

Вінницький національний технічний університет
Факультет інформаційних технологій та комп'ютерної інженерії
Кафедра обчислювальної техніки

БАКАЛАВРСЬКА КВАЛІФІКАЦІЙНА РОБОТА

на тему:

«Комп'ютерна мережа компанії з моніторингом перевантаження»

08-54.БКР.010.00.000 ПЗ

Виконав студент 4 курсу, групи ІСП-216
спеціальності 123 — Комп'ютерна інженерія
освітня програма «Системне програмування»

 — Олійник І. І.

Керівник к.т.н., доц. каф. ОТ

 — Войцеховська О. В.

Рецензент: к.т.н., доц. каф. ПЗ

 — Чехместрук Р. Ю.

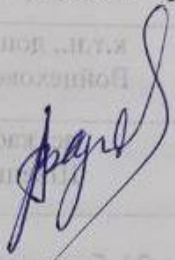


Допущено до захисту
Завідувач кафедри ОТ
д.т.н., проф. Азаров О.Д.

«10» 06 2025 р.

Вінниця ВНТУ — 2025 рік

Вінницький національний технічний університет
Факультет інформаційних технологій та комп'ютерної інженерії
Кафедра обчислювальної техніки
Освітньо-кваліфікаційний рівень бакалавр
Галузь знань – 12 – Інформаційні технології
Спеціальність – 123 – Комп'ютерна інженерія
Освітньо-професійна програма – Системне програмування


ЗАТВЕРДЖУЮ

Завідувач кафедри ОТ

д.т.н., проф. Азаров О. Д

“ 21 ” березня 2025 року

ЗАВДАННЯ НА БАКАЛАВРСЬКУ КВАЛІФІКАЦІЙНУ РОБОТУ СТУДЕНТУ

Олійнику Івану Івановичу

1 Тема роботи: «Комп'ютерна мережа компанії з моніторингом перевантаження», керівник роботи к.тех.н., доц. каф. ОТ Войцеховська О.В. затверджені наказом ВНТУ від «20» березня 2025 року №97;

2 Термін подання студентом роботи: 13.05.2025.

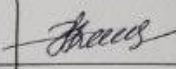
3 Вихідні дані до роботи: призначення — забезпечення надійної роботи мережі компанії інтернет-провайдера за допомогою системи моніторингу перевантаження та балансування трафіку; топологія — тип «квітка»; механізм моніторингу — протокол зв'язку SNMP у системі Zabbix; балансування трафіку — протокол зв'язку ECMP для мультиканальності; безпека — застосування ACL від несанкціонованого доступу; обладнання — маршрутизатори, L2/L3 комутатори, сервер моніторингу; підтримується розширення мережі без зміни архітектури, середовище розробки — Cisco Packet Tracer.

4 Зміст розрахунково-пояснювальної записки: вступ, аналіз сучасних технологій побудови комп'ютерних мереж, розробка структури корпоративної мережі компанії, розробка корпоративної мережі компанії з моніторингом перевантаження та тестування, висновки, список використаних джерел, додатки

5 Перелік ілюстративного матеріалу: фізична схема компанії; логічна мережі компанії.

6 Консультанти розділів роботи приведені в таблиці 1

Таблиця 1 — Консультанти роботи

Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		завдання видав	завдання прийняв
1-3	к.т.н., доц. каф. ОТ Войцеховська О. В.	21.03.25 	13.05.25 
Нормоконтроль	ас. каф. ОТ Швець С. І. 	14.05.2025	30.05.25

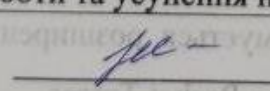
7 Дата видачі завдання 21 березня 2025 року

8 Календарний план наведений у таблиці 2.

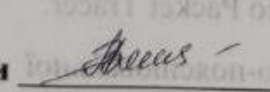
Таблиця 2 — Календарний план

№ з/п	Назва етапів дипломної роботи	Строк виконання	Підп
1	Постановка задачі роботи	21.03.25	
2	Пошук матеріалів та інформаційних джерел	21.03.25— 30.03.25	
3	Аналіз існуючих топологій, обладнання та визначення вимог до створюваної системи моніторингу	21.03.25— 1.04.25	
4	Обґрунтування архітектури мережі та принципів роботи системи моніторингу	1.04.25— 14.04.25	
5	Розробка й тестування мережі та системи моніторингу	14.04.25— 2.05.25	
7	Оформлення пояснювальної записки та ілюстративного матеріалу	3.05.25— 10.05.25	
8	Перевірка якості виконання бакалаврської кваліфікаційної роботи та усунення недоліків	13.05.25	

Студент

 Іван ОЛІЙНИК

Керівник роботи

 к.т.н., доц. Олена ВОЙЦЕХОВСЬКА

АНОТАЦІЯ

УДК 621.3

Олійник І.І. Комп'ютерна мережа компанії з моніторингом перевантаження. Бакалаврська кваліфікаційна робота зі спеціальності 123 «Комп'ютерна інженерія», освітня програма «Системне програмування». Вінниця: ВНТУ, 2025. 92 с.

На укр. мові. Бібліогр.: 30 назв; рис.: 58; табл. 7.

В бакалаврській кваліфікаційній роботі було розроблено комп'ютерну мережу компанії із системою моніторингу, яка аналізує та спостерігає за станом мережі та у разі виникнення перевантаження мережі потенційно виявляє вузькі місця, де буде застосовуватися механізм балансування трафіку. Комп'ютерна мережа компанії складається з адміністративних відділів, серверних та трьох вузлів агрегації із вузлами доступу для хостів в мережі.

Ключові слова: комп'ютерна мережа компанії; моніторинг стану мережі; механізм балансування трафіку; перевантаження мережі; вузли агрегації; вузли доступу.

ABSTRACT

Oliynyk I.I. Computer network of the company with monitoring of overload. Bachelor's qualification work in the specialty 123 "Computer Engineering", educational program "System Programming". Vinnytsia: VNTU, 2025. 92 p.

In Ukrainian. Bibliography: 30 titles; Fig.: 58; Table. 7.

In the bachelor's qualification work, a computer network of the company was developed, capable of analyzing and monitoring the state of the network and, in the event of network overload, potentially identifying bottlenecks where the traffic balancing mechanism will be applied. The company's computer network consists of administrative departments, server and three aggregation nodes with access nodes for hosts in the network.

Keywords: computer network of the company; monitoring of the state of the network; traffic balancing mechanism; network overload; aggregation nodes; access nodes.

ПЕРЕЛІК УМОВНИХ ПОЗНАЧЕНЬ, СКОРОЧЕНЬ, ТЕРМІНІВ

ACL — Access Control List

BGP — Border Gateway Protocol

CLI — Command Line Interface

CoS — Class of Service

DHCP — Dynamic Host Configuration Protocol

DPI — Deep Packet Inspection

DSCP — Differentiated Services Code Point

DUAL — Diffusing Update Algorithm

ECMP — Equal-Cost Multi-Path

EIGRP — Enhanced Interior Gateway Routing Protocol

GRE — Generic Routing Encapsulation

HSRP — Hot Standby Router Protocol

ICMP — Internet Control Message Protocol

IGP — Interior Gateway Protocol

IoT — Internet of Things

IP — Internet Protocol

LAN — Local Area Network

LLDP — Link Layer Discovery Protocol

L2TP — Layer 2 Tunneling Protocol

MPLS — Multiprotocol Label Switching

NAT — Network Address Translation

PoE — Power over Ethernet

QoS — Quality of Service

RIP — Routing Information Protocol

SDN — Software Defined Networking

SNMP — Simple Network Management Protocol

VPN — Virtual Private Network

VLAN — Virtual LAN

WAN — Wide Area Network

ЗМІСТ

ВСТУП.....	3
1 АНАЛІЗ СУЧАСНИХ ТЕХНОЛОГІЙ КОМП'ЮТЕРНИХ МЕРЕЖ	5
1.1 Технології побудови корпоративних комп'ютерних мереж	5
1.2 Аналіз топологій комп'ютерних мереж.....	10
1.3 Огляд протоколів зв'язку для передачі даних	15
1.4 Технології моніторингу стану комп'ютерних мереж.....	17
1.5 Проблеми перевантаження мережі та методи їх вирішення	20
2 РОЗРОБКА СТРУКТУРИ КОРПОРАТИВНОЇ МЕРЕЖІ КОМПАНІЇ	23
2.1 Розробка загальної структури компанії	23
2.2 Проектування логічної схеми корпоративної мережі компанії інтернет-провайдера	26
2.3 Розрахунок адресного простору мережі компанії	30
2.4 Вибір мережного обладнання	33
2.5 Використання віртуальних локальних мереж (VLAN)	41
2.6 Реалізація механізму балансування трафіку	44
2.7 Розробка системи моніторингу та керування перевантаженням	45
3 РОЗРОБКА КОРПОРАТИВНОЇ МЕРЕЖІ КОМПАНІЇ З МОНІТОРИНГОМ ПЕРЕВАНТАЖЕННЯ ТА ТЕСТУВАННЯ.....	49
3.1 Налаштування логічної схеми мережі	49
3.2 Налаштування системи моніторингу та перевантаження	61
3.3 Перевірка працездатності розробленої мережі	64
3.4 Тестування системи перевантаження.....	69
ВИСНОВКИ	73
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	75
ДОДАТОК А Технічне завдання	78
ДОДАТОК Б Протокол перевірки навчальної (кваліфікаційної) роботи	82
ДОДАТОК В Ілюстративна частина	83
ДОДАТОК Г Лістинг конфігураційних файлів	86

ВСТУП

Корпоративні комп'ютерні мережі відіграють ключову роль у забезпеченні стабільної роботи великих компаній. Ефективність їхньої роботи залежить не лише від швидкості передачі даних, а й від здатності виявляти потенційні перевантаження та своєчасно реагувати на них. У зв'язку з цим все більшої актуальності набувають рішення, спрямовані на моніторинг навантаження мережі та балансування трафіку.

Сучасні інформаційні системи мають справу з дедалі більшими обсягами даних, що циркулюють в мережі. У таких умовах стає критично важливою здатність мережі не лише ефективно передавати дані, а й виявляти «вузькі місця» у структурі, що можуть стати причиною затримок або відмов. Проблема перевантажень мережних пристроїв, надмірного трафіку потребує впровадження ефективних механізмів моніторингу та оптимізації.

Актуальність теми зумовлена зростанням потреби в створенні високонадійних комп'ютерних мереж, які забезпечують безперебійну комунікацію між структурними підрозділами компанії. Особливо важливою є здатність таких мереж до автоматичного виявлення “вузьких місць” та оптимального розподілу трафіку, що зменшує ймовірність перевантаження та сприяє підвищенню загальної продуктивності інформаційної інфраструктури підприємства.

У цьому контексті все більшого значення набуває застосування протоколів моніторингу, таких як SNMP, а також технологій балансування навантаження, наприклад, ECMP, які дозволяють не лише забезпечити стабільну роботу системи, а й гарантувати її відмовостійкість. Інструмент мережного моделювання такий як Cisco Packet Tracer, дає можливість для аналізу, тестування та оптимізації архітектури мережі. Це особливо важливо для провайдерів, які працюють із великим числом клієнтів і повинні гарантувати якісний зв'язок у будь-який час.

Метою роботи є проєктування та налаштування надійної, масштабованої і відмовостійкої комп'ютерної мережі інтернет-провайдера, яка забезпечує ефективний моніторинг перевантаження та балансування трафіку.

Така мережа має гарантувати стабільну роботу навіть за умов високих навантажень, забезпечувати централізоване управління та дозволяти швидко

виявляти і локалізувати проблемні ділянки. У роботі також передбачено застосування сучасних технологій, зокрема SNMP та ECMP, що дозволяють підвищити гнучкість і надійність мережної інфраструктури.

Для досягнення поставленої мети потрібно розв'язати такі **задачі**:

- провести аналіз сучасних технологій побудови та моніторингу стану корпоративних мереж та вибрати оптимальні рішення для досягнення високої продуктивності та надійності;
- розробити архітектуру корпоративної мережі з урахуванням топології та протоколів маршрутизації;
- вибрати мережне обладнання, необхідне для реалізації розробленої мережі;
- виконати розподіл адресного простору та провести налаштування мережного обладнання;
- створити систему моніторингу перевантажень для виявлення вузьких місць та автоматичного балансування трафіку;
- створити механізм балансування трафіку для рівномірного розподілу навантаження та підвищення стійкості мережі до пікових навантажень;
- протестувати мережу та систему моніторингу перевантаження в реальних умовах для перевірки її стабільності та надійності.

Об'єктом дослідження є процес функціонування комп'ютерної мережі.

Предметом дослідження є логічна структура та конфігурація комп'ютерної мережі компанії — інтернет провайдера.

Апробація результатів роботи здійснена під час доповіді на конференції Молодь в науці: дослідження, проблеми, перспективи (МН-2025) факультету інформаційних технологій та комп'ютерної інженерії [1].

Матеріали: Використання протоколів маршрутизації в мережі інтернет-провайдера / Олійник І. І., Войцеховська О. В., Клепко Д. Д. / Матеріали конференції Молодь в науці: дослідження, проблеми, перспективи (МН-2025) факультету інформаційних технологій та комп'ютерної інженерії. [Електронний ресурс] — Режим доступу: <https://conferences.vntu.edu.ua/index.php/mn/mn2025/paper/viewFile/25229/20873>

1 АНАЛІЗ СУЧАСНИХ ТЕХНОЛОГІЙ КОМП'ЮТЕРНИХ МЕРЕЖ

1.1 Технології побудови корпоративних комп'ютерних мереж

У сучасному інформаційному суспільстві корпоративні комп'ютерні мережі є критично важливими інфраструктурами для функціонування підприємств, установ і організацій. Вони забезпечують ефективний обмін даними між користувачами, централізоване управління ресурсами, підвищення рівня інформаційної безпеки та відмовостійкості роботи бізнес-процесів. Побудова таких мереж вимагає врахування ряду технологічних і архітектурних рішень, зокрема вибору відповідної топології, типу мережного обладнання, протоколів зв'язку та засобів моніторингу [2].

Одним із ключових етапів проектування є вибір архітектури мережі, яка може бути централізованою, децентралізованою або гібридною. У більшості сучасних корпоративних мереж застосовують ієрархічну трирівневу модель, що включає рівні доступу, агрегації та ядра (рисунок 1.1).

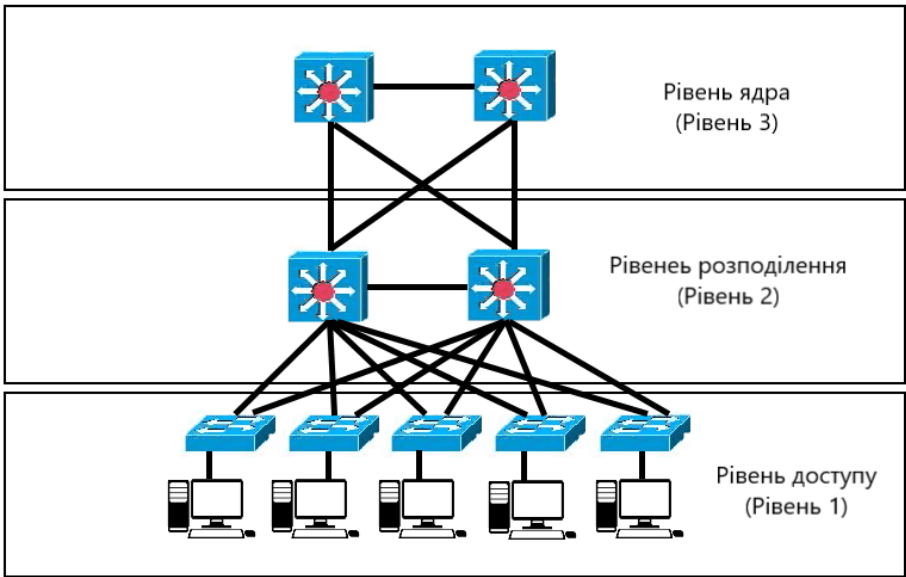


Рисунок 1.1 — Ієрархічна трирівнева модель мережі

Такий підхід дозволяє оптимізувати продуктивність мережі, спростити адміністрування та забезпечити масштабованість. При цьому рівень доступу безпосередньо відповідає за з'єднання кінцевих пристроїв із мережею, рівень агрегації об'єднує декілька сегментів та оптимізує потоки трафіку, а рівень ядра відповідає за високошвидкісну маршрутизацію між великими частинами мережі.

При побудові мережі особливу увагу приділяють вибору фізичної та логічної топології. До найпоширеніших фізичних топологій належать зірка, кільце, шина та сітка, але у корпоративних мережах найчастіше застосовується модифікована зірка або часткова сітка завдяки їхній гнучкості та можливості резервування. Логічна топологія відображає спосіб організації потоків даних у межах мережі і часто не збігається з фізичним розміщенням пристроїв [2].

Ще однією важливою тенденцією в побудові корпоративних мереж є використання програмно-визначених мереж (SDN — Software Defined Networking). У SDN управління мережею відокремлюється від апаратної частини: централізований контролер керує перемиканням і маршрутизацією на основі глобального бачення стану мережі (рисунок 1.2). Це дозволяє зменшити складність управління великими інфраструктурами, прискорити впровадження нових політик і покращити автоматизацію мережних процесів [2].

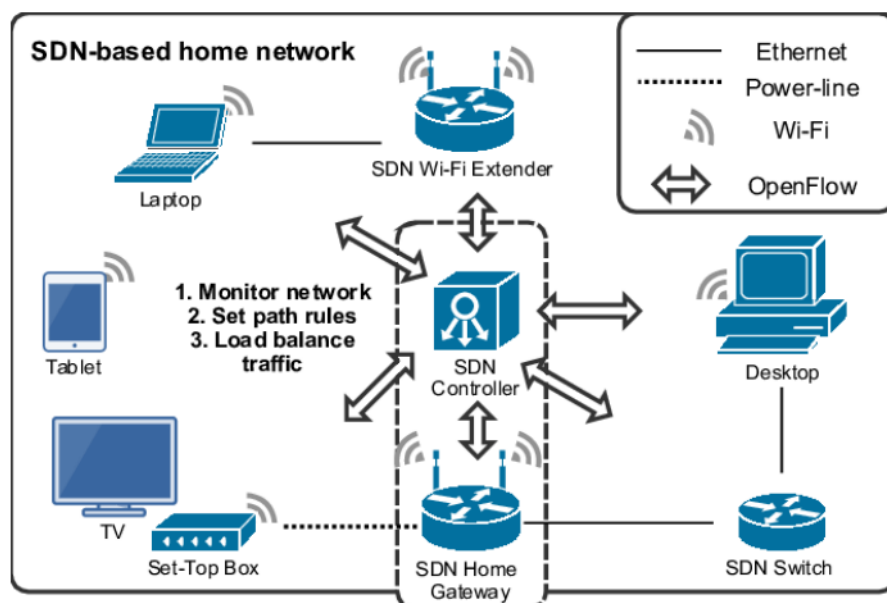


Рисунок 1.2 — Домашня мережа на основі гібридної архітектури SDN

Ключовими технологіями побудови корпоративних мереж є комутація (switching), маршрутизація (routing), а також впровадження віртуальних локальних мереж (VLAN). Комутація на рівні доступу дозволяє забезпечити швидкий та ефективний обмін даними між пристроями в межах одного сегменту мережі. Маршрутизація забезпечує з'єднання між різними сегментами або підмережами, дозволяючи реалізувати складну багаторівневу мережну структуру. Віртуальні

локальні мережі (VLAN) дозволяють логічно розділяти трафік різних підрозділів компанії без фізичної ізоляції мережного обладнання, підвищуючи рівень безпеки та гнучкість управління [2].

Важливою складовою сучасних мереж є впровадження протоколів динамічної маршрутизації, зокрема OSPF (Open Shortest Path First), EIGRP (Enhanced Interior Gateway Routing Protocol) та BGP (Border Gateway Protocol). Ці протоколи дозволяють автоматично адаптувати маршрути залежно від стану мережі та забезпечують стійкість до відмов шляхом повторного обчислення маршрутів у разі змін у топології.

З метою підвищення ефективності роботи мережі застосовуються також технології балансування навантаження (наприклад, ECMP — Equal-Cost Multi-Path), які дають змогу розподіляти трафік між декількома маршрутами з однаковою вартістю (рисунок 1.3). Це особливо важливо для зниження ризику перевантаження вузлів та підвищення продуктивності при пікових навантаженнях.

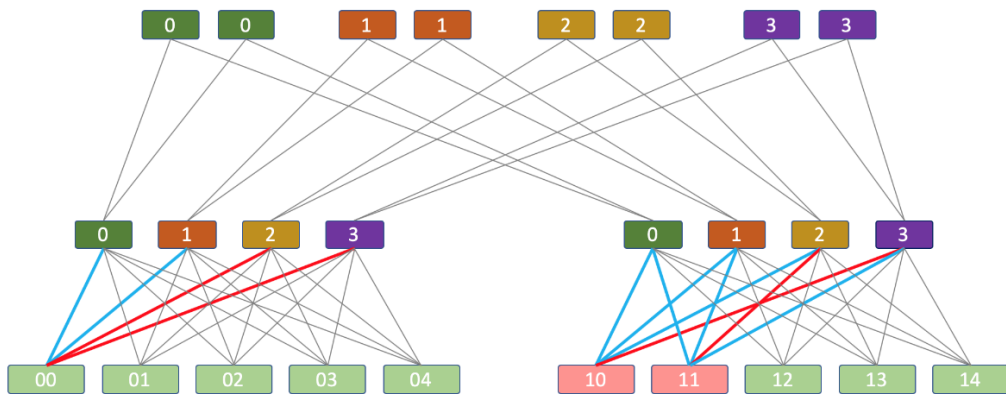


Рисунок 1.3 — Технології балансування навантаження з трьома ECMP групами

Не менш важливою складовою побудови корпоративної мережі є питання безпеки. Тут використовуються як програмні, так і апаратні рішення: фаєрволи, системи виявлення вторгнень (IDS), шифрування даних, впровадження політик доступу (ACL — Access Control List), а також багатофакторна автентифікація. Впровадження безпечної архітектури Zero Trust, що передбачає постійний контроль доступу незалежно від розташування користувача, також дедалі більше використовується у корпоративному середовищі.

Варто зазначити, що нехтування елементарними заходами безпеки може

призвести до катастрофічних наслідків. Наприклад, у 2025 році британські компанії Marks & Spencer та Co-op стали жертвами кібератаки, здійсненої хакерською групою Scattered Spider. Зловмисники використали методи соціальної інженерії, зокрема, видавали себе за співробітників компаній, щоб переконати IT-підтримку скинути паролі облікових записів з привілейованим доступом. Це дозволило їм отримати несанкціонований доступ до критичних частин мережі. Наслідки включали порушення роботи онлайн-сервісів, дефіцит товарів та витік персональних даних співробітників. Відновлення систем вимагало значних ресурсів і часу, а Національний центр кібербезпеки Великої Британії закликав компанії посилити процедури перевірки особистості під час скидання паролів, особливо для адміністративних користувачів [3].

Ще одним важливим аспектом є використання віртуалізації не лише на рівні логічного поділу мережі, а й у побудові інфраструктури дата-центрів. Завдяки технологіям, таким як VMware NSX або Cisco ACI, можлива віртуалізація не лише серверів, а й самої мережі, що забезпечує високу гнучкість у масштабуванні, міграції робочих навантажень, підвищення безпеки та ефективного управління ресурсами (рисунок 1.4).

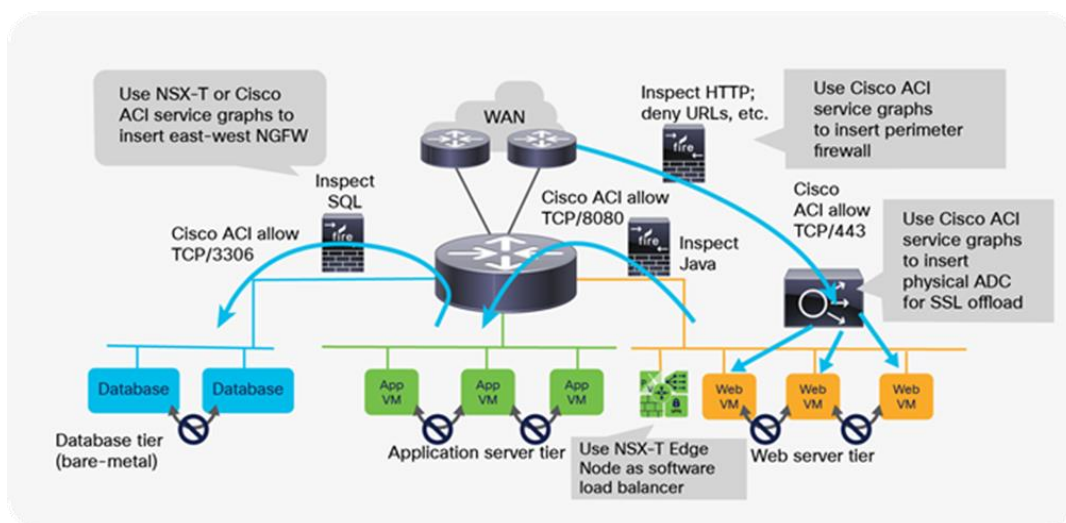


Рисунок 1.4 — Віртуалізація мережі VMware NSX з Cisco ACI

Іншим прикладом є кібератака на Equifax у 2017 році, в результаті якої було скомпрометовано особисті дані приблизно 145 мільйонів американців. Хакери скористалися відомою вразливістю в Apache Struts, яку Equifax не виправила

вчасно, незважаючи на наявність відповідного оновлення. Після проникнення в систему зломисники отримали доступ до внутрішніх серверів та баз даних, використовуючи облікові дані з дефолтними логінами та паролями. Відсутність двофакторної автентифікації та належного моніторингу дозволила їм залишатися непоміченими протягом 76 днів. Цей інцидент підкреслює критичну важливість своєчасного застосування оновлень безпеки, використання складних паролів та впровадження багатофакторної автентифікації для захисту корпоративних мереж [3].

Корпоративні мережі також часто інтегруються з хмарними інфраструктурами, що породжує нові виклики і в той же час нові можливості. Застосування гібридних хмарних рішень дозволяє частину обчислювальних ресурсів переносити в публічні хмари (AWS, Microsoft Azure, Google Cloud), тоді як критично важливі системи залишаються в межах локальної інфраструктури (рисунок 1.5) [3].



Рисунок 1.5 — Основні функції гібридної хмари

Це вимагає реалізації додаткових засобів забезпечення безпеки, швидкого з'єднання (наприклад, через VPN або спеціалізовані L2/L3 канали) і надійного механізму синхронізації даних.

Також варто згадати про інструменти централізованого управління мереженою інфраструктурою, такі як Cisco DNA Center, Aruba Central. Вони дозволяють адміністратору відстежувати стан мережі в реальному часі, автоматизувати оновлення конфігурацій, виявляти аномалії в роботі та оперативно

реагувати на загрози або несправності [4].

Із розвитком технологій корпоративні мережі стали основою для реалізації концепцій Інтернету речей (IoT), індустрії 4.0, систем розумних будівель та автоматизації виробництва. Це розширює кількість пристроїв, підключених до мережі, що вимагає ще вищого рівня масштабованості, сегментації та безпеки (рисунок 1.6). У зв'язку з цим активно впроваджуються нові стандарти бездротового зв'язку (Wi-Fi 6/6E/7), мережі з малими затримками (low-latency) та підтримкою великої кількості підключень [5].

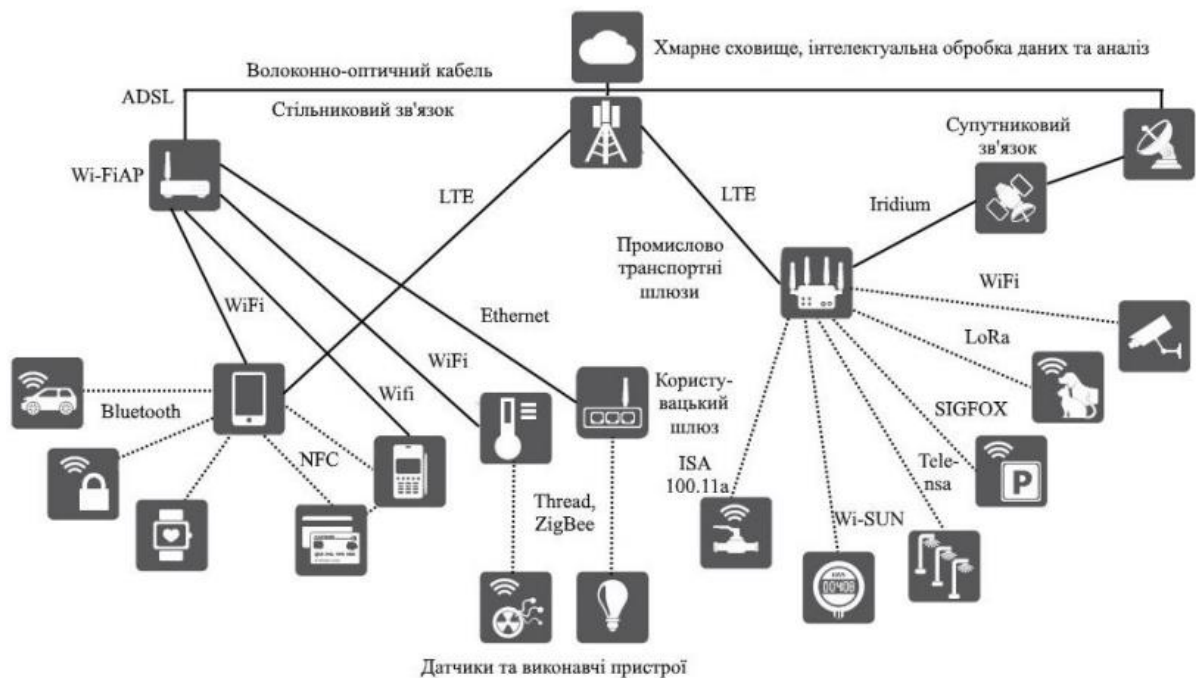


Рисунок 1.6 — Мережа з численними підключеними IoT-пристроями

У підсумку, сучасні технології побудови корпоративних комп'ютерних мереж є багатокомпонентною системою, що включає як фізичні, так і програмні елементи, які повинні працювати як єдине ціле, забезпечуючи безперебійну, безпечну та масштабовану інфраструктуру для функціонування організацій усіх рівнів.

1.2 Аналіз топологій комп'ютерних мереж

На етапі проєктування комп'ютерної мережі надзвичайно важливим є вибір відповідної топології, яка визначає фізичну та логічну структуру з'єднання між усіма її елементами: робочими станціями, серверами, маршрутизаторами,

комутаторами тощо. Від цього вибору залежить надійність, продуктивність, масштабованість та здатність мережі справлятися з високими навантаженнями й запобігати перевантаженням. У корпоративному середовищі, особливо там, де важливим є моніторинг навантаження й своєчасне реагування на пікові ситуації, ці параметри стають критично важливими [6].

Однією з найпростіших топологій є шинна (рисунок 1.7). У такій структурі всі пристрої підключені до єдиної магістралі, якою передаються дані. Всі пакети даних надсилаються в один канал і досягають кожного вузла, але приймає їх лише адресат. Основна перевага цієї топології — простота реалізації і мінімальні витрати на кабельну інфраструктуру. Проте з технічної точки зору вона має серйозні обмеження: при високому трафіку виникає колізія даних, що значно уповільнює роботу мережі. Крім того, у разі пошкодження магістралі — вся мережа виходить з ладу. Через це топологія шини практично не застосовується в сучасних мережах середнього та великого масштабу.

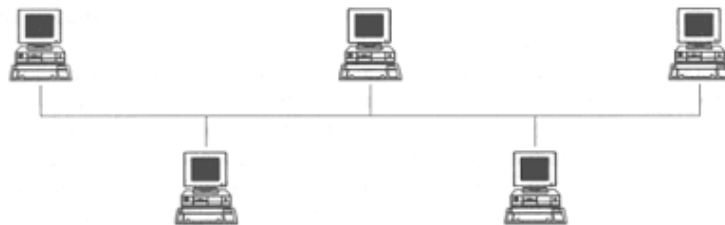


Рисунок 1.7 — Схема топології “Шина”

Кільцева топологія передбачає з’єднання пристроїв в замкнутий контур, де дані передаються послідовно від одного вузла до іншого (рисунок 1.8).

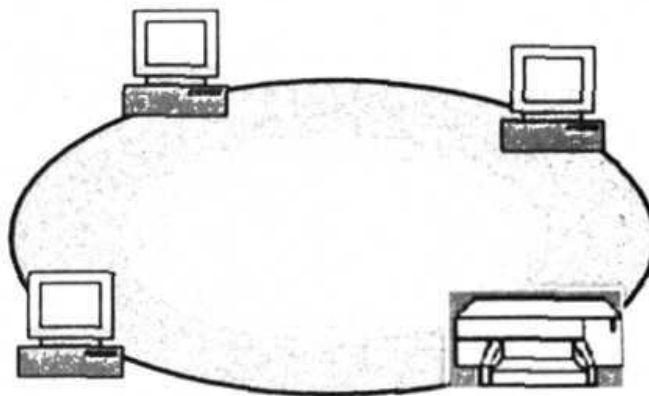


Рисунок 1.8 — Схема топології “Кільце”

Вона зменшує ймовірність конфліктів у передачі даних, оскільки кожен вузол має чітко визначену послідовність передачі. Однак у такій мережі існує серйозна вразливість: вихід з ладу одного вузла чи з'єднання призводить до повного припинення роботи мережі. Також у кільцевих структурах виникають затримки, пов'язані з необхідністю проходження даних через усі попередні вузли, що є критичним у системах із моніторингом навантаження в реальному часі [6].

Однією з найпоширеніших є зіркоподібна топологія, яка передбачає підключення всіх вузлів до одного центрального елемента — комутатора або маршрутизатора (рисунок 1.9).

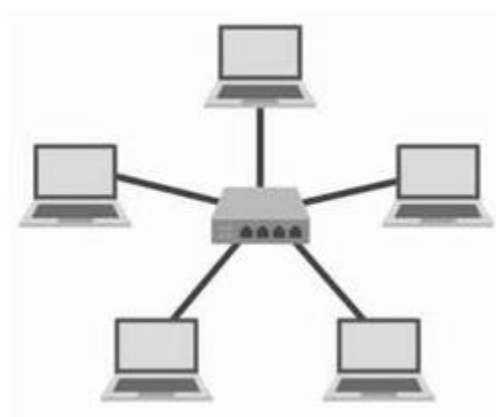


Рисунок 1.9 — Схема топології “Зірка”

Це дозволяє легко керувати трафіком, локалізувати проблемні ділянки, швидко виявляти перевантаження й здійснювати централізований моніторинг. Збої в одному з периферійних вузлів не впливають на інші, що підвищує надійність. Основною вразливістю залишається центральний вузол: при його виході з ладу вся мережа перестає функціонувати. Незважаючи на це, топологія зірки добре підходить для середніх корпоративних мереж з помірним трафіком [7].

Більш складною структурою є деревоподібна топологія, яка фактично є ієрархічним розширенням зірки (рисунок 1.10). У цій топології кілька зіркових сегментів поєднуються між собою через центральні вузли вищого рівня. Такий підхід дозволяє логічно організувати підрозділи компанії, оптимізувати розподіл трафіку, відокремити сегменти мережі та встановити пріоритети обробки даних. Ієрархія також сприяє зручному впровадженню системи моніторингу: можна окремо відслідковувати навантаження на кожному рівні мережі. Однак при

перевантаженні або несправності вузлів вищого рівня може порушитися робота значної частини мережі.

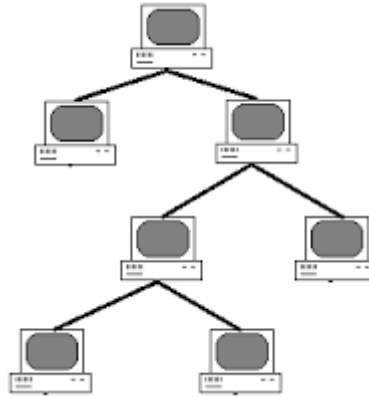


Рисунок 1.10 — Схема топології “Дерево”

Mesh-топологія (сітка) характеризується наявністю кількох шляхів з'єднання між усіма або окремими вузлами (рисунок 1.1).

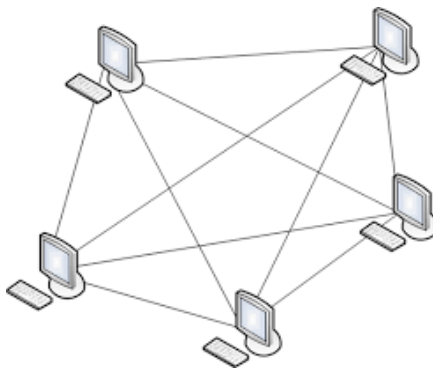


Рисунок 1.11 — Схема топології “Сітка”

Це забезпечує високу надійність, оскільки в разі виходу з ладу одного з'єднання — дані автоматично перенаправляються альтернативним маршрутом. Такі мережі здатні ефективно балансувати навантаження та уникати перевантажень. Проте mesh-топологія вимагає значних ресурсів для побудови, складного адміністрування й масштабного моніторингу, що обмежує її застосування у середніх компаніях [7].

У сучасних системах іноді використовують гібридну топологію, яка поєднує елементи кількох базових структур. Наприклад, можна поєднати зіркоподібну та деревоподібну топології для оптимального розподілу трафіку в офісній мережі з декількома відділами. Гібридна структура дозволяє врахувати особливості

архітектури компанії, однак її складність створює додаткове навантаження на систему моніторингу та адміністрування.

Окремо варто розглянути менш поширену, але перспективну топологію, яка належить до гібридних топологій — “квітка”, яка є варіацією зіркоподібної структури з кількома центральними вузлами, з’єднаними між собою в логічну “пелюстку” (рисунок 1.12).

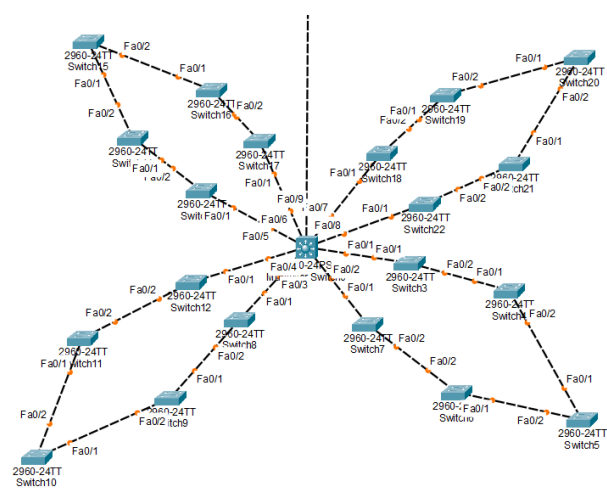


Рисунок 1.12 — Схема топології “Квітка”

У такій конфігурації кожен підрозділ компанії підключається до окремого центрального вузла, а самі центральні вузли з’єднані між собою кільцем або сіткою. Це дозволяє поєднати переваги зірки (локалізація вузлів), кільця (надійність між центрами) та сітки (резервування маршрутів). Топологія "квітка" особливо ефективна для впровадження системи моніторингу перевантажень, оскільки дозволяє відслідковувати навантаження як у локальних підмережах, так і між їх центральними вузлами. Така структура підвищує відмовостійкість, дає можливість балансувати трафік між пелюстками та швидко локалізувати проблемні ділянки (таблиця 1.1) [8].

У контексті розробки комп’ютерної мережі компанії з фокусом на моніторинг перевантажень, найбільш доцільним буде використання топології "квітка". Вона забезпечить гнучку ієрархічну побудову мережі, поєднує централізоване керування з високою надійністю, дозволяє реалізувати точний моніторинг трафіку на всіх рівнях мереженої структури та ефективно запобігати перевантаженням як у локальних сегментах, так і в міжсегментних з’єднаннях [8].

Таблиця 1.1 — Порівняльна таблиця топологій

Топологія	Складність реалізації	Надійність	Масштабованість	Моніторинг перевантаження	Відмовостійкість	Придатність для корпоративної мережі
Шина	Низька	Низька	Обмежена	Слабка	Низька	Низька
Кільцева	Середня	Середня	Обмежена	Посередня	Низька (при відсутності резерву)	Низька
Зіркоподібна	Середня	Висока (за винятком центру)	Висока	Висока	Середня	Висока
Деревоподібна	Висока	Середня–висока	Висока	Висока	Середня	Висока
Mesh (сітка)	Висока	Дуже висока	Складна у масштабуванні	Висока	Висока	Обмежено (через вартість і складність)
Гібридна	Висока	Висока	Висока	Висока	Висока	Висока
Квітка	Середня–висока	Висока	Висока	Дуже висока	Висока	Оптимальна

1.3 Огляд протоколів зв'язку для передачі даних

У сучасних комп'ютерних мережах ключову роль у забезпеченні ефективної та надійної передачі інформації відіграють протоколи маршрутизації. Саме вони дозволяють динамічно обирати оптимальні шляхи для пакета між пристроями у мережі, адаптуючись до змін у топології або навантаженні [9].

Протоколи маршрутизації поділяються на внутрішньомережної (IGP), які функціонують у межах однієї автономної системи, та міжмережні (EGP), призначені для обміну маршрутами між різними автономними системами. Найвідомішими представниками IGP є RIP, OSPF, IS-IS та EIGRP, тоді як BGP — це основний протокол для міжмережної маршрутизації [9].

Також протоколи класифікують за принципом роботи алгоритмів: на основі вектору відстаней, стану каналу, гібридні та шляхові. Протоколи вектору відстаней, такі як RIP, базуються на простому обміні таблицями маршрутів між сусідніми маршрутизаторами, де метрикою є кількість проміжних вузлів. Хоча RIP був одним із перших протоколів такого типу, він має суттєві обмеження: максимальна кількість стрибків — лише 15, а повна конвергенція мережі після зміни маршруту може займати значний час. Ці недоліки зумовлюють його обмежене використання

у сучасних мережах.

Іншим типом є протоколи стану каналу. Сюди належать OSPF та IS-IS. Вони забезпечують кожному маршрутизатору повне уявлення про структуру мережі, дозволяючи на основі зібраної карти самостійно обчислювати найкоротші шляхи за алгоритмом Дейкстри. OSPF широко застосовується в корпоративних мережах завдяки відкритому стандарту, високій швидкості конвергенції, підтримці ієрархічної побудови мережі та ефективному балансуванню навантаження. Подібним за структурою є протокол IS-IS, який більше поширений у мережах інтернет-провайдерів, хоча має потенціал для використання і в корпоративному сегменті [10].

Гібридним підходом вирізняється EIGRP — власна розробка компанії Cisco. Цей протокол поєднує переваги векторного та каналного підходів, використовуючи власний алгоритм DUAL для швидкого реагування на зміни у мережі. EIGRP забезпечує майже миттєву конвергенцію і дозволяє використовувати декілька маршрутів одночасно, навіть якщо їхні метрики різні. Проте його пропрієтарність історично обмежувала його застосування в мережах, де використовується обладнання інших виробників [10].

Окрему категорію становить BGP — протокол міжмереженої маршрутизації, який є основою функціонування Інтернету. Його особливістю є те, що він оперує не метриками у класичному сенсі, а політиками — через передачу маршрутів разом із атрибутами, зокрема шляхом проходження через автономні системи. Це дозволяє адміністраторам мереж гнучко керувати маршрутами, визначаючи пріоритети та правила обміну трафіком між організаціями чи провайдерами. Незважаючи на свою складність і повільну конвергенцію, BGP є незамінним для реалізації зв'язків між великими мережами та побудови відмовостійких схем доступу до Інтернету [10].

У корпоративних мережах протоколи маршрутизації забезпечують автоматизацію вибору маршрутів, що дозволяє оперативно реагувати на збої або перевантаження. Наприклад, OSPF та EIGRP, використовуючи тригерні оновлення та резервні шляхи, сприяють високій доступності послуг. Вони також дозволяють балансувати навантаження між кількома маршрутами, що є критично важливим при проектуванні мереж з урахуванням моніторингу перевантаження.

Хоча більшість протоколів не здійснюють прямого вимірювання навантаження, вони можуть враховувати пропускну здатність або затримку як метрики для визначення пріоритетного маршруту. Окрім того, у практиці застосовуються зовнішні системи моніторингу — SNMP, NetFlow, а також системи динамічного керування маршрутизацією, що допомагають адаптувати мережу до змінних умов [11].

Таблиця 1.2 — Порівняльна таблиця протоколів маршрутизації

Характеристика	RIP	OSPF	EIGRP	IS-IS	BGP
Тип протоколу	IGP, вектор відстаней	IGP, стан каналу	IGP, гібридний	IGP, стан каналу	EGP, шляховий вектор
Алгоритм маршрутизації	Bellman-Ford	Дейкстра	DUAL	Дейкстра	Політика на основі атрибутів
Максимальна кількість стрибків	15	Обмежень немає	224	Обмежень немає	Обмежень немає
Швидкість конвергенції	Низька	Висока	Дуже висока	Висока	Низька
Масштабованість	Низька	Висока	Середня	Висока	Дуже висока
Підтримка ієрархії	Ні	Так (area 0, інші області)	Обмежено (через summarizatio)	Так	Ні (підтримка політик дуже гнучка)
Балансування навантаження	Обмежене (рівні маршрут)	Так (ECMP)	Так (unequal load balancing)	Так	Обмежене
Використання у реальних мережах	Локальні або застарілі системи	Корпоративн і мережі	Мережі на обладнанні Cisco	Великі мережі	Інтернет, між провайдерські з'єднання
Підтримка відкритих стандартів	Так	Так	Частково	Так	Так
Надійність	Низька	Висока	Висока	Висока	Дуже висока

Отже, протоколи маршрутизації є основою у побудові гнучких, масштабованих та надійних мережних інфраструктур. Кожен із розглянутих протоколів має свої сильні сторони та недоліки, і їхній вибір залежить від конкретних потреб підприємства, структури мережі та вимог до швидкості відновлення зв'язку при збоях (таблиця 1.2).

1.4 Технології моніторингу стану комп'ютерних мереж

Надійність функціонування комп'ютерних мереж значною мірою

визначається здатністю адміністратора оперативно виявляти відхилення в роботі інфраструктури.

Для цього застосовуються спеціалізовані технології моніторингу, які забезпечують безперервне спостереження за технічним станом мережі та її компонентів у режимі реального або наближеного до реального часу.

Ключовим протоколом, що лежить в основі більшості систем моніторингу, є SNMP (Simple Network Management Protocol) (рисунок 1.13).

Завдяки йому адміністратор може централізовано отримувати інформацію від мережного обладнання: інтерфейсні лічильники, показники завантаженості процесора, статистику по пам'яті, температуру, стан каналів тощо. Інформація організовується у вигляді MIB (Management Information Base), що містить ієрархічно структуровані змінні, які можуть бути як доступні тільки для читання, так і для запису. Найновіша версія SNMPv3 дозволяє використовувати шифрування та автентифікацію, що критично важливо для захищених середовищ [12].

Інструментом глибшого аналізу мережного трафіку виступає технологія NetFlow. Вона забезпечує агрегацію даних про трафік у вигляді потоків (flows), що дозволяє не лише визначати кількість переданих байтів, а й аналізувати поведінку користувачів або сервісів у мережі. Це особливо ефективно при виявленні перевантаження, прихованих каналів витоку даних, несанкціонованих підключень чи сплесків активності.

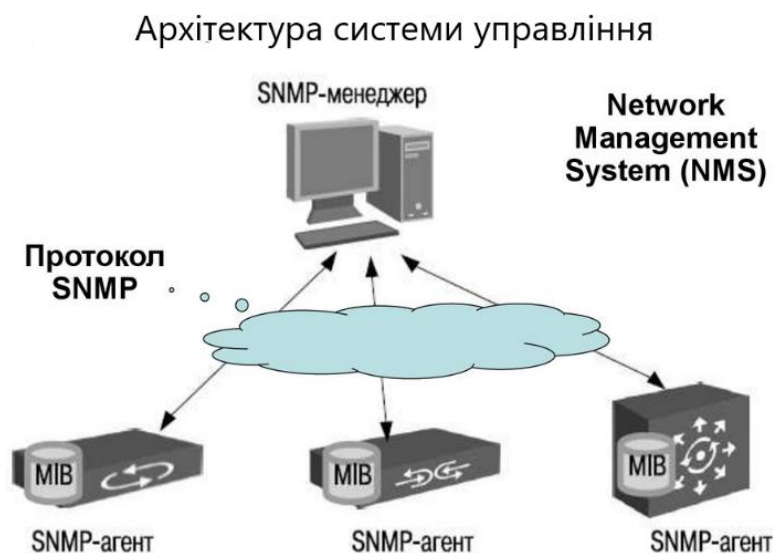


Рисунок 1.13 — Принцип роботи протоколу SNMP

Альтернативним підходом є sFlow — технологія вибіркового семплювання трафіку, яка дозволяє досягти високої масштабованості за рахунок зменшення обсягу оброблюваної інформації. На відміну від NetFlow, який формує потоки за визначеним шаблоном, sFlow передає випадкові зразки пакетів разом зі статистикою інтерфейсів, що робить його придатним для використання в середовищах з великою кількістю високошвидкісних з’єднань [12].

У базових сценаріях моніторингу стану мережі застосовуються ICMP-запити, які дозволяють перевірити доступність вузлів (ping) та простежити маршрут проходження трафіку (traceroute). Ці інструменти, попри свою простоту, є незамінними при побудові базових механізмів оповіщення про втрату зв’язку або зростання затримок.

На рівні програмного забезпечення широке розповсюдження отримали такі системи як Zabbix, Nagios, Prometheus, PRTG та інші. Вони дозволяють централізовано збирати дані, обробляти метрики, створювати тригери, генерувати повідомлення про помилки та візуалізувати стан мережі [13].

Zabbix, наприклад, підтримує активні та пасивні агенти, сценарії перевірок, а також інтеграцію з системами керування подіями (рисунок 1.14). Prometheus у поєднанні з Grafana дозволяє будувати гнучкі дашборди з глибокою аналітикою історичних даних.



Рисунок 1.14 — Графіки навантаження мережі у Zabbix

Окремим напрямком є візуалізація результатів моніторингу. Графіки навантаження, мапи зв’язності вузлів, індикатори доступності та кольорові статуси

значно спрощують роботу адміністратора. Сучасні системи часто підтримують інтеграцію з зовнішніми системами автоматизації та API, що дозволяє побудувати повноцінну екосистему для реактивного або навіть активного адміністрування [13].

Таким чином, сучасні технології моніторингу забезпечують не лише збирання статистики, а й трансформацію її у корисну інформацію, яка дозволяє приймати обґрунтовані рішення щодо підтримки, оптимізації або модернізації мережної інфраструктури.

1.5 Проблеми перевантаження мережі та методи їх вирішення

Перевантаження комп'ютерної мережі є однією з критичних проблем, яка призводить до деградації якості обслуговування, збільшення затримок, втрати пакетів і в кінцевому підсумку — до зниження загальної продуктивності ІТ-систем. Це явище виникає тоді, коли обсяг переданих даних перевищує доступну пропускну здатність мережних каналів або обчислювальні ресурси мережних пристроїв. Типовими причинами перевантаження є висока щільність трафіку в пікові періоди, неефективна маршрутизація, надлишкова ретрансляція широкомовних запитів, циклічні петлі, помилки в конфігурації QoS або DDoS-атаки [14].

Найбільш поширеним наслідком перевантаження є збільшення часу затримки (latency), що критично впливає на чутливі до затримок сервіси — наприклад, IP-телефонію, відеоконференції або онлайн-ігри. Окрім цього, при перевищенні буферної ємності маршрутизаторів або комутаторів виникає черга пакетів, яка зрештою призводить до їх відкидання. Це провокує повторну передачу в TCP-з'єднаннях, що, у свою чергу, ще більше підвищує навантаження [14].

Для виявлення перевантаження зазвичай застосовується поєднання активного моніторингу (через SNMP, NetFlow, sFlow) та пасивного аналізу логів пристроїв. Завдяки цьому можна визначити не лише сам факт перевантаження, а й його джерело — конкретний сегмент мережі або тип трафіку, що його спричинив.

До основних методів вирішення проблеми перевантаження належить впровадження механізмів якості обслуговування (QoS). QoS дозволяє класифікувати трафік, призначаючи пріоритети різним типам даних (рисунок 1.15).

Наприклад, критичний для бізнесу трафік отримує вищий пріоритет порівняно з фоновими завантаженнями чи резервним копіюванням. Також впроваджуються механізми контролю черг, формування трафіку (traffic shaping), маркування (DSCP, CoS) та обмеження швидкості (rate limiting).

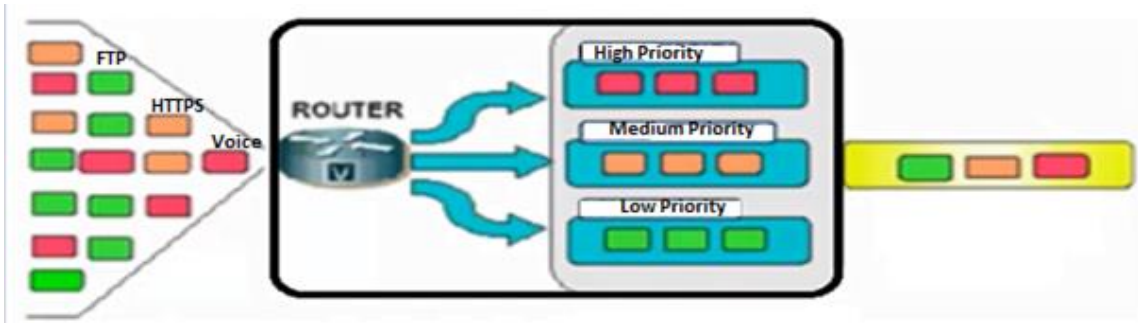


Рисунок 1.15 — Приклад роботи пріоритетності трафіку

Другий ефективний підхід — балансування навантаження між кількома маршрутами або каналами (рисунок 1.16).

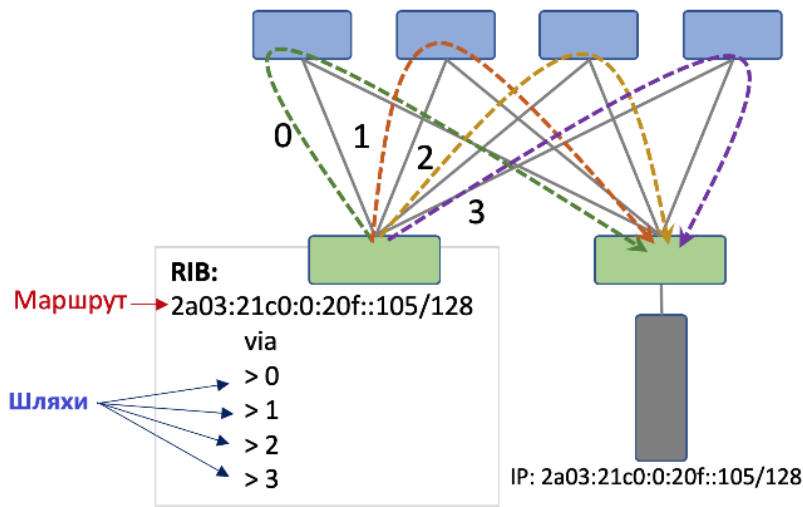


Рисунок 1.16 — Принцип роботи ECMP з декількома рівнозначними маршрутами

Зокрема, технологія ECMP (Equal-Cost Multi-Path) дозволяє одночасно використовувати декілька маршрутів з однаковою метрикою для передачі трафіку, що підвищує загальну пропускну здатність і знижує ризик локальних перевантажень. У складніших архітектурах використовується також динамічне перенаправлення потоків на основі поточної завантаженості каналів [15].

Окрім того, при проєктуванні масштабованих мереж критично важливо

впроваджувати ієрархічну структуру топології (наприклад, із розподілом на ядро, розподільчий і доступний рівні), що дозволяє локалізувати перевантаження і мінімізувати їхній вплив на загальну інфраструктуру.

У середовищах з інтенсивною взаємодією хостів доцільним є також обмеження широкомовного трафіку, застосування VLAN-ів та механізмів IGMP Snooping, які запобігають надмірній трансляції небажаного трафіку на всі порти мережі [15].

Важливо, що виявлення та усунення перевантажень неможливе без постійного моніторингу. Тому ключовим супровідним інструментом завжди є системи візуалізації та аналізу трафіку, які дають змогу оперативно фіксувати зміну патернів навантаження і відповідно до цього адаптувати політики маршрутизації або QoS.

У підсумку перевантаження мережі є складною проблемою, що потребує комплексного підходу для вирішення. Ефективність її усунення досягається завдяки поєднанню правильного проєктування мережі, використання адаптивних алгоритмів маршрутизації, впровадження механізмів QoS та постійного моніторингу стану мережі для своєчасного виявлення і реагування на зміни в навантаженні [15]

2 РОЗРОБКА СТРУКТУРИ КОРПОРАТИВНОЇ МЕРЕЖІ КОМПАНІЇ

2.1 Розробка загальної структури компанії

Як комп'ютерну мережу компанії яка буде розроблятися було обрано мережу інтернет-провайдера. Для розгортання масштабної мережі інтернет-провайдера потрібний великий фізичний простір та логічно організована структура мережі, тому на етапі розробки загальної архітектури мережі ключовим завданням є створення логічної та фізичної моделі, що враховує організаційну структуру підприємства, кількість відділів та їх взаємодію між собою, а також безпеку самої будівлі організації, доступність до обладнання та логістику для можливості легкого масштабування мережі. Мережна архітектура має бути побудована таким чином, щоб забезпечити ефективний, зручний, швидкий і захищений обмін даними між усіма підрозділами організації [16].

Фізична модель компанії (рисунок В.1) має центральний корпус, ліворуч від якого знаходяться служба безпеки, центр дзвінків і ІТ-відділ. У центрі розташована серверна кімната, що межує праворуч із відділом монтажників і відділом обслуговування обладнання. Нижче — адміністративний офіс і відділ програмістів. По обидва боки корпусу додатково розміщено ще по одній кімнаті служби безпеки.

Підстанція 1 має службу безпеки зліва, під нею — відділ обслуговування обладнання, а праву частину займає серверна. Аналогічно і з підстанцією 2.

Компанія має комплексну внутрішню структуру, де кожен відділ виконує визначену функцію у межах єдиної мережної системи. Усі підрозділи — включаючи серверну кімнату, службу безпеки, ІТ-відділ, адміністративний офіс керівництва, відділ програмістів, технічний відділ, відділ монтажників, відділ підключення абонентів та центр прийняття дзвінків — взаємодіють між собою через централізовану мережу, що побудована за ієрархічним принципом. Така структура дозволяє ефективно координувати дії всіх підрозділів, забезпечуючи надійний обмін даними між технічними та адміністративними компонентами системи [17]. Окремо розташовані підстанції, що також включають ключові функціональні блоки, інтегруються до загальної мережі через головну станцію, яка виступає ядром управління, обслуговування та адміністрування всієї інфраструктури компанії інтернет-провайдера (рисунок 2.1).

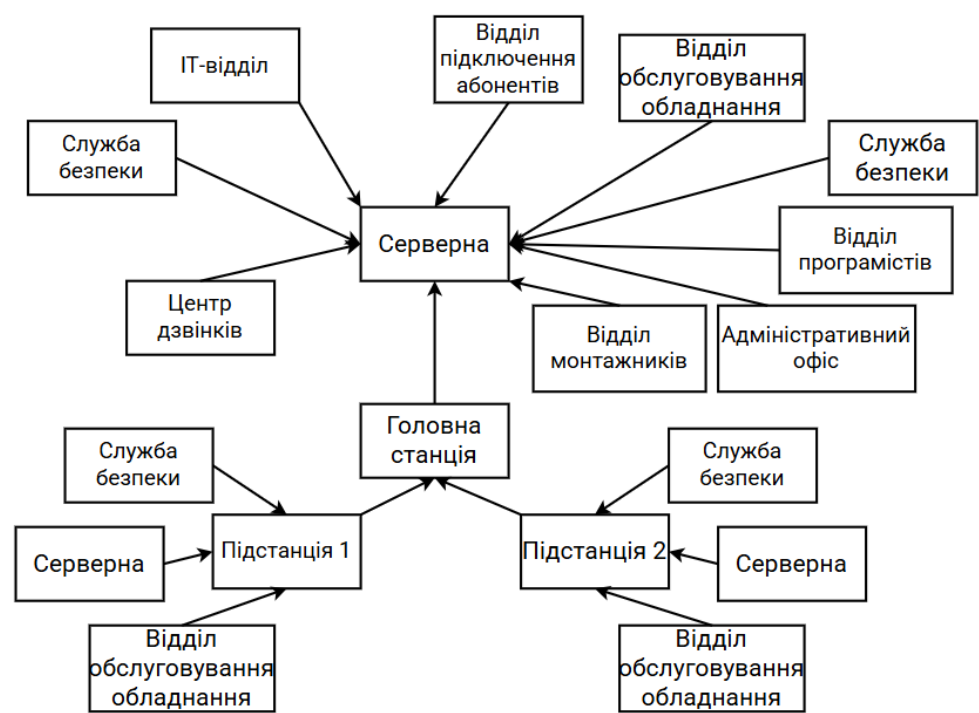


Рисунок 2.1 — Структурна схема компанії інтернет-провайдера

Кожна підстанція включатиме серверні кімнати, кімнати служби безпеки та відділ обслуговування обладнання. Керування здійснюватиметься безпосередньо із основної станції, де розташовані всі ключові відділи керування, адміністрування та обслуговування.

Центральним елементом мережі виступає серверна кімната, в якій розміщується критичне обладнання: маршрутизатори, комутатори ядра, сервери баз даних, файлові сервери, поштові та вебсервери, а також засоби резервного копіювання. До серверної підключаються всі інші сегменти мережі через основний комутаційний вузол, що забезпечує високу пропускну здатність.

Служба безпеки підприємства забезпечує фізичний захист та запобігає несанкціонованому доступу до апаратних компонентів мережі сторонніми особами. Вона забезпечує охорону приміщень, у яких розміщується критичне мережне обладнання, зокрема серверна кімната, комутаційні шафи, склад з обладнанням.

ІТ-відділ також розміщується в окремій кімнаті та має доступ до інструментів адміністрування, управління користувачами, оновлення ПЗ, віддаленого доступу та моніторингу. Для забезпечення гнучкості керування мережею, їхній доступ розширено до серверної інфраструктури та систем резервного копіювання даних.

Адміністративна частина підприємства являє собою офіс, де сидить керівництво. Цей офіс повинен мати пріоритетний доступ до ключових бізнес-додатків і аналітичних систем. Для цих відділів організовано окремий виділений канал зв'язку в мережі із підвищеним рівнем QoS, що забезпечує швидку передачу даних та окремий доступ до системи моніторингу мережі.

Відділ програмістів вимагає високої пропускної здатності для доступу до репозиторіїв коду, тестових середовищ, а також можливості запуску віртуальних серверів. Цей сегмент ізольовано для безпеки, але при потребі забезпечено обмежену комунікацію з іншими підрозділами.

Відділ монтажників відповідає за початкове фізичне встановлення мережної інфраструктури. До його завдань входить прокладання основних магістральних кабельних трас, через які буде проходити основний трафік в мережі, а також монтаж вузлів доступу, тобто комутаторів для підключення абонентів. Працівники цього відділу також займаються маркуванням кабелів і тестуванням ліній після прокладки. Відділ працює у тісній взаємодії з IT-відділом, дотримуючись вимог щодо стандартів побудови структурованих кабельних систем.

Відділ обслуговування обладнання здійснює технічну підтримку вже встановленої мережної інфраструктури. Основні завдання включають діагностику та ремонт несправного мережного обладнання, його заміну, а також регулярне технічне обслуговування. Важливою складовою є ревізія та профілактичне обслуговування джерел безперебійного живлення (UPS) та інверторів, що забезпечують стабільну роботу критичних елементів мережі в умовах тимчасових відключень електропостачання. Цей відділ забезпечує вчасне усунення збоїв і підтримку працездатності мережної інфраструктури відповідно вимогам надійності та безпеки мережі.

Відділ підключення абонентів виконує запити на підключення абонентів, тестування з'єднань та конфігурацію обладнання клієнтів. Для нього передбачено доступ до баз даних клієнтів. Він також зв'язується із IT-відділом та відділом монтажників у разі виникнення проблем підключення абонентів.

Центр дзвінків обробляє вхідні та вихідні дзвінки, має доступ до IP-телефонії, внутрішнього інформаційного порталу та системи обліку для перегляду даних про

абонентів. Цей відділ тісно співпрацює із ІТ-відділом, для врахування технічних аспектів при обробці запитів на підключення абонентів.

Структура компанії інтернет-провайдера організована, так що загальна архітектура базується на зручній фізичній організації приміщень та логічній взаємодії між ключовими відділами, що дозволяє забезпечити стабільну, масштабовану та ефективну роботу мережі в умовах постійного її розширення.

2.2 Проектування логічної схеми корпоративної мережі компанії інтернет-провайдера

В основу логічної схеми покладено ієрархічну структуру з центральним мультисвічем, до якого під’єднані всі основні сегменти мережі, зокрема серверні ресурси, користувацькі зони, підстанції та мережі агрегації (рисунок 2.2).

Логічна топологія комп’ютерної мережі компанії інтернет-провайдера, що розглядається в даній роботі, побудована за схемою «квітка», де центральним вузлом є головна станція компанії, вузлами розповсюдження є дві віддалені підстанції, а самим вузлами агрегації є комутатори L3 рівня з вузлами доступу, в яких знаходяться вузли доступу. Така організація мережі дозволяє досягти централізованого контролю, надійної маршрутизації та ефективного керування всіма сегментами системи [17].

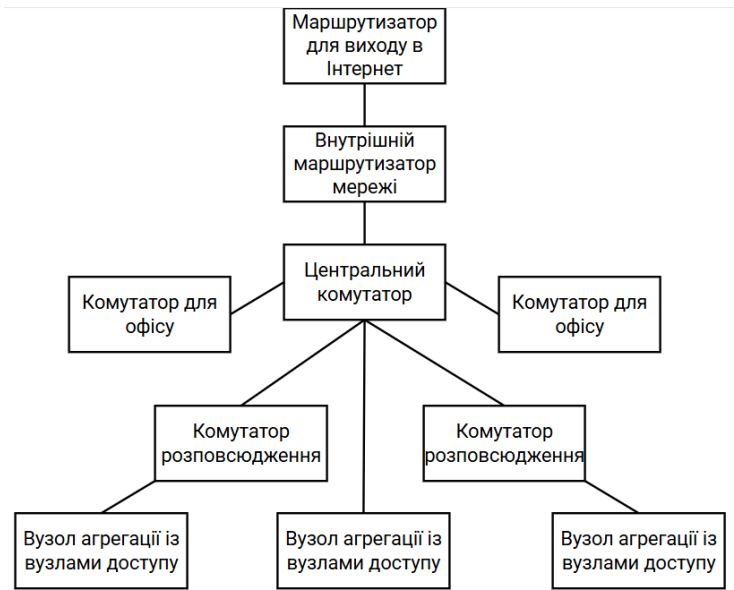


Рисунок 2.2 — Схема підключення обладнання

Головна станція виконує роль ядра мережі, саме тут розміщуються ключові структурні підрозділи компанії, такі як серверна кімната, служба безпеки, ІТ-відділ, адміністративний офіс керівництва, відділ програмістів, технічний відділ, відділ монтажників, відділ підключення абонентів та центр прийняття дзвінків. Усі дані з периферійних підстанцій надходять до цієї центральної точки, де відбувається їх обробка, аналіз і розподіл відповідно до заданих правил маршрутизації та безпеки (рисунк 2.3).

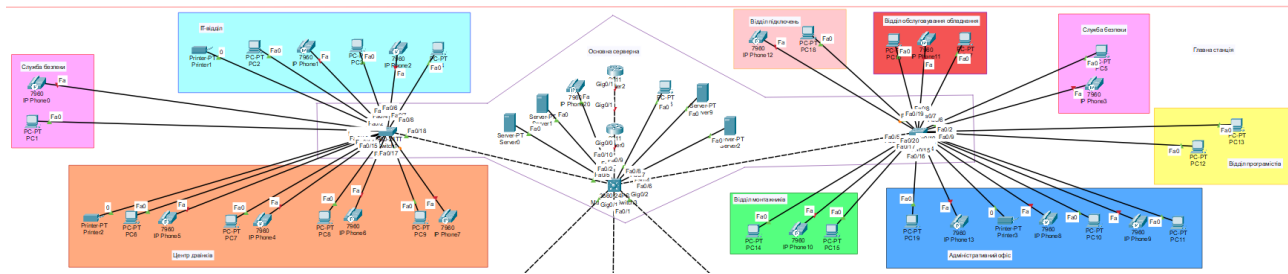


Рисунок 2.3 — Логічна схема головної станції

Дві підстанції, які є незалежними сегментами, розташовані окремо від основної компанії. Вони мають власні серверні кімнати, приміщення служби безпеки та відділи обслуговування обладнання (рисунк 2.4).

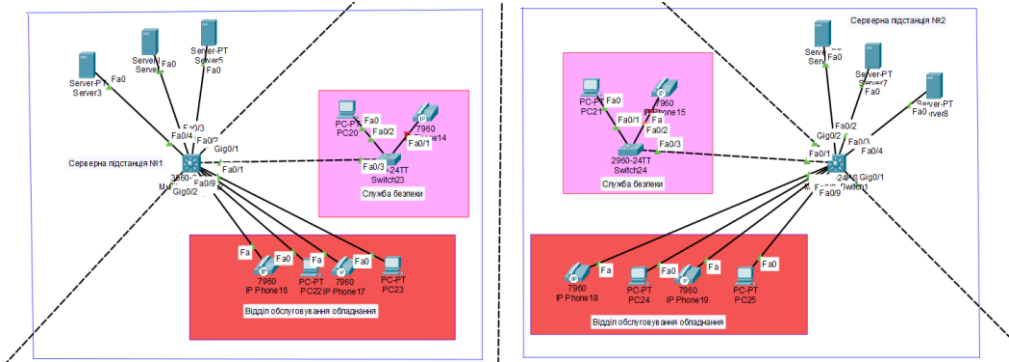


Рисунок 2.4 — Логічна схема двох підстанцій

Попри фізичну віддаленість, логічно ці підстанції тісно інтегровані в загальну інфраструктуру через комутатори та основний маршрутизатор, що забезпечує єдиний інформаційний простір та стабільне з'єднання між усіма елементами мережі.

Передача даних у мережі відбувається через центральний маршрутизатор, до якого підключені всі комутатори основної станції та підстанцій. Така структура

дозволяє організувати надійний трафік даних між усіма підрозділами компанії, незалежно від їхнього фізичного розташування.

Квітка представляє собою вузол агрегації, який знаходиться у центрі самої квітки, до якого під’єднуються вузли доступу по колу, таким чином що вони утворюють пелюстку.

Логічна топологія у формі квітки забезпечує чітке централізоване управління, розподілене підключення та ефективне функціонування всієї мережі інтернет-провайдера (рисунок 2.5).

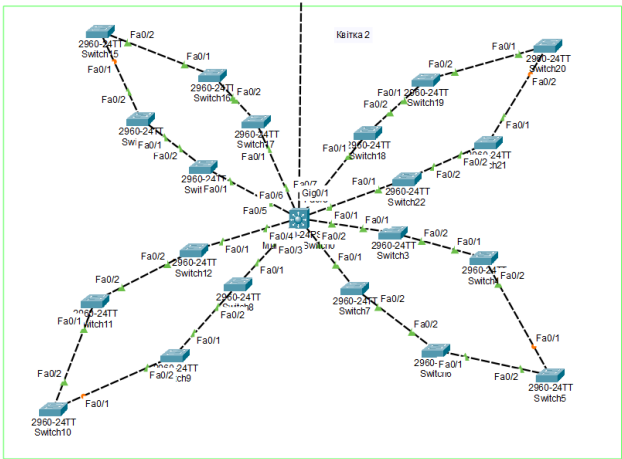


Рисунок 2.5 — Логічна схема “квітки” з вузлом агрегації та вузлами доступу

Кожна її пелюстка працює в рамках загальної структури з можливістю відстеження стану навантаження, автоматичного реагування на перевантаження та підтримки стабільної роботи в умовах зростання трафіку.

В результаті від основного мултисвіча на головній станції відходять дві незалежні підстанції, логічно інтегровані до мережі через виділені канали зв’язку з основним маршрутизатором. Кожна підстанція має власну серверну кімнату, службу безпеки та відділ обслуговування обладнання. Вони виконують роль вузлів розповсюдження і забезпечують локальну обробку трафіку, підвищуючи надійність і масштабованість мережі.

До підстанцій підключаються вузли агрегації — потужні комутатори, від яких відгалужуються пелюстки, що складаються з вузлів доступу. Саме через ці пелюстки відбувається підключення кінцевих користувачів — абонентів провайдера (рисунок 2.6).

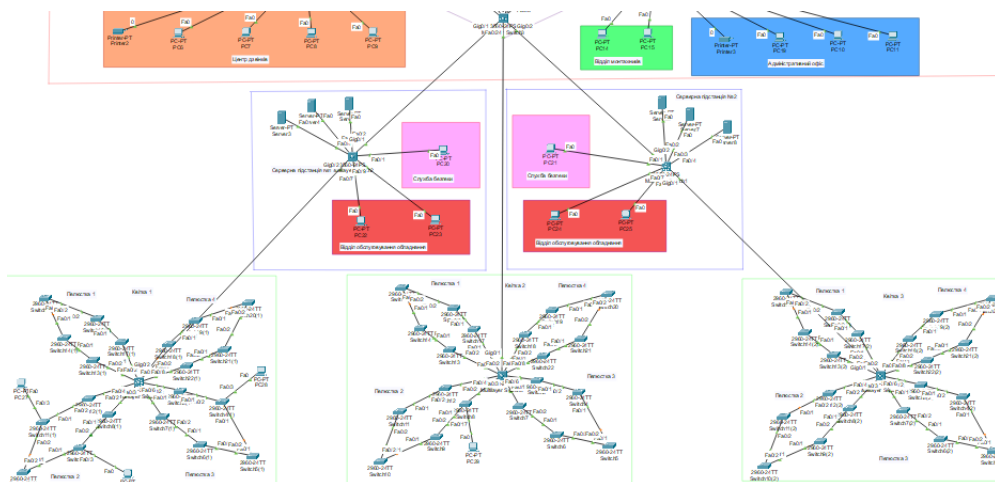


Рисунок 2.6 — Схема зв'язків головного корпусу, підстанції та квітки

Це дозволить ефективно обслуговувати великі географічні зони з мінімальними затримками, забезпечуючи централізований контроль і швидке реагування на зміни навантаження.

Повна схема комп'ютерної мережі компанії із підключенням всіх сегментів мережі наведена на рисунку В.2.

Важливо врахувати не лише фізичне з'єднання елементів та логічну топологію, а й ефективне управління трафіком та станом інфраструктури в режимі реального часу. Тому для забезпечення високої продуктивності та відмовостійкості, до архітектури мережі впроваджено систему моніторингу навантаження. Вона використовує сучасні програмні засоби на основі протоколу SNMP та програмного забезпечення Zabbix і дозволяє в реальному часі спостерігати за роботою мережних пристроїв, виявляти перевантаження, втрати пакетів чи збої.

У разі перевантаження мережних вузлів система моніторингу ініціює механізми балансування навантаження. Завдяки цьому, дані динамічно перенаправляються на менш завантажені канали чи пристрої, що дозволяє уникнути затримок, підвищити швидкість обробки запитів і забезпечити рівномірне використання ресурсів. Технологія балансування також позитивно впливає на надійність мережі, оскільки мінімізує ризик простоїв у випадку збоїв окремих елементів мережі.

Вихід у мережу Інтернет здійснюється через центральний маршрутизатор, на

якому налаштовано мережну адресацію NAT, що дозволяє всім пристроям внутрішньої мережі використовувати одну публічну IP-адресу для зовнішнього з'єднання. Для доступу до Інтернету використовуються публічні DNS- і Web-сервери. DHCP-сервер, що також працює на центральному вузлі, автоматично призначає IP-адреси всім пристроям локальної мережі. Для підвищення рівня безпеки застосовуються списки контролю доступу (ACL), які визначають дозволені маршрути та заборонені з'єднання між сегментами мережі.

2.3 Розрахунок адресного простору мережі компанії

Під час побудови корпоративної мережі було проведено детальний розрахунок адресного простору з урахуванням потреб масштабування, резервування та ефективної маршрутизації.

У серверній кімнаті розміщено чотири сервери. Для їх об'єднання виділена окрема підмережа з IP-адресою 10.0.1.0/27, що забезпечує діапазон для підключення хостів від 10.0.1.1 до 10.0.1.30. Така маска дозволяє використовувати до 30 хостів і залишає запас для додавання серверів або інфраструктурних пристроїв у майбутньому. Обрана маска є фіксованою для серверного сегмента з урахуванням стабільного навантаження.

До центрального мультисвіча також підключено два мультисвіча користувацького доступу. Перший обслуговує вісім комп'ютерів, сім IP-телефонів і два принтери. Для цього сегмента використано підмережу 10.0.1.32/27 з діапазоном для хостів IP-адрес від 10.0.1.33 до 10.0.1.62. Це дозволяє розмістити до 30 пристроїв з достатнім резервом. Маска /27 тут також є фіксованою та обрана на основі наявної кількості пристроїв з допуском на ріст.

Другий користувацький мультисвіч забезпечує роботу одинадцяти комп'ютерів і семи телефонів. Для нього виділена підмережа 10.0.1.64/27 (діапазон IP-адрес для хостів від 10.0.1.65 до 10.0.1.94). Як і в попередньому випадку, маска /27 забезпечує до 30 IP-адрес — цього вистачає для поточних потреб і майбутніх підключень.

Окремий сегмент займають підстанції, яких у мережі дві. Для першої підстанції виділена підмережа 10.0.2.0/27, що забезпечує IP-адреси в діапазоні

10.0.2.1–10.0.2.30. Цього вистачає для трьох серверів, трьох комп'ютерів і трьох телефонів, а також додаткових пристроїв у разі розширення. Аналогічно, для другої підстанції призначено підмережу 10.0.2.32/27 з IP-адресами 10.0.2.33–10.0.2.62.

До кожної підстанції та центрального сегмента підключено окремі мультисвічі агрегації, що ведуть до квіткової топології. У кожній квітці — чотири пелюстки, по п'ять комутаторів у кожній. Для кожної пелюстки центрального мультисвіча виділено по окремій підмережі:

- пелюстка 1: 10.0.0.0/27 (IP-адреси: 10.0.0.1–10.0.0.30);
- пелюстка 2: 10.0.0.32/27 (IP-адреси: 10.0.0.33–10.0.0.62);
- пелюстка 3: 10.0.0.64/27 (IP-адреси: 10.0.0.65–10.0.0.94);
- пелюстка 4: 10.0.0.96/27 (IP-адреси: 10.0.0.97–10.0.0.126).

Для квітки підстанції 1 використано чотири підмережі:

- пелюстка 1: 10.0.3.0/27 (10.0.3.1–10.0.3.30);
- пелюстка 2: 10.0.3.32/27 (10.0.3.33–10.0.3.62);
- пелюстка 3: 10.0.3.64/27 (10.0.3.65–10.0.3.94);
- пелюстка 4: 10.0.3.96/27 (10.0.3.97–10.0.3.126).

Для квітки підстанції 2 виділено чотири підмережі:

- пелюстка 1: 10.0.4.0/27 (10.0.3.1–10.0.3.30);
- пелюстка 2: 10.0.4.32/27 (10.0.3.33–10.0.3.62);
- пелюстка 3: 10.0.4.64/27 (10.0.3.65–10.0.3.94);
- пелюстка 4: 10.0.4.96/27 (10.0.3.97–10.0.3.126).

У разі збільшення кількості пелюсток у квітці адресний простір має можливість для розширення.

Кожна з пелюсток розрахована на підключення до 22 комп'ютерів, тому діапазон у 30 адрес є виправданим і залишає достатній резерв. Усі ці підмережі використовують фіксовану маску /27 (255.255.255.224), яка дозволяє оптимально використовувати простір приватної адресації без надлишкових витрат.

Для організації з'єднання між маршрутизаторами, а також підключення до провайдера, використано маску /30 (255.255.255.252), яка дозволяє використовувати лише дві адреси, що є достатнім для point-to-point каналів:

Для з'єднання між маршрутизаторами виділено підмережу 192.168.0.0/30

(діапазон доступних IP-адрес для адресації хостів: 192.168.0.1–192.168.0.2)

Під з'єднання з провайдером виділено підмережу: 192.168.0.4/30 (діапазон доступних IP-адрес для адресації хостів: 192.168.0.5–192.168.0.6).

Розподіл адресного простору та застосування ESMР, і впровадження BGP, розроблена мережна архітектура забезпечує надійність, масштабованість і гнучкість у керуванні як внутрішнім, так і зовнішнім трафіком.

Використання мережі класу А (10.0.0.0/8) повністю відповідає потребам інтернет-провайдера. Цей діапазон дозволяє не лише покрити існуючу інфраструктуру з сотнями комутаторів, серверів і кінцевих пристроїв, а й передбачає значний резерв під динамічне розширення — як територіальне, так і функціональне тому що має більше 16 мільйонів адрес, з яких використано лише частину.

Застосування масок /27 для локальних сегментів гарантує оптимальну щільність адресного простору з достатнім резервом, що критично важливо у розгалуженій мережі з великою кількістю точок підключення. Для point-to-point з'єднань між маршрутизаторами і з провайдером була застосована маска /30, що мінімізує втрати адрес і відповідає стандартам організації WAN-зв'язків.

Весь розрахований адресний простір для підмереж в мережі представлений на таблиці 2.1

Послідовне з'єднання маршрутизаторів дозволяє впровадити розмежування зон відповідальності, спростити адміністрування та реалізувати повноцінне резервування зовнішнього підключення. Застосування внутрішнього BGP (iBGP) між маршрутизаторами разом із зовнішнім BGP (eBGP) для взаємодії з провайдером дозволяє забезпечити контроль над оголошенням префіксів, ефективний розподіл маршрутів та швидке перемикання у випадку відмови одного з ланцюгів.

Загальна логічна модель побудована з урахуванням фізичної ізоляції сегментів, чіткої топології у вигляді квіток для агрегації та використання мультисвічів як базових елементів кожного рівня. Такий підхід дає змогу підтримувати велику кількість пристроїв, реалізувати балансування навантаження, зберігати стабільну продуктивність при зростанні кількості клієнтів, а також бути

готовими до інтеграції нових сервісів у майбутньому.

Таблиця 2.1 — Розрахований адресний простір для мережі

Net	Адреса мережі	Перша IP	Остання IP	Широкомовна адреса	Маска	Призначення
1	10.0.0.0	10.0.0.1	10.0.0.30	10.0.0.31	255.255.255.224	Пелюстка 1 центральної квітки
2	10.0.0.32	10.0.0.33	10.0.0.62	10.0.0.63	255.255.255.224	Пелюстка 2 центральної квітки
3	10.0.0.64	10.0.0.65	10.0.0.94	10.0.0.95	255.255.255.224	Пелюстка 3 центральної квітки
4	10.0.0.96	10.0.0.97	10.0.0.126	10.0.0.127	255.255.255.224	Пелюстка 4 центральної квітки
5	10.0.1.0	10.0.1.1	10.0.1.30	10.0.1.31	255.255.255.224	Серверна кімната
6	10.0.1.32	10.0.1.33	10.0.1.62	10.0.1.63	255.255.255.224	Користувацький сегмент 1
7	10.0.1.64	10.0.1.65	10.0.1.94	10.0.1.95	255.255.255.224	Користувацький сегмент 2
8	10.0.2.0	10.0.2.1	10.0.2.30	10.0.2.31	255.255.255.224	Підстанція 1
9	10.0.2.32	10.0.2.33	10.0.2.62	10.0.2.63	255.255.255.224	Підстанція 2
10	10.0.3.0	10.0.3.1	10.0.3.30	10.0.3.31	255.255.255.224	Квітка підстанції 1, пелюстка 1
11	10.0.3.32	10.0.3.33	10.0.3.62	10.0.3.63	255.255.255.224	Квітка підстанції 1, пелюстка 2
12	10.0.3.64	10.0.3.65	10.0.3.94	10.0.3.95	255.255.255.224	Квітка підстанції 1, пелюстка 3
13	10.0.3.96	10.0.3.97	10.0.3.126	10.0.3.127	255.255.255.224	Квітка підстанції 1, пелюстка 4
14	10.0.4.0	10.0.3.1	10.0.3.30	10.0.3.31	255.255.255.224	Квітка підстанції 2, пелюстка 1
15	10.0.4.32	10.0.3.33	10.0.3.62	10.0.3.63	255.255.255.224	Квітка підстанції 2, пелюстка 2
16	10.0.4.64	10.0.3.65	10.0.3.94	10.0.3.95	255.255.255.224	Квітка підстанції 2, пелюстка 3
17	10.0.4.96	10.0.3.97	10.0.3.126	10.0.3.127	255.255.255.224	Квітка підстанції 2, пелюстка 4
18	192.168.0.0	192.168.0.1	192.168.0.2	192.168.0.3	255.255.255.252	З'єднання між маршрутизаторами
19	192.168.0.4	192.168.0.5	192.168.0.6	192.168.0.7	255.255.255.252	З'єднання з провайдером

2.4 Вибір мережного обладнання

Важливу роль у забезпеченні надійного з'єднання, високої пропускної

здатності та масштабованості відіграє правильно підібране мережне обладнання.

Мережна архітектура реалізована за логічною топологією типу «квітка», де центральним елементом є головна станція з вузлом агрегації, що виконує роль ядра системи. До цього центру сходяться зв'язки від двох підстанцій, які виступають у ролі вузлів розповсюдження, а комутатори доступу на «пелюстках» відповідають за підключення локальних користувачів. Такий підхід дозволяє централізувати керування, зменшити затримки трафіку та чітко розмежувати зони відповідальності між сегментами мережі.

Особливої уваги потребують вузли агрегації, до яких надходить трафік від численних вузлів доступу. З огляду на вимоги до продуктивності, відмовостійкості та підтримки протоколів балансування навантаження, для цих цілей було обрано обладнання від двох провідних виробників — Juniper та Extreme Networks. Їхні пристрої орієнтовані на роботу в мережі інтернет-провайдера, де критично важлива безперебійність і стабільна маршрутизація великих обсягів трафіку [18].

Комутатори серії Juniper EX4300 вирізняються широкими функціональними можливостями, зокрема підтримкою стекування у режимі Virtual Chassis, що дозволяє об'єднувати декілька фізичних пристроїв в один логічний блок. Це забезпечує централізоване керування та зменшує складність адміністрування (рисунок 2.7) [19].

Завдяки операційній системі Junos OS забезпечується високий рівень гнучкості при налаштуванні протоколів маршрутизації, політик якості обслуговування (QoS), безпеки доступу та системи моніторингу.



Рисунок 2.7 — Зовнішній вигляд комутатора L3 Juniper EX4300

Крім того, EX4300 підтримує такі технології, як ECMP, що дозволяє здійснювати балансування трафіку між кількома маршрутами з однаковою метрикою, тим самим знижуючи навантаження на окремі канали і запобігаючи їх перевантаженню. Завдяки таким характеристикам, комутатори Juniper EX4300 забезпечують високу стабільність і масштабованість, що робить їх відмінним вибором для вузлів агрегації в мережах інтернет-провайдера.

У свою чергу, обладнання Extreme Networks, зокрема серія X460-G2, також має потужні функції для забезпечення надійної агрегації трафіку (рисунок 2.8).

Підтримка стекування, протоколів EAPS і MLAG дозволяє уникнути виникнення єдиної точки відмови, а також забезпечити швидке переключення у разі збоїв.

Це обладнання широко застосовується в транспортних мережах операторів зв'язку, що підкреслює його відповідність критичним навантаженням і високим вимогам до відмовостійкості [20].

Проте в даному випадку, враховуючи специфіку майбутньої мережі, зроблено вибір на користь Juniper EX4300, оскільки цей вибір є більш обґрунтованим, з огляду на додаткові можливості по інтеграції з іншими компонентами та підтримку більш широких масштабів для роботи з великою кількістю вузлів агрегації (таблиця 2.2).



Рисунок 2.8 — Зовнішній вигляд комутатора L3 Extreme X460-G2

На рівні вузлів доступу, які знаходяться в пелюстках квітки та є кінцевими точками приєднання абонентського обладнання та внутрішніх офісних пристроїв, доцільним є використання рішень, орієнтованих на високу керованість,

енергоефективність та простоту в обслуговуванні. Комутатори Cisco серії C1000 забезпечують базову підтримку керованих мереж на рівні 2, підтримують VLAN, STP, SNMP і Power over Ethernet, що особливо важливо для живлення IP-телефонів або точок доступу (рисунок 2.9) [21].

Таблиця 2.2 — Порівняльна таблиця комутаторів для вузлів агрегації

Параметр	Juniper EX4300	Extreme Networks X460-G2
Основне призначення	Агреація трафіку, маршрутизація, високі навантаження	Агреація трафіку, високі вимоги до відмовостійкості
Підтримка стекування	Virtual Chassis (об'єднання кількох пристроїв в один блок)	EAPS, MLAG (запобігання єдиній точці відмови)
Протоколи маршрутизації	OSPF, BGP	OSPF, BGP
Балансування навантаження	ECMP (Equal-Cost Multi-Path)	Підтримка балансування навантаження через MLAG
Операційна система	Junos OS	ExtremeXOS
Відмовостійкість	Висока (завдяки стекуванню та підтримці протоколів)	Висока (EAPS, MLAG)
Основні технології	Стекування, QoS, моніторинг	EAPS, MLAG, резервування, швидке переключення
Сфера застосування	Оператори зв'язку, провайдери мереж	Оператори зв'язку, транспортні мережі



Рисунок 2.9 — Зовнішній вигляд комутатора L2 Cisco серії C1000

Вони ідеально підходять для офісних відділів і технічних служб з невеликим навантаженням. Додатково у вузлах доступу можуть бути застосовані комутатори Edge-Core ECS4100, які забезпечують стабільну передачу даних, мають розширене керування через SNMP, LLDP і вебінтерфейс, а також підтримують базові засоби захисту (рисунок 2.10) [22].



Рисунок 2.10 — Зовнішній вигляд комутатора L2 Edge-Core ECS4100

Їхня архітектура ідеально відповідає вимогам до надійності та економічної ефективності при обслуговуванні середніх сегментів мережі.

Ще одним ефективним рішенням є комутатори Zyxel GS1920-24HPv2, які мають розширені можливості керування, включаючи підтримку VLAN, розподіл трафіку, пріоритетизацію пакетів і PoE на всіх портах (рисунок 2.11) [23].



Рисунок 2.11 — Зовнішній вигляд комутатора L2 Zyxel GS1920-24HPv2

Також можлива інтеграція з хмарною системою управління Nebula у майбутньому, вона забезпечить можливість централізованого моніторингу та швидкого виявлення перевантажень або аномальної активності без необхідності фізичного втручання.

Така гнучкість критично важлива в умовах розгалуженої інфраструктури, де вузли доступу розташовані на різних територіально віддалених ділянках.

Характеристики розглянутого обладнання наведено (таблиця 2.3).

Остаточний вибір між різними постачальниками обладнання для вузлів агрегації та доступу є ключовим для стабільності і ефективності майбутньої мережі.

Таблиця 2.3 — Порівняльна таблиця комутаторів для вузлів доступу

Параметр	Cisco C1000	Edge-Core ECS4100	Zyxel GS1920-24HPv2
Основне призначення	Вузли доступу для офісів, підключення користувачів	Вузли доступу для середніх сегментів мережі	Вузли доступу для середніх і великих мереж
Підтримка PoE	Так	Так	Так
Керування мережами	VLAN, STP, SNMP	SNMP, LLDP, веб-інтерфейс	VLAN, SNMP, Nebula Cloud Management
Підтримка VLAN	Так	Так	Так
Пріоритетизація трафіку	Обмежена (підтримка базових функцій)	Основні функції захисту і керування	Розширене керування пріоритетами трафіку
Моніторинг та керування	SNMP, базова підтримка керування	SNMP, веб-інтерфейс	Централізоване керування через Nebula
Енергоефективність	Підтримка PoE для підключення пристроїв	Підтримка PoE	Підтримка PoE на всіх портах
Призначення для трафіку	Відділи з невеликим навантаженням	Середні сегменти мережі	Середні та великі мережі

Комутатор Juniper EX4300 є найкращим варіантом для вузлів агрегації, оскільки для них важливі висока продуктивність, масштабованість і гнучкість, що в свою чергу забезпечує стабільну роботи мережі на рівні провайдера.

Для вузлів доступу у свою чергу найкращим варіантом є комутатори Zyxel GS1920-24HPv2 завдяки їхнім можливостям централізованого моніторингу, простоті в налаштуванні і підтримці PoE, що ідеально підходить для розгалужених і економічних інфраструктур.

Окрім комутаторів, важливим елементом інфраструктури є серверний маршрутизатор, який виконує роль центрального вузла маршрутизації, а також точку стику з магістральною мережею провайдера чи іншими сегментами. Цей пристрій має забезпечувати високу продуктивність, підтримку сучасних протоколів, можливості шифрування, фільтрації трафіку, резервування та інтеграції з системами керування. Для таких задач є два відповідних варіанти: MikroTik CCR1036-8G-2S+ та Cisco ISR 4331.

Маршрутизатор MikroTik CCR1036-8G-2S+ побудований на

високопродуктивному 36-ядерному процесорі Tiler Tile-Gx і спеціалізований на багатопотоковій маршрутизації (рисунок 2.12).



Рисунок 2.12 — Зовнішній вигляд маршрутизатора MikroTik CCR1036-8G-2S+

Він оснащений вісьмома гігабітними Ethernet-портами та двома SFP+ для 10GbE підключення. Працює під керуванням RouterOS, яка забезпечує широкі функціональні можливості — підтримку BGP, OSPF, MPLS, тунелів VPN, IPSec, L2TP, GRE, VLAN, а також систему керування трафіком і моніторингу у реальному часі. Пристрій оптимально підходить для навантажених магістральних або прикордонних сегментів мережі і вирізняється порівняно низькою вартістю за високу продуктивність [24].

Cisco ISR 4331 — це маршрутизатор корпоративного класу, побудований на модульній архітектурі, яка дозволяє додавати сервіси, порти без заміни пристрою (рисунок 2.13).



Рисунок 2.13 — Зовнішній вигляд Cisco ISR 4331

Основна перевага цього рішення — розширена функціональність рівня Enterprise. Підтримується інтеграція з Cisco IOS XE — сучасною ОС, що дозволяє керувати конфігураціями, політиками безпеки, здійснювати моніторинг через Cisco Prime, SNMP, NetFlow. У стандартній конфігурації пристрій має два порти Gigabit

Ethernet, але може бути розширений додатковими модулями WAN/LAN, включаючи DSL, 4G, 10GbE та сервіси безпеки. Завдяки підтримці функцій Cisco Zone-Based Firewall, IPS/IDS, advanced QoS, DPI, SSL VPN і можливості інтеграції з Cisco Umbrella, маршрутизатор ідеально підходить для корпоративних мереж, що мають суворі вимоги до безпеки, фільтрації контенту та SLA [25].

Незважаючи на те, що MikroTik CCR1036-8G-2S+ демонструє хорошу продуктивність і має привабливе співвідношення ціни та функцій, він поступається в аспектах розширюваності, безпеки корпоративного рівня, що є критичними у роботі з потенцією масштабування.

Маршрутизатор Cisco ISR 4331 забезпечує розширене керування трафіком, підтримку сучасних протоколів маршрутизації ECMP, BGP, OSPF, MPLS, вбудовані засоби безпеки: IPS, DPI, Zone-Based Firewall, а також можливість централізованого моніторингу через Cisco Prime та NetFlow.

Його модульна архітектура дозволяє додавати нові інтерфейси та сервіси без заміни пристрою, що особливо важливо для довгострокової роботи (таблиця 2.3).

Таблиця 2.3 — Порівняльна таблиця серверних маршрутизаторів

Параметр	MikroTik CCR1036-8G-2S+	Cisco ISR 4331
Процесор	36-ядерний Tiler Tile-Gx	4-ядерний x86
Продуктивність	До 24 млн пакетів/сек	До 100 Мбіт/с (до 300 Мбіт/с із ліцензіями)
Інтерфейси	8 x 1GbE, 2 x 10GbE SFP+	2 x 1GbE, розширення до 10GbE через модулі
Операційна система	RouterOS	Cisco IOS XE
Підтримка BGP/OSPF/MPLS	Так	Так
Безпека	Firewall, NAT, IPSec, DoS-захист	Zone-Based Firewall, IPS/IDS, DPI, SSL VPN
Моніторинг	SNMP, NetFlow, Web GUI, CLI	SNMP, NetFlow, CLI, Cisco Prime
Підтримка VPN	L2TP, PPTP, GRE, IPSec	SSL VPN, IPSec, FlexVPN
Масштабування	Немає, фіксована конфігурація	Так, через ліцензії та модулі
Інтеграція з хмарою	Обмежена (через API або зовнішній моніторинг)	Підтримка Cisco Umbrella, Cisco DNA Center
Призначення	ISP/провайдерські мережі, агрегація, прикордонний маршрутизатор	Корпоративні мережі, філії, центри обробки даних

З урахуванням особливостей проєктованої інфраструктури, ключовими критеріями вибору серверного маршрутизатора стали гнучкість у налаштуванні

протоколів маршрутизації, високий рівень безпеки, інтеграція з системами моніторингу та можливість масштабування в майбутньому, що найкраще реалізує Cisco ISR 4331. Тому Cisco ISR 4331 обраний як серверний маршрутизатор, який найкраще відповідає вимогам стабільності, надійності й адаптивності майбутньої мережі.

Узгоджена робота обладнання агрегаційного рівня та рівня доступу, реалізація резервування, застосування протоколів ECMP, BGP та VLAN, а також інтеграція з системами моніторингу Zabbix та NetFlow, дозволяють ефективно обслуговувати трафік, виявляти та реагувати на перевантаження, а також забезпечити високу стабільність функціонування мережі навіть в умовах максимального навантаження. Такий підхід повністю відповідає потребам в умовах топології типу «квітка».

2.5 Використання віртуальних локальних мереж (VLAN)

Для логічного розділення трафіку, підвищення рівня безпеки, керованості та масштабованості мережі прийнято рішення щодо впровадження віртуальних локальних мереж (VLAN). Незважаючи на те, що фізична топологія залишається незмінною, VLAN дозволяють організувати логічні сегменти мережі відповідно до функціонального призначення підключених пристроїв.

Враховуючи структуру побудованої корпоративної мережі, впровадження VLAN забезпечує ефективне розділення різних зон мережі, таких як серверна кімната, користувацькі зони, підстанції, квіткові топології агрегації, а також точки підключення до провайдера і маршрутизаторів (рисунк 2.14). Наприклад, серверна кімната, яка використовує підмережу 10.0.1.0/27, призначена для критично важливих серверів, таких як файлові, баз даних, контролери домену чи веб-сервери.

Всі пристрої, підключені до цієї VLAN, ізольовані від решти мережі, що дозволяє ефективніше застосовувати політики доступу та моніторинг трафіку.

Користувацькі зони, представлені підмережами 10.0.1.32/27 та 10.0.1.64/27, розділяються між собою окремими VLAN, що дає змогу створити незалежні політики безпеки, обмежити широкомовний трафік та запровадити контроль доступу до ресурсів.

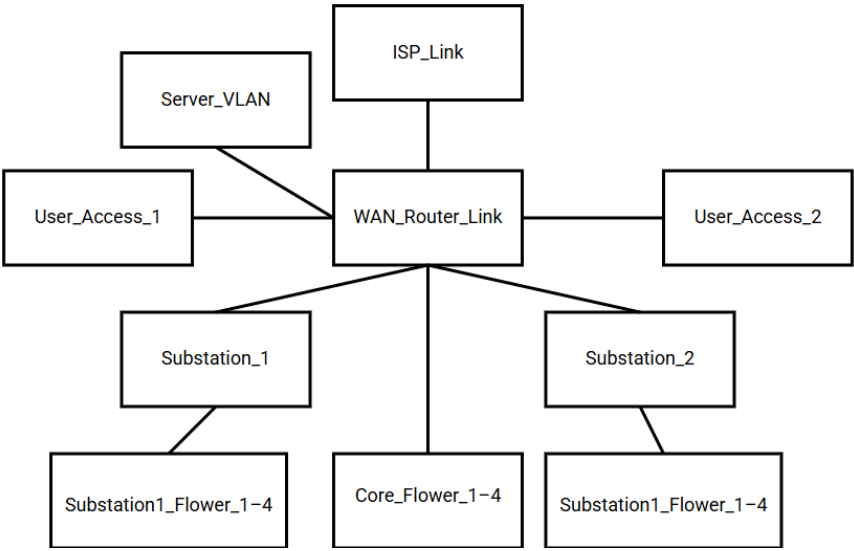


Рисунок 2.14 — Загальна схема логічного поділу мережі на VLAN

Кожен з цих сегментів, що має власну підмережу згідно з адресним планом, отримує окремий ідентифікатор VLAN (таблиця 2.4).

Таблиця 2.4 — Таблиця поділу мережі на VLAN

VLAN ID	Назва сегмента	Підмережа	Призначення
10	Server_VLAN	10.0.1.0/27	Серверна кімната
20	User_Access_1	10.0.1.32/27	Користувацький сегмент 1
21	User_Access_2	10.0.1.64/27	Користувацький сегмент 2
30	Substation_1	10.0.2.0/27	Підстанція 1
31	Substation_2	10.0.2.32/27	Підстанція 2
40–43	Core_Flower_1–4	10.0.0.0/27–10.0.0.96/27	Пелюстки центральної квітки
50–53	Substation1_Flower_1–4	10.0.3.0/27–10.0.3.96/27	Пелюстки квітки підстанції 1
60–63	Substation2_Flower_1–4	10.0.4.0/27–10.0.4.96/27	Пелюстки квітки підстанції 2
100	WAN_Router_Link	192.168.0.0/30	З’єднання між маршрутизаторами

Всі комп’ютери й принтери, що підключені до кожної з цих зон, можуть бути логічно згруповані у свої VLAN відповідно до ролі користувача чи відділу. Це особливо важливо у великих організаціях, де інформаційна безпека відіграє критичну роль.

Підстанції, які мають окремі підмережі — 10.0.2.0/27 та 10.0.2.32/27 — так само ізольовані у VLAN. Це дає можливість незалежно адмініструвати

інфраструктуру кожної з підстанцій, навіть якщо вони фізично підключені до одного мультисвіча. За допомогою VLAN можливо реалізувати обмеження доступу, наприклад, надати право взаємодії з серверами лише певним пристроям із підстанції або ж дозволити доступ до центрального ядра мережі лише адміністративним системам.

Особливу увагу заслуговують квіткові топології агрегації, де кожна пелюстка в межах центрального сегмента, підстанції 1 та підстанції 2 має свою власну підмережу і, відповідно, свою VLAN. Ці VLAN призначені для об'єднання комутаторів нижчого рівня, які обслуговують до 22 пристроїв у кожній пелюстці. Використання VLAN у цьому випадку дозволяє централізовано керувати маршрутизацією, забезпечити безперервність обслуговування при оновленнях, мінімізувати міжпелюстковий трафік, а також підтримувати баланс навантаження між кількома ядрами мережі.

Використання VLAN також стало критично важливим при побудові з'єднань між маршрутизаторами та провайдером. Хоча фізично такі з'єднання реалізовані як point-to-point з використанням маски /30, логічно вони теж відносяться до окремих VLAN. Це дозволяє забезпечити логічну сегментацію між внутрішнім маршрутизатором, граничним маршрутизатором та точкою взаємодії з провайдером, що значно спрощує впровадження політик безпеки, обробку пріоритетного трафіку та підтримку високого рівня ізоляції.

Кожен мультисвіч у мережі підтримує trunk-з'єднання, що передають трафік усіх VLAN до центрального комутатора або маршрутизатора. Такий підхід дозволяє уникнути дублювання фізичних каналів для кожної VLAN, забезпечуючи централізований контроль на рівні маршрутизатора з підтримкою міжвланової маршрутизації. Всі ці VLAN маршрутизуються за допомогою Layer 3 комутаторів або маршрутизаторів, які підтримують логіку міжвланового доступу за суворо визначеними правилами. Наприклад, VLAN користувачів не має прямого доступу до VLAN серверів, якщо це не передбачено політиками, реалізованими через ACL.

Важливо також, що впровадження VLAN створює передумови для гнучкого масштабування мережі. У разі розширення або реорганізації відділів, змін у кількості пристроїв чи їх переміщення, немає потреби у фізичному

перепідключенні кабельної інфраструктури — достатньо перевизначити VLAN на порту комутатора. Такий підхід істотно скорочує витрати на адміністрування та знижує ймовірність помилок при фізичному втручанні.

Отже, впровадження VLAN у даній мережній архітектурі не лише відповідає сучасним стандартам побудови корпоративних мереж, але й забезпечує необхідний рівень безпеки, гнучкості, масштабованості та централізованого управління. У поєднанні з детально продуманим адресним простором та квітковою топологією, VLAN формують основу ефективної, стабільної та легко розширюваної мережної інфраструктури.

2.6 Реалізація механізму балансування трафіку

Забезпечення відмовостійкості в мережі інтернет-провайдера є критично важливим для гарантованого надання безперервних послуг абонентам. У представлений мережній архітектурі застосовано комплексний підхід до підвищення надійності, який охоплює як фізичний, так і логічний рівень.

На фізичному рівні основна станція провайдера об'єднана з двома підстанціями через резервовані канали зв'язку. Усі критично важливі вузли мають дублювання, наприклад: дубльовані комутатори агрегації та резервні маршрутизатори, які автоматично активуються у разі збоїв основного обладнання [26].

Для оптимізації розподілу трафіку в мережі реалізовано підтримку ESMR — технології, що дозволяє використовувати кілька маршрутів з однаковою метрикою для передачі даних.

Це не лише підвищує швидкість доставки інформації, а й забезпечує автоматичне перенаправлення трафіку у разі недоступності одного з маршрутів (рисунок 2.15) [27].

На логічному рівні використано низку мережних протоколів і технологій для забезпечення відмовостійкості та балансування навантаження. Зокрема, реалізовано протокол EtherChannel, який забезпечує об'єднання декількох фізичних з'єднань в один логічний канал для балансування навантаження і підвищення пропускної здатності.

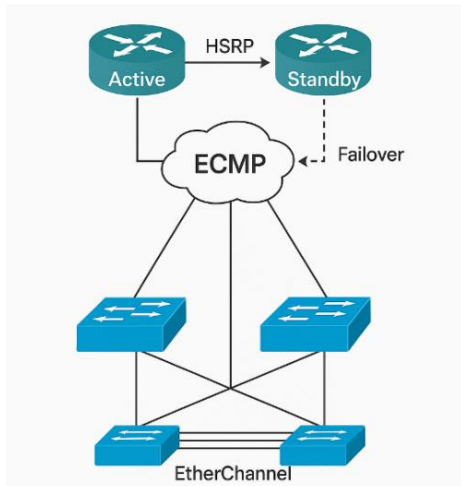


Рисунок 2.15 — Логічна схема використання HSRP, EtherChannel ECMP

У випадках, коли потрібна взаємодія з іншими автономними системами, можлива базова реалізація протоколу BGP. Цей протокол можна використовувати для демонстрації зовнішньої маршрутизації та основ обміну маршрутами з іншими провайдерами.

Така функціональність особливо актуальна для резервного доступу до інтернету через альтернативні канали (рисунок 2.16).

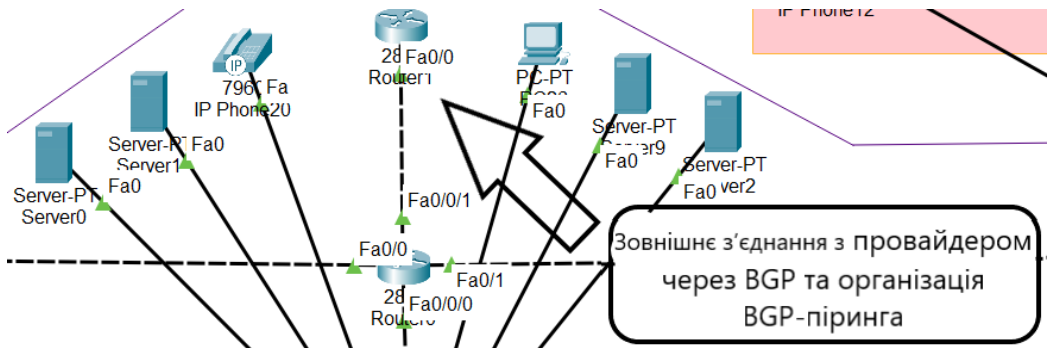


Рисунок 2.16 — Організація з'єднання з провайдером через BGP

Реалізація зазначених технологій у мережі дозволяє провайдеру досягати високого рівня надійності, забезпечуючи безперервне обслуговування клієнтів навіть у випадках часткових збоїв обладнання або перевантаження окремих сегментів мережі.

2.7 Розробка системи моніторингу та керування перевантаженням

Для стабільної роботи мережі інтернет-провайдера критично важливо

впровадження ефективного моніторингу та керування перевантаженням. Основна мета — забезпечення безперебійного функціонування мережних сервісів, своєчасне виявлення перевантажених ділянок та автоматизоване балансування навантаження.

Система моніторингу мережі розробляється з використанням стандартних протоколів, зокрема SNMP, який дозволяє отримувати інформацію про стан пристроїв, завантаження інтерфейсів, кількість активних підключень, помилки та інші метрики (рисунок 2.17).



Рисунок 2.17— Схема роботи SNMP протоколу

Забезпечення ефективного розподілу трафіку та уникнення перевантаження використовуються технології балансування навантаження, такі як ECMP та EtherChannel (рисунок 2.18).

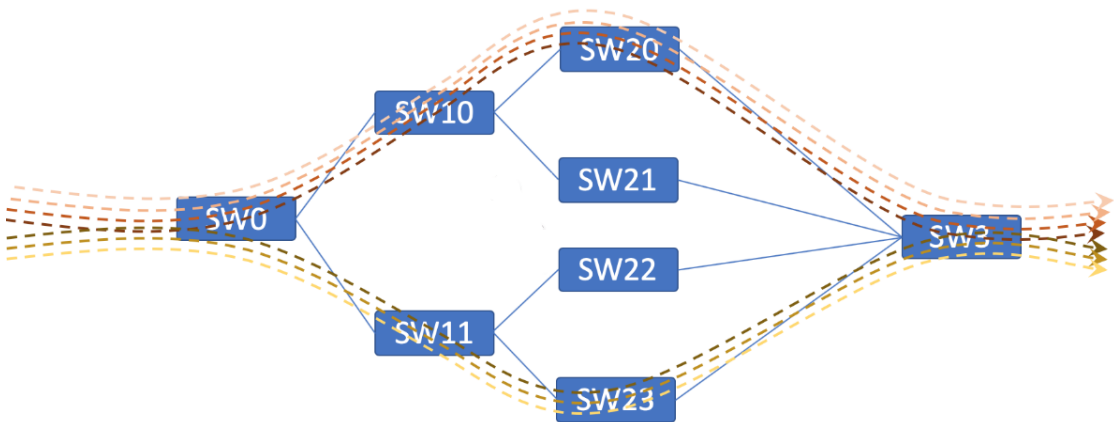


Рисунок 2.18 — Принцип роботи ECMP балансування

ECMP дозволяє рівномірно розподіляти трафік між декількома маршрутами з однаковою вартістю, що значно підвищує продуктивність і відмовостійкість

мережі. У випадку недоступності одного з маршрутів система автоматично переключає трафік на альтернативний шлях без втрати зв'язку. EtherChannel, у свою чергу, дозволяє об'єднувати кілька фізичних каналів у логічний, підвищуючи загальну пропускну здатність і забезпечуючи гнучкість у керуванні потоками даних.

У мережі встановлюються SNMP-менеджери на ключових маршрутизаторах і комутаторах, які передають дані до центрального моніторингового сервера. На основі цих даних можливо відслідковувати трафік в реальному часі, будувати графіки навантаження та налаштовувати системи оповіщення

Для централізованого керування розроблюваною мережею компанії інтернет-провайдера та візуалізації перевантажень використовується SNMP-менеджер Zabbix, який дозволяє у реальному часі збирати та аналізувати дані з мережного обладнання за допомогою протоколу SNMP. Zabbix забезпечує гнучку побудову графіків навантаження (рисунок 2.19), дозволяє налаштовувати тригери на перевищення порогових значень трафіку, втрат пакетів або завантаження інтерфейсів, а також надсилає сповіщення адміністратору про потенційні перевантаження чи відмови у роботі мережі (рисунок 2.20) [28].

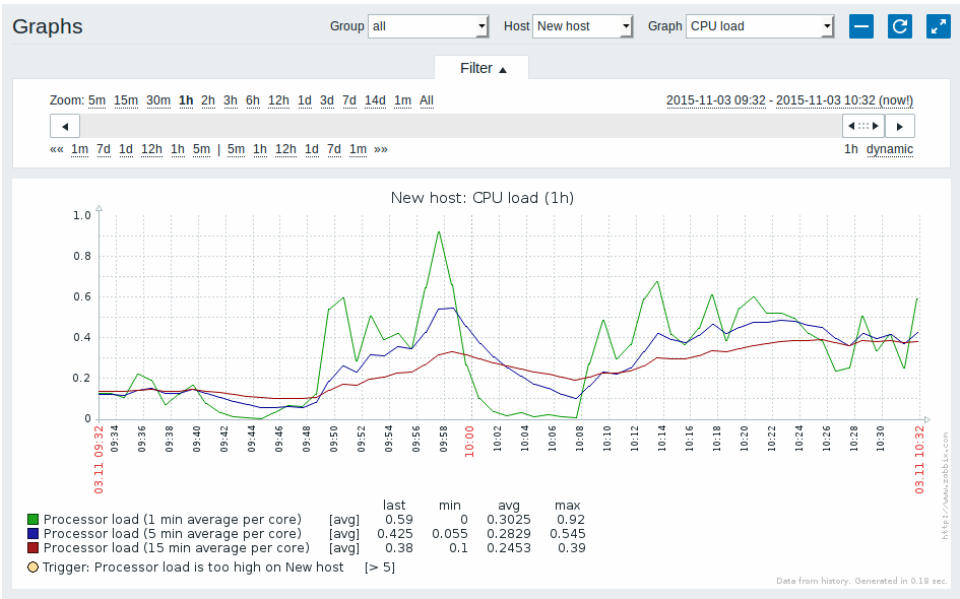


Рисунок 2.19 — Приклад графіку навантаження на обладнання

Крім того, у системі передбачено механізм трафік-шейпінгу та

пріоритетизації даних. Це дозволяє надавати вищий пріоритет трафіку критичних сервісів порівняно з менш важливим трафіком, таким як оновлення програмного забезпечення чи резервне копіювання.



System status						
HOST GROUP	DISASTER	HIGH	AVERAGE	WARNING	INFORMATION	NOT CLASSIFIED
Clouds	0	0	0	0	0	0
Database servers	0	0	0	0	0	0
Discovered hosts	0	0	0	1	1	0
JB applications	0	0	0	0	0	0
Linux servers	0	1	0	0	1	0
Network devices	0	0	0	0	0	0
SNMP hosts	0	0	0	0	0	0
Virtual machines	0	0	0	0	0	0
Web servers	0	0	0	0	0	0
Windows servers	0	0	0	0	0	0
Zabbix servers	0	0	0	1	1	0

Рисунок 2.20— Панель оповіщення Zabbix про зміни в мережі

Реалізація такого моніторингу та механізму балансування дає змогу не тільки виявляти проблеми, а й запобігати їм заздалегідь, оптимізуючи розподіл ресурсів та гарантуючи стабільну якість послуг навіть під час пікового навантаження

3 РОЗРОБКА КОРПОРАТИВНОЇ МЕРЕЖІ КОМПАНІЇ З МОНІТОРИНГОМ ПЕРЕВАНТАЖЕННЯ ТА ТЕСТУВАННЯ

3.1 Налаштування логічної схеми мережі

Для побудови корпоративної мережі компанії було використано симуляційне середовище Cisco Packet Tracer, що дозволяє моделювати реальні сценарії роботи мережної інфраструктури. Розроблена логічна схема включає центральний офіс, три регіональні філії, а також окремі сегменти для різних підрозділів компанії.

Мережа побудована за ієрархічною моделлю з центральним маршрутизатором, до якого підключено всі основні сегменти.

До складу центрального офісу входять такі підрозділи: IT-відділ, служба безпеки, центр дзвінків, відділ підключень, відділ монтажників, відділ обслуговування обладнання, адміністративний офіс, відділ програмістів.

Кожен підрозділ має власний комутатор та окрему IP-підмережу, що дозволяє реалізувати логічну сегментацію за допомогою VLAN.

На рисунку В.2 представлено логічну схему комп'ютерної мережі компанії, яка забезпечує ефективну взаємодію між усіма підрозділами підприємства. Мережна інфраструктура побудована за ієрархічною моделлю, що включає рівень ядра, рівень розподілу, рівень доступу та рівень кінцевих пристроїв. Така модель дозволяє досягти масштабованості, гнучкості, зручності в адмініструванні та високої продуктивності мережі.

Основною метою створення комп'ютерної системи є забезпечення доступу співробітників різних структурних підрозділів до загальних ресурсів, централізованих управлінських та інформаційно-комунікаційних систем, організація ефективної спільної роботи, документообігу, електронної пошти, а також надання доступу до зовнішніх ресурсів та взаємодії з іншими інформаційними системами (державними, фінансовими, партнерськими тощо). Враховано можливість доступу як в онлайн, так і в офлайн режимах.

На рівні ядра мережі функціонує головний маршрутизатор, який забезпечує маршрутизацію між усіма підмережами підприємства та здійснює

підключення до зовнішньої мережі Інтернет. Це єдина точка централізованого керування та контролю міжзонального трафіку.

Рівень розподілу реалізований за допомогою декількох маршрутизаторів, кожен з яких відповідає за певну географічну або логічну зону підприємства. На схемі зображено три основні WAN-зони, які складаються з внутрішніх підмереж, об'єднаних у логічні сегменти, що дозволяє реалізувати гнучку маршрутизацію та ефективну організацію трафіку в межах кожної зони.

Рівень доступу до мережі реалізований за допомогою комутаторів, які з'єднують кінцеві пристрої (персональні комп'ютери, принтери, сервери) з відповідними маршрутизаторами. Для підвищення безпеки та сегментації трафіку використовуються віртуальні локальні мережі (VLAN), що дозволяє ізолювати підрозділи один від одного на рівні доступу.

Кінцеві пристрої включають персональні комп'ютери працівників (з індивідуальними IP-адресами), мережні принтери, файлові сервери, веб-сервери, DNS-сервери. Зокрема, веб-сервер відповідає за обробку HTTP-запитів та забезпечує доступ до корпоративного сайту та внутрішніх веб-сервісів. DNS-сервер реалізує розв'язання імен для внутрішніх та зовнішніх адрес.

Після побудови структурної схеми розподілимо адресний простір. У запропонованій мережній інфраструктурі було застосовано протокол динамічного призначення IP-адрес (DHCP — Dynamic Host Configuration Protocol), що суттєво спрощує процес налаштування та адміністрування мережі, особливо у випадках великої кількості пристроїв.

Використання DHCP дозволяє автоматизувати процес призначення IP-адрес, маски підмережі, шлюзу за замовчуванням, а також DNS-серверів без необхідності ручного втручання адміністратора в кожен пристрій. Це не лише прискорює розгортання мережі, але й зменшує ризик помилок, пов'язаних з дублюванням IP-адрес або неправильним введенням параметрів.

У запропонованій топології було налаштовано два DHCP-сервери на маршрутизаторах, де кожний відноситься до своєї зони. До них передаються запити від клієнтських пристроїв, після чого сервери динамічно надають їм

відповідні мережні параметри згідно з визначеним пулом IP-адрес. Таким чином, приєднання нового пристрою до мережі відбувається автоматично без потреби втручання адміністратора, що особливо актуально в умовах динамічного росту кількості підключень.

На практиці перший DHCP-сервер було налаштовано безпосередньо на маршрутизаторі Router4. Спочатку було зарезервовано певну IP-адресу, яка не повинна призначатися динамічно. Це, як правило, адреса самого маршрутизатора або іншого важливого пристрою. У даному випадку виключено з пулу адресу 10.1.1.62, що, використовується як шлюз за замовчуванням (рисунок 3.1).

```
Router>en
Router#conf t
Enter configuration commands, one per line. End with CNTL/Z.
Router(config)#ip dhcp excluded-address 10.1.1.62
Router(config)#!
Router(config)#ip dhcp pool User_Access_2_R4
Router(dhcp-config)# network 10.1.1.32 255.255.255.224
Router(dhcp-config)# default-router 10.1.1.62
Router(dhcp-config)#
```

Рисунок 3.1 — Налаштування DHCP для Router4 у CLI

Далі було створено DHCP-пул із назвою User_Access_2_R4, який обслуговує підмережу 10.1.1.32/27. Ця підмережа дозволяє розмістити до 30 активних пристроїв. Для коректної роботи клієнтів також було визначено адресу шлюзу — 10.1.1.62, яку DHCP-сервер автоматично надає підключеним пристроям як маршрут за замовчуванням (рисунок 3.2).

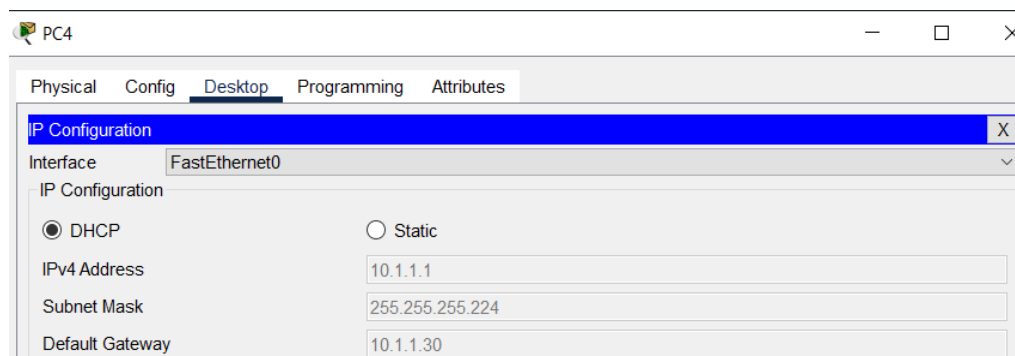


Рисунок 3.2 — Вікно налаштування IP-адресації кінцевого пристрою мережі 10.1.1.32/27

Аналогічні налаштування з другим DHCP-сервером було виконано для Router3 тільки з мережею 10.1.1.0/27 та зарезервованою адресою 10.1.1.30 (рисунок 3.3).

```
ip dhcp excluded-address 10.1.1.30
!
ip dhcp pool User_Access_1_R3
network 10.1.1.0 255.255.255.224
default-router 10.1.1.30
```

Рисунок 3.3 — Налаштування DHCP для Router3 в CLI

Сервер присвоїв пристрою IP-адресу, як показано на рисунку 3.4.

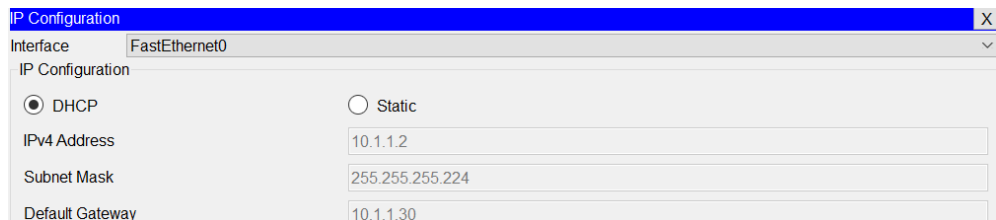


Рисунок 3.4 — Вікно налаштування IP-адресації кінцевого пристрою мережі 10.1.1.0/27

На центральному маршрутизаторі Router2 також організовано динамічний розподіл IP-адрес для пристроїв на підстанціях та в самих квітках (рисунок 3.5).

Для цього додатково було виділено окрему мережу 10.1.2.0/27 в якій для серверів виділено діапазон IP-адрес 10.1.2.2–10.1.2.12 крім 10.1.2.10, яка виділена для ПК в серверній.

```
ip dhcp excluded-address 10.1.2.1
ip dhcp excluded-address 10.1.2.2
ip dhcp excluded-address 10.1.2.3
ip dhcp excluded-address 10.1.2.4
ip dhcp excluded-address 10.1.2.5
ip dhcp excluded-address 10.1.2.6
ip dhcp excluded-address 10.1.2.7
ip dhcp excluded-address 10.1.2.8
ip dhcp excluded-address 10.1.2.9
ip dhcp excluded-address 10.1.2.10
ip dhcp excluded-address 10.1.2.11
ip dhcp excluded-address 10.1.2.12
!
ip dhcp pool User_Access_3_R2
network 10.1.2.0 255.255.255.224
default-router 10.1.2.1
```

Рисунок 3.5 — Налаштування DHCP для Router2 в CLI

Присвоєння сервером IP-адреси для пристрою в мережі 10.1.2.0/27 показано на рисунку 3.6.

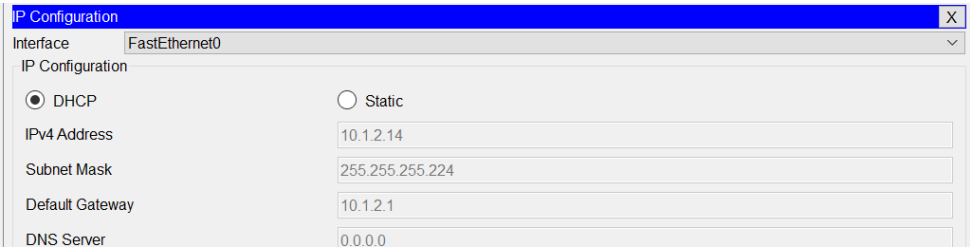


Рисунок 3.6 — Вікно налаштування IP-адресації кінцевого пристрою мережі 10.1.2.0

У межах побудови логічної схеми мережі інтернет-провайдера VLAN відіграють ключову роль у сегментації мережі для ізоляції трафіку, підвищення безпеки та оптимізації використання мережних ресурсів. У даній реалізації кожна група користувачів або окремі підрозділи, зокрема офіс клієнтів, служба підтримки, call-центр або адміністративний офіс, були віднесені до окремих підмереж, яким відповідно призначено VLAN.

Хоча фізично пристрої можуть бути підключені до одного комутатора або маршрутизатора, логічно вони розділені на різні VLAN, кожен з яких має свою IP-підмережу (рисунок 3.7).

VLAN	Name	Status	Ports
1	default	active	Gig0/1, Gig0/2
10	Server_VLAN	active	
20	User_Access_1	active	
21	User_Access_2	active	
30	Substation_1	active	
31	Substation_2	active	
40	Core_Flower_1	active	
41	Core_Flower_2	active	
42	Core_Flower_3	active	
43	Core_Flower_4	active	
50	Substation1_Flower_1	active	Fa0/3, Fa0/4, Fa0/5, Fa0/6 Fa0/7, Fa0/8, Fa0/9, Fa0/10 Fa0/11, Fa0/12, Fa0/13, Fa0/14 Fa0/15, Fa0/16, Fa0/17, Fa0/18 Fa0/19, Fa0/20, Fa0/21, Fa0/22 Fa0/23, Fa0/24
51	Substation1_Flower_2	active	
52	Substation1_Flower_3	active	
53	Substation1_Flower_4	active	
60	Substation2_Flower_1	active	
61	Substation2_Flower_2	active	
62	Substation2_Flower_3	active	
63	Substation2_Flower_4	active	

Рисунок 3.7 — Використання VLAN на інтерфейсах комутатора в петлюстці 3 на центральній квітці

Це дозволяє зменшити широкомовний трафік, обмежити доступ між сегментами через списки доступу та забезпечити масштабованість мережі. Трафік між VLAN курсує через маршрутизатори з підтримкою міжвланової маршрутизації, що також забезпечує контроль над комунікацією між логічно розділеними зонами.

У межах реалізованої мережі провайдера було створено низку VLAN для логічного поділу трафіку між різними сегментами, функціональними блоками та пристроями. Кожен VLAN виконує окрему роль, відповідаючи за специфічні напрямки роботи мережної інфраструктури.

VLAN 10, під назвою `Server_VLAN`, використовується для ізоляції серверної частини мережі, забезпечуючи безпечний доступ до центральних ресурсів, таких як бази даних, файлові сервери або внутрішні системи обліку. Це дозволяє захистити критичні вузли від стороннього доступу.

VLAN 20, названий `User_Access_1`, об'єднує робочі відділи компанії першої групи мережі, зокрема — працівників call-центру, служби безпеки, IT-відділу. Він забезпечує контрольований доступ до сервісів і сегментує трафік цієї групи окремо від решти користувачів.

VLAN 21, або `User_Access_2`, служить для другої групи робочих відділів компанії, наприклад — співробітників служби безпеки, адміністрації, працівників обслуговування обладнання, відділу підключень та відділу програмістів дозволяючи ефективно розділяти навантаження та керувати правами доступу.

VLAN 30, з ім'ям `Substation_1`, об'єднує пристрої першої підстанції, що включає сервери та робочі ПК, розташовані в географічно віддаленій точці. Аналогічну роль виконує VLAN 31 — `Substation_2` — для другої підстанції, розділяючи трафік окремого об'єкта.

VLAN 40, 41, 42 і 43 — `Core_Flower_1`, `Core_Flower_2`, `Core_Flower_3`, `Core_Flower_4` слугують для підключень до основного ядра центральної квітки мережі пелюсток — окремих сегментів в середині квітки.

VLAN 50, 51, 52 та 53 — `Substation1_Flower_1` до `Substation1_Flower_4` — призначені для організації каналів між підстанцією 1 та основним ядром квітки

яка з'єднана з підстанцією 1 мережі. Кожен з цих VLAN може забезпечувати окремий напрямок трафіку або тип даних, що передаються між підстанцією та центром.

Аналогічно, VLAN 60, 61, 62 та 63 — Substation2_Flower_1 до Substation2_Flower_4 реалізовані для підключення другої квітки до підстанції 2.

Останній VLAN — це VLAN 100 з назвою WAN_Router_Link — зарезервований для з'єднання з зовнішніми маршрутизаторами. Через нього проходить весь трафік, що прямує в Інтернет або надходить із зовнішньої мережі, тому він є ключовим для взаємодії з глобальною мережею.

Кожний пристрій знає всі можливі VLAN в мережі для забезпечення безперервної маршрутизації, динамічного призначення портів та підтримки міжвузлової взаємодії через trunk-з'єднання (рисунк 3.8).

VLAN	Name	Status
1	default	active
10	Server_VLAN	active
20	User_Access_1	active
21	User_Access_2	active
30	Substation_1	active
31	Substation_2	active
40	Core_Flower_1	active
41	Core_Flower_2	active
42	Core_Flower_3	active
43	Core_Flower_4	active
50	Substation1_Flower_1	active
51	Substation1_Flower_2	active
52	Substation1_Flower_3	active
53	Substation1_Flower_4	active
60	Substation2_Flower_1	active
61	Substation2_Flower_2	active
62	Substation2_Flower_3	active
63	Substation2_Flower_4	active
100	WAN_Router_Link	active

Рисунок 3.8 — База даних VLAN кожного пристрою в мережі

Також на кожному комутаторі в пелюстці інтерфейси FastEthernet1-2 були переведені в режим trunk для проходження нетегованого трафіку між комутаторами. Інтерфейси FastEthernet3-24 були призначенні відповідні VLAN.

Це дозволяє комутаторам та маршрутизаторам коректно обробляти трафік з різних VLAN, пересилати його між сегментами мережі та реалізовувати

централізоване управління доступом, безпекою й трафіком. Завдяки цьому підтримується гнучкість і масштабованість інфраструктури при збереженні ізольованості логічних зон.

Таким чином, VLAN-архітектура забезпечує не лише структуровану організацію мережі, але й підвищує її керованість, безпеку та ефективність.

При налаштуванні маршрутизації для розроблюваної мережі найдоцільнішим варіантом буде обрати протокол OSPF. По-перше, OSPF є відкритим стандартом, що дозволяє застосовувати його в обладнанні різних виробників без втрати сумісності. По-друге, завдяки використанню алгоритму найкоротшого шляху, OSPF ефективно знаходить оптимальні маршрути, що мінімізує затримки та підвищує продуктивність мережі.

Також OSPF підтримує масштабування мережі за рахунок ієрархічної структури з областями, що важливо для великих корпоративних мереж із численними підмережами.

Крім того, протокол швидко адаптується до змін топології — наприклад, при виході з ладу окремих ліній чи пристроїв маршрутизації — і автоматично переналаштовує маршрути, забезпечуючи безперервність роботи мережі.

Першим етапом налаштування протоколу OSPF є визначення областей маршрутизації, що дозволяє об'єднати маршрутизатори та підмережі у логічні групи для ефективнішого управління ресурсами мережі.

Області спрощують процес обчислення маршрутів і зменшують навантаження на мережний трафік.

Далі потрібно вказати мережі, які маршрутизатори будуть анонсувати в рамках протоколу OSPF.

Це здійснюється шляхом налаштування мережних інтерфейсів маршрутизаторів з відповідними IP-адресами та масками підмереж.

Після цього задається унікальний ідентифікатор маршрутизатора командою:

router-id [х.х.х.х], де [х.х.х.х] — це IP-адреса або інше унікальне значення, що ідентифікує маршрутизатор у мережі.

Потім налаштовують OSPF для кожної підмережі, що підключена до маршрутизатора, командою: `network [адреса підмережі] [зворотня маска] area 0`, де `area 0` — це основна область маршрутизації (backbone area), яка зазвичай використовується за замовчуванням. Ця команда додає підмережу до зони маршрутизації OSPF.

Конфігурування OSPF для Router4 наведено нижче на лістингу 3.1

Лістинг 3.1 — Конфігурування OSPF для Router4

```
router-id 3.3.3.3
log-adjacency-changes
redistribute bgp 65004 subnets
network 10.1.1.96 0.0.0.7 area 0
network 10.1.1.32 0.0.0.31 area 2
distance 10 0.0.0.0 255.255.255.255
```

Аналогічно було виконано налаштування маршрутизатора Router3 для мережі 10.1.1.0 за допомогою команд, які наведені в лістингу 3.2.

Лістинг 3.2 — Конфігурування OSPF для Router3

```
router ospf 1
router-id 1.1.1.1
log-adjacency-changes
redistribute bgp 65003 subnets
network 10.1.1.0 0.0.0.31 area 1
network 10.1.1.64 0.0.0.7 area 0
distance 10 0.0.0.0 255.255.255.255
```

Для налаштування маршрутизації центрального внутрішнього роутера було створено дві проміжні мережі 10.1.1.64 та 10.1.1.96, які були призначенні area 0. Команди для налаштувань наведено в лістингу 3.3.

Лістинг 3.3 — Конфігурування OSPF для Router2

```
router-id 2.2.2.2
log-adjacency-changes
redistribute bgp 65002 subnets
network 10.1.1.96 0.0.0.7 area 0
network 10.1.1.64 0.0.0.7 area 0
network 10.1.2.0 0.0.0.31 area 4
distance 10 0.0.0.0 255.255.255.255
```

Виконаємо налаштування BGP для виходу за межі нашої мережі в зовнішню мережу Інтернет для організації пірингових угод та постачання послуг для користувачів. Для Router1 присвоєно номер 65001. Цей маршрутизатор виступає зовнішнім вузлом, що з'єднується з ядром мережі інтернет-провайдера через інтерфейс з IP-адресою 192.168.12.1 для виходу в Інтернет.

Визначено основного BGP-сусіда з яким встановлюється з'єднання, — Router2, що має номер автономної системи (AS 65002). Команда `router bgp 65001` ініціює BGP. В параметрі `neighbor 192.168.12.2 remote-as 65002` — номер AS 65002, Для моніторингу стану з'єднання застосовано команду `bgp log-neighbor-changes`, що дозволяє фіксувати зміну станів BGP-пірів, що є важливим для стабільності з'єднання.

Відключення синхронізації з внутрішніми маршрутизаторами здійснюється за допомогою `no synchronization`, що дозволяє уникнути необхідності узгодження з IGP-протоколами, тому що на маршрутизаторі налаштовано OSPF та використовується для внутрішньої маршрутизації.

Router1 оголошує дві мережі: 10.1.1.64/29 та 10.1.1.96/29. Відповідно, ці підмережі стають видимими для інших BGP-партнерів через Router1.

Команди які були введені в CLI на лістингу 3.4.

Лістинг 3.4 — Конфігурування BGP для Router1

```
router bgp 65001
bgp log-neighbor-changes
no synchronization
neighbor 192.168.12.2 remote-as 65002
network 10.1.1.64 mask 255.255.255.248
network 10.1.1.96 mask 255.255.255.248
```

На центральному внутрішньому маршрутизаторі мережі провайдера — Router2 було присвоєно номер 65002. Він відповідає за з'єднання як із зовнішніми вузлами, так і з внутрішніми клієнтами та з офісом, що потребують доступу до ресурсів Інтернету. Конфігурація `router bgp 65002` запускає BGP-процес. Параметри `neighbor` задають зв'язки з сусідніми маршрутизаторами: 192.168.12.1 (AS 65001) — це зовнішній клієнтський вузол, тоді як 10.1.1.65 (AS

65003) і 10.1.1.97 (AS 65004) — інші маршрутизатори, що також підключені до ядра.

Router2 оголошує три мережі: 10.1.1.0/27, 10.1.1.64/29 і 10.1.1.96/29. Це дозволяє зробити їх доступними для інших BGP-партнерів. Команда `redistribute ospf 1` також додає до BGP-маршрутів ті маршрути, що динамічно навчаються через внутрішній протокол OSPF, що забезпечує повноцінний обмін інформацією між пристроями всередині мережі.

Налаштування здійснювалося введення таких команд в CLI, які наведені на лістингу 3.5.

Лістинг 3.5 — Конфігурування BGP для Router2

```
router bgp 65002
bgp log-neighbor-changes
no synchronization
neighbor 192.168.12.1 remote-as 65001
neighbor 10.1.1.32 remote-as 65002
neighbor 10.1.1.65 remote-as 65003
neighbor 10.1.1.97 remote-as 65004
network 10.1.1.0 mask 255.255.255.224
network 10.1.1.64 mask 255.255.255.248
network 10.1.1.96 mask 255.255.255.248
redistribute ospf 1
```

Router3, якому присвоюється номер 65003 в мережі за протоколом BGP, здійснює обмін інформацією між мережею компанії інтернет-провайдера та мережею Інтернет, до якої звертатимуться офіс центру дзвінків, IT-відділ та кімната служби безпеки через цей маршрутизатор.

Команда `router bgp 65003` ініціює BGP-сесію для цієї AS. Зв'язок із ядром мережі забезпечується через параметр `neighbor 10.1.1.66 remote-as 65002`, де вказано адресу Router2.

За допомогою `network 10.1.1.0 mask 255.255.255.192` анонсується внутрішня мережа цього вузла, яка буде передаватися через BGP в глобальну мережу Інтернет.

Команди для налаштування Router3 наведені в лістингу 3.6.

Лістинг 3.6 — Конфігурування BGP для Router3

```
router bgp 65003
bgp log-neighbor-changes
no synchronization
neighbor 10.1.1.66 remote-as 65002
network 10.1.1.0 mask 255.255.255.192
redistribute ospf 1
```

Router4, якому присвоєно номер AS 65004, відповідає за забезпечення зв'язку та надає доступу до Інтернету для відділу підключень, відділу обслуговування обладнання, адміністративному офісу, відділу монтажників, відділу програмістів та кімнати служби безпеки.

Далі було введено команди для налаштування BGP на Router4, які наведені на лістингу 3.7.

Лістинг 3.7 — Конфігурування BGP для Router4

```
router bgp 65004
bgp log-neighbor-changes
no synchronization
neighbor 10.1.1.98 remote-as 65002
network 10.1.1.32 mask 255.255.255.224
redistribute ospf 1
```

Для забезпечення контролю доступу до мережних ресурсів у топології провайдера були застосовані списки доступу (ACL), що дозволяють обмежити або дозволити передавання трафіку на основі IP-адрес та протоколів. У наведеній конфігурації використано розширений ACL [29].

У мережі було реалізовано розширені ACL (Access Control Lists), які регулюють доступ до окремих ресурсів та сервісів, відповідно до політики безпеки. Основна мета — обмежити небажаний трафік

Перший ACL під номером 100 був налаштований для повного блокування будь-яких спроб доступу з підмережі 10.1.2.0/27 такими командами:

Лістинг 3.8 — Конфігурування ACL для Router2

```
access-list 100 deny ip 10.1.2.0 0.0.0.31 any
access-list 100 permit ip any any
ip access-group 100 in
```

Це правило означає, що будь-який трафік з підмережі 10.1.2.0, незалежно від протоколу чи порту, буде відкинутий, а весь інший трафік дозволений. Такий підхід часто використовується для ізоляції окремих сегментів мережі, наприклад, гостьових або підозрілих зон.

Другий ACL — під номером 110 — налаштовано для дозволу лише специфічного трафіку: HTTP (порт 80) та DHCP (порти 67 та 68). Цей список забезпечує функціонування базових сервісів, залишаючи при цьому інші типи трафіку заблокованими. Такий метод налаштування ACL чітко регламентує мережну взаємодію відповідно до вимог безпеки та функціональності.

ACL під номером 110 був налаштований такими командами, які наведені на лістингу 3.9.

Лістинг 3.9 — Конфігурування ACL для Router3 та Router4

```
access-list 110 permit tcp any any eq 80
access-list 110 permit udp any eq 67 any eq 68
access-list 110 permit udp any eq 68 any eq 67
access-list 110 deny ip any any
interface FastEthernet0/1
ip access-group 110 in
```

Ці правила дозволяють веб-доступ та динамічне отримання IP-адрес через DHCP, що є критичним для користувачів, але блокує весь інший IP-трафік. Це допомагає обмежити потенційні загрози та небажану активність у мережі.

Таким чином, реалізація ACL дозволяє ефективно сегментувати трафік і захистити внутрішню інфраструктуру без шкоди для необхідної функціональності мережі компанії інтернет-провайдера.

3.2 Налаштування системи моніторингу та перевантаження

Підключення SNMP у мережі здійснюється для забезпечення моніторингу стану маршрутизаторів та збирання діагностичної інформації. На кожному маршрутизаторі було налаштовано базову конфігурацію, яка дозволяє зовнішнім системам моніторингу здійснювати запити до пристроїв у режимі лише для читання. Для цього задається SNMP-ключ доступу з назвою «public» і правами

read-only, що дозволяє отримувати дані без можливості їх змінювати.

Після введення команди `snmp-server community public RO` з'являється повідомлення «SNMP-5-WARMSTART: SNMP agent on host Router is undergoing a warm start», що сигналізує про теплий старт агента і це свідчить про перезапуск SNMP без повного перезавантаження маршрутизатора.

Додатково було вказано логічне розташування пристрою в мережі “Sever_room”.

Команда яка була введена в CLI: `snmp-server location Sever_room`.

Також було добавлено контактну інформацію відповідального адміністратора.

Для цього було введено команду: `snmp-server contact admin.Ivan@example.com`.

Після застосування цих налаштувань SNMP-менеджер, такий як Zabbix, може успішно отримувати інформацію про навантаження інтерфейсів, стан з'єднань, активність маршрутів і загальну працездатність обладнання.

Таким чином, SNMP забезпечить централізований контроль і покращить можливості для оперативного реагування на потенційні проблеми в мережній інфраструктурі.

У мережу було впроваджено ESMR для реалізації ефективного механізму балансування трафіку та забезпечення високої доступності між маршрутизаторами. Особливо важливу роль ця технологія відіграє на Router2 — центральному вузлі, який пов'язує ядро мережі з периферійними маршрутизаторами Router3 та Router4. Обидва ці маршрутизатори підключені до Router2 через підмережі 10.1.1.64/29 та 10.1.1.96/29 відповідно, і обидва надають доступ до однієї й тієї ж мережі клієнтів — 10.1.2.0/27.

ESMR дозволяє Router2 бачити до мережі 10.1.2.0/27 два рівноцінні шляхи: один через Router3, інший — через Router4. Обидва маршрути мають однакову метрику і адміністративну відстань, що дозволяє використовувати їх паралельно для передачі даних. У результаті, трафік рівномірно розподіляється між цими маршрутами, зменшуючи навантаження на кожен окремий канал і підвищуючи

загальну пропускну здатність системи.

Щоб забезпечити функціонування ECMP, на Router2 у BGP було застосовано команду `maximum-paths 2`, яка дозволяє використовувати два маршрути з однаковими параметрами одночасно. OSPF як внутрішній протокол маршрутизації поширює інформацію про мережу 10.1.2.0/27 від обох маршрутизаторів, що забезпечує стабільність у випадку відмови одного з них.

Це підтверджується таблицею маршрутизації Router2 на лістингу 3.10.

Лістинг 3.10 — Таблиця Router2 в CLI для перевірки налаштування ECMP

```
Router2# show ip route 10.1.2.0
Routing entry for 10.1.2.0/27
Known via "bgp 65002", distance 20, metric 0
10.1.1.65, from 10.1.1.32 (Router2)
10.1.1.65, from 10.1.1.65 (Router3)
10.1.1.97, from 10.1.1.97 (Router4)
```

Тут видно, що до однієї й тієї ж мережі ведуть два маршрути — через 10.1.1.65 і 10.1.1.97. Це дозволяє системі як рівномірно розподіляти навантаження, так і миттєво переключатись на доступний маршрут у разі аварії. Таким чином, ECMP забезпечує одночасно масштабованість, резервування, стійкість до відмов та надійного механізму балансування, що потрібно для критично важливої інфраструктури мережі компанії інтернет-провайдера.

Систему моніторингу та перевантаження компанії було розроблено з впровадженням протоколу SNMP та сервіс-менеджера Zabbix. Система моніторингу забезпечить централізований моніторинг мережного обладнання з можливістю оперативного реагування на збої. Конфігурація з використанням «public» у режимі read-only гарантує безпечний доступ до діагностичної інформації. Визначення логічного розташування пристроїв та контактної інформації адміністратора покращує структуру управління.

Завдяки ECMP реалізовано балансування навантаження і підвищено доступність маршрутизаторів. Такий підхід сприяє стабільному та ефективному курсуванню трафіку в мережі.

3.3 Перевірка працездатності розробленої мережі

Після розробки та реалізації логічної структури мережі компанії було проведено комплексне тестування процесів передачі даних. Основною метою цього етапу є перевірка коректної маршрутизації, забезпечення зв'язності всіх підмереж, визначення надійності каналів зв'язку та перевірка відсутності затримок, які можуть впливати на якість обслуговування кінцевих користувачів.

Для проведення тестування передачі даних у корпоративній мережі було заплановано перевірити такі основні параметри:

- пропускну здатність мережі;
- затримки при передачі пакетів;
- рівень втрат пакетів;
- стабільність з'єднання при тривалій роботі.

У процесі тестування використовувались як стандартні інструменти командного рядка, так і спеціалізоване програмне забезпечення:

- ICMP (Ping) — для перевірки доступності мережних вузлів;
- Cisco Packet Tracer (у разі емуляційного середовища) — для перевірки конфігурацій у віртуальному середовищі.

Розглянемо таблиці маршрутизації на кожному з маршрутизаторів нашої мережі, щоб підтвердити коректність маршрутизації, роботу динамічних протоколів (OSPF, BGP) та забезпечення зв'язності між підмережами.

На Router1 видно, що маршрутизатор має динамічні маршрути до трьох підмереж — 10.1.1.64/29, 10.1.1.96/29 та 10.1.2.0/27. Усі вони отримані через BGP від сусіднього Router2 по каналу з адресою 192.168.12.2. Крім того, локально підключена мережа 192.168.12.0/24, зокрема інтерфейс 192.168.12.1 (рисунк 3.9).

```

10.0.0.0/8 is variably subnetted, 3 subnets, 2 masks
B    10.1.1.64/29 [20/0] via 192.168.12.2, 00:00:00
B    10.1.1.96/29 [20/0] via 192.168.12.2, 00:00:00
B    10.1.2.0/27 [20/20] via 192.168.12.2, 00:00:00
192.168.12.0/24 is variably subnetted, 2 subnets, 2 masks
C    192.168.12.0/24 is directly connected, FastEthernet0/1
L    192.168.12.1/32 is directly connected, FastEthernet0/1

```

Рисунок 3.9 — Вигляд таблиці маршрутизації Router 1 у CLI

Router2 виступає центральним вузлом і має безпосереднє підключення до підмереж 10.1.1.0/27, 10.1.1.30/32, 10.1.1.64/29, а також маршрут до вузлів Router3 і Router4: 10.1.1.32/27, 10.1.1.96/29, 10.1.2.0/27 — через інтерфейс 10.1.1.66. Це свідчить про коректну роботу ESRP та динамічної маршрутизації OSPF. Підмережа 192.168.12.0/27 маршрутизується до зовнішньої частини мережі (рисунок 3.10).

```

10.0.0.0/8 is variably subnetted, 7 subnets, 3 masks
C    10.1.1.0/27 is directly connected, FastEthernet0/0
L    10.1.1.30/32 is directly connected, FastEthernet0/0
O IA 10.1.1.32/27 [10/3] via 10.1.1.66, 00:02:28, FastEthernet0/1
C    10.1.1.64/29 is directly connected, FastEthernet0/1
L    10.1.1.65/32 is directly connected, FastEthernet0/1
O    10.1.1.96/29 [10/2] via 10.1.1.66, 00:02:48, FastEthernet0/1
O IA 10.1.2.0/27 [10/2] via 10.1.1.66, 00:02:38, FastEthernet0/1
    192.168.12.0/27 is subnetted, 1 subnets
S    192.168.12.0/27 [1/0] via 10.1.1.66

```

Рисунок 3.10 — Вигляд таблиці маршрутизації Router 2 у CLI

На Router3 налаштовані маршрути до 10.1.1.0/27 та 10.1.1.64/29 — вони безпосередньо підключені через інтерфейси FastEthernet0/0 та FastEthernet0/1 відповідно. Доступ до підмережі 192.168.12.0/27 реалізовано через маршрутизатор за адресою 10.1.1.66 — тобто, через Router2 (рисунок 3.11).

```

10.0.0.0/8 is variably subnetted, 4 subnets, 3 masks
C    10.1.1.0/27 is directly connected, FastEthernet0/0
L    10.1.1.30/32 is directly connected, FastEthernet0/0
C    10.1.1.64/29 is directly connected, FastEthernet0/1
L    10.1.1.65/32 is directly connected, FastEthernet0/1
    192.168.12.0/27 is subnetted, 1 subnets
S    192.168.12.0/27 [1/0] via 10.1.1.66

```

Рисунок 3.11 — Вигляд таблиці маршрутизації Router3 у CLI

Router4 має аналогічну ситуацію — безпосередній доступ до 10.1.1.32/27 та 10.1.1.96/29, а зв'язок із зовнішньою мережею (192.168.12.0/27) здійснюється через 10.1.1.98. Це свідчить про те, що весь зовнішній трафік проходить через централізовану точку доступу — Router2, як і передбачалося в логічній схемі (рисунок 3.12).

```

10.0.0.0/8 is variably subnetted, 4 subnets, 3 masks
C    10.1.1.32/27 is directly connected, FastEthernet0/0
L    10.1.1.62/32 is directly connected, FastEthernet0/0
C    10.1.1.96/29 is directly connected, FastEthernet0/1
L    10.1.1.97/32 is directly connected, FastEthernet0/1
192.168.12.0/27 is subnetted, 1 subnets
S    192.168.12.0/27 [1/0] via 10.1.1.98

```

Рисунок 3.12 — Вигляд таблиці маршрутизації Router 4 у CLI

Таким чином, таблиці маршрутизації на всіх маршрутизаторах підтверджують повну зв'язність між усіма сегментами мережі. Протоколи BGP та OSPF успішно виконують свої функції, а маршрути встановлені як для локальних, так і для віддалених підмереж. Це забезпечує стабільний обмін даними між усіма частинами корпоративної мережі.

Далі розглянемо кілька сценаріїв для перевірки ефективності роботи мережі. Перший сценарій включає внутрішню комунікацію. В рамках однієї підмережі перевіримо здатність між різними ПК обмінюватися даними без затримок та втрат пакетів (рисунок 3.13).

PDU List Window										
Fire	Last Status	Source	Destination	Type	Color	Time(sec)	Periodic	Num	Edit	Delete
	Successful	PC4	PC9	ICMP		0.000	N	0	(edit)	
	Successful	PC3	PC8	ICMP		0.000	N	1	(edit)	

Рисунок 3.13 — Перевірка сигналу між відділами однієї підмережі

Результат підтверджує коректну роботу комутатора, ARP, а також правильність налаштування IP-адрес.

Другий сценарій включає міжмережний зв'язок через маршрутизатори. Тестування взаємодії пристроїв з різних VLAN, що потребувало маршрутизації між підмережами. Перевірялася коректність маршрутизованого доступу та функціональність протоколу OSPF (рисунок 3.14).

PDU List Window										
Fire	Last Status	Source	Destination	Type	Color	Time(sec)	Periodic	Num	Edit	Delete
	Successful	PC2	PC18	ICMP		0.000	N	0	(edit)	
	Successful	PC7	PC15	ICMP		0.000	N	1	(edit)	

Рисунок 3.14 — Перевірка обміну пакетами між відділами різних підмереж

Розглянемо третій сценарій — зовнішнє з'єднання. Відправлення запитів на зовнішні ресурси (емуляція Інтернету), правил доступу (ACL), та зворотної маршрутизації (рисунок 3.15).

Fire	Last Status	Source	Destination	Type	Color	Time(sec)	Periodic	Num	Edit	Delete
	Successful	PC27	Router1	ICMP		0.000	N	0	(edit)	

Рисунок 3.15 — Відправка сигналу у зовнішню мережу Інтернет

Перевірка BGP протоколу для зв'язку із зовнішньою мережею Інтернет пройшла успішно і сигнал доходить до маршрутизатора.

Четвертий сценарій передбачає тестування з'єднання між офісом та квіткою (рисунок 3.16).

Fire	Last Status	Source	Destination	Type	Color	Time(sec)	Periodic	Num
	Failed	PC0	PC4	ICMP		0.000	N	0
	Failed	PC28	PC4	ICMP		0.000	N	1
	Failed	PC29	PC4	ICMP		0.000	N	2

Рисунок 3.16 — Заблоковані запити ICMP пакетів з квітки на офіс

Дана перевірка підтвердила, що сигнал не проходить між офісом та кінцевим пристроєм у квітці. Це пояснюється тим, що тут застосовуються ACL для перешкоджання несанкціонованому доступу з середини мережі

Далі протестуємо з'єднання між серверними підстанцій та серверною головною станцією (рисунок 3.17).

Fire	Last Status	Source	Destination	Type	Color	Time(sec)	Periodic	Num
	Successful	Server3	Server7	ICMP		0.000	N	0
	Successful	Server5	Server0	ICMP		0.000	N	1

Рисунок 3.17 — Успішний обмін пакетами між серверними

Як бачимо сигнал проходить та все працює коректно, що означає виключення втрати даних клієнтів.

Шостий сценарій передбачає перевірку сигналу для коректної взаємодії між підстанціями (рисунок 3.18).

Fire	Last Status	Source	Destination	Type	Color	Time(sec)	Periodic	Num
	Successful	PC20	PC21	ICMP		0.000	N	0
	Successful	PC22	PC25	ICMP		0.000	N	1

Рисунок 3.18 — Успішне проходження сигналу між підстанціями

Сигнал проходить між підстанціями, що характеризує про чітку та надійну комунікацію між ними.

Щоб краще оцінити, наскільки добре мережа захищена від небажаного трафіку, проведено тестування списків контролю доступу (рисунок 3.19). Це також необхідно для того щоб кінцеві користувачі в мережі не могли отримати доступ до сервісів керування мережею та запобігти несанкціонованому доступу для запобігання витоку інформації компанії.

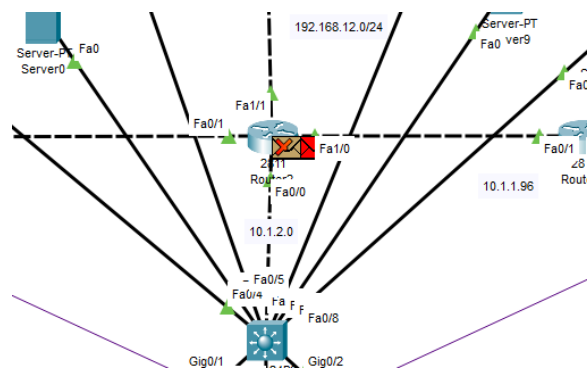


Рисунок 3.19 — Перевірка роботи ACL

ACL з номером 100 був налаштований з метою перешкоджання надходженню трафіку з мережі 10.1.2.0 і в при отриманні пакетів з цієї мережі він блокує їх. Перевірка показала що він працює коректно оскільки вони не пропускають жодного пакету до офісу компанії, що перешкоджає недоброчесним діям зі стороною кінцевих користувачів мережі.

В результаті проведеного тестування було підтверджено повну функціональність розробленої логічної структури мережі. Всі маршрутизатори успішно здійснюють обмін маршрутною інформацією за допомогою протоколів BGP та OSPF, а таблиці маршрутизації містять коректні маршрути до всіх підмереж.

Верифікація комунікації в межах VLAN, між VLAN, а також між офісом, підстанціями та зовнішнім середовищем підтвердила доступність усіх критичних сегментів мережі. Виявлені обмеження трафіку через ACL також відповідають вимогам безпеки.

Таким чином, тестування показало стабільну роботу мережі з мінімальними затримками, відсутністю втрат пакетів і високим рівнем доступності всіх сервісів, безпеки з'єднання та запобігання несанкціонованому впливу із середини мережі компанії інтернет-провайдера.

3.4 Тестування системи перевантаження

Одним із ключових завдань було тестування системи моніторингу перевантаження мережі яка є важливою для оперативного виявлення неполадок і підтримки стабільної роботи. Для цього використано протокол SNMP, який надає можливість не лише віддалено збирати статистику та інформацію про роботу мережних пристроїв, але й робить це з високим рівнем безпеки завдяки підтримці автентифікації та шифрування.

Було здійснено перевірку працездатності механізму моніторингу мережних пристроїв за допомогою SNMP у програмі Zabbix. Для цього з адміністративного комп'ютера введемо команду `snmpwalk -v2c -c public 10.1.1.1` що робить опитування маршрутизаторів за допомогою SNMP-менеджера, що

сповістить про спробу несанкціонованого доступу до одного з маршрутизаторів. В результаті SNMP-менеджер отримав повідомлення про помилку автентифікації, що підтвердило активацію механізму сповіщення (рисунок 3.20).

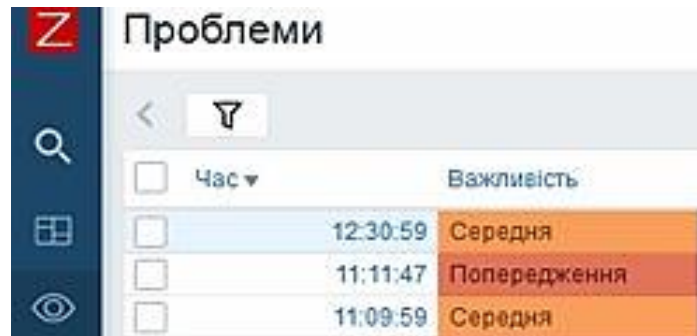


Рисунок 3.20 — Вікно оповіщення у Zabbix про несанкціонований доступ до мережі адміністратору мережі

Для моніторингу доступності мережних пристроїв та оцінки їхнього стану в реальному часі використовувався Zabbix. Одним із основних методів перевірки було виконання ICMP-пінгу (ping) до контрольованих вузлів (рисунок 3.21).

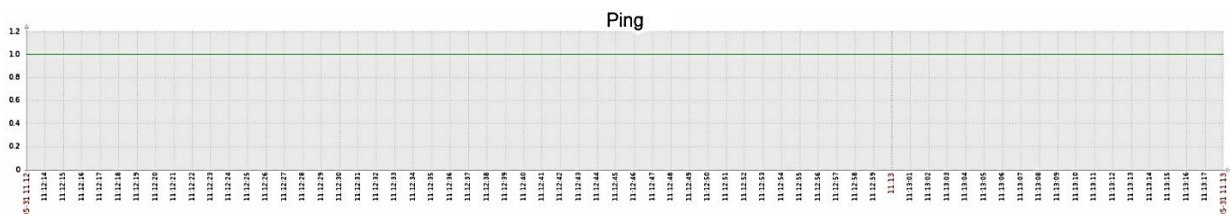


Рисунок 3.21 — Перевірка сигналу у Zabbix

Як бачимо з графіку сигнал стабільний і не виникає ніяких збоїв у мережі.

Для перевірки реального розподілу трафіку провели паралельні сеанси передачі даних, а також виконали команду `debug ip packet` в CLI. Після чого спостерігалось розподілення трафіку між двома маршрутами, що підтверджує коректне функціонування ESMR. Це видно з виводу команд в CLI на лістингу 3.12.

Лістинг 3.12 — Результат виконання команди `debug ip packet` що підтверджує застосування ESMR

IP: tableid=0, s=10.1.1.97 (FastEthernet1/0), d=10.1.1.98 (FastEthernet1/0),
routed via RIB
 IP: s=10.1.1.97 (FastEthernet1/0), d=10.1.1.98 (FastEthernet1/0), len 96, rcvd
 3
 IP: s=10.1.1.98 (local), d=224.0.0.5 (FastEthernet1/0), len 20, sending
 broad/multicast
 IP: tableid=0, s=10.1.1.97 (FastEthernet1/0), d=10.1.1.98 (FastEthernet1/0),
routed via RIB
 IP: s=10.1.1.97 (FastEthernet1/0), d=10.1.1.98 (FastEthernet1/0), len 96, rcvd
 3
 IP: s=10.1.2.1 (local), d=224.0.0.5 (FastEthernet0/0), len 20, sending
 broad/multicast
 IP: s=10.1.1.66 (local), d=224.0.0.5 (FastEthernet0/1), len 20, sending
 broad/multicast
 IP: s=10.1.1.65 (FastEthernet0/1), d=224.0.0.5 len 20, rcvd 2
 IP: tableid=0, s=10.1.1.66 (FastEthernet0/1), d=10.1.1.65 (FastEthernet0/1),
routed via RIB
 IP: tableid=0, s=10.1.1.65 (FastEthernet0/1), d=10.1.1.66 (FastEthernet0/1),
routed via RIB

Видно, що трафік автоматично перенаправляється за альтернативним маршрутом, демонструючи надійність механізму балансування трафіку.

Протестуємо працездатність мережі, який обсяг даних вона може пропускати та чи не виникає збоїв при навантаженні мережі, тобто при проходженні через великого обсягу трафіку.

Для цього надсилаємо великий потік пакетів з одної квітки на іншу та проводимо виміри трафіку, і спостерігаємо за змінами в мережі (рисунок 3.22).

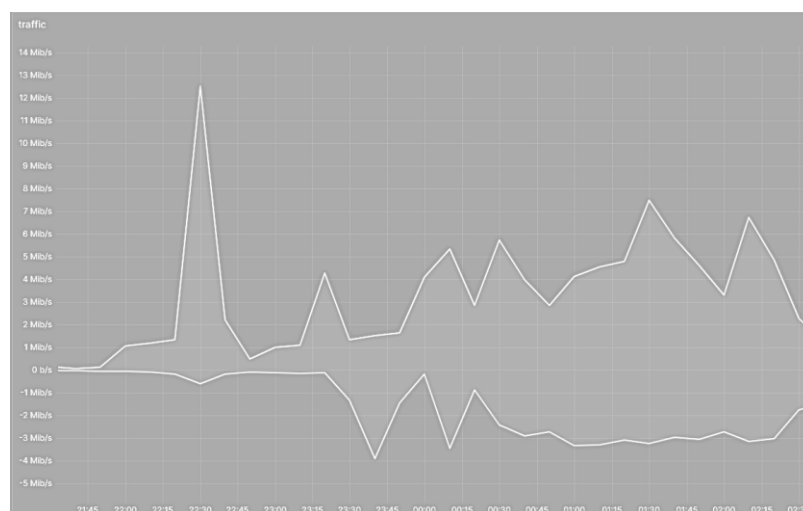


Рисунок 3.22 — Графік трафіку який проходить через сервер

З графіку видно, що при проходженні трафіку через сервер спостерігається незначне навантаження на мережу, що характеризує систему моніторингу як надійну систему, яка забезпечує швидке реагування на проблеми.

Проведене тестування системи перевантаження показало, що механізми балансування навантаження працюють стабільно та ефективно, забезпечуючи розподіл трафіку між кількома маршрутами навіть при відмові окремих інтерфейсів.

Моніторинг через SNMP підтвердив своє коректне функціонування та оперативне надсилання сповіщень про проблеми. Перевірка ACL показала, що списки контролю доступу надійно блокують небажаний трафік. Також моніторинг пінгу через Zabbix дозволяє оперативно виявляти проблеми з доступністю мережних пристроїв, що підвищує загальну надійність мережі.

В цілому, система добре справляється з перевантаженнями та забезпечує стабільну роботу мережі.

ВИСНОВКИ

У бакалаврській кваліфікаційній роботі розроблено комп'ютерну мережу компанії інтернет-провайдера з урахуванням сучасних технологій побудови, моніторингу та управління інфраструктурою. Розроблена мережа орієнтована на забезпечення високої продуктивності, масштабованості, доступності, а також безпеки передачі даних та сервісів.

У першому розділі проведено аналіз сучасних технологій побудови корпоративних мереж, проаналізовано топології, протоколи передачі даних, методи моніторингу та балансування навантаження. Також було проаналізовано топологію типу «квітка» та протоколи SNMP та ECMP, що стало основою для подальшої розробки мережі інтернет-провайдера.

У другому розділі було здійснено розробку самої структури комп'ютерної мережі компанії інтернет-провайдера: визначено загальну архітектуру, обрано відповідне мережне обладнання, побудовано логічну схему, розраховано IP-адресний простір. Також розроблено сегментацію за допомогою VLAN, реалізовано механізм балансування трафіку з протоколом ECMP та систему моніторингу перевантаження на основі SNMP і Zabbix. Розроблені рішення дозволили створити ефективну, керовану та стійку до перевантажень мережу, що відповідає потребам сучасної інфраструктури інтернет-провайдера.

В третьому розділі проведено тестування працездатності та пропускну здатності мережі, яке показало наявність зв'язку між усіма компонентами мережі, що підтвердило правильність налаштування протоколів OSPF та BGP, а також розподілення адресного простору та сегментацію мережі по VLAN. Система моніторингу перевантаження разом із механізмом балансування трафіку у ході тестування показали високу продуктивність та надійність. Списки контролю доступу продемонстрували свою надійність та забезпечили блокування шкідливого трафіку.

Топологія «квітка», яка забезпечує багатоканальне резервування між вузлами, зменшує кількість критичних точок відмови й сприяє підвищенню надійності. Такий тип побудови дозволяє масштабувати мережу без порушення її цілісності та зниження продуктивності.

В розробленій комп'ютерній мережі виконана інтеграція механізму балансування трафіку, що базується на технології ESMR. Цей підхід забезпечив рівномірний розподіл навантаження між кількома маршрутами з однаковою метрикою, дозволив зменшити затримки, уникнути перевантажень і оптимізувати використання доступних мережних ресурсів. Завдяки цьому досягнуто підвищення швидкодії та стійкості мережі до пікових навантажень.

З метою цілодобового моніторингу та управління станом мережної інфраструктури організовано та запущено SNMP-менеджер Zabbix, який за допомогою SNMP забезпечує збір діагностичної інформації в реальному часі. Така система моніторингу мережі дозволяє оперативно реагувати на перевантаження та відмови обладнання, що є критично важливим для стабільної роботи мережі провайдера.

Результатом виконання бакалаврської кваліфікаційної роботи стала реалізована модель комп'ютерної мережі компанії інтернет-провайдера, яка поєднує сучасні технології побудови, балансування трафіку та моніторингу перевантаження мережі [30]. Розроблена інфраструктура відповідає вимогам високої надійності, масштабованості та безпеки, а проведені тестування підтвердили її надійність.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Використання протоколів маршрутизації в мережі інтернет-провайдера. Молодь в науці: дослідження, проблеми, перспективи (МН-2025) м. Вінниця, 18.05.2025. ВНТУ, 2025. Режим доступу: <https://conferences.vntu.edu.ua/index.php/mn/mn2025/paper/viewFile/25229/20873>
2. Network Load Balancer: Як працює балансувальник навантаження? [Електронний ресурс] — Режим доступу: <https://blog.colobridge.net/uk/2024/01/network-load-balancer-ua/>
3. What Types of Load Balancing Are Available: A Comprehensive Guide [Електронний ресурс] — Режим доступу: <https://proxyelite.info/uk/what-types-of-load-balancing-are-available-a-comprehensive-guide/>
4. Cisco Digital Network Architecture [Електронний ресурс] — Режим доступу: <https://www.telesphera.net/blog/cisco-digital-network-architecture.html>
5. Aruba Central: Огляд рішень для управління мережею [Електронний ресурс] — Режим доступу: <https://www.elko.ua/novosti2/aruba-central-review>
6. Кременецкий Блог — Проблеми побудови мереж [Електронний ресурс] — Режим доступу: <https://kremenetskyu.blogspot.com/2017/10/blog-post.html>
7. Топологія комп'ютерних мереж [Електронний ресурс] — Режим доступу: https://stud.com.ua/53329/informatika/topologiya_kompyuternih_merezh
8. Протоколи та моделі комп'ютерних мереж [Електронний ресурс] — Режим доступу: <https://studfile.net/preview/5083544/page:4/>
9. Типи мережних протоколів і їх призначення [Електронний ресурс] — Режим доступу: https://deltahost.ua/ua/tipi-merezhevix-protokoliv-i-ih-priznachennya-http-ip-ssh-ftp-pop3mac.html?srsId=AfmBOooOE7ApLzzic3ILC36ryI27E5Ro6Qi32isB_6A0e9A_k93a-bkh
10. Протоколи передачі даних: огляд мережних протоколів [Електронний ресурс] — Режим доступу: <https://klaster.ua/ua/stati-i->

obzory/protokoly-peredachi-dannykh-obzor-setevykh-protokolov/?srsId=AfmBOopWD1ZS310fvSHsHTWZFsUqXE7n1aO5vYkGeGfJC FdG1bra4uOv

11. Протоколи передачі даних: їх типи та особливості [Електронний ресурс] — Режим доступу: <https://highload.tech/uk/protokoli-peredavannya-danih-yih-tipi-ta-osoblivosti/>

12. Кордяк В., Дронюк І., Федевич О. Інформаційна технологія моніторингу та аналізу трафіку у комп'ютерних мережах. — Національний університет «Львівська політехніка», 2015. [Електронний ресурс] — Режим доступу: <http://ena.lp.edu.ua:8080/bitstream/ntb/31297/1/07-35-42.pdf>

13. Взаємозв'язок між протоколами SNMP версії 1, версії 2 та версії 3 стандартної комп'ютерної мережі [Електронний ресурс] — Режим доступу: <https://datatracker.ietf.org/doc/rfc3584/>

14. Zhang Y., Wang H. Congestion-Aware Routing in Software Defined Networks // IEEE Access. 2021. Vol. 9. P. 4501–4515. DOI: 10.1109/ACCESS.2021.3050000

15. Биков В. В., Шишкіна М. П. Актуальні проблеми оптимізації комп'ютерних мереж: теорія та практика. — Київ: КНУ, 2020. — 224 с.

16. Tanenbaum A. S., Wetherall D. J. Computer Networks (5th Edition). — 2010. — 960 p.

17. Архітектура мережі або мережна архітектура [Електронний ресурс] — Режим доступу: <https://nettech.ua/news/arxitektura-seti-ili-setevaja-arxitektura>

18. Основні мережні пристрої: принцип дії та відмінності [Електронний ресурс] — Режим доступу: <https://comtrade.ua/ua/blog/osnovnye-setevye-ustroystva-printsip-deystviya-i-otlichiya/>

19. Тестування комутатора Juniper [Електронний ресурс] — Режим доступу: <https://itbiz.ua/statti-ta-obzori/testuvannya-komutatora-juniper-qfx5200-32c-i-optichnix-moduliv-100-gbps-itbqs28-dp-x20/>

20. Комутатор Extreme X460-G2 [Електронний ресурс] — Режим доступу: <https://omnilink.ua/komutator-extreme-x460-g2/>

21. Cisco Catalyst 1000 Series Switches Data Sheet [Електронний ресурс]— Режим доступу: <https://www.cisco.com/c/en/us/products/collateral/switches/catalyst-1000-series-switches/nb-06-cat1k-ser-switch-ds-cte-en.html>
22. Комутатор Edge-Core ECS4100-28T [Електронний ресурс] — Режим доступу: <https://ua.setevoe.com.ua/aktivka/switches/edge-core-ecs4100-28t.html>
23. Комутатор керований L2 ZyXEL GS1920-24HP V2 [Електронний ресурс] — Режим доступу: <https://e-server.com.ua/uk/aktivne-obladnannja/komutatori/komutatori-shho-nalashtovujutsja/komutator-kerovanij-l2-zyxel-gs1920-24hp-v2-gs192024hpv2-eu0101f-detail?srsId=AfmBOooG4NlTVXiRJRJni-1tVjWROcYYgS8xbXr2ais4iZRweIH2Yshq>
24. MikroTik CCR1036-8G-2S+ [Електронний ресурс] — Режим доступу: <https://mikrotik.com/product/CCR1036-8G-2Splus>
25. Cisco 4000 Family Integrated Services Router Data Sheet [Електронний ресурс] — Режим доступу: https://www.cisco.com/c/en/us/products/collateral/routers/4000-series-integrated-services-routers-isr/data_sheet-c78-732542.html
26. Network Warrior by Gary A. Donahue 2007 — 595р. [Електронний ресурс] — Режим доступу: <https://theswissbay.ch/pdf/Gentoomen%20Library/Networking/O%27Reilly%20Network%20Warrior.pdf>
27. What Network Load Balancing Is And How It Works[Електронний ресурс] — Режим доступу: <https://blog.colobridge.net/en/2024/02/network-load-balancer-en/>
28. Моніторинг інфраструктури та сервісів [Електронний ресурс] — Режим доступу: <https://pronet.ua/monitoring-infrastrukturi-i-servisiv/>
29. Access Control List (ACL) [Електронний ресурс] — Режим доступу: <https://www.wallarm.com/what/whats-the-access-control-list-acl>
30. Методичні вказівки до виконання бакалаврських кваліфікаційних робіт для студентів спеціальності 123 «Комп’ютерна інженерія» (освітня програма «Системне програмування») [Електронний ресурс] / уклад. О. Д. Азаров, С. В. Богомолів, С. І. Швець. – Вінниця : ВНТУ, 2023. – 104 с.

ДОДАТОК А

Технічне завдання

Міністерство освіти і науки України

Вінницький національний технічний університет

Факультет інформаційних технологій та комп'ютерної інженерії

Кафедра обчислювальної техніки

ЗАТВЕРДЖУЮ

Завідувач кафедри ОТ

проф., д.т.н.. Азаров О.Д.

« 21 » березня 2025 року

ТЕХНІЧНЕ ЗАВДАННЯ

на виконання бакалаврської кваліфікаційної роботи

«Комп'ютерна мережа компанії з моніторингом перевантаження»

08-54.БКР.010.00.000 ТЗ

Науковий керівник: к.т.н., доц.

Войцеховська О.В.

Студент групи ІСП-216

Олійник І.І.

1 Підставою для виконання бакалаврської кваліфікаційної роботи

1.1 Важливим є актуальність дослідження у сфері побудови надійної корпоративної комп'ютерної мережі з механізмами моніторингу та балансування навантаження. Сучасні компанії, зокрема інтернет-провайдери, потребують стабільної інфраструктури, здатної ефективно адаптуватися до зростання обсягів трафіку та підвищених вимог до безпеки, що зумовлює необхідність впровадження сучасних протоколів, топологічних рішень та систем керування. Підставою для розробки є наказ про затвердження теми БКР.

1.2 Наказ про затвердження теми БКР.

2 Мета БКР і призначення розробки

2.1 Мета роботи — проєктування та налаштування надійної, масштабованої і відмовостійкої комп'ютерної мережі інтернет-провайдера, яка забезпечує ефективний моніторинг перевантаження та балансування трафіку.

2.2 Призначення розробки — корпоративна комп'ютерна мережа компанії інтернет-провайдера призначена для забезпечення надійної, стабільної та безпечної передачі даних з можливістю виявлення перевантажень мережних вузлів, балансування трафіку та оперативного реагування на критичні навантаження за допомогою інтегрованої системи моніторингу.

3 Вихідні дані для виконання БКР

3.1 Проведення аналізу сучасних принципів побудови корпоративних комп'ютерних мереж, методів балансування трафіку та моніторингу перевантаження.

3.2 Розробка загальної архітектури мережі з урахуванням масштабованості, безпеки та надійності, побудованої за обраною топологією.

3.3 Вибір мережного обладнання, що підтримує необхідні протоколи для моніторингу та для балансування трафіку.

3.4 Побудова логічної та фізичної структур мережі з впровадженням безпеки та механізму розподілу навантаження.

3.5 Організація моніторингу перевантаження.

3.6 Врахування вимог нормативних документів, стандартів і технічних регламентів щодо проєктування, експлуатації та захисту мережної інфраструктури.

4 Вимоги до виконання БКР

Основною вимогою до бакалаврської кваліфікаційної роботи є створення надійної, масштабованої та захищеної мережі інтернет-провайдера з активним моніторингом перевантажень. Потрібно обґрунтувати архітектуру, реалізувати логічну адресацію, інтегрувати механізм балансування трафіку, впровадити сучасні засоби моніторингу до технічних вимог і стандартів оформлення.

5 Етапи БКР та очікувані результати

Етапи роботи та очікувані результати приведено в Таблиці А.1.

Таблиця А.1 — Етапи БКР

№	Назва етапу	Термін виконання		Очікувані результати
		початок	кінець	
1	Аналіз літературних джерел	12.03.25	24.03.25	1 розділ
2	Огляд технологій моніторингу та проблем перевантаження мережі	24.03.25	29.03.25	1 розділ
3	Розробка архітектури мережі	29.03.25	07.04.25	2 розділ
4	Вибір обладнання, створення схем адресації та VLAN	07.04.25	15.04.25	2 розділ
5	Організація моніторингу та механізму балансування навантаження	15.04.25	21.04.25	2 розділ
6	Розробка та налаштування схеми та моніторингу перевантаження	21.04.25	28.04.25	3 розділ
8	Тестування мережі та моніторингу перевантаження	28.04.25	05.05.25	3 розділ
9	Підготовка матеріалів та опис комп'ютерної мережі	05.05.25	09.05.25	Тези БКР
10	Аналіз роботи, висновки, додатки	09.05.25	11.05.25	Готова ПЗ
11	Перевірка якості виконання бакалаврської роботи та усунення недоліків	12.05.25	13.05.25	Готова робота та презентація

6 Матеріали, що подаються до захисту БКП

Перед захистом подаються такі матеріали як: пояснювальна записка до БКР, графічна та ілюстративна частини, готовий протокол попереднього захисту БКР на кафедрі, відгук наукового керівника, рецензія від рецензента, протоколи складання державних екзаменів, анотації до БКР українською та іноземною мовами.

7 Контроль виконання та захисту БКР

Виконання графічної та розрахункової документації БКР відбувається під контролем наукового керівника відповідно до визначених строків. Захист БКР проходить на засіданні Екзаменаційної комісії, яка затверджується наказом ректора.

8 Вимоги до оформлювання та порядок виконання БКР

8.1 При оформлювання БКП використовуються:

— ДСТУ 3008: 2015 «Звіти в сфері науки і техніки. Структура та правила оформлювання»;

— ДСТУ 8302: 2015 «Бібліографічні посилання. Загальні положення та правила складання»;

— методичні вказівки до виконання бакалаврських кваліфікаційних робіт зі спеціальності 123 «Комп'ютерна інженерія» (освітня програма «Системне програмування»). Кафедра обчислювальної техніки ВНТУ 2022;

— документами на які посилаються у вище вказаних.

8.2 Порядок виконання БКР викладено в «Положення про кваліфікаційні роботи на першому (бакалаврському) рівні вищої освіти СУЯ ВНТУ-03.02.02-П.001.01:21».

Керівник роботи _____ Войцеховська О. В.
(підпис) (прізвище, ініціали)

ДОДАТОК В

Ілюстративна частина

**КОМП'ЮТЕРНА МЕРЕЖА КОМПАНІЇ З МОНІТОРИНГОМ
ПЕРЕВАНТАЖЕННЯ**

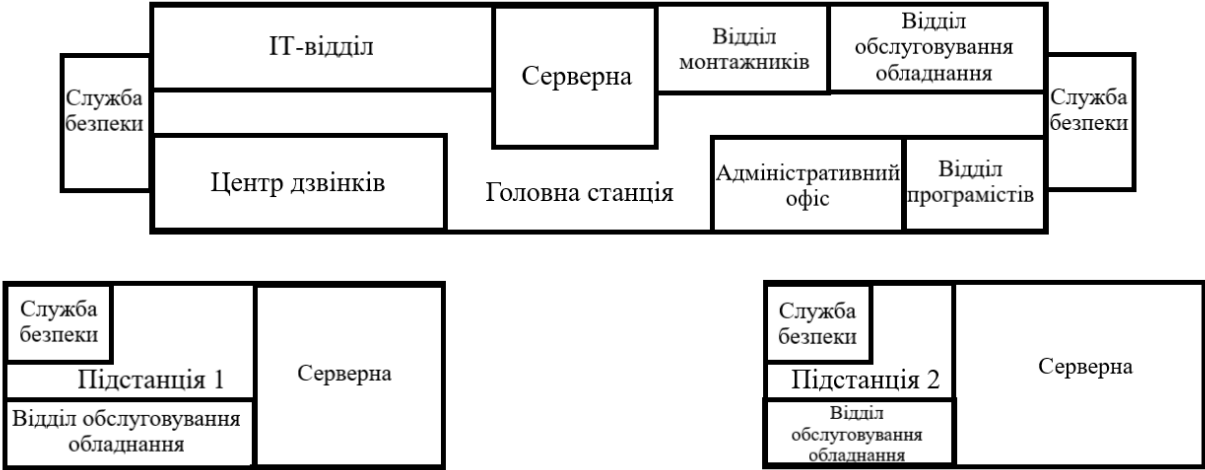


Рисунок В.1 — Фізична схема компанії

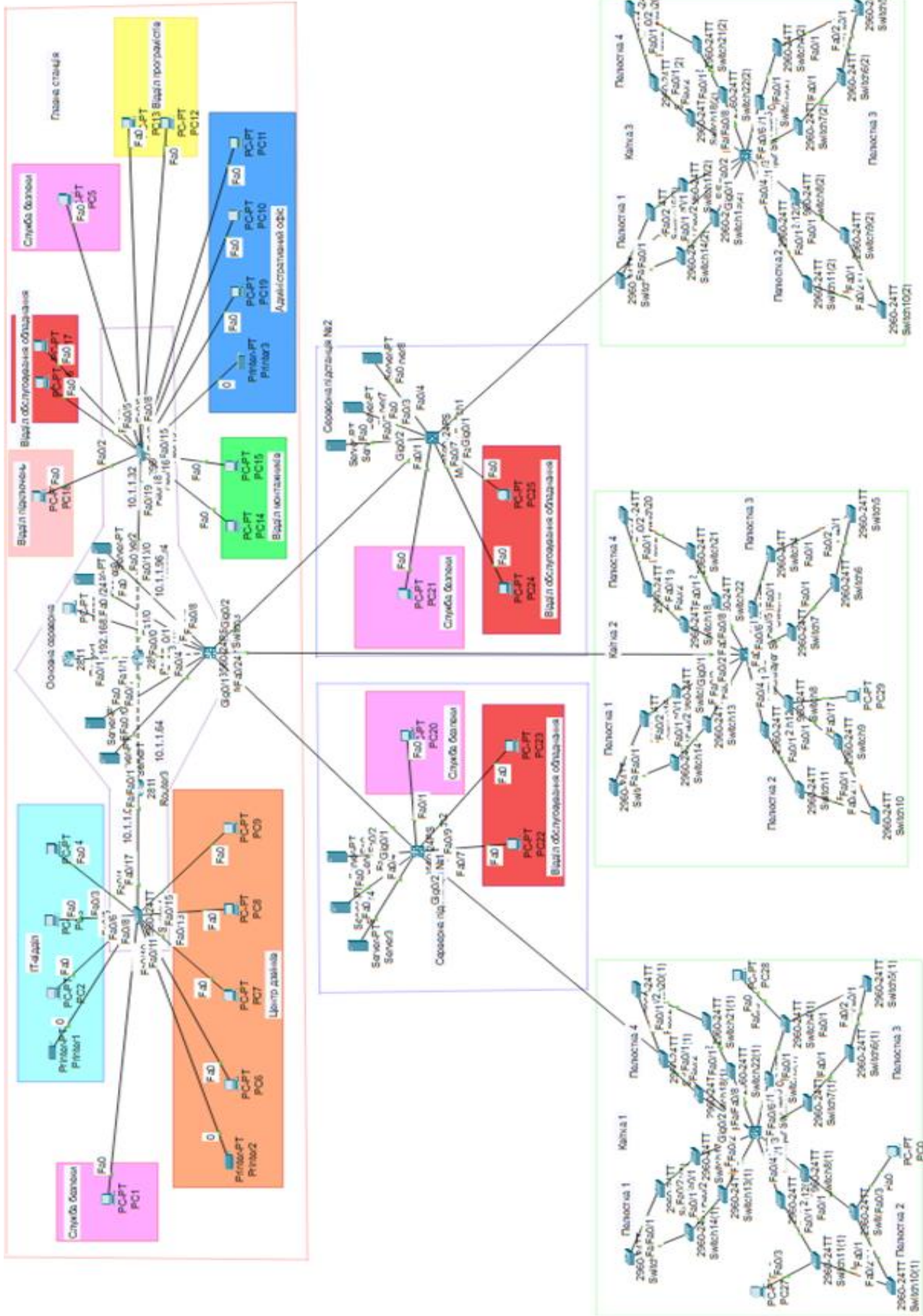


Рисунок В.2 — Логічна схема мережі компанії

ДОДАТОК Г

Лістинг конфігураційних файлів

Лістинг Г.1 — Лістинг конфігураційного файлу Router1

```
version 15.1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
hostname Router
ip cef
no ipv6 cef
license udi pid CISCO2811/K9 sn FTX1017M0HK-
spanning-tree mode pvst
interface FastEthernet0/0
no ip address
duplex auto
speed auto
interface FastEthernet0/1
ip address 192.168.12.1 255.255.255.0
duplex auto
speed auto
interface Vlan1
no ip address
shutdown
router bgp 65001
bgp log-neighbor-changes
no synchronization
neighbor 192.168.12.2 remote-as 65002
network 10.1.1.64 mask 255.255.255.248
network 10.1.1.96 mask 255.255.255.248
ip classless
ip flow-export version 9
line con 0
line aux 0
line vty 0 4
login
end
```

Лістинг Г.2 — Лістинг конфігураційного файлу Router2

```
version 15.1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
hostname Router
ip dhcp excluded-address 10.1.2.1
ip dhcp excluded-address 10.1.2.2
ip dhcp excluded-address 10.1.2.3
ip dhcp excluded-address 10.1.2.4
ip dhcp excluded-address 10.1.2.5
```



```
ip dhcp excluded-address 10.1.2.6
ip dhcp excluded-address 10.1.2.7
ip dhcp excluded-address 10.1.2.8
ip dhcp excluded-address 10.1.2.9
ip dhcp excluded-address 10.1.2.10
ip dhcp excluded-address 10.1.2.11
ip dhcp excluded-address 10.1.2.12
ip dhcp pool User_Access_3_R2
  network 10.1.2.0 255.255.255.224
  default-router 10.1.2.1
no ip cef
no ipv6 cef
license udi pid CISCO2811/K9 sn FTX1017W051-
spanning-tree mode pvst
interface FastEthernet0/0
  ip address 10.1.2.1 255.255.255.224
  ip access-group 100 in
  duplex auto
  speed auto
interface FastEthernet0/1
  ip address 10.1.1.66 255.255.255.248
  duplex auto
  speed auto
interface FastEthernet1/0
  ip address 10.1.1.98 255.255.255.248
  duplex auto
  speed auto
interface FastEthernet1/1
  ip address 192.168.12.2 255.255.255.0
  duplex auto
  speed auto
interface Vlan1
  no ip address
  shutdown
router ospf 1
  router-id 2.2.2.2
  log-adjacency-changes
  redistribute bgp 65002 subnets
  network 10.1.1.96 0.0.0.7 area 0
  network 10.1.1.64 0.0.0.7 area 0
  network 10.1.2.0 0.0.0.31 area 4
  distance 10 0.0.0.0 255.255.255.255
router bgp 65002
  bgp log-neighbor-changes
  no synchronization
  neighbor 192.168.12.1 remote-as 65001
  neighbor 10.1.1.65 remote-as 65003
  neighbor 10.1.1.97 remote-as 65004
  network 10.1.1.0 mask 255.255.255.224
  network 10.1.1.64 mask 255.255.255.248
  network 10.1.1.96 mask 255.255.255.248
```

```

redistribute ospf 1
router rip
ip classless
ip flow-export version 9
access-list 110 permit tcp any any eq www
access-list 110 permit udp any eq bootps any eq bootpc
access-list 110 permit udp any eq bootpc any eq bootps
access-list 110 deny ip any any
snmp-server community monitor RO
line con 0
line aux 0
line vty 0 4
  login
end

```

Лістинг Г.3 — Лістинг конфігураційного файлу Router3

```

version 15.1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
ip dhcp excluded-address 10.1.1.30
ip dhcp pool User_Access_1_R3
  network 10.1.1.0 255.255.255.224
  default-router 10.1.1.30
no ip cef
no ipv6 cef
license udi pid CISCO2811/K9 sn FTX1017F383-
spanning-tree mode pvst
interface FastEthernet0/0
  ip address 10.1.1.30 255.255.255.224
  ip access-group 102 in
interface FastEthernet0/1
  ip address 10.1.1.65 255.255.255.248
  ip access-group 110 in
interface Vlan1
  no ip address
  shutdown
router ospf 1
  router-id 1.1.1.1
  log-adjacency-changes
  redistribute bgp 65003 subnets
  network 10.1.1.0 0.0.0.31 area 1
  network 10.1.1.64 0.0.0.7 area 0
  distance 10 0.0.0.0 255.255.255.255
router bgp 65003
  bgp log-neighbor-changes
  no synchronization
  neighbor 10.1.1.66 remote-as 65002
  network 10.1.1.0 mask 255.255.255.192
  redistribute ospf 1
  ip classless

```

```

ip route 192.168.12.0 255.255.255.224 10.1.1.66
ip flow-export version 9
access-list 100 deny ip 10.1.2.0 0.0.0.31 any
access-list 100 permit ip any any
access-list 110 permit tcp any any eq www
access-list 110 permit udp any eq bootps any eq bootpc
access-list 110 permit udp any eq bootpc any eq bootps
access-list 110 deny ip any any
line con 0
line aux 0
line vty 0 4
login
end

```

Лістинг Г.4 — Лістинг конфігураційного файлу Router4

```

version 15.1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
ip dhcp excluded-address 10.1.1.62
ip dhcp pool User_Access_2_R4
network 10.1.1.32 255.255.255.224
default-router 10.1.1.62
no ip cef
no ipv6 cef
license udi pid CISCO2811/K9 sn FTX1017MAUW-
spanning-tree mode pvst
interface FastEthernet0/0
ip address 10.1.1.62 255.255.255.224
ip access-group 102 in
interface FastEthernet0/1
ip address 10.1.1.97 255.255.255.248
ip access-group 110 in
interface Vlan1
no ip address
shutdown
router ospf 1
router-id 3.3.3.3
log-adjacency-changes
redistribute bgp 65004 subnets
network 10.1.1.96 0.0.0.7 area 0
network 10.1.1.32 0.0.0.31 area 2
distance 10 0.0.0.0 255.255.255.255
router bgp 65004
bgp log-neighbor-changes
no synchronization
neighbor 10.1.1.98 remote-as 65002
network 10.1.1.32 mask 255.255.255.224
redistribute ospf 1
ip classless
ip route 192.168.12.0 255.255.255.224 10.1.1.98

```

```

ip flow-export version 9
access-list 100 deny ip 10.1.2.0 0.0.0.31 any
access-list 100 permit ip any any
access-list 110 permit tcp any any eq www
access-list 110 permit udp any eq bootps any eq bootpc
access-list 110 permit udp any eq bootpc any eq bootps
access-list 110 deny ip any any
line con 0
line aux 0
line vty 0 4
  login
end

```

Лістинг Г.5 — Лістинг конфігураційного файлу центрального комутатора

```

version 12.2(37)SE1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
ip dhcp excluded-address 10.1.2.1
ip dhcp excluded-address 10.1.2.2
ip dhcp excluded-address 10.1.2.3
ip dhcp excluded-address 10.1.2.4
ip dhcp excluded-address 10.1.2.5
ip dhcp excluded-address 10.1.2.6
ip dhcp excluded-address 10.1.2.7
ip dhcp excluded-address 10.1.2.8
ip dhcp excluded-address 10.1.2.9
ip dhcp excluded-address 10.1.2.10
ip dhcp excluded-address 10.1.2.11
ip dhcp excluded-address 10.1.2.12
ip dhcp pool User_Access_3_R2
  network 10.1.2.0 255.255.255.224
  default-router 10.1.2.1
ip routing
spanning-tree mode pvst
interface FastEthernet0/1-2
  switchport trunk encapsulation dot1q
  switchport mode trunk
interface FastEthernet0/3-5, 7-8
  switchport access vlan 10
no switchport
ip address 10.1.2.2 255.255.255.224
duplex auto
speed auto
switchport access vlan 10
interface FastEthernet0/9-23
interface FastEthernet0/24
  switchport trunk encapsulation dot1q
  switchport mode trunk
interface GigabitEthernet0/1-2
  switchport trunk encapsulation dot1q

```

```

switchport mode trunk
interface Vlan1
  no ip address
  shutdown
ip classless
ip flow-export version 9
line con 0
line aux 0
line vty 0 4
login
end

```

Лістинг Г.6 — Лістинг конфігураційного файлу комутатора підстанції 1

```

version 12.2(37)SE1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
hostname Switch
ip dhcp excluded-address 10.1.1.113
ip routing
spanning-tree mode pvst
interface FastEthernet0/1
  switchport mode access
interface FastEthernet0/2
  switchport access vlan 10
  switchport mode access
interface FastEthernet0/3
  switchport access vlan 10
  switchport mode access
interface FastEthernet0/4
  switchport access vlan 10
  switchport mode access
interface FastEthernet0/5
interface FastEthernet0/6
interface FastEthernet0/7
interface FastEthernet0/8
interface FastEthernet0/9
interface FastEthernet0/10
interface FastEthernet0/11
interface FastEthernet0/12
interface FastEthernet0/13
interface FastEthernet0/14
interface FastEthernet0/15
interface FastEthernet0/16
interface FastEthernet0/17
interface FastEthernet0/18
interface FastEthernet0/19
interface GigabitEthernet0/1
  switchport trunk encapsulation dot1q
  switchport mode trunk
interface GigabitEthernet0/2

```

```

interface Vlan1
  no ip address
  shutdown
ip classless
ip flow-export version 9
line con 0
line aux 0
line vty 0 4
  login
end

```

Лістинг Г.7 — Лістинг конфігураційного файлу комутатора вузла агрегації 1

```

version 12.2(37)SE1
no service timestamps log datetime msec
no service timestamps debug datetime msec
no service password-encryption
hostname Switch
spanning-tree mode pvst
interface FastEthernet0/1-24
  switchport trunk encapsulation dot1q
  switchport mode trunk
switchport mode trunk
interface GigabitEthernet0/1
  switchport trunk encapsulation dot1q
  switchport mode trunk
interface GigabitEthernet0/2
  switchport trunk encapsulation dot1q
  switchport mode trunk
interface Vlan1
  no ip address
  shutdown
ip classless
ip flow-export version 9
line con 0
line aux 0
line vty 0 4
  login
end

```