

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
Міжнародний гуманітарний університет
Кафедра комп'ютерної інженерії та інноваційних технологій

Йона Л.Г., Педяш В.В., Мазур Г.Д.

КОМП'ЮТЕРНІ МЕРЕЖІ

Конспект лекцій для здобувачів вищої освіти зі спеціальностей:
A4.09 Середня освіта (Інформатика),
F2 Інженерія програмного забезпечення,
F3 Комп'ютерні науки,
F7 Комп'ютерна інженерія,
F5 Кібербезпека та захист інформації,
G5 Електроніка, електронні комунікації, приладобудування та
радіотехніка

Йона Л.Г., Педяш В.В., Мазур Г.Д. Комп'ютерні мережі: Конспект лекцій [Електронне видання] / Йона Л.Г., Педяш В.В., Мазур Г.Д. — Кафедра комп'ютерної інженерії та інноваційних технологій Міжнародного гуманітарного університету. — Одеса, 2025. - 161 с.

Конспект лекцій з дисципліни «Комп'ютерні мережі» розроблено відповідно до навчального плану підготовки здобувачів вищої освіти. Посібник охоплює 14 тем лекцій, спрямованих на освоєння теоретичних основ побудови, архітектури, протоколів та безпеки комп'ютерних мереж. Матеріали містять систематизований виклад від визначення та класифікації мереж, еталонних моделей OSI та TCP/IP, до налаштування маршрутизації (статичної, динамічної IGP/EGP), механізмів безпеки (ACL, AAA, IPsec VPN), служб адресації (NAT, DHCP) та засобів централізованого моніторингу інфраструктури.

Посібник призначено для здобувачів першого (бакалаврського) рівня вищої освіти спеціальностей: А4.09 Середня освіта (Інформатика), F2 Інженерія програмного забезпечення, F3 Комп'ютерні науки, F7 Комп'ютерна інженерія, F5 Кібербезпека та захист інформації, G5 Електроніка, електронні комунікації, приладобудування та радіотехніка.

Матеріал може використовуватися як для аудиторної роботи під час лекцій, так і для самостійного опрацювання студентами. Теоретичні знання, отримані з посібника, можуть бути закріплені під час виконання практичних завдань у навчальному симуляторі Cisco Packet Tracer, у більш просунутих середовищах (наприклад, GNS3) або на реальному мережевому обладнанні, доступному в лабораторіях факультету кібербезпеки, програмної інженерії та комп'ютерних наук. Такий підхід забезпечує фундаментальну теоретичну підготовку та сприяє формуванню практичних інженерних навичок, необхідних для подальшого вивчення сучасних мережевих технологій, проходження міжнародних сертифікацій (зокрема CCNA) та виконання кваліфікаційних робіт у галузі інформаційно-комунікаційних технологій.

ЗМІСТ

Вступ.....	4
------------	---

Змістовий модуль 1. Основи комп'ютерних мереж та їх архітектурна організація

Тема 1 - Визначення, класифікація та архітектурні рівні комп'ютерних мереж...	5
Тема 2 - Еталонні моделі мережної взаємодії OSI та TCP/IP.....	20
Тема 3 - Середовища передачі, технології та пристрої фізичного і канального рівнів.....	30
Тема 4 - Протоколи мережного рівня, маршрутизація та IP-адресація.....	44
Тема 5 - Протоколи транспортного та прикладного рівнів.....	52
Тема 6 - Механізми безпеки мережевого обладнання.....	60
Тема 7 - Проектування та моделювання комп'ютерних мереж.....	73

Змістовий модуль 2. Технології маршрутизації та адміністрування комп'ютерних мереж

Тема 1 - Протоколи внутрішньої динамічної маршрутизації (IGP)	82
Тема 2 - Протоколи зовнішньої динамічної маршрутизації (EGP).....	103
Тема 3 - Механізми трансляції адрес та автоматичного призначення IP.....	113
Тема 4 - Централізована аутентифікація та система AAA.....	124
Тема 5 - Тунельні технології захисту на основі IPsec VPN.....	135
Тема 6 - Централізований моніторинг та діагностика мережі.....	145
Тема 7 - Архітектура та експлуатаційні характеристики серверного обладнання в інтегрованій IT-інфраструктурі.....	153

Рекомендована література.....	161
-------------------------------	-----

ВСТУП

У сучасному цифровому світі комп'ютерні мережі виступають фундаментальною основою функціонування глобального інформаційного простору, забезпечуючи взаємодію мільярдів пристроїв у режимі реального часу. Стрімкий розвиток інформаційно-комунікаційних технологій, поширення хмарних обчислень, Інтернету речей (IoT) та технологій мобільного зв'язку п'ятого покоління (5G) зумовлює критичну потребу у фахівцях, які глибоко розуміють архітектуру, протоколи та механізми безпеки мережевої інфраструктури. Комп'ютерні мережі є не лише засобом передачі даних, а й ключовим елементом економічної стабільності, національної безпеки та соціальної взаємодії. В умовах зростання кіберзагроз та ускладнення мережевих топологій, якісне опанування дисципліни «Комп'ютерні мережі» стає необхідною умовою для підготовки конкурентоспроможних інженерів та IT-фахівців, здатних проектувати, експлуатувати та захищати сучасні телекомунікаційні системи.

Даний конспект лекцій розроблено кафедрою комп'ютерної інженерії та інноваційних технологій Міжнародного гуманітарного університету відповідно до навчального плану підготовки здобувачів вищої освіти. Посібник містить систематизований виклад матеріалу, охоплюючи 14 тематичних модулів, що забезпечують комплексне вивчення дисципліни. Тематика лекцій послідовно розкриває визначення та класифікацію мереж, еталонні моделі взаємодії OSI та TCP/IP, характеристики середовищ передачі даних та пристроїв фізичного і каналного рівнів. Значна увага приділяється протоколам мережного рівня, методам IP-адресації та алгоритмам маршрутизації (статичної, динамічної IGP/EGP). Окремі розділи присвячені протоколам транспортного та прикладного рівнів, механізмам безпеки мережевого обладнання (ACL, AAA, IPsec VPN), службам адресації (NAT, DHCP), а також засобам централізованого моніторингу та діагностики інфраструктури. Завершує курс тема щодо архітектури та експлуатаційних характеристик серверного обладнання в інтегрований IT-інфраструктурі.

Метою використання цього посібника є забезпечення фундаментальної теоретичної підготовки та формування практичних інженерних навичок у здобувачів першого (бакалаврського) рівня вищої освіти спеціальностей: «Середня освіта (Інформатика)», «Інженерія програмного забезпечення», «Комп'ютерні науки», «Комп'ютерна інженерія», «Кібербезпека та захист інформації», «Електроніка, електронні комунікації, приладобудування та радіотехніка». Матеріал призначений як для аудиторної роботи під час лекцій, так і для самостійного опрацювання. Теоретичні знання, отримані з посібника, рекомендовано закріплювати під час виконання практичних завдань у навчальних симуляторах (Cisco Packet Tracer, GNS3) або на реальному мережевому обладнанні. Такий підхід сприяє підготовці студентів до проходження міжнародних сертифікацій (зокрема CCNA), виконання кваліфікаційних робіт та успішної професійної діяльності у галузі інформаційно-комунікаційних технологій.

Тема 1. Визначення, класифікація та архітектурні рівні комп'ютерних мереж

1.1. Поняття комп'ютерної мережі та класифікація за географічним охопленням

Комп'ютерна мережа являє собою сукупність автономних обчислювальних пристроїв, які з'єднані між собою каналами передачі даних та здатні обмінюватися інформацією за допомогою спільних протоколів комунікації. Основною метою створення комп'ютерних мереж є забезпечення спільного доступу до ресурсів, таких як файли, програми, периферійні пристрої, а також організація ефективної взаємодії між користувачами незалежно від їх географічного розташування. Сучасні комп'ютерні мережі функціонують на основі стандартизованих технологій та протоколів, що дозволяє інтегрувати обладнання різних виробників у єдину інфраструктуру. Розвиток мережевих технологій призвів до формування глобального інформаційного простору, в якому мільярди пристроїв можуть взаємодіяти між собою в режимі реального часу.

Принциповою особливістю комп'ютерної мережі є наявність принаймні двох або більше взаємопов'язаних пристроїв, здатних передавати та приймати дані. Кожен елемент мережі повинен мати унікальний ідентифікатор, що дозволяє однозначно визначити джерело та призначення інформаційного потоку. Комунікація в мережі здійснюється за допомогою спеціалізованих протоколів, які визначають правила форматування, передачі, маршрутизації та обробки даних. Важливим аспектом функціонування мережі є забезпечення надійності передачі інформації, що досягається через механізми виявлення та виправлення помилок, підтвердження доставки пакетів та повторної передачі у разі втрати даних. Сучасна комп'ютерна мережа є складною системою, яка включає фізичні компоненти, програмне забезпечення, протоколи взаємодії та адміністративні політики управління.

Класифікація комп'ютерних мереж за територіальним принципом є фундаментальним способом систематизації різноманітних мережевих архітектур та визначає основні характеристики їх побудови і функціонування. Відповідно до географічного охоплення, мережі поділяються на чотири основні категорії: персональні мережі (PAN), локальні мережі (LAN), міські мережі (MAN) та глобальні мережі (WAN). Кожна з цих категорій характеризується специфічними технологіями передачі даних, пропускну здатністю каналів зв'язку, методами організації доступу до середовища передачі та протоколами маршрутизації. Розуміння особливостей різних типів мереж є критично важливим для проектування ефективної мережевої інфраструктури, яка відповідає конкретним потребам організації або приватних користувачів.

Класифікація комп'ютерних мереж за територіальним принципом є фундаментальним способом систематизації різноманітних мережевих архітектур та визначає основні характеристики їх побудови і функціонування. Відповідно до географічного охоплення, мережі поділяються на чотири основні

категорії: персональні мережі (PAN), локальні мережі (LAN), міські мережі (MAN) та глобальні мережі (WAN). Кожна з цих категорій характеризується специфічними технологіями передачі даних, пропускнуою здатністю каналів зв'язку, методами організації доступу до середовища передачі та протоколами маршрутизації. Розуміння особливостей різних типів мереж є критично важливим для проектування ефективної мережевої інфраструктури, яка відповідає конкретним потребам організації або приватних користувачів.

Персональна мережа (Personal Area Network, PAN) представляє собою найменшу за масштабом категорію комп'ютерних мереж, що охоплює простір безпосередньо навколо окремого користувача, зазвичай в межах кількох метрів. Типовим прикладом PAN є з'єднання смартфона з бездротовою гарнітурою через Bluetooth, підключення ноутбука до мишки або клавіатури, синхронізація носимих пристроїв з мобільним телефоном. Технології, що використовуються в персональних мережах, включають Bluetooth, Zigbee, інфрачервоний зв'язок IrDA та технологію ближнього поля NFC. Характерною особливістю PAN є мінімальне енергоспоживання пристроїв, невисока пропускна здатність каналів та відсутність складної інфраструктури для організації зв'язку. Розвиток концепції Інтернету речей сприяє активному розширенню застосування персональних мереж для об'єднання різноманітних побутових та носимих електронних пристроїв.

Локальна мережа (Local Area Network, LAN) охоплює обмежену територію, як правило, в межах одного будинку, офісу, навчального закладу або групи будівель в межах одного кампусу. Розміри LAN зазвичай не перевищують кількох кілометрів, що дозволяє організувати високошвидкісну передачу даних з мінімальними затримками. Типова локальна мережа будується на основі технології Ethernet та забезпечує швидкість передачі даних від 100 Мбіт/с до 100 Гбіт/с та вище, залежно від використаного обладнання і середовища передачі. Архітектура LAN може включати дротові з'єднання через вите пару або оптичне волокно, а також бездротові сегменти на базі стандарту Wi-Fi. Важливою перевагою локальних мереж є можливість централізованого управління, високий рівень безпеки завдяки фізичній ізоляції від зовнішнього середовища та низька вартість обслуговування порівняно з географічно розподіленими мережами.

Міська мережа (Metropolitan Area Network, MAN) займає проміжне положення між локальними та глобальними мережами і охоплює територію міста або великого мегаполісу з діаметром до кількох десятків кілометрів. MAN зазвичай використовуються телекомунікаційними операторами для об'єднання корпоративних локальних мереж, розташованих у різних районах міста, або для надання широкосмугового доступу до мережі Інтернет великій кількості абонентів. Технологічною основою міських мереж часто виступають оптоволоконні магістралі, технології Metro Ethernet, бездротові стандарти WiMAX або мікрохвильові лінії зв'язку. Типовим прикладом MAN є мережа кабельного телебачення з можливістю надання послуг інтернет-доступу, мережа операторів мобільного зв'язку в межах міста або корпоративна розподілена мережа великої організації з численними філіями. Швидкість

передачі даних в міських мережах може досягати десятків і сотень гігабіт на секунду, що забезпечує ефективне надання різноманітних телекомунікаційних послуг.

Глобальна мережа (Wide Area Network, WAN) об'єднує комп'ютери та локальні мережі, розташовані на значних відстанях один від одного, в межах країни, континенту або по всьому світу. Найбільшим та найвідомішим прикладом глобальної мережі є Інтернет, який з'єднує мільярди пристроїв по всій планеті. WAN зазвичай будуються з використанням орендованих каналів зв'язку телекомунікаційних операторів, супутникових ліній, підводних оптоволоконних кабелів та інших технологій дальнього зв'язку. Характерними особливостями глобальних мереж є відносно низька пропускна здатність порівняно з LAN, більші затримки передачі даних, вища вартість каналів зв'язку та підвищені вимоги до надійності та відмовостійкості. Для організації WAN використовуються спеціалізовані протоколи маршрутизації, технології віртуальних приватних мереж VPN, механізми оптимізації трафіку та забезпечення якості обслуговування QoS, що дозволяє ефективно керувати передачею різнотипної інформації через географічно розподілену інфраструктуру.

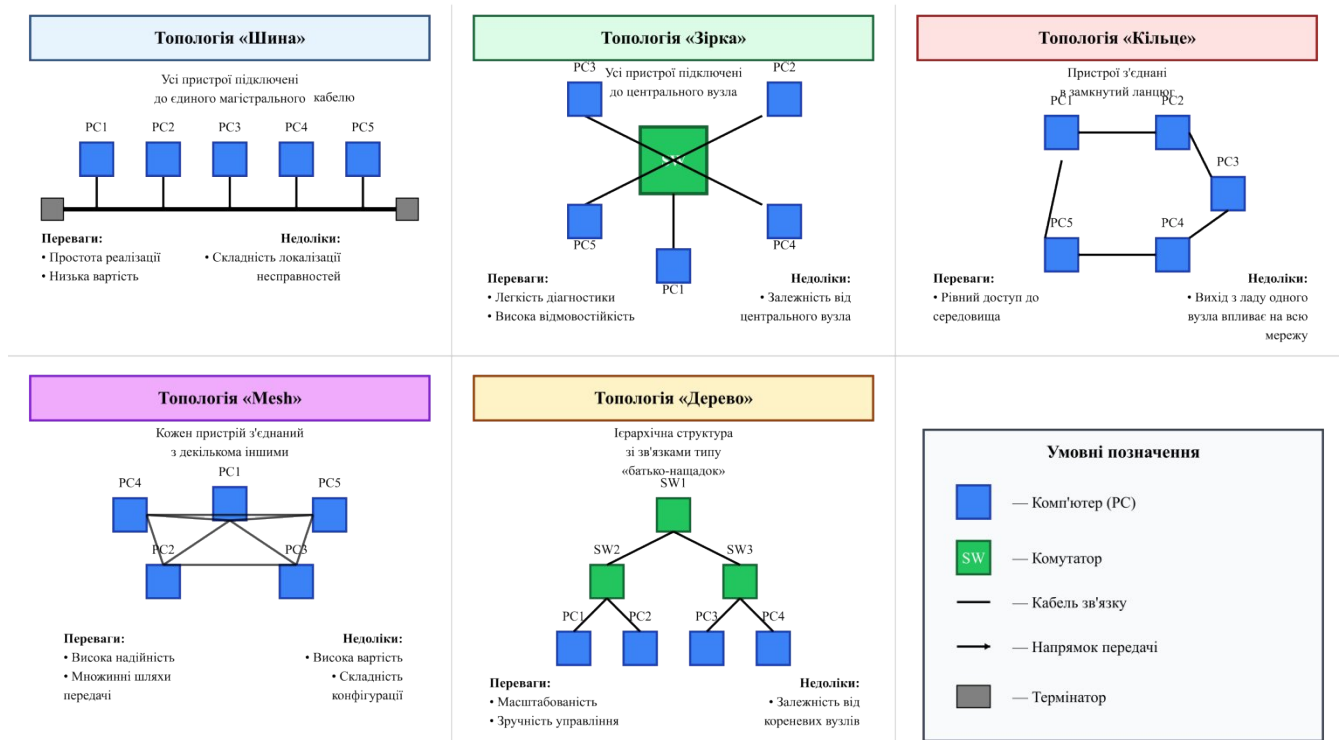
1.2. Топології мереж та їх характеристики

Топологія комп'ютерної мережі визначає спосіб організації з'єднань між вузлами мережі та шляхи передачі даних між ними, при цьому розрізняють фізичну та логічну топології, які можуть суттєво відрізнятися одна від одної. Фізична топологія описує реальне географічне розташування мережевих пристроїв та фактичну схему прокладання кабельних з'єднань між ними, тоді як логічна топологія визначає шляхи проходження сигналів та даних через мережу незалежно від фізичного розміщення компонентів. Розуміння різниці між цими двома аспектами топології є критично важливим для проектування, діагностики та оптимізації мережевої інфраструктури. У таблиці 1.1 наведено порівняльну характеристику основних типів мережевих топологій з їх ключовими особливостями.

Таблиця 1.1 - Порівняльна характеристика мережевих топологій

Тип топології	Опис	Переваги	Недоліки
Шина	Усі пристрої підключені до єдиного магістрального кабелю	Простота реалізації, низька вартість	Складність локалізації несправностей, низька відмовостійкість
Зірка	Усі пристрої підключені до центрального вузла	Легкість діагностики, висока відмовостійкість	Залежність від центрального вузла
Кільце	Пристрої з'єднані в замкнутий ланцюг	Рівний доступ до середовища, передбачувана затримка	Вихід з ладу одного вузла впливає на всю мережу

Mesh	Кожен пристрій з'єднаний з декількома іншими	Висока надійність, множинні шляхи передачі	Висока вартість, складність конфігурації
Дерево	Ієрархічна структура зі зв'язками типу "батько-нащадок"	Масштабованість, зручність управління	Залежність від корневих вузлів



Фізична топологія зірка є найбільш поширеною в сучасних локальних мережах і передбачає підключення всіх кінцевих пристроїв до центрального комутатора або концентратора за допомогою окремих кабельних сегментів. Перевагою такої архітектури є простота додавання нових пристроїв до мережі, легкість виявлення та усунення несправностей на окремих з'єднаннях, а також висока відмовостійкість, оскільки вихід з ладу одного кабелю впливає лише на один комп'ютер. Водночас логічна топологія в такій мережі може відрізнятися від фізичної: наприклад, при використанні концентратора логічна топологія залишається шиною, тоді як при застосуванні комутатора створюється логічна топологія типу "точка-точка" між кожною парою взаємодіючих пристроїв. Центральний елемент мережі стає критичною точкою відмови, тому в професійних інсталяціях часто застосовується резервування комутаторів для забезпечення безперервності роботи.

Топологія шина історично використовувалася в ранніх версіях Ethernet та передбачала підключення всіх пристроїв до єдиного магістрального коаксіального кабелю за допомогою Т-подібних роз'ємів. Дані, передані будь-яким вузлом, поширювалися вздовж усього кабелю і приймалися всіма підключеними пристроями, але оброблялися лише тим вузлом, якому вони були адресовані. Основним недоліком такої топології була низька відмовостійкість: обрив або пошкодження магістрального кабелю призводили до повного

припинення функціонування всієї мережі. Крім того, складність локалізації місця несправності та обмеження пропускну здатності через спільне використання середовища передачі зробили цю топологію непридатною для сучасних високошвидкісних мереж, внаслідок чого вона практично не використовується в нових інсталяціях.

Кільцева топологія реалізує послідовне з'єднання мережевих пристроїв у замкнутий контур, де кожен вузол з'єднаний рівно з двома сусідніми пристроями. Дані передаються по кільцю в одному напрямку від вузла до вузла, при цьому кожен проміжний пристрій регенерує та ретранслює сигнал далі по ланцюгу. Прикладом мереж з кільцевою топологією є технології Token Ring та FDDI, де використовувався механізм маркера для керування доступом до середовища передачі. Для підвищення відмовостійкості в критичних системах застосовується подвійне кільце, де дані можуть передаватися в обох напрямках, що дозволяє зберегти працездатність мережі навіть при обриві одного з кабелів. Сучасні Metro Ethernet мережі часто використовують логічну кільцеву топологію для організації захищених і надійних магістральних з'єднань між вузлами міської мережі.

Топологія mesh або повнозв'язна топологія характеризується наявністю прямих з'єднань між кожною парою вузлів мережі, що забезпечує максимальну надійність та множинні альтернативні шляхи передачі даних. У повній mesh-топології кількість з'єднань зростає квадратично відносно кількості вузлів, що робить її економічно недоцільною для великих мереж через високу вартість кабельної інфраструктури та портів обладнання. Натомість часто використовується часткова mesh-топологія, де критично важливі вузли з'єднані декількома резервними каналами, а решта пристроїв мають обмежену кількість з'єднань. Така архітектура типова для побудови магістралей глобальних мереж, дата-центрів та систем високої доступності, де надійність і безперервність роботи є пріоритетними вимогами. Mesh-топологія також активно використовується в бездротових мережах на базі стандарту IEEE 802.11s, де кожна точка доступу може виконувати функції маршрутизатора для інших вузлів.

1.3. Категорії мережевого обладнання та їх функціональне призначення

Мережева інфраструктура складається з різноманітного спеціалізованого обладнання, кожен тип якого виконує специфічні функції на певному рівні моделі OSI та забезпечує ефективну передачу, комутацію та маршрутизацію даних. Класифікація мережевих пристроїв базується на їх функціональному призначенні, рівні роботи в стеку протоколів та впливі на процес обробки мережевого трафіку. До основних категорій мережевого обладнання належать повторювачі, концентратори, комутатори, маршрутизатори, точки доступу та міжмережеві екрани, кожен з яких має унікальні характеристики та сферу застосування. Правильний вибір і конфігурація мережевого обладнання є

критично важливим фактором для забезпечення необхідної продуктивності, безпеки та надійності мережевої інфраструктури.

Повторювач представляє собою найпростіший пристрій фізичного рівня, призначений для регенерації та підсилення електричних або оптичних сигналів з метою збільшення максимальної довжини мережевого сегмента. В процесі передачі сигналу по кабелю відбувається його загасання та спотворення через активний опір провідників, електромагнітні перешкоди та дисперсію, що обмежує відстань надійної комунікації між пристроями. Повторювач приймає ослаблений сигнал, відновлює його форму, амплітуду та часові характеристики, після чого передає регенований сигнал далі по мережевому сегменту. Важливо розуміти, що повторювач не аналізує зміст даних і не приймає рішень щодо маршрутизації, а лише виконує фізичну ретрансляцію сигналу, тому всі пристрої, з'єднані через повторювач, залишаються в межах одного колізійного домену. Сучасні мережі рідко використовують автономні повторювачі, оскільки їх функції інтегровані в більш складне активне обладнання.

Концентратор або хаб є багатопортовим повторювачем, який об'єднує декілька мережевих сегментів у єдиний колізійний домен і працює на фізичному рівні моделі OSI. При надходженні сигналу на будь-який порт концентратора, пристрій регенує цей сигнал і передає його одночасно на всі інші порти, незалежно від реального призначення даних, що значно знижує ефективну пропускну здатність мережі при збільшенні кількості активних вузлів. Всі пристрої, підключені до концентратора, спільно використовують доступну пропускну здатність і повинні застосовувати алгоритм CSMA/CD для уникнення та вирішення колізій при одночасній передачі даних. Через низьку продуктивність та неефективне використання мережевих ресурсів концентратори практично повністю витіснені з сучасних мереж комутаторами, які забезпечують значно вищу пропускну здатність та ізоляцію трафіку.

Комутатор є інтелектуальним пристроєм канального рівня, який аналізує MAC-адреси в кадрах Ethernet і приймає рішення про пересилання даних лише на той порт, до якого підключений пристрій-призначення. На відміну від концентратора, комутатор створює окремі колізійні домени для кожного порту, що дозволяє встановлювати одночасні незалежні з'єднання між різними парами пристроїв і забезпечує повнодуплексний режим роботи з подвоєною сумарною пропускну здатністю. Комутатор підтримує внутрішню таблицю комутації або САМ-таблицю, в якій зберігаються відповідності між MAC-адресами та номерами портів, причому ця таблиця формується автоматично шляхом аналізу адрес джерел у кадрах, що надходять. Існують комутатори другого рівня, які працюють виключно з MAC-адресами, та комутатори третього рівня, які додатково можуть маршрутизувати трафік між різними IP-підмережами, об'єднуючи функції комутації та маршрутизації в одному пристрої для підвищення продуктивності.

Маршрутизатор функціонує на мережевому рівні моделі OSI та призначений для з'єднання різних мереж або підмереж шляхом аналізу IP-адрес призначення в пакетах і вибору оптимального шляху їх пересилання. Основною відмінністю маршрутизатора від комутатора є робота з логічною адресацією

третього рівня замість фізичних MAC-адрес, що дозволяє створювати складні ієрархічні мережеві структури з можливістю оптимізації маршрутів на основі метрик відстані, пропускну здатності, затримки або політик безпеки. Маршрутизатор підтримує таблицю маршрутизації, яка містить інформацію про доступні мережі та шляхи до них, причому записи в цій таблиці можуть бути створені вручну адміністратором або автоматично за допомогою протоколів динамічної маршрутизації RIP, OSPF, EIGRP або BGP. Кожен інтерфейс маршрутизатора належить до окремої IP-підмережі і формує межу широкомовного домену, що запобігає поширенню широкомовного трафіку між сегментами мережі та покращує загальну продуктивність і безпеку інфраструктури.

Точка доступу є спеціалізованим мережевим пристроєм, який забезпечує бездротове з'єднання клієнтських пристроїв з дротовою мережевою інфраструктурою на основі стандартів IEEE 802.11. Точка доступу працює на фізичному та канальному рівнях, перетворюючи дротовий Ethernet-трафік у бездротові радіосигнали та навпаки, при цьому вона може функціонувати в різних режимах, включаючи інфраструктурний режим, режим моста для з'єднання віддалених мережевих сегментів або режим ретранслятора для розширення зони покриття. Сучасні точки доступу підтримують множинні SSID для створення декількох логічних бездротових мереж з різними політиками безпеки, механізми автентифікації та шифрування WPA2/WPA3, можливості керування потужністю передавача та автоматичного вибору радіоканалів для мінімізації інтерференції. Корпоративні точки доступу часто керуються централізовано через контролер бездротових мереж, що спрощує конфігурацію, моніторинг та забезпечення безпеки в масштабних розгортаннях.

Міжмережевий екран або файрвол є засобом забезпечення мережевої безпеки, який контролює вхідний і вихідний трафік на основі визначених правил безпеки та дозволяє або блокує передачу даних залежно від IP-адрес, протоколів, портів та інших параметрів. Міжмережеві екрани можуть функціонувати на різних рівнях моделі OSI: пакетні фільтри працюють на мережевому рівні і аналізують заголовки IP-пакетів, шлюзи сеансового рівня відстежують стан TCP-з'єднань, а прикладні шлюзи виконують глибокий аналіз змісту на рівні протоколів додатків. Сучасні міжмережеві екрани нового покоління поєднують традиційну фільтрацію з функціями виявлення та запобігання вторгненням, антивірусного сканування, фільтрації веб-контенту та інспекції шифрованого трафіку для комплексного захисту мережевого периметру. Правильна конфігурація міжмережевого екрана передбачає реалізацію принципу найменших привілеїв, коли за замовчуванням весь трафік блокується, а явно дозволяються лише необхідні для роботи з'єднання.

На рисунку 1.1 наведено типову структуру корпоративної мережі з використанням різних категорій мережевого обладнання та їх ієрархічним розташуванням.

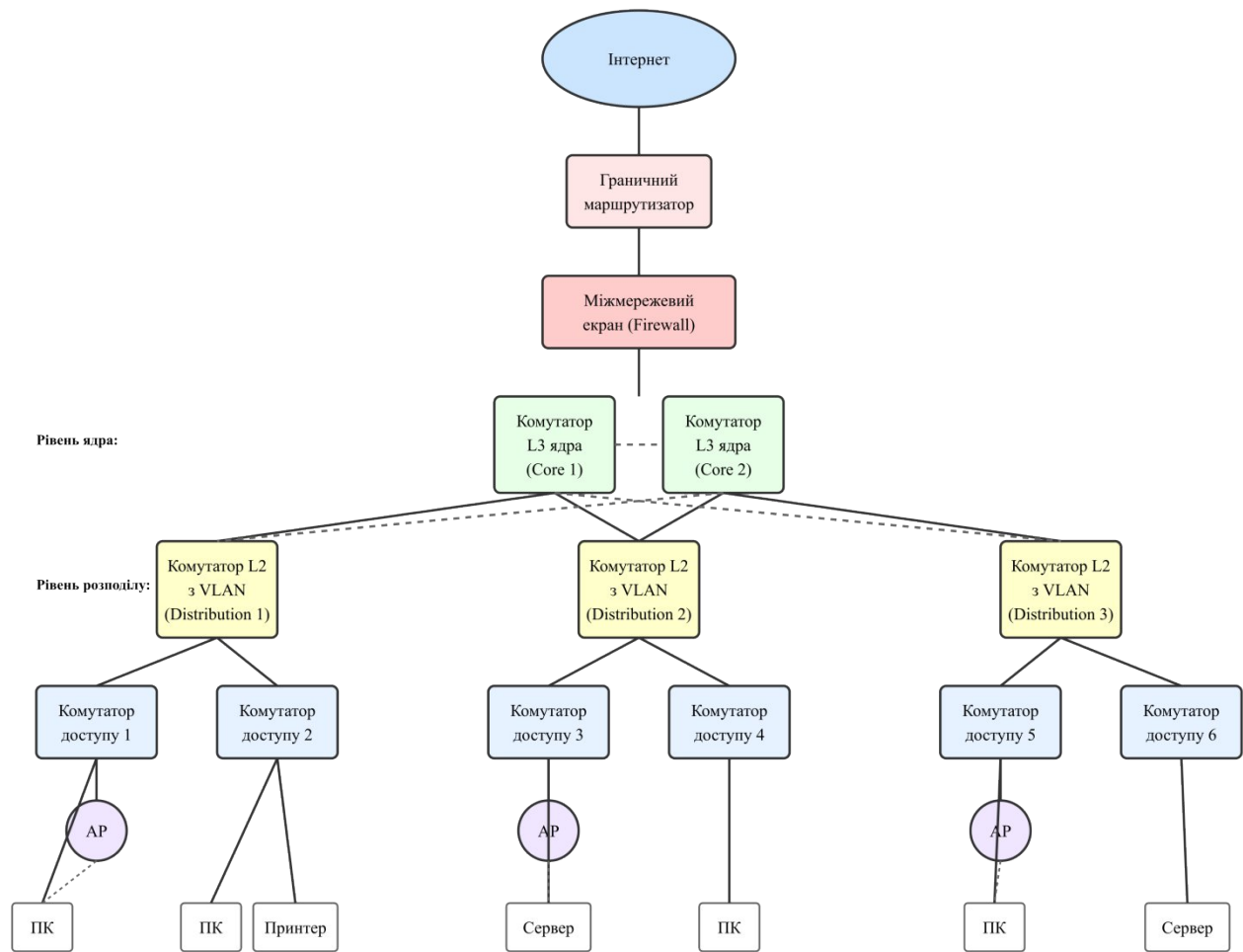


Рисунок 1.1 - Архітектура корпоративної мережі з різними типами обладнання

Рисунок демонструє трирівневу ієрархічну модель мережі, де на верхньому рівні розташований граничний маршрутизатор, який з'єднує корпоративну мережу з мережею провайдера та Інтернетом. Між граничним маршрутизатором та внутрішньою мережею встановлено міжмеревий екран для контролю трафіку та забезпечення безпеки периметра. На рівні ядра мережі розміщено високопродуктивні комутатори третього рівня, які забезпечують швидку маршрутизацію між різними сегментами внутрішньої мережі та підтримують резервні з'єднання для відмовостійкості. Рівень розподілу представлений комутаторами другого рівня з підтримкою віртуальних локальних мереж VLAN, які агрегують трафік від комутаторів рівня доступу і реалізують політики безпеки та якості обслуговування. На рівні доступу розташовані комутатори, до яких безпосередньо підключаються кінцеві пристрої користувачів, сервери відділів та принтери. У кожному сегменті мережі також присутні точки бездротового доступу, які під'єднані до комутаторів рівня доступу і забезпечують мобільність користувачів. Така архітектура забезпечує чітке розмежування функцій, масштабованість, відмовостійкість через резервування критичних компонентів та спрощує адміністрування завдяки модульній структурі.

Різні категорії мережевого обладнання функціонують на специфічних рівнях еталонної моделі OSI, що визначає їх можливості з обробки мережевого трафіку та взаємодії з іншими компонентами інфраструктури. Розуміння відповідності між типами пристроїв та рівнями моделі OSI є фундаментальним для ефективного проектування мереж, діагностики проблем та оптимізації продуктивності. Повторювачі та концентратори працюють виключно на фізичному рівні і не мають можливості аналізувати структуру кадрів або пакетів, натомість вони лише регенерують електричні або оптичні сигнали без жодної обробки даних. Комутатори функціонують переважно на канальному рівні, аналізуючи MAC-адреси для прийняття рішень про пересилання кадрів, хоча існують комутатори третього рівня, які додатково виконують функції маршрутизації на основі IP-адрес.

Маршрутизатори є типовими представниками пристроїв мережевого рівня і спеціалізуються на аналізі логічної адресації для визначення оптимальних шляхів доставки пакетів між різними мережами. Робота на третьому рівні дозволяє маршрутизаторам приймати рішення на основі глобальної топології мережі, метрик маршрутизації та політик безпеки, що робить їх ключовими елементами для побудови складних розподілених інфраструктур. Міжмережеві екрани можуть функціонувати на декількох рівнях одночасно: базові пакетні фільтри працюють на мережевому рівні, екрани з відстеженням стану аналізують інформацію транспортного рівня, а системи глибокої інспекції вмісту обробляють дані на прикладному рівні для виявлення шкідливого трафіку.

Важливою характеристикою мережевого обладнання є поняття колізійного та широкомовного доменів, які визначають область поширення відповідних типів трафіку. Колізійний домен охоплює всі пристрої, які конкурують за доступ до спільного фізичного середовища передачі, і в межах якого можливе виникнення колізій при одночасній передачі даних декількома вузлами. Повторювачі та концентратори не розділяють колізійні домени, тому всі підключені до них пристрої залишаються в межах одного домену з обмеженою сумарною пропускну здатністю. Комутатори створюють окремий колізійний домен для кожного порту при роботі в повнодуплексному режимі, що усуває можливість колізій і дозволяє максимально ефективно використовувати пропускну здатність кожного з'єднання.

Широкомовний домен визначає область поширення широкомовних кадрів, які адресовані всім пристроям у мережі і повинні бути оброблені кожним вузлом незалежно від того, чи призначена йому ця інформація. Комутатори другого рівня не обмежують поширення широкомовного трафіку і всі порти комутатора або група портів у одному VLAN належать до єдиного широкомовного домену. Маршрутизатори формують межі широкомовних доменів, оскільки не пересилають широкомовні пакети між різними інтерфейсами, що дозволяє ізолювати сегменти мережі та зменшити непродуктивне навантаження на кінцеві пристрої. Правильна сегментація мережі на широкомовні домени за допомогою маршрутизаторів або VLAN є

важливим аспектом оптимізації продуктивності великих корпоративних інфраструктур.

Спеціалізоване мережеве обладнання також включає пристрої, орієнтовані на вирішення специфічних завдань, таких як балансувальники навантаження, які розподіляють вхідні запити між декількома серверами для забезпечення високої доступності та продуктивності веб-додатків. Системи запобігання вторгненням IPS активно блокують підозрілий або шкідливий трафік на основі сигнатурного аналізу та поведінкових аномалій, доповнюючи функції традиційних міжмережевих екранів. Контролери бездротових мереж централізовано керують множиною точок доступу, автоматизують конфігурацію, забезпечують роумінг клієнтів між точками доступу та реалізують політики безпеки в масштабних корпоративних Wi-Fi мережах. Мережеві комутатори можуть також підтримувати функцію Power over Ethernet, яка дозволяє передавати електроживлення разом з даними по одному кабелю витой пари, спрощуючи інсталяцію пристроїв у місцях, де складно організувати окреме електропостачання.

Вибір відповідного мережевого обладнання є критичним етапом проектування інфраструктури і повинен базуватися на комплексному аналізі вимог до продуктивності, масштабованості, надійності, безпеки та бюджетних обмежень конкретного проекту. При виборі комутаторів необхідно враховувати кількість і типи портів, швидкість комутації, підтримку VLAN та протоколів керування, наявність механізмів якості обслуговування QoS, можливості моніторингу трафіку та резервування критичних компонентів. Для маршрутизаторів ключовими параметрами є продуктивність обробки пакетів, кількість одночасних маршрутів у таблиці маршрутизації, підтримка протоколів динамічної маршрутизації, можливості VPN та фільтрації трафіку. Точки бездротового доступу вибираються з урахуванням стандартів Wi-Fi, кількості одночасних клієнтів, наявності механізмів керування потужністю та автоматичного вибору каналів, інтеграції з системами автентифікації та можливості централізованого управління.

Конфігурація мережевого обладнання зазвичай здійснюється через інтерфейс командного рядка CLI або веб-інтерфейс, причому для професійного обладнання переважає використання CLI через більшу гнучкість та можливість автоматизації. Початкова конфігурація передбачає призначення IP-адрес інтерфейсам, встановлення паролів доступу, налаштування протоколів маршрутизації або комутації, створення VLAN та визначення політик безпеки. Важливим аспектом є документування всіх змін конфігурації та регулярне створення резервних копій налаштувань для можливості швидкого відновлення після збоїв обладнання. Моніторинг стану мережеских пристроїв здійснюється за допомогою протоколу SNMP, системних журналів syslog та спеціалізованих систем мережевого менеджменту, які збирають статистику завантаження портів, помилок передачі та доступності пристроїв для проактивного виявлення проблем.

Сучасні тенденції в розвитку мережевого обладнання включають програмно-конфігуровані мережі SDN, де площина управління відокремлена

від площини даних і централізовано керується програмним контролером, що забезпечує динамічну адаптацію мережі до змінних умов навантаження. Віртуалізація мережевих функцій NFV дозволяє реалізувати функціонал маршрутизаторів, міжмережевих екранів та балансувальників навантаження у вигляді програмних компонентів, які виконуються на універсальних серверах замість спеціалізованого апаратного забезпечення. Такі підходи підвищують гнучкість інфраструктури, зменшують капітальні витрати та спрощують масштабування мережевих сервісів відповідно до поточних потреб організації.

1.4. Продуктивність, стандартизація та сучасні тенденції розвитку мережевих технологій

Продуктивність мережевого обладнання визначається комплексом параметрів, які впливають на швидкість та ефективність обробки мережевого трафіку в різних умовах навантаження. Пропускна здатність комутаторів та маршрутизаторів вимірюється в пакетах за секунду або в бітах за секунду і характеризує максимальну кількість даних, яку пристрій здатен обробити без втрат та затримок. Для комутаторів критичним параметром є швидкість комутації або *switching fabric capacity*, яка визначає внутрішню пропускну здатність між портами та повинна перевищувати сумарну швидкість всіх портів для уникнення блокування трафіку. Затримка комутації або *forwarding delay* характеризує час між надходженням кадру на вхідний порт та його передачею на вихідний порт, причому існують різні методи комутації: *cut-through* з мінімальною затримкою починає передачу відразу після отримання адреси призначення, *store-and-forward* зберігає весь кадр для перевірки цілісності, а *fragment-free* є компромісом між швидкістю та надійністю.

Маршрутизатори оцінюються за продуктивністю пересилання пакетів, яка залежить від складності аналізу таблиці маршрутизації, кількості правил ACL та необхідності фрагментації або дефрагментації пакетів. Апаратні маршрутизатори використовують спеціалізовані інтегральні схеми ASIC для високошвидкісної обробки заголовків пакетів з мінімальною затримкою, тоді як програмні маршрутизатори на базі універсальних процесорів забезпечують більшу гнучкість конфігурації за рахунок зниженої продуктивності. Важливим показником є розмір таблиці маршрутизації, яку пристрій здатен підтримувати в оперативній пам'яті, оскільки магістральні маршрутизатори інтернет-провайдерів повинні зберігати сотні тисяч маршрутів для забезпечення глобальної зв'язності.

Відмовостійкість мережевого обладнання забезпечується через резервування критичних компонентів, таких як блоки живлення, вентилятори охолодження, процесорні модулі та інтерфейсні карти. Високоступні системи використовують протоколи резервування першого переходу, такі як HSRP, VRRP або GLBP, які дозволяють декільком маршрутизаторам спільно обслуговувати один шлюз за замовчуванням з автоматичним перемиканням у разі відмови активного пристрою. Комутатори підтримують протокол Spanning Tree для уникнення петель комутації при наявності резервних з'єднань між

пристроями, автоматично блокуючи надлишкові шляхи в нормальному режимі та активуючи їх при виході з ладу основного каналу. Час відновлення після збою або failover time є критичним параметром для систем, які потребують безперервної роботи, і може варіюватися від десятків мілісекунд до декількох секунд залежно від використаних механізмів резервування.

Масштабованість мережевої інфраструктури передбачає можливість нарощування кількості підключених пристроїв, пропускної здатності каналів та функціональних можливостей без фундаментальної зміни архітектури. Модульна конструкція корпоративних комутаторів та маршрутизаторів дозволяє додавати інтерфейсні модуля з різними типами портів, розширювати оперативну пам'ять та процесорну потужність відповідно до зростаючих вимог. Технологія стекування комутаторів об'єднує декілька фізичних пристроїв в один логічний комутатор з єдиною площиною управління та високошвидкісними міжз'єднаннями між модулями, що спрощує адміністрування та підвищує надійність порівняно з окремими незалежними пристроями. Віртуальна комутація або chassis virtualization дозволяє керувати групою комутаторів як єдиною системою з можливістю створення портових каналів між різними фізичними пристроями та синхронізації конфігурації.

Енергоефективність стає все більш важливим критерієм вибору мережевого обладнання в контексті зростаючих вимог до екологічної відповідальності та скорочення операційних витрат на електроенергію та охолодження. Сучасні комутатори підтримують стандарт Energy Efficient Ethernet, який автоматично знижує споживання енергії портами при низькому навантаженні або в період неактивності шляхом переходу в режим низького енергоспоживання. Технологія Power over Ethernet Plus дозволяє передавати до тридцяти ват електроживлення на підключений пристрій, що достатньо для живлення IP-телефонів, точок доступу та камер відеоспостереження без необхідності окремих джерел живлення. Інтелектуальне керування вентиляторами охолодження адаптує їх швидкість обертання залежно від температури компонентів, зменшуючи енергоспоживання та рівень шуму в періоди низького навантаження.

Взаємна сумісність мережевого обладнання різних виробників забезпечується дотриманням міжнародних стандартів, розроблених організаціями IEEE, IETF, ITU-T та ISO. Стандарти IEEE 802.3 визначають фізичні характеристики та протоколи канального рівня для технології Ethernet, включаючи специфікації для різних швидкостей передачі від десяти мегабіт до ста гігабіт на секунду та типи середовищ передачі. Протоколи маршрутизації, такі як OSPF та BGP, стандартизовані в RFC документах IETF, що гарантує можливість побудови гетерогенних мереж з обладнанням різних виробників. Стандарт IEEE 802.1Q визначає механізм тегування VLAN для передачі трафіку декількох віртуальних мереж через один фізичний канал, забезпечуючи сумісність комутаторів при реалізації складних сегментованих інфраструктур.

Проприетарні розширення стандартних протоколів розробляються виробниками для надання додаткових можливостей їх обладнанню, але можуть призводити до проблем сумісності при інтеграції з пристроями інших

виробників. Наприклад, протокол EIGRP компанії Cisco забезпечує швидшу збіжність маршрутизації порівняно зі стандартним OSPF, але не підтримується обладнанням інших виробників, що обмежує можливості побудови мультивендорних мереж. Протокол VTP для автоматичної синхронізації конфігурації VLAN між комутаторами також є проприетарною розробкою і вимагає однорідної інфраструктури з обладнанням одного виробника. При проектуванні мереж необхідно балансувати між перевагами проприетарних технологій та вимогами до гнучкості і незалежності від постачальника обладнання.

Сертифікація та тестування мережевого обладнання проводиться незалежними організаціями для підтвердження відповідності стандартам та оцінки продуктивності в різних сценаріях використання. Програми сертифікації виробників, такі як Cisco CCIE або Juniper JNCIE, підтверджують кваліфікацію інженерів у конфігурації та підтримці складних мережесих інфраструктур на основі обладнання конкретного виробника. Участь у процесах стандартизації дозволяє виробникам впливати на розробку нових специфікацій та забезпечити підтримку майбутніх технологій у своїх продуктах. Відкриті специфікації та програмні інтерфейси API стають все більш важливими в контексті програмно-конфігурованих мереж, де обладнання повинно динамічно керуватися централізованими контролерами незалежно від виробника.

Еволюція мережесих технологій характеризується постійним зростанням швидкості передачі даних, впровадженням інтелектуальних механізмів управління трафіком та конвергенцією різних типів комунікацій на єдиній IP-інфраструктурі. Перехід до Ethernet зі швидкістю сто гігабіт на секунду та розробка стандартів для чотирьохсот гігабітного Ethernet відкривають нові можливості для побудови високопродуктивних дата-центрів та магістральних мереж операторів зв'язку. Технології оптичної комутації дозволяють передавати терабіти даних на великі відстані без електронної регенерації сигналу, що значно знижує затримки та енергоспоживання в транспортних мережах. Розвиток бездротових технологій п'ятого покоління 5G забезпечує швидкість передачі даних, порівнянну з дротовими з'єднаннями, та відкриває можливості для масового впровадження Інтернету речей з підключенням мільярдів сенсорів та пристроїв.

Штучний інтелект та машинне навчання інтегруються в системи управління мережами для автоматичного виявлення аномалій, прогнозування збоїв обладнання та оптимізації маршрутизації трафіку на основі аналізу історичних даних. Самонавчальні системи можуть автоматично адаптувати конфігурацію мережі до змінних умов навантаження, виявляти нові типи кібератак та рекомендувати заходи для усунення виявлених проблем без втручання адміністратора. Технології intent-based networking дозволяють описувати бажаний стан мережі на рівні бізнес-вимог, після чого система автоматично генерує необхідну конфігурацію обладнання та постійно перевіряє відповідність реального стану заданим політикам. Такий підхід суттєво спрощує управління складними розподіленими інфраструктурами та зменшує

ймовірність помилок конфігурації, які є основною причиною збоїв в корпоративних мережах.

Контейнеризація мережевих сервісів на базі технологій Docker та Kubernetes дозволяє швидко розгортати та масштабувати мережеві функції в хмарних середовищах з мінімальними накладними витратами. Мікросервісна архітектура розбиває монолітні мережеві додатки на набір незалежних компонентів, які можуть окремо оновлюватися та масштабуватися залежно від навантаження на конкретні функції. Edge computing переносить обчислювальні ресурси ближче до джерел даних для зменшення затримок та навантаження на магістральні канали зв'язку, що особливо важливо для додатків реального часу, таких як автономні транспортні засоби або доповнена реальність. Розвиток квантових комунікацій обіцяє революцію в області захисту даних завдяки фізичній неможливості перехоплення квантових ключів шифрування без виявлення факту компрометації каналу зв'язку.

Ефективна експлуатація мережевого обладнання вимагає систематичного моніторингу стану пристроїв, проактивного обслуговування та дотримання найкращих практик конфігурації та безпеки. Системи моніторингу на базі протоколу SNMP періодично опитують мережеві пристрої для збору статистики завантаження інтерфейсів, кількості помилок передачі, температури компонентів та стану резервних блоків живлення. Централізовані системи збору журналів syslog агрегують повідомлення про події від усього мережевого обладнання, що дозволяє аналізувати інциденти безпеки, виявляти нестабільні з'єднання та діагностувати проблеми продуктивності. Автоматичні системи оповіщення інформують адміністраторів про критичні події, такі як вихід з ладу інтерфейсу, перевищення порогів завантаження або недоступність пристрою, що дозволяє швидко реагувати на проблеми до того, як вони вплинуть на користувачів.

Регулярне оновлення мікропрограмного забезпечення мережевого обладнання є критично важливим для усунення вразливостей безпеки, виправлення помилок та отримання доступу до нових функціональних можливостей. Процедура оновлення повинна включати попереднє тестування нової версії firmware на тестовому обладнанні, створення резервних копій поточної конфігурації та мікропрограми, планування робіт в період мінімального навантаження мережі та підготовку процедури швидкого відкату у разі виникнення проблем. Документування топології мережі, конфігурацій пристроїв та схем IP-адресації спрощує діагностику проблем, планування змін та передачу знань новим членам команди адміністрування. Впровадження системи керування змінами забезпечує контроль модифікацій інфраструктури, дозволяє відстежувати історію конфігурації та швидко відкочувати невдалі зміни.

Фізична безпека мережевого обладнання передбачає розміщення комутаторів, маршрутизаторів та серверів у спеціалізованих серверних приміщеннях з контрольованим доступом, системами охолодження та безперебійного живлення. Організація резервного електропостачання через джерела безперебійного живлення UPS та дизель-генератори забезпечує

безперервну роботу критичної інфраструктури під час відключень електромережі. Системи пожежогасіння в серверних приміщеннях повинні використовувати безпечні для обладнання засоби, такі як газове або аерозольне гасіння, замість традиційних водяних спринклерів. Контроль параметрів середовища, таких як температура, вологість та запиленість, запобігає передчасному виходу з ладу електронних компонентів та забезпечує стабільну роботу мережевої інфраструктури протягом тривалого періоду експлуатації.

Планування життєвого циклу мережевого обладнання включає оцінку моменту, коли пристрої повинні бути замінені через моральне застарівання, вичерпання можливостей модернізації або закінчення терміну підтримки виробником. Типовий термін служби комерційного мережевого обладнання становить п'ять-сім років, після чого підтримка може бути припинена, а ймовірність апаратних відмов суттєво зростає. Поступова заміна обладнання дозволяє розподілити капітальні витрати в часі та мінімізувати ризики масових відмов застарілих пристроїв. Утилізація виведеного з експлуатації обладнання повинна враховувати вимоги екологічного законодавства та забезпечувати безпечне знищення даних, які можуть зберігатися в енергонезалежній пам'яті пристроїв для запобігання витоку конфіденційної інформації про конфігурацію та структуру мережі.

Тема 2. Еталонні моделі мережної взаємодії OSI та TCP/IP

Розвиток комп'ютерних мереж у другій половині XX століття супроводжувався появою численних несумісних технологій та протоколів від різних виробників, що створювало значні труднощі при побудові гетерогенних мережевих інфраструктур. Для вирішення проблеми сумісності та стандартизації мережевої взаємодії міжнародні організації зі стандартизації розробили еталонні моделі, які визначають функціональну декомпозицію мережевих протоколів на рівні архітектури. Найбільш відомими та широко використовуваними є семирівнева еталонна модель відкритих систем OSI, розроблена Міжнародною організацією зі стандартизації, та чотирирівнева модель TCP/IP, що виникла в результаті практичної реалізації мережі ARPANET. Розуміння цих моделей є фундаментальним для проектування, налаштування, діагностики та усунення несправностей у сучасних комп'ютерних мережах, оскільки вони забезпечують концептуальну основу для розуміння взаємодії між різними компонентами мережевої інфраструктури.

2.1 Семирівнева модель OSI та функції її рівнів

Еталонна модель взаємодії відкритих систем, відома як модель OSI (Open Systems Interconnection), була розроблена Міжнародною організацією зі стандартизації ISO на початку 1980-х років з метою створення universal framework для мережевих протоколів. Модель складається з семи ієрархічних рівнів, кожен з яких виконує чітко визначені функції та взаємодіє лише з сусідніми рівнями через стандартизовані інтерфейси. Основна ідея модульної архітектури полягає в тому, що кожен рівень надає послуги вищому рівню, приховуючи деталі реалізації, та використовує послуги нижчого рівня для виконання своїх функцій. Така структура дозволяє незалежно розробляти та модифікувати протоколи на різних рівнях без впливу на функціонування інших рівнів, що значно спрощує процес стандартизації та впровадження нових технологій. Модель OSI, хоча і не була повністю реалізована в чистому вигляді, стала концептуальною основою для розуміння мережевої архітектури та широко використовується в освітніх цілях і для систематизації знань про мережеві протоколи.

Фізичний рівень, найнижчий у моделі OSI, відповідає за передачу необроблених бітів даних через фізичне середовище зв'язку. Цей рівень визначає електричні, механічні, процедурні та функціональні специфікації для активації, підтримки та деактивації фізичних з'єднань між мережевими пристроями. До функцій фізичного рівня належить визначення характеристик фізичного середовища передачі, таких як тип кабелю, роз'єми, рівні напруги та форми сигналів, а також методи кодування бітів у фізичні сигнали. Канальний рівень забезпечує надійну передачу кадрів даних між двома безпосередньо з'єднаними вузлами мережі, виявляючи та виправляючи помилки, що можуть виникнути на фізичному рівні. Він відповідає за формування кадрів з бітового потоку, адресацію на рівні локальної мережі через MAC-адреси, контроль

доступу до середовища передачі та управління потоком даних для запобігання переповненню приймача.

Мережний рівень забезпечує логічну адресацію та маршрутизацію пакетів через складні міжмережеві структури, що складаються з множини проміжних вузлів. Основні функції цього рівня включають визначення оптимальних шляхів передачі даних від джерела до призначення, фрагментацію та збірку пакетів відповідно до максимального розміру передавальної одиниці різних мережевих сегментів, а також логічну адресацію, що дозволяє ідентифікувати пристрої незалежно від їх фізичного розташування в мережі. Транспортний рівень забезпечує наскрізну передачу даних між кінцевими вузлами, гарантуючи або надійну доставку з встановленням з'єднання, або швидку передачу без гарантій доставки залежно від вимог застосунку. Цей рівень відповідає за сегментацію великих блоків даних на менші сегменти, контроль помилок наскрізної передачі, управління потоком для узгодження швидкостей відправника та одержувача, а також мультиплексування декількох з'єднань від різних застосунків через єдине мережеве з'єднання.

Сеансовий рівень керує встановленням, підтримкою та завершенням сеансів зв'язку між застосунками, що виконуються на різних вузлах мережі. Він забезпечує синхронізацію діалогу між взаємодіючими процесами, управління чергуванням передачі даних у напівдуплексних з'єднаннях та встановлення контрольних точок для відновлення сеансу після збоїв. Рівень представлення відповідає за перетворення даних між форматами, які використовуються застосунками, та стандартизованим мережевим форматом, забезпечуючи незалежність від специфічних способів представлення даних в різних системах. До функцій цього рівня належить шифрування та дешифрування даних для забезпечення конфіденційності, стиснення та розпакування для ефективного використання пропускну здатності, а також перетворення символічних кодувань та форматів даних. Прикладний рівень, найвищий у моделі OSI, надає мережеві послуги безпосередньо кінцевим застосункам користувача, визначаючи протоколи для конкретних задач, таких як передача файлів, електронна пошта, віддалений доступ до ресурсів та веб-перегляд.

Для систематизації інформації про функції рівнів моделі OSI розглянемо таблицю 2.1, яка узагальнює основні характеристики кожного рівня.

Таблиця 2.1 - Функції та одиниці даних рівнів моделі OSI

Ріве нь	Назва	Одиниця даних	Основні функції
7	Прикладни й	Дані (Data)	Мережеві сервіси для застосунків
6	Представле ння	Дані (Data)	Перетворення, шифрування, стиснення
5	Сеансовий	Дані (Data)	Управління сеансами зв'язку
4	Транспортн ий	Сегмент (Segment)	Наскрізна передача, надійність
3	Мережний	Пакет (Packet)	Логічна адресація,

			маршрутизація
2	Канальний	Кадр (Frame)	Фізична адресація, доступ до середовища
1	Фізичний	Біти (Bits)	Передача бітів через фізичне середовище

Як видно з таблиці 2.1, кожен рівень моделі OSI оперує специфічною одиницею даних і виконує чітко визначений набір функцій, що забезпечує логічне розділення відповідальності між різними компонентами мережевого стеку. Верхні три рівні, часто називані прикладними рівнями, фокусуються на наданні послуг безпосередньо кінцевим застосункам та обробці даних у форматах, зрозумілих користувачеві. Транспортний рівень є проміжним, що забезпечує перехід від орієнтованих на застосунок функцій до функцій, орієнтованих на мережеву передачу. Нижні три рівні, або мережеві рівні, відповідають за фактичну передачу даних через мережеву інфраструктуру, включаючи маршрутизацію, контроль доступу до середовища та фізичну передачу сигналів.

2.2 Процеси інкапсуляції та декапсуляції даних

Фундаментальним принципом функціонування багаторівневих мережевих архітектур є процес інкапсуляції даних при передачі від вищих рівнів до нижчих та відповідний процес декапсуляції при прийманні даних і їх просуванні від нижчих рівнів до вищих. Інкапсуляція полягає в тому, що кожен рівень додає до блоку даних, отриманого від вищого рівня, власний заголовок, а в деяких випадках і кінцевик, що містять службову інформацію, необхідну для виконання функцій цього рівня. Таким чином, дані прикладного рівня послідовно "загортаються" в оболонки кожного наступного нижчого рівня, подібно до вкладених одна в одну матрьошок, де кожна оболонка несе специфічну для свого рівня керуючу інформацію. Процес декапсуляції на приймальній стороні є зворотнім: кожен рівень аналізує та обробляє заголовок свого рівня, потім видаляє його та передає блок даних вищому рівню для подальшої обробки. Ця модель забезпечує прозорість взаємодії між рівнями та дозволяє кожному рівню функціонувати незалежно, оперуючи лише тією інформацією, яка релевантна для його функцій.

На прикладному рівні користувацький застосунок генерує дані, які потрібно передати іншому застосунку в мережі. Ці дані можуть мати різну природу: текстовий документ, зображення, відеопотік, запит до бази даних або будь-яка інша інформація, що обробляється застосунком. Прикладний рівень може додати власні заголовки, специфічні для протоколу прикладного рівня, наприклад, HTTP-заголовки для веб-трафіку або SMTP-команди для електронної пошти. На рівнях представлення та сеансу дані можуть зазнавати перетворень, таких як шифрування або стиснення, і до них додаються відповідні службові поля, хоча в сучасних мережевих стеках функції цих рівнів часто інтегровані в інші рівні. Транспортний рівень отримує блок даних від

вищих рівнів, при необхідності сегментує його на менші частини та додає транспортний заголовок, що містить номери портів джерела та призначення для ідентифікації застосунків, а також інформацію про порядкові номери, контрольні суми та параметри керування потоком для протоколів з встановленням з'єднання. Результуючий блок даних називається сегментом для TCP або датаграмою для UDP.

Мережний рівень приймає сегмент від транспортного рівня та додає мережевий заголовок, що містить логічні адреси джерела та призначення, зазвичай IP-адреси, а також інформацію про час життя пакета, тип протоколу верхнього рівня, контрольну суму заголовка та інші поля, необхідні для маршрутизації пакета через міжмережу. Утворена структура називається пакетом або IP-датаграмою. Канальний рівень інкапсулює пакет у кадр, додаючи заголовок канального рівня з фізичними адресами джерела та призначення, зазвичай MAC-адресами для Ethernet, а також поле типу протоколу або довжини даних. До кінця кадру додається кінцевик, що містить послідовність перевірки кадру, яка використовується для виявлення помилок при передачі через фізичне середовище. Нарешті, фізичний рівень перетворює кадр у послідовність фізичних сигналів, придатних для передачі через конкретне фізичне середовище: електричні імпульси для мідних кабелів, світлові імпульси для оптичного волокна або радіохвилі для бездротових з'єднань.

Для ілюстрації процесу інкапсуляції розглянемо рисунок 2.1, який схематично відображає додавання заголовків на кожному рівні моделі OSI. На рисунку зображено вертикальний стек з семи рівнів моделі OSI, розташованих зверху вниз від прикладного до фізичного. Від кожного рівня відходить горизонтальний блок, що представляє структуру даних на цьому рівні. На прикладному рівні показано початковий блок даних користувача темно-синім кольором. На рівні представлення цей блок залишається незмінним або доповнюється незначною службовою інформацією, позначеною світло-синім відтінком. На сеансовому рівні аналогічно може додаватися мінімальна керуюча інформація. На транспортному рівні зліва від блоку даних додається транспортний заголовок, позначений зеленим кольором, і утворений блок підписаний як "Сегмент". На мережевому рівні до сегмента зліва додається мережевий заголовок жовтого кольору, і структура позначена як "Пакет". На канальному рівні до пакета додаються як заголовок зліва, так і кінцевик справа, обидва позначені помаранчевим кольором, і вся структура підписана як "Кадр". На фізичному рівні кадр представлений у вигляді послідовності бітів або хвильової форми сигналу сірого кольору. Стрілка вниз вздовж лівого краю рисунка вказує напрямком інкапсуляції від прикладного рівня до фізичного.

При прийманні даних на стороні отримувача відбувається зворотний процес декапсуляції, коли дані просуваються від фізичного рівня до прикладного, і на кожному рівні відповідний заголовок аналізується та видаляється. Фізичний рівень приймаючого пристрою перетворює фізичні сигнали на цифрові біти та передає їх канальному рівню. Канальний рівень перевіряє цілісність кадру за допомогою послідовності перевірки кадру,

виявляючи можливі помилки передачі, аналізує MAC-адресу призначення, щоб визначити, чи адресований кадр саме цьому пристрою, та після успішної перевірки видаляє заголовок і кінцевик канального рівня, передаючи пакет мережевому рівню. Мережевий рівень аналізує IP-адресу призначення, перевіряє контрольну суму заголовка та інші параметри, визначає протокол верхнього рівня з відповідного поля заголовка, видаляє мережевий заголовок і передає сегмент транспортному рівню відповідно до вказаного протоколу.

Транспортний рівень використовує номер порту призначення для ідентифікації цільового застосунку, перевіряє контрольну суму для виявлення помилок у даних, у випадку протоколів з встановленням з'єднання керує порядковими номерами для забезпечення правильної послідовності та повноти даних, відправляє підтвердження отримання та за необхідності запитує повторну передачу втрачених сегментів. Після обробки всіх службових функцій транспортний заголовок видаляється, і дані передаються вищим рівням. Сеансовий та рівень представлення, якщо їх функції реалізовані окремо, виконують відповідно управління сеансом і перетворення даних, такі як дешифрування або розпакування. Нарешті, прикладний рівень отримує дані в форматі, зрозумілому застосунку, та надає їх кінцевій програмі користувача для відображення, обробки або зберігання. Таким чином, завершується повний цикл передачі даних від застосунку відправника через всі рівні мережевого стеку вниз, через фізичне середовище, і назад вгору через рівні мережевого стеку отримувача до цільового застосунку.

2.3 Чотирирівнева модель TCP/IP та її співвідношення з моделлю OSI

Модель TCP/IP, також відома як модель Інтернету або модель DoD (Department of Defense), виникла в результаті практичної реалізації мережі ARPANET наприкінці 1960-х та на початку 1970-х років і стала фундаментом для сучасного Інтернету. На відміну від моделі OSI, яка була розроблена як теоретична основа до впровадження конкретних протоколів, модель TCP/IP еволюціонувала разом з розробкою та впровадженням реальних мережевих протоколів, що забезпечили міжмережеву взаємодію в гетерогенному середовищі. Модель TCP/IP складається з чотирьох рівнів: мережевого інтерфейсу (або канального), міжмережевого, транспортного та прикладного, кожен з яких об'єднує функції декількох рівнів моделі OSI. Така більш практична та менш формалізована структура відображає реальну архітектуру протоколів TCP/IP та спрощує розуміння їх взаємодії, хоча іноді за рахунок деталізації теоретичних аспектів. Модель TCP/IP зосереджена на забезпеченні гнучкої, масштабованої та надійної міжмережевої комунікації, що дозволило їй стати домінуючою архітектурою в глобальних та локальних мережах.

Рівень мережевого інтерфейсу, найнижчий у моделі TCP/IP, відповідає за фізичну передачу даних через конкретне мережеве середовище та об'єднує функції фізичного та канального рівнів моделі OSI. Цей рівень визначає методи доступу до фізичного середовища, формування кадрів, фізичну адресацію та інші аспекти, специфічні для конкретної технології локальної мережі, такої як

Ethernet, Token Ring, Wi-Fi або інші. Модель TCP/IP не встановлює жорстких вимог до технологій рівня мережевого інтерфейсу, що забезпечує її гнучкість та можливість роботи поверх різноманітних фізичних та каналних технологій. Міжмережевий рівень є центральним у моделі TCP/IP і відповідає мережевому рівню моделі OSI, забезпечуючи логічну адресацію, маршрутизацію пакетів через множинні мережі та фрагментацію пакетів при необхідності. Основним протоколом цього рівня є IP (Internet Protocol), який визначає формат IP-пакетів, схему IP-адресації та правила маршрутизації. Крім IP, на міжмережевому рівні функціонують допоміжні протоколи, такі як ICMP для діагностики мережі та повідомлення про помилки, ARP для відображення IP-адрес на MAC-адреси та IGMP для керування групами багатоадресної розсилки.

Транспортний рівень моделі TCP/IP відповідає транспортному рівню моделі OSI і забезпечує наскрізну передачу даних між застосунками на кінцевих вузлах мережі. На цьому рівні функціонують два основні протоколи з різними характеристиками: TCP (Transmission Control Protocol) та UDP (User Datagram Protocol). Протокол TCP є орієнтованим на з'єднання та забезпечує надійну, упорядковану доставку даних з механізмами керування потоком, контролю перевантажень та повторної передачі втрачених сегментів, що робить його придатним для застосунків, які вимагають гарантованої доставки, таких як передача файлів, електронна пошта та веб-перегляд. Протокол UDP є простішим, не встановлює з'єднання та не гарантує доставки або порядку пакетів, що забезпечує меншу затримку та overhead, роблячи його оптимальним для застосунків реального часу, таких як потокове відео, VoIP та онлайн-ігри, де своєчасність доставки важливіша за абсолютну надійність. Прикладний рівень моделі TCP/IP є найвищим і об'єднує функції прикладного, представлення та сеансового рівнів моделі OSI, надаючи інтерфейс між мережевими застосунками та транспортним рівнем.

На прикладному рівні моделі TCP/IP функціонує велика кількість протоколів, кожен з яких призначений для певного типу мережеских послуг. Серед найбільш поширених можна назвати HTTP та HTTPS для веб-сервісів, FTP та SFTP для передачі файлів, SMTP, POP3 та IMAP для електронної пошти, DNS для перетворення доменних імен на IP-адреси, DHCP для автоматичного призначення IP-адрес, Telnet та SSH для віддаленого доступу до командного рядка, SNMP для управління мережевими обладнаннями та багато інших. Ці протоколи визначають формати повідомлень, послідовність обміну повідомленнями та правила взаємодії між клієнтами та серверами для реалізації конкретних мережеских сервісів. Прикладний рівень також відповідає за представлення даних у форматах, зрозумілих застосункам, включаючи кодування символів, формати файлів та методи шифрування, що в моделі OSI відноситься до окремого рівня представлення.

Для порівняння структури моделей OSI та TCP/IP доцільно розглянути таблицю 2.2, яка демонструє співвідношення рівнів цих двох архітектурних підходів.

Таблиця 2.2 - Співвідношення рівнів моделей OSI та TCP/IP

Модель OSI	Модель TCP/IP	Приклади протоколів
Прикладний Представлення Сеансовий	Прикладний	HTTP, HTTPS, FTP, SMTP, DNS, DHCP, SSH, Telnet, SNMP
Транспортний	Транспортний	TCP, UDP
Мережний	Міжмережевий	IP, ICMP, ARP, IGMP
Канальний Фізичний	Мережевий інтерфейс	Ethernet, Wi-Fi, PPP, HDLC

З таблиці 2.2 видно, що модель TCP/IP має більш компактну структуру порівняно з моделлю OSI, об'єднуючи функції декількох рівнів OSI в одному рівні. Прикладний рівень TCP/IP інкапсулює функції трьох верхніх рівнів OSI, що спрощує архітектуру, але може ускладнювати аналіз специфічних функцій представлення даних або керування сеансами. Транспортний та міжмережевий рівні в обох моделях відповідають один одному майже повністю, оскільки ці рівні виконують критичні функції наскрізної передачі та міжмережевої маршрутизації, які є фундаментальними для будь-якої мережевої архітектури. Рівень мережевого інтерфейсу TCP/IP об'єднує два нижні рівні OSI, залишаючи деталі фізичної та каналної передачі поза межами основної специфікації моделі, що відображає філософію незалежності від конкретних технологій локальних мереж.

2.4 Порівняльний аналіз моделей OSI та TCP/IP

Модель OSI та модель TCP/IP, будучи двома найбільш впливовими підходами до структурування мережевих протоколів, мають як спільні риси, так і суттєві відмінності, що відображають різні філософії їх розробки та призначення. Обидві моделі базуються на принципі ієрархічної організації функцій у вигляді рівнів або шарів, де кожен рівень надає послуги вищому рівню та використовує послуги нижчого рівня, забезпечуючи модульність та незалежність компонентів. Обидві моделі використовують концепцію інкапсуляції даних, при якій кожен рівень додає власні керуючі поля до даних, отриманих від вищого рівня, та видаляє їх при просуванні даних вгору по стеку на приймальній стороні. Обидві моделі визначають транспортний та мережний рівні з аналогічними функціями, що підкреслює універсальність цих функцій для будь-якої мережевої архітектури. Проте відмінності між моделями є не менш значущими та впливають з різних контекстів їх створення.

Модель OSI була розроблена як теоретична основа до впровадження конкретних протоколів, з наміром створити всеосяжний стандарт, що

забезпечить сумісність між різними системами від різних виробників. Її семирівнева структура забезпечує детальну декомпозицію мережевих функцій, чітко розділяючи відповідальність між рівнями та полегшуючи розуміння складних мережевих процесів у навчальних цілях. Модель OSI визначає окремі рівні для представлення даних та управління сеансами, що теоретично дозволяє більш точно специфікувати ці функції, хоча на практиці багато сучасних протоколів інтегрують ці функції в інші рівні. Однак складність та формалізованість моделі OSI, а також відсутність на момент її розробки зрілих та широко впроваджених протоколів, що точно відповідають її структурі, призвели до того, що вона не стала домінуючою в практичних реалізаціях, залишившись переважно освітнім та концептуальним інструментом.

Модель TCP/IP, навпаки, виникла з практичної необхідності забезпечити міжмережеву взаємодію в гетерогенному середовищі та еволюціонувала разом з розробкою та впровадженням реальних протоколів, таких як IP, TCP та UDP, що стали основою Інтернету. Її чотирирівнева структура є більш гнучкою та прагматичною, об'єднуючи функції декількох рівнів OSI там, де це не заважає практичній реалізації. Модель TCP/IP не встановлює жорстких вимог до технологій нижніх рівнів, дозволяючи працювати поверх різноманітних фізичних та каналних технологій, що забезпечує її універсальність та масштабованість. Протоколи стеку TCP/IP були широко впроваджені та стандартизовані, що призвело до їх домінування в сучасних мережах, включаючи глобальний Інтернет та більшість корпоративних та локальних мереж. Модель TCP/IP зосереджена на забезпеченні міжмережевої сумісності та надійної передачі даних, а не на абстрактній повноті опису всіх можливих мережевих функцій.

З точки зору практичного застосування та діагностики мережевих процесів, обидві моделі надають корисні інструменти для систематизації знань та локалізації проблем. Модель OSI з її детальною структурою дозволяє більш точно ідентифікувати, на якому саме рівні виникла проблема, що може бути корисним при навчанні та при аналізі складних мережевих конфігурацій. Наприклад, розуміння різниці між функціями каналного та мережного рівнів допомагає чітко розмежувати проблеми локальної мережі, такі як конфлікти MAC-адрес або збої комутації, від проблем міжмережевої маршрутизації або конфігурації IP-адресації. Модель TCP/IP, маючи менше рівнів, може здаватися менш деталізованою для аналітичних цілей, але її тісна відповідність реальним протоколам робить її більш безпосередньо застосовною при налаштуванні та діагностиці конкретних мережевих систем. Розуміння того, як працюють IP, TCP та UDP у контексті моделі TCP/IP, безпосередньо допомагає в налаштуванні маршрутизаторів, міжмережевих екранів та мережевих застосунків.

У сучасній практиці мережеві фахівці зазвичай використовують обидві моделі залежно від контексту: модель OSI для структурованого розуміння та навчання основ мережевої взаємодії, та модель TCP/IP для практичної роботи з реальними протоколами та обладнанням. Наприклад, при діагностиці проблем зв'язку можна починати аналіз з найнижчого рівня моделі OSI, перевіряючи

фізичне з'єднання та статус каналу, потім переходити до перевірки IP-адресації та маршрутизації на мережному рівні, далі до перевірки доступності портів та стану TCP-з'єднань на транспортному рівні, і нарешті до аналізу роботи конкретних застосунків на прикладному рівні. Така методика, що поєднує елементи обох моделей, дозволяє систематично та ефективно локалізувати та усувати мережеві проблеми, використовуючи концептуальну чіткість моделі OSI та практичну релевантність моделі TCP/IP.

Важливо також розуміти, що жодна з моделей не є абсолютно точним відображенням усіх можливих мережевих архітектур та протоколів. Реальні протоколи можуть охоплювати функції декількох рівнів або реалізовувати функції в іншій послідовності, ніж передбачено моделями. Наприклад, протокол ARP, який відображає IP-адреси на MAC-адреси, функціонує на межі мережного та каналного рівнів і не вписується чітко в один рівень жодної з моделей. Протоколи шифрування, такі як SSL/TLS, можуть розглядатися як функціонуючі на рівні представлення в моделі OSI, але в моделі TCP/IP вони зазвичай розташовуються між транспортним та прикладним рівнями. MPLS, технологія комутації пакетів з міткою, працює на рівні, який часто називають рівнем 2.5, оскільки він виконує функції, проміжні між каналним та мережним рівнями. Ці приклади показують, що еталонні моделі є корисними концептуальними інструментами, але не повинні розглядатися як жорсткі рамки, які обмежують розуміння реальних мережевих технологій.

Розглянемо рисунок 2.2, який ілюструє процес передачі даних через моделі OSI та TCP/IP паралельно, показуючи інкапсуляцію на стороні відправника та декапсуляцію на стороні отримувача. Рисунок складається з трьох вертикальних секцій: ліва секція показує стек моделі OSI на пристрої відправника, центральна секція представляє фізичне середовище передачі у вигляді горизонтальної лінії з хвильовою формою сигналу, права секція показує стек моделі OSI на пристрої отримувача. У лівій секції рівні моделі OSI розташовані зверху вниз: прикладний, представлення, сеансовий, транспортний, мережний, каналний, фізичний. Поряд з кожним рівнем зображено блок даних з відповідними заголовками: на прикладному рівні показано блок користувацьких даних, на транспортному рівні до даних додається заголовок TCP або UDP, на мережному рівні додається IP-заголовок, на каналному рівні додаються Ethernet-заголовок та кінцевик. Стрілки вказують напрямок інкапсуляції вниз по стеку. У правій секції процес відбувається у зворотному порядку: дані просуваються вгору по стеку, і на кожному рівні відповідний заголовок видаляється. Під стеками OSI показано відповідні рівні моделі TCP/IP з позначенням, які рівні OSI відповідають кожному рівню TCP/IP: рівень мережевого інтерфейсу TCP/IP охоплює фізичний та каналний рівні OSI, міжмережний рівень відповідає мережному рівню OSI, транспортний рівень співпадає в обох моделях, прикладний рівень TCP/IP об'єднує три верхні рівні OSI.

Розуміння еталонних моделей OSI та TCP/IP є фундаментальним для роботи в галузі комп'ютерних мереж, оскільки вони надають концептуальну основу для вивчення протоколів, проектування мережевих архітектур,

налаштування мережевого обладнання та діагностики проблем. Модель OSI залишається важливим освітнім інструментом, що допомагає структурувати розуміння складних мережевих процесів, тоді як модель TCP/IP відображає реальну архітектуру протоколів, що домінують у сучасних мережах. Комбіноване використання обох моделей дозволяє мережевим фахівцям ефективно аналізувати, проектувати та підтримувати комп'ютерні мережі різного масштабу та складності, від невеликих локальних мереж до глобальних розподілених систем. Принципи ієрархічної організації, модульності, інкапсуляції та чіткого розділення функцій між рівнями, закладені в цих моделях, продовжують визначати розвиток мережевих технологій і залишаються актуальними навіть у контексті сучасних тенденцій, таких як програмно-конфігуровані мережі, віртуалізація мережевих функцій та хмарні обчислення.

Тема 3. Середовища передачі, технології та пристрої фізичного і каналного рівнів

3.1 Середовища передачі даних фізичного рівня

Фізичний рівень моделі OSI відповідає за передачу бітів даних через фізичне середовище, перетворюючи цифрові сигнали у відповідні електричні, оптичні або радіоімпульси. Вибір середовища передачі є критично важливим рішенням при проектуванні мережі, оскільки воно визначає такі ключові параметри як максимальна пропускна здатність, довжина сегмента без проміжного підсилення сигналу, стійкість до електромагнітних перешкод та вартість розгортання інфраструктури. Кожне середовище передачі має свої унікальні характеристики, переваги та обмеження, що робить його більш або менш придатним для конкретних умов експлуатації.

Вита пара представляє собою один з найпоширеніших типів кабельних середовищ у локальних обчислювальних мережах завдяки оптимальному співвідношенню вартості, продуктивності та простоти монтажу. Конструктивно кабель складається з чотирьох пар мідних провідників, скручених між собою за певним кроком, що значно зменшує електромагнітні наведення та перехресні перешкоди між сусідніми парами. Неекранована вита пара (UTP - Unshielded Twisted Pair) не має додаткового екранування і використовується переважно в офісних приміщеннях, де рівень зовнішніх електромагнітних перешкод є відносно низьким. Екранована вита пара (STP - Shielded Twisted Pair) містить металеву оплетку або фольгу навколо всіх пар або кожної пари окремо, що забезпечує кращий захист від зовнішніх електромагнітних впливів у промислових умовах або поблизу потужного електрообладнання.

Категорії витой пари стандартизовані міжнародними організаціями та визначають максимальну пропускну здатність і частотний діапазон роботи кабелю. Категорія 5e (Enhanced Category 5) підтримує швидкість передачі до 1 Гбіт/с на частоті до 100 МГц і довжину сегмента до 100 метрів, що робить її мінімальним стандартом для сучасних мережевих інсталяцій. Категорія 6 розширює частотний діапазон до 250 МГц та забезпечує стабільну роботу на швидкостях 1-10 Гбіт/с при довжині сегмента до 55 метрів для 10-гігабітного Ethernet. Категорія 6a (Augmented Category 6) підтримує 10 Гбіт/с на повній відстані 100 метрів при частоті до 500 МГц завдяки покращеній конструкції та додатковим заходам захисту від перешкод. Для високопродуктивних центрів обробки даних використовуються категорії 7 та 8, які підтримують швидкості до 40 Гбіт/с на частотах до 2000 МГц, але мають суттєво вищу вартість та складність монтажу.

Коаксіальний кабель історично був одним з перших середовищ передачі в комп'ютерних мережах, але в сучасних LAN його значною мірою витіснила вита пара через нижчу вартість та простішу інсталяцію останньої. Конструктивно коаксіальний кабель складається з центрального мідного провідника, діелектричного ізолятора, металевій оплетки та зовнішньої захисної оболонки, що забезпечує відмінний захист від електромагнітних

перешкод. Завдяки цій структурі коаксіальний кабель може передавати сигнали на значно більші відстані без проміжного підсилення порівняно з витю парою, підтримуючи довжину сегмента до 500 метрів для товстого коаксіалу (10Base5) та до 185 метрів для тонкого (10Base2). В сучасних мережах коаксіальний кабель знаходить застосування переважно в кабельному телебаченні, системах відеоспостереження та в деяких спеціалізованих промислових інсталяціях, де потрібна висока стійкість до перешкод.

3.2 Оптичне волокно та бездротові технології передачі

Оптоволоконний кабель являє собою найсучасніше та найпродуктивніше середовище передачі даних, яке використовує світлові імпульси замість електричних сигналів для транспортування інформації. Основою оптоволокна є тонке скляне або пластикове волокно діаметром від 8 до 62.5 мікрометрів, покрите оболонкою з матеріалу з нижчим коефіцієнтом заломлення, що забезпечує повне внутрішнє відбиття світла та його поширення вздовж волокна. Оптичне волокно повністю імунне до електромагнітних перешкод, не створює власного електромагнітного випромінювання, забезпечує набагато більшу пропускну здатність та довжину сегмента порівняно з мідними кабелями, а також практично неможливе для несанкціонованого прослуховування без фізичного доступу до кабелю.

Одномодове волокно (Single-Mode Fiber, SMF) має дуже малий діаметр серцевини близько 8-10 мікрометрів, що дозволяє поширюватися лише одному променю світла вздовж осі волокна без відбиттів від стінок. Такий принцип роботи мінімізує дисперсію сигналу та забезпечує можливість передачі даних на відстані до 40-80 кілометрів без проміжного підсилення при швидкостях 10-100 Гбіт/с та вище. Одномодове волокно використовується переважно для магістральних з'єднань між будівлями, в мережах операторів зв'язку та в центрах обробки даних для міжкластерних з'єднань. Багатомодове волокно (Multi-Mode Fiber, MMF) має більший діаметр серцевини 50 або 62.5 мікрометрів, що дозволяє світловим променям поширюватися під різними кутами та відбиватися від стінок волокна. Ця властивість призводить до модальної дисперсії та обмежує максимальну довжину сегмента до 300-550 метрів залежно від типу волокна та швидкості передачі, але робить багатомодове волокно значно дешевшим через використання менш точних світлодіодних джерел замість лазерних.

Основні характеристики різних середовищ передачі наведено в таблиці 3.1, яка дозволяє порівняти їх ключові параметри та визначити оптимальний вибір для конкретних умов розгортання мережі. Як видно з таблиці, оптичне волокно забезпечує найвищу пропускну здатність та найбільшу довжину сегмента, але має вищу вартість монтажу та вимагає спеціалізованого обладнання для з'єднання.

Таблиця 3.1 - Порівняльні характеристики середовищ передачі даних

Середовище	Макс. пропускна здатність	Макс. довжина сегмента	Стійкість до перешкод	Відносна вартість
UTP Cat 5e	1 Гбіт/с	100 м	Низька	Низька
UTP Cat 6	10 Гбіт/с (55 м)	100 м (1 Гбіт/с)	Середня	Низька
UTP Cat 6a	10 Гбіт/с	100 м	Середня	Середня
STP	10 Гбіт/с	100 м	Висока	Середня
Коаксіальний	10 Мбіт/с	185-500 м	Висока	Середня
MMF 50/125	100 Гбіт/с	300-550 м	Дуже висока	Висока
SMF 9/125	100+ Гбіт/с	40-80 км	Дуже висока	Дуже висока

Бездротові технології передачі даних використовують радіохвилі або інфрачервоне випромінювання для з'єднання мережевих пристроїв без фізичних кабелів, що забезпечує мобільність користувачів та спрощує розгортання мережі в складних архітектурних умовах. Стандарти IEEE 802.11 (Wi-Fi) визначають різні частотні діапазони, методи модуляції та швидкості передачі для бездротових локальних мереж. Wi-Fi 4 (802.11n) працює в діапазонах 2.4 ГГц та 5 ГГц, підтримує швидкості до 600 Мбіт/с та використовує технологію MIMO (Multiple Input Multiple Output) для підвищення пропускної здатності. Wi-Fi 5 (802.11ac) працює виключно в діапазоні 5 ГГц, досягає швидкостей до 6.9 Гбіт/с при використанні широких каналів та багатопотокової передачі, але має менший радіус покриття через вищі частоти. Wi-Fi 6 (802.11ax) підтримує обидва діапазони, забезпечує швидкості до 9.6 Гбіт/с та впроваджує технології OFDMA та MU-MIMO для ефективнішої роботи в щільних мережах з великою кількістю пристроїв.

Бездротові мережі мають специфічні обмеження та виклики, які відрізняють їх від провідних середовищ передачі. Радіосигнал послаблюється при проходженні через стіни, меблі та інші перешкоди, що суттєво зменшує ефективний радіус покриття та швидкість передачі на віддалених від точки доступу ділянках. Діапазон 2.4 ГГц забезпечує кращу проникність через перешкоди та більший радіус покриття, але страждає від перевантаженості через велику кількість пристроїв та інтерференції з мікрохвильовими печами, бездротовими телефонами та Bluetooth-пристроями. Діапазон 5 ГГц має більше неперехресних каналів та менше перешкод, але гірше проникає через стіни та має менший радіус покриття. Безпека бездротових мереж вимагає обов'язкового використання шифрування WPA2 або WPA3, оскільки

радіосигнал може бути перехоплений будь-яким пристроєм в радіусі покриття без фізичного доступу до кабельної інфраструктури.

3.3 Технологія Ethernet та протоколи канального рівня

Ethernet є домінуючою технологією канального рівня для побудови локальних обчислювальних мереж, стандартизованою IEEE під номером 802.3 та постійно еволюціонуючою для підтримки все вищих швидкостей передачі. Історично Ethernet використовував метод доступу до середовища CSMA/CD (Carrier Sense Multiple Access with Collision Detection), який дозволяв множинним пристроям спільно використовувати одне фізичне середовище шляхом прослуховування каналу перед передачею та виявлення колізій при одночасній передачі. Сучасні Ethernet-мережі побудовані на комутаторах, які забезпечують виділені повнодуплексні з'єднання між портами, усуваючи проблему колізій та необхідність в механізмі CSMA/CD. Різні варіанти Ethernet стандартизують специфічні комбінації швидкості передачі та типу фізичного середовища: 10Base-T використовує виту пару категорії 3 для 10 Мбіт/с, 100Base-TX працює на категорії 5 зі швидкістю 100 Мбіт/с, 1000Base-T забезпечує гігабітну швидкість на категорії 5e, а 10GBase-T підтримує 10 Гбіт/с на категорії 6a.

Структура Ethernet-кадру визначає формат, в якому дані інкапсулюються на канальному рівні для передачі через мережу, та складається з кількох обов'язкових полів з фіксованими розмірами. Кадр починається з преамбули довжиною 7 байтів та delimiter-байту, які використовуються для синхронізації приймача з потоком бітів відправника та не вважаються частиною самого кадру. Поле призначення (Destination MAC Address) довжиною 6 байтів містить апаратну адресу пристрою-одержувача, визначаючи кому призначений даний кадр в межах локального сегмента мережі. Поле джерела (Source MAC Address) також займає 6 байтів та ідентифікує відправника кадру, дозволяючи одержувачу надіслати відповідь або підтвердження отримання. Поле Type/Length розміром 2 байти може вказувати або тип протоколу верхнього рівня (наприклад, 0x0800 для IPv4, 0x86DD для IPv6), або довжину поля даних залежно від використовуваного формату кадру. Поле даних (Payload) має змінну довжину від 46 до 1500 байтів та містить інкапсульовану інформацію з вищих рівнів мережевого стеку, причому мінімальний розмір забезпечується доповненням нулями при необхідності. Завершує кадр поле контрольної суми FCS (Frame Check Sequence) довжиною 4 байти, яке містить результат CRC32-обчислення для виявлення пошкоджень при передачі.

MAC-адреса (Media Access Control address) є унікальним 48-бітним ідентифікатором мережевого інтерфейсу, який призначається виробником обладнання та зберігається в енергонезалежній пам'яті пристрою. Структурно MAC-адреса складається з двох частин: перші 24 біти (3 байти) формують OUI (Organizationally Unique Identifier), який ідентифікує виробника та призначається IEEE, а наступні 24 біти визначають серійний номер конкретного пристрою в межах продукції цього виробника. MAC-адреси записуються в шістнадцятковому форматі та можуть бути представлені в різних нотаціях: через дефіси (01-23-45-67-89-AB), через двокрапки (01:23:45:67:89:AB) або через крапки після кожних двох байтів (0123.4567.89AB). Перший байт MAC-

адреси містить два спеціальні біти: I/G (Individual/Group) біт вказує чи є адреса унікальною (0) або груповою/broadcast (1), а U/L (Universal/Local) біт визначає чи адреса призначена глобально виробником (0) або локально адміністратором (1).

Типова структура Ethernet-кадру з детальним описом призначення кожного поля представлена на рисунку 3.1, який демонструє послідовність та розмір всіх компонентів кадру. Рисунок показує горизонтальну схему кадру зліва направо, де кожне поле представлено прямокутником з підписаною назвою зверху та розміром в байтах знизу. Преамбула та SFD займають загалом 8 байтів на початку, після чого йдуть Destination MAC (6 байтів) та Source MAC (6 байтів), виділені яскравішим кольором для акцентування важливості адресних полів. Далі розташоване поле Type/Length розміром 2 байти, за яким слідує найбільше змінне поле Data від 46 до 1500 байтів, показане розширеним прямокутником зі стрілками, що вказують на можливість варіації розміру. Завершує структуру поле FCS розміром 4 байти, виділене окремим кольором для позначення контрольної функції. Під схемою кадру стрілками вказані групування полів: MAC Header включає всі поля від Destination до Type/Length, а Ethernet Frame охоплює всі поля від Destination до FCS.

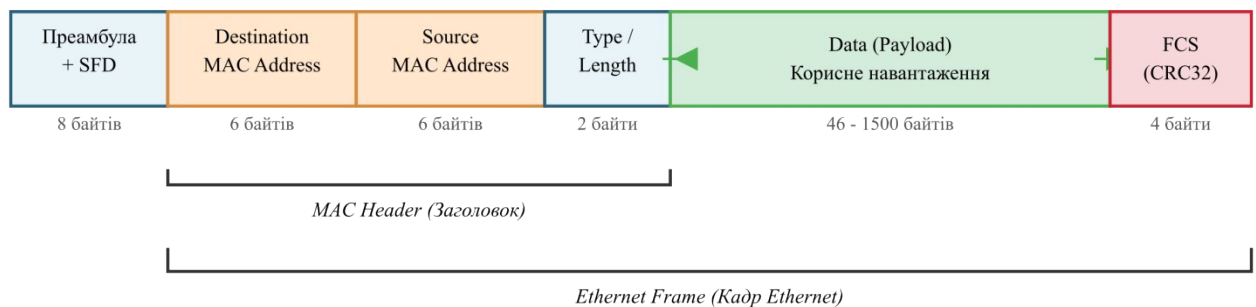


Рисунок 3.1 - Структура Ethernet-кадру стандарту IEEE 802.3

Протокол Spanning Tree Protocol (STP) стандартизований як IEEE 802.1D та призначений для запобігання формуванню петель в комутованих мережах з надлишковими з'єднаннями. Коли в мережевій топології існують множинні шляхи між комутаторами для забезпечення відмовостійкості, без додаткових механізмів координації broadcast-кадри можуть циркулювати в мережі нескінченно, спричиняючи broadcast-шторм та повну недоступність мережі. STP автоматично виявляє топологію мережі, обирає один корневий комутатор (Root Bridge) на основі найменшого Bridge ID, та обчислює найкоротші шляхи від кожного комутатора до кореня. Порти комутаторів переводяться в різні стани: forwarding для активної передачі трафіку, blocking для запобігання петлям, learning для вивчення MAC-адрес без пересилання кадрів, та listening для участі в обчисленні топології. При виході з ладу активного з'єднання STP автоматично переводить заблоковані порти в активний стан, відновлюючи зв'язність мережі, хоча традиційний STP потребує 30-50 секунд для конвергенції, що може бути неприйнятним для критичних додатків.

Удосконалені версії протоколу Spanning Tree вирішують проблему повільної конвергенції та додають нові функціональні можливості для складних корпоративних мереж. Rapid Spanning Tree Protocol (RSTP, IEEE 802.1w) зменшує час конвергенції до 1-2 секунд через використання нових типів портів edge та point-to-point, а також через активну синхронізацію станів між сусідніми комутаторами замість пасивного очікування таймерів. Multiple Spanning Tree Protocol (MSTP, IEEE 802.1s) дозволяє створювати окремі екземпляри spanning tree для різних груп VLAN, забезпечуючи балансування навантаження між надлишковими каналами та більш ефективне використання мережевої інфраструктури. Протоколи виявлення сусідів CDP (Cisco Discovery Protocol) та LLDP (Link Layer Discovery Protocol) працюють на каналному рівні та дозволяють мережевим пристроям автоматично виявляти безпосередньо підключених сусідів, обмінюючись інформацією про модель обладнання, версію програмного забезпечення, можливості портів та налаштування VLAN.

3.4 Комутатори та механізми їх функціонування

Комутатор (switch) є центральним пристроєм каналного рівня в сучасних Ethernet-мережах, який забезпечує інтелектуальну пересилку кадрів між підключеними портами на основі MAC-адрес призначення. На відміну від концентратора (hub), який просто повторює всі отримані біти на всі інші порти створюючи єдиний домен колізій, комутатор аналізує адресну інформацію в кадрах та пересилає їх лише на той порт, до якого підключений пристрій з відповідною MAC-адресою. Така селективна пересилка створює окремі домени колізій для кожного порту, усуває конкуренцію за доступ до середовища, дозволяє одночасну передачу даних між різними парами портів та суттєво підвищує загальну пропускну здатність мережі. Комутатори підтримують повнодуплексний режим роботи, коли порт може одночасно приймати та передавати дані, фактично подвоюючи ефективну пропускну здатність порту порівняно з напівдуплексним режимом, який вимагав механізму CSMA/CD.

Комутатори класифікуються за рівнем OSI, на якому вони приймають рішення про пересилку, що визначає їх функціональність та застосування в мережевій інфраструктурі. Комутатори другого рівня (L2 switches) приймають рішення виключно на основі MAC-адрес та працюють в межах одного широкомовного домену без можливості маршрутизації між різними IP-мережами. Такі комутатори є основою для побудови локальних сегментів мережі, підтримують VLAN для логічної сегментації, та забезпечують високу швидкість комутації завдяки простоті алгоритмів пересилки. Комутатори третього рівня (L3 switches) поєднують функції комутації та маршрутизації, здатні приймати рішення на основі IP-адрес призначення та виконувати міжмережеву маршрутизацію на апаратному рівні з високою продуктивністю. L3 комутатори використовуються для побудови ядра корпоративної мережі, де потрібна швидка маршрутизація між множинними VLAN та IP-підмережами, забезпечуючи при цьому латентність набагато нижчу ніж у традиційних програмних маршрутизаторів. Деякі високопродуктивні комутатори

підтримують функції четвертого рівня (L4), аналізуючи номери TCP/UDP портів для балансування навантаження та класифікації трафіку додатків.

Таблиця комутації або САМ-таблиця (Content Addressable Memory) є основою роботи комутатора та містить відповідність між MAC-адресами пристроїв та номерами портів, до яких ці пристрої підключені. САМ-таблиця реалізована на спеціалізованій асоціативній пам'яті, яка дозволяє швидкий паралельний пошук за адресою призначення в апаратному забезпеченні комутатора, забезпечуючи wire-speed пересилку без затримок на обробку. Процес формування САМ-таблиці відбувається автоматично та динамічно через аналіз адрес джерела у вхідних кадрах: коли комутатор отримує кадр на певному порті, він додає або оновлює запис, що пов'язує MAC-адресу джерела з номером цього порту та поточним часом. Кожен запис в САМ-таблиці має таймер життя (зазвичай 300 секунд або 5 хвилин), який скидається при кожному отриманні кадру з відповідною адресою джерела, а записи, які не оновлювалися протягом таймауту, автоматично видаляються для звільнення пам'яті.

Алгоритм пересилки кадрів комутатором включає кілька послідовних кроків прийняття рішення на основі інформації в САМ-таблиці та типу адреси призначення. Коли кадр надходить на порт комутатора, спочатку адреса джерела використовується для навчання: якщо ця адреса відсутня в таблиці, створюється новий запис, а якщо присутня на іншому порті, запис оновлюється. Після навчання комутатор аналізує адресу призначення кадру: якщо це unicast-адреса, яка присутня в САМ-таблиці, кадр пересилається лише на відповідний порт; якщо unicast-адреса відсутня в таблиці (unknown unicast), кадр розсилається на всі порти крім вхідного, як broadcast; якщо це broadcast або multicast адреса, кадр також пересилається на всі порти крім вхідного. Деякі комутатори підтримують різні режими пересилки: store-and-forward приймає весь кадр, перевіряє FCS та потім пересилає, забезпечуючи найкращу фільтрацію помилок але додаючи латентність; cut-through починає пересилку одразу після прочитання адреси призначення, мінімізуючи затримку але пропускаючи пошкоджені кадри; fragment-free читає перші 64 байти перед пересилкою, усуваючи фрагменти колізій але дозволяючи інші помилки.

Технологія Virtual LAN (VLAN) дозволяє логічно сегментувати фізичну мережу на множинні ізольовані широкомовні домени без необхідності фізичного розділення обладнання. VLAN створюють окремі логічні мережі на одному комутаторі або наборі комутаторів, де пристрої в різних VLAN не можуть безпосередньо обмінюватися кадрами навіть якщо підключені до одного комутатора, що покращує безпеку та продуктивність мережі. Кожен VLAN ідентифікується унікальним номером від 1 до 4094 згідно зі стандартом IEEE 802.1Q, який визначає механізм маркування кадрів 4-байтовим VLAN-тегом між полями Source MAC та Type/Length. Порти комутатора конфігуруються в одному з двох режимів: access-порти належать одному VLAN та обслуговують кінцеві пристрої, автоматично додаючи VLAN-тег до вихідних кадрів та видаляючи його з вхідних; trunk-порти передають трафік множинних

VLAN між комутаторами, зберігаючи теги для правильної ідентифікації належності кадрів.

Механізми забезпечення якості обслуговування (Quality of Service, QoS) на комутаторах другого рівня використовують поле Priority Code Point (PCP) в 802.1Q тезі для класифікації та пріоритезації трафіку. Поле PCP займає 3 біти та дозволяє визначити 8 рівнів пріоритету від 0 (найнижчий) до 7 (найвищий), які можуть використовуватися для надання переважної обробки критичному трафіку типу VoIP або відеоконференцій. Комутатори створюють окремі апаратні черги для кадрів різних пріоритетів та використовують алгоритми планування типу Strict Priority або Weighted Round Robin для обслуговування черг в порядку пріоритету. Стандарт IEEE 802.1p визначає рекомендовані відповідності між значеннями пріоритету та типами трафіку: 0 для best-effort даних, 1 для фонових трафіку, 2 для стандартних даних, 3 для важливих бізнес-додатків, 4 для потокового відео, 5 для інтерактивного голосу, 6 для критичного управлінського трафіку, 7 для мережевого контролю. Механізми обмеження швидкості (rate limiting) та формування трафіку (traffic shaping) дозволяють контролювати пропускну здатність для окремих портів, VLAN або класів трафіку, запобігаючи монополізації ресурсів та забезпечуючи справедливий розподіл смуги пропускання.

Основні параметри конфігурації комутатора та їх вплив на функціонування мережі систематизовані в таблиці 3.2, яка служить довідником для мережевих адміністраторів при плануванні налаштувань обладнання. Таблиця демонструє взаємозв'язок між різними налаштуваннями та їх вплив на безпеку, продуктивність та відмовостійкість мережевої інфраструктури.

Таблиця 3.2 - Ключові параметри конфігурації комутаторів

Параметр	Призначення	Типові значення	Вплив на мережу
CAM таймер старіння	Час зберігання MAC-записів	300 секунд	Менше - більше навантаження на CPU, більше - застарілі записи
Розмір CAM-таблиці	Кількість MAC-адрес	8K-128K записів	Обмежує кількість пристроїв в мережі
Режим STP	Версія протоколу запобігання петель	STP/RST P/MSTP	Швидкість конвергенції при збоях
Native VLAN	VLAN для нетегованих кадрів	VLAN 1	Потенційна вразливість безпеки
Port Security	Обмеження MAC-адрес на порті	1 - 132 адреси	Захист від MAC flooding атак
Flow Control	Управління перевантаженням	Enabled/Disabled	Запобігає втратам пакетів при перевантаженні
Jumbo Frames	Максимальний розмір кадру	1500-9000 байтів	Зменшення overhead, підвищення throughput
Storm Control	Обмеження broadcast/multicast	Threshold в %	Захист від broadcast-штормів

Безпека комутаторів включає широкий спектр механізмів захисту від несанкціонованого доступу та атак на каналному рівні, які є критичними для забезпечення цілісності мережевої інфраструктури. Port Security обмежує кількість MAC-адрес, які можуть навчатися на конкретному порті, та визначає дії при порушенні цього обмеження: режим protect тихо відкидає кадри з невідомих адрес, режим restrict генерує SNMP-повідомлення та збільшує лічильник порушень, а режим shutdown автоматично переводить порт в стан err-disabled при виявленні порушення. DHCP Snooping створює базу даних довірених відповідей між MAC-адресами, IP-адресами, портами та VLAN на основі спостереження за DHCP-трафіком, що дозволяє запобігти атакам типу DHCP spoofing та підміни IP-адрес. Dynamic ARP Inspection (DAI) використовує DHCP Snooping базу для валідації ARP-пакетів та відкидання підроблених ARP-відповідей, які могли б перенаправити трафік на зловмисника через ARP poisoning атаку. IP Source Guard перевіряє відповідність IP та MAC адрес джерела в кадрах та відкидає пакети, які не відповідають записам в DHCP Snooping базі, запобігаючи підміні IP-адрес на рівні доступу.

Атака MAC flooding є поширеною загрозою безпеці комутуваних мереж, коли зловмисник надсилає велику кількість кадрів з випадковими адресами джерела для переповнення CAM-таблиці комутатора. Після переповнення таблиці комутатор переходить в режим hub-like поведінки та починає розсилати всі кадри на всі порти, оскільки не може знайти адресу призначення в заповненій таблиці, що дозволяє зловмиснику перехоплювати чужий трафік. Захист від MAC flooding реалізується через Port Security з обмеженням максимальної кількості MAC-адрес на порті до розумного значення (зазвичай 1-3 для портів кінцевих пристроїв), а також через моніторинг швидкості навчання нових адрес та автоматичне блокування портів при аномальній активності. VLAN hopping атака експлуатує неправильну конфігурацію trunk-портів для отримання доступу до трафіку інших VLAN: в атаці double-tagging зловмисник надсилає кадр з двома VLAN-тегами, де перший тег відповідає native VLAN, що призводить до видалення першого тегу та пересилання кадру в цільовий VLAN за другим тегом. Захист від VLAN hopping включає зміну native VLAN на невикористовуване значення, явне присвоєння всіх портів до конкретних VLAN, вимкнення DTP (Dynamic Trunking Protocol) на access-портах та використання VLAN Access Control Lists для фільтрації між-VLAN трафіку.

Моніторинг та діагностика роботи комутатора є невід'ємною частиною підтримки працездатності мережі та швидкого виявлення проблем до того, як вони вплинуть на користувачів. Протокол SNMP (Simple Network Management Protocol) дозволяє централізовано збирати статистику використання портів, стан STP-топології, завантаженість CPU та пам'яті, температуру обладнання та інші метрики для аналізу трендів та виявлення аномалій. Syslog-повідомлення надсилаються комутатором на централізований сервер при важливих подіях типу зміни стану порту, порушення безпеки, зміни STP-топології або переповнення таблиць, що забезпечує аудит-трейл для розслідування інцидентів. SPAN (Switched Port Analyzer) або port mirroring дозволяє копіювати весь трафік з одного або кількох портів на спеціальний порт

моніторингу, до якого підключений аналізатор трафіку або система виявлення вторгнень для детального інспектування пакетів. NetFlow або sFlow експортують узагальнену інформацію про потоки трафіку (пари адрес джерела/призначення, порти, протоколи) на колектор для аналізу патернів використання мережі, виявлення аномалій та планування ємності.

Високодоступні конфігурації комутаторів використовують протоколи резервування для забезпечення безперервної роботи мережі при відмові обладнання або каналів зв'язку. EtherChannel або Link Aggregation (IEEE 802.3ad LACP) об'єднує множинні фізичні канали між двома пристроями в один логічний канал, що збільшує пропускну здатність та забезпечує автоматичне перемикавання на резервні канали при збої. Протоколи узгодження типу LACP (Link Aggregation Control Protocol) або PAgP (Port Aggregation Protocol від Cisco) автоматично виявляють та активують члени групи агрегації, моніторять їх стан та динамічно додають або видаляють канали при зміні конфігурації. Virtual Switching System (VSS) або StackWise дозволяють об'єднати два або більше фізичних комутаторів в один логічний пристрій з єдиною площиною управління, що спрощує конфігурацію, усуває необхідність в STP для міжкомутаторних з'єднань та забезпечує швидке перемикавання при відмові одного з членів стека. Multi-Chassis Link Aggregation (MLAG або vPC) розширює концепцію EtherChannel на два окремих комутаторів, дозволяючи кінцевому пристрою використовувати активні з'єднання до обох комутаторів одночасно для максимальної відмовостійкості та балансування навантаження.

Продуктивність комутатора характеризується кількома ключовими метриками, які визначають його придатність для конкретних сценаріїв застосування в мережевій інфраструктурі. Throughput або пропускну здатність комутатора вимірюється в пакетах за секунду (pps) або в бітах за секунду (bps) та показує максимальну кількість даних, яку комутатор може обробити без втрат, причому для досягнення wire-speed на всіх портах одночасно потрібна комутаційна матриця (switching fabric) з пропускну здатністю рівною подвоєній сумі швидкостей всіх портів через повнодуплексну роботу. Латентність визначає затримку між отриманням першого біту кадру на вхідному порті та початком передачі на вихідному порті, типово складаючи 5-50 мікросекунд для L2 комутаторів та 50-500 мікросекунд для L3 комутаторів через додаткову обробку заголовків IP. Розмір буферів пакетів впливає на здатність комутатора справлятися з короточасними сплесками трафіку без втрат: малі буфери (кілька мегабайт) достатні для рівномірного трафіку, але для з'єднань на великі відстані або трафіку з високою варіативністю потрібні більші буфери (десятки мегабайт) для поглинання пульсацій. Packet buffer architecture може бути спільною (shared) коли всі порти використовують загальний пул пам'яті з динамічним розподілом, або виділеною (dedicated) коли кожен порт має фіксований буфер, причому спільна архітектура забезпечує кращу утилізацію ресурсів але потенційно дозволяє одному порту монополізувати весь буфер.

Енергоефективні технології в сучасних комутаторах знижують споживання електроенергії та експлуатаційні витрати через інтелектуальне

управління живленням портів та компонентів. Energy Efficient Ethernet (IEEE 802.3az) переводить порти в режим низького споживання під час пауз в передачі даних, коли немає активного трафіку, зменшуючи споживання на 50-80% при типових офісних навантаженнях з періодичною активністю. Auto-power down автоматично відключає живлення портів, до яких не підключені пристрої або з'єднання знаходиться в стані link down, економлячи кілька ват на порт та суттєво зменшуючи загальне споживання в мережах з неповним заповненням портів. Power over Ethernet (PoE) стандартизований як IEEE 802.3af (PoE), 802.3at (PoE+) та 802.3bt (PoE++) дозволяє передавати електроживлення разом з даними по витій парі, забезпечуючи до 15.4, 30 або 60-100 Вт на порт відповідно для живлення IP-телефонів, бездротових точок доступу, камер спостереження та інших пристроїв без окремих кабелів живлення. Комутатори з підтримкою PoE мають суттєво більший бюджет живлення та вимагають відповідного планування електропостачання, оскільки повністю завантажений 48-портовий PoE+ комутатор може споживати понад 1500 Вт при одночасному живленні всіх портів максимальною потужністю.

Майбутні тенденції розвитку технологій комутації включають підвищення швидкостей портів, впровадження програмно-конфігурованих мереж та інтеграцію функцій безпеки та аналітики безпосередньо в комутаційну інфраструктуру. Стандарти 25/50/100/200/400 Gigabit Ethernet поступово впроваджуються в центрах обробки даних для забезпечення зростаючих вимог до пропускної здатності від хмарних сервісів, віртуалізації та додатків з штучним інтелектом. Software-Defined Networking (SDN) відокремлює площину управління від площини пересилки даних, дозволяючи централізовано програмувати поведінку комутаторів через відкриті API типу OpenFlow для динамічної адаптації до змінних вимог додатків. Білі коробкові комутатори (white-box switches) використовують стандартне апаратне забезпечення з відкритими операційними системами типу SONiC або Cumulus Linux, знижуючи вартість та прив'язку до конкретного постачальника, але вимагають більшої експертизи для інтеграції та підтримки. Інтегрована безпека включає вбудовані можливості виявлення та запобігання вторгненням, аналіз зашифрованого трафіку через метадані без розшифрування, та автоматичну сегментацію на основі ідентифікації користувачів та пристроїв для реалізації концепції zero-trust мережевої архітектури.

Вибір комутатора для конкретного проекту вимагає ретельного аналізу поточних та майбутніх вимог мережі з урахуванням множини технічних та економічних факторів. Кількість та типи портів визначаються числом підключених пристроїв, їх потребами в пропускній здатності та очікуваним зростанням: зазвичай рекомендується планувати 20-30% резервних портів для майбутнього розширення та тимчасових підключень. Вимоги до функціональності третього рівня залежать від необхідності міжVLAN маршрутизації та складності мережевої топології: невеликі мережі можуть обійтися L2 комутаторами з централізованим маршрутизатором, тоді як великі розподілені мережі потребують L3 комутаторів для локальної маршрутизації та оптимізації трафіку. Бюджет PoE має покривати сумарну потужність всіх

пристроїв, які отримуватимуть живлення через Ethernet, з 20% запасом для стабільності та можливості апгрейду, враховуючи що реальне споживання пристроїв часто менше номінального максимуму. Рівень шуму та габарити важливі для розміщення в офісних приміщеннях, де безвентиляторні або тихі моделі можуть бути критичними для комфорту працівників, на відміну від центрів обробки даних де пріоритет має максимальна продуктивність та щільність портів. Підтримка стандартів та сумісність з існуючим обладнанням, особливо для функцій типу агрегації каналів або стекування, може обмежити вибір конкретними виробниками або моделями для забезпечення безпроблемної інтеграції в наявну інфраструктуру.

Концепція ієрархічного дизайну мережі використовує комутатори різних класів та можливостей на трьох логічних рівнях для забезпечення масштабованості, продуктивності та відмовостійкості корпоративної інфраструктури. Рівень доступу (Access Layer) складається з комутаторів, до яких безпосередньо підключаються кінцеві пристрої користувачів, і характеризується високою щільністю портів 1/10 Gigabit Ethernet, підтримкою PoE, Port Security та базовими функціями QoS для класифікації трафіку. Рівень розподілу (Distribution Layer) агрегує з'єднання від множинних комутаторів доступу, виконує міжVLAN маршрутизацію, застосовує політики безпеки та QoS, та забезпечує надмірність через подвійні uplink з'єднання від комутаторів доступу, вимагаючи L3 комутаторів з високою продуктивністю та можливостями policy-based routing. Ядро мережі (Core Layer) забезпечує швидкісний транспорт між рівнями розподілу різних секцій мережі або локацій, оптимізується виключно для максимальної пропускної здатності та мінімальної латентності без додаткових функцій обробки трафіку, використовуючи найшвидші комутатори з портами 40/100 Gigabit та резервними шляхами для відмовостійкості. В невеликих мережах рівень розподілу може бути об'єднаний з ядром в collapsed core архітектуру, що зменшує складність та вартість за рахунок дещо меншої масштабованості порівняно з повністю тришаровою моделлю.

Процедури введення в експлуатацію нового комутатора включають послідовність кроків конфігурації та тестування для забезпечення безпечної та надійної роботи в продуктивному середовищі. Базова конфігурація починається з призначення IP-адреси управління, hostname для ідентифікації пристрою, налаштування часової зони та синхронізації часу через NTP для точних міток в логах, та конфігурації доступу через SSH з вимкненням небезпечного Telnet. Налаштування VLAN включає створення необхідних VLAN, призначення портів до відповідних VLAN в access режимі або налаштування trunk для міжкомутаторних з'єднань з явним переліком дозволених VLAN замість "all" для безпеки. Конфігурація STP встановлює пріоритети для визначення root bridge, активує PortFast на access-портах для швидкого переходу в forwarding стан, та налаштовує BPDU Guard для захисту від підключення неавторизованих комутаторів. Безпека посилюється через Port Security з відповідними лімітами MAC-адрес, DHCP Snooping на untrusted портах, DAI та IP Source Guard де необхідно, а також через налаштування локального або RADIUS/TACACS+

автентифікації для адміністративного доступу. Тестування перевіряє зв'язність на всіх рівнях через ping між VLAN після налаштування маршрутизації, коректність STP-топології та відсутність петель, функціонування PoE на відповідних портах, та збір метрик продуктивності під типовим навантаженням для baseline майбутнього моніторингу.

Тема 4. Протоколи мережного рівня, маршрутизація та IP-адресація

4.1. Функції та призначення мережного рівня в моделі OSI

Мережний рівень посідає третю позицію в еталонній моделі OSI і виконує критично важливі функції для забезпечення комунікації між пристроями в різних мережах. Основне призначення цього рівня полягає в забезпеченні логічної адресації та визначенні оптимальних шляхів передачі даних між кінцевими вузлами незалежно від їх фізичного розташування. На відміну від канального рівня, який оперує фізичними MAC-адресами і працює лише в межах одного сегмента мережі, мережний рівень здатний забезпечувати міжмережеву комунікацію завдяки використанню ієрархічної системи логічних адрес. Протоколи мережного рівня відповідають за інкапсуляцію сегментів транспортного рівня в пакети, додаючи заголовки з інформацією про адреси відправника та одержувача, що дозволяє маршрутизаторам приймати рішення щодо пересилання трафіку.

Однією з ключових функцій мережного рівня є маршрутизація – процес визначення найкращого шляху для доставки пакетів від джерела до призначення через один або кілька проміжних маршрутизаторів. Маршрутизатори аналізують адресу призначення в заголовку пакета та консультуються зі своєю таблицею маршрутизації для прийняття рішення про наступний крок пересилання. Цей процес повторюється на кожному маршрутизаторі вздовж шляху, поки пакет не досягне мережі призначення. Ефективність маршрутизації залежить від актуальності інформації в таблицях маршрутизації, які можуть формуватися статично адміністратором або динамічно за допомогою протоколів маршрутизації. Важливо розуміти, що кожен маршрутизатор приймає незалежне рішення про пересилання на основі поточної інформації, що робить мережу гнучкою та адаптивною до змін топології.

Фрагментація пакетів є ще однією важливою функцією мережного рівня, яка вирішує проблему передачі даних через мережі з різними максимальними розмірами кадрів (MTU – Maximum Transmission Unit). Коли пакет, що прямує через мережу, виявляється занадто великим для передачі наступним сегментом мережі, маршрутизатор може розділити його на менші фрагменти. Кожен фрагмент отримує власний заголовок мережного рівня з інформацією про фрагментацію, що дозволяє пристрою призначення правильно зібрати оригінальний пакет. У протоколі IPv4 фрагментацію може виконувати будь-який маршрутизатор на шляху проходження пакета, тоді як у IPv6 фрагментацію виконує лише вузол-джерело, що спрощує обробку пакетів проміжними маршрутизаторами та підвищує продуктивність мережі.

Логічна адресація на мережному рівні забезпечує унікальну ідентифікацію кожного пристрою в глобальній мережі та містить інформацію про мережеву приналежність вузла. На відміну від плоских MAC-адрес, IP-адреси мають ієрархічну структуру, що дозволяє ефективно агрегувати маршрути та зменшувати розмір таблиць маршрутизації. Адреса складається з

двох частин: мережевої частини, яка ідентифікує конкретну мережу, та вузлової частини, яка ідентифікує окремий пристрій у цій мережі. Така структура дозволяє маршрутизаторам приймати рішення про пересилання на основі лише мережевої частини адреси, не аналізуючи вузлову частину, що значно спрощує та прискорює процес маршрутизації в великих мережах. Крім того, мережний рівень підтримує різні типи адресації, включаючи одноадресну (unicast), широкомовну (broadcast) та групову (multicast) передачі.

4.2. Протокол IPv4: структура, адресація та класифікація

Протокол IPv4 (Internet Protocol version 4) є основним протоколом мережного рівня, який забезпечує передачу пакетів у сучасних мережах TCP/IP. Розроблений у 1981 році, IPv4 використовує 32-бітні адреси, що теоретично дозволяє адресувати близько 4,3 мільярда унікальних пристроїв. Заголовок IPv4-пакета містить критично важливу інформацію для маршрутизації та обробки даних, включаючи версію протоколу, довжину заголовка, тип сервісу, загальну довжину пакета, ідентифікатор фрагмента, прапорці фрагментації, зміщення фрагмента, час життя (TTL), протокол вищого рівня, контрольну суму заголовка, адреси джерела та призначення. Мінімальний розмір заголовка IPv4 становить 20 байтів, але може збільшуватися до 60 байтів при використанні додаткових опцій. Поле TTL (Time To Live) відіграє особливу роль у запобіганні нескінченному циркулюванню пакетів у мережі: кожен маршрутизатор зменшує значення TTL на одиницю, і якщо воно досягає нуля, пакет відкидається.

Version 4 біти	IHL 4 біти	Type of Service (ToS) 8 бітів
Total Length 16 бітів		
Identification 16 бітів	Flags 3 біти	Fragment Offset 13 бітів
Time to Live 8 бітів	Protocol 8 бітів	Header Checksum 16 бітів
Source Address 32 біти (4 байти)		
Destination Address 32 біти (4 байти)		
Options (необов'язково) 0-320 бітів (0-40 байтів) Додаткові опції (макс. 60 байтів заголовка)		
Padding Доповнення до 32-бітного слова		

Умовні позначення	
	— Обов'язкові поля заголовка
	— Критичні поля маршрутизації
	— Необов'язкові поля
	— Дані корисного навантаження
Опис полів:	
• Version - версія протоколу IP (4 для IPv4)	
• IHL - довжина заголовка в 32-бітних словах	
• Type of Service - поле диференційованого сервісу	
• Total Length - загальна довжина пакета в байтах	
• Identification - унікальний ідентифікатор пакета	
• Flags/Fragment Offset - параметри фрагментації	
• TTL - лічильник "часу життя" пакета	
• Protocol - ідентифікатор протоколу вищого рівня	
• Header Checksum - контрольна сума для перевірки цілісності	
• Options - додаткові опції (необов'язково)	

Структура IPv4-адреси традиційно представляється у десятковому форматі з крапковим роздільником, де кожен з чотирьох октетів записується як десяткове число від 0 до 255, наприклад 192.168.1.1. Внутрішньо кожен октет представлений 8 бітами, що в сумі дає 32-бітну адресу. Для визначення меж між мережевою та вузловою частинами адреси використовується маска підмережі – 32-бітне число, в якому послідовні одиниці зліва позначають мережеву частину, а нулі справа – вузлову частину. Маска підмережі може записуватися як у повному форматі (наприклад, 255.255.255.0), так і в нотації

CIDR (Classless Inter-Domain Routing), де вказується кількість бітів мережевої частини після слешу (наприклад, /24). Застосування бітової операції логічного І між IP-адресою та маскою підмережі дозволяє визначити адресу мережі, до якої належить певний вузол, що є фундаментальною операцією для процесу маршрутизації.

Історично IPv4-адреси класифікувалися за класовою системою, яка визначала фіксовані межі між мережевою та вузловою частинами адреси. У таблиці 4.1 представлено основні характеристики класів IPv4-адрес.

Таблиця 4.1 – Класифікація IPv4-адрес за класовою системою

лас	Діапазон першого октета	Маска за замовчуванням	Кількість мереж	Кількість вузлів на мережу	Призначення
	1-126	255.0.0.0 (/8)	126	16 777 214	Великі мережі
	128-191	255.255.0.0 (/16)	16 384	65 534	Середні мережі
	192-223	255.255.255.0 (/24)	2 097 152	254	Малі мережі
	224-239	Не застосовується	Не застосовується	Не застосовується	Групова адресація
	240-255	Не застосовується	Не застосовується	Не застосовується	Зарезервовано

Класова система адресації, представлена в таблиці 4.1, виявилася неефективною через нераціональне використання адресного простору та жорсткість структури, що призвело до розробки безкласової адресації CIDR. При використанні класової адресації організація, якій потрібно підключити 300 вузлів, змушена була отримувати адресу класу В з можливістю адресації 65 534 вузлів, що призводило до марнування більше 65 тисяч адрес. CIDR дозволяє використовувати маски змінної довжини, що забезпечує більш гнучкий поділ адресного простору та можливість агрегації маршрутів. Наприклад, замість оголошення окремих маршрутів для чотирьох мереж класу С (192.168.0.0/24, 192.168.1.0/24, 192.168.2.0/24, 192.168.3.0/24), можна оголосити один агрегований маршрут 192.168.0.0/22, що значно зменшує розмір таблиць маршрутизації в магістральних маршрутизаторах Інтернету.

В IPv4 існують спеціальні типи адрес, які виконують особливі функції в мережевій комунікації. Адреса мережі визначається встановленням всіх бітів вузлової частини в нуль і використовується для ідентифікації самої мережі в таблицях маршрутизації. Широкомовна адреса мережі отримується встановленням всіх бітів вузлової частини в одиницю і використовується для одночасної передачі даних усім вузлам у конкретній мережі. Петлеві адреси (loopback), які належать до діапазону 127.0.0.0/8, використовуються для тестування мережевого стека на локальному пристрої без фактичної передачі даних в мережу, при цьому найчастіше використовується адреса 127.0.0.1. Приватні адреси, визначені в RFC 1918, включають діапазони 10.0.0.0/8,

172.16.0.0/12 та 192.168.0.0/16 і призначені для використання у внутрішніх мережах без прямого доступу до Інтернету. Адреса автоматичної конфігурації APIPA (Automatic Private IP Addressing) з діапазону 169.254.0.0/16 автоматично призначається пристрою, коли він не може отримати адресу від DHCP-сервера, що дозволяє здійснювати локальну комунікацію без участі адміністратора.

4.3. Протокол IPv6 та механізми переходу від IPv4

Вичерпання адресного простору IPv4 та зростаючі потреби в підключенні мільярдів пристроїв до Інтернету зумовили розробку нового протоколу IPv6 (Internet Protocol version 6). Основна перевага IPv6 полягає у використанні 128-бітних адрес, що забезпечує приблизно $3,4 \times 10^{38}$ унікальних адрес – кількість, достатню для призначення унікальних адрес кожному атому на поверхні Землі. IPv6-адреса записується у шістнадцятковому форматі як вісім груп по чотири шістнадцяткові цифри, розділених двокрапками, наприклад 2001:0db8:85a3:0000:0000:8a2e:0370:7334. Для спрощення запису існують правила скорочення: провідні нулі в кожній групі можуть бути опущені, а послідовність груп, що складаються виключно з нулів, може бути замінена подвійною двокрапкою, але лише один раз в адресі. Таким чином, адреса 2001:0db8:0000:0000:0000:0000:0000:0001 може бути скорочена до 2001:db8::1, що значно полегшує читання та запис адрес.

Структура заголовка IPv6 була суттєво спрощена порівняно з IPv4, що покращує продуктивність обробки пакетів маршрутизаторами. Основний заголовок IPv6 має фіксований розмір 40 байтів і містить поля версії, класу трафіку, мітки потоку, довжини корисного навантаження, наступного заголовка, ліміту переходів, адреси джерела та адреси призначення. На відміну від IPv4, у IPv6 відсутні поля контрольної суми заголовка та опцій, що зменшує обчислювальне навантаження на маршрутизатори. Функції, які раніше реалізовувалися через опції IPv4, в IPv6 винесені в окремі розширені заголовки, які розміщуються між основним заголовком та даними вищого рівня. Поле ліміту переходів (Hop Limit) виконує аналогічну функцію полю TTL в IPv4, запобігаючи нескінченному циркулюванню пакетів. Важливою особливістю є те, що фрагментація в IPv6 виконується виключно вузлом-джерелом за допомогою спеціального розширеного заголовка фрагментації, а проміжні маршрутизатори не мають права фрагментувати пакети.

В IPv6 визначено три основні типи адрес, кожен з яких має специфічне призначення та область застосування. Одноадресні (unicast) адреси ідентифікують окремий мережевий інтерфейс і забезпечують доставку пакета точно одному одержувачу. Групові (multicast) адреси ідентифікують групу інтерфейсів, і пакет, надісланий на таку адресу, доставляється всім членам групи, що є ефективним механізмом для розповсюдження інформації множині одержувачів одночасно. Адреси anycast призначаються групі інтерфейсів, але пакет доставляється лише найближчому члену групи згідно з метрикою маршрутизації, що використовується для забезпечення відмовостійкості та балансування навантаження. На відміну від IPv4, в IPv6 відсутня концепція

широкомовної (broadcast) адресації, її функції виконуються спеціальними групові адресами, що забезпечує більш контрольовану та ефективну доставку пакетів.

Version 4 біти	Traffic Class 8 бітів	Flow Label 20 бітів
Payload Length 16 бітів		
Next Header 8 бітів	Hop Limit 8 бітів	
Source Address 128 бітів (16 байтів) IPv6-адреса джерела		
Destination Address 128 бітів (16 байтів) IPv6-адреса призначення		

Поля заголовка IPv6

Основні поля:

- Version (4 біти) — версія протоколу IP (6)
- Traffic Class (8 бітів) — клас трафіку, QoS
- Flow Label (20 бітів) — ідентифікатор потоку
- Payload Length (16 бітів) — довжина корисного навантаження в байтах
- Next Header (8 бітів) — тип наступного заголовка
- Hop Limit (8 бітів) — ліміт стрибків (аналог TTL)
- Source Address (128 бітів) — адреса джерела
- Destination Address (128 бітів) — адреса призначення

Розширені заголовки:

- Hop-by-Hop Options
- Routing Header
- Fragment Header
- Authentication Header (AH)
- Encapsulating Security Payload (ESP)
- Destination Options

У таблиці 4.2 представлено основні типи одноадресних IPv6-адрес та їх характеристики.

Таблиця 4.2 – Типи одноадресних IPv6-адрес

Тип адреси	П рефікс	Обла сть дії	Призначення
Глобальна одноадресна	2 000::/3	Глоб альна	Маршрутизація в Інтернеті
Унікальна локальна	f c00::/7	Лока льна	Внутрішні мережі організацій
Локальна канальна	f e80::/10	Кана л	Автоконфігурація, взаємодія в сегменті
Петлева	:: 1/128	Вузо л	Тестування локального стека
Невизначена	:: /128	Неви значена	Ініціалізація інтерфейсу

Як видно з таблиці 4.2, глобальні одноадресні адреси призначені для використання в публічному Інтернеті та є функціональним еквівалентом публічних IPv4-адрес. Унікальні локальні адреси виконують роль, аналогічну приватним адресам IPv4, дозволяючи організаціям створювати внутрішні мережі з власним адресним простором. Локальні канальні адреси автоматично генеруються на кожному IPv6-інтерфейсі та використовуються для комунікації в межах одного сегмента мережі, зокрема для роботи протоколів виявлення сусідів (Neighbor Discovery Protocol). Ці адреси не маршрутизуються за межі локального каналу, що забезпечує їх ізоляцію та безпеку.

Перехід від IPv4 до IPv6 є складним процесом, який потребує тривалого періоду співіснування обох протоколів через величезну кількість існуючої IPv4-інфраструктури. Для забезпечення сумісності та поступової міграції розроблено кілька механізмів переходу. Подвійний стек (Dual Stack) передбачає одночасну підтримку обох протоколів на мережевих пристроях, що дозволяє їм комунікувати як з IPv4, так і з IPv6 вузлами залежно від можливостей одержувача. Тунелювання дозволяє інкапсулювати IPv6-пакети в IPv4-пакети для передачі через мережі, які не підтримують IPv6, при цьому існують різні технології тунелювання, такі як 6to4, ISATAP та Teredo. Трансляція адрес

(NAT64) забезпечує можливість взаємодії між вузлами з IPv6-адресами та серверами з IPv4-адресами через спеціальний шлюз, який виконує трансляцію пакетів та адрес між двома протоколами. Вибір конкретного механізму переходу залежить від топології мережі, наявного обладнання та стратегічних цілей організації щодо впровадження IPv6.

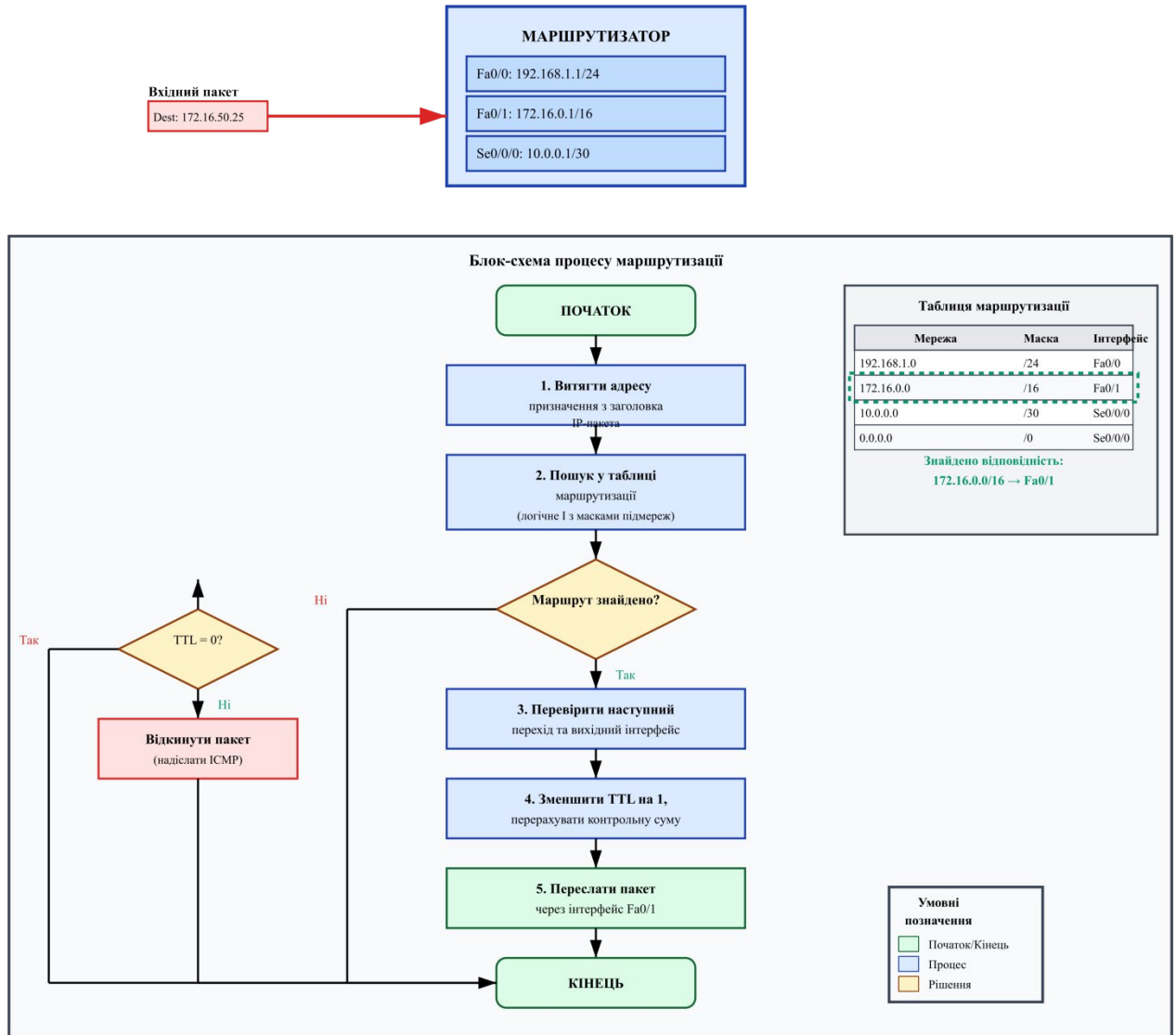
4.4. Принципи маршрутизації та таблиці маршрутизації

Маршрутизація є фундаментальним процесом мережного рівня, який визначає шлях проходження пакетів від джерела до призначення через взаємопов'язані мережі. Маршрутизатор виконує дві основні функції: визначення оптимального шляху для кожного пакета на основі адреси призначення та пересилання пакета на відповідний вихідний інтерфейс. Процес прийняття рішення про маршрутизацію базується на таблиці маршрутизації – структурованій базі даних, яка містить інформацію про доступні мережі та шляхи до них. Коли маршрутизатор отримує пакет, він витягує адресу призначення з заголовка пакета, виконує пошук найбільш специфічного відповідного маршруту в таблиці маршрутизації та перенаправляє пакет на інтерфейс, вказаний у знайденому маршруті. Якщо в таблиці відсутній маршрут до мережі призначення і не налаштовано маршрут за замовчуванням, пакет відкидається, а відправнику може бути надіслано повідомлення ICMP про недосяжність мережі.

Таблиця маршрутизації містить кілька критично важливих компонентів для кожного маршруту. Адреса мережі призначення та маска підмережі визначають діапазон IP-адрес, для яких застосовується цей маршрут. Адреса наступного переходу (next-hop) вказує IP-адресу сусіднього маршрутизатора, якому слід передати пакет для подальшого просування до мережі призначення. Вихідний інтерфейс ідентифікує фізичний або логічний інтерфейс маршрутизатора, через який пакет має бути відправлений. Метрика відображає “вартість” використання певного маршруту та застосовується для вибору найкращого шляху, коли до однієї мережі існує кілька альтернативних маршрутів. Адміністративна відстань (Administrative Distance) визначає пріоритет джерела інформації про маршрут, дозволяючи маршрутизатору вибирати між маршрутами, отриманими з різних джерел. Залежно від способу отримання інформації про мережі, маршрути класифікуються як безпосередньо підключені, статичні та динамічні, кожен тип має свої переваги та сфери застосування.

Рисунок 4.1 ілюструє процес прийняття рішення про маршрутизацію в типовому маршрутизаторі. На схемі зображено маршрутизатор з трьома інтерфейсами: Fa0/0 з адресою 192.168.1.1/24, Fa0/1 з адресою 172.16.0.1/16 та Se0/0/0 з адресою 10.0.0.1/30. До маршрутизатора надходить пакет з адресою призначення 172.16.50.25. Блок-схема демонструє послідовність операцій: спочатку маршрутизатор витягує адресу призначення з заголовка IP-пакета, потім здійснює пошук у таблиці маршрутизації, застосовуючи операцію логічного І між адресою призначення та масками підмережі для кожного

запису. Знаходиться відповідність з маршрутом 172.16.0.0/16, після чого маршрутизатор перевіряє наступний перехід та вихідний інтерфейс, зменшує TTL пакета на одиницю, перераховує контрольну суму заголовка та пересилає пакет через інтерфейс Fa0/1. Якщо TTL дорівнює нулю або відповідний маршрут відсутній, пакет відкидається з відправленням ICMP-повідомлення джерелу.



Безпосередньо підключені маршрути автоматично з'являються в таблиці маршрутизації, коли на інтерфейсі маршрутизатора налаштовується IP-адреса та інтерфейс переходить у робочий стан (up/up). Ці маршрути мають найвищий пріоритет (адміністративна відстань 0) та нульову метрику, оскільки мережа доступна безпосередньо без проміжних маршрутизаторів. Статичні маршрути налаштовуються адміністратором вручну за допомогою команд конфігурації та використовуються для забезпечення передбачуваного та контрольованого маршрутування трафіку. Перевагами статичних маршрутів є відсутність накладних витрат на обмін маршрутною інформацією, повний контроль адміністратора над шляхами проходження трафіку та можливість налаштування резервних маршрутів з різними метриками. Недоліками є необхідність ручного

оновлення при зміні топології мережі та складність підтримки в великих динамічних мережах. Статичні маршрути мають адміністративну відстань 1 за замовчуванням, що робить їх пріоритетнішими за динамічні маршрути. Особливим випадком статичного маршруту є маршрут за замовчуванням (default route) з адресою призначення 0.0.0.0/0, який застосовується для пересилання пакетів до всіх мереж, для яких немає більш специфічних маршрутів у таблиці.

Динамічні маршрути автоматично вивчаються та підтримуються за допомогою протоколів динамічної маршрутизації, які дозволяють маршрутизаторам обмінюватися інформацією про топологію мережі та автоматично адаптуватися до змін. Протоколи маршрутизації можуть бути класифіковані за різними критеріями. За областю застосування розрізняють протоколи внутрішньої маршрутизації (IGP – Interior Gateway Protocol), які працюють в межах однієї автономної системи (RIP, OSPF, EIGRP), та протоколи зовнішньої маршрутизації (EGP – Exterior Gateway Protocol), які забезпечують обмін маршрутною інформацією між автономними системами (BGP). За алгоритмом роботи протоколи поділяються на дистанційно-векторні, які приймають рішення на основі відстані до мережі призначення та напрямку пересилання (RIP, EIGRP), та протоколи стану каналу, які будують повну топологічну карту мережі та обчислюють найкоротші шляхи за алгоритмом Дейкстри (OSPF, IS-IS). Кожен протокол має свою метрику для оцінки якості маршрутів: RIP використовує кількість переходів, OSPF – вартість на основі пропускну здатності, EIGRP – комбіновану метрику з урахуванням пропускну здатності, затримки, навантаження та надійності.

Процес вибору найкращого маршруту при наявності кількох альтернатив відбувається за чітким алгоритмом. Спочатку маршрутизатор шукає найбільш специфічне відповідність, тобто маршрут з найдовшим префіксом (longest prefix match), що відповідає адресі призначення. Наприклад, якщо в таблиці присутні маршрути 172.16.0.0/16 та 172.16.50.0/24, для пакета з адресою призначення 172.16.50.25 буде обрано другий маршрут як більш специфічний. Якщо існує кілька маршрутів з однаковим префіксом, але з різними джерелами (наприклад, статичний та динамічний), перевага надається маршруту з найменшою адміністративною відстанню. Якщо маршрути мають однаковий префікс та адміністративну відстань, рішення приймається на основі метрики – обирається маршрут з найменшою метрикою. У випадку повної рівності всіх параметрів, маршрутизатор може здійснювати балансування навантаження, розподіляючи трафік між кількома рівноцінними шляхами, що підвищує пропускну здатність та відмовостійкість мережі.

Тема 5. Протоколи транспортного та прикладного рівнів

5.1. Призначення та функції транспортного рівня

Транспортний рівень займає четверту позицію в еталонній моделі OSI та забезпечує наскрізну передачу даних між додатками, що виконуються на різних вузлах мережі. Основним завданням цього рівня є створення логічних каналів зв'язку між процесами на кінцевих станціях, незалежно від фізичної структури мережі та кількості проміжних маршрутизаторів. Транспортний рівень відповідає за сегментацію даних, отриманих від рівня додатків, на частини прийнятної розміру для передачі через мережу, а також за зворотну операцію – збирання сегментів у вихідний потік даних на приймальній стороні. Важливою функцією є мультиплексування, яке дозволяє одночасно підтримувати декілька мережевих з'єднань від різних додатків на одному вузлі, використовуючи для їх розрізнення номери портів. Демультиплексування виконує зворотну задачу – розподіляє вхідні сегменти між відповідними додатками на основі номерів портів призначення.

Транспортний рівень може забезпечувати різні рівні якості обслуговування залежно від вимог додатків. Деякі протоколи цього рівня гарантують надійну доставку даних, встановлюючи логічне з'єднання перед передачею, контролюючи цілісність даних та здійснюючи повторну передачу втрачених сегментів. Інші протоколи орієнтовані на швидкість передачі та мінімальні затримки, відмовляючись від механізмів підтвердження та повторної передачі. Вибір конкретного протоколу транспортного рівня визначається характером додатку та його вимогами до швидкості, надійності та порядку доставки даних. Крім того, транспортний рівень може виконувати керування потоком даних, запобігаючи переповненню буферів приймача занадто інтенсивним потоком від відправника, а також керування перевантаженням мережі, адаптуючи швидкість передачі до поточного стану мережевої інфраструктури.

5.2. Протокол TCP: структура та механізми надійної передачі

Протокол TCP (Transmission Control Protocol) є одним з основних протоколів транспортного рівня та забезпечує надійну, орієнтовану на з'єднання передачу даних між додатками. Структура TCP-сегмента включає заголовок фіксованого розміру та поле даних змінної довжини. Заголовок TCP містить поля номера порту джерела та призначення, кожне розміром 16 біт, що дозволяє адресувати до 65536 різних портів на кожному вузлі. Послідовний номер (Sequence Number) розміром 32 біти ідентифікує позицію першого байта даних сегмента в загальному потоці, що передається від відправника до одержувача. Номер підтвердження (Acknowledgment Number) також займає 32 біти та вказує на номер наступного очікуваного байта, тим самим підтверджуючи отримання всіх попередніх байтів. Поле довжини заголовка

визначає розмір заголовка TCP у 32-бітних словах, що необхідно через змінну довжину поля опцій.

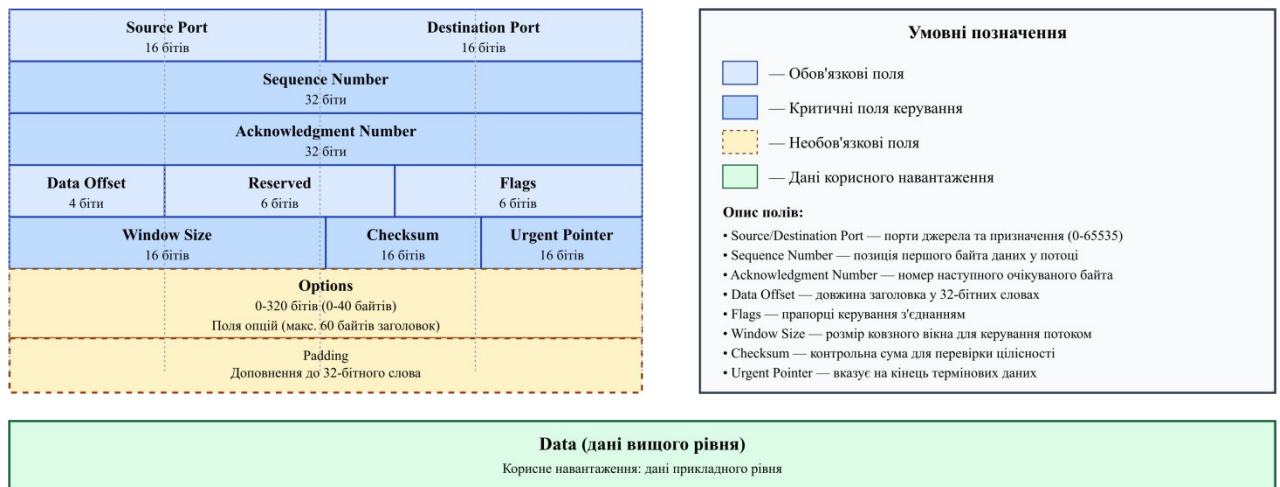


Рисунок 5.1 - Структура TCP-сегмента

Механізм встановлення з'єднання в TCP реалізується через процедуру триразового рукостискання (three-way handshake), яка гарантує синхронізацію обох сторін перед початком передачі даних. Ініціатор з'єднання відправляє сегмент з встановленим прапорцем SYN та випадковим початковим послідовним номером, інформуючи одержувача про намір встановити з'єднання. Одержувач відповідає сегментом з встановленими прапорцями SYN та ACK, підтверджуючи отримання запиту та повідомляючи власний початковий послідовний номер. На завершальному етапі ініціатор надсилає сегмент з прапорцем ACK, підтверджуючи отримання відповіді від одержувача, після чого з'єднання вважається встановленим та може використовуватися для передачі даних. Триразове рукостискання запобігає проблемам, що можуть виникнути через затримки в мережі або дублювання запитів на встановлення з'єднання, забезпечуючи коректну синхронізацію послідовних номерів на обох сторонах.

Керування потоком у TCP здійснюється за допомогою механізму ковзного вікна (sliding window), який дозволяє регулювати кількість даних, що можуть бути передані без отримання підтвердження. Розмір вікна вказується в заголовку TCP у полі Window Size та визначає кількість байтів, які одержувач готовий прийняти в свій буфер. Відправник не може передати більше даних, ніж дозволяє поточний розмір вікна, що запобігає переповненню буферів одержувача. Коли одержувач обробляє дані та звільняє місце в буфері, він збільшує розмір вікна в наступних сегментах підтвердження, дозволяючи відправнику надіслати додаткові дані. Механізм ковзного вікна є динамічним та адаптується до змін у швидкості обробки даних одержувачем, забезпечуючи ефективне використання мережевих ресурсів без ризику втрати даних через переповнення буферів.

Надійність передачі в TCP забезпечується через систему підтверджень та повторних передач, яка гарантує доставку всіх сегментів у правильному

порядку. Кожен отриманий сегмент підтверджується одержувачем шляхом відправлення сегмента з номером підтвердження, що вказує на наступний очікуваний байт. Якщо відправник не отримує підтвердження протягом певного часу (тайм-аут), він повторно передає сегмент, припускаючи, що оригінальний сегмент або підтвердження були втрачені в мережі. TCP використовує адаптивний алгоритм розрахунку тайм-ауту, який враховує поточний час кругової подорожі (RTT) та його варіацію, що дозволяє ефективно працювати в мережах з різними характеристиками затримки. Крім того, TCP підтримує вибірккові підтвердження (SACK), які дозволяють одержувачу повідомити відправнику про діапазони успішно отриманих байтів, навіть якщо деякі проміжні сегменти були втрачені, що значно підвищує ефективність повторних передач.

У таблиці 5.1 наведено основні прапорці TCP-заголовка та їх призначення в процесі управління з'єднанням та передачею даних.

Таблиця 5.1 - Прапорці TCP-заголовка та їх функції

Прапорець	Назва	Призначення
SYN	Synchronize	Ініціація встановлення з'єднання, синхронізація послідовних номерів
ACK	Acknowledgment	Підтвердження отримання даних, поле номера підтвердження є валідним
FIN	Finish	Ініціація коректного завершення з'єднання
RST	Reset	Негайне скидання з'єднання через помилку або неочікуваний сегмент
PSH	Push	Запит негайної передачі даних додатку без буферизації
URG	Urgent	Наявність терміново важливих даних, вказаних в полі Urgent Pointer

Завершення TCP-з'єднання також є контрольованим процесом, що включає обмін сегментами з прапорцем FIN. Сторона, яка ініціює закриття, відправляє сегмент з встановленим FIN, повідомляючи, що більше не буде передавати дані, але продовжує приймати дані від іншої сторони. Одержувач підтверджує отримання FIN сегментом з ACK та може продовжувати передавати власні дані. Коли одержувач також завершує передачу, він надсилає власний FIN, який підтверджується ініціатором закриття, після чого з'єднання повністю закривається. Цей процес чотириразового рукостискання забезпечує коректне завершення обох напрямків передачі даних та дозволяє кожній стороні упевнитися, що всі дані були успішно доставлені та оброблені перед остаточним закриттям з'єднання.

5.3. Протокол UDP та порівняння з TCP

Протокол UDP (User Datagram Protocol) є альтернативним протоколом транспортного рівня, який забезпечує простий механізм передачі даних без встановлення з'єднання та гарантій доставки. Структура UDP-датаграми є значно простішою порівняно з TCP-сегментом та включає заголовок фіксованого розміру лише 8 байтів. Заголовок містить поля номера порту джерела та призначення розміром по 16 біт кожне, поле довжини датаграми, що вказує загальний розмір заголовка та даних у байтах, та поле контрольної суми для перевірки цілісності заголовка та даних. Відсутність додаткових полів для послідовних номерів, підтверджень та керування потоком робить UDP надзвичайно ефективним з точки зору накладних витрат, дозволяючи передавати дані з мінімальною затримкою та використанням мережевих ресурсів.

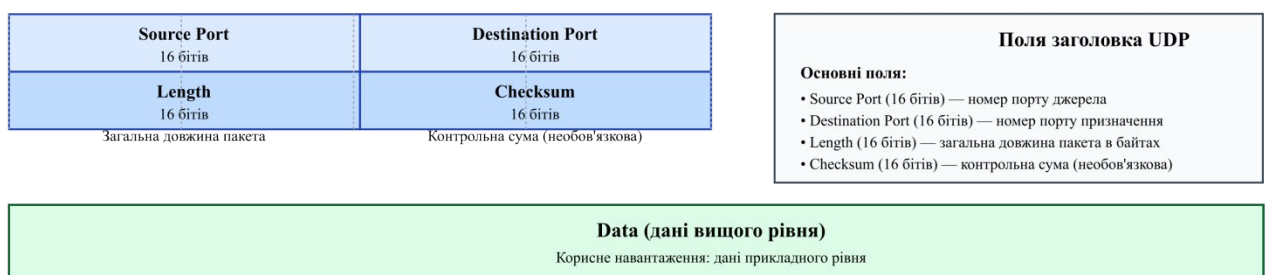


Рисунок 5.2 - UDP-датаграми

UDP не встановлює з'єднання перед передачею даних, що означає відсутність фази рукоштовування та можливість негайного відправлення датаграм одразу після їх формування додатком. Кожна UDP-датаграма обробляється незалежно від інших, навіть якщо вони належать до одного потоку даних від того самого додатку. Протокол не гарантує доставку датаграм до одержувача, не контролює їх порядок отримання та не виконує повторну передачу втрачених пакетів. Відповідальність за виявлення втрат, дублювання або неправильного порядку датаграм повністю покладається на додаток, який використовує UDP, що надає розробникам гнучкість у реалізації необхідних механізмів відповідно до специфічних вимог програми. Така простота робить UDP придатним для додатків, де швидкість передачі є критичною, а втрата окремих пакетів не призводить до катастрофічних наслідків для функціонування системи.

Основні відмінності між TCP та UDP визначають сфери їх застосування в різних типах мережевих додатків. TCP забезпечує надійну, послідовну доставку всіх переданих даних, що робить його ідеальним для додатків, де цілісність інформації є критичною, таких як передача файлів, електронна пошта, веб-браузер та віддалений доступ до систем. UDP, з іншого боку, мінімізує затримки та накладні витрати, що є перевагою для додатків реального часу, таких як потокове відео та аудіо, онлайн ігри, передача голосу через IP-мережі та системи відеоконференцій. У цих сценаріях короткочасна втрата кількох пакетів є прийнятною, оскільки повторна передача застарілих даних не має

сенсу для потокових медіа, а затримка від механізмів підтвердження та повторної передачі TCP могла б значно погіршити користувацький досвід.

Вибір між TCP та UDP також залежить від характеру взаємодії між клієнтом та сервером. Для запит-відповідь додатків з короткими повідомленнями, таких як DNS-запити, UDP може бути більш ефективним, оскільки уникнення накладних витрат на встановлення з'єднання скорочує загальний час обслуговування запиту. Однак для додатків з тривалими сесіями та великими обсягами даних TCP забезпечує кращу продуктивність завдяки механізмам керування потоком та перевантаженням, які оптимізують використання доступної пропускної здатності. Деякі сучасні додатки використовують гібридний підхід, застосовуючи UDP як базовий транспорт та реалізуючи власні механізми надійності на рівні додатку, що дозволяє точно налаштувати баланс між швидкістю та надійністю відповідно до специфічних вимог.

На рисунку 5.1 зображено порівняльну схему структури TCP-сегмента та UDP-датаграми. TCP-сегмент складається з заголовка розміром мінімум 20 байтів (може бути більшим при використанні опцій) та поля даних. Заголовок містить такі поля розташовані послідовно зверху вниз: порт джерела (16 біт), порт призначення (16 біт), послідовний номер (32 біти), номер підтвердження (32 біти), довжина заголовка та зарезервовані біти (8 біт), прапорці керування (8 біт включаючи URG, ACK, PSH, RST, SYN, FIN), розмір вікна (16 біт), контрольна сума (16 біт), вказівник терміновості (16 біт) та опції змінної довжини. Нижче заголовка розміщується поле даних змінного розміру. UDP-датаграма має значно простішу структуру з заголовком лише 8 байтів, що містить: порт джерела (16 біт), порт призначення (16 біт), довжину датаграми (16 біт) та контрольну суму (16 біт), після чого слідує поле даних. Візуально схема демонструє, що TCP-заголовок приблизно втричі більший за UDP-заголовок, що ілюструє різницю в складності та функціональності цих протоколів.

5.4. Мультиплексування, демюльтиплексування та протоколи прикладного рівня

Концепція мультиплексування на транспортному рівні дозволяє декільком додаткам на одному вузлі одночасно використовувати мережеві послуги, розрізняючи їх з'єднання за допомогою номерів портів. Номер порту є 16-бітним цілим числом, що забезпечує діапазон від 0 до 65535, який поділяється на три категорії: добре відомі порти (0-1023), зареєстровані порти (1024-49151) та динамічні або приватні порти (49152-65535). Добре відомі порти призначаються IANA (Internet Assigned Numbers Authority) для стандартних сервісів та протоколів, що забезпечує їх однакову інтерпретацію в усіх мережевих системах. Коли додаток відправляє дані, транспортний рівень додає до них заголовок з номером порту джерела, який ідентифікує додаток-відправник на локальному вузлі, та номером порту призначення, який вказує на додаток-одержувач на віддаленому вузлі.

Демультіплексування виконує зворотну функцію на приймальній стороні, розподіляючи вхідні сегменти між відповідними додатками на основі інформації в заголовках. Коли мережний рівень передає IP-пакет транспортному рівню, останній аналізує номер порту призначення в заголовку TCP або UDP та направляє дані до відповідного додатку або процесу, який очікує на з'єднання через цей порт. Для UDP демультіплексування є безз'єднанним процесом, що базується виключно на парі адреса IP призначення та порт призначення. TCP використовує чотирикомпонентний кортеж для ідентифікації з'єднання: адреса IP джерела, порт джерела, адреса IP призначення та порт призначення, що дозволяє підтримувати множинні з'єднання від різних клієнтів до одного серверного порту.

Протокол HTTP (Hypertext Transfer Protocol) є основним протоколом прикладного рівня для передачі гіпертекстових документів у Всесвітній павутині та використовує TCP-порт 80 для незахищеного з'єднання або порт 443 для HTTPS з шифруванням TLS/SSL. HTTP працює за моделлю запит-відповідь, де клієнт (зазвичай веб-браузер) ініціює з'єднання та надсилає запит на отримання ресурсу, а сервер обробляє запит та повертає відповідь з запитаним ресурсом або кодом стану, що вказує на результат обробки. Запит HTTP складається з рядка запиту, що містить метод (GET, POST, PUT, DELETE), URI ресурсу та версію протоколу, заголовків запиту з додатковою інформацією та необов'язкового тіла повідомлення для методів, що передають дані на сервер. Відповідь сервера включає рядок стану з кодом відповіді (200 OK, 404 Not Found, 500 Internal Server Error), заголовки відповіді та тіло з запитаним вмістом, таким як HTML-документ, зображення або дані в форматі JSON.

Протокол DNS (Domain Name System) забезпечує розподілену службу перетворення доменних імен у IP-адреси та працює переважно через UDP-порт 53, хоча може використовувати TCP для великих відповідей або передачі зон між серверами. DNS-запит містить ім'я домену, яке потрібно розв'язати, тип запитуваного ресурсного запису (A для IPv4-адреси, AAAA для IPv6, MX для поштових серверів, CNAME для канонічних імен) та клас запиту, зазвичай IN для інтернет-адрес. DNS-сервер обробляє запит, перевіряючи спочатку власний кеш, потім рекурсивно опитуючи інші DNS-сервери в ієрархії, починаючи від кореневих серверів через сервери доменів верхнього рівня до авторитативних серверів конкретного домену. Відповідь містить один або більше ресурсних записів з запитаною інформацією та значенням TTL (Time To Live), що вказує, скільки часу запис може зберігатися в кеші перед необхідністю повторного запиту, що зменшує навантаження на DNS-інфраструктуру та прискорює наступні запити до тих самих доменів.

DHCP (Dynamic Host Configuration Protocol) автоматизує процес призначення IP-адрес та інших мережних параметрів вузлам у локальній мережі, використовуючи UDP-порти 67 для сервера та 68 для клієнта. Процес отримання конфігурації DHCP складається з чотирьох етапів: виявлення (DHCP Discover), пропозиція (DHCP Offer), запит (DHCP Request) та підтвердження (DHCP Acknowledge). Клієнт без налаштованої адреси надсилає широкомовне

повідомлення DHCP Discover для пошуку доступних DHCP-серверів у мережі. Сервери відповідають повідомленням DHCP Offer, пропонуючи вільну IP-адресу з свого пулу разом з додатковими параметрами, такими як маска підмережі, адреса шлюзу за замовчуванням, адреси DNS-серверів та час оренди адреси. Клієнт вибирає одну з отриманих пропозицій та надсилає DHCP Request, інформуючи обраний сервер про прийняття його пропозиції та відхилення інших. Сервер завершує процес повідомленням DHCP Acknowledge, підтверджуючи призначення адреси та передаючи остаточну конфігурацію, після чого клієнт може використовувати отримані параметри для мережевої взаємодії протягом періоду оренди.

У таблиці 5.2 представлено основні протоколи прикладного рівня з вказівкою їх номерів портів та транспортних протоколів, які вони використовують.

Таблиця 5.2 - Основні протоколи прикладного рівня та їх характеристики

Протокол	Номер порту	Транспортний протокол	Призначення
HTTP	80	TCP	Передача гіпертекстових документів у веб
HTTPS	443	TCP	Захищена передача веб-контенту з шифруванням TLS/SSL
FTP (дані)	20	TCP	Передача файлів між клієнтом та сервером
FTP (керування)	21	TCP	Керуючі команди для FTP-сесії
SSH	22	TCP	Захищений віддалений доступ до командного рядка
Telnet	23	TCP	Незахищений віддалений доступ до командного рядка
SMTP	25	TCP	Відправлення електронної пошти між поштовими серверами
DNS	53	UDP/TCP	Розв'язання доменних імен у IP-адреси
DHCP (сервер)	67	UDP	Автоматичне призначення IP-конфігурації клієнтам
DHCP (клієнт)	68	UDP	Отримання IP-конфігурації від DHCP-сервера
TFTP	69	UDP	Проста передача файлів без автентифікації
POP3	110	TCP	Отримання електронної пошти з поштового сервера
IMAP	143	TCP	Управління електронною поштою на поштовому сервері

Протокол FTP (File Transfer Protocol) забезпечує передачу файлів між клієнтом та сервером через TCP-з'єднання, використовуючи дві окремі сесії: керуючу (порт 21) для передачі команд та відповідей, та сесію даних (порт 20 в активному режимі) для власне передачі вмісту файлів. FTP підтримує два

режими роботи: активний, де сервер ініціює з'єднання даних до клієнта, та пасивний, де клієнт ініціює обидва з'єднання, що важливо для роботи через міжмережеві екрани та механізми NAT. Протокол дозволяє автентифікацію користувачів за логіном та паролем, підтримує навігацію по файловій системі сервера, створення та видалення каталогів, а також передачу файлів у текстовому або бінарному режимі. SMTP (Simple Mail Transfer Protocol) використовує TCP-порт 25 для передачі електронної пошти між поштовими серверами та працює як простий текстовий протокол на основі команд та відповідей. Клієнт встановлює з'єднання з сервером, ідентифікує себе командою HELO або EHLO, вказує відправника командою MAIL FROM, одержувачів командою RCPT TO та передає вміст повідомлення після команди DATA, завершуючи рядком з одиничною крапкою, після чого сервер приймає повідомлення до черги доставки або повідомляє про помилку.

Розуміння взаємодії між транспортним та прикладним рівнями є критичним для проектування та діагностики мережевих додатків. Правильний вибір транспортного протоколу та налаштування його параметрів безпосередньо впливають на продуктивність та надійність додатків прикладного рівня. Адміністратори мереж повинні розуміти, які порти використовують різні сервіси, щоб коректно налаштовувати міжмережеві екрани, списки контролю доступу та механізми трансляції адрес. Знання структури протоколів та їх поведінки дозволяє ефективно використовувати засоби моніторингу та аналізу трафіку для виявлення проблем з продуктивністю, безпекою або неправильною конфігурацією мережевих компонентів, що є невід'ємною частиною професійної роботи з комп'ютерними мережами.

ТЕМА 6. МЕХАНІЗМИ БЕЗПЕКИ МЕРЕЖЕВОГО ОБЛАДНАННЯ

6.1. Захист доступу до інтерфейсу командного рядка мережевих пристроїв

Забезпечення безпеки мережевого обладнання починається з організації надійного захисту доступу до інтерфейсу командного рядка (CLI - Command Line Interface), через який здійснюється конфігурування маршрутизаторів, комутаторів та інших активних пристроїв мережі. Неавторизований доступ до CLI надає зловмиснику повний контроль над мережевим обладнанням, що може призвести до порушення конфіденційності, цілісності та доступності мережевих сервісів. Тому методи захисту доступу до командного рядка є фундаментальним елементом загальної стратегії мережевої безпеки та мають бути реалізовані на всіх критично важливих пристроях корпоративної інфраструктури.

Доступ до CLI мережевого обладнання може здійснюватися через різні канали зв'язку, кожен з яких потребує специфічних механізмів захисту. Консольний доступ передбачає фізичне підключення до пристрою через послідовний порт за допомогою консольного кабелю, що робить цей метод найбезпечнішим з точки зору захисту від віддалених атак, оскільки зловмисник повинен мати безпосередній фізичний доступ до обладнання. Проте навіть консольний доступ вимагає налаштування автентифікації, адже фізичний доступ до серверної кімнати або комунікаційної шафи може отримати не лише уповноважений персонал. Віддалений доступ до CLI здійснюється через протоколи Telnet або SSH (Secure Shell), що дозволяє адміністраторам керувати мережевим обладнанням з будь-якого місця в мережі, але водночас створює додаткові вектори атак.

Базовим методом захисту консольного доступу є встановлення пароля на консольну лінію, що вимагає від користувача проходження автентифікації перед отриманням доступу до CLI. Конфігурація захисту консолі на обладнанні Cisco передбачає перехід до режиму налаштування лінії консолі, встановлення пароля та активацію вимоги входу в систему. Хоча консольний пароль забезпечує певний рівень захисту, він має суттєві обмеження, зокрема пароль зберігається у конфігураційному файлі в незашифрованому вигляді, якщо не застосовано додаткових заходів шифрування. Крім того, консольна автентифікація зазвичай не передбачає розмежування прав доступу різних користувачів, що ускладнює аудит дій адміністраторів.

Віддалений доступ через протокол Telnet був широко поширеним методом управління мережевим обладнанням на ранніх етапах розвитку комп'ютерних мереж, але сьогодні вважається небезпечним через передачу всіх даних, включаючи облікові дані користувача, у відкритому вигляді. Будь-який зловмисник, який має можливість перехоплювати мережевий трафік, може отримати логін та пароль адміністратора, просто аналізуючи пакети Telnet-сесії. Незважаючи на очевидні недоліки безпеки, Telnet все ще використовується в ізольованих мережевих сегментах, де загроза перехоплення трафіку мінімальна,

наприклад, в окремих VLAN управління, які фізично відокремлені від загальної корпоративної мережі. Конфігурація Telnet-доступу передбачає налаштування віртуальних терміналів (VTY-ліній), встановлення паролів та визначення протоколів транспортування, які будуть дозволені для підключення.

Протокол SSH є сучасним стандартом захищеного віддаленого доступу до мережевого обладнання, який забезпечує шифрування всього трафіку сесії, включаючи процес автентифікації. SSH використовує криптографічні алгоритми для створення захищеного каналу між клієнтом та сервером, що унеможливорює перехоплення облікових даних або модифікацію команд під час передачі. Протокол підтримує різні методи автентифікації, включаючи парольну автентифікацію та автентифікацію на основі криптографічних ключів, причому останній метод вважається значно безпечнішим, оскільки виключає можливість атак перебору паролів. Для роботи SSH на мережевому пристрої необхідно згенерувати криптографічні ключі RSA, що використовуються для встановлення захищеного з'єднання, при цьому розмір ключа безпосередньо впливає на рівень криптографічної стійкості.

Порівняння протоколів Telnet та SSH з точки зору безпеки наведено в таблиці 6.1, яка демонструє ключові відмінності між цими технологіями. Як видно з таблиці 6.1, SSH має переваги практично за всіма параметрами безпеки, що робить його єдиним прийнятним варіантом для організації віддаленого доступу в сучасних корпоративних мережах.

Таблиця 6.1 - Порівняльна характеристика протоколів Telnet та SSH

Параметр	Telnet	SSH
Шифрування трафіку	Відсутнє	Повне шифрування сесії
Захист облікових даних	Передача у відкритому вигляді	Шифрування логіна та пароля
Стійкість до прослуховування	Відсутня	Висока
Методи автентифікації	Лише пароль	Пароль, криптографічні ключі
Номер порту за замовчуванням	23	22
Цілісність даних	Не гарантується	Перевірка цілісності кожного пакета
Рекомендації до використання	Не рекомендується	Обов'язково для продуктивних мереж

Налаштування SSH-доступу на обладнанні Cisco вимагає виконання декількох обов'язкових кроків, які забезпечують правильне функціонування протоколу. Спочатку необхідно встановити доменне ім'я пристрою, оскільки воно використовується як частина процесу генерації криптографічних ключів. Далі генеруються RSA-ключі з вказанням бажаного розміру модуля, при цьому рекомендується використовувати ключі довжиною не менше 2048 біт для забезпечення достатнього рівня криптографічної стійкості. Після генерації ключів створюються облікові записи локальних користувачів з паролями, які

будуть використовуватися для автентифікації SSH-сесій. На завершальному етапі конфігуруються VTY-лінії з вказівкою використання локальної бази користувачів для автентифікації та обмеженням дозволених протоколів транспортування лише до SSH, що виключає можливість підключення через незахищений Telnet.

Версія протоколу SSH також має значення для забезпечення безпеки, оскільки існують дві основні версії: SSH версії 1 та SSH версії 2. Перша версія протоколу містить відомі вразливості та не рекомендується до використання в сучасних мережах, тому на всіх пристроях слід примусово активувати використання SSH версії 2. Додатковим заходом безпеки є налаштування таймаутів неактивних сесій, що автоматично розриває SSH-з'єднання після певного періоду бездіяльності, зменшуючи ризик несанкціонованого використання залишеної активної сесії. Також доцільно обмежити кількість одночасних VTY-сесій та налаштувати списки контролю доступу, які визначають, з яких IP-адрес дозволено встановлювати SSH-підключення до мережевого обладнання.

6.2. Механізми автентифікації та управління привілеями

Автентифікація є процесом перевірки ідентичності користувача або пристрою, що намагається отримати доступ до ресурсів мережі, і становить першу лінію захисту від несанкціонованого доступу. На мережевому обладнанні автентифікація може здійснюватися на основі локальних облікових записів, які зберігаються безпосередньо в конфігурації пристрою, або через централізовані сервери AAA (Authentication, Authorization, Accounting), які забезпечують єдину точку управління обліковими записами для всієї мережевої інфраструктури. Локальна автентифікація є простішою в налаштуванні та не вимагає додаткової інфраструктури, але ускладнює централізоване управління обліковими записами в великих мережах, де кількість активних пристроїв може сягати сотень або тисяч одиниць.

Створення локальних облікових записів користувачів передбачає визначення імені користувача, встановлення пароля та опціональне призначення рівня привілеїв, який визначає набір команд та операцій, доступних цьому користувачу. Паролі локальних користувачів можуть зберігатися в конфігурації в різних форматах: у відкритому вигляді, що є абсолютно неприйнятним з точки зору безпеки, у форматі шифрування типу 7, який забезпечує лише базове обфускування та легко піддається дешифруванню, або у форматі MD5-хешування типу 5, який є криптографічно стійким односпрямованим алгоритмом. При створенні нового користувача рекомендується використовувати опцію “secret” замість “password”, що автоматично застосовує MD5-хешування до пароля перед збереженням його в конфігурації.

Рівні привілеїв на обладнанні Cisco представлені числовими значеннями від 0 до 15, причому кожен рівень визначає набір команд, які користувач з відповідним рівнем може виконувати. За замовчуванням використовуються три

основні рівні: рівень 0 передбачає дуже обмежений набір команд, включаючи лише базові команди виходу з системи, рівень 1 відповідає звичайному режиму користувача (user EXEC mode) з мінімальними правами перегляду інформації, а рівень 15 надає повний доступ до всіх команд та можливостей конфігурування, що еквівалентно привілейованому режиму (privileged EXEC mode). Адміністратори можуть створювати власні проміжні рівні привілеїв від 2 до 14, призначаючи кожному рівню специфічний набір дозволених команд, що дозволяє реалізувати гранульоване розмежування прав доступу відповідно до організаційних ролей персоналу.

Пароль привілейованого режиму (enable password) є критичним елементом безпеки, оскільки він контролює доступ до режиму, в якому можна вносити зміни в конфігурацію пристрою. Існує два методи встановлення цього пароля: команда “enable password” встановлює пароль, який зберігається у відкритому вигляді або з простим обфускуванням, тоді як команда “enable secret” створює MD5-хеш пароля, що забезпечує значно вищий рівень захисту. У випадку, коли в конфігурації присутні обидва типи паролів, система надає пріоритет “enable secret”, ігноруючи менш безпечний “enable password”. Важливо розуміти, що MD5-хешування є односпрямованою функцією, що означає неможливість відновлення оригінального пароля з хешу математичними методами, проте сучасні техніки атак через попередньо обчислені таблиці (rainbow tables) все ж створюють певні ризики для слабких паролів.

Для підвищення безпеки всіх паролів, що зберігаються в конфігурації, можна активувати глобальне шифрування паролів командою “service password-encryption”, що застосовує шифрування типу 7 до всіх незахищених паролів у конфігураційному файлі. Хоча шифрування типу 7 не є криптографічно стійким та може бути зламане за допомогою широко доступних інструментів, воно все ж забезпечує базовий рівень захисту від випадкового розголошення паролів при перегляді конфігурації через плече адміністратора або при випадковій демонстрації екрану. Необхідно пам'ятати, що це шифрування не може замінити використання “enable secret” або паролів типу “secret” для користувачів, які забезпечують значно надійніший захист через застосування криптографічних хеш-функцій.

Управління привілеями також включає можливість тимчасового підвищення рівня доступу користувача без необхідності повторної автентифікації. Це реалізується через механізм “enable”, який дозволяє користувачу з обмеженими правами перейти до привілейованого режиму, ввівши відповідний пароль. Для різних рівнів привілеїв можна встановити окремі паролі, що дозволяє створити багаторівневу систему контролю доступу, де різні групи адміністраторів мають доступ до різних наборів команд відповідно до їхніх функціональних обов'язків. Така диференціація прав доступу є важливою складовою принципу найменших привілеїв, згідно з яким кожен користувач повинен мати лише той мінімальний набір прав, який необхідний для виконання його безпосередніх службових обов'язків.

Централізована автентифікація через AAA-сервери (TACACS+ або RADIUS) представляє більш масштабоване рішення для великих корпоративних мереж, де кількість мережевих пристроїв та адміністраторів робить локальне управління обліковими записами непрактичним. AAA-сервери забезпечують не лише автентифікацію, але й авторизацію, яка визначає, які дії може виконувати автентифікований користувач, та облік (accounting), який реєструє всі дії користувачів для подальшого аудиту та розслідування інцидентів безпеки. Конфігурація AAA на мережевому пристрої передбачає визначення адреси AAA-сервера, встановлення спільного секретного ключа для шифрування комунікації між пристроєм та сервером, та налаштування методів автентифікації, які визначають послідовність спроб автентифікації через різні джерела.

У таблиці 6.2 наведено порівняння різних методів зберігання паролів на мережевому обладнанні, що допомагає зрозуміти переваги та недоліки кожного підходу. З таблиці 6.2 видно, що використання типу 5 (MD5-хеш) або типу 8/9 (SHA-256) забезпечує найкращий захист паролів, тоді як інші методи мають значні вразливості.

Таблиця 6.2 - Методи зберігання паролів на мережевому обладнанні Cisco

тип	Метод шифрування/хешування	Команда для встановлення	Рівень безпеки	Можливість відновлення оригінального пароля
	Відкритий текст	password	Дуже низький	Так, безпосередньо
	Vigenère cipher (обфускація)	password + service password-encryption	Низький	Так, легко дешифрується
	MD5 хеш	secret	Високий	Ні, лише атака перебором
	PBKDF2 з SHA-256	algorithm-type sha256 secret	Дуже високий	Ні, стійкий до rainbow tables
	Scrypt	algorithm-type scrypt secret	Найвищий	Ні, максимальна стійкість

Політика паролів є важливим аспектом забезпечення безпеки автентифікації та повинна визначати мінімальні вимоги до складності паролів, періодичність їх зміни, правила повторного використання попередніх паролів та процедури блокування облікових записів після декількох невдалих спроб входу. Хоча базове обладнання Cisco має обмежені можливості примусового застосування політик паролів, адміністратори повинні встановлювати організаційні правила, які вимагають від користувачів створювати складні паролі довжиною не менше 12 символів з комбінацією великих та малих літер, цифр та спеціальних символів. Регулярна зміна паролів, хоча й викликає певні дискусії в спільноті фахівців з безпеки, все ще рекомендується для критично важливих облікових записів, особливо якщо є підозра на можливу компрометацію облікових даних.

6.3. Списки контролю доступу як механізм фільтрації мережевого трафіку

Списки контролю доступу (ACL - Access Control Lists) є фундаментальним інструментом фільтрації трафіку на маршрутизаторах та комутаторах третього рівня, що дозволяє адміністраторам визначати правила, які контролюють проходження пакетів через мережеві інтерфейси на основі різноманітних критеріїв. ACL можуть використовуватися для реалізації політик безпеки, обмеження доступу до певних мережевих ресурсів, фільтрації небажаного трафіку, захисту від мережевих атак та оптимізації використання пропускну здатності каналів зв'язку. Правильне проектування та імплементація списків контролю доступу вимагає глибокого розуміння потоків трафіку в мережі, оскільки неправильно налаштовані ACL можуть заблокувати легітимний трафік або, навпаки, дозволити проходження небажаних пакетів, створюючи вразливості в системі безпеки.

Списки контролю доступу на обладнанні Cisco поділяються на два основні типи: стандартні ACL та розширені ACL, які відрізняються за набором критеріїв, що можуть використовуватися для прийняття рішення про дозвіл або заборону проходження пакета. Стандартні списки контролю доступу (номери від 1 до 99 та від 1300 до 1999 у випадку нумерованих ACL) можуть фільтрувати трафік виключно на основі IP-адреси джерела пакета, що робить їх простими у конфігуруванні, але обмеженими у функціональності. Розширені списки контролю доступу (номери від 100 до 199 та від 2000 до 2699) надають значно більше можливостей, дозволяючи фільтрувати пакети на основі IP-адреси джерела, IP-адреси призначення, протоколу транспортного або мережного рівня, номерів портів джерела та призначення, прапорців TCP, типів ICMP-повідомлень та інших параметрів.

Кожен список контролю доступу складається з послідовності правил (access control entries - ACE), які обробляються пристроєм послідовно зверху вниз до моменту, коли пакет відповідає умовам певного правила. Коли знайдено відповідність, до пакета застосовується дія, вказана в цьому правилі (permit для дозволу проходження або deny для блокування), і обробка ACL припиняється без перевірки наступних правил. Наприкінці кожного списку контролю доступу неявно присутнє правило "deny any", яке блокує весь трафік, що не відповідав жодному з явно визначених правил, тому адміністратори повинні обов'язково включати дозвільні правила для легітимного трафіку, інакше ACL заблокує всі пакети. Порядок розміщення правил у списку має критичне значення, оскільки більш специфічні правила повинні розташовуватися перед більш загальними, щоб уникнути ситуацій, коли загальне правило перехоплює трафік, який повинен був би оброблятися більш специфічним правилом.

Стандартні списки контролю доступу зазвичай використовуються для простих задач фільтрації, наприклад, обмеження доступу до VTY-ліній маршрутизатора лише з певних IP-адрес адміністративної мережі або фільтрації оновлень протоколів маршрутизації. Синтаксис стандартного ACL є відносно

простим: після ідентифікатора списку (номера або імені) вказується дія (permit або deny), IP-адреса джерела та маска підстановки (wildcard mask), яка визначає, які біти IP-адреси мають збігатися. Маска підстановки є інверсією звичайної маски підмережі: біт зі значенням 0 в масці підстановки означає, що відповідний біт IP-адреси має точно збігатися, тоді як біт зі значенням 1 означає, що відповідний біт ігнорується при перевірці. Наприклад, маска підстановки 0.0.0.255 означає, що перші три октети IP-адреси повинні збігатися точно, а четвертий октет може мати будь-яке значення, що фактично визначає цілу підмережу класу C.

Розширені списки контролю доступу надають значно більшу гнучкість та дозволяють реалізовувати складні політики безпеки з детальним контролем різних типів трафіку. Крім IP-адрес джерела та призначення, розширені ACL можуть фільтрувати пакети на основі протоколу, наприклад, дозволяючи TCP та UDP, але блокуючи ICMP, або навпаки. Для протоколів транспортного рівня можна вказувати номери портів або діапазони портів, що дозволяє створювати правила, які блокують певні сервіси, наприклад, заборону вихідних з'єднань на порт 23 (Telnet) для попередження використання незахищеного протоколу. Можливість перевірки прапорців TCP дозволяє створювати правила, які блокують спроби встановлення нових з'єднань (пакети з встановленим прапорцем SYN), але дозволяють пакети існуючих з'єднань, що є корисним для реалізації простого stateful-фільтрування на пристроях, які не підтримують повноцінні міжмережеві екрани.

Застосування списків контролю доступу до інтерфейсів маршрутизатора здійснюється шляхом прив'язки ACL у вхідному або вихідному напрямку, що визначає, коли буде відбуватися фільтрація відносно процесу маршрутизації. Вхідний ACL (inbound) обробляє пакети безпосередньо після їх надходження на інтерфейс, до прийняття рішення про маршрутизацію, що дозволяє заблокувати небажаний трафік на самому ранньому етапі та зменшити навантаження на процесор маршрутизатора. Вихідний ACL (outbound) застосовується до пакетів після прийняття рішення про маршрутизацію, безпосередньо перед їх відправленням через інтерфейс, що є корисним для фільтрації трафіку, який покидає певний мережевий сегмент. Вибір напрямку застосування ACL залежить від специфічних вимог політики безпеки та топології мережі, але загальна рекомендація полягає в розміщенні стандартних ACL якомога ближче до призначення трафіку, а розширених ACL - якомога ближче до джерела, що мінімізує обсяг небажаного трафіку в мережі.

У таблиці 6.3 представлено основні відмінності між стандартними та розширеними списками контролю доступу, що допомагає обрати відповідний тип ACL для конкретних задач фільтрації. Аналіз таблиці 6.3 показує, що вибір типу ACL має базуватися на складності вимог до фільтрації та необхідності гранулярного контролю трафіку.

Таблиця 6.3 - Порівняння типів списків контролю доступу

Характеристика	Стандартні ACL	Розширені ACL
Діапазон номерів	1-99, 1300-1999	100-199, 2000-2699
Критерії фільтрації	Лише IP-адреса джерела	IP джерела/призначення, протокол, порти
Місце розміщення	Близько до призначення	Близько до джерела
Складність конфігурації	Низька	Середня/Висока
Гнучкість	Обмежена	Висока
Типові застосування	Обмеження доступу до VTY, фільтрація маршрутів	Детальна фільтрація трафіку, політики безпеки
Вплив на продуктивність	Мінімальний	Помірний (залежить від складності правил)

Іменовані списки контролю доступу (named ACL) представляють альтернативу нумерованим ACL та надають додаткові переваги в управлінні складними конфігураціями. Використання описових імен замість номерів робить конфігурацію більш зрозумілою та самодокументованою, що особливо важливо в великих мережах з численними списками контролю доступу. Іменовані ACL також надають можливість редагувати окремі правила без необхідності видалення та повторного створення всього списку, що було обов'язковим для нумерованих ACL у старіших версіях операційної системи. Кожне правило в іменованому ACL може мати порядковий номер (sequence number), що дозволяє вставляти нові правила в довільні позиції списку, змінювати порядок правил або видаляти окремі записи без впливу на решту конфігурації.

Продуктивність обробки списків контролю доступу є важливим фактором, особливо на високонавантажених маршрутизаторах, де через інтерфейси проходять мільйони пакетів на секунду. Кожен пакет повинен бути перевірений на відповідність всім правилам ACL, поки не буде знайдено відповідність або досягнуто кінця списку, тому надмірно довгі списки з сотнями правил можуть суттєво збільшити затримку обробки пакетів. Для оптимізації продуктивності рекомендується розміщувати найбільш часто спрацьовуючі правила на початку списку, щоб мінімізувати кількість перевірок для типового трафіку, об'єднувати множинні правила в більш загальні, де це можливо без втрати необхідної гранулярності контролю, та використовувати апаратну підтримку ACL на платформах, які надають таку можливість через спеціалізовані ASIC або TCAM.

Моніторинг та аудит списків контролю доступу є невід'ємною частиною процесу управління безпекою мережі, оскільки неправильно налаштовані або

застарілі ACL можуть створювати вразливості або блокувати легітимний бізнес-трафік. Операційна система Cisco IOS надає команди для перегляду налаштованих ACL, перевірки кількості пакетів, що відповідали кожному правилу (hit counts), та тестування відповідності конкретних пакетів правилам списку без фактичного пропускання трафіку через інтерфейс. Лічильники спрацювань правил є особливо корисними для виявлення неактивних правил, які ніколи не спрацьовували та можуть бути видалені, або для підтвердження, що певні правила блокування дійсно перехоплюють зловмисний трафік. Регулярний аудит конфігурацій ACL має включати перевірку відповідності налаштувань актуальним політикам безпеки організації, видалення застарілих правил, оптимізацію порядку правил для покращення продуктивності та документування призначення кожного списку контролю доступу.

6.4. Практичні аспекти реалізації політик безпеки на мережевому обладнанні

Комплексна стратегія захисту мережевої інфраструктури вимагає інтеграції всіх розглянутих механізмів безпеки в єдину систему, що забезпечує багаторівневий захист від різноманітних загроз. Концепція захисту в глибину (defense in depth) передбачає використання множинних незалежних рівнів захисту, так щоб компрометація одного рівня не призводила до повного порушення безпеки системи, а зловмисник стикався з додатковими перешкодами на кожному етапі атаки. На рівні мережевого обладнання це означає комбінування захищеного доступу до CLI через SSH, багатофакторної автентифікації через AAA-сервери, гранулярного розмежування прав через рівні привілеїв та комплексної фільтрації трафіку через списки контролю доступу.

Проектування політик безпеки повинно базуватися на аналізі ризиків, що передбачає ідентифікацію критичних активів мережевої інфраструктури, оцінку потенційних загроз та вразливостей, визначення ймовірності реалізації загроз та потенційного впливу успішних атак. Критичними активами зазвичай є маршрутизатори граничної мережі, які контролюють підключення до Інтернету, комутатори ядра, через які проходить весь внутрішній трафік організації, та пристрої, що забезпечують критичні сервіси, такі як DNS, DHCP або контролери бездротових мереж. Для кожного класу активів мають бути визначені специфічні вимоги безпеки, що враховують як технічні можливості обладнання, так і організаційні процедури управління доступом та реагування на інциденти.

Сегментація мережі за допомогою VLAN у поєднанні зі списками контролю доступу створює ефективний механізм ізоляції трафіку різних груп користувачів та обмеження поширення атак у межах одного сегмента. Типова корпоративна мережа може включати окремі VLAN для різних департаментів організації, гостьової мережі Wi-Fi, мережі управління інфраструктурою, серверного сегменту, мережі IP-телефонії та інших функціональних зон. Між цими VLAN налаштовуються ACL, які дозволяють лише необхідний

міжсегментний трафік, наприклад, дозволяючи користувацьким VLAN доступ до серверів DNS та корпоративних веб-додатків, але блокуючи прямі з'єднання між робочими станціями різних департаментів. Особливої уваги потребує ізоляція мережі управління, доступ до якої повинен бути обмежений виключно адміністративними станціями через надійно автентифіковані VPN-з'єднання.

Захист площини управління (management plane) мережевого обладнання включає не лише обмеження доступу до CLI, але й захист протоколів управління, таких як SNMP, NTP, Syslog та інших сервісів, які використовуються для моніторингу та конфігурування пристроїв. Протокол SNMP у версіях 1 та 2c передає community strings у відкритому вигляді та не забезпечує автентифікації, тому для продуктивних мереж обов'язковим є використання SNMP версії 3 з криптографічною автентифікацією та шифруванням. Списки контролю доступу повинні обмежувати SNMP-запити лише з IP-адрес авторизованих систем моніторингу, а community strings мають бути складними та регулярно змінюватися. Синхронізація часу через NTP є критично важливою для точності логів та розслідування інцидентів, тому маршрутизатори повинні бути налаштовані на використання внутрішніх або зовнішніх довірених NTP-серверів з автентифікацією, що попереджає атаки модифікації часу.

Логування подій безпеки та централізований збір логів на виділеному syslog-сервері забезпечує можливість моніторингу активності в мережі, виявлення аномалій та розслідування інцидентів безпеки постфактум. Мережеве обладнання має бути налаштоване на генерацію логів для всіх значущих подій, включаючи успішні та невдалі спроби автентифікації, зміни конфігурації, спрацювання правил ACL (за необхідності), зміни стану інтерфейсів та протоколів маршрутизації, помилки та аварійні ситуації. Рівні важливості логів (від 0-Emergency до 7-Debugging) дозволяють фільтрувати повідомлення та концентрувати увагу на найбільш критичних подіях, хоча для повноцінного аудиту безпеки рекомендується зберігати логи всіх рівнів. Централізовані системи управління логами (SIEM - Security Information and Event Management) можуть автоматично аналізувати потоки логів з усіх мережевих пристроїв, виявляти підозрілі патерни активності та генерувати оповіщення для команди реагування на інциденти.

На рисунку 6.1 наведено схему типової конфігурації безпеки корпоративного маршрутизатора, що ілюструє інтеграцію різних механізмів захисту. Рисунок демонструє граничний маршрутизатор, підключений з одного боку до Інтернету через інтерфейс GigabitEthernet0/0, а з іншого боку до внутрішньої корпоративної мережі через інтерфейс GigabitEthernet0/1. На зовнішньому інтерфейсі застосовано вхідний розширений ACL під назвою "INTERNET-INBOUND", який блокує весь вхідний трафік, за винятком відповідей на з'єднання, ініційовані з внутрішньої мережі, та дозволених сервісів, таких як вхідна електронна пошта та веб-трафік до DMZ. На внутрішньому інтерфейсі застосовано вихідний стандартний ACL "INTERNAL-ADMIN", який обмежує доступ до VTY-ліній маршрутизатора лише з підмережі адміністративних станцій 10.10.100.0/24. Маршрутизатор

налаштовано на використання SSH версії 2 для віддаленого доступу з обов'язковою автентифікацією через зовнішній TACACS+ сервер з резервним використанням локальних облікових записів у випадку недоступності AAA-сервера. Всі події безпеки логуються на виділений Syslog-сервер 10.10.200.50, а конфігурації регулярно резервуються на TFTP-сервер 10.10.200.51.

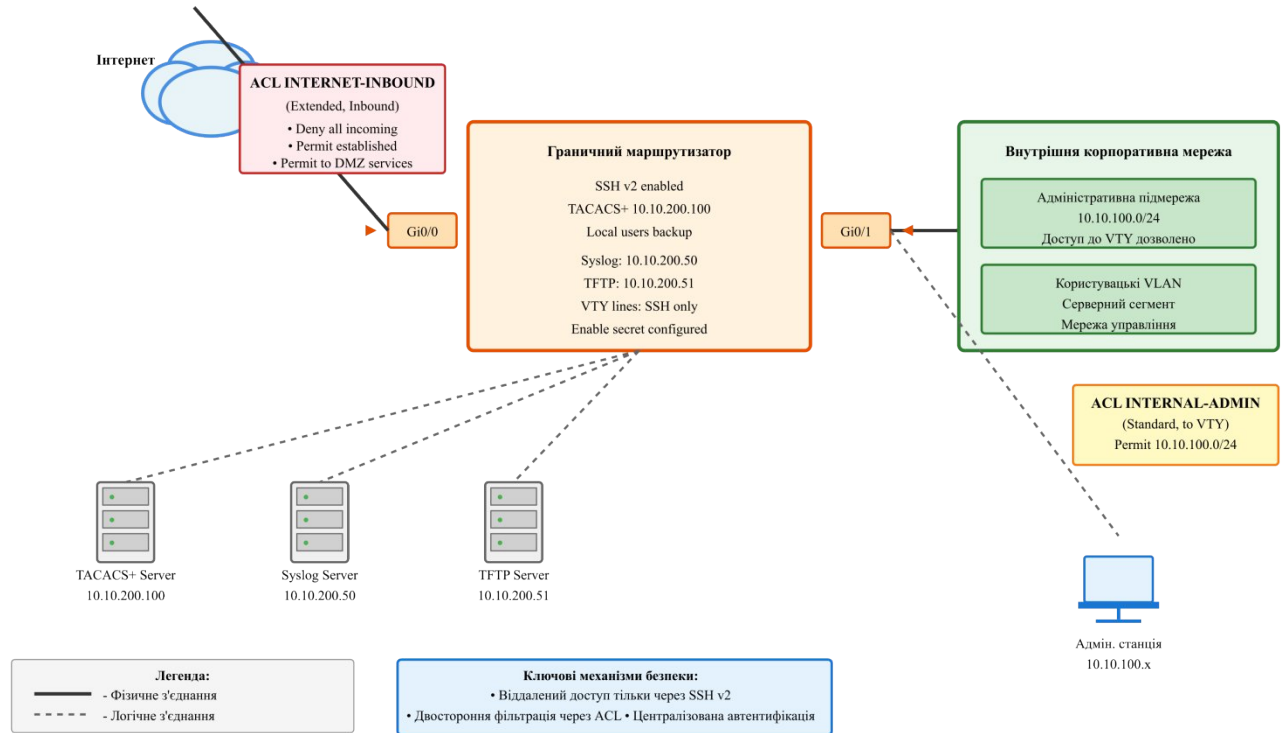


Рисунок 6.1 - Схема конфігурації безпеки корпоративного граничного маршрутизатора

Регулярний аудит безпеки мережевої інфраструктури є необхідною практикою для виявлення неправильних конфігурацій, застарілих налаштувань та потенційних вразливостей. Аудит повинен включати перевірку сили паролів та методів їх зберігання, аналіз налаштувань доступу до CLI з верифікацією, що Telnet вимкнено на всіх продуктивних пристроях, перегляд конфігурацій ACL на предмет надмірно дозвільних правил або логічних помилок у послідовності обробки, валідацію налаштувань AAA та перевірку функціонування механізмів логування. Автоматизовані інструменти аудиту можуть сканувати конфігурації множинних пристроїв та генерувати звіти про відхилення від базових налаштувань безпеки (security baselines), що значно спрощує процес аудиту в великих мережах.

Управління змінами (change management) є критично важливим процесом, що забезпечує контрольоване впровадження модифікацій у конфігурації мережевого обладнання з мінімізацією ризиків порушення роботи мережі або створення вразливостей безпеки. Кожна зміна конфігурації, включаючи додавання або модифікацію ACL, створення нових облікових записів користувачів або зміну паролів, повинна проходити через формальний процес запиту змін, оцінки впливу, схвалення уповноваженою особою,

тестування в лабораторному середовищі (якщо можливо), впровадження у продуктивну мережу та верифікації результатів. Всі зміни мають документуватися з вказівкою дати, часу, особи, яка впроваджувала зміну, причини модифікації та результатів впровадження, що створює аудиторський слід для подальшого аналізу.

Резервне копіювання конфігурацій є обов'язковою практикою, що забезпечує можливість швидкого відновлення налаштувань обладнання у випадку збою або випадкової шкідливої модифікації. Конфігурації повинні резервуватися автоматично після кожної зміни та зберігатися на віддалених серверах з належним захистом доступу, оскільки файли конфігурацій містять чутливу інформацію, включаючи хеші паролів, спільні ключі для криптографічних протоколів та іншу інформацію, яка може бути використана зловмисниками. Системи управління версіями можуть застосовуватися для відстеження історії змін конфігурацій, порівняння різних версій та відкату до попередніх стабільних налаштувань у випадку виявлення проблем після впровадження змін.

Тестування процедур відновлення після збоїв та інцидентів безпеки має проводитися регулярно для підтвердження, що резервні копії конфігурацій є коректними та можуть бути успішно відновлені, механізми автентифікації функціонують належним чином, включаючи резервні методи доступу на випадок недоступності основних систем, списки контролю доступу забезпечують очікувану фільтрацію без блокування критичного трафіку, а персонал володіє необхідними навичками для виконання процедур відновлення в умовах стресу. Результати тестувань мають документуватися з ідентифікацією виявлених проблем та розробкою планів їх усунення, що забезпечує постійне вдосконалення процесів забезпечення безпеки та готовності до реагування на інциденти.

Тема 7. Проектування та моделювання комп'ютерних мереж

7.1. Методологія проектування комп'ютерних мереж

Проектування комп'ютерної мережі являє собою складний багатоетапний процес, який вимагає системного підходу та врахування численних технічних, економічних та організаційних факторів. Якість проектних рішень безпосередньо впливає на продуктивність, масштабованість, безпеку та загальну ефективність мережевої інфраструктури організації. На відміну від інтуїтивного підходу, методологічне проектування передбачає послідовне виконання чітко визначених етапів, кожен з яких має власні цілі, завдання та критерії успішності. Сучасні методології проектування базуються на принципах ієрархічного моделювання, модульності та можливості подальшого розширення мережі без кардинальної перебудови її архітектури.

Перший етап проектування полягає у детальному аналізі вимог замовника та специфіки його бізнес-процесів. Проектувальник повинен з'ясувати кількість користувачів мережі, їх географічне розташування, типи застосунків та сервісів, які планується використовувати, а також специфічні вимоги до продуктивності, доступності та безпеки. Важливо враховувати не лише поточні потреби, але й перспективи розвитку організації на найближчі три-п'ять років, оскільки мережева інфраструктура створюється з розрахунком на тривалу експлуатацію. На цьому етапі також визначаються бюджетні обмеження проекту, терміни його реалізації та наявність кваліфікованого персоналу для подальшого обслуговування мережі. Результатом аналізу вимог стає технічне завдання, яке формалізує всі ключові параметри майбутньої мережі та слугує основою для наступних етапів проектування.

Визначення масштабу мережі безпосередньо пов'язане з її географічним охопленням та кількістю вузлів. Локальна мережа підприємства може охоплювати один офіс, декілька поверхів будівлі або навіть комплекс будівель у межах одного кампусу. Для мереж типу LAN характерна висока швидкість передачі даних, використання технологій Ethernet на швидкостях від 100 Мбіт/с до 100 Гбіт/с, відносно низька вартість каналів зв'язку та можливість централізованого управління. Якщо організація має територіально розподілені філії, виникає потреба у побудові мережі масштабу MAN або WAN, що передбачає використання каналів провайдерів телекомунікаційних послуг, організацію VPN-з'єднань та вирішення специфічних завдань забезпечення якості обслуговування на великих відстанях.

Вибір топології мережі визначається як технічними, так і економічними міркуваннями. В таблиці 7.1 наведено порівняльний аналіз основних топологій комп'ютерних мереж.

Таблиця 7.1 - Порівняння топологій комп'ютерних мереж

Топологія	Переваги	Недоліки	Область застосування
Шина	Простота реалізації, низька вартість	Складність локалізації несправностей, обмежена довжина	Застарілі малі мережі
Зірка	Централізоване управління, простота діагностики	Залежність від центрального вузла	Сучасні LAN на базі комутаторів
Кільце	Рівномірний доступ до середовища, відсутність колізій	Відмова одного вузла може паралізувати всю мережу	Token Ring, FDDI (рідко)
Дерево подібна	Ієрархічна структура, масштабованість	Складність проектування, множинні точки відмови	Корпоративні мережі середнього розміру
Повнозв'язна	Максимальна відмовостійкість	Висока вартість, складність управління	Критичні системи, ядра мереж

Фізична топологія визначає реальне розташування кабельних трас, мережевого обладнання та кінцевих пристроїв у просторі будівлі. Логічна топологія описує шляхи передачі даних між вузлами мережі незалежно від фізичної реалізації. Наприклад, сучасні мережі Ethernet фізично реалізуються за топологією “зірка” з використанням комутаторів як центральних вузлів, однак логічно вони можуть бути організовані у вигляді плоскої мережі або ієрархічної структури з сегментацією за допомогою віртуальних локальних мереж. Розуміння відмінності між фізичною та логічною топологіями є критично важливим для правильного проектування та усунення несправностей у мережі.

7.2. IP-адресація та підмережі в проектуванні мереж

Планування IP-адресації є одним з найважливіших етапів проектування мережі, оскільки від правильності його виконання залежить ефективність використання адресного простору, можливість подальшого масштабування та зручність адміністрування. Традиційна класова адресація, яка передбачала використання фіксованих масок підмереж для класів А, В та С, виявилася неефективною через марнотратне витрачання адресного простору. Для вирішення цієї проблеми було запроваджено технологію VLSM (Variable Length Subnet Mask), яка дозволяє використовувати маски підмереж змінної довжини в межах однієї мережі. Завдяки VLSM проектувальник може створювати підмережі різного розміру відповідно до реальних потреб кожного сегмента мережі, мінімізуючи втрати невикористаних IP-адрес.

Процес розрахунку IP-адресації методом VLSM починається з визначення кількості хостів у кожному сегменті мережі. Спочатку всі сегменти впорядковуються за спаданням кількості вузлів, що дозволяє раціонально розподілити адресний простір, виділяючи спершу найбільші підмережі. Для

кожного сегмента визначається мінімальна довжина префікса, яка забезпечить достатню кількість адрес для всіх хостів з урахуванням резерву на майбутнє зростання. При розрахунках необхідно пам'ятати, що перша адреса підмережі є мережевою адресою, остання - широкомовною адресою, тому кількість адрес для хостів дорівнює $2^{(32-n)} - 2$, де n - довжина префікса маски підмережі. Наприклад, якщо сегмент містить 50 робочих станцій, мінімальний префікс становитиме /26, що забезпечує 62 адреси для хостів.

Після визначення розмірів підмереж виконується їх послідовне виділення з наявного адресного простору. Припустимо, організації виділено приватну мережу 172.16.0.0/16, і необхідно створити підмережі для декількох відділів. Найбільший відділ налічує 200 комп'ютерів, що потребує підмережі з префіксом /24 (254 адреси для хостів). Йому можна виділити підмережу 172.16.0.0/24. Наступний відділ з 80 комп'ютерами отримає підмережу 172.16.1.0/25 (126 адрес). Відділ з 30 комп'ютерами може використовувати 172.16.1.128/27 (30 адрес). Такий підхід забезпечує ефективне використання кожної адреси та дозволяє зберегти значну частину адресного простору для майбутніх потреб. Важливо також врахувати адреси для з'єднань точка-точка між маршрутизаторами, для яких достатньо підмереж з префіксом /30, що надає лише дві корисні адреси.

Документування схеми IP-адресації є обов'язковим елементом проектної документації. В таблиці 7.2 наведено приклад розподілу адресного простору для типової корпоративної мережі.

Таблиця 7.2 - Приклад розподілу IP-адресного простору

Підме режа	Мережев а адреса/префікс	Маска підмережі	Діапазо н адрес хостів	К ількість хостів	Приз начення
Відділ продажів	172.16.0. 0/24	255.255. 255.0	172.16.0 .1 - 172.16.0.254	2 54	Робо чі станції, принтери
Відділ розробки	172.16.1. 0/25	255.255. 255.128	172.16.1 .1 - 172.16.1.126	1 26	Робо чі станції
Адміні страція	172.16.1. 128/27	255.255. 255.224	172.16.1 .129 - 172.16.1.158	3 0	Робо чі станції керівників
Серве рна ферма	172.16.2. 0/26	255.255. 255.192	172.16.2 .1 - 172.16.2.62	6 2	Серв ери, системи зберігання
DMZ	172.16.3. 0/27	255.255. 255.224	172.16.3 .1 - 172.16.3.30	3 0	Публі чні сервери
Управ ління	192.168. 100.0/24	255.255. 255.0	192.168. 100.1 - 192.168.100.254	2 54	Інтер фейси управління обладнання м

Суммаризація маршрутів є важливою технікою оптимізації таблиць маршрутизації, яка особливо актуальна при проектуванні великих мереж з ієрархічною структурою. Суть суммаризації полягає у представленні декількох суміжних підмереж однією агрегованою мережевою адресою з коротшим префіксом. Наприклад, якщо маршрутизатор має інформацію про чотири підмережі 172.16.4.0/24, 172.16.5.0/24, 172.16.6.0/24 та 172.16.7.0/24, їх можна представити єдиним маршрутом 172.16.4.0/22, що зменшує розмір таблиці маршрутизації та прискорює процес пошуку маршруту. Коректна суммаризація потребує ретельного планування адресного простору ще на етапі проектування, щоб суміжні підмережі логічно групувалися у більші блоки.

7.3. Віртуальні локальні мережі та принципи сегментації

Планування віртуальних локальних мереж (VLAN) є ключовим аспектом логічної організації сучасної корпоративної мережі. VLAN дозволяють розділити одну фізичну мережу на декілька логічно ізольованих сегментів, що забезпечує покращену безпеку, спрощує управління широкомовним трафіком та підвищує загальну продуктивність мережі. Кожна VLAN утворює окремий домен широкомовної передачі, тому широкомовні кадри, згенеровані в одній VLAN, не поширюються на інші віртуальні мережі. Це особливо важливо у великих мережах, де надмірний широкомовний трафік може суттєво знижувати ефективну пропускну здатність каналів та створювати додаткове навантаження на процесори мережевих пристроїв.

Проектування структури VLAN зазвичай здійснюється з урахуванням організаційної структури підприємства, функціонального призначення обладнання та вимог безпеки. Типовий підхід передбачає створення окремих VLAN для різних відділів або підрозділів організації, що дозволяє застосовувати диференційовані політики безпеки та якості обслуговування. Наприклад, VLAN відділу розробки може мати необмежений доступ до Інтернету та внутрішніх серверів розробки, тоді як VLAN гостьового доступу буде обмежена лише виходом в Інтернет без можливості звернення до внутрішніх ресурсів. Окрема VLAN часто виділяється для IP-телефонії, що дозволяє застосовувати механізми пріоритезації голосового трафіку та забезпечити йому гарантовану якість обслуговування незалежно від завантаженості мережі іншими типами трафіку.

Для забезпечення взаємодії між VLAN використовується маршрутизація, яка може виконуватися на маршрутизаторі або на комутаторі третього рівня. Класична схема “router-on-a-stick” передбачає підключення маршрутизатора до комутатора через один фізичний інтерфейс, який конфігурується з використанням стандарту IEEE 802.1Q для передачі теґованого трафіку декількох VLAN. На маршрутизаторі створюються логічні субінтерфейси, кожен з яких асоціюється з окремою VLAN та має власну IP-адресу, що виконує роль шлюзу за замовчуванням для цієї VLAN. Хоча цей підхід простий у реалізації, він створює вузьке місце продуктивності, оскільки весь міжвлановий трафік проходить через один фізичний канал. У великих мережах

перевагу надають комутаторам третього рівня, які здатні виконувати маршрутизацію між VLAN на апаратному рівні з мінімальними затримками.

Стандарт IEEE 802.1Q, також відомий як Dot1Q, визначає механізм додавання тегів VLAN до Ethernet-кадрів. Тег VLAN являє собою 4-байтове поле, яке вставляється у кадр після адреси джерела та містить ідентифікатор VLAN (12 біт), пріоритет (3 біти) та інші службові поля. Комутатори, що підтримують 802.1Q, можуть передавати через один фізичний порт трафік декількох VLAN шляхом маркування кадрів відповідними тегами. Порти комутатора конфігуруються як access-порти, які належать до однієї VLAN та передають нетеговані кадри, або як trunk-порти, які передають теговані кадри декількох VLAN. Trunk-порти зазвичай використовуються для з'єднання комутаторів між собою або для підключення маршрутизаторів та серверів, які повинні мати доступ до декількох VLAN одночасно.

Рисунок 7.1 ілюструє типову схему організації VLAN у корпоративній мережі. На рисунку зображено три поверхи офісної будівлі, на кожному з яких розміщено комутатор доступу. Комутатори доступу з'єднані trunk-каналами з комутатором розподілу, який виконує функції маршрутизації між VLAN. На першому поверху в VLAN 10 знаходяться робочі станції відділу продажів, на другому поверсі VLAN 20 містять комп'ютери бухгалтерії, на третьому поверсі VLAN 30 належить відділу розробки. Серверна ферма підключена до окремої VLAN 100 через trunk-порт. Всі VLAN мають власні IP-підмережі: VLAN 10 використовує 172.16.10.0/24, VLAN 20 - 172.16.20.0/24, VLAN 30 - 172.16.30.0/24, VLAN 100 - 172.16.100.0/24. Trunk-канали між комутаторами позначені подвійними лініями та переносять трафік всіх VLAN з тегами 802.1Q. Access-порти, до яких підключені робочі станції, показані одинарними лініями та належать до відповідних VLAN без використання тегів на кінцевих пристроях.

Безпека VLAN потребує особливої уваги на етапі проектування. Існує декілька типових векторів атак, специфічних для мереж з VLAN, включаючи VLAN hopping та атаки на протокол динамічного транкінгу DTP. VLAN hopping дозволяє зловмиснику отримати доступ до трафіку інших VLAN шляхом подвійного тегування кадрів або експлуатації неправильно сконфігурованих trunk-портів. Для запобігання таким атакам рекомендується відключати невикористовувані порти комутатора та переводити їх у неактивну VLAN, явно конфігурувати режим портів (access або trunk) замість використання автоматичного узгодження, використовувати окрему нативну VLAN для trunk-портів та забороняти її використання для передачі користувацького трафіку. Також важливо застосовувати приватні VLAN (PVLAN) у сегментах з підвищеними вимогами до ізоляції, наприклад, у серверних фермах або в мережах надавачів хостингу.

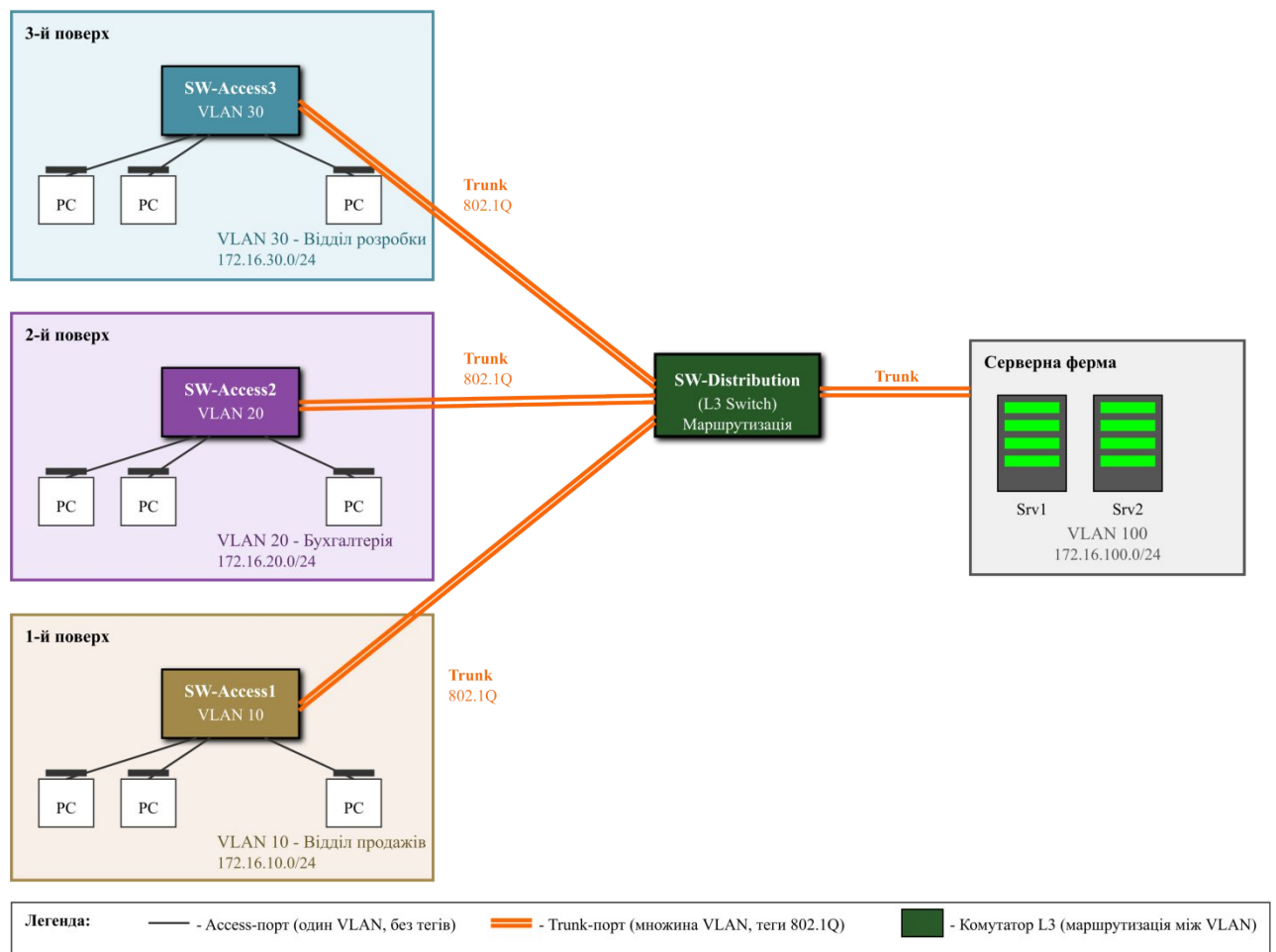


Рисунок 7.1 - Організація VLAN у триповерховій офісній будівлі

7.4. Модульність, надмірність та моделювання мереж

Принцип модульності в проектуванні мереж передбачає розбиття складної мережевої інфраструктури на логічно відокремлені функціональні блоки, кожен з яких виконує специфічні завдання та має чітко визначені інтерфейси взаємодії з іншими модулями. Класична ієрархічна модель Cisco описує три основних рівні корпоративної мережі: рівень доступу (access layer), рівень розподілу (distribution layer) та ядро мережі (core layer). Рівень доступу забезпечує безпосереднє підключення кінцевих пристроїв користувачів та виконує базові функції комутації, застосування політик безпеки на рівні портів, класифікації трафіку. Обладнання рівня доступу зазвичай не потребує надвисокої продуктивності, але повинно забезпечувати достатню щільність портів та підтримку необхідних функцій Layer 2.

Рівень розподілу виконує функції агрегації трафіку від комутаторів доступу, маршрутизації між VLAN, застосування політик безпеки та якості обслуговування, контролю меж широкомовних доменів. На цьому рівні часто реалізуються списки контролю доступу, які фільтрують трафік між різними сегментами мережі, механізми резервування шляхів та балансування навантаження. Комутатори розподілу повинні мати високу продуктивність маршрутизації, підтримку протоколів резервування типу HSRP або VRRP,

можливість агрегації каналів за технологією EtherChannel. Ядро мережі відповідає виключно за швидку передачу великих обсягів трафіку між різними модулями рівня розподілу з мінімальними затримками. Обладнання ядра характеризується максимальною пропускною здатністю, дублюванням критичних компонентів та простотою конфігурації, оскільки складні функції типу фільтрації трафіку можуть знижувати продуктивність.

Надмірність (redundancy) є критично важливим принципом проектування відмовостійких мереж. Повна відсутність надмірних елементів означає, що вихід з ладу будь-якого компонента паралізує роботу частини або всієї мережі. З іншого боку, надмірне дублювання призводить до невиправданого подорожчання проекту та ускладнення експлуатації. Рациональний підхід полягає у визначенні критичності кожного сегмента мережі та забезпеченні рівня надмірності, адекватного бізнес-вимогам. Для критичних сегментів, відмова яких призводить до повної зупинки бізнес-процесів, доцільно застосовувати повне дублювання активного обладнання, множинні канали зв'язку, резервні джерела живлення. Менш критичні сегменти можуть обмежуватися резервуванням лише окремих компонентів або взагалі працювати без надмірності, якщо допустимий певний час простою для відновлення.

Технологія EtherChannel дозволяє об'єднати декілька фізичних каналів Ethernet у один логічний канал з сумарною пропускною здатністю. Це забезпечує як збільшення доступної смуги пропускання між комутаторами, так і резервування на випадок відмови окремих каналів. Протоколи LACP (Link Aggregation Control Protocol) та PAgP (Port Aggregation Protocol) автоматизують процес створення та підтримки агрегованих каналів, відстежуючи стан окремих фізичних інтерфейсів та перерозподіляючи навантаження при виході з ладу одного з них. При проектуванні топології з використанням EtherChannel важливо враховувати, що всі порти агрегованого каналу повинні мати однакову швидкість, дуплексний режим та конфігурацію VLAN.

Протоколи резервування шлюзу за замовчуванням, такі як HSRP (Hot Standby Router Protocol), VRRP (Virtual Router Redundancy Protocol) та GLBP (Gateway Load Balancing Protocol), забезпечують відмовостійкість на рівні шлюзів для кінцевих пристроїв. Ці протоколи дозволяють двом або більше маршрутизаторам спільно використовувати одну віртуальну IP-адресу, яка конфігурується на робочих станціях як шлюз за замовчуванням. В нормальному режимі один маршрутизатор виконує роль активного шлюзу та обробляє весь трафік, тоді як резервні маршрутизатори перебувають у режимі очікування та постійно моніторять стан активного пристрою. При виході з ладу активного маршрутизатора один з резервних автоматично бере на себе його функції з мінімальною перервою у обслуговуванні, що часто залишається непомітним для кінцевих користувачів.

Використання засобів моделювання є невід'ємною частиною сучасного процесу проектування мереж. Cisco Packet Tracer являє собою безкоштовний симулятор мережевого обладнання, який дозволяє створювати віртуальні моделі мереж довільної складності, конфігурувати маршрутизатори та комутатори, перевіряти зв'язність між вузлами та аналізувати проходження

пакетів через мережу. На відміну від роботи з реальним обладнанням, моделювання не вимагає значних фінансових витрат, дозволяє швидко експериментувати з різними варіантами топології та конфігурації, забезпечує можливість імітації несправностей та перевірки поведінки мережі у нештатних ситуаціях. Packet Tracer підтримує широкий спектр мережевих пристроїв та протоколів, включаючи статичну та динамічну маршрутизацію, VLAN, EtherChannel, протоколи резервування, списки контролю доступу та багато інших технологій.

Процес моделювання зазвичай починається зі створення фізичної топології шляхом розміщення необхідних пристроїв на робочому полі та з'єднання їх відповідними типами каналів. Після побудови топології виконується логічна конфігурація пристроїв, яка включає призначення IP-адрес інтерфейсам, налаштування маршрутизації, створення VLAN та конфігурацію trunk-портів. Packet Tracer надає зручний інтерфейс командного рядка, максимально наближений до реального обладнання Cisco, що дозволяє практикувати навички роботи з IOS. Важливою функцією є режим симуляції, який дозволяє покроково відстежувати проходження пакетів через мережу, спостерігати процеси інкапсуляції та декапсуляції на різних рівнях моделі OSI, аналізувати вміст кадрів та пакетів на кожному етапі їх пересування.

Верифікація проектних рішень у середовищі моделювання передбачає систематичну перевірку всіх аспектів функціонування мережі. Спочатку перевіряється базова зв'язність між всіма вузлами мережі за допомогою утиліти ping, яка використовує протокол ICMP для виявлення доступності віддаленого хоста. Далі аналізуються таблиці маршрутизації на всіх маршрутизаторах, переконуючись, що вони містять маршрути до всіх підмереж, які мають бути доступними. CAM-таблиці комутаторів перевіряються для підтвердження правильності вивчення MAC-адрес. Списки контролю доступу тестуються шляхом спроб встановлення з'єднань, які мають бути заблоковані відповідно до політики безпеки. Механізми резервування перевіряються шляхом симуляції відмов обладнання або каналів зв'язку та спостереження за процесом автоматичного перемикавання на резервні шляхи.

В таблиці 7.3 наведено основні етапи верифікації проектних рішень у Packet Tracer.

Таблиця 7.3 - Етапи верифікації мережевого проекту

етап	Об'єкт перевірки	Методи перевірки	Критерії успішності
	Фізична зв'язність	Перевірка стану інтерфейсів, індикація портів	Всі активні інтерфейси в стані up/up
	IP-конфігурація	Команди show ip interface brief, show ipv6 interface	Коректні адреси, відсутність конфліктів
	Таблиці маршрутизації	show ip route, traceroute	Наявність маршрутів до всіх підмереж, оптимальні шляхи
	VLAN	show vlan, show	Правильне призначення

		interfaces trunk	портів, коректна передача тегів
	Міжвлановий трафік	ping між різними VLAN	Успішна маршрутизація між VLAN
	Списки доступу	Тестові з'єднання з різних джерел	Блокування заборонених з'єднань, дозвіл легітимних
	Резервування	Вимкнення інтерфейсів, пристроїв	Автоматичне перемикання, мінімальний час відновлення
	Продуктивність	Складний трафік, статистика інтерфейсів	Відсутність перевантажень, прийнятні затримки

Документування результатів моделювання є важливою складовою проектної документації. Рекомендується зберігати файли проектів Packet Tracer з різними сценаріями конфігурації, створювати знімки екрану з ключовими таблицями маршрутизації та результатами тестування, фіксувати всі виявлені проблеми та способи їх вирішення. Така документація стане цінним ресурсом на етапі реалізації проекту в реальному середовищі та значно спростить процес усунення несправностей під час експлуатації мережі. Моделювання також дозволяє підготувати детальні інструкції з конфігурації обладнання, які можуть бути використані технічним персоналом при розгортанні мережі, зменшуючи ймовірність помилок та скорочуючи час введення системи в експлуатацію.

Завершальним етапом процесу проектування є підготовка комплексної технічної документації, яка включає логічні та фізичні схеми мережі, таблиці IP-адресації, специфікації обладнання та матеріалів, кошторисні розрахунки, інструкції з монтажу та налаштування, плани тестування та приймання системи. Якість документації безпосередньо впливає на успішність реалізації проекту, оскільки вона є основним засобом комунікації між проектувальниками, монтажниками та експлуатаційним персоналом. Схеми повинні бути виконані з використанням стандартизованих графічних позначень, які зрозумілі всім учасникам процесу. Таблиці адресації мають бути організовані таким чином, щоб спростити процес призначення адрес при монтажі та полегшити пошук інформації при подальшому обслуговуванні мережі.

ТЕМА 1. ПРОТОКОЛИ ВНУТРІШНЬОЇ ДИНАМІЧНОЇ МАРШРУТИЗАЦІЇ (IGP)

1.1 Класифікація протоколів динамічної маршрутизації та їх роль у сучасних мережах

Сучасні комп'ютерні мережі характеризуються складною топологією, великою кількістю вузлів та динамічними змінами конфігурації, що робить статичну маршрутизацію неефективною для середовищ із десятками та сотнями маршрутизаторів. Протоколи динамічної маршрутизації дозволяють мережевим пристроям автоматично обмінюватися інформацією про доступні маршрути, адаптуватися до змін топології та забезпечувати оптимальну доставку пакетів без постійного ручного втручання адміністраторів. У контексті загальної ієрархії інтернет-маршрутизації всі протоколи динамічної маршрутизації поділяються на дві основні категорії: протоколи внутрішньої маршрутизації (Interior Gateway Protocols, IGP) та протоколи зовнішньої маршрутизації (Exterior Gateway Protocols, EGP). Протоколи IGP призначені для функціонування всередині однієї автономної системи (Autonomous System, AS) – сукупності мереж, що перебувають під єдиним адміністративним управлінням та використовують спільну політику маршрутизації.

Автономна система являє собою логічну межу, всередині якої адміністратор має повний контроль над конфігурацією маршрутизації, вибором протоколів та налаштуванням метрик. Типовими прикладами автономних систем є корпоративна мережа підприємства, мережа університетського кампусу або інфраструктура регіонального інтернет-провайдера. Протоколи IGP оптимізовані для роботи в таких середовищах, де всі маршрутизатори знаходяться під контролем однієї організації і мають можливість швидко обмінюватися детальною топологічною інформацією. На відміну від IGP, протоколи EGP використовуються для маршрутизації між різними автономними системами, забезпечуючи взаємодію між мережами різних організацій та провайдерів у глобальному масштабі інтернету. Найпоширенішим представником класу EGP є протокол BGP (Border Gateway Protocol), який дозволяє реалізувати складні політики маршрутизації на основі бізнес-вимог та домовленостей між організаціями.

У таблиці 1.1 наведено порівняльну характеристику основних відмінностей між протоколами внутрішньої та зовнішньої маршрутизації, що дозволяє краще зрозуміти їхні функціональні особливості та сфери застосування.

Таблиця 1.1 – Порівняння протоколів IGP та EGP

Характеристика	IGP	EGP
Область застосування	Всередині автономної системи	Між автономними системами
Приклади протоколів	RIP, OSPF, EIGRP, IS-IS	BGP

Основний критерій вибору маршруту	Технічні метрики (пропускна здатність, затримка, кількість переходів)	Політики маршрутизації, бізнес-домовленості
Швидкість збіжності	Висока (секунди-хвилини)	Повільніша (хвилини)
Масштабованість	Обмежена розміром AS (зазвичай до кількох сотень маршрутизаторів)	Глобальна (десятки тисяч AS)
Довіра між учасниками	Повна (єдине адміністративне управління)	Обмежена (різні організації)

Серед протоколів IGP найбільш широко використовуються RIP (Routing Information Protocol), OSPF (Open Shortest Path First), EIGRP (Enhanced Interior Gateway Routing Protocol) та IS-IS (Intermediate System to Intermediate System). Протокол RIP є одним із найстаріших протоколів маршрутизації, що базується на алгоритмі Беллмана-Форда та використовує кількість переходів (hop count) як метрику відстані до мережі призначення. OSPF представляє сучасний стандартизований протокол на основі алгоритму Дейкстри, який будує повну топологічну карту мережі та обчислює найкоротші шляхи з урахуванням вартості каналів зв'язку. EIGRP є власницьким протоколом компанії Cisco, який поєднує переваги дистанційно-векторного та link-state підходів, забезпечуючи швидку збіжність та ефективне використання мережевих ресурсів. IS-IS використовується переважно в мережах великих телекомунікаційних операторів та має архітектуру, подібну до OSPF, але працює безпосередньо на канальному рівні моделі OSI.

1.2 Архітектурні підходи до реалізації протоколів IGP

Протоколи внутрішньої маршрутизації можна класифікувати за архітектурним підходом до обміну інформацією та обчислення маршрутів на три основні типи: дистанційно-векторні (distance-vector), протоколи стану каналів (link-state) та гібридні протоколи. Дистанційно-векторні протоколи базуються на алгоритмі Беллмана-Форда і реалізують розподілене обчислення найкоротших шляхів шляхом періодичного обміну векторами відстаней між сусідніми маршрутизаторами. Кожен маршрутизатор зберігає таблицю маршрутизації, що містить інформацію про відстань (метрику) до кожної відомої мережі та наступний перехід (next hop) на шляху до неї. Маршрутизатори періодично надсилають своїм безпосереднім сусідам копії своєї таблиці маршрутизації або інформацію про зміни в ній, а сусіди використовують отриману інформацію для оновлення власних таблиць згідно з принципом “дізнався від сусіда про мережу – додай метрику каналу до сусіда”. Класичним представником цього підходу є протокол RIP, який використовує кількість переходів як метрику і має обмеження максимальної відстані у 15 переходів.

Основною перевагою дистанційно-векторних протоколів є простота реалізації та низькі вимоги до обчислювальних ресурсів маршрутизатора,

оскільки кожен пристрій зберігає лише результат обчислень (таблицю маршрутів), а не повну топологічну інформацію про мережу. Проте цей підхід має суттєві недоліки, серед яких повільна збіжність при змінах топології, можливість виникнення петель маршрутизації та проблема “рахування до нескінченності” (counting to infinity). Для мінімізації цих проблем у дистанційно-векторних протоколах застосовуються спеціальні механізми: split horizon (заборона анонсування маршруту назад у напрямку, звідки його було отримано), route poisoning (встановлення нескінченної метрики для недоступних мереж), poison reverse (явне анонсування недоступних маршрутів із нескінченною метрикою) та hold-down таймери (тимчасове ігнорування оновлень про маршрут після отримання інформації про його недоступність). Незважаючи на ці механізми, дистанційно-векторні протоколи залишаються менш ефективними для великих та динамічних мереж порівняно з протоколами інших архітектурних типів.

Протоколи стану каналів (link-state protocols) реалізують принципово інший підхід до маршрутизації, базуючись на алгоритмі Дейкстри для пошуку найкоротших шляхів у графі. Замість обміну готовими векторами відстаней, кожен маршрутизатор збирає детальну інформацію про стан усіх каналів зв'язку в мережі та будує повну топологічну карту у вигляді графа, де вершинами є маршрутизатори та мережі, а ребрами – канали зв'язку з відповідними метриками. Процес роботи link-state протоколу можна описати послідовністю етапів: встановлення сусідства з іншими маршрутизаторами через спеціальні Hello-повідомлення, створення пакетів оголошення стану каналів (Link-State Advertisements, LSA), що містять інформацію про безпосередньо підключені мережі та їхні метрики, розповсюдження LSA за принципом flooding всім маршрутизаторам в області маршрутизації, побудова бази даних топології (Link-State Database, LSDB) на основі отриманих LSA від усіх маршрутизаторів, та обчислення найкоротших шляхів до всіх мереж за допомогою алгоритму SPF (Shortest Path First) Дейкстри. Типовими представниками протоколів link-state є OSPF та IS-IS, які широко використовуються в корпоративних та операторських мережах завдяки своїй ефективності та масштабованості.

Перевагами протоколів стану каналів є швидка збіжність після змін топології, відсутність петель маршрутизації (завдяки наявності повної топологічної карти), підтримка різних метрик та можливість гнучкого налаштування вартості каналів, а також ефективне використання пропускну здатності мережі через передачу оновлень лише при зміні топології, а не періодично. До недоліків можна віднести вищу складність реалізації, більші вимоги до процесорної потужності та обсягу пам'яті маршрутизатора (для зберігання повної LSDB та виконання SPF-обчислень), а також потребу в ієрархічному дизайні мережі для забезпечення масштабованості в дуже великих середовищах. Зокрема, OSPF використовує концепцію областей (areas) для логічного поділу мережі, де кожна область має власну LSDB, а SPF-обчислення виконуються незалежно всередині кожної області, що значно зменшує навантаження на маршрутизатори.

Гібридні протоколи поєднують характеристики дистанційно-векторного та link-state підходів, намагаючись використати переваги обох архітектур та мінімізувати їхні недоліки. Найвідомішим представником цього класу є протокол EIGRP компанії Cisco, який формально класифікується як “розширений дистанційно-векторний” протокол, але має багато рис протоколів стану каналів. EIGRP підтримує швидку збіжність завдяки алгоритму DUAL (Diffusing Update Algorithm), який гарантує відсутність петель та дозволяє локально обчислювати альтернативні маршрути без необхідності глобального перерахунку. На відміну від класичних дистанційно-векторних протоколів, EIGRP не надсилає повні таблиці маршрутизації періодично, а передає лише інкрементальні оновлення при зміні топології, що схоже на поведінку link-state протоколів. Протокол використовує складену метрику, що враховує пропускну здатність каналу, затримку, навантаження та надійність, забезпечуючи більш гнучкий вибір оптимальних шляхів порівняно з простим підрахунком переходів у RIP.

У таблиці 1.2 наведено детальне порівняння характеристик різних архітектурних підходів до реалізації протоколів IGP, що дозволяє оцінити їхню придатність для конкретних мережових середовищ.

Таблиця 1.2 – Порівняння архітектурних підходів протоколів IGP

Характеристика	Дистанційно-векторні (RIP)	Link-State (OSPF, IS-IS)	Гібридні (EIGRP)
Алгоритм обчислення маршрутів	Беллмана-Форда	Дейкстри (SPF)	DUAL
Топологічна інформація	Лише наступний перехід та відстань	Повна карта топології	Часткова (сусіди та маршрути)
Частота оновлень	Періодична (30 сек для RIP)	При зміні топології	При зміні топології
Швидкість збіжності	Повільна (хвилини)	Швидка (секунди)	Дуже швидка (секунди)
Використання пропускну здатності	Високе (регулярні повні оновлення)	Низьке (тільки зміни)	Низьке (інкрементальні оновлення)
Вимоги до ресурсів	Низькі	Високі (CPU, RAM)	Середні
Масштабованість	Обмежена (до 15 переходів)	Висока (з ієрархією)	Висока
Підтримка VLSM	RIPv1 – ні, RIPv2 – так	Так	Так
Запобігання петлям	Split horizon, poison reverse, hold-down	Топологічна база даних	Алгоритм DUAL

1.3 Основні вимоги до протоколів внутрішньої маршрутизації

Ефективний протокол внутрішньої маршрутизації повинен задовольняти ряд критичних вимог, що визначають його придатність для використання в сучасних корпоративних та операторських мережах. Першою та найважливішою вимогою є швидка збіжність (convergence) – здатність мережі швидко досягати стану, коли всі маршрутизатори мають узгоджену та коректну інформацію про топологію після будь-яких змін у мережі. Збіжність включає два аспекти: виявлення змін топології (наприклад, відмова каналу або маршрутизатора) та розповсюдження оновленої інформації всім учасникам протоколу маршрутизації. Час збіжності критично впливає на доступність мережесервісів, оскільки під час перехідного періоду можуть виникати втрати пакетів, петлі маршрутизації або недоступність окремих мереж. Сучасні протоколи, такі як OSPF та EIGRP, забезпечують збіжність протягом декількох секунд завдяки використанню ефективних алгоритмів та механізмів швидкого виявлення відмов.

Швидкість збіжності залежить від багатьох факторів, включаючи тип протоколу, розмір мережі, налаштування таймерів, топологію та продуктивність маршрутизаторів. Для прискорення виявлення відмов використовуються механізми безпосереднього контролю стану каналів на фізичному рівні, а також протоколи швидкої детекції відмов, такі як BFD (Bidirectional Forwarding Detection), що дозволяють виявляти проблеми в підсекундному діапазоні. Важливим аспектом є також мінімізація перехідних станів, коли різні маршрутизатори мають несумісні уявлення про топологію – в таких ситуаціях можливе утворення тимчасових петель або чорних дір маршрутизації. Протоколи link-state краще справляються з цим завданням завдяки синхронізації повної топологічної бази даних між усіма маршрутизаторами області, тоді як дистанційно-векторні протоколи потребують додаткових механізмів стабілізації, таких як hold-down таймери.

Другою ключовою вимогою є масштабованість – здатність протоколу ефективно працювати в мережах різного розміру, від невеликих офісних мереж з десятком маршрутизаторів до великих корпоративних або операторських інфраструктур із сотнями та тисячами пристроїв. Масштабованість визначається кількома параметрами: максимальною кількістю маршрутизаторів в одній області маршрутизації, обсягом маршрутної інформації, що може бути оброблений протоколом, споживанням пропускну здатності для обміну маршрутною інформацією, та навантаженням на процесор і пам'ять маршрутизаторів. Протоколи link-state, такі як OSPF, забезпечують високу масштабованість через використання ієрархічного дизайну з поділом мережі на області, де маршрутизатори в одній області знають детальну топологію лише своєї області, а про інші області – лише агреговану інформацію через граничні маршрутизатори (Area Border Routers, ABR). Це дозволяє значно зменшити розмір LSDB та обсяг SPF-обчислень на кожному маршрутизаторі.

Третьою важливою вимогою є підтримка масок змінної довжини (Variable Length Subnet Mask, VLSM) та безкласової адресації (Classless Inter-Domain Routing, CIDR). В епоху дефіциту IPv4-адрес можливість використання підмереж різного розміру всередині однієї мережі класу є критичною для

ефективного використання адресного простору. Протоколи, що підтримують VLSM, передають у своїх оновленнях не тільки мережеву адресу, але й маску підмережі, що дозволяє точно визначити розмір кожної мережі. Старіші версії протоколів, такі як RIPv1 та IGRP, не підтримують VLSM і використовують класову адресацію, що робить їх непридатними для сучасних мереж. Сучасні протоколи RIPv2, OSPF, IS-IS та EIGRP повністю підтримують VLSM та CIDR, дозволяючи гнучко розподіляти адресний простір та виконувати агрегацію маршрутів для зменшення розміру таблиць маршрутизації.

Четвертою вимогою є гнучкість метрики маршрутизації, яка визначає, наскільки точно протокол може відображати характеристики реальних каналів зв'язку при виборі оптимального шляху. Прості протоколи, такі як RIP, використовують примітивну метрику – кількість переходів (hop count), що не враховує пропускну здатність, затримку чи надійність каналів і може призводити до вибору неоптимальних маршрутів. Більш досконалі протоколи використовують складні метрики, що базуються на реальних характеристиках каналів: OSPF використовує вартість (cost), що зазвичай обчислюється як опорна пропускна здатність поділена на пропускну здатність інтерфейсу, EIGRP використовує композитну метрику, що враховує пропускну здатність, затримку, навантаження та надійність каналу. Гнучка метрика дозволяє адміністратору впливати на вибір маршрутів через налаштування вартості каналів відповідно до бізнес-вимог та технічних характеристик інфраструктури.

Рисунок 1.1 схематично ілюструє процес збіжності протоколу маршрутизації після відмови одного з каналів зв'язку в мережі. На діаграмі зображено мережеву топологію з п'ятьма маршрутизаторами (R1, R2, R3, R4, R5), з'єднаними в кільцеву структуру з додатковими резервними зв'язками. У початковому стані всі маршрутизатори мають узгоджені таблиці маршрутизації, що відображені однаковим кольором фону на діаграмі. Після відмови каналу між R2 та R3 (позначено червоним хрестиком) відбувається процес виявлення відмови маршрутизаторами R2 та R3, які детектують втрату сусідства. Далі показано розповсюдження інформації про зміну топології через LSA-повідомлення (для link-state протоколів) або оновлення таблиць (для distance-vector протоколів), що поширюються хвилями від точки відмови. На фінальному етапі всі маршрутизатори перераховують свої таблиці маршрутизації, обираючи альтернативні шляхи через резервні канали, і мережа досягає нового стану збіжності, що знову відображено однаковим кольором фону.

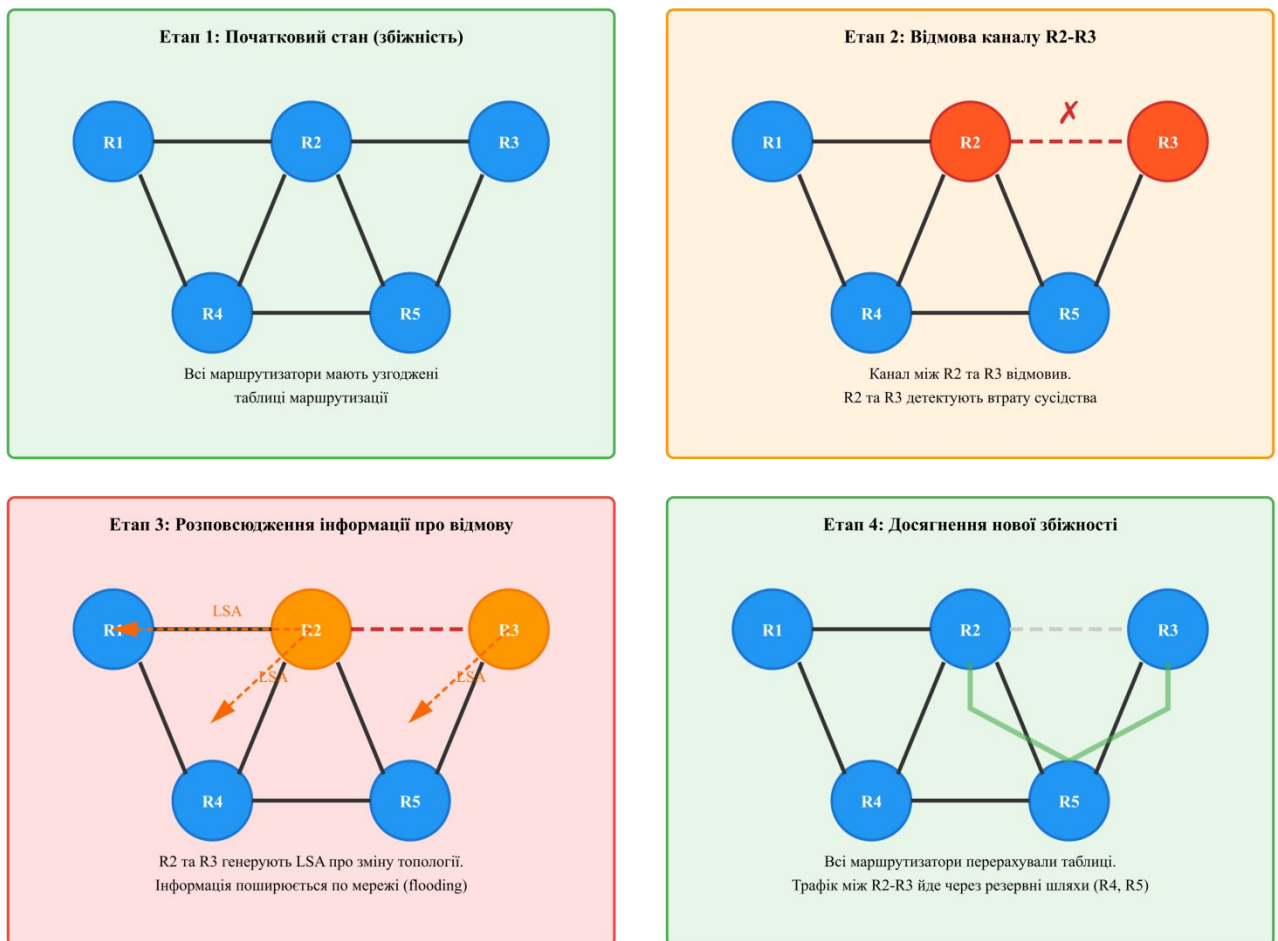


Рисунок 1.1 – Процес збіжності протоколу маршрутизації після відмови каналу

Додатковими важливими вимогами до протоколів IGP є підтримка автентифікації для захисту від несанкціонованого впровадження маршрутної інформації, можливість балансування навантаження між декількома рівноцінними шляхами (Equal-Cost Multi-Path, ECMP), підтримка різних типів мереж (broadcast, point-to-point, NBMA), та зворотна сумісність з попередніми версіями протоколу для плавної міграції. Сучасні реалізації протоколів IGP також включають розширені функції, такі як traffic engineering для оптимізації використання мережевих ресурсів, швидке переключення на резервні маршрути (Fast Reroute), підтримка IPv6 паралельно з IPv4, та інтеграція з системами моніторингу та управління мережею.

1.4 Механізми обміну маршрутною інформацією та запобігання петлям

Ефективний обмін маршрутною інформацією між маршрутизаторами є фундаментальним процесом, що забезпечує функціонування протоколів динамічної маршрутизації. Механізми обміну суттєво відрізняються в залежності від архітектурного типу протоколу, але всі вони мають спільну мету – своєчасно інформувати всі маршрутизатори мережі про доступні шляхи до

мереж призначення та зміни в топології. У дистанційно-векторних протоколах обмін здійснюється шляхом періодичної передачі таблиць маршрутизації або їх частин безпосереднім сусідам, які, в свою чергу, використовують отриману інформацію для оновлення власних таблиць та передачі далі своїм сусідам. Протокол RIP, наприклад, кожні тридцять секунд надсилає всю свою таблицю маршрутизації в широкомовному повідомленні на адресу 255.255.255.255 (RIPv1) або на multicast-адресу 224.0.0.9 (RIPv2), що дозволяє всім маршрутизаторам у сегменті мережі отримати оновлену інформацію.

Протоколи link-state використовують принципово інший підхід до обміну інформацією, що складається з кількох етапів та типів повідомлень. Спочатку маршрутизатори встановлюють сусідські відносини (adjacency) через обмін Hello-пакетами, що надсилаються періодично на спеціальні multicast-адреси для виявлення інших маршрутизаторів у мережі та підтримки інформації про їхню доступність. Після встановлення сусідства відбувається синхронізація баз даних топології через обмін описами LSA (Database Description packets), що містять список усіх записів у LSDB. Якщо виявляється, що сусід має LSA, яких немає в локальній базі, або більш свіжі версії існуючих записів, маршрутизатор запитує конкретні LSA через Link-State Request пакети та отримує їх у відповідь. Після початкової синхронізації подальші оновлення надсилаються лише при зміні топології – коли змінюється стан каналу, маршрутизатор генерує новий LSA та розповсюджує його всім маршрутизаторам області через механізм flooding.

Механізм flooding забезпечує надійне та швидке розповсюдження топологічної інформації по всій області маршрутизації. Коли маршрутизатор отримує LSA від одного зі своїх сусідів, він порівнює номер послідовності отриманого LSA з тим, що зберігається в його базі даних – якщо отриманий LSA новіший, маршрутизатор зберігає його в базі та негайно пересилає копію всім своїм сусідам, окрім того, від якого отримав цей LSA. Цей процес повторюється на кожному маршрутизаторі, забезпечуючи швидке поширення інформації по всій мережі. Для запобігання нескінченній циркуляції LSA використовуються номери послідовності та час життя (Age) кожного запису – маршрутизатор не пересилає LSA, який він уже бачив (має такий самий або новіший номер послідовності в базі), а застарілі записи автоматично видаляються після закінчення часу життя. OSPF використовує 32-бітні номери послідовності та максимальний час життя LSA в одну годину, після чого запис повинен бути оновлений або видалений з бази.

Важливою проблемою протоколів маршрутизації є можливість виникнення петель маршрутизації – ситуацій, коли пакет циркулює між двома або більше маршрутизаторами, ніколи не досягаючи призначення, що призводить до марного споживання пропускну здатності та ресурсів обладнання. Петлі виникають у перехідних станах, коли різні маршрутизатори мають неузгоджену інформацію про топологію, особливо після відмови каналу чи маршрутизатора. Уявімо ситуацію, коли маршрутизатор R1 безпосередньо підключений до мережі 192.168.1.0/24, маршрутизатор R2 має маршрут до цієї мережі через R1, а маршрутизатор R3 – через R2. Після відмови з'єднання між

R1 та мережею, R1 швидко видаляє цей маршрут зі своєї таблиці, але інформація про відмову ще не досягла R2 та R3. Якщо R1 отримує пакет для 192.168.1.0/24 та не має маршруту, він може запитати R2, який все ще має застарілий маршрут через R1, що створює петлю між R1 та R2.

Для запобігання петлям дистанційно-векторні протоколи використовують комплекс механізмів, кожен з яких вирішує певний аспект проблеми. Механізм split horizon забороняє маршрутизатору анонсувати маршрут назад в інтерліфейс, через який цей маршрут було отримано – це запобігає прямим петлям між двома маршрутизаторами. Розширення цього механізму, split horizon with poison reverse, передбачає явне анонсування маршруту назад із нескінченною метрикою (16 для RIP), що прискорює видалення недійсних маршрутів. Механізм route poisoning встановлює метрику недоступної мережі в нескінченність і негайно розповсюджує цю інформацію, сигналізуючи всім маршрутизаторам про необхідність видалення маршруту. Hold-down таймери запобігають прийняттю нових маршрутів до того ж призначення протягом певного часу після отримання інформації про недоступність – це дає час інформації про відмову розповсюдитися по всій мережі до того, як будуть прийняті потенційно некоректні альтернативні маршрути.

Протоколи link-state мають природний захист від петель завдяки наявності повної топологічної карти в кожного маршрутизатора. Алгоритм Дейкстри будує дерево найкоротших шляхів від поточного маршрутизатора до всіх мереж, і за визначенням в дереві не може бути циклів. Навіть у перехідні періоди після зміни топології, коли різні маршрутизатори можуть мати різні версії LSDB, петлі є короткочасними та швидко усуваються після синхронізації баз даних. Додатковим механізмом захисту є зменшення значення TTL (Time To Live) в IP-заголовку кожного пакета при проходженні через маршрутизатор – якщо петля все ж виникає, пакет буде відкинутий після певної кількості переходів. EIGRP використовує алгоритм DUAL, який математично гарантує відсутність петель через концепцію Feasibility Condition – маршрутизатор приймає новий маршрут тільки якщо звітувана відстань (reported distance) від наступного переходу менша за поточну найкращу відстань (feasible distance).

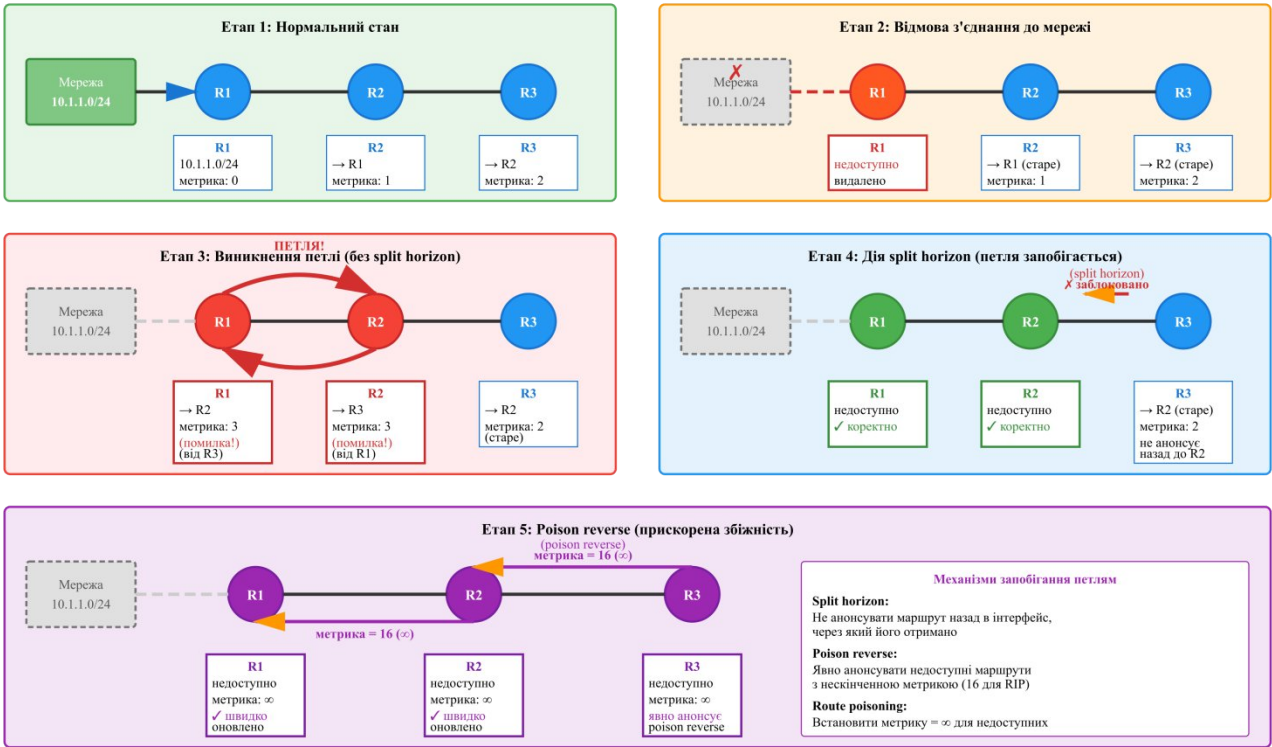
У таблиці 1.3 узагальнено основні механізми запобігання петлям у різних типах протоколів IGP, що демонструє еволюцію підходів від простих таймерів до складних алгоритмічних гарантій.

Таблиця 1.3 – Механізми запобігання петлям маршрутизації

П протокол	Основні механізми	Час виявлення петлі	Ефективність
RI P	Split horizon, poison reverse, hold-down таймери, TTL	Хвили ни	Низька (можливі тимчасові петлі)
O SPF	Повна топологічна база даних, алгоритм Дейкстри, синхронізація LSDB	Секун ди	Висока (петлі дуже рідкісні)
EI GRP	Алгоритм DUAL, Feasibility Condition, запити та відповіді	Секун ди	Дуже висока (гарантована відсутність)

			петель)
-IS	IS	Повна топологічна база даних, алгоритм SPF, flooding LSP	Секунди
			Висока (аналогічно OSPF)

Рисунок 1.2 ілюструє приклад виникнення петлі маршрутизації в дистанційно-векторному протоколі та її усунення через механізми split horizon та poison reverse. На діаграмі показано три маршрутизатори R1, R2, R3, з'єднані послідовно, де R1 безпосередньо підключений до мережі 10.1.1.0/24. У нормальному стані (етап 1) таблиці маршрутизації показують, що R1 має прямий маршрут до 10.1.1.0/24 з метрикою 0, R2 має маршрут через R1 з метрикою 1, а R3 – через R2 з метрикою 2. На етапі 2 показано момент відмови з'єднання між R1 та мережею – R1 видаляє маршрут, але R2 та R3 ще мають старі записи. Якщо R2 оновлює свою таблицю на основі інформації від R3 без split horizon, виникає петля (етап 3), де R2 думає, що може досягти мережі через R3 з метрикою 3, а R3 – через R2 з метрикою 2. На етапі 4 показано, як механізм split horizon запобігає цій ситуації – R3 не анонсує маршрут назад до R2, тому R2 коректно визначає мережу як недоступну. На етапі 5 демонструється poison reverse, де R3 явно інформує R2 про недоступність мережі (метрика 16), прискорюючи збіжність.



Порівняння результатів			
Сценарій	Результат	Час збіжності	Коментар
Без захисних механізмів	Петля маршрутизації (counting to infinity)	Дуже повільно (хвилини)	Пакети циркулюють між R1 та R2 до досягнення макс. метрики
Зі split horizon	Петлі запобіжено коректне видалення	Повільно (до 180 сек)	R3 не надсилає маршрут назад, запобігає помилковому оновленню
З poison reverse	Петлі запобіжено швидке видалення	Швидко (30-60 сек)	Явне оголошення метрики ∞ прискорює розповсюдження інформації
Link-state (OSPF, IS-IS)	Природний захист	Дуже швидко (1-5 сек)	Повна топологічна база даних, алгоритм Дейкстри

Рисунок 1.2 – Виникнення петлі маршрутизації та механізми її запобігання

Сучасні реалізації протоколів IGP включають додаткові механізми підвищення стабільності та швидкодії. Протокол BFD (Bidirectional Forwarding Detection) забезпечує швидке виявлення відмов каналів на рівні від мілісекунд до секунд, незалежно від таймерів самого протоколу маршрутизації, що дозволяє досягти часу збіжності менше секунди. Механізм Graceful Restart дозволяє маршрутизатору продовжувати пересилання пакетів навіть під час перезавантаження процесу маршрутизації, зберігаючи forwarding table та інформуючи сусідів про тимчасову недоступність control plane. Технологія Loop-Free Alternate (LFA) та Remote LFA в OSPF та IS-IS дозволяють попередньо обчислити резервні маршрути, які гарантовано не створюють петель, та миттєво перемикаються на них при відмові основного шляху без очікування збіжності протоколу маршрутизації.

ТЕМА 2. ПРОТОКОЛИ ЗОВНІШНЬОЇ ДИНАМІЧНОЇ МАРШРУТИЗАЦІЇ (EGP)

2.1 Архітектура міжавтономної маршрутизації та концепція автономних систем

Глобальна мережа Інтернет являє собою складну ієрархічну структуру, яка об'єднує тисячі незалежних мереж, що належать різним організаціям, провайдерам та корпораціям. Кожна така мережа має власну адміністративну політику, технічні вимоги та бізнес-цілі, що робить неможливим застосування єдиного централізованого механізму маршрутизації для всього Інтернету. Саме тому було запроваджено концепцію автономних систем (Autonomous System, AS), яка дозволяє розділити глобальний простір маршрутизації на керовані сегменти зі збереженням можливості їх взаємодії.

Автономна система визначається як набір IP-мереж та маршрутизаторів, які перебувають під єдиним адміністративним управлінням і використовують спільну політику маршрутизації для зовнішнього світу. Всередині автономної системи можуть застосовуватися різноманітні протоколи внутрішньої маршрутизації (IGP), такі як OSPF, EIGRP або RIP, проте для зовнішнього світу вся AS виглядає як єдина логічна одиниця з унікальним ідентифікатором. Цей ідентифікатор, відомий як номер автономної системи (ASN), призначається регіональними інтернет-реєстрами (RIR) і є обов'язковим для участі в глобальній системі маршрутизації. Спочатку ASN були 16-бітними числами, що дозволяло мати до 65536 унікальних номерів, але зростання Інтернету призвело до розширення простору ASN до 32-бітних значень, що забезпечує понад 4 мільярди можливих номерів.

Міжавтономна маршрутизація суттєво відрізняється від внутрішньої маршрутизації як за цілями, так і за механізмами реалізації. Якщо протоколи IGP зосереджені на пошуку оптимальних шляхів всередині однорідного адміністративного домену з використанням технічних метрик, то протоколи зовнішньої маршрутизації (EGP) мають враховувати політичні, економічні та організаційні аспекти взаємодії між різними автономними системами. Основним завданням EGP є не стільки знаходження найкоротшого шляху з технічної точки зору, скільки забезпечення дотримання угод про транзит трафіку, реалізація бізнес-політик маршрутизації та контроль над розповсюдженням інформації про мережі. Це призводить до того, що рішення про вибір маршруту в EGP приймаються на основі складних наборів атрибутів та політик, а не простих числових метрик.

Історично перший протокол зовнішньої маршрутизації, також названий EGP (Exterior Gateway Protocol), був розроблений у 1980-х роках, але швидко виявив свої обмеження у зростаючому Інтернеті. Головною проблемою цього протоколу була відсутність механізмів запобігання петлям маршрутизації та підтримки складних топологій з множинними шляхами між автономними системами. На зміну йому прийшов протокол Border Gateway Protocol (BGP), який був вперше стандартизований у 1989 році і з того часу пройшов кілька

ревізій, при цьому сучасна версія BGP-4 залишається єдиним протоколом зовнішньої маршрутизації, який використовується в Інтернеті. BGP забезпечує необхідну гнучкість для реалізації складних політик маршрутизації, підтримує безкласову адресацію (CIDR) та має масштабованість, достатню для обробки сотень тисяч маршрутів у глобальній таблиці маршрутизації Інтернету.

2.2 Протокол BGP: основні принципи роботи та типи сесій

Border Gateway Protocol є path-vector протоколом, що означає передачу не просто інформації про доступність мережі, а повного шляху автономних систем, через які проходить маршрут до цієї мережі. Цей підхід принципово відрізняється від distance-vector протоколів, які передають лише інформацію про відстань до мережі, та від link-state протоколів, які будують повну топологічну карту. Передача повного AS-шляху дозволяє маршрутизатору самостійно виявляти потенційні петлі маршрутизації, оскільки якщо в отриманому анонсі міститься номер власної автономної системи, це однозначно вказує на наявність петлі і такий маршрут відкидається.

BGP функціонує на основі встановлення довгострокових TCP-з'єднань між сусідніми маршрутизаторами, які називаються BGP-пірами або сусідами. Використання TCP як транспортного протоколу (порт 179) забезпечує надійну доставку повідомлень маршрутизації, автоматичне керування потоком даних та виявлення обривів зв'язку між сусідами. Після встановлення TCP-з'єднання BGP-сусіди обмінюються повідомленнями OPEN, які містять параметри сесії, включаючи номер автономної системи, BGP-ідентифікатор маршрутизатора та час утримання сесії. Успішний обмін повідомленнями OPEN переводить сесію в стан Established, після чого починається обмін повною таблицею маршрутизації через повідомлення UPDATE. Надалі сусіди підтримують сесію живою шляхом періодичного обміну невеликими повідомленнями KEEPALIVE та передають лише інкрементальні зміни в маршрутизації, що значно зменшує навантаження на мережу та процесор маршрутизатора.

Розрізняють два основні типи BGP-сесій: зовнішні (eBGP) та внутрішні (iBGP), які мають суттєві відмінності в поведінці та призначенні. Сесії eBGP встановлюються між маршрутизаторами, які належать до різних автономних систем, зазвичай прямо з'єднаних один з одним на фізичному або каналному рівні. При передачі маршрутів через eBGP-сесію маршрутизатор додає номер своєї автономної системи до атрибуту AS_PATH, що дозволяє відстежувати повний шлях проходження анонсу через різні AS. За замовчуванням, eBGP-сесії встановлюються між безпосередньо підключеними інтерфейсами, а значення TTL в IP-пакетах встановлюється в 1, що є додатковим захисним механізмом проти випадкового встановлення сесій з віддаленими маршрутизаторами. Для сесій eBGP також типовою є зміна атрибуту NEXT_HOP на адресу власного інтерфейсу при передачі маршруту сусідній AS, що забезпечує коректну маршрутизацію трафіку.

На відміну від eBGP, сесії iBGP встановлюються між маршрутизаторами всередині однієї автономної системи для розповсюдження інформації про

зовнішні маршрути між граничними маршрутизаторами. Важливою особливістю iBGP є правило split-horizon, яке забороняє маршрутизатору передавати маршрути, отримані від одного iBGP-сусіда, іншим iBGP-сусідам. Це правило запроваджено для запобігання петлям маршрутизації всередині AS, але має важливий наслідок: всі маршрутизатори, які беруть участь у BGP всередині автономної системи, повинні мати повнозв'язну топологію iBGP-сесій, тобто кожен маршрутизатор повинен мати сесії з усіма іншими. В великих мережах це призводить до проблеми масштабованості, оскільки кількість необхідних сесій зростає квадратично з кількістю маршрутизаторів, що вирішується через використання механізмів Route Reflector або AS Confederation. При передачі маршрутів через iBGP атрибут AS_PATH не змінюється, оскільки анонс залишається всередині тієї самої автономної системи, а атрибут NEXT_HOP зазвичай зберігається без змін, що вимагає забезпечення доступності адреси NEXT_HOP через протоколи IGP.

Таблиця 2.1 - Порівняння характеристик eBGP та iBGP сесій

Характеристика	eBGP	iBGP
Розташування сусідів	Різні автономні системи	Одна автономна система
Типова топологія	Прямо підключені інтерфейси	Логічні з'єднання через IGP
Значення TTL за замовчуванням	1 (один хоп)	255 (множинні хопи)
Зміна атрибуту AS_PATH	Додається власний ASN	Не змінюється
Зміна атрибуту NEXT_HOP	Змінюється на власну адресу	Зберігається оригінальне значення
Правило split-horizon	Не застосовується	Застосовується (потрібна full mesh)
Administrative Distance	20	200

2.3 Атрибути BGP-маршрутів та механізм вибору найкращого шляху

Кожен маршрут у BGP супроводжується набором атрибутів, які описують різноманітні характеристики цього маршруту та використовуються для прийняття рішень про вибір найкращого шляху та застосування політик маршрутизації. Атрибути класифікуються на обов'язкові, які повинні присутні в кожному анонсі (well-known mandatory), необов'язкові, які розпізнаються всіма реалізаціями BGP але можуть бути відсутні (well-known discretionary), та опціональні, які не обов'язково підтримуються всіма маршрутизаторами (optional). Серед найважливіших атрибутів виділяються AS_PATH, NEXT_HOP, LOCAL_PREF, MED та ORIGIN, кожен з яких відіграє специфічну роль у процесі маршрутизації.

Атрибут AS_PATH містить упорядкований список номерів автономних систем, через які пройшов даний маршрутний анонс від джерела до поточного маршрутизатора. Цей атрибут є фундаментальним для BGP, оскільки виконує

дві критичні функції: запобігання петлям маршрутизації та участь у виборі найкращого шляху. Кожен маршрутизатор, який передає анонс через eBGP-сесію, додає номер своєї AS на початок списку AS_PATH, таким чином при отриманні маршруту можна прослідкувати весь його шлях. Якщо маршрутизатор отримує анонс, в AS_PATH якого міститься номер його власної автономної системи, такий маршрут автоматично відкидається як такий, що містить петлю. При виборі між декількома маршрутами до однієї мережі BGP віддає перевагу маршруту з найкоротшим AS_PATH, виходячи з припущення, що менша кількість транзитних автономних систем означає кращу якість маршруту. Проте адміністратори можуть маніпулювати вибором шляхів через механізм AS_PATH prepending, коли до списку штучно додаються додаткові копії власного номера AS для зменшення привабливості певного маршруту.

Атрибут NEXT_HOP вказує IP-адресу наступного маршрутизатора на шляху до мережі призначення, через який слід пересилати пакети для досягнення анонсованої мережі. Для маршрутів, отриманих через eBGP, NEXT_HOP зазвичай встановлюється в адресу сусіднього eBGP-маршрутизатора на безпосередньо підключеному інтерфейсі. Важливою особливістю є те, що при розповсюдженні маршруту через iBGP атрибут NEXT_HOP за замовчуванням не змінюється і продовжує вказувати на оригінальний eBGP-сусід, навіть якщо цей маршрутизатор знаходиться за межами прямого підключення. Це означає, що всі маршрутизатори всередині автономної системи повинні мати можливість досягти адреси NEXT_HOP через протокол внутрішньої маршрутизації (IGP), інакше маршрут буде вважатися недосяжним незважаючи на його присутність в BGP-таблиці. У деяких сценаріях, наприклад при використанні Route Reflector або при наявності проміжних мереж, може знадобитися примусова зміна NEXT_HOP через застосування команди next-hop-self, яка змушує маршрутизатор встановлювати власну адресу як NEXT_HOP для маршрутів, що передаються iBGP-сусідам.

Атрибут LOCAL_PREF (Local Preference) є well-known discretionary атрибутом, який використовується виключно всередині автономної системи для індикації переваги виходу трафіку з AS через певний граничний маршрутизатор. Цей атрибут передається лише через iBGP-сесії і не виходить за межі автономної системи, що робить його ідеальним інструментом для реалізації внутрішньої політики вибору шляхів без впливу на зовнішні AS. Більше значення LOCAL_PREF вказує на вищий пріоритет маршруту, і за замовчуванням всі маршрути отримують значення 100. Адміністратор може встановлювати різні значення LOCAL_PREF для маршрутів, отриманих від різних провайдерів або через різні канали зв'язку, таким чином контролюючи вихідний трафік. Наприклад, якщо організація має два підключення до Інтернету через різних провайдерів, можна встановити вище значення LOCAL_PREF для маршрутів від основного провайдера, що зробить цей канал переважним для виходу трафіку, зберігаючи другий канал як резервний.

Атрибут MED (Multi-Exit Discriminator), також відомий як метрика BGP, використовується для впливу на вхідний трафік у випадках, коли до сусідньої автономної системи існує кілька точок підключення. На відміну від

LOCAL_PREF, який контролює вихідний трафік і діє локально, MED передається в сусідню AS і підказує їй, який з кількох граничних маршрутизаторів є переважним для доставки трафіку. Менше значення MED вказує на вищу перевагу маршруту, і за замовчуванням всі маршрути мають MED рівний 0. Важливою особливістю MED є те, що він порівнюється лише для маршрутів, отриманих від однієї автономної системи, і не передається далі третім AS. Це обмеження означає, що MED є рекомендацією, а не вимогою, і сусідня AS може ігнорувати його або застосовувати власні політики. Типовим сценарієм використання MED є ситуація, коли організація має два канали до одного провайдера в різних географічних точках: встановлюючи нижче значення MED на маршрутах, анонсованих через географічно ближчу точку підключення, можна підказати провайдеру переважний шлях доставки трафіку.

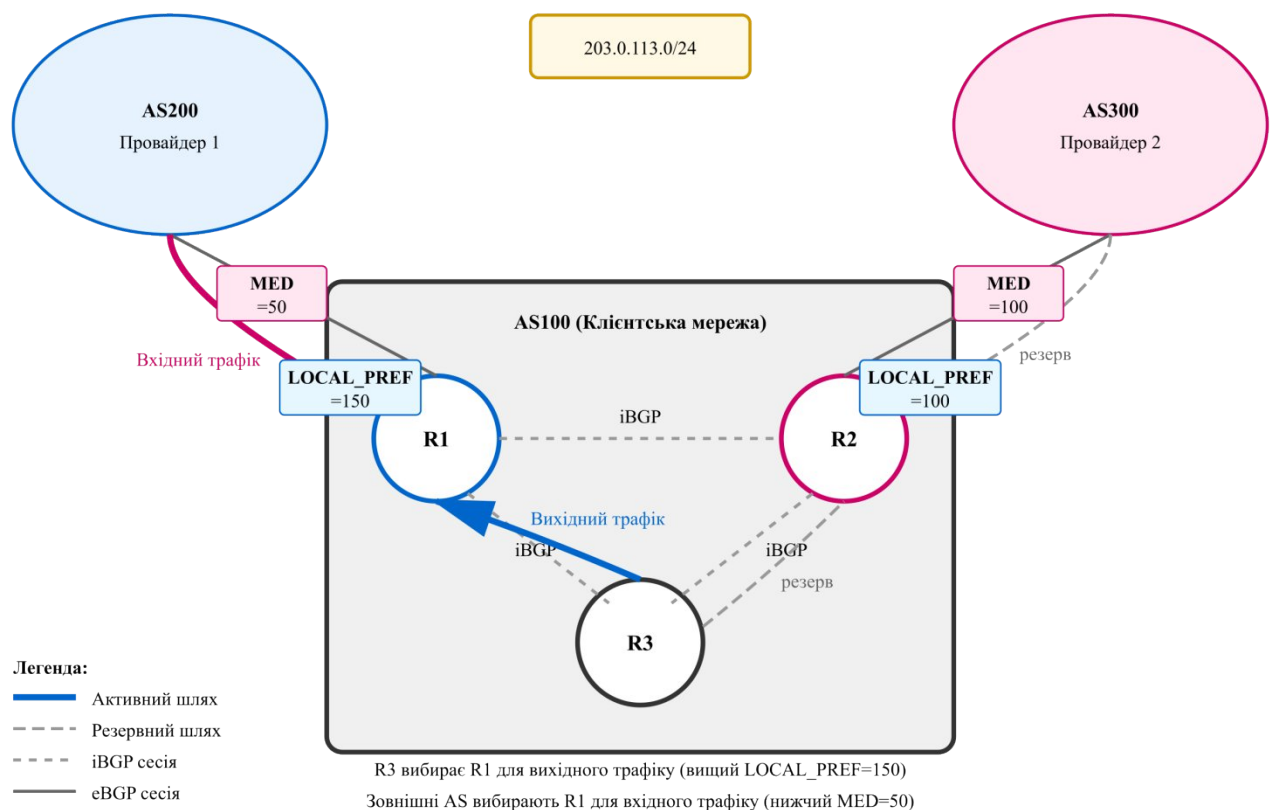


Рисунок 2.1 - Схема використання атрибутів LOCAL_PREF та MED для керування трафіком

На рисунку зображена автономна система AS100, яка має два граничні маршрутизатори R1 та R2, підключені до різних провайдерів AS200 та AS300 відповідно. Внутрішній маршрутизатор R3 отримує маршрути до зовнішньої мережі 203.0.113.0/24 через обидва граничні маршрутизатори. R1 встановлює LOCAL_PREF=150 для маршрутів від AS200, а R2 встановлює LOCAL_PREF=100 для маршрутів від AS300. Це призводить до того, що R3 вибирає маршрут через R1 для вихідного трафіку. Для вхідного трафіку R1 анонсує мережі AS100 з MED=50, а R2 з MED=100, що підказує зовнішнім системам використовувати R1 як переважну точку входу. Стрілки на схемі

показують напрямок потоків трафіку, товсті лінії позначають активні шляхи, а пунктирні - резервні.

Процес вибору найкращого маршруту в BGP є детермінованим алгоритмом, який послідовно застосовує набір правил для порівняння всіх доступних маршрутів до конкретної мережі призначення. BGP-маршрутизатор зберігає всі отримані маршрути в таблиці BGP (Adj-RIB-In), але в таблицю маршрутизації (RIB) встановлює лише один найкращий маршрут, який потім може бути переданий сусідам. Алгоритм вибору найкращого шляху BGP включає наступні кроки, які застосовуються послідовно до виключення всіх маршрутів крім одного: перевірка доступності NEXT_HOP, перевага вищого значення LOCAL_PREF, перевага локально згенерованих маршрутів, перевага коротшого AS_PATH, перевага кращого ORIGIN (IGP краще за EGP, EGP краще за Incomplete), перевага нижчого MED, перевага eBGP над iBGP, перевага маршруту через найближчий IGP-сусід до NEXT_HOP, і нарешті, якщо всі попередні критерії однакові, вибір маршруту з найнижчим BGP Router ID. Цей багатоступеневий процес забезпечує стабільний та передбачуваний вибір маршрутів при збереженні гнучкості для адміністраторів через налаштування атрибутів.

2.4 Політики маршрутизації, фільтрація та агрегація маршрутів у BGP

Одним з найважливіших аспектів протоколу BGP є можливість реалізації складних політик маршрутизації, які дозволяють автономним системам контролювати прийом, відбір та розповсюдження маршрутної інформації відповідно до технічних, організаційних та комерційних вимог. На відміну від протоколів IGP, які зазвичай прагнуть до досягнення оптимального шляху за єдиною метрикою, BGP надає адміністраторам повний контроль над процесом прийняття рішень через систему політик. Політики маршрутизації реалізуються за допомогою механізмів фільтрації та модифікації атрибутів маршрутів, які застосовуються як до вхідних (inbound), так і до вихідних (outbound) оновлень на кожній BGP-сесії. Правильно сконфігуровані політики є критично важливими не лише для оптимізації потоків трафіку, але й для забезпечення безпеки мережі та відповідності комерційним угодам між провайдерами та клієнтами.

Фільтрація маршрутів у BGP виконується за допомогою списків доступу (access-list), префіксних списків (prefix-list) або route-map, які дозволяють визначати умови для прийняття чи відхилення маршрутів на основі їх префіксів, довжини маски підмережі, значень атрибутів або комбінації цих параметрів. Найпростіша форма фільтрації використовує prefix-list для дозволу або заборони анонсування конкретних діапазонів адрес. Наприклад, провайдер може фільтрувати анонси від клієнта, приймаючи лише ті префікси, які належать виділеному клієнту адресному простору, відкидаючи будь-які інші анонси для захисту від помилкової конфігурації або зловмисних дій. Більш складні сценарії фільтрації можуть включати перевірку довжини AS_PATH для

виявлення аномально довгих шляхів, які можуть вказувати на петлі або атаки route hijacking, або фільтрацію на основі наявності певних автономних систем в AS_PATH для реалізації політик транзиту. Фільтрація застосовується як до вхідних маршрутів для контролю того, що приймається від сусідів, так і до вихідних для контролю того, які маршрути анонсуються зовні, що дає повний контроль над потоком маршрутної інформації.

Route-map є найпотужнішим інструментом для реалізації політик BGP, оскільки дозволяє не лише фільтрувати маршрути, але й модифікувати їх атрибути умовно на основі різноманітних критеріїв. Route-map складається з послідовності записів, кожен з яких містить умови відповідності (match) та дії (set), аналогічно до конструкції if-then в програмуванні. Умови відповідності можуть перевіряти префікс мережі, довжину маски, значення AS_PATH, присутність певних community-тегів та багато інших параметрів. Якщо всі умови в записі виконуються, застосовуються визначені дії, які можуть включати встановлення значень атрибутів LOCAL_PREF, MED, AS_PATH prepending, модифікацію NEXT_HOP або додавання community-тегів. Якщо жоден запис у route-map не відповідає маршруту, за замовчуванням маршрут відхиляється, що відповідає неявному deny в кінці кожної route-map. Ця поведінка вимагає явного дозволу всіх бажаних маршрутів, що підвищує безпеку але потребує ретельного планування.

BGP community є опціональним транзитивним атрибутом, який дозволяє групувати маршрути з подібними характеристиками для спрощення застосування політик. Community представляє собою 32-бітне значення, яке зазвичай записується у форматі ASN:NN, де перша частина вказує номер автономної системи, а друга є довільним числом, визначеним адміністратором. Маршрути можуть мати множинні community-теги, і ці теги передаються між автономними системами, дозволяючи реалізувати складні розподілені політики маршрутизації. Типовими застосуваннями community є маркування географічного походження маршрутів, індикація типу підключення клієнта або сигналізація бажаного способу обробки маршруту в транзитних AS. Деякі community мають зарезервоване значення і розпізнаються всіма реалізаціями BGP: NO_EXPORT вказує, що маршрут не повинен анонсуватися за межі AS, NO_ADVERTISE забороняє анонсування маршруту будь-яким сусідам, а NO_EXPORT_SUBCONFED обмежує розповсюдження в межах AS-конфедерації.

Таблиця 2.2 - Типові сценарії використання BGP-політик

Сценарій	Застосовуваний механізм	Мета
Запобігання анонсуванню приватних адрес	Prefix-list на вихідних оновленнях	Безпека: блокування RFC1918 адрес
Контроль клієнтських анонсів	Prefix-list з точним переліком дозволених мереж	Захист від hijacking та помилкової конфігурації
Вибір переважного	Route-map з	Оптимізація вихідного трафіку

провайдера	встановленням LOCAL_PREF	
Балансування вхідного трафіку	Route-map з модифікацією MED	Розподіл навантаження між каналами
Обмеження глибини транзиту	AS_PATH filter за довжиною	Запобігання петлям та атакам
Маркування типу клієнта	Встановлення BGP community	Спрощення політик на магістральних маршрутизаторах

Агрегація маршрутів є критичним механізмом для підтримання масштабованості глобальної таблиці маршрутизації Інтернету, оскільки дозволяє об'єднувати множину більш специфічних префіксів в один суммарний маршрут, зменшуючи кількість записів в таблиці BGP. Без агрегації кожна невелика підмережа потребувала б окремого запису, що призвело б до надмірного зростання таблиць маршрутизації та виснаження пам'яті маршрутизаторів. BGP підтримує автоматичну агрегацію через команду `aggregate-address`, яка створює суммарний маршрут для заданого префіксу, якщо в таблиці BGP присутній хоча б один більш специфічний компонент цього суммарного маршруту. За замовчуванням агрегований маршрут анонсується одночасно з компонентами, але опція `summary-only` дозволяє придушити анонсування конкретних префіксів, залишаючи лише агрегат. При створенні агрегованого маршруту AS_PATH формується як об'єднання всіх унікальних номерів AS з компонентних маршрутів, що може призводити до втрати інформації про точний шлях, проте атрибут AS_SET зберігає цю інформацію для запобігання петлям. Адміністратори повинні обережно підходити до агрегації, оскільки надмірна агрегація може призвести до неоптимальної маршрутизації або до анонсування адресного простору, частина якого насправді не використовується.

Особливу роль у сучасному Інтернеті відіграють системи перевірки походження маршрутів, такі як RPKI (Resource Public Key Infrastructure), які спрямовані на підвищення безпеки BGP через криптографічну верифікацію права автономної системи анонсувати певний адресний простір. RPKI використовує ієрархію сертифікатів, видану регіональними інтернет-реєстрами, для створення ROA (Route Origin Authorization) об'єктів, які однозначно вказують, яка AS має право анонсувати певний префікс та максимальну дозволену довжину маски. Маршрутизатори, які підтримують RPKI, можуть перевіряти отримані BGP-анонси на відповідність ROA, класифікуючи маршрути як Valid (відповідає існуючому ROA), Invalid (суперечить ROA або перевищує максимальну довжину), або Unknown (немає відповідного ROA). На основі цієї класифікації можна застосовувати політики, наприклад, повністю відкидаючи Invalid маршрути або знижуючи їх LOCAL_PREF для зменшення ймовірності вибору, що значно підвищує захищеність від атак route hijacking та випадкових помилок конфігурації.

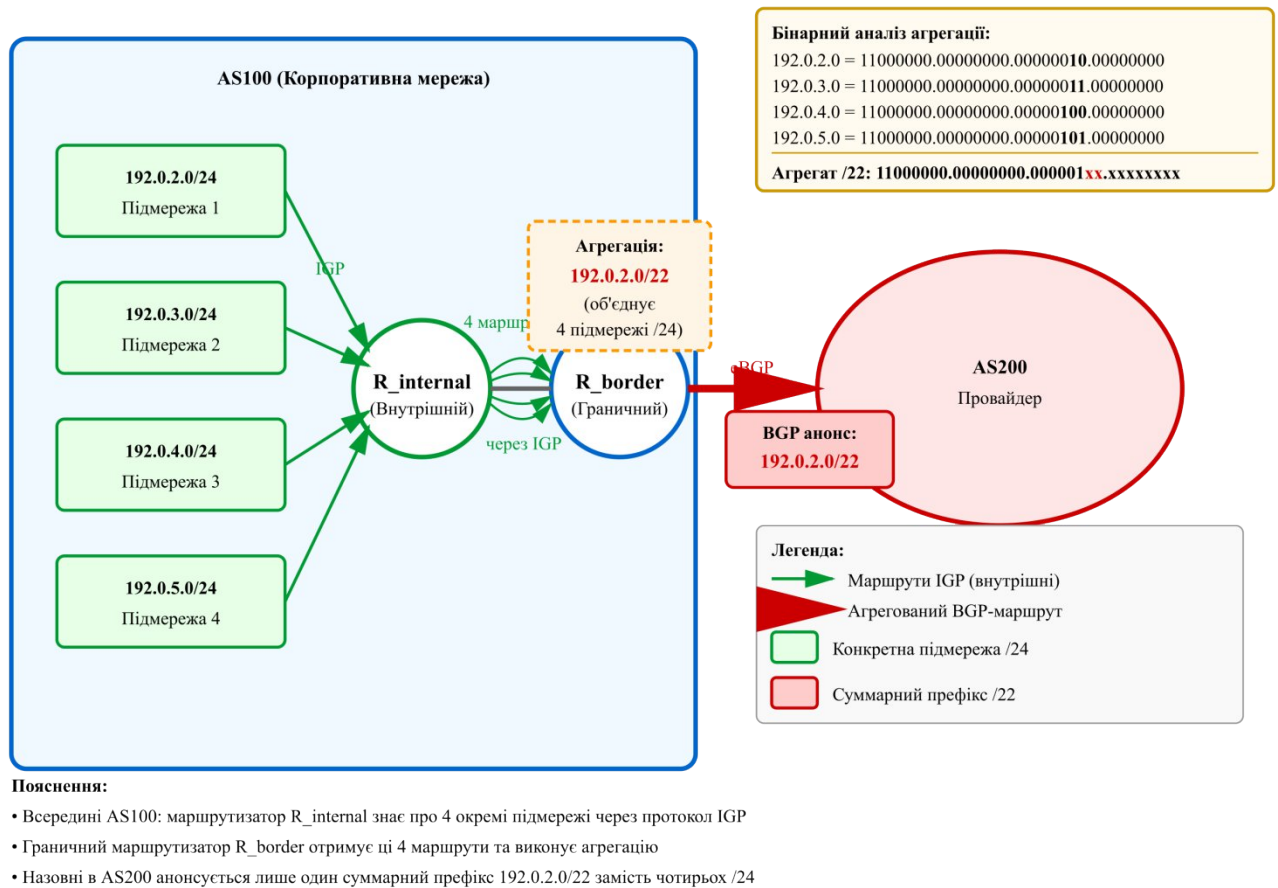


Рисунок 2.2 - Приклад агрегації маршрутів у BGP

На рисунку показана автономна система AS100, яка має внутрішню адресну структуру з чотирма підмережами класу C: 192.0.2.0/24, 192.0.3.0/24, 192.0.4.0/24 та 192.0.5.0/24. Внутрішній маршрутизатор R_internal отримує всі чотири маршрути через протокол IGP та встановлює їх в свою таблицю маршрутизації. Граничний маршрутизатор R_border, який підключений до зовнішнього провайдера AS200, налаштований на агрегацію цих чотирьох /24 префіксів в один суммарний /22 префікс: 192.0.2.0/22. Замість анонсування чотирьох окремих маршрутів в BGP, R_border анонсує лише один агрегований префікс, зменшуючи навантаження на таблицю маршрутизації провайдера. На схемі суцільними стрілками показані конкретні маршрути всередині AS, а товстою стрілкою - агрегований анонс назовні. Така практика є стандартною для корпоративних мереж та провайдерів, оскільки дозволяє зменшити розмір глобальної таблиці маршрутизації Інтернету.

Успішна реалізація політик BGP вимагає ретельного планування та тестування, оскільки помилки в конфігурації можуть призвести до серйозних наслідків, включаючи повну втрату зв'язності, створення чорних дір для трафіку або ненавмисне анонсування чужого адресного простору. Рекомендується документувати всі політики маршрутизації, використовувати послідовні схеми нумерації для community-тегів, регулярно перевіряти таблиці BGP на наявність аномалій та застосовувати максимальні обмеження (max-prefix) на кількість маршрутів, прийнятих від кожного сусіда, для захисту від випадкових або зловмисних флудів маршрутною інформацією. Моніторинг та

логування змін в BGP-сесіях також є критично важливими для швидкого виявлення та реагування на інциденти, пов'язані з маршрутизацією.

ТЕМА 2. ПРОТОКОЛИ ЗОВНІШНЬОЇ ДИНАМІЧНОЇ МАРШРУТИЗАЦІЇ (EGP)

2.1 Архітектура міжавтономної маршрутизації та концепція автономних систем

Глобальна мережа Інтернет являє собою складну ієрархічну структуру, яка об'єднує тисячі незалежних мереж, що належать різним організаціям, провайдерам та корпораціям. Кожна така мережа має власну адміністративну політику, технічні вимоги та бізнес-цілі, що робить неможливим застосування єдиного централізованого механізму маршрутизації для всього Інтернету. Саме тому було запроваджено концепцію автономних систем (Autonomous System, AS), яка дозволяє розділити глобальний простір маршрутизації на керовані сегменти зі збереженням можливості їх взаємодії.

Автономна система визначається як набір IP-мереж та маршрутизаторів, які перебувають під єдиним адміністративним управлінням і використовують спільну політику маршрутизації для зовнішнього світу. Всередині автономної системи можуть застосовуватися різноманітні протоколи внутрішньої маршрутизації (IGP), такі як OSPF, EIGRP або RIP, проте для зовнішнього світу вся AS виглядає як єдина логічна одиниця з унікальним ідентифікатором. Цей ідентифікатор, відомий як номер автономної системи (ASN), призначається регіональними інтернет-реєстрами (RIR) і є обов'язковим для участі в глобальній системі маршрутизації. Спочатку ASN були 16-бітними числами, що дозволяло мати до 65536 унікальних номерів, але зростання Інтернету призвело до розширення простору ASN до 32-бітних значень, що забезпечує понад 4 мільярди можливих номерів.

Міжавтономна маршрутизація суттєво відрізняється від внутрішньої маршрутизації як за цілями, так і за механізмами реалізації. Якщо протоколи IGP зосереджені на пошуку оптимальних шляхів всередині однорідного адміністративного домену з використанням технічних метрик, то протоколи зовнішньої маршрутизації (EGP) мають враховувати політичні, економічні та організаційні аспекти взаємодії між різними автономними системами. Основним завданням EGP є не стільки знаходження найкоротшого шляху з технічної точки зору, скільки забезпечення дотримання угод про транзит трафіку, реалізація бізнес-політик маршрутизації та контроль над розповсюдженням інформації про мережі. Це призводить до того, що рішення про вибір маршруту в EGP приймаються на основі складних наборів атрибутів та політик, а не простих числових метрик.

Історично перший протокол зовнішньої маршрутизації, також названий EGP (Exterior Gateway Protocol), був розроблений у 1980-х роках, але швидко виявив свої обмеження у зростаючому Інтернеті. Головною проблемою цього протоколу була відсутність механізмів запобігання петлям маршрутизації та підтримки складних топологій з множинними шляхами між автономними системами. На зміну йому прийшов протокол Border Gateway Protocol (BGP), який був вперше стандартизований у 1989 році і з того часу пройшов кілька

ревізій, при цьому сучасна версія BGP-4 залишається єдиним протоколом зовнішньої маршрутизації, який використовується в Інтернеті. BGP забезпечує необхідну гнучкість для реалізації складних політик маршрутизації, підтримує безкласову адресацію (CIDR) та має масштабованість, достатню для обробки сотень тисяч маршрутів у глобальній таблиці маршрутизації Інтернету.

2.2 Протокол BGP: основні принципи роботи та типи сесій

Border Gateway Protocol є path-vector протоколом, що означає передачу не просто інформації про доступність мережі, а повного шляху автономних систем, через які проходить маршрут до цієї мережі. Цей підхід принципово відрізняється від distance-vector протоколів, які передають лише інформацію про відстань до мережі, та від link-state протоколів, які будують повну топологічну карту. Передача повного AS-шляху дозволяє маршрутизатору самостійно виявляти потенційні петлі маршрутизації, оскільки якщо в отриманому анонсі міститься номер власної автономної системи, це однозначно вказує на наявність петлі і такий маршрут відкидається.

BGP функціонує на основі встановлення довгострокових TCP-з'єднань між сусідніми маршрутизаторами, які називаються BGP-пірами або сусідами. Використання TCP як транспортного протоколу (порт 179) забезпечує надійну доставку повідомлень маршрутизації, автоматичне керування потоком даних та виявлення обривів зв'язку між сусідами. Після встановлення TCP-з'єднання BGP-сусіди обмінюються повідомленнями OPEN, які містять параметри сесії, включаючи номер автономної системи, BGP-ідентифікатор маршрутизатора та час утримання сесії. Успішний обмін повідомленнями OPEN переводить сесію в стан Established, після чого починається обмін повною таблицею маршрутизації через повідомлення UPDATE. Надалі сусіди підтримують сесію живою шляхом періодичного обміну невеликими повідомленнями KEEPALIVE та передають лише інкрементальні зміни в маршрутизації, що значно зменшує навантаження на мережу та процесор маршрутизатора.

Розрізняють два основні типи BGP-сесій: зовнішні (eBGP) та внутрішні (iBGP), які мають суттєві відмінності в поведінці та призначенні. Сесії eBGP встановлюються між маршрутизаторами, які належать до різних автономних систем, зазвичай прямо з'єднаних один з одним на фізичному або каналному рівні. При передачі маршрутів через eBGP-сесію маршрутизатор додає номер своєї автономної системи до атрибуту AS_PATH, що дозволяє відстежувати повний шлях проходження анонсу через різні AS. За замовчуванням, eBGP-сесії встановлюються між безпосередньо підключеними інтерфейсами, а значення TTL в IP-пакетах встановлюється в 1, що є додатковим захисним механізмом проти випадкового встановлення сесій з віддаленими маршрутизаторами. Для сесій eBGP також типовою є зміна атрибуту NEXT_HOP на адресу власного інтерфейсу при передачі маршруту сусідній AS, що забезпечує коректну маршрутизацію трафіку.

На відміну від eBGP, сесії iBGP встановлюються між маршрутизаторами всередині однієї автономної системи для розповсюдження інформації про

зовнішні маршрути між граничними маршрутизаторами. Важливою особливістю iBGP є правило split-horizon, яке забороняє маршрутизатору передавати маршрути, отримані від одного iBGP-сусіда, іншим iBGP-сусідам. Це правило запроваджено для запобігання петлям маршрутизації всередині AS, але має важливий наслідок: всі маршрутизатори, які беруть участь у BGP всередині автономної системи, повинні мати повнозв'язну топологію iBGP-сесій, тобто кожен маршрутизатор повинен мати сесії з усіма іншими. В великих мережах це призводить до проблеми масштабованості, оскільки кількість необхідних сесій зростає квадратично з кількістю маршрутизаторів, що вирішується через використання механізмів Route Reflector або AS Confederation. При передачі маршрутів через iBGP атрибут AS_PATH не змінюється, оскільки анонс залишається всередині тієї самої автономної системи, а атрибут NEXT_HOP зазвичай зберігається без змін, що вимагає забезпечення доступності адреси NEXT_HOP через протоколи IGP.

Таблиця 2.1 - Порівняння характеристик eBGP та iBGP сесій

Характеристика	eBGP	iBGP
Розташування сусідів	Різні автономні системи	Одна автономна система
Типова топологія	Прямо підключені інтерфейси	Логічні з'єднання через IGP
Значення TTL за замовчуванням	1 (один хоп)	255 (множинні хопи)
Зміна атрибуту AS_PATH	Додається власний ASN	Не змінюється
Зміна атрибуту NEXT_HOP	Змінюється на власну адресу	Зберігається оригінальне значення
Правило split-horizon	Не застосовується	Застосовується (потрібна full mesh)
Administrative Distance	20	200

2.3 Атрибути BGP-маршрутів та механізм вибору найкращого шляху

Кожен маршрут у BGP супроводжується набором атрибутів, які описують різноманітні характеристики цього маршруту та використовуються для прийняття рішень про вибір найкращого шляху та застосування політик маршрутизації. Атрибути класифікуються на обов'язкові, які повинні присутні в кожному анонсі (well-known mandatory), необов'язкові, які розпізнаються всіма реалізаціями BGP але можуть бути відсутні (well-known discretionary), та опціональні, які не обов'язково підтримуються всіма маршрутизаторами (optional). Серед найважливіших атрибутів виділяються AS_PATH, NEXT_HOP, LOCAL_PREF, MED та ORIGIN, кожен з яких відіграє специфічну роль у процесі маршрутизації.

Атрибут AS_PATH містить упорядкований список номерів автономних систем, через які пройшов даний маршрутний анонс від джерела до поточного маршрутизатора. Цей атрибут є фундаментальним для BGP, оскільки виконує

дві критичні функції: запобігання петлям маршрутизації та участь у виборі найкращого шляху. Кожен маршрутизатор, який передає анонс через eBGP-сесію, додає номер своєї AS на початок списку AS_PATH, таким чином при отриманні маршруту можна прослідкувати весь його шлях. Якщо маршрутизатор отримує анонс, в AS_PATH якого міститься номер його власної автономної системи, такий маршрут автоматично відкидається як такий, що містить петлю. При виборі між декількома маршрутами до однієї мережі BGP віддає перевагу маршруту з найкоротшим AS_PATH, виходячи з припущення, що менша кількість транзитних автономних систем означає кращу якість маршруту. Проте адміністратори можуть маніпулювати вибором шляхів через механізм AS_PATH prepending, коли до списку штучно додаються додаткові копії власного номера AS для зменшення привабливості певного маршруту.

Атрибут NEXT_HOP вказує IP-адресу наступного маршрутизатора на шляху до мережі призначення, через який слід пересилати пакети для досягнення анонсованої мережі. Для маршрутів, отриманих через eBGP, NEXT_HOP зазвичай встановлюється в адресу сусіднього eBGP-маршрутизатора на безпосередньо підключеному інтерфейсі. Важливою особливістю є те, що при розповсюдженні маршруту через iBGP атрибут NEXT_HOP за замовчуванням не змінюється і продовжує вказувати на оригінальний eBGP-сусід, навіть якщо цей маршрутизатор знаходиться за межами прямого підключення. Це означає, що всі маршрутизатори всередині автономної системи повинні мати можливість досягти адреси NEXT_HOP через протокол внутрішньої маршрутизації (IGP), інакше маршрут буде вважатися недосяжним незважаючи на його присутність в BGP-таблиці. У деяких сценаріях, наприклад при використанні Route Reflector або при наявності проміжних мереж, може знадобитися примусова зміна NEXT_HOP через застосування команди next-hop-self, яка змушує маршрутизатор встановлювати власну адресу як NEXT_HOP для маршрутів, що передаються iBGP-сусідам.

Атрибут LOCAL_PREF (Local Preference) є well-known discretionary атрибутом, який використовується виключно всередині автономної системи для індикації переваги виходу трафіку з AS через певний граничний маршрутизатор. Цей атрибут передається лише через iBGP-сесії і не виходить за межі автономної системи, що робить його ідеальним інструментом для реалізації внутрішньої політики вибору шляхів без впливу на зовнішні AS. Більше значення LOCAL_PREF вказує на вищий пріоритет маршруту, і за замовчуванням всі маршрути отримують значення 100. Адміністратор може встановлювати різні значення LOCAL_PREF для маршрутів, отриманих від різних провайдерів або через різні канали зв'язку, таким чином контролюючи вихідний трафік. Наприклад, якщо організація має два підключення до Інтернету через різних провайдерів, можна встановити вище значення LOCAL_PREF для маршрутів від основного провайдера, що зробить цей канал переважним для виходу трафіку, зберігаючи другий канал як резервний.

Атрибут MED (Multi-Exit Discriminator), також відомий як метрика BGP, використовується для впливу на вхідний трафік у випадках, коли до сусідньої автономної системи існує кілька точок підключення. На відміну від

LOCAL_PREF, який контролює вихідний трафік і діє локально, MED передається в сусідню AS і підказує їй, який з кількох граничних маршрутизаторів є переважним для доставки трафіку. Менше значення MED вказує на вищу перевагу маршруту, і за замовчуванням всі маршрути мають MED рівний 0. Важливою особливістю MED є те, що він порівнюється лише для маршрутів, отриманих від однієї автономної системи, і не передається далі третім AS. Це обмеження означає, що MED є рекомендацією, а не вимогою, і сусідня AS може ігнорувати його або застосовувати власні політики. Типовим сценарієм використання MED є ситуація, коли організація має два канали до одного провайдера в різних географічних точках: встановлюючи нижче значення MED на маршрутах, анонсованих через географічно ближчу точку підключення, можна підказати провайдеру переважний шлях доставки трафіку.

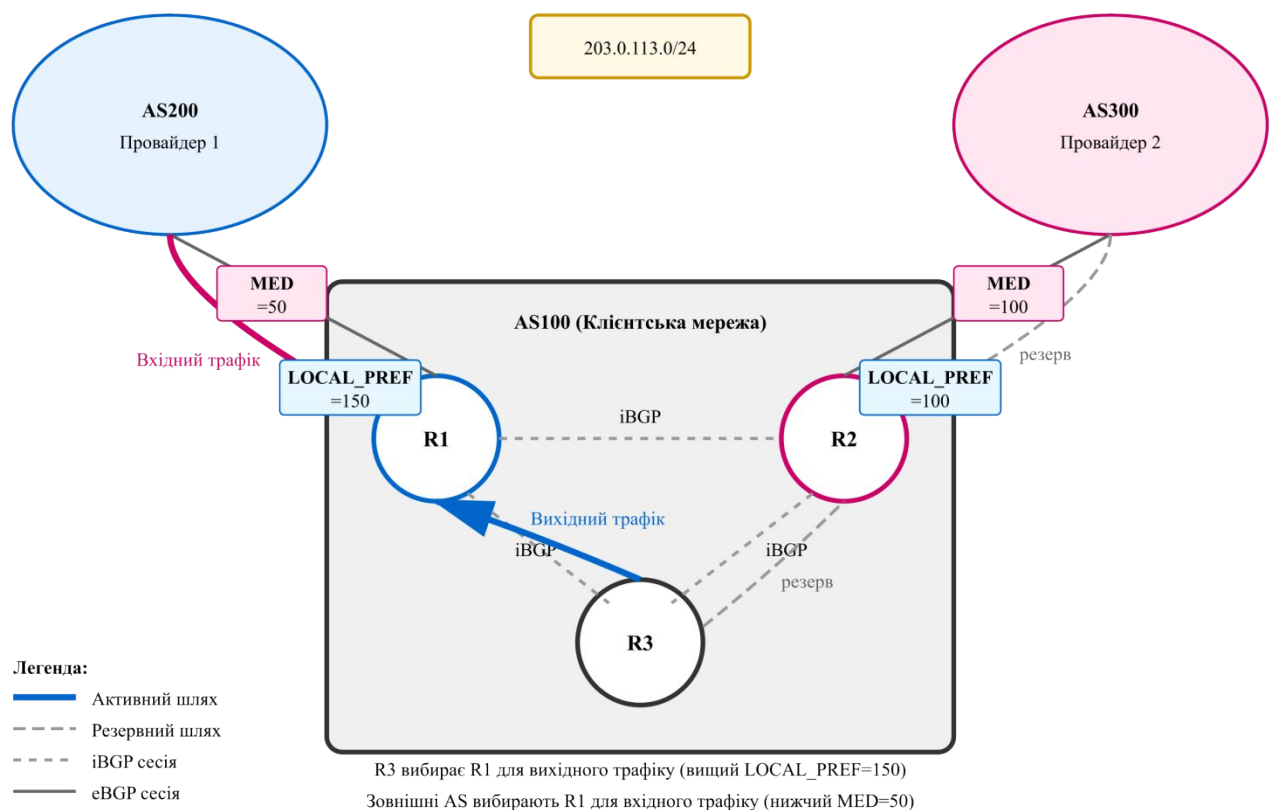


Рисунок 2.1 - Схема використання атрибутів LOCAL_PREF та MED для керування трафіком

На рисунку зображена автономна система AS100, яка має два граничні маршрутизатори R1 та R2, підключені до різних провайдерів AS200 та AS300 відповідно. Внутрішній маршрутизатор R3 отримує маршрути до зовнішньої мережі 203.0.113.0/24 через обидва граничні маршрутизатори. R1 встановлює LOCAL_PREF=150 для маршрутів від AS200, а R2 встановлює LOCAL_PREF=100 для маршрутів від AS300. Це призводить до того, що R3 вибирає маршрут через R1 для вихідного трафіку. Для вхідного трафіку R1 анонсує мережі AS100 з MED=50, а R2 з MED=100, що підказує зовнішнім системам використовувати R1 як переважну точку входу. Стрілки на схемі

показують напрямок потоків трафіку, товсті лінії позначають активні шляхи, а пунктирні - резервні.

Процес вибору найкращого маршруту в BGP є детермінованим алгоритмом, який послідовно застосовує набір правил для порівняння всіх доступних маршрутів до конкретної мережі призначення. BGP-маршрутизатор зберігає всі отримані маршрути в таблиці BGP (Adj-RIB-In), але в таблицю маршрутизації (RIB) встановлює лише один найкращий маршрут, який потім може бути переданий сусідам. Алгоритм вибору найкращого шляху BGP включає наступні кроки, які застосовуються послідовно до виключення всіх маршрутів крім одного: перевірка доступності NEXT_HOP, перевага вищого значення LOCAL_PREF, перевага локально згенерованих маршрутів, перевага коротшого AS_PATH, перевага кращого ORIGIN (IGP краще за EGP, EGP краще за Incomplete), перевага нижчого MED, перевага eBGP над iBGP, перевага маршруту через найближчий IGP-сусід до NEXT_HOP, і нарешті, якщо всі попередні критерії однакові, вибір маршруту з найнижчим BGP Router ID. Цей багатоступеневий процес забезпечує стабільний та передбачуваний вибір маршрутів при збереженні гнучкості для адміністраторів через налаштування атрибутів.

2.4 Політики маршрутизації, фільтрація та агрегація маршрутів у BGP

Одним з найважливіших аспектів протоколу BGP є можливість реалізації складних політик маршрутизації, які дозволяють автономним системам контролювати прийом, відбір та розповсюдження маршрутної інформації відповідно до технічних, організаційних та комерційних вимог. На відміну від протоколів IGP, які зазвичай прагнуть до досягнення оптимального шляху за єдиною метрикою, BGP надає адміністраторам повний контроль над процесом прийняття рішень через систему політик. Політики маршрутизації реалізуються за допомогою механізмів фільтрації та модифікації атрибутів маршрутів, які застосовуються як до вхідних (inbound), так і до вихідних (outbound) оновлень на кожній BGP-сесії. Правильно сконфігуровані політики є критично важливими не лише для оптимізації потоків трафіку, але й для забезпечення безпеки мережі та відповідності комерційним угодам між провайдерами та клієнтами.

Фільтрація маршрутів у BGP виконується за допомогою списків доступу (access-list), префіксних списків (prefix-list) або route-map, які дозволяють визначати умови для прийняття чи відхилення маршрутів на основі їх префіксів, довжини маски підмережі, значень атрибутів або комбінації цих параметрів. Найпростіша форма фільтрації використовує prefix-list для дозволу або заборони анонсування конкретних діапазонів адрес. Наприклад, провайдер може фільтрувати анонси від клієнта, приймаючи лише ті префікси, які належать виділеному клієнту адресному простору, відкидаючи будь-які інші анонси для захисту від помилкової конфігурації або зловмисних дій. Більш складні сценарії фільтрації можуть включати перевірку довжини AS_PATH для

виявлення аномально довгих шляхів, які можуть вказувати на петлі або атаки route hijacking, або фільтрацію на основі наявності певних автономних систем в AS_PATH для реалізації політик транзиту. Фільтрація застосовується як до вхідних маршрутів для контролю того, що приймається від сусідів, так і до вихідних для контролю того, які маршрути анонсуються зовні, що дає повний контроль над потоком маршрутної інформації.

Route-map є найпотужнішим інструментом для реалізації політик BGP, оскільки дозволяє не лише фільтрувати маршрути, але й модифікувати їх атрибути умовно на основі різноманітних критеріїв. Route-map складається з послідовності записів, кожен з яких містить умови відповідності (match) та дії (set), аналогічно до конструкції if-then в програмуванні. Умови відповідності можуть перевіряти префікс мережі, довжину маски, значення AS_PATH, присутність певних community-тегів та багато інших параметрів. Якщо всі умови в записі виконуються, застосовуються визначені дії, які можуть включати встановлення значень атрибутів LOCAL_PREF, MED, AS_PATH prepending, модифікацію NEXT_HOP або додавання community-тегів. Якщо жоден запис у route-map не відповідає маршруту, за замовчуванням маршрут відхиляється, що відповідає неявному deny в кінці кожної route-map. Ця поведінка вимагає явного дозволу всіх бажаних маршрутів, що підвищує безпеку але потребує ретельного планування.

BGP community є опціональним транзитивним атрибутом, який дозволяє групувати маршрути з подібними характеристиками для спрощення застосування політик. Community представляє собою 32-бітне значення, яке зазвичай записується у форматі ASN:NN, де перша частина вказує номер автономної системи, а друга є довільним числом, визначеним адміністратором. Маршрути можуть мати множинні community-теги, і ці теги передаються між автономними системами, дозволяючи реалізувати складні розподілені політики маршрутизації. Типовими застосуваннями community є маркування географічного походження маршрутів, індикація типу підключення клієнта або сигналізація бажаного способу обробки маршруту в транзитних AS. Деякі community мають зарезервоване значення і розпізнаються всіма реалізаціями BGP: NO_EXPORT вказує, що маршрут не повинен анонсуватися за межі AS, NO_ADVERTISE забороняє анонсування маршруту будь-яким сусідам, а NO_EXPORT_SUBCONFED обмежує розповсюдження в межах AS-конфедерації.

Таблиця 2.2 - Типові сценарії використання BGP-політик

Сценарій	Застосовуваний механізм	Мета
Запобігання анонсуванню приватних адрес	Prefix-list на вихідних оновленнях	Безпека: блокування RFC1918 адрес
Контроль клієнтських анонсів	Prefix-list з точним переліком дозволених мереж	Захист від hijacking та помилкової конфігурації
Вибір переважного	Route-map з	Оптимізація вихідного трафіку

провайдера	встановленням LOCAL_PREF	
Балансування вхідного трафіку	Route-map з модифікацією MED	Розподіл навантаження між каналами
Обмеження глибини транзиту	AS_PATH filter за довжиною	Запобігання петлям та атакам
Маркування типу клієнта	Встановлення BGP community	Спрощення політик на магістральних маршрутизаторах

Агрегація маршрутів є критичним механізмом для підтримання масштабованості глобальної таблиці маршрутизації Інтернету, оскільки дозволяє об'єднувати множину більш специфічних префіксів в один суммарний маршрут, зменшуючи кількість записів в таблиці BGP. Без агрегації кожна невелика підмережа потребувала б окремого запису, що призвело б до надмірного зростання таблиць маршрутизації та виснаження пам'яті маршрутизаторів. BGP підтримує автоматичну агрегацію через команду `aggregate-address`, яка створює суммарний маршрут для заданого префіксу, якщо в таблиці BGP присутній хоча б один більш специфічний компонент цього суммарного маршруту. За замовчуванням агрегований маршрут анонсується одночасно з компонентами, але опція `summary-only` дозволяє придушити анонсування конкретних префіксів, залишаючи лише агрегат. При створенні агрегованого маршруту AS_PATH формується як об'єднання всіх унікальних номерів AS з компонентних маршрутів, що може призводити до втрати інформації про точний шлях, проте атрибут AS_SET зберігає цю інформацію для запобігання петлям. Адміністратори повинні обережно підходити до агрегації, оскільки надмірна агрегація може призвести до неоптимальної маршрутизації або до анонсування адресного простору, частина якого насправді не використовується.

Особливу роль у сучасному Інтернеті відіграють системи перевірки походження маршрутів, такі як RPKI (Resource Public Key Infrastructure), які спрямовані на підвищення безпеки BGP через криптографічну верифікацію права автономної системи анонсувати певний адресний простір. RPKI використовує ієрархію сертифікатів, видану регіональними інтернет-реєстрами, для створення ROA (Route Origin Authorization) об'єктів, які однозначно вказують, яка AS має право анонсувати певний префікс та максимальну дозволена довжину маски. Маршрутизатори, які підтримують RPKI, можуть перевіряти отримані BGP-анонси на відповідність ROA, класифікуючи маршрути як Valid (відповідає існуючому ROA), Invalid (суперечить ROA або перевищує максимальну довжину), або Unknown (немає відповідного ROA). На основі цієї класифікації можна застосовувати політики, наприклад, повністю відкидаючи Invalid маршрути або знижуючи їх LOCAL_PREF для зменшення ймовірності вибору, що значно підвищує захищеність від атак route hijacking та випадкових помилок конфігурації.

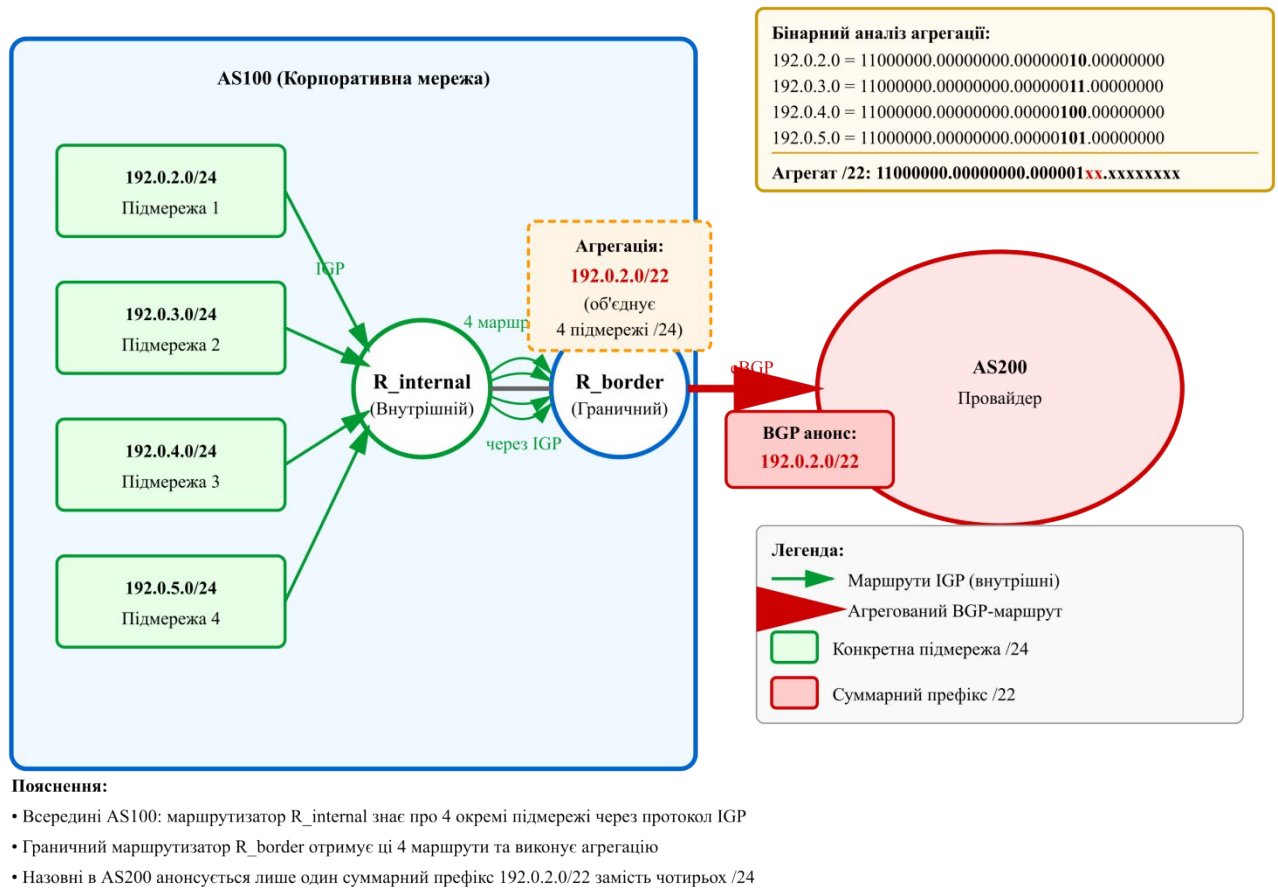


Рисунок 2.2 - Приклад агрегації маршрутів у BGP

На рисунку показана автономна система AS100, яка має внутрішню адресну структуру з чотирма підмережами класу C: 192.0.2.0/24, 192.0.3.0/24, 192.0.4.0/24 та 192.0.5.0/24. Внутрішній маршрутизатор R_internal отримує всі чотири маршрути через протокол IGP та встановлює їх в свою таблицю маршрутизації. Граничний маршрутизатор R_border, який підключений до зовнішнього провайдера AS200, налаштований на агрегацію цих чотирьох /24 префіксів в один суммарний /22 префікс: 192.0.2.0/22. Замість анонсування чотирьох окремих маршрутів в BGP, R_border анонсує лише один агрегований префікс, зменшуючи навантаження на таблицю маршрутизації провайдера. На схемі суцільними стрілками показані конкретні маршрути всередині AS, а товстою стрілкою - агрегований анонс назовні. Така практика є стандартною для корпоративних мереж та провайдерів, оскільки дозволяє зменшити розмір глобальної таблиці маршрутизації Інтернету.

Успішна реалізація політик BGP вимагає ретельного планування та тестування, оскільки помилки в конфігурації можуть призвести до серйозних наслідків, включаючи повну втрату зв'язності, створення чорних дір для трафіку або ненавмисне анонсування чужого адресного простору. Рекомендується документувати всі політики маршрутизації, використовувати послідовні схеми нумерації для community-тегів, регулярно перевіряти таблиці BGP на наявність аномалій та застосовувати максимальні обмеження (max-prefix) на кількість маршрутів, прийнятих від кожного сусіда, для захисту від випадкових або зловмисних флудів маршрутної інформації. Моніторинг та

логування змін в BGP-сесіях також є критично важливими для швидкого виявлення та реагування на інциденти, пов'язані з маршрутизацією.

Тема 3. Механізми трансляції адрес та автоматичного призначення IP

3.1. Протоколи трансляції мережевих адрес NAT та PAT

Протокол трансляції мережевих адрес (Network Address Translation, NAT) є ключовим механізмом сучасних комп'ютерних мереж, який дозволяє вирішити проблему обмеженості адресного простору IPv4 та забезпечити додатковий рівень безпеки для внутрішньої мережі. Основна ідея NAT полягає у відображенні приватних IP-адрес внутрішньої мережі на одну або кілька публічних адрес, які використовуються для зв'язку з зовнішніми мережами та Інтернетом. Цей механізм став особливо актуальним з моменту вичерпання вільних IPv4-адрес та необхідності економного використання наявного адресного простору. Впровадження NAT дозволило організаціям використовувати приватні діапазони адрес, визначені стандартом RFC 1918, для побудови внутрішніх мереж будь-якого розміру, при цьому зберігаючи можливість доступу до глобальних ресурсів через обмежену кількість публічних адрес.

Процес трансляції адрес виконується спеціалізованим пристроєм, зазвичай граничним маршрутизатором або міжмережевим екраном, який розташований на межі між внутрішньою та зовнішньою мережами. Коли пакет від внутрішнього вузла прямує до зовнішнього призначення, NAT-пристрій модифікує поле адреси відправника в IP-заголовку, замінюючи приватну адресу на публічну. Одночасно пристрій створює запис у таблиці трансляцій, який містить інформацію про відповідність між внутрішньою та зовнішньою адресами, а також додаткові дані про з'єднання, такі як номери портів для транспортного рівня. Коли відповідь від зовнішнього сервера повертається, NAT-пристрій використовує цю таблицю для зворотної трансляції, замінюючи публічну адресу призначення на відповідну приватну адресу внутрішнього вузла. Цей процес є прозорим для кінцевих вузлів і не вимагає жодних змін у їх конфігурації або програмному забезпеченні.

Існують три основні типи NAT, кожен з яких має свої особливості застосування та відрізняється механізмом відображення адрес, як показано в таблиці 3.1.

Таблиця 3.1 - Порівняльна характеристика типів NAT

Тип NAT	Відображення адрес	Напрямок ініціації	Типове застосування	Обмеження
Статичний NAT	Одна приватна ↔ одна публічна (постійне)	Обидва напрямки	Публікація серверів	Кількість внутрішніх вузлів = кількості публічних адрес
Динамічний NAT	Приватна ↔ публічна з пулу (тимчасове)	Зсереди назовні	Надання доступу групі користувачів	Пул публічних адрес може вичерпатися

PAT (NAT Overload)	Множина приватних ↔ одна публічна (різні порти)	Зсер єдини назовні	Типова корпоративна мережа	Обмеження одночасних з'єднань (~65000)
--------------------------	--	--------------------------	----------------------------------	--

Статичний NAT забезпечує постійне взаємооднозначне відображення між приватною та публічною адресами, що робить його ідеальним рішенням для випадків, коли внутрішній ресурс повинен бути доступним з зовнішньої мережі під постійною адресою. Такий підхід часто використовується для публікації веб-серверів, поштових серверів або інших служб, які повинні приймати вхідні з'єднання з Інтернету. Однак статичний NAT не дозволяє економити публічні адреси, оскільки кожному внутрішньому вузлу відповідає окрема зовнішня адреса. Динамічний NAT вирішує цю проблему шляхом створення тимчасових відображень з пулу доступних публічних адрес на момент встановлення з'єднання. Коли внутрішній вузол ініціює з'єднання з зовнішньою мережею, NAT-пристрій вибирає вільну адресу з пулу, створює запис у таблиці трансляцій і використовує цю адресу до завершення сесії, після чого адреса повертається до пулу для повторного використання.

Найбільш ефективним з точки зору економії адресного простору є механізм PAT (Port Address Translation), також відомий як NAT Overload. Цей метод дозволяє відображати велику кількість внутрішніх адрес на одну зовнішню адресу шляхом використання різних номерів портів транспортного рівня для розрізнення окремих з'єднань. При обробці вихідного пакета PAT-пристрій не лише замінює адресу відправника, але й модифікує номер порту джерела, присвоюючи йому унікальне значення з діапазону доступних портів. У таблиці трансляцій зберігається комбінація внутрішньої адреси, внутрішнього порту, зовнішньої адреси, зовнішнього порту та протоколу, що дозволяє однозначно ідентифікувати кожне з'єднання при обробці зворотного трафіку. Завдяки цьому механізму одна публічна адреса може обслуговувати тисячі одночасних з'єднань, що робить PAT стандартним рішенням для більшості корпоративних та домашніх мереж.

Рисунок 3.1 ілюструє процес роботи PAT в типовій корпоративній мережі. На схемі зображено внутрішню мережу з адресним простором 192.168.1.0/24, в якій розташовано три робочі станції з адресами 192.168.1.10, 192.168.1.11 та 192.168.1.12. Ці станції підключені до внутрішнього інтерфейсу граничного маршрутизатора, який виконує функції PAT. Зовнішній інтерфейс маршрутизатора має публічну адресу 203.0.113.5 і підключений до мережі Інтернет. На схемі показано три активні з'єднання: перша робоча станція звертається до веб-сервера 198.51.100.20 з локального порту 52341, друга станція також звертається до того ж сервера з порту 51892, а третя станція встановлює SSH-з'єднання до сервера 203.0.113.50 з порту 49123. Маршрутизатор виконує трансляцію, замінюючи внутрішні адреси на публічну 203.0.113.5 та призначаючи унікальні зовнішні порти: 1024, 1025 та 1026 відповідно. Таблиця трансляцій на маршрутизаторі містить записи, що зв'язують кожну комбінацію внутрішньої адреси та порту з відповідною

комбінацією зовнішньої адреси та порту, дозволяючи коректно маршрутизувати зворотній трафік до відповідних внутрішніх вузлів.

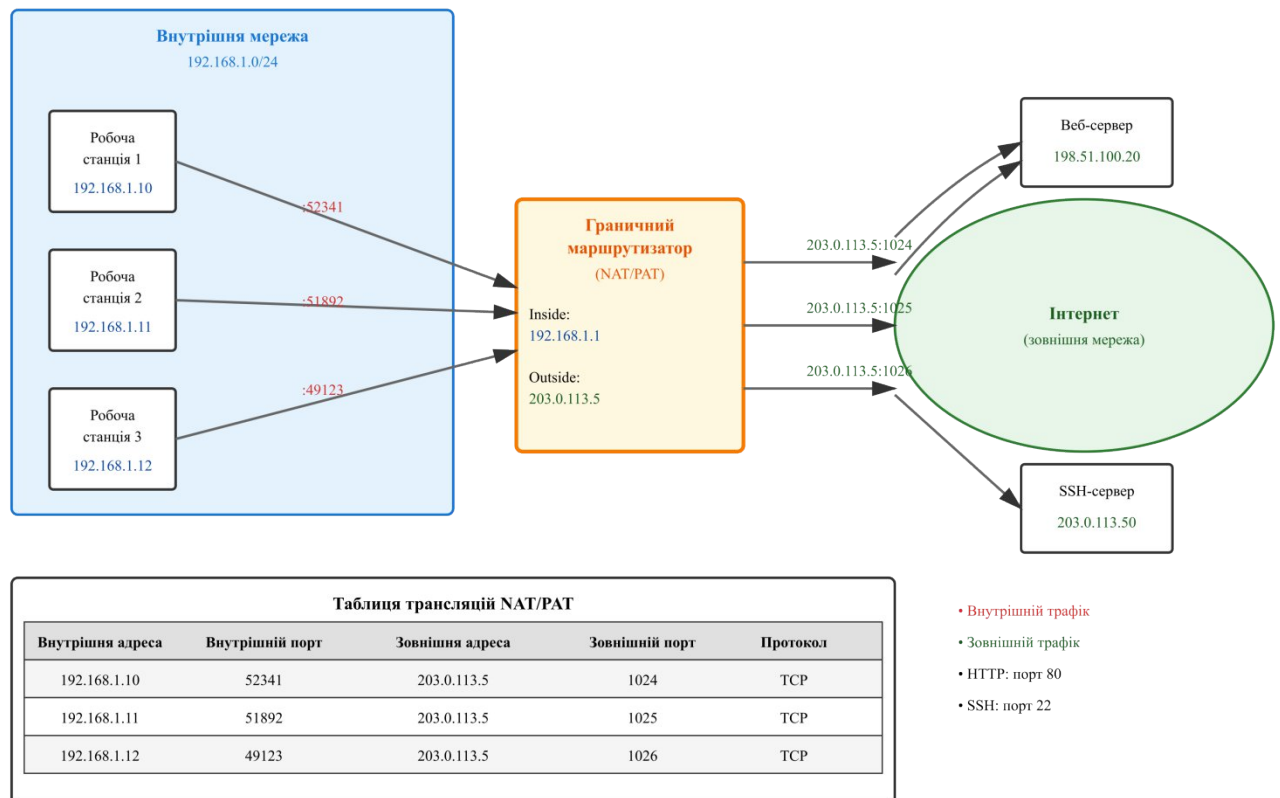


Рисунок 3.1 - Схема роботи PAT в корпоративній мережі

Впровадження технологій NAT/PAT має не лише переваги, але й певні обмеження та особливості, які необхідно враховувати при проектуванні та експлуатації мережевої інфраструктури. Однією з основних переваг є економія публічних IPv4-адрес, що дозволяє організації використовувати лише невелику кількість зовнішніх адрес незалежно від розміру внутрішньої мережі. Додатковою перевагою є підвищення рівня безпеки через приховування внутрішньої топології мережі та адресації від зовнішніх спостерігачів, що ускладнює проведення розвідувальних атак та сканування внутрішніх ресурсів. Механізм трансляції адрес також забезпечує певний рівень фільтрації трафіку, оскільки вхідні з'єднання блокуються за замовчуванням, якщо не налаштовано явні правила статичної трансляції або переадресації портів. Проте використання NAT створює проблеми для деяких протоколів прикладного рівня, особливо тих, що вбудовують IP-адреси в корисне навантаження пакетів, як от FTP, SIP або деякі протоколи VPN, що вимагає застосування додаткових механізмів обходу NAT або використання Application Layer Gateway (ALG).

3.2. Протокол автоматичного призначення IP-адрес DHCP

Протокол динамічної конфігурації вузла (Dynamic Host Configuration Protocol, DHCP) є стандартним механізмом автоматичного призначення мережевих параметрів клієнтським пристроям в IP-мережах. Основне

призначення DHCP полягає в централізованому управлінні розподілом IP-адрес та супутніх параметрів конфігурації, таких як маска підмережі, адреса шлюзу за замовчуванням, адреси DNS-серверів та інші опції, необхідні для коректної роботи вузла в мережі. До появи DHCP адміністратори були змушені вручну налаштовувати мережеві параметри на кожному пристрої, що створювало значні труднощі при управлінні великими мережами, призводило до помилок конфігурації та конфліктів адрес. Автоматизація цього процесу через DHCP суттєво спростила адміністрування мереж, зменшила кількість помилок та дозволила ефективно використовувати обмежений пул IP-адрес через механізм тимчасової оренди адрес.

DHCP побудований на архітектурі клієнт-сервер, де DHCP-сервер централізовано управляє пулами адрес та політиками їх розподілу, а DHCP-клієнти на робочих станціях, мобільних пристроях та іншому обладнанні запитують конфігурацію при підключенні до мережі. Протокол використовує транспортний протокол UDP, причому сервер очікує запити на порту 67, а клієнт отримує відповіді на порту 68. Така асиметрія портів дозволяє клієнту, який ще не має власної IP-адреси, коректно отримати відповідь від сервера. DHCP-сервер підтримує три режими роботи з адресами: ручне призначення, коли адміністратор створює постійну прив'язку MAC-адреси до IP-адреси, автоматичне призначення, коли сервер видає адресу назавжди з доступного пулу, та динамічне призначення, найбільш поширений режим, в якому адреси видаються на певний час оренди з можливістю подовження або вивільнення для повторного використання.

Процес отримання конфігурації клієнтом від сервера відбувається через послідовність з чотирьох повідомлень, відому як процедура DORA (Discovery, Offer, Request, Acknowledgment), детальний опис якої наведено в таблиці 3.2.

Таблиця 3.2 - Етапи процесу DORA в протоколі DHCP

етап	Повідомлення	Відправник	Адреса джерела	Адреса призначення	Зміст	Мета
	DHCP Discover	Клієнт	0.0.0.0	255.255.255.255	Запит на конфігурацію	Виявлення доступних DHCP-серверів
	DHCP Offer	Сервер(и)	IP сервера	255.255.255.255 або MAC клієнта	Пропозиція адреси та параметрів	Надання варіанту конфігурації
	DHCP Request	Клієнт	0.0.0.0	255.255.255.255	Прийняття обраної пропозиції	Запит на видачу адреси від обраного сервера

	DHCP Acknowledgment	Сер вер	I Р сервера	IP клієнта	Підтвердження та фінальні параметри	Завдання процесу призначення
--	------------------------	------------	-------------------	---------------	--	------------------------------------

Перший етап процесу починається з відправлення клієнтом широкомовного повідомлення DHCP Discover, яке використовує адресу джерела 0.0.0.0, оскільки клієнт ще не має власної IP-адреси, та адресу призначення 255.255.255.255 для досягнення всіх вузлів в локальному сегменті мережі. Це повідомлення містить унікальний ідентифікатор транзакції, MAC-адресу клієнта та може включати додаткові параметри запиту, такі як бажана IP-адреса або список необхідних опцій конфігурації. Всі DHCP-сервери в мережі, які отримують це повідомлення, аналізують запит і, якщо мають доступні адреси в своїх пулах, готують відповіді. На другому етапі кожен сервер відправляє клієнту повідомлення DHCP Offer, яке містить пропозицію IP-адреси, маску підмережі, час оренди та інші параметри конфігурації згідно з політикою сервера. Якщо в мережі присутні кілька DHCP-серверів, клієнт може отримати декілька пропозицій, з яких він обирає найбільш придатну, зазвичай ту, що надійшла першою.

Третій етап включає відправлення клієнтом широкомовного повідомлення DHCP Request, яке інформує всі сервери про те, пропозицію якого сервера клієнт вирішив прийняти. Це повідомлення містить ідентифікатор обраного сервера та запропоновану адресу, що дозволяє іншим серверам звільнити зарезервовані для цього клієнта адреси та повернути їх до пулу доступних. Обраний сервер, отримавши DHCP Request, виконує фінальну перевірку доступності адреси та відправляє підтверджуюче повідомлення DHCP Acknowledgment, яке завершує процес призначення конфігурації. Після отримання цього підтвердження клієнт конфігурує свій мережевий інтерфейс з отриманими параметрами та починає використовувати призначену IP-адресу. У випадку, якщо адреса з якихось причин не може бути призначена, наприклад, через конфлікт адрес або помилку в конфігурації, сервер відправляє повідомлення DHCP Negative Acknowledgment (NAK), після отримання якого клієнт повинен почати процес заново з етапу Discovery.

DHCP підтримує широкий набір опцій конфігурації, що дозволяє передавати клієнтам різноманітні параметри мережевого середовища понад базові налаштування IP-адресації. Найбільш часто використовувані опції включають адресу шлюзу за замовчуванням (опція 3), яка визначає маршрутизатор для доступу до віддалених мереж, адреси DNS-серверів (опція 6) для розв'язання доменних імен, доменне ім'я (опція 15) для автоматичного формування повних доменних імен вузлів, та адреси NTP-серверів (опція 42) для синхронізації часу. Більш специфічні опції можуть включати параметри для IP-телефонії, такі як адреси TFTP-серверів для завантаження конфігурацій VoIP-телефонів, параметри завантаження безфайлових робочих станцій (PXE boot), налаштування для специфічних протоколів та служб. Механізм оренди

адрес є критично важливим аспектом роботи DHCP, оскільки він дозволяє ефективно використовувати обмежений пул адрес в динамічних середовищах, де пристрої регулярно підключаються та відключаються від мережі.

Типовий час оренди налаштовується адміністратором відповідно до характеристик мережевого середовища і може варіюватися від кількох годин до декількох днів або тижнів. Для мереж з високою мобільністю користувачів, таких як публічні Wi-Fi точки доступу, зазвичай встановлюють короткий час оренди в межах кількох годин, що дозволяє швидко вивільняти адреси від вузлів, що відключилися. В корпоративних мережах з відносно стабільним складом обладнання можна використовувати більш тривалі терміни оренди, що зменшує навантаження на DHCP-сервер від частих операцій подовження. Коли проходить половина терміну оренди (T1 timer), клієнт автоматично намагається подовжити оренду, відправляючи повідомлення DHCP Request безпосередньо на адресу сервера, що видав адресу. Якщо подовження успішне, термін оренди скидається, і клієнт продовжує використовувати ту ж саму адресу. У випадку недоступності основного сервера при досягненні 87.5% терміну оренди (T2 timer) клієнт робить широкомовний запит, дозволяючи іншим DHCP-серверам підхопити оренду та забезпечити безперервність роботи.

Рисунок 3.2 демонструє повний життєвий цикл DHCP-оренди від початкового отримання до вивільнення адреси. Часова діаграма показує клієнта, що виконує початковий процес DORA для отримання IP-адреси 192.168.1.100 з часом оренди 24 години. Після отримання адреси в момент T=0 клієнт успішно використовує мережу протягом 12 годин, після чого досягає моменту T1 (50% від часу оренди). У цей момент клієнт відправляє одноадресне повідомлення DHCP Request безпосередньо на DHCP-сервер з метою подовження оренди. Сервер відповідає повідомленням DHCP Acknowledgment, скидаючи таймер оренди до початкового значення 24 години. Клієнт продовжує роботу ще 12 годин до наступного моменту T1, коли знову успішно подовжує оренду. На діаграмі також показано альтернативний сценарій, де в момент T2 (87.5% від часу оренди) початковий сервер недоступний, тому клієнт виконує широкомовний запит, на який відповідає резервний DHCP-сервер, забезпечуючи безперервність роботи. Нарешті, коли клієнт коректно завершує роботу або відключається від мережі, він відправляє повідомлення DHCP Release, яке інформує сервер про готовність повернути адресу до пулу для наступного використання іншими клієнтами.

3.3. Інтеграція NAT/PAT та DHCP на маршрутизаторі Cisco

Сучасні граничні маршрутизатори Cisco поєднують функціональність NAT/PAT та DHCP-сервера, що дозволяє реалізувати повноцінне рішення для управління адресацією в корпоративних мережах на базі єдиного пристрою. Така інтеграція забезпечує спрощення архітектури мережі, зменшення кількості необхідного обладнання та централізоване управління механізмами трансляції та розподілу адрес. Маршрутизатор може одночасно виконувати роль DHCP-сервера для внутрішньої мережі, автоматично призначаючи IP-адреси

клієнтським пристроям, та здійснювати трансляцію цих приватних адрес у публічні при пересиланні трафіку до зовнішніх мереж. Конфігурація таких служб на маршрутизаторах Cisco виконується через інтерфейс командного рядка (CLI) з використанням специфічних команд операційної системи Cisco IOS.

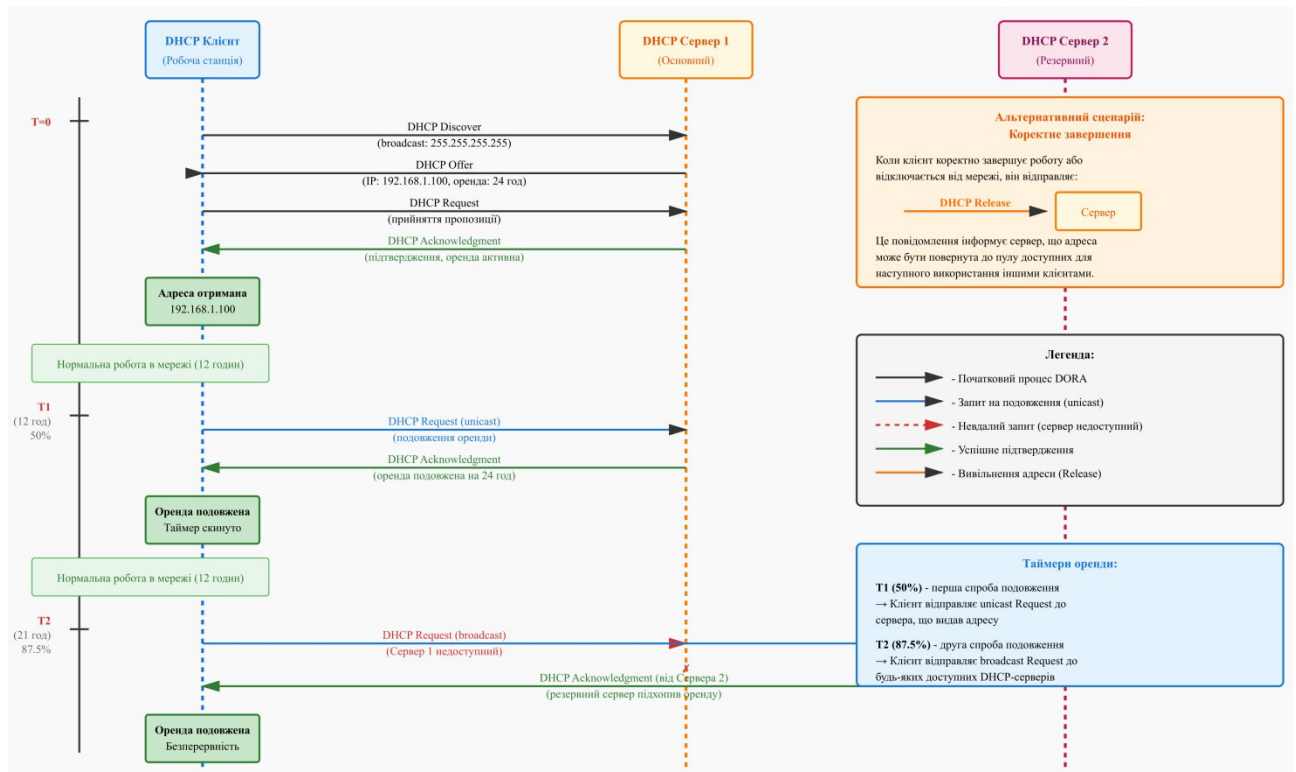


Рисунок 3.2 - Життєвий цикл DHCP-оренди IP-адреси

Налаштування DHCP-сервера на маршрутизаторі Cisco починається зі створення пулу адрес, який визначає діапазон IP-адрес для автоматичного розподілу, а також асоційовані параметри конфігурації. Команда “ip dhcp pool” з ім’ям пулу переводить маршрутизатор у режим конфігурації DHCP-пулу, де адміністратор може визначити мережу та маску підмережі командою “network”, вказати адресу шлюзу за замовчуванням командою “default-router”, налаштувати DNS-сервери через “dns-server” та встановити інші необхідні опції. Окремою командою “ip dhcp excluded-address” можна виключити певні адреси або діапазони з автоматичного розподілу, що дозволяє зарезервувати їх для статичного призначення серверам, принтерам або мережевому обладнанню. Після створення пулу та визначення всіх необхідних параметрів DHCP-служба автоматично активується на маршрутизаторі і починає відповідати на запити клієнтів, що надходять на внутрішні інтерфейси.

Конфігурація NAT на маршрутизаторі Cisco вимагає визначення внутрішніх та зовнішніх інтерфейсів, а також правил трансляції адрес відповідно до обраного типу NAT. Для позначення інтерфейсів використовуються команди “ip nat inside” на інтерфейсах, підключених до внутрішньої мережі, та “ip nat outside” на інтерфейсі, що веде до зовнішньої мережі або провайдера. Статична трансляція налаштовується командою “ip nat

inside source static”, яка створює постійне відображення між внутрішньою та зовнішньою адресами. Для динамічного NAT спочатку створюється пул зовнішніх адрес командою “ip nat pool” з вказівкою діапазону публічних адрес, потім визначається список доступу (ACL), який ідентифікує внутрішні адреси, що підлягають трансляції, і нарешті встановлюється зв’язок між списком доступу та пулом через команду “ip nat inside source list”. Найпоширеніша конфігурація PAT реалізується командою “ip nat inside source list overload”, яка вказує на використання перевантаження портів для одночасної трансляції множини внутрішніх адрес на одну або кілька зовнішніх адрес.

Таблиця 3.3 наводить порівняння основних команд конфігурації для різних типів NAT на маршрутизаторах Cisco IOS.

Таблиця 3.3 - Команди конфігурації NAT на маршрутизаторах Cisco

Тип NAT	Основні команди конфігурації	Додаткові вимоги	Приклад використання
Статичний	ip nat inside source static	Визначення inside/outside інтерфейсів	Публікація веб-сервера 192.168.1.10 під адресою 203.0.113.10
Динамічний	ip nat pool <ім’я> netmask ip nat inside source list pool <ім’я>	ACL для визначення внутрішніх адрес	Трансляція мережі 192.168.1.0/24 в пул 203.0.113.10-20
PAT	ip nat inside source list interface overload	ACL для визначення внутрішніх адрес	Трансляція всієї внутрішньої мережі на адресу зовнішнього інтерфейсу

Інтеграція DHCP та NAT на одному пристрої створює певні особливості взаємодії цих механізмів, які необхідно враховувати при проектуванні конфігурації. Важливо забезпечити, щоб діапазон адрес, який використовується DHCP-сервером для автоматичного розподілу, повністю покривався списком доступу, що визначає адреси для NAT-трансляції. Якщо деякі динамічно призначені адреси не будуть включені до NAT ACL, відповідні клієнти не зможуть отримати доступ до зовнішніх мереж. Також рекомендується виключати з DHCP-пулу адреси, які використовуються для статичного NAT або призначені мережевому обладнанню, щоб уникнути конфліктів адресації. В середовищах з високими вимогами до доступності доцільно розглянути використання окремого виділеного DHCP-сервера замість вбудованої функціональності маршрутизатора, оскільки це дозволяє розподілити навантаження, забезпечити резервування та отримати більш розширені можливості управління.

Моніторинг та діагностика роботи NAT і DHCP на маршрутизаторах Cisco виконується за допомогою спеціалізованих команд перегляду стану та статистики. Команда “show ip nat translations” відображає поточну таблицю активних трансляцій NAT, показуючи відповідності між внутрішніми та зовнішніми адресами і портами для кожного активного з’єднання, що дозволяє відстежувати використання пулів адрес та діагностувати проблеми з’єднань.

Команда “show ip nat statistics” надає агреговану статистику роботи NAT, включаючи загальну кількість активних трансляцій, кількість промахів та влучень в таблицю, та інформацію про використання налаштованих пулів адрес. Для перевірки роботи DHCP-сервера використовується команда “show ip dhcp binding”, яка відображає список всіх активних орендованих адрес з інформацією про MAC-адреси клієнтів, призначені IP-адреси, терміни оренди та стан кожного призначення. Команда “show ip dhcp server statistics” показує статистику обробки DHCP-повідомлень, включаючи кількість отриманих запитів кожного типу та відправлених відповідей, що допомагає виявляти аномалії в роботі служби.

При виникненні проблем з роботою NAT або DHCP адміністратори можуть використовувати додаткові діагностичні інструменти та команди для детального аналізу поведінки системи. Команда “debug ip nat” активує детальне журналювання всіх операцій NAT в режимі реального часу, дозволяючи відстежувати створення та видалення записів трансляції, але слід пам’ятати, що активація режиму debug може суттєво навантажити процесор маршрутизатора, тому її використання в продуктивних середовищах має бути обмеженим. Аналогічно, команда “debug ip dhcp server events” виводить детальну інформацію про обробку DHCP-запитів, включаючи всі етапи процесу DORA та прийняті рішення щодо призначення адрес. Для очищення застарілих або проблемних записів адміністратор може використовувати команди “clear ip nat translation”, яка дозволяє видалити окремі або всі активні NAT-трансляції, та “clear ip dhcp binding”, що вивільняє орендовані адреси та дозволяє клієнтам отримати нову конфігурацію при наступному запиті.

3.4. Роль NAT та DHCP в архітектурі сучасних мереж

Протоколи NAT та DHCP відіграють фундаментальну роль у побудові сучасних IP-мереж, забезпечуючи критично важливі функції управління адресацією та зв’язності між приватними та публічними сегментами інфраструктури. Поширення цих технологій було зумовлено як технічними обмеженнями IPv4, зокрема вичерпанням доступного адресного простору, так і еволюцією вимог до мережевої безпеки та зручності адміністрування. NAT фактично став проміжним рішенням, яке дозволило відтермінувати глобальний перехід на IPv6 та продовжити використання існуючої інфраструктури IPv4 навіть в умовах експоненційного зростання кількості підключених пристроїв. DHCP, у свою чергу, вирішив проблему масштабованості управління мережевою конфігурацією, дозволивши організаціям ефективно адмініструвати мережі з тисячами та десятками тисяч кінцевих пристроїв без необхідності ручного втручання для кожного вузла.

Вплив NAT на мережеву безпеку є значним, хоча і часто неоднозначним з точки зору класичних принципів побудови мереж. З одного боку, NAT створює природний бар’єр між внутрішньою та зовнішньою мережею, оскільки приховує реальну топологію та адресацію внутрішньої інфраструктури від зовнішніх спостерігачів, що суттєво ускладнює проведення розвідки та

цільових атак на конкретні внутрішні вузли. Механізм PAT додатково підвищує рівень захисту, оскільки за замовчуванням блокує всі несанкціоновані вхідні з'єднання, дозволяючи проходження лише трафіку, що є відповіддю на ініційовані зсередини запити. Ця властивість робить NAT/PAT функціональним аналогом простого stateful firewall для вихідного трафіку. З іншого боку, критики зазначають, що NAT порушує базовий принцип end-to-end connectivity в IP-мережах, ускладнює роботу деяких протоколів та створює додаткові точки відмови в архітектурі мережі, а покладання на NAT як на основний механізм безпеки може призвести до недооцінки необхідності повноцінних систем захисту периметра.

Взаємодія NAT з різними протоколами прикладного рівня створює специфічні виклики, які вимагають застосування додаткових технологій обходу або адаптації протоколів. Протоколи, що вбудовують IP-адреси або номери портів у корисне навантаження пакетів на рівні додатків, можуть некоректно функціонувати при проходженні через NAT без спеціальної обробки. Класичним прикладом є протокол FTP в активному режимі, де сервер ініціює з'єднання назад до клієнта, використовуючи адресу та порт, які клієнт повідомляє в командному каналі, але якщо ця інформація містить приватну адресу, зовнішній сервер не зможе встановити з'єднання. Для вирішення таких проблем на NAT-пристроях застосовуються модулі Application Layer Gateway (ALG), які виконують глибоку інспекцію пакетів на рівні додатків, розпізнають специфічні протоколи та автоматично модифікують відповідні поля в корисному навантаженні, замінюючи приватні адреси на публічні та налаштовуючи необхідні додаткові трансляції портів. Сучасні маршрутизатори зазвичай включають вбудовані ALG для популярних протоколів, таких як FTP, SIP, H.323, TFTP, але адміністратори повинні бути обізнані про потенційні проблеми сумісності та можливість необхідності додаткової конфігурації.

В контексті переходу на IPv6 роль та майбутнє NAT є предметом активних дискусій у мережевій спільноті. Протокол IPv6 був спроектований з урахуванням уроків IPv4 і надає достатній адресний простір для унікальної глобальної адресації кожного пристрою на планеті, теоретично усуваючи необхідність у трансляції адрес для економії адресного простору. Автори стандартів IPv6 передбачали повернення до моделі end-to-end connectivity, де кожен вузол має глобально маршрутизовану адресу і може безпосередньо взаємодіяти з будь-яким іншим вузлом в Інтернеті без посередників у вигляді NAT. Проте на практиці багато організацій продовжують використовувати концепції, схожі на NAT, навіть в IPv6-мережах, мотивуючи це міркуваннями безпеки, контролю та збереження звичних архітектурних патернів. Для IPv6 розроблено технологію NAT66 або NPTv6 (Network Prefix Translation), яка виконує трансляцію префіксів адрес замість повних адрес, дозволяючи приховати внутрішню структуру адресації та спростити міграцію при зміні провайдерів чи злитті мереж, хоча використання таких механізмів суперечить початковим принципам дизайну IPv6 та критикується частиною експертної спільноти.

Протокол DHCP також еволюціонував для підтримки IPv6 у вигляді DHCPv6, який надає аналогічну функціональність автоматичної конфігурації для нового протоколу, але з урахуванням специфіки IPv6-адресації. На відміну від IPv4, де DHCP є фактично єдиним стандартним механізмом автоконфігурації, в IPv6 існує альтернативний механізм Stateless Address Autoconfiguration (SLAAC), який дозволяє вузлам генерувати власні адреси на основі інформації, отриманої від маршрутизаторів через повідомлення Router Advertisement, без необхідності централізованого сервера. Це створює цікаву ситуацію, де організації можуть обирати між використанням DHCPv6 для повного контролю над призначенням адрес, використанням SLAAC для спрощеної та децентралізованої автоконфігурації, або комбінованим підходом, де SLAAC використовується для призначення адрес, а DHCPv6 надає додаткові параметри конфігурації, такі як адреси DNS-серверів. Вибір між цими підходами залежить від конкретних вимог організації до контролю, відстеження клієнтів та інтеграції з існуючими системами управління мережею.

Інтеграція DHCP з системами безпеки та управління мережею відкриває додаткові можливості для контролю доступу та моніторингу активності в мережі. DHCP-сервери можуть взаємодіяти з системами RADIUS або іншими AAA-серверами для перевірки автентичності клієнтів перед призначенням IP-адрес, забезпечуючи додатковий рівень контролю над тим, які пристрої можуть отримати мережевий доступ. Механізми DHCP snooping на комутаторах дозволяють захистити мережу від несанкціонованих DHCP-серверів та атак типу DHCP starvation або spoofing шляхом фільтрації DHCP-повідомлень на основі довірених портів та валідації зв'язків між MAC-адресами, IP-адресами та портами комутатора. Логи DHCP-сервера надають цінну інформацію для аудиту та розслідування інцидентів безпеки, дозволяючи встановити, який пристрій (ідентифікований MAC-адресою) використовував певну IP-адресу в конкретний момент часу, що є критично важливим для forensics та відповідності регуляторним вимогам. Сучасні рішення для управління IP-адресами (IPAM) інтегрують функціональність DHCP та DNS з централізованим веб-інтерфейсом, автоматизацією, аналітикою та інтеграцією з системами моніторингу, забезпечуючи комплексний підхід до управління адресним простором в складних корпоративних середовищах.

Тема 4. Централізована аутентифікація та система AAA

Сучасні корпоративні мережі характеризуються значною кількістю мережевого обладнання, яке потребує постійного адміністрування та контролю доступу. Традиційний підхід до управління обліковими записами, коли кожен маршрутизатор або комутатор налаштовується окремо з локальними обліковими записами адміністраторів, призводить до ряду серйозних проблем. По-перше, це складність централізованого управління паролями та правами доступу, адже зміна пароля або прав одного адміністратора вимагає внесення змін на кожному пристрої окремо. По-друге, відсутність єдиної точки обліку дій адміністраторів ускладнює аудит безпеки та розслідування інцидентів. По-третє, при звільненні співробітника або зміні його посадових обов'язків необхідно вручну видаляти або змінювати облікові записи на всіх пристроях, що створює ризики несанкціонованого доступу.

Для вирішення цих проблем було розроблено концепцію централізованої аутентифікації та авторизації, яка реалізується через модель AAA (Authentication, Authorization, Accounting). Ця модель передбачає винесення функцій перевірки ідентичності користувачів, визначення їх прав доступу та ведення журналу дій на окремі спеціалізовані сервери. Мережеве обладнання при цьому виступає клієнтом AAA-сервера і делегує йому повноваження щодо прийняття рішень про надання або заборону доступу. Такий підхід дозволяє адміністраторам мережі управляти обліковими записами з єдиної точки, застосовувати гнучкі політики безпеки, вести детальний облік дій користувачів та оперативно реагувати на зміни в організаційній структурі підприємства.

4.1 Модель AAA та її компоненти

Модель AAA складається з трьох взаємопов'язаних компонентів, кожен з яких відповідає за певний аспект контролю доступу до мережевих ресурсів. Перший компонент – аутентифікація (Authentication) – відповідає за перевірку ідентичності користувача, який намагається отримати доступ до системи. Процес аутентифікації передбачає надання користувачем облікових даних, найчастіше у вигляді комбінації імені користувача та пароля, хоча можуть використовуватися і більш складні методи, такі як цифрові сертифікати, токени або біометричні дані. Система перевіряє надані облікові дані з інформацією, що зберігається в базі даних, і приймає рішення про дійсність ідентичності користувача. Важливо розуміти, що аутентифікація лише підтверджує, що користувач є тим, за кого себе видає, але не визначає, які дії йому дозволено виконувати в системі.

Другий компонент моделі AAA – авторизація (Authorization) – визначає, які ресурси та операції доступні аутентифікованому користувачу. Після успішної аутентифікації система звертається до профілю користувача або до політик доступу, щоб визначити рівень привілеїв, які йому належать. Наприклад, один адміністратор може мати право лише переглядати конфігурацію маршрутизаторів, тоді як інший – повні права на зміну

налаштувань, включно з можливістю перезавантаження обладнання або зміни протоколів маршрутизації. Авторизація може бути деталізована до рівня окремих команд інтерфейсу командного рядка (CLI), що дозволяє реалізувати принцип найменших привілеїв – кожен користувач отримує тільки ті права, які необхідні для виконання його службових обов'язків. Рішення про авторизацію приймається на основі атрибутів користувача, груп, до яких він належить, та політик безпеки організації.

Третій компонент – облік (Accounting) – забезпечує реєстрацію всіх дій користувачів в системі, створюючи аудиторський слід, який може бути використаний для моніторингу, аналізу та розслідування інцидентів безпеки. Система обліку фіксує такі події, як час входу та виходу з системи, виконані команди, спроби доступу до заборонених ресурсів, зміни конфігурації обладнання. Ці дані зберігаються у вигляді журналів, які можуть аналізуватися як в режимі реального часу для виявлення підозрілої активності, так і ретроспективно для аудиту дотримання політик безпеки. Облік також дозволяє відстежувати використання мережевих ресурсів різними користувачами, що може бути корисним для планування потужностей, оптимізації використання обладнання та розподілу витрат між підрозділами організації. У разі виникнення проблем з безпекою журнали обліку стають критично важливим джерелом інформації для з'ясування обставин інциденту та визначення відповідальних осіб.

Взаємодія між компонентами моделі AAA відбувається в чітко визначеній послідовності, яка забезпечує комплексний контроль доступу. Спочатку користувач проходить аутентифікацію, і лише після підтвердження його ідентичності система переходить до етапу авторизації. Якщо аутентифікація не пройшла успішно, процес зупиняється, і користувачу відмовляється в доступі. Після успішної авторизації користувач отримує доступ до дозволених ресурсів, і всі його подальші дії реєструються компонентом обліку. Важливо відзначити, що система може працювати і в неповному режимі – наприклад, використовувати лише аутентифікацію та авторизацію без обліку, або навіть лише аутентифікацію. Однак для забезпечення високого рівня безпеки рекомендується активувати всі три компоненти моделі AAA. Таблиця 4.1 демонструє порівняння характеристик компонентів моделі AAA.

Таблиця 4.1 – Компоненти моделі AAA та їх функції

Компонент	Основна функція	Типові методи реалізації
Authentication (Аутентифікація)	Перевірка ідентичності користувача, підтвердження, що користувач є тим, за кого себе видає	Пароль, цифровий сертифікат, токен, біометрія, одноразові паролі (OTP)
Authorization (Авторизація)	Визначення прав доступу та дозволених операцій для аутентифікованого користувача	Рольова модель (RBAC), списки контролю доступу (ACL), атрибутна авторизація (ABAC)

Accounting (Облік)	Реєстрація всіх дій користувача в системі, створення аудиторського сліду для аналізу та розслідування	Журнали подій (logs), SIEM-системи, бази даних обліку, системи аналітики
-----------------------	---	--

Реалізація моделі AAA в корпоративному середовищі передбачає розгортання спеціалізованих AAA-серверів, які виконують функції централізованого сховища облікових даних та політик безпеки. Найбільш поширеними протоколами для взаємодії мережевого обладнання з AAA-серверами є RADIUS (Remote Authentication Dial-In User Service) та TACACS+ (Terminal Access Controller Access-Control System Plus). Обидва протоколи підтримують всі три компоненти моделі AAA, однак мають певні відмінності в архітектурі та функціональних можливостях. RADIUS є стандартизованим протоколом, визначеним в RFC 2865 та RFC 2866, і широко використовується не тільки для аутентифікації адміністраторів мережевого обладнання, але й для надання доступу кінцевим користувачам до мережевих сервісів, таких як VPN або бездротові мережі. TACACS+, розроблений компанією Cisco, забезпечує більш детальний контроль над процесом авторизації, дозволяючи розділяти аутентифікацію та авторизацію, а також надає можливість детальної авторизації на рівні окремих команд CLI. Вибір між RADIUS та TACACS+ залежить від конкретних вимог організації до безпеки та функціональності системи контролю доступу.

4.2 Протокол RADIUS: архітектура та принципи роботи

Протокол RADIUS (Remote Authentication Dial-In User Service) є одним з найбільш поширених протоколів для централізованої аутентифікації, авторизації та обліку в комп'ютерних мережах. Розроблений спочатку для забезпечення аутентифікації користувачів комутованого доступу (dial-in), протокол згодом еволюціонував і знайшов широке застосування в різноманітних сценаріях, включаючи аутентифікацію адміністраторів мережевого обладнання, контроль доступу до бездротових мереж, надання доступу до VPN та багато іншого. RADIUS визначений низкою RFC-документів, серед яких основними є RFC 2865 (Authentication) та RFC 2866 (Accounting), що забезпечує стандартизацію та сумісність між реалізаціями різних виробників. Архітектура RADIUS базується на моделі клієнт-сервер, де клієнтом виступає мережевий пристрій, який потребує аутентифікації користувача, а сервером – спеціалізована служба, яка зберігає облікові дані та політики доступу. Взаємодія між клієнтом і сервером відбувається за допомогою обміну повідомленнями по протоколу UDP, що забезпечує високу швидкість обробки запитів, але водночас вимагає реалізації механізмів забезпечення надійності на рівні додатка.

Основою функціонування протоколу RADIUS є обмін повідомленнями між клієнтом та сервером, які інкапсулюються в UDP-дейтаграми. Для процесу аутентифікації та авторизації використовується UDP-порт 1812, а для процесу

обліку – UDP-порт 1813, хоча в деяких застарілих реалізаціях можуть використовуватися порти 1645 та 1646 відповідно. Кожне повідомлення RADIUS складається з заголовка фіксованого розміру та змінної кількості атрибутів, які несуть інформацію про користувача, запитувані сервіси та параметри доступу. Заголовок повідомлення містить поля Code, Identifier, Length та Authenticator. Поле Code визначає тип повідомлення: Access-Request для запиту аутентифікації, Access-Accept для підтвердження успішної аутентифікації, Access-Reject для відмови в доступі, Accounting-Request для надсилання облікової інформації. Поле Identifier використовується для співставлення запитів і відповідей, а поле Authenticator містить криптографічні дані для перевірки цілісності повідомлення та аутентифікації джерела. Атрибути RADIUS являють собою структури типу Type-Length-Value (TLV), де Type визначає тип атрибута, Length – його довжину, а Value – власне значення, яке може бути рядком, числом або структурованими даними.

Процес аутентифікації в RADIUS розпочинається з того, що користувач намагається отримати доступ до мережевого пристрою, наприклад, підключається до маршрутизатора через SSH або telnet. RADIUS-клієнт (мережевий пристрій) запитує у користувача ім'я та пароль і формує повідомлення Access-Request, яке містить атрибути User-Name та User-Password, а також додаткові атрибути, що описують параметри запиту, такі як NAS-IP-Address (IP-адреса мережевого пристрою), NAS-Port (порт, через який користувач підключився) та Service-Type (тип запитуваного сервісу). Пароль користувача не передається у відкритому вигляді – він шифрується за допомогою алгоритму MD5 з використанням спільного секрету (shared secret), який заздалегідь налаштований як на клієнті, так і на сервері. RADIUS-сервер отримує повідомлення Access-Request, дешифрує пароль, перевіряє облікові дані користувача в своїй базі даних та приймає рішення про надання або відмову в доступі. Якщо облікові дані вірні і користувачу дозволено доступ, сервер надсилає повідомлення Access-Accept, яке може містити атрибути авторизації, такі як рівень привілеїв, дозволені команди або час сесії. У разі невірних облікових даних або порушення політик доступу сервер надсилає повідомлення Access-Reject, і користувачу відмовляється в доступі до пристрою.

Важливим аспектом протоколу RADIUS є механізм забезпечення безпеки передачі даних між клієнтом і сервером. Оскільки взаємодія відбувається по відкритій мережі, необхідно захистити облікові дані та іншу чутливу інформацію від перехоплення та несанкціонованого доступу. Для цього RADIUS використовує спільний секрет – текстовий рядок, який відомий як клієнту, так і серверу, але не передається в мережі. Спільний секрет використовується для обчислення криптографічного хешу MD5, який включається в поле Authenticator повідомлення і дозволяє серверу перевірити, що повідомлення дійсно надійшло від довіреного клієнта і не було змінено в процесі передачі. Також спільний секрет використовується для шифрування пароля користувача, який передається в атрибуті User-Password. Алгоритм шифрування базується на обчисленні MD5-хешу від комбінації спільного

секрету, Authenticator та інших даних, після чого пароль піддається операції XOR з отриманим хешем. Незважаючи на ці механізми, стандартний протокол RADIUS не забезпечує шифрування всього трафіку між клієнтом і сервером, що створює потенційні ризики для конфіденційності. Для підвищення безпеки рекомендується використовувати додаткові засоби захисту, такі як VPN-тунелі або протокол RADSEC, який інкапсулює RADIUS-трафік в TLS.

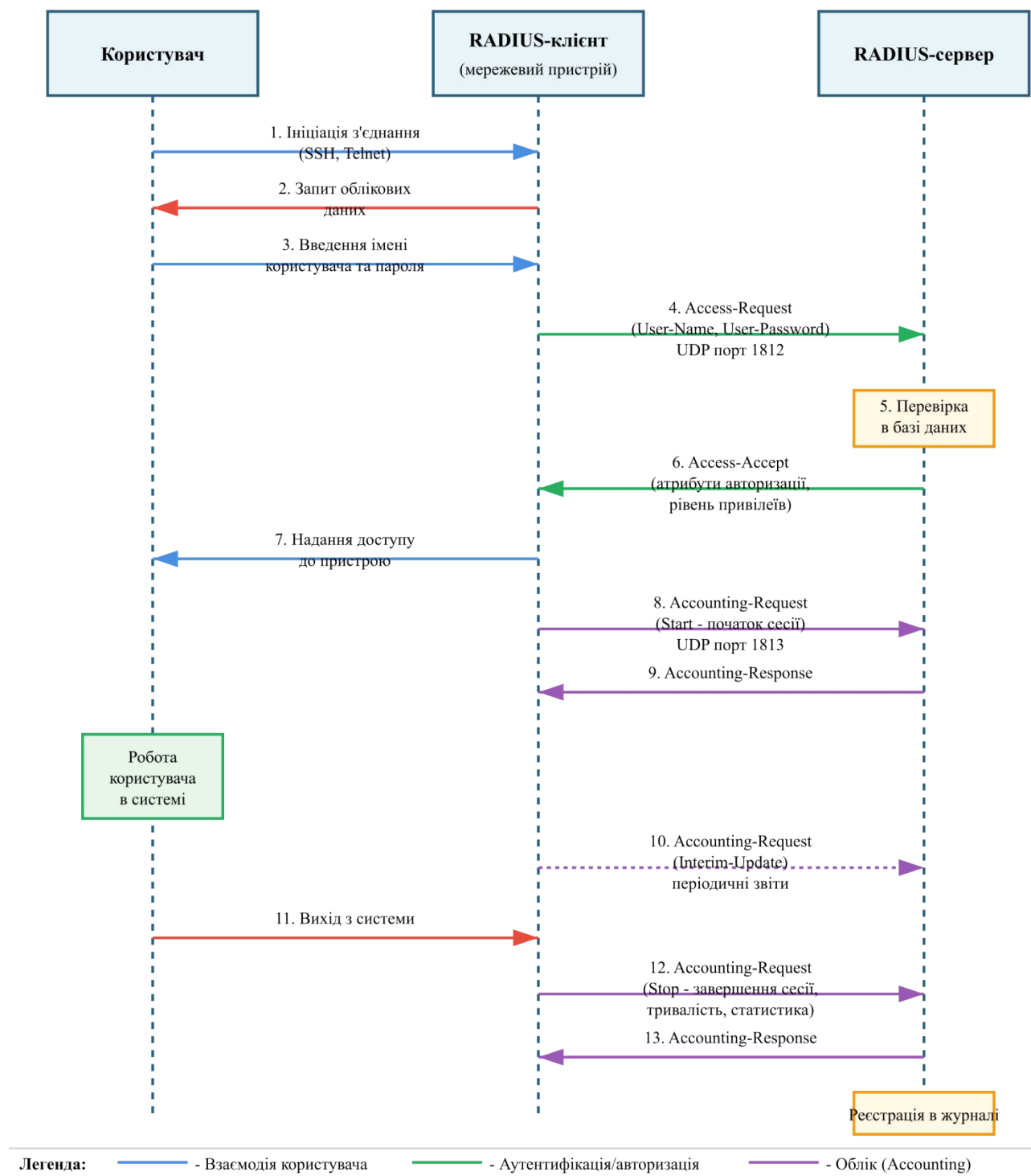


Рисунок 4.1 – Послідовність обміну повідомленнями в протоколі RADIUS

На рисунку 4.1 схематично зображено послідовність взаємодії між користувачем, RADIUS-клієнтом та RADIUS-сервером в процесі аутентифікації. Спочатку користувач ініціює з'єднання з мережевим пристроєм, після чого RADIUS-клієнт запитує облікові дані. Користувач вводить ім'я та пароль, які RADIUS-клієнт упаковує в повідомлення Access-Request та надсилає на RADIUS-сервер по UDP-порту 1812. Сервер перевіряє облікові дані у своїй базі даних і надсилає відповідь Access-Accept або Access-Reject назад до клієнта. У разі успішної аутентифікації користувач отримує доступ до пристрою, а клієнт може надіслати на сервер повідомлення Accounting-Request для реєстрації початку сесії. Протягом роботи сесії клієнт періодично надсилає проміжні облікові повідомлення (Interim-Update), а після завершення сесії надсилає фінальне повідомлення Accounting-Request зі статусом Stop, що дозволяє серверу зафіксувати тривалість сесії та виконані операції. Ця схема ілюструє базовий механізм роботи RADIUS, який може бути розширений додатковими атрибутами та функціями залежно від конкретних вимог інфраструктури.

Таблиця 4.2 – Основні типи повідомлень протоколу RADIUS

Код	Тип повідомлення	Призначення
1	Access-Request	Запит аутентифікації від клієнта до сервера, містить облікові дані користувача
2	Access-Accept	Підтвердження успішної аутентифікації, може містити атрибути авторизації
3	Access-Reject	Відмова в доступі через невірні облікові дані або порушення політик
4	Accounting-Request	Надсилання облікової інформації про сесію користувача (початок, оновлення, завершення)
5	Accounting-Response	Підтвердження отримання облікового повідомлення від сервера
11	Access-Challenge	Запит додаткової інформації для багатфакторної аутентифікації або зміни пароля

Атрибути RADIUS є ключовим механізмом передачі інформації між клієнтом і сервером, дозволяючи гнучко налаштовувати параметри аутентифікації, авторизації та обліку. Протокол визначає стандартний набір атрибутів, кожен з яких має унікальний числовий ідентифікатор та специфічний формат даних. Серед найбільш важливих атрибутів можна виділити User-Name (атрибут 1), який містить ім'я користувача, User-Password (атрибут 2) з зашифрованим паролем, NAS-IP-Address (атрибут 4), що ідентифікує IP-адресу RADIUS-клієнта, Service-Type (атрибут 6) для вказання типу запитуваного сервісу, та Framed-IP-Address (атрибут 8) для призначення IP-адреси користувачу в разі успішної аутентифікації. Для процесу авторизації використовуються атрибути, що визначають права доступу користувача, такі як Filter-Id для застосування списків контролю доступу або Vendor-Specific

Attributes (VSA, атрибут 26), які дозволяють виробникам обладнання розширювати стандартний набір атрибутів власними специфічними параметрами. Облік реалізується за допомогою атрибутів Acct-Status-Type, який вказує на тип облікового повідомлення (Start, Stop, Interim-Update), Acct-Session-Id для ідентифікації сесії, Acct-Session-Time для фіксації тривалості сесії, та Acct-Input-Octets і Acct-Output-Octets для обліку переданого трафіку. Така гнучкість дозволяє адаптувати RADIUS до широкого спектру сценаріїв використання, від простої аутентифікації адміністраторів до складних систем контролю доступу з детальним обліком використання ресурсів.

4.3 Інтеграція маршрутизатора з RADIUS-сервером

Впровадження централізованої аутентифікації на базі RADIUS в корпоративній мережі вимагає належного налаштування як RADIUS-сервера, так і мережевого обладнання, яке виступатиме в ролі RADIUS-клієнта. На прикладі маршрутизаторів Cisco можна розглянути типову послідовність кроків, необхідних для інтеграції обладнання з AAA-інфраструктурою. Перш за все, необхідно активувати функціональність AAA на маршрутизаторі за допомогою глобальної команди конфігурації, яка вмикає підтримку моделі AAA та дозволяє подальше налаштування методів аутентифікації, авторизації та обліку. Після активації AAA адміністратор визначає параметри підключення до RADIUS-сервера, вказуючи IP-адресу або доменне ім'я сервера, порти для аутентифікації та обліку, а також спільний секрет, який використовуватиметься для криптографічного захисту обміну даними. Важливо забезпечити синхронізацію спільного секрету між маршрутизатором та RADIUS-сервером, оскільки будь-яка невідповідність призведе до неможливості успішної аутентифікації.

Після визначення параметрів RADIUS-сервера необхідно сконфігурувати методи аутентифікації для різних типів доступу до маршрутизатора. Cisco IOS підтримує концепцію груп серверів AAA, які дозволяють об'єднати кілька RADIUS-серверів в логічну групу для забезпечення надмірності та балансування навантаження. Адміністратор створює групу серверів, присвоює їй ім'я та додає до неї один або більше RADIUS-серверів з відповідними параметрами підключення. Далі налаштовуються списки методів аутентифікації, які визначають послідовність використання різних механізмів перевірки ідентичності. Наприклад, можна налаштувати список методів для аутентифікації через консольний порт, через лінії VTY (віртуальні термінальні лінії для Telnet та SSH) або для аутентифікації в привілейованому режимі (enable). Кожен список методів може включати кілька методів у певному порядку, таких як RADIUS, локальна база даних маршрутизатора, або дозвіл доступу без аутентифікації. Така гнучкість дозволяє реалізувати резервні механізми аутентифікації, коли в разі недоступності RADIUS-сервера маршрутизатор може звернутися до локальної бази даних, забезпечуючи таким чином безперервність роботи системи адміністрування.

Налаштування авторизації на маршрутизаторі Cisco дозволяє контролювати, які команди може виконувати аутентифікований користувач в різних режимах роботи обладнання. Списки методів авторизації визначають, чи має користувач право виконувати певні команди або переходити в певні режими конфігурації. Наприклад, можна налаштувати авторизацію виконання команд (exec authorization), яка перевіряє права користувача при кожній спробі виконання команди в режимі EXEC, або авторизацію команд конфігурації (commands authorization), яка контролює доступ до команд глобальної конфігурації та конфігурації інтерфейсів. RADIUS-сервер повертає атрибути авторизації, такі як рівень привілеїв (privilege level), які визначають, до яких команд має доступ користувач. В Cisco IOS існує шістнадцять рівнів привілеїв (від 0 до 15), де рівень 0 надає мінімальний доступ лише до базових команд перегляду, рівень 1 є стандартним користувацьким режимом, а рівень 15 надає повний адміністративний доступ до всіх команд конфігурації. RADIUS-сервер може призначати користувачу конкретний рівень привілеїв на основі його ролі в організації, забезпечуючи таким чином гранулярний контроль доступу відповідно до принципу найменших привілеїв.

Компонент обліку в AAA налаштовується для реєстрації дій користувачів та генерації облікових записів, які надсилаються на RADIUS-сервер для централізованого зберігання та аналізу. Списки методів обліку визначають, які події мають реєструватися та куди надсилатися облікова інформація. Типовими подіями для обліку є початок та завершення сесії користувача (exec accounting), виконання окремих команд (commands accounting), а також зміни конфігурації системи (configuration accounting). При налаштуванні обліку адміністратор вказує рівень деталізації реєстрованих подій. Наприклад, можна налаштувати облік всіх виконаних команд або лише команд певного рівня привілеїв. Маршрутизатор формує облікові повідомлення RADIUS, які містять інформацію про користувача, час події, виконану команду, статус виконання та інші релевантні дані, і надсилає їх на RADIUS-сервер. Сервер зберігає отримані дані в базі даних або пересилає їх в системи централізованого збору логів, такі як Syslog або SIEM-системи, де вони можуть аналізуватися для виявлення аномалій, порушень політик безпеки або для проведення аудиту дотримання нормативних вимог.

Таблиця 4.3 – Приклад базової конфігурації AAA на маршрутизаторі Cisco

Команда конфігурації	Призначення
<code>aaa new-model</code>	Активация функціональності AAA на маршрутизаторі
<code>radius server RAD-SRV address ipv4 192.168.1.100 auth-port 1812 acct-port 1813 key MySecret123</code>	Визначення RADIUS-сервера з ім'ям RAD-SRV, IP-адресою 192.168.1.100 та спільним секретом
<code>aaa group server radius RADIUS-GROUP server name RAD-SRV</code>	Створення групи RADIUS-серверів та додавання до неї сервера RAD-SRV

aaa authentication login default group RADIUS-GROUP local	Налаштування аутентифікації через RADIUS з резервним використанням локальної бази
aaa authorization exec default group RADIUS-GROUP local	Авторизація виконання команд через RADIUS з локальним резервом
aaa accounting exec default start-stop group RADIUS-GROUP	Облік сесій користувачів з реєстрацією початку та завершення
line vty 0 4 login authentication default transport input ssh	Застосування налаштованого методу аутентифікації до віртуальних терміналів

Важливим аспектом інтеграції є забезпечення відмовостійкості системи аутентифікації шляхом налаштування резервних механізмів та множинних RADIUS-серверів. Рекомендується розгортати принаймні два RADIUS-сервери в різних сегментах мережі для забезпечення доступності сервісу навіть у разі відмови одного з серверів або проблем з мережевою зв'язністю. Маршрутизатор може бути налаштований для послідовного звертання до серверів у групі, при цьому якщо перший сервер не відповідає протягом заданого таймауту, клієнт автоматично звертається до наступного сервера в списку. Також важливо налаштувати локальну базу даних користувачів на маршрутизаторі як останній резервний механізм, який спрацює в разі повної недоступності всіх RADIUS-серверів. Це гарантує, що адміністратори зможуть отримати доступ до обладнання навіть при повному виході з ладу AAA-інфраструктури, що критично важливо для відновлення роботи системи після збоїв. Однак необхідно забезпечити належний рівень захисту локальних облікових записів, використовуючи надійні паролі та обмежуючи кількість локальних користувачів виключно адміністраторами найвищого рівня.

4.4 Переваги централізованого управління доступом

Впровадження централізованої системи аутентифікації та авторизації на базі RADIUS приносить корпоративним мережам численні переваги, які суттєво підвищують рівень безпеки, спрощують адміністрування та покращують контроль над доступом до критичних мережевих ресурсів. Однією з найбільш очевидних переваг є централізоване управління обліковими записами адміністраторів, яке дозволяє ІТ-департаменту підтримувати єдину базу даних користувачів замість розподіленого налаштування облікових записів на кожному мережевому пристрої окремо. Це значно спрощує процеси створення нових облікових записів при прийомі на роботу нових співробітників, зміни прав доступу при переведенні на іншу посаду та видалення облікових записів при звільненні. Замість необхідності вносити зміни на десятках або сотнях маршрутизаторів і комутаторів, адміністратор може виконати всі необхідні операції в одній точці – на RADIUS-сервері. Такий підхід не лише економить час, але й мінімізує ризик помилок, пов'язаних з

людським фактором, коли облікові записи можуть бути випадково не видалені з деяких пристроїв, створюючи потенційні вразливості безпеки.

Централізована система також значно полегшує впровадження та підтримку політик безпеки паролів в організації. Адміністратори можуть налаштувати на RADIUS-сервері єдині вимоги до складності паролів, терміни їх дії, правила блокування облікових записів після невдалих спроб входу, що автоматично застосовуються до всіх користувачів при доступі до будь-якого мережевого пристрою. Це забезпечує однорідність політик безпеки в усій мережевій інфраструктурі та дозволяє швидко реагувати на нові загрози шляхом оновлення параметрів безпеки в одному місці. Крім того, централізована система дозволяє інтегруватися з корпоративними службами каталогів, такими як Microsoft Active Directory або LDAP, що дає можливість використовувати єдині облікові записи для доступу як до мережевого обладнання, так і до інших корпоративних ресурсів. Така інтеграція підвищує зручність для користувачів, які можуть використовувати один набір облікових даних для доступу до різних систем, та одночасно спрощує управління ідентичністю користувачів в масштабах всієї організації.

Функція обліку в системі AAA надає організації потужний інструмент для аудиту та моніторингу дій адміністраторів, що є критично важливим як з точки зору безпеки, так і для дотримання нормативних вимог. Детальні журнали обліку фіксують інформацію про кожне підключення адміністратора до мережевого обладнання, включаючи час входу та виходу, IP-адресу, з якої виконувалося підключення, виконані команди та результати їх виконання. Ця інформація дозволяє проводити ретроспективний аналіз дій адміністраторів при розслідуванні інцидентів безпеки, визначати, хто та коли вніс зміни в конфігурацію критичних систем, виявляти несанкціоновані спроби доступу або підозрілу активність. У багатьох індустріях, таких як фінансовий сектор, охорона здоров'я або державні установи, ведення детальних журналів доступу до критичних систем є обов'язковою вимогою регуляторів, і система AAA з функцією обліку допомагає організаціям демонструвати відповідність цим вимогам під час аудитів. Централізоване зберігання логів на RADIUS-сервері також захищає їх від можливої маніпуляції з боку злоумисників, які отримали доступ до окремого мережевого пристрою.

Гнучкість в налаштуванні рівнів доступу є ще однією важливою перевагою централізованої системи AAA. Використовуючи механізми авторизації RADIUS, організація може реалізувати деталізовану рольову модель доступу, де різні групи адміністраторів отримують різні рівні привілеїв відповідно до їх посадових обов'язків. Наприклад, адміністратори першої лінії підтримки можуть мати права лише на перегляд стану обладнання та виконання базових діагностичних команд, не маючи можливості вносити зміни в конфігурацію. Мережеві інженери можуть отримати права на зміну конфігурації маршрутизації та комутації, але без доступу до критичних параметрів безпеки. Старші адміністратори безпеки можуть мати повний доступ до всіх функцій обладнання. Така градація прав доступу реалізує принцип найменших привілеїв, коли кожен користувач отримує тільки ті права,

які необхідні для виконання його робочих завдань, що мінімізує потенційні наслідки помилок або зловживань.

Ще однією перевагою централізованої системи є можливість швидкого реагування на інциденти безпеки шляхом оперативної зміни прав доступу або блокування облікових записів. У разі виявлення підозрілої активності або компрометації облікового запису адміністратор безпеки може негайно заблокувати цей обліковий запис на RADIUS-сервері, що автоматично унеможливить подальший доступ користувача до всього мережевого обладнання. Це набагато ефективніше, ніж необхідність послідовно блокувати облікові записи на десятках окремих пристроїв, особливо в умовах активного інциденту, коли час реакції є критичним. Централізована система також дозволяє швидко змінювати паролі всіх адміністраторів у випадку виявлення витоку облікових даних або як частину планових заходів з підвищення безпеки. Інтеграція RADIUS з системами моніторингу безпеки (SIEM) дозволяє автоматизувати процес виявлення аномалій, таких як спроби входу з незвичайних локацій, множинні невдалі спроби аутентифікації або підключення в нетипові години, що дає можливість проактивно виявляти та запобігати інцидентам безпеки ще на ранніх стадіях.

Таким чином, впровадження централізованої аутентифікації на базі протоколу RADIUS та моделі AAA є важливим кроком у напрямку підвищення безпеки та ефективності управління корпоративною мережевою інфраструктурою. Ця технологія забезпечує єдину точку контролю над доступом адміністраторів до мережевого обладнання, дозволяє реалізувати гранулярні політики авторизації відповідно до ролей та відповідальності співробітників, веде детальний облік всіх дій для аудиту та розслідування інцидентів. Хоча впровадження централізованої системи вимагає первинних інвестицій в обладнання серверів, програмне забезпечення та час на налаштування, довгострокові переваги у вигляді підвищеної безпеки, спрощеного адміністрування та відповідності нормативним вимогам значно переважають ці витрати. В умовах постійно зростаючих кіберзагроз та все більш жорстких вимог регуляторів централізована система AAA стає не просто бажаною опцією, а необхідним компонентом сучасної захищеної мережевої інфраструктури будь-якої організації, яка серйозно ставиться до питань інформаційної безпеки.

Тема 5. Тунельні технології захисту на основі IPsec VPN

5.1. Архітектура IPsec та основні компоненти

Сучасні корпоративні мережі потребують захищеного обміну даними між географічно розподіленими філіалами, мобільними користувачами та партнерськими організаціями. Традиційні методи побудови приватних мереж на основі виділених каналів зв'язку є економічно не вигідними та не забезпечують необхідної гнучкості. Технологія IPsec (Internet Protocol Security) надає можливість створення захищених віртуальних приватних мереж (VPN) поверх публічної інфраструктури Інтернет, забезпечуючи конфіденційність, цілісність та автентичність переданих даних. IPsec являє собою набір протоколів та алгоритмів, які працюють на мережевому рівні моделі OSI, що дозволяє захищати будь-який тип трафіку без модифікації прикладних програм.

Архітектура IPsec базується на модульному підході, де кожен компонент виконує специфічні функції безпеки. Основу складають два протоколи захисту: Authentication Header (AH) та Encapsulating Security Payload (ESP), які можуть використовуватися окремо або у комбінації залежно від вимог безпеки. Протокол AH забезпечує автентифікацію джерела даних та цілісність IP-пакетів шляхом обчислення криптографічної хеш-функції над незмінними полями IP-заголовок та корисним навантаженням. Цей протокол додає до пакета додатковий заголовок, що містить значення ICV (Integrity Check Value), яке дозволяє приймаючій стороні переконатися, що пакет не був модифікований під час передачі. Однак AH не забезпечує шифрування даних, тому конфіденційність трафіку залишається незахищеною.

Протокол ESP надає більш широкий спектр послуг безпеки, включаючи не тільки автентифікацію та цілісність, але й конфіденційність даних через симетричне шифрування. ESP додає до пакета заголовок та трейлер, між якими розміщується зашифроване корисне навантаження. Залежно від режиму роботи, ESP може шифрувати лише дані користувача або весь оригінальний IP-пакет разом із заголовком. Автентифікація в ESP є опціональною, але на практиці майже завжди використовується для забезпечення комплексного захисту. Сучасні реалізації IPsec віддають перевагу протоколу ESP через його універсальність та можливість одночасного забезпечення конфіденційності й цілісності.

Крім протоколів захисту, архітектура IPsec включає протокол Internet Key Exchange (IKE), який відповідає за встановлення захищених з'єднань та управління ключами шифрування. IKE працює у двох фазах, кожна з яких виконує специфічні завдання. Перша фаза (IKE Phase 1) встановлює захищений канал між пірамі IPsec, що називається ISAKMP SA (Internet Security Association and Key Management Protocol Security Association). У цій фазі сторони автентифікують одна одну, узгоджують алгоритми шифрування та хешування, обмінюються криптографічним матеріалом за допомогою алгоритму Diffie-Hellman. Результатом першої фази є двонаправлений

захищений тунель, який використовується для безпечного обміну повідомленнями другої фази.

Друга фаза IKE (IKE Phase 2) відповідає за створення IPsec SA, які безпосередньо захищають користувацький трафік. У цій фазі використовується вже встановлений захищений канал з першої фази для узгодження параметрів захисту даних: вибору протоколу (AH або ESP), алгоритмів шифрування та автентифікації, часу життя ключів. Важливо зазначити, що IKE Phase 1 SA є двонаправленою, тоді як IPsec SA є однонаправленими, тому для повноцінної двосторонньої комунікації потрібна пара SA. Періодична реноціація ключів, передбачена в IKE Phase 2, підвищує безпеку системи, обмежуючи обсяг даних, зашифрованих одним ключем.

Таблиця 5.1 наведена нижче демонструє порівняльні характеристики протоколів AH та ESP в контексті забезпечення різних аспектів безпеки. Як видно з таблиці, ESP є більш універсальним рішенням для сучасних VPN-мереж, оскільки забезпечує повний набір функцій безпеки.

Таблиця 5.1 - Порівняння протоколів AH та ESP

Характеристика	AH	ESP
Автентифікація джерела	Так	Так (опційно)
Цілісність даних	Так	Так (опційно)
Конфіденційність	Ні	Так
Захист від повторів	Так	Так
Номер протоколу IP	51	50
Сумісність з NAT	Обмежена	Краща (з NAT-T)
Розмір додаткових даних	24 байти	25-57+ байтів

5.2. Режими роботи IPsec та концепція Security Association

IPsec може функціонувати у двох основних режимах: транспортному (transport mode) та тунельному (tunnel mode), які відрізняються обсягом захищених даних та сценаріями застосування. У транспортному режимі IPsec захищає тільки корисне навантаження оригінального IP-пакета, залишаючи оригінальний IP-заголовок незмінним. Цей режим застосовується переважно для захисту з'єднань кінець-до-кінця між двома хостами або між хостом та шлюзом безпеки. При використанні ESP в транспортному режимі заголовок ESP вставляється між оригінальним IP-заголовком та заголовком транспортного рівня (TCP або UDP), після чого дані транспортного рівня шифруються. Транспортний режим економить пропускну здатність каналу, оскільки не додає додаткового IP-заголовка, але він розкриває інформацію про кінцеві точки комунікації.

Тунельний режим IPsec забезпечує більш високий рівень безпеки, інкапсулюючи весь оригінальний IP-пакет у новий IP-пакет з іншими адресами відправника та отримувача. У цьому режимі оригінальний пакет повністю шифрується та автентифікується, після чого додається новий зовнішній IP-

заголовок, який вказує адреси VPN-шлюзів. Тунельний режим є стандартним вибором для побудови site-to-site VPN між корпоративними мережами, оскільки він приховує внутрішню топологію мережі та дозволяє використовувати приватні IP-адреси у захищеному тунелі. Зовнішній спостерігач бачить лише комунікацію між VPN-шлюзами, не маючи інформації про реальні джерела та призначення трафіку всередині корпоративних мереж.

Концепція Security Association (SA) є фундаментальною для роботи IPsec та представляє собою логічне з'єднання, яке визначає всі параметри безпеки для потоку даних в одному напрямку. SA ідентифікується трьома параметрами: Security Parameter Index (SPI), IP-адресою призначення та протоколом безпеки (AH або ESP). SPI є 32-бітовим ідентифікатором, який включається в заголовок AH або ESP і дозволяє приймаючому пристрою визначити, які параметри безпеки використовувати для обробки пакета. Кожна SA містить інформацію про алгоритми шифрування та автентифікації, ключі, режим роботи (транспортний або тунельний), час життя SA та лічильник послідовності для захисту від атак повтору.

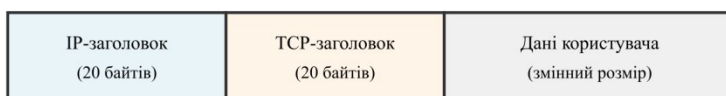
Управління SA здійснюється через базу даних Security Association Database (SAD), яка зберігається на кожному IPsec-пристрої. SAD містить записи про всі активні SA з усією необхідною інформацією для обробки вхідних та вихідних пакетів. Коли пристрій отримує IPsec-пакет, він використовує SPI, IP-адресу та протокол для пошуку відповідного SA в SAD і застосовує необхідні операції дешифрування та перевірки автентичності. Для вихідного трафіку пристрій спочатку консультується з Security Policy Database (SPD), яка визначає, який трафік потребує захисту IPsec, після чого вибирає або створює відповідний SA для обробки пакетів.

Час життя SA обмежується або за часом (наприклад, 3600 секунд), або за обсягом переданих даних (наприклад, 100 МБ), залежно від політики безпеки організації. Після закінчення терміну дії SA автоматично запускається процес реноціації через IKE Phase 2, що дозволяє створити нові ключі шифрування без переривання потоку даних. Механізм “soft lifetime” та “hard lifetime” забезпечує плавний перехід на нові ключі: коли досягається soft lifetime, ініціюється створення нового SA, але старий продовжує використовуватися до досягнення hard lifetime або до успішного встановлення нового SA. Такий підхід забезпечує безперервність сервісу та підвищує стійкість до криптоаналітичних атак, обмежуючи кількість даних, зашифрованих одним ключем.

Рисунок 5.1 ілюструє структуру IPsec-пакета в тунельному режимі з використанням протоколу ESP. На рисунку зображено оригінальний IP-пакет, що складається з IP-заголовка (20 байтів), TCP-заголовка (20 байтів) та даних користувача. Після застосування ESP в тунельному режимі структура змінюється наступним чином: спочатку йде новий зовнішній IP-заголовок з адресами VPN-шлюзів, далі ESP-заголовок (8 байтів) з SPI та порядковим номером, потім зашифрована частина, яка включає оригінальний IP-заголовок, TCP-заголовок, дані користувача та ESP-трейлер з padding, і завершується ESP-автентифікацією (12-24 байти). Зашифрована область на рисунку виділена

сірим кольором, а область, що охоплюється автентифікацією, позначена пунктирною лінією над пакетом.

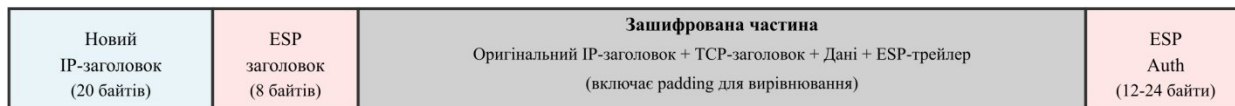
Оригінальний IP-пакет:



Застосування ESP в тунельному режимі

Автентифікація (Authentication)

IPsec-пакет (ESP Tunnel Mode):



Деталізація зашифрованої частини:



Області захисту:

Шифрування (Encryption)

Легенда:



- зашифрована область



- ESP заголовки



- IP заголовки

Рисунок 5.1 - Структура IPsec-пакета в тунельному режимі з ESP

5.3. Налаштування IPsec VPN на маршрутизаторах Cisco

Практична реалізація site-to-site IPsec VPN між граничними маршрутизаторами Cisco потребує послідовного виконання кількох етапів конфігурації, кожен з яких відповідає за специфічний аспект захищеного з'єднання. Перший етап полягає у визначенні цікавого трафіку (interesting traffic), тобто тих IP-потоків, які повинні захищатися через IPsec-тунель. Це досягається створенням розширеного списку контролю доступу (extended ACL), який специфікує адреси джерел та призначень у локальних мережах за VPN-шлюзами. Важливо розуміти, що цей ACL використовується не для фільтрації трафіку, а виключно для класифікації пакетів, які підлягають обробці IPsec.

Другий етап передбачає конфігурацію параметрів IKE Phase 1 через створення політики ISAKMP (Internet Security Association and Key Management Protocol). ISAKMP-політика визначає набір алгоритмів та параметрів, які будуть використовуватися для встановлення захищеного каналу між пірами: алгоритм шифрування (DES, 3DES, AES), алгоритм хешування (MD5, SHA-1, SHA-256), метод автентифікації (pre-shared keys або RSA-підписи), групу Diffie-Hellman для обміну ключами та час життя SA. На маршрутизаторах Cisco можна створити декілька ISAKMP-політик з різними пріоритетами, і під час переговорів пристрої автоматично вибирають найбільш пріоритетну політику, параметри якої підтримуються обома сторонами. Для простих сценаріїв зазвичай використовується автентифікація на основі попередньо розподілених

ключів (pre-shared keys), коли на обох маршрутизаторах конфігурується однаковий секретний ключ.

Третій етап включає налаштування параметрів IKE Phase 2 через створення transform set та crypto map. Transform set визначає комбінацію протоколу безпеки (ESP або AH), алгоритму шифрування та алгоритму автентифікації для захисту користувацького трафіку. Один transform set може містити до чотирьох різних наборів перетворень, і під час переговорів Phase 2 сторони вибирають перший спільний набір. Crypto map є ключовою конструкцією, яка зв'язує всі компоненти IPsec-конфігурації: вона асоціює ACL цікавого трафіку з transform set, вказує IP-адресу віддаленого піра та визначає режим роботи (transport або tunnel). Після створення crypto map її необхідно прив'язати до зовнішнього інтерфейсу маршрутизатора, через який здійснюється VPN-з'єднання.

Таблиця 5.2 демонструє приклад команд конфігурації IPsec VPN на маршрутизаторі Cisco для з'єднання двох філіалів компанії. У таблиці наведено послідовність команд з поясненнями їхнього призначення.

Таблиця 5.2 - Приклад конфігурації site-to-site IPsec VPN на Cisco

Команда	Призначення
crypto isakmp policy 10	Створення ISAKMP-політики з пріоритетом 10
encryption aes 256	Використання AES-256 для шифрування Phase 1
hash sha256	Використання SHA-256 для автентифікації повідомлень
authentication pre-share	Автентифікація на основі pre-shared key
group 14	Використання DH-групи 14 (2048-біт)
crypto isakmp key MySecretKey123 address 203.0.113.2	Визначення pre-shared key для віддаленого піра
crypto ipsec transform-set MYSET esp-aes 256 esp-sha256-hmac	Створення transform set з ESP-AES-256 та HMAC-SHA-256
mode tunnel	Встановлення тунельного режиму для transform set
access-list 101 permit ip 192.168.1.0 0.0.0.255 192.168.2.0 0.0.0.255	ACL для визначення цікавого трафіку між мережами
crypto map VPNMAP 10 ipsec-isakmp	Створення crypto map з порядковим номером 10
set peer 203.0.113.2	Вказання IP-адреси віддаленого VPN-шлюзу
set transform-set MYSET	Прив'язка transform set до crypto map
match address 101	Асоціація ACL з crypto map
interface GigabitEthernet0/0	Вхід у режим конфігурації зовнішнього інтерфейсу
crypto map VPNMAP	Застосування crypto map до

Четвертий етап конфігурації стосується додаткових параметрів для підвищення надійності та сумісності VPN-з'єднання. Зокрема, часто виникає проблема з Maximum Transmission Unit (MTU), коли додавання IPsec-заголовків збільшує розмір пакета понад MTU каналу, що призводить до фрагментації або втрати пакетів. Для вирішення цієї проблеми рекомендується налаштувати TCP MSS (Maximum Segment Size) adjustment на VPN-інтерфейсах, що примушує TCP-сесії використовувати менший розмір сегментів. Типове значення MSS для IPsec VPN становить 1360 байтів замість стандартних 1460 байтів. Також важливо переконатися, що на міжмережевих екранах та проміжних маршрутизаторах дозволений прохід трафіку UDP порт 500 (ISAKMP) та протоколу ESP (IP протокол 50).

П'ятий етап включає налаштування механізмів підтримки з'єднання та автоматичного відновлення після збоїв. Keepalive-механізми дозволяють швидко виявляти недоступність віддаленого піра та ініціювати процес відновлення тунелю. На маршрутизаторах Cisco можна використовувати Dead Peer Detection (DPD), який періодично відправляє запити для перевірки доступності віддаленої сторони. Також корисно налаштувати криптографічні lifetime значення, які визначають, як часто IPsec SA повинні реногоціюватися. Занадто короткий lifetime створює надмірне навантаження на процесор через часту реногоціацію, тоді як надто довгий lifetime знижує безпеку. Типові значення становлять 3600 секунд (1 година) для ISAKMP SA та 1800 секунд (30 хвилин) для IPsec SA.

Верифікація працездатності IPsec VPN здійснюється через низку діагностичних команд, які надають інформацію про стан тунелю, статистику трафіку та можливі помилки. Команда “show crypto isakmp sa” відображає активні Phase 1 SA з інформацією про піра, стан з'єднання та використані алгоритми. Команда “show crypto ipsec sa” надає детальну інформацію про Phase 2 SA, включаючи лічильники зашифрованих та дешифрованих пакетів, кількість помилок автентифікації та поточний SPI. Для відлагодження проблем корисна команда “debug crypto isakmp” та “debug crypto ipsec”, які виводять детальну інформацію про процес переговорів та обробки пакетів. При аналізі проблем важливо перевіряти синхронізацію часу між пірами, оскільки великі розбіжності можуть спричинити відмову автентифікації при використанні сертифікатів.

5.4. Розширені функції та інтеграція IPsec у корпоративну інфраструктуру

Розгортання IPsec VPN у реальному корпоративному середовищі потребує врахування додаткових аспектів, таких як масштабованість, інтеграція з існуючою інфраструктурою безпеки та забезпечення високої доступності. Одним з критичних питань є сумісність IPsec з технологією Network Address Translation (NAT), яка широко використовується в корпоративних мережах для

економії публічних IP-адрес. Традиційний IPsec, особливо протокол AH, має фундаментальні проблеми з NAT, оскільки зміна IP-адрес у заголовку пакета порушує криптографічні хеші, що обчислюються для перевірки цілісності. Навіть ESP може мати проблеми, коли NAT виконується між IPsec-пірами, оскільки UDP-інкапсуляція стандартного IPsec не підтримується деякими NAT-пристроями.

Для вирішення проблеми сумісності з NAT було розроблено розширення NAT Traversal (NAT-T), описане в RFC 3947 та RFC 3948, яке дозволяє IPsec-трафіку проходити через NAT-пристрої без порушення цілісності захисту. NAT-T працює шляхом інкапсуляції ESP-пакетів у UDP-дейтаграми з портом призначення 4500, що дозволяє NAT-пристроєм коректно обробляти IPsec-трафік та виконувати трансляцію портів. Під час встановлення з'єднання IKE автоматично визначає наявність NAT-пристроїв на шляху між пірами, обмінюючись спеціальними хеш-значеннями IP-адрес та портів. Якщо виявлено NAT, обидві сторони автоматично перемикаються на використання NAT-T режиму. Сучасні маршрутизатори Cisco підтримують NAT-T автоматично без додаткової конфігурації, але адміністратор повинен переконатися, що міжмережеві екрани пропускають UDP трафік на порту 4500.

Інтеграція IPsec VPN з корпоративними системами моніторингу та управління є критично важливою для забезпечення операційної стабільності та швидкого реагування на інциденти. IPsec-шлюзи можна налаштувати для відправки Syslog-повідомлень про важливі події, такі як встановлення та розрив тунелів, помилки автентифікації, вичерпання lifetime SA та інші аномалії. Через SNMP можна збирати статистику про використання тунелів, завантаженість процесора на криптографічні операції та стан апаратних акселераторів шифрування. Централізований збір цієї інформації дозволяє адміністраторам виявляти проблеми на ранніх стадіях, планувати масштабування інфраструктури та проводити ретроспективний аналіз інцидентів безпеки.

Високодоступні архітектури IPsec VPN використовують кілька підходів для забезпечення безперервності сервісу при відмові обладнання або каналів зв'язку. Один з методів полягає у використанні резервних VPN-шлюзів з автоматичним перемиканням за допомогою протоколів динамічної маршрутизації, таких як OSPF або BGP. У такій конфігурації кожна сторона має два або більше VPN-шлюзів, між якими встановлюються множинні IPsec-тунелі. Динамічна маршрутизація через тунелі дозволяє автоматично перенаправляти трафік на резервний шлюз при відмові основного. Альтернативний підхід використовує Hot Standby Router Protocol (HSRP) або Virtual Router Redundancy Protocol (VRRP) для створення віртуального VPN-шлюзу, за яким стоять кілька фізичних маршрутизаторів, що забезпечують прозоре перемикання при збоях.

Оптимізація продуктивності IPsec VPN є важливим аспектом при роботі з високими обсягами трафіку або ресурсомістких додатків. Криптографічні операції шифрування та автентифікації створюють значне навантаження на процесор маршрутизатора, що може стати вузьким місцем у мережі. Для підвищення продуктивності використовуються апаратні криптографічні

акселератори, які виконують операції шифрування на спеціалізованих чіпах, розвантажуючи центральний процесор. Сучасні корпоративні маршрутизатори Cisco серій ASR та ISR мають вбудовані криптографічні модулі, які можуть обробляти VPN-трафік на швидкостях до кількох гігабіт за секунду. При виборі обладнання важливо враховувати не тільки номінальну пропускну здатність інтерфейсів, але й криптографічну продуктивність, яка зазвичай значно нижча.

Управління ключами в великих IPsec-розгортаннях з десятками або сотнями тунелів представляє серйозну адміністративну проблему при використанні pre-shared keys. Кожен новий тунель потребує конфігурації унікального секретного ключа на обох кінцях, що ускладнює масштабування та створює ризики безпеки через необхідність ручного розповсюдження ключів. Для вирішення цієї проблеми в корпоративних середовищах використовується автентифікація на основі цифрових сертифікатів та інфраструктури відкритих ключів (PKI). У такій схемі кожен VPN-шлюз отримує цифровий сертифікат від довіреного центру сертифікації (CA), який використовується для автентифікації під час IKE Phase 1. Це дозволяє встановлювати тунелі без попереднього розповсюдження секретних ключів, спрощує адміністрування та забезпечує кращий контроль над життєвим циклом облікових даних.

Рисунок 5.2 ілюструє типову топологію корпоративної IPsec VPN-мережі типу hub-and-spoke для підключення віддалених філіалів до центрального офісу. На рисунку зображено центральний офіс (hub) з двома резервними VPN-концентраторами для забезпечення високої доступності, підключеними до внутрішньої корпоративної мережі 10.0.0.0/16 та до Інтернет через різних провайдерів. Три віддалених філії (spoke sites) з приватними мережами 192.168.1.0/24, 192.168.2.0/24 та 192.168.3.0/24 підключаються до центрального офісу через IPsec-тунелі, позначені на рисунку пунктирними лініями з замками. Кожен філіал має граничний маршрутизатор з функціями VPN, який захищає трафік між локальною мережею філіалу та центральним офісом. На рисунку також показано сервери корпоративних служб (DHCP, DNS, контролери домену) у центральному офісі, доступ до яких отримують користувачі з філіалів через захищені VPN-тунелі.

Інтеграція IPsec VPN з системами виявлення вторгнень (IDS/IPS) потребує особливої уваги, оскільки шифрування трафіку унеможливорює інспектування вмісту пакетів на проміжних пристроях безпеки. Для забезпечення комплексного захисту IDS/IPS повинні розміщуватися у точках, де трафік вже дешифрований: або на самих VPN-шлюзах після обробки IPsec, або на окремих сенсорах безпеки у внутрішній мережі за VPN-шлюзом. Деякі сучасні платформи об'єднують функції VPN-концентратора та системи виявлення вторгнень в одному пристрої, що дозволяє аналізувати трафік одразу після дешифрування. При проектуванні архітектури безпеки важливо враховувати, що VPN захищає дані під час передачі каналом, але не гарантує безпеку кінцевих точок, тому необхідні додаткові засоби контролю доступу та моніторингу всередині мереж.

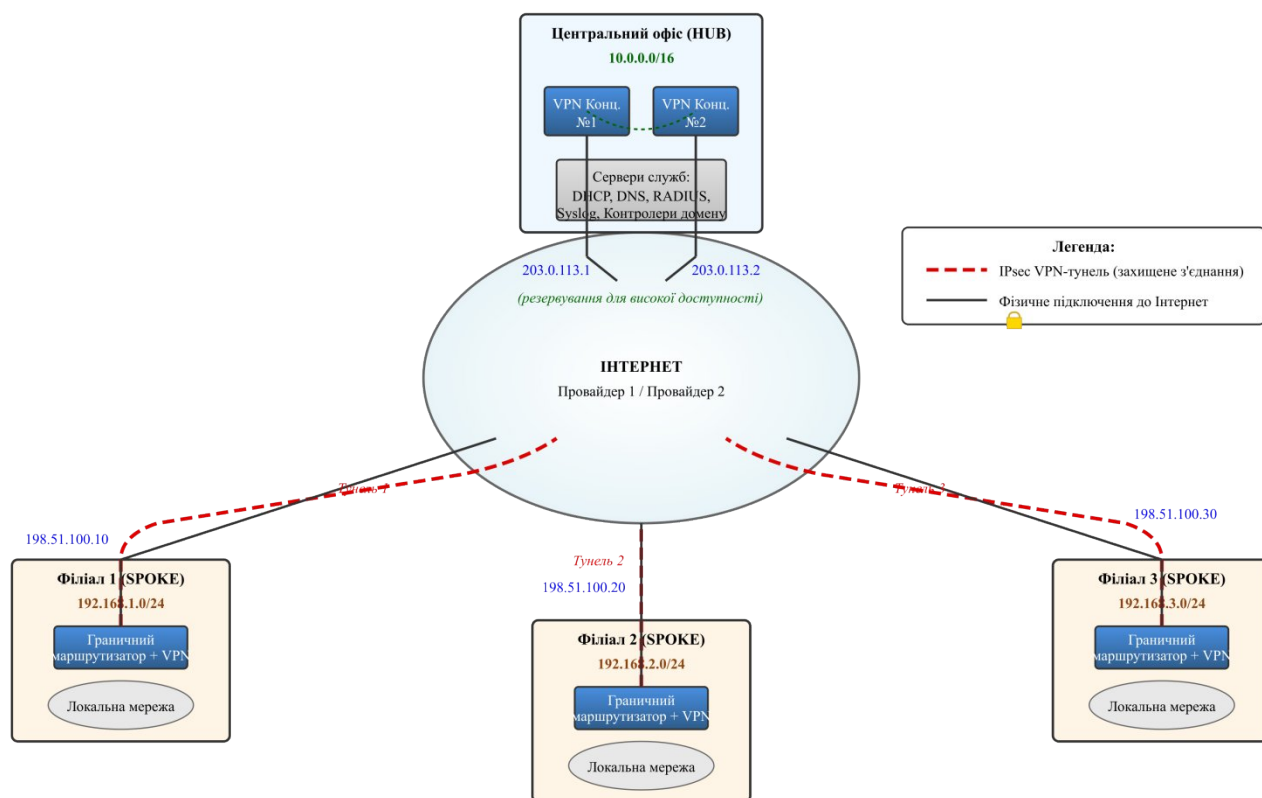


Рисунок 5.2 - Топологія корпоративної IPsec VPN-мережі hub-and-spoke

Таблиця 5.3 узагальнює основні рекомендації щодо вибору криптографічних алгоритмів для IPsec VPN з урахуванням балансу між безпекою та продуктивністю. Наведені параметри відповідають сучасним стандартам безпеки та рекомендаціям NIST.

Таблиця 5.3 - Рекомендовані криптографічні параметри для IPsec VPN

Параметр	Мінімальна конфігурація	Рекомендована конфігурація	Високий рівень безпеки
Шифрування Phase 1	AES-128	AES-256	AES-256
Хешування Phase 1	SHA-1 (застаріло)	SHA-256	SHA-384 або SHA-512
ДН-група Phase 1	Група 2 (1024-біт)	Група 14 (2048-біт)	Група 19 або 20 (ECC)
Шифрування Phase 2	AES-128	AES-256	AES-256-GCM
Автентифікація Phase 2	HMAC-SHA-1	HMAC-SHA-256	Включено в GCM
Lifetime ISAKMP SA	86400 сек (24 год)	28800 сек (8 год)	3600 сек (1 год)
Lifetime IPsec SA	3600 сек (1 год)	1800 сек (30 хв)	900 сек (15 хв)
Perfect Forward Secrecy	Відключено	Група 14	Група 19 або 20

Майбутній розвиток технології IPsec пов'язаний з переходом на нове покоління протоколу Internet Key Exchange версії 2 (IKEv2), яке забезпечує значні покращення порівняно з класичним IKEv1. IKEv2 характеризується спрощеним процесом встановлення з'єднання, кращою підтримкою мобільних пристроїв через механізм MOBIKE (Mobility and Multihoming Protocol), вбудованою підтримкою NAT-T та більш надійним виявленням збоїв. Встановлення IPsec-тунелю через IKEv2 вимагає менше обмінів повідомленнями, що знижує затримки та споживання ресурсів. Також IKEv2 краще захищений від атак типу denial-of-service через вбудовані механізми протидії, такі як cookie-челенджі та обов'язкова перевірка стану під час переговорів. Поступовий перехід корпоративних мереж на IKEv2 дозволяє підвищити ефективність та безпеку VPN-інфраструктури при збереженні сумісності з існуючими системами через підтримку обох версій протоколу на сучасному обладнанні.

Тема 6. Централізований моніторинг та діагностика мережі

6.1. Основні концепції централізованого моніторингу мережевої інфраструктури

Сучасні комп'ютерні мережі характеризуються високою складністю архітектури, наявністю великої кількості активного обладнання та розподіленою топологією, що охоплює географічно віддалені сегменти. За таких умов традиційні методи адміністрування, що передбачають локальне підключення до кожного пристрою для перевірки його стану, стають неефективними та ресурсовитратними. Централізований моніторинг та діагностика мережі представляють собою комплексний підхід до збору, обробки та аналізу інформації про функціонування мережевих компонентів з єдиного управляючого центру. Цей підхід дозволяє адміністраторам оперативно виявляти несправності обладнання, відстежувати зміни конфігурації, аналізувати завантаження каналів зв'язку та своєчасно реагувати на інциденти інформаційної безпеки без необхідності фізичного доступу до кожного мережевого вузла.

Основними завданнями систем централізованого моніторингу є збір статистичних даних про роботу мережевих інтерфейсів, відстеження стану портів комутаторів та маршрутизаторів, контроль використання ресурсів процесора та оперативної пам'яті на активному обладнанні, а також реєстрація подій безпеки. Для вирішення цих завдань використовуються спеціалізовані протоколи обміну інформацією між мережевими пристроями та серверами моніторингу, серед яких найбільш поширеними є протокол простого управління мережею SNMP та протокол системного журналювання Syslog. Ці протоколи мають різну архітектуру та призначення, проте доповнюють один одного в контексті побудови комплексної системи спостереження за станом мережі. SNMP орієнтований на регулярний опитування пристроїв для отримання числових метрик та можливість віддаленої зміни параметрів конфігурації, тоді як Syslog призначений для асинхронної передачі повідомлень про події різного рівня критичності від мережевого обладнання до централізованого сервера журналювання.

Впровадження централізованого моніторингу суттєво підвищує ефективність експлуатації мережевої інфраструктури шляхом скорочення часу реагування на інциденти, зменшення тривалості простоїв обладнання та оптимізації використання пропускної здатності каналів зв'язку. Системи моніторингу дозволяють не лише виявляти проблеми в режимі реального часу, але й накопичувати історичні дані для аналізу трендів, прогнозування навантаження та планування модернізації мережі. Крім того, централізоване збереження журналів подій забезпечує виконання вимог регуляторних стандартів щодо аудиту інформаційних систем та розслідування інцидентів безпеки. Архітектура типової системи моніторингу включає агентів або вбудовані механізми на керованих пристроях, мережеві протоколи передачі

даних, сервери збору та зберігання інформації, а також засоби візуалізації та оповіщення адміністраторів про критичні події.

6.2. Протокол SNMP як основа інтелектуального управління мережею

Протокол простого управління мережею SNMP є стандартом прикладного рівня моделі OSI, призначеним для моніторингу та управління мережевими пристроями в IP-мережах. Архітектура SNMP базується на взаємодії двох основних компонентів: менеджера управління, який виконує роль централізованого сервера моніторингу, та агентів, що функціонують на керованих пристроях і надають доступ до інформації про їх стан. Менеджер ініціює запити до агентів для отримання значень певних параметрів, таких як завантаження процесора, кількість переданих пакетів через інтерфейс, або поточна температура обладнання, а також може змінювати конфігураційні параметри шляхом відправлення команд встановлення значень. Агенти, у свою чергу, можуть ініціювати відправлення сповіщень-трапів до менеджера при виникненні певних подій, наприклад, при зміні стану інтерфейсу з активного на неактивний або при перевищенні порогового значення завантаження каналу.

Інформаційна модель SNMP організована у вигляді ієрархічної бази управляючої інформації MIB, яка представляє собою структуровану сукупність об'єктів, що описують параметри та характеристики керованого пристрою. Кожен об'єкт в MIB ідентифікується унікальним об'єктним ідентифікатором OID, який являє собою послідовність чисел, розділених крапками, що визначає шлях від кореня ієрархічного дерева до конкретного об'єкта. Стандартизовані MIB описують загальні параметри мережеских пристроїв, такі як інформація про інтерфейси в MIB-II, таблиці маршрутизації, параметри протоколів TCP та UDP, тоді як виробники обладнання можуть визначати власні розширення MIB для специфічних функцій своїх пристроїв. Для ефективної роботи з SNMP адміністратор повинен знати структуру MIB керованих пристроїв та OID потрібних параметрів, що дозволяє формувати точні запити для отримання необхідної інформації або виконання управляючих операцій.

Протокол SNMP існує в декількох версіях, що відрізняються функціональністю та рівнем безпеки. Перша версія SNMPv1 використовує простий механізм аутентифікації на основі текстових рядків community strings, що передаються відкритим текстом і мають обмежені можливості захисту від несанкціонованого доступу. Друга версія SNMPv2c зберігає механізм community strings, але вносить покращення в ефективність протоколу та розширює набір типів даних. Третя версія SNMPv3 суттєво посилює механізми безпеки шляхом впровадження криптографічної аутентифікації, шифрування повідомлень та контролю доступу на основі моделі користувачів і груп. У корпоративних мережах рекомендується використовувати SNMPv3 для захисту управляючої інформації від перехоплення та модифікації зловмисниками, особливо якщо трафік моніторингу передається через незахищені сегменти мережі.

Таблиця 6.1 наводить порівняльну характеристику основних операцій протоколу SNMP та їх призначення в контексті управління мережевими пристроями.

Таблиця 6.1 - Основні операції протоколу SNMP

Операція	Оп ерация	Напря мок	Призначення	Особливості використання
T	GE	Менеджер → Агент	Отримання значення одного OID	Використовується для запиту конкретного параметра
T-NEXT	GE	Менеджер → Агент	Отримання наступного OID в ієрархії	Дозволяє послідовно обходити дерево MIB
T-BULK	GE	Менеджер → Агент	Отримання множини OID за один запит	Ефективне для читання великих таблиць (SNMPv2/v3)
T	SE	Менеджер → Агент	Встановлення значення OID	Змінює конфігураційні параметри пристрою
AP	TR	Агент → Менеджер	Асинхронне сповіщення про подію	Відправляється агентом без запиту менеджера
FORM	IN	Агент → Менеджер	Сповіщення з підтвердженням доставки	Гарантує отримання повідомлення менеджером (SNMPv2/v3)

Практична реалізація моніторингу на основі SNMP передбачає налаштування на мережевих пристроях параметрів доступу для менеджера, включаючи визначення community strings для SNMPv1/v2c або облікових записів користувачів для SNMPv3, обмеження списку IP-адрес, з яких дозволені запити, та вибір версії протоколу. На стороні менеджера конфігурується перелік пристроїв для моніторингу з вказанням їх IP-адрес та параметрів доступу, визначаються OID для регулярного опитування, встановлюються інтервали між запитами та порогові значення для генерації сповіщень адміністраторів. Сучасні системи моніторингу, такі як Nagios, Zabbix, PRTG або Cisco Prime Infrastructure, автоматизують процес збору даних через SNMP, надають графічну візуалізацію метрик у вигляді графіків та діаграм, зберігають історичні дані для аналізу трендів та генерують оповіщення при виявленні аномалій у роботі мережі.

6.3. Протокол Syslog для централізованого журналювання подій

Протокол Syslog є стандартним механізмом передачі повідомлень про події від мережевих пристроїв та серверів до централізованого сервера журналювання, що забезпечує консолідоване зберігання логів для полегшення аналізу, аудиту та розслідування інцидентів. На відміну від SNMP, який орієнтований на опитування числових метрик, Syslog призначений для асинхронної передачі текстових повідомлень про різноманітні події, що відбуваються на пристрої, включаючи зміни стану інтерфейсів, перезавантаження системи, помилки конфігурації, спроби несанкціонованого доступу та діагностичні повідомлення від мережевих протоколів. Архітектура

Syslog є простою та ефективною: пристрій-джерело генерує повідомлення та відправляє їх на сервер Syslog за допомогою протоколу UDP на стандартний порт 514, хоча сучасні реалізації також підтримують TCP для гарантованої доставки та TLS для захищеної передачі.

Кожне повідомлення Syslog містить кілька ключових компонентів, що дозволяють ідентифікувати джерело події, визначити її важливість та зрозуміти суть проблеми. Повідомлення включає мітку часу генерації події, ім'я або IP-адресу пристрою-відправника, facility - категорію системи або підсистеми, що згенерувала повідомлення, severity - рівень важливості події, та текст повідомлення з детальним описом події. Facility дозволяє класифікувати повідомлення за функціональними підсистемами, наприклад, окремо обробляти події від підсистеми безпеки, від протоколів маршрутизації або від системи управління конфігурацією. Рівень severity визначає критичність події за вісімбальною шкалою: від найвищого рівня Emergency, що сигналізує про повну непрацездатність системи, до найнижчого рівня Debug, що використовується для детальної діагностичної інформації під час розробки та налагодження.

Таблиця 6.2 систематизує рівні важливості повідомлень Syslog та надає характеристику типових ситуацій для кожного рівня.

Таблиця 6.2 - Рівні важливості повідомлень протоколу Syslog

Рівень	Назва	Числове значення	Опис та типові приклади використання
0	Emergency	0	Система повністю непрацездатна (kernel panic, повна відмова обладнання)
1	Alert	1	Необхідні негайні дії (критична помилка RAID, відсутність зв'язку з основним каналом)
2	Critical	2	Критичний стан (перегрів обладнання, відмова резервного джерела живлення)
3	Error	3	Помилки, що впливають на функціональність (інтерфейс перейшов в down, відмова автентифікації)
4	Warning	4	Попередження про потенційні проблеми (високе завантаження CPU, заповнення пам'яті)
5	Notice	5	Нормальні, але значущі події (зміна конфігурації, встановлення BGP-сесії)
6	Informational	6	Інформаційні повідомлення (успішна автентифікація, статистика інтерфейсів)
7	Debug	7	Детальна діагностична інформація для налагодження протоколів та процесів

Налаштування маршрутизатора або комутатора Cisco для відправлення повідомлень на Syslog-сервер включає визначення IP-адреси сервера, вибір мінімального рівня важливості повідомлень для передачі та конфігурацію додаткових параметрів, таких як включення міток часу з високою точністю та додавання порядкових номерів повідомлень. За замовчуванням Cisco IOS відправляє повідомлення рівня Informational та вище, проте адміністратор може налаштувати фільтрацію для передачі на сервер лише повідомлень певного

рівня критичності, щоб уникнути перевантаження сервера надмірною кількістю несуттєвих записів. Важливо синхронізувати час на всіх мережевих пристроях за допомогою протоколу NTP, оскільки точні мітки часу критично важливі для кореляції подій між різними пристроями під час аналізу інцидентів та побудови хронології атак.

Централізований Syslog-сервер здійснює прийом повідомлень від множини мережевих пристроїв, зберігає їх у файлах або базах даних з можливістю категоризації за джерелами, facility та severity, надає засоби пошуку та фільтрації для швидкого виявлення подій певного типу. Сучасні рішення для управління логами, такі як Splunk, Graylog, ELK Stack (Elasticsearch, Logstash, Kibana) або rsyslog, розширюють базову функціональність Syslog додатковими можливостями парсингу повідомлень, виявлення аномалій на основі машинного навчання, кореляції подій між різними джерелами та автоматичного оповіщення адміністраторів при виявленні критичних ситуацій. Інтеграція Syslog з системами інформаційної безпеки дозволяє виявляти спроби вторгнень, відстежувати підозрілі зміни конфігурації та забезпечувати відповідність вимогам регуляторних стандартів щодо зберігання аудиторських журналів протягом визначеного періоду часу.

6.4. Інтеграція моніторингу в операційні процеси адміністрування мережі

Ефективне використання систем централізованого моніторингу вимагає інтеграції механізмів SNMP та Syslog у комплексну операційну модель управління мережевою інфраструктурою, що охоплює не лише технічні аспекти збору та зберігання даних, але й організаційні процедури реагування на інциденти, планування ресурсів та оптимізації продуктивності мережі. Побудова такої системи починається з інвентаризації всього активного мережевого обладнання та визначення критично важливих параметрів для моніторингу на кожному типі пристроїв, з урахуванням їх ролі в мережевій архітектурі та вимог до доступності сервісів. Для граничних маршрутизаторів пріоритетними метриками є стан BGP-сесій з провайдерами, завантаження зовнішніх інтерфейсів та кількість відкинутих пакетів через політики безпеки, для комутаторів ядра мережі важливі показники використання портів агрегації, стабільність протоколів резервування HSRP або VRRP, а для комутаторів рівня доступу критичними є виявлення петель у топології STP та моніторинг автентифікації користувачів через 802.1X.

Проектування архітектури серверів моніторингу має враховувати масштаб мережевої інфраструктури, географічну розподіленість пристроїв та вимоги до надійності системи спостереження. У великих корпоративних мережах доцільно розгортати ієрархічну структуру з регіональними серверами збору даних у кожному датацентрі або великому офісі, які виконують первинну обробку та агрегацію інформації, та центральним сервером для консолідованого зберігання, аналізу та звітності. Така архітектура зменшує навантаження на канали зв'язку між віддаленими локаціями, підвищує

відмовостійкість системи моніторингу через локальне збереження даних навіть при недоступності центрального сервера, та покращує продуктивність обробки запитів завдяки розподіленню навантаження. Резервування серверів моніторингу через кластерні рішення або активно-пасивні конфігурації забезпечує неперервність спостереження навіть у випадку відмови основного сервера, що критично важливо для виконання угод про рівень сервісу SLA та дотримання регуляторних вимог.

Рисунок 6.1 ілюструє типову архітектуру інтегрованої системи моніторингу корпоративної мережі з використанням протоколів SNMP та Syslog. На схемі зображено три основні сегменти: сегмент керованих пристроїв, що включає маршрутизатори, комутатори та точки доступу Wi-Fi з активованими агентами SNMP та клієнтами Syslog; сегмент серверів моніторингу, що містить SNMP-менеджер для збору метрик, Syslog-сервер для централізованого журналювання та сервер NTP для синхронізації часу; сегмент адміністрування з робочими станціями операторів, на яких встановлені консолі управління та засоби візуалізації.

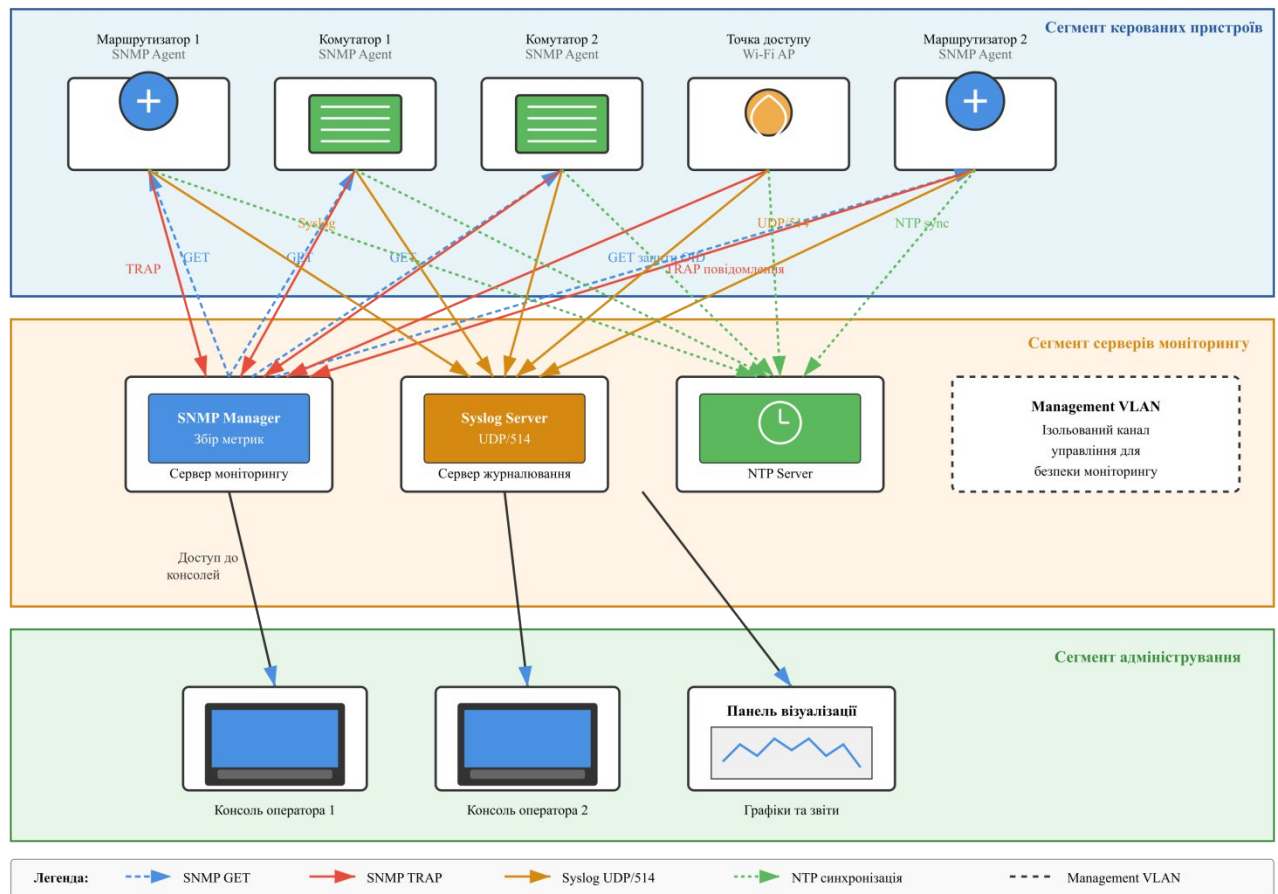


Рисунок 6.1 - Архітектура інтегрованої системи моніторингу корпоративної мережі

Стрілки на схемі показують напрямки комунікації: від менеджера до агентів йдуть GET-запити для опитування OID, від агентів до менеджера надходять TRAP-повідомлення про критичні події, від усіх мережевих

пристроїв до Syslog-сервера передаються журнальні повідомлення через UDP/514, та від усіх компонентів до NTP-сервера відправляються запити синхронізації часу. Додатково на схемі позначено окремий канал управління management VLAN, ізольований від користувацького трафіку для підвищення безпеки та стабільності системи моніторингу.

Налаштування базового моніторингу на маршрутизаторі Cisco IOS включає активацію SNMP-агента з визначенням community string для read-only доступу та обмеженням джерел запитів списком ACL, конфігурацію відправлення SNMP trap до менеджера при критичних подіях, таких як зміна стану інтерфейсу або перезавантаження системи, налаштування передачі Syslog-повідомлень на сервер з визначенням мінімального рівня severity та включенням детальних міток часу, а також синхронізацію системного часу з NTP-сервером для забезпечення точності часових міток в журналах. Рекомендується також активувати локальне буферне журналювання на пристрої для зберігання останніх повідомлень у оперативній пам'яті, що дозволяє адміністратору швидко переглянути недавні події при підключенні до пристрою навіть якщо зв'язок з централізованим сервером тимчасово недоступний. Важливим аспектом конфігурації є балансування між повнотою моніторингу та навантаженням на мережу: надто часте опитування через SNMP або передача повідомлень Syslog з рівнем Debug можуть створювати значний трафік та споживати ресурси процесора на керованих пристроях.

Використання даних моніторингу для проактивного управління мережею передбачає регулярний аналіз трендів завантаження каналів зв'язку для планування розширення пропускної здатності до виникнення перевантажень, відстеження статистики помилок на інтерфейсах для раннього виявлення деградації фізичного середовища передачі або проблем з кабельною системою, моніторинг стабільності сесій динамічної маршрутизації для виявлення нестабільності топології, що може вказувати на проблеми з якістю каналів або некоректну конфігурацію. Кореляція подій між різними пристроями дозволяє швидко локалізувати джерело проблеми: наприклад, одночасне отримання Syslog-повідомлень про недоступність шлюзу за замовчуванням від декількох комутаторів доступу однозначно вказує на проблему з маршрутизатором ядра мережі, а не з окремими комутаторами. Автоматизація реакції на певні типи подій через скрипти або системи оркестрації може суттєво скоротити час відновлення сервісів: система моніторингу може автоматично перемикати трафік на резервний канал при виявленні критичного завантаження основного, перезавантажувати проблемні інтерфейси або тимчасово блокувати порти з підозрілою активністю до втручання адміністратора.

Документування конфігурації моніторингу, створення стандартизованих процедур інтерпретації типових повідомлень та побудова ескалаційних ланцюжків для різних категорій інцидентів є невід'ємною частиною операційної моделі управління мережею. Адміністратори повинні регулярно переглядати налаштування порогових значень для сповіщень, адаптуючи їх до реальних умов експлуатації мережі та уникаючи як надмірної кількості хибних спрацювань, що призводить до ігнорування сповіщень, так і занадто високих

порогів, що можуть пропустити реальні проблеми. Періодичне тестування системи моніторингу через симуляцію відмов обладнання або навантажувальні тести дозволяє переконатися в коректності налаштувань сповіщень та готовності команди адміністрування до реагування на різні сценарії інцидентів, а аналіз збережених історичних даних надає цінну інформацію для вдосконалення архітектури мережі та обґрунтування інвестицій в модернізацію інфраструктури.

Тема 7. Архітектура та експлуатаційні характеристики серверного обладнання в інтегрованій ІТ-інфраструктурі

7.1 Класифікація та архітектурні особливості серверного обладнання

Сучасне серверне обладнання є основою корпоративних інформаційних систем і забезпечує виконання критично важливих обчислювальних завдань у розподілених мережових середовищах. Сервер являє собою спеціалізований комп'ютер, оптимізований для надання сервісів множині клієнтів одночасно, з підвищеними вимогами до надійності, продуктивності та безперервності роботи. На відміну від робочих станцій, серверне обладнання проектується з урахуванням можливості довготривалої експлуатації під високим навантаженням та підтримки механізмів резервування критичних компонентів.

Класифікація серверів за форм-фактором визначає їхні фізичні характеристики та специфіку розгортання в інфраструктурі. Rack-сервери являють собою найпоширенішу категорію обладнання для центрів обробки даних, що монтується в стандартні 19-дюймові стійки та вимірюється в юнітах висоти (1U = 44.45 мм). Tower-сервери за зовнішнім виглядом нагадують настільні комп'ютери великого розміру і використовуються переважно в малих організаціях або філіях, де відсутні серверні приміщення зі спеціалізованим обладнанням. Blade-сервери представляють високощільне рішення, де кілька серверних модулів (блейдів) встановлюються в спільний корпус (chassis), що забезпечує централізоване живлення, охолодження та мережеву зв'язність.

Hyper-converged infrastructure (HCI) являє собою сучасний підхід до організації серверної інфраструктури, де обчислювальні ресурси, системи зберігання даних та мережеві компоненти об'єднуються в єдину програмно-визначену платформу. У такій архітектурі віртуалізація ресурсів здійснюється на рівні кожного вузла, а централізоване управління забезпечується спеціалізованим програмним забезпеченням. Основною перевагою HCI є спрощення розгортання та масштабування інфраструктури за рахунок уніфікованих апаратних блоків та автоматизації управління ресурсами. Така архітектура особливо ефективна для віртуалізованих середовищ, приватних хмар та розподілених філіальних мереж, де потрібна гнучкість та мінімізація операційних витрат на адміністрування.

Вибір форм-фактора серверного обладнання залежить від множини чинників, включаючи масштаб інфраструктури, бюджетні обмеження, вимоги до щільності розміщення та специфіку навантаження. Таблиця 7.1 наведена нижче систематизує основні типи серверів та їхні характеристики для полегшення процесу вибору оптимального рішення.

Таблиця 7.1 – Порівняльна характеристика форм-факторів серверного обладнання

Форм-фактор	Переваги	Недоліки	Типові сценарії
Rack	Стандартизація, ефективне використання простору ЦОД, централізоване живлення	Вимагає серверної стійки, обмежена розширюваність у межах корпусу	Центри обробки даних, великі корпоративні мережі
Tower	Не потребує спеціального обладнання, легке обслуговування, низька вартість входу	Займає багато площі, обмежена масштабованість, підвищений рівень шуму	Малі офіси, філії, спеціалізовані робочі станції
Blade	Максимальна щільність, спільна інфраструктура, централізоване управління	Висока вартість шасі, прив'язка до вендора, складність міграції	Великі віртуалізовані середовища, хмарні провайдери
HCI	Простота масштабування, програмно-визначена архітектура, швидке розгортання	Вища вартість на етапі входу, обмеження у виборі компонентів	Приватні хмари, віддалені офіси, швидко масштабовані інфраструктури

7.2 Компонентна архітектура та основні підсистеми сервера

Архітектура сучасного сервера являє собою складну інтеграцію апаратних компонентів, кожен з яких відповідає за виконання специфічних функцій у загальній системі. Центральний процесор (CPU) визначає обчислювальну потужність сервера та здатність до паралельної обробки завдань. Серверні процесори відрізняються від настільних наявністю підтримки багатопроесорних конфігурацій, розширеними кешами третього рівня, підвищеною кількістю ядер та потоків, а також спеціалізованими інструкціями для віртуалізації та криптографічних операцій. Архітектури на базі процесорів Intel Xeon Scalable та AMD EPYC забезпечують до 64 ядер на один сокет, що дозволяє ефективно розподіляти навантаження між віртуальними машинами та контейнерами.

Підсистема оперативної пам'яті в серверному обладнанні використовує спеціалізовані модулі з кодом корекції помилок (ECC), що виявляють та виправляють однобітові помилки і детектують багатобітові збої, критично важливі для підтримки цілісності даних у довготривалих обчислювальних процесах. Сучасні сервери підтримують архітектуру пам'яті DDR5 з пропускнуою здатністю до 6400 MT/s та можливістю встановлення до кількох терабайт ємності в межах одного сервера. Технології NUMA (Non-Uniform Memory Access) та інтерлівінг каналів пам'яті дозволяють оптимізувати доступ

процесорів до модулів RAM, мінімізуючи затримки та максимізуючи сумарну пропускну здатність шини.

Системи зберігання даних у серверному середовищі представлені жорсткими дисками (HDD) та твердотільними накопичувачами (SSD) різних інтерфейсів підключення. HDD залишаються актуальними для зберігання великих обсягів холодних даних завдяки вищій питомій ємності та нижчій вартості за гігабайт, тоді як SSD на базі флеш-пам'яті NAND забезпечують істотно вищі показники швидкості читання-запису та IOPS (операцій введення-виведення за секунду). Інтерфейси NVMe (Non-Volatile Memory Express) через шину PCIe дозволяють SSD досягати пропускну здатності понад 7000 МБ/с на читання та мінімальних затримок на рівні мікросекунд. RAID-масиви (Redundant Array of Independent Disks) об'єднують множину фізичних дисків у логічні томи з різними рівнями резервування та продуктивності, забезпечуючи захист від відмов окремих накопичувачів. Контролери RAID можуть бути інтегрованими на материнській платі або представлені окремими адаптерами з власним кешем та резервним живленням для захисту даних у буфері при аварійних ситуаціях.

Мережеві інтерфейси серверів забезпечують з'єднання з локальною мережею та можуть включати кілька портів Ethernet різних швидкостей від 1 Gbps до 100 Gbps. Багатопортіві мережеві карти дозволяють реалізувати агрегацію каналів (LACP) для підвищення пропускну здатності та резервування зв'язку. Технології SR-IOV (Single Root I/O Virtualization) дають можливість віртуалізувати мережеві адаптери, забезпечуючи прямий доступ віртуальних машин до фізичного обладнання без втрат продуктивності через гіпервізор.

Baseboard Management Controller (BMC) являє собою спеціалізований мікроконтролер, вбудований у серверну материнську плату, що забезпечує віддалене управління та моніторинг апаратного стану сервера незалежно від основної операційної системи. BMC функціонує на базі стандарту IPMI (Intelligent Platform Management Interface) та надає можливості віддаленого включення/вимкнення живлення, перезавантаження, доступу до консолі через мережу, моніторингу температур, напруг, обертів вентиляторів та інших сенсорів. Сучасні реалізації BMC від різних виробників (iDRAC у Dell, iLO у HPE, IMM у Lenovo) розширюють базовий функціонал IPMI додатковими можливостями віртуальних медіа, агрегації логів апаратних подій та інтеграції з системами централізованого управління інфраструктурою. Наявність BMC критично важлива для безперебійної експлуатації серверів у віддалених локаціях та центрах обробки даних, де фізичний доступ до обладнання обмежений.

Структурна схема типового rack-сервера з ключовими компонентами представлена на рисунку 7.1. На схемі відображено взаємозв'язок основних підсистем: двопроцесорна конфігурація з роздільними банками пам'яті DDR5 ECC, підключеними за архітектурою NUMA; масив дисків з контролером RAID; мережеві адаптери з підтримкою SR-IOV; BMC з виділеним мережевим портом для управління; блоки живлення з резервуванням за схемою N+1. Шини

PCIe забезпечують високошвидкісне з'єднання процесорів з контролерами дисків, мережевими адаптерами та іншими периферійними пристроями розширення.

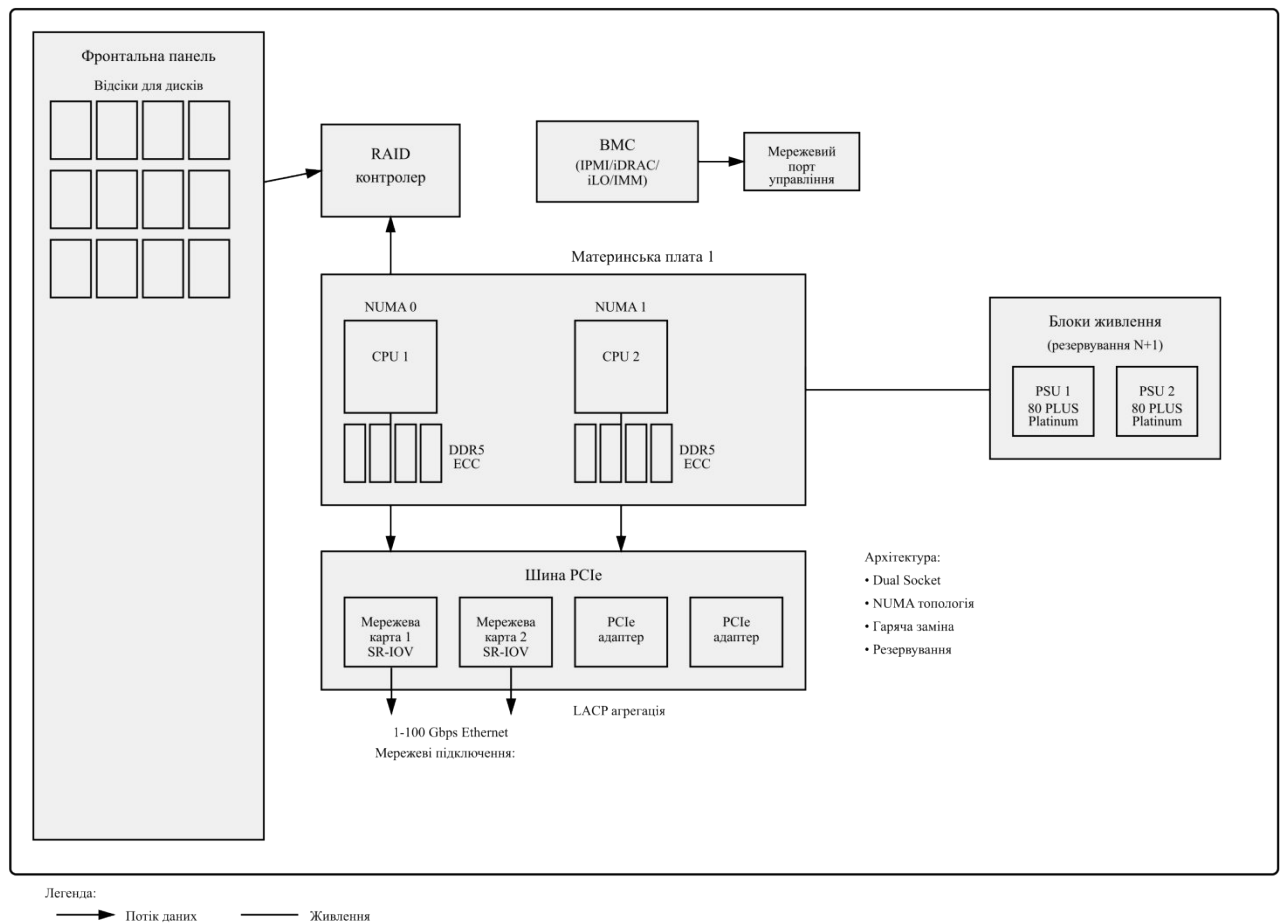


Рисунок 7.1 – Структурна схема компонентів rack-сервера

На схемі зображено прямокутний корпус rack-сервера з фронтальною панеллю, на якій розміщені відсіки для гарячої заміни дисків (показані як вертикальні сегменти). Усередині корпусу знаходяться дві материнські плати з двома процесорними сокетом та модулями пам'яті, розташованими навколо кожного процесора. У центральній частині розміщено RAID-контролер, від якого лінії з'єднання йдуть до дискових відсіків. У правій частині схеми показано два блоки живлення з позначкою резервування. У нижній частині материнської плати розташовано чотири слоти PCIe з підключеними мережевими картами. Окремо виділено BMC-контролер з лінією до виділеного мережевого порту управління. Стрілками позначено напрямки передачі даних між компонентами: від процесорів до пам'яті, від процесорів через PCIe до периферії, від RAID-контролера до дисків.

7.3 Експлуатаційні характеристики та метрики надійності серверного обладнання

Оцінка придатності серверного обладнання для конкретних завдань базується на аналізі комплексу експлуатаційних характеристик, що визначають продуктивність, надійність, енергоефективність та економічну доцільність інвестицій. Продуктивність сервера вимірюється набором показників, специфічних для типу навантаження: кількість операцій з плаваючою комою (FLOPS) для наукових обчислень, транзакцій за секунду для баз даних (TPS), запитів за секунду для веб-серверів (RPS). Стандартизовані бенчмарки, такі як SPECint для цілочисельних обчислень, SPECfp для операцій з плаваючою комою, TPC-C для OLTP-систем, дозволяють об'єктивно порівнювати різні апаратні платформи.

Надійність серверного обладнання кількісно характеризується показником MTBF (Mean Time Between Failures) – середнім часом між відмовами, що вимірюється в годинах та відображає очікувану тривалість безвідмовної роботи компонента або системи в цілому. Типові значення MTBF для серверів корпоративного класу складають від 500 000 до 1 000 000 годин, що відповідає статистичній ймовірності відмови менше 0.2% на рік експлуатації. Комплементарним показником є MTTR (Mean Time To Repair) – середній час відновлення, який включає діагностику несправності та заміну компонента. Співвідношення MTBF до суми MTBF та MTTR визначає коефіцієнт готовності системи (availability), критично важливий для служб, що вимагають високої доступності.

Механізми підвищення надійності в серверному обладнанні включають резервування критичних компонентів за схемами N+1 або 2N, де N – мінімально необхідна кількість елементів для функціонування системи. Блоки живлення з гарячою заміною та резервуванням забезпечують безперервність роботи при виході з ладу одного з модулів. Вентилятори з підвищеною швидкістю обертання компенсують відмову сусідніх блоків охолодження. ЕСС-пам'ять виявляє та виправляє помилки без зупинки обчислювального процесу. RAID-масиви забезпечують збереження даних при відмові одного або кількох дисків залежно від обраного рівня. Таблиця 7.2 нижче систематизує основні рівні RAID та їхні характеристики.

Таблиця 7.2 – Порівняння рівнів RAID для серверних систем зберігання

Рівень RAID	Принцип організації	Ємність	Відмовостійкість	Типове застосування
RAID 0	Striping без резервування	100% (n дисків)	Відсутня	Тимчасові дані, кеші
RAID 1	Дзеркалювання даних	50% (n/2)	1 диск	ОС, критичні БД
RAID 5	Striping з розподіленою парністю	$(n-1)/n \times 100\%$	1 диск	Файлові сервери

D 6	RAI Striping з подвійною парністю	$(n-2)/n \times 100\%$	2 диски	Архів и, великі масиви
D 10	RAI Комбінація striping і mirroring	50% (n/2)	1 диск на пару	OLTP БД, високі IOPS

Енергоефективність серверного обладнання набуває критичного значення в умовах експлуатації великих дата-центрів, де витрати на електроенергію та охолодження можуть становити до 40% від сукупної вартості володіння інфраструктурою протягом п'ятирічного циклу. Метрика PUE (Power Usage Effectiveness) визначається як відношення загального енергоспоживання дата-центру до енергоспоживання безпосередньо ІТ-обладнання, при цьому значення PUE = 1.0 означає ідеальну ефективність без витрат на додаткову інфраструктуру. Сучасні дата-центри досягають показників PUE на рівні 1.2-1.3 завдяки оптимізації систем охолодження, використанню вільного повітряного охолодження (free cooling) та впровадженню гарячих/холодних коридорів для керування повітряними потоками.

Серверні процесори та компоненти підтримують динамічне регулювання частоти та напруги (DVFS – Dynamic Voltage and Frequency Scaling), що дозволяє знижувати енергоспоживання у періоди низького навантаження без впливу на здатність швидко відповідати на зростання запитів. Сертифікація блоків живлення за стандартом 80 PLUS гарантує ефективність перетворення змінного струму в постійний не менше 80% при різних рівнях навантаження, при цьому найвищі рівні Platinum та Titanium досягають ефективності 94-96%. Можливості масштабування серверної інфраструктури визначаються як вертикальним нарощуванням ресурсів окремих серверів (scale-up), так і горизонтальним додаванням нових вузлів до кластера (scale-out), при цьому вибір стратегії залежить від архітектури додатків та характеру навантаження.

7.4 Роль серверного обладнання в інтегрованій ІТ-інфраструктурі

Серверне обладнання в сучасних корпоративних мережах виступає фундаментальною платформою для розгортання критично важливих мережевих служб, що забезпечують функціонування інформаційних систем підприємства. Служба DHCP (Dynamic Host Configuration Protocol) централізовано управляє автоматичним призначенням IP-адрес клієнтським пристроям, елімінуючи необхідність ручного конфігурування та мінімізуючи ризики конфліктів адрес у мережі. DNS (Domain Name System) забезпечує перетворення символічних імен хостів у відповідні IP-адреси, формуючи основу для зручної навігації користувачів у корпоративних мережевих ресурсах та інтернеті.

RADIUS-сервери (Remote Authentication Dial-In User Service) реалізують централізовану аутентифікацію, авторизацію та облік підключень користувачів до мережевої інфраструктури, включаючи доступ через VPN, бездротові точки доступу та адміністративні сесії на мережевому обладнанні. Інтеграція RADIUS

з корпоративними каталогами Active Directory або LDAP дозволяє уніфікувати управління обліковими записами та політиками доступу, знижуючи адміністративні витрати та підвищуючи рівень безпеки. Syslog-сервери агрегують журнали подій від розподіленого мережевого обладнання, маршрутизаторів, комутаторів, міжмережевих екранів, формуючи централізоване сховище для аналізу інцидентів, виявлення аномалій та підтримки вимог регуляторних стандартів щодо збереження аудит-логів.

Служба NTP (Network Time Protocol) синхронізує годинники на всіх пристроях мережі з точними джерелами часу, що є критичним для коректного функціонування розподілених систем, протоколів безпеки (Kerberos), корелювання подій між різними системами та забезпечення юридичної значущості електронних документів з часовими мітками. Серверні платформи для систем моніторингу на базі SNMP (Simple Network Management Protocol) збирають метрики продуктивності та доступності з мережевої інфраструктури, надаючи адміністраторам інструменти для проактивного виявлення проблем, планування ємності та оптимізації використання ресурсів.

Взаємодія серверів з мережевою інфраструктурою потребує ретельного планування топології та налаштування механізмів відмовостійкості. VLAN-сегментація дозволяє логічно ізолювати різні типи трафіку, розділяючи управлінський доступ до BMC, продуктивний трафік додатків, резервне копіювання та міграцію віртуальних машин. Агрегація мережевих інтерфейсів за протоколом LACP забезпечує як підвищення пропускної здатності, так і відмовостійкість каналів зв'язку між серверами та комутаторами доступу. Протоколи резервування на третьому рівні моделі OSI, такі як VRRP або HSRP, гарантують безперервну доступність мережевих служб при відмові окремих серверів через автоматичне перемикання віртуальної IP-адреси на резервний вузол.

Інтеграція серверного обладнання з системами управління інфраструктурою передбачає використання стандартизованих API та протоколів для автоматизації розгортання, конфігурування та моніторингу. Інструменти Infrastructure as Code (IaC) на базі Ansible, Terraform, Puppet дозволяють декларативно описувати бажаний стан серверів та мережевого обладнання, забезпечуючи відтворюваність конфігурацій та швидке відновлення після збоїв. SNMP та Redfish API надають програмний доступ до апаратних параметрів серверів, включаючи сенсори температур, стан RAID-масивів, завантаження процесорів та пам'яті.

Механізми безпеки на рівні серверної інфраструктури включають захист фізичного доступу до обладнання через контроль доступу в серверні приміщення, шифрування дисків для захисту даних у стані спокою (data at rest), мережеву сегментацію для обмеження латерального руху потенційних зловмисників, регулярне оновлення мікропрограм (firmware) BMC та BIOS для усунення вразливостей. Технології Trusted Platform Module (TPM) забезпечують апаратне зберігання криптографічних ключів та підтримку механізмів довіреного завантаження (Secure Boot), що запобігає виконанню неавторизованого коду на етапі ініціалізації системи. Ізоляція мережі

управління ВМС від продуктивних мереж є критично важливою практикою, оскільки компрометація контролера управління надає зловмиснику повний контроль над апаратною платформою незалежно від налаштувань операційної системи.

Планування життєвого циклу серверного обладнання охоплює етапи проектування інфраструктури, закупівлі обладнання, розгортання та конфігурування, експлуатації з регулярним обслуговуванням, та виведення з експлуатації з безпечним знищенням даних. Типовий термін амортизації серверів складає 3-5 років, після чого обладнання може бути замінено на більш сучасні платформи з вищою енергоефективністю та продуктивністю. Підтримка виробником у формі гарантійного обслуговування та доступності запасних частин є критичним фактором при виборі серверних платформ для корпоративних середовищ.

Таким чином, серверне обладнання становить невід'ємну складову інтегрованої ІТ-інфраструктури, забезпечуючи обчислювальні ресурси для мережеслужб, систем управління, моніторингу та безпеки. Розуміння архітектурних особливостей, експлуатаційних характеристик та принципів інтеграції серверів з мережевою інфраструктурою є фундаментальною компетенцією для фахівців у галузі комп'ютерних мереж, необхідною для проектування надійних, продуктивних та масштабованих інформаційних систем підприємства.

РЕКОМЕНДОВАНА ЛІТЕРАТУРА

1. Жураковський Б. Ю. Комп'ютерні мережі. Частина 1. [навчальний посібник] / Б. Ю. Жураковський, І.О. Зенів. – Київ : КПІ ім. Ігоря Сікорського, 2020. – 336 с.
2. Б. Ю. Жураковський. Комп'ютерні мережі. Частина 2. [навчальний посібник] / Б. Ю. Жураковський, І.О. Зенів. – Електронні текстові дані. – Київ : КПІ ім. Ігоря Сікорського, 2020. – 372 с.
3. Азаров О. Д. Комп'ютерні мережі : підручник / О. Д. Азаров, С. М. Захарченко, О. В. Кадук та ін. – Вінниця : ВНТУ, 2020. – 378 с.
4. Карпенко М. Ю. Конспект лекцій з курсу «Комп'ютерні мережі» / М. Ю. Карпенко, Н. В. Макогон; Харків. нац. ун-т міськ. госп-ва ім. О. М. Бекетова. – Харків : ХНУМГ ім. О. М. Бекетова, 2019. – 99 с.
5. Блозва А.І. Комп'ютерні мережі [навчальний посібник] / А.І.Блозва, Ю.В.Матус, В.В.Смолій, Б.С.Гусев, Д.Ю.Касаткін, Т.Ю.Осипова, Я.А.Савицька // - К.: Компрінт, 2017.- 821с.
6. Тарбаєв С.І. Проектування інфокомунікаційних мереж. Навчальний посібник. – Київ: ННІТІ ДУТ, 2019. – 186 ст.
7. Задерейко О. В. Комп'ютерні мережі : навчальний посібник [Електронне видання] / О. В. Задерейко, Н. І. Логінова, А. А. Толокнов. – Одеса : Фенікс, 2022. – 249 с.
8. Comer D. E. Computer Networks and Internets. Global Edition : textbook / D. E. Comer. – 6th ed. – Boston : Pearson Education, 2016. – 672 p.
9. Kurose J. F. Computer Networking: A Top-Down Approach : textbook / J. F. Kurose, K. W. Ross. – 8th ed. – Boston : Pearson Education, 2021. – 864 p.
10. Peterson L. L. Computer Networks: A Systems Approach : textbook / L. L. Peterson, B. S. Davie. – 6th ed. – Amsterdam : Morgan Kaufmann, 2021. – 896 p.
11. Goralski W. The Illustrated Network: How TCP/IP Works in a Modern Network : textbook / W. Goralski. – 2nd ed. – Amsterdam : Morgan Kaufmann, 2017. – 936 p.
12. Odom W. CCNA 200-301 Official Cert Guide. Volume 1 / W. Odom. – 2nd ed. – Indianapolis : Cisco Press, 2024. – 1024 p.
13. Odom W. CCNA 200-301 Official Cert Guide. Volume 2 / W. Odom. – 2nd ed. – Indianapolis : Cisco Press, 2024. – 992 p.
14. Sequeira A. J. CCNA 200-301 Hands-on Mastery with Packet Tracer : practical guide / A. J. Sequeira, R. Wong. – 2nd ed. – Indianapolis : Cisco Press, 2024. – 704 p.
15. Sanders C. Practical Packet Analysis: Using Wireshark to Solve Real-World Network Problems : textbook / C. Sanders. – 3rd ed. – San Francisco : No Starch Press, 2021. – 432 p.
16. Limoncelli T. A. The Practice of Network and System Administration : textbook / T. A. Limoncelli, C. J. Hogan, S. R. Chalup. – 2nd ed. – Boston : Addison-Wesley, 2020. – 1040 p.
17. Hucaby D. CCNA 200-301 Portable Command Guide : reference guide / D. Hucaby, W. Odom. – 2nd ed. – Indianapolis : Cisco Press, 2024. – 352 p.