kubernetes.io.
<https://kubernetes.io/docs/concepts/security/>
35. OWASP Container Security Top 10. (2024). owasp.org.
<https://owasp.org/www-project-top-ten/>
36. DevOps with Kubernetes. (2022). Viktor Farcic. Packt Publishing.
37. Continuous Delivery. (2010). Humble, J., & Farley, D. Addison-Wesley.
38. Bass, L., Weber, I., & Zhu, L. (2015). DevOps: A Software Architect's Perspective. Addison-Wesley.
39. Kubernetes Performance and Scalability. (2024). CNCF Reports.
<https://www.cncf.io/reports/>
40. Sharma, R., & Kumar, S. (2023). Container Orchestration for Scalable Applications. Journal of Cloud Computing, 12(4), 55–72.
<https://doi.org/10.1186/s13677-023-00456-2>

БІБЛІОГРАФІЧНА ДОВІДКА

Тема бакалаврської роботи бакалавра: "Використання контейнеризації для масштабування додатків "

Обсяг пояснювальної записки: 68 аркушів

Дата закінчення дипломної роботи 10 червня 2026р.

Підпис студента _____

