

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ
НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ ТЕХНОЛОГІЙ
КАФЕДРА КОМП'ЮТЕРНОЇ ІНЖЕНЕРІЇ**

КВАЛІФІКАЦІЙНА РОБОТА
на тему: «КОМБІНАЦІЇ ТОПОЛОГІЙ ДЛЯ НАДІЙНОСТІ
ФУНКЦІОНУВАННЯ КОМП'ЮТЕРНИХ МЕРЕЖ»

на здобуття освітнього ступеня магістра
зі спеціальності 123 Комп'ютерна інженерія
освітньо-професійної програми Комп'ютерні системи та мережі

*Кваліфікаційна робота містить результати власних досліджень.
Використання ідей, результатів і текстів інших авторів мають посилання на
відповідне джерело*

(підпис)

Антон ГУБА
Ім'я, ПРІЗВИЩЕ здобувача

Виконав: здобувач вищої освіти групи
КСДМ-62 Антон ГУБА

Ім'я, ПРІЗВИЩЕ

Керівник: Анастасія ВСЧЕРКОВСЬКА

Ім'я, ПРІЗВИЩЕ
к.т.н., доцент

Рецензент: _____
Ім'я, ПРІЗВИЩЕ
науковий ступінь,
вчене звання

**ДЕРЖАВНИЙ УНІВЕРСИТЕТ
ІНФОРМАЦІЙНО-КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ**

Навчально-науковий інститут інформаційних технологій

Кафедра Комп'ютерна інженерія

Ступінь вищої освіти Магістр

Спеціальність 123 Комп'ютерна інженерія

Освітньо-професійна програма «Комп'ютерні системи та мережі»

ЗАТВЕРДЖУЮ

Завідувач кафедри

Комп'ютерної інженерії

_____ Наталія ЛАЩЕВСЬКА

«_____» _____ 2024 р.

**ЗАВДАННЯ
НА КВАЛІФІКАЦІЙНУ РОБОТУ**

Губі Антону Ігоровичу

1. Тема кваліфікаційної роботи: : «Комбінації топологій для надійності функціонування комп'ютерних мереж»

керівник кваліфікаційної роботи Вечерковська А.С., к. т. н.

затверджені наказом Державного університету інформаційно-комунікаційних технологій від «15» жовтня 2024 р. № 320.

2. Строк подання кваліфікаційної роботи «20» грудня 2024 р.

3. Вихідні дані до кваліфікаційної роботи: науково-технічна література, посібники компанії Cisco щодо роботи і налаштування їх обладнання, документація команд для ОС Debian.

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

1. Аналіз методів агрегація і резервування каналів та топологій

2. Розробка методу тестування впливу топологій на протоколи LACP, GLBP, MSTP

3. Розробка топологій для кожної топології та проведення над них тестів

4. Аналіз отриманих результатів

5. Перелік ілюстративного матеріалу: *презентація*

1. Мета, об'єкт, предмет дослідження

2. Актуальність роботи

3. Протоколи LACP, MSTP, GLBP
4. Створенні та налаштуванні тестові топології
5. Проведені експерименти
6. Результати тестування протоколу LACP
7. Результати тестування протоколу MSTP
8. Результати тестування протоколу GLBP
9. Аналіз результатів
10. Висновки

6. Дата видачі завдання «1» вересня 2024 р.

КАЛЕНДАРНИЙ ПЛАН

№ з/п	Назва етапів бакалаврської роботи	Строк виконання етапів роботи	Примітка
1	Збір даних	01.09.2024 – 10.09.2024	Виконано
2	Обробка матеріалу	11.09.2024 – 15.09.2024	Виконано
3	Розробка теоретичної частини	15.09.2024 – 25.09.2024	Виконано
4	Розробка методології тестування	26.09.2024 – 30.09.2024	Виконано
5	Побудова топологій	01.10.2024 – 15.10.2024	Виконано
6	Виконання тестування топологій	15.10.2024 – 29.10.2024	Виконано
7	Аналіз отриманих даних	30.10.2024 – 31.10.2024	Виконано
9	Створення презентації	11.11.2024 – 08.12.2024	Виконано

Студент _____ Губа А.І
(підпис) (прізвище та ініціали)

Керівник роботи _____ Вечерковська А.С.,
(підпис) (прізвище та ініціали)

РЕФЕРАТ

Текстова частина кваліфікаційної роботи на здобуття освітнього ступеня магістра: 74 стор., 6 табл., 38 рис., 13 джерел.

Мета роботи – дослідити вплив мережевих топологій: зірка, кільце, дерево, сітка, точка-точка на роботу протоколів: LACP, MSTP та GLBP. А саме на надійність роботи мережі, ключові метрики якої доступність, затримка, втрати пакетів та час відновлення після аварій. Параметри параметри оцінюють надійність роботу мережі через особливість роботи обраних протоколів, що розроблені для швидкого відновлення після аварійних ситуацій для забезпечення надійності функціонування комп'ютерних мереж.

Об'єкт дослідження – мережеві топології та методи та засоби резервування та агрегації каналів у комп'ютерних мережах.

Предмет дослідження – Вплив різних топологій мереж на ефективність роботи протоколів LACP, MSTP та GLBP.

У роботі використано такі методи, як моделювання, експеримент вимірювання, аналіз.

Проведено аналіз сучасних методів агрегації каналів , усунення петель у топологіях та резервування маршрутизаторів із балансуванням навантаження .

Розроблено тестові мережі з різними топологіями, що враховують особливості роботи протоколів резервування та агрегації каналів.

Проведено експерименти для визначення впливу топологій на ключові мережеві метрики, а також оцінки стійкості мережі до відмов вузлів та каналів.

Проаналізовано отримані результати та наведені висновки.

КЛЮЧОВІ СЛОВА: Комп'ютерні мережі, надійність, топологія, резервування каналів, агрегація каналів, протокол LACP, протокол MSTP, протокол GLBP, емуляція, GNS3.

ABSTRACT

The text part of the qualification work for the degree of master's degree: 74 pages, 6 tables, 38 figures, 13 sources.

The purpose of the work is to study the impact of network topologies: star, ring, tree, mesh, point-to-point on the operation of protocols: LACP, MSTP and GLBP. Namely, on network reliability, the key metrics of which are availability, delay, packet loss, and disaster recovery time. The parameters evaluate the reliability of the network due to the peculiarities of the selected protocols, which are designed for rapid recovery from emergencies to ensure the reliability of computer networks.

Object of research - network topologies and methods and means of redundancy and channel aggregation in computer networks.

The subject of the study is the influence of different network topologies on the efficiency of LACP, MSTP and GLBP protocols.

Methods used in the study include modelling, measurement experiment, and analysis.

The analysis of modern methods of channel aggregation, elimination of loops in topologies and redundant routers with load balancing is carried out.

Test networks with different topologies that take into account the peculiarities of the channel redundancy and aggregation protocols are developed.

Experiments were conducted to determine the impact of topologies on key network metrics, as well as to assess the network's resilience to node and link failures.

The results are analysed and conclusions are drawn.

KEYWORDS: Computer networks, reliability, topology, channel redundancy, channel aggregation, LACP protocol, MSTP protocol, GLBP protocol, emulation, GNS3.

ЗМІСТ

ВСТУП.....	9
1 Огляд теоретичних аспектів Технологій і топології для підвищення надійності мережі	11
1.1 Агрегація каналів	11
1.1.1 Вступ до агрегації	11
1.1.2 Агрегування каналів на обладнанні Cisco	11
1.1.3 Балансування трафіку	14
1.1.4 Агрегація і надійність	17
1.2 STP	18
1.3 RSTP	21
1.3.1 Введення	21
1.3.2 Переваги та зміни порівняно з STP	21
1.4 MSTP	28
1.4.1 Введення	28
1.4.2 Функціонування	29
1.4.3 Особливості налаштування	32
1.4.4 Взаємодія з зовнішнім світом	34
1.4.5 Висновки	36
1.5 Методи резервування на мережевому рівні.	38
1.5.1 FHRP	38
1.5.2 HRSP	39
1.5.3 GLBP	41
1.6 Мережеві топології	44
2 Дослідження комбінацій топологій і технологій для підвищення надійності мережі	51
2.1 Методологія дослідження	51
2.1 Дослідження протоколу LACP	55
2.2 Дослідження протоколу MSTP	60
2.3 Дослідження протоколу GLBP	68
3 Аналіз результатів дослідження	70
ВИСНОВКИ.....	73
ПЕРЕЛІК ПОСИЛАНЬ	75
Додаток А Демонстраційний матеріал. Презентація	75

ВСТУП

З огляду на стрімкий розвиток інформаційних технологій та їх інтеграцію у більшість сфер нашого життя, забезпечення стабільної роботи комп'ютерних мереж стає одним із ключових завдань сучасного світу. Перебої у функціонуванні мережі можуть спричинити значні труднощі: від зупинки бізнес-процесів у корпоративному середовищі до втрати доступу до важливих ресурсів для звичайних користувачів. У сучасних реаліях особливої уваги потребує надійність мереж, яка напряду залежить від правильного вибору технологій резервування та агрегації каналів.

Проблеми забезпечення надійності мереж, такі як автоматичне переключення трафіку у разі відмови каналів чи вузлів, та балансування навантаження між декількома каналами, є критично важливими для корпоративних мереж та інфраструктури провайдерів. Саме ці аспекти враховуються при проектуванні мереж для забезпечення максимальної доступності та мінімізації часу простою у разі аварійних ситуацій.

Об'єктом дослідження в роботі виступає комбінування мережевих топологій та а саме технології резервування та агрегації каналів. Предметом дослідження є вплив різних топологій мережі на роботу технологій резервування та агрегації каналів, таких як LACP, MSTP та GLBP.

Метою роботи є визначення впливу мережевих топологій (зірка, кільце, дерево, сітка, точка-точка) на надійність роботи мережевих протоколів шляхом проведення симуляцій у середовищі GNS3 та оцінки ключових мережевих метрик: доступності, втрат пакетів, затримки та часу відновлення після відмови.

У дослідженні використані такі методи, як аналіз, моделювання та експеримент. Експериментальне тестування дозволить перевірити роботу протоколів у різних топологіях, оцінити їх поведінку у разі відмови вузлів або каналів, а також підтвердити незалежність технологій від типу мережевої топології.

Практична значущість роботи полягає у розробці рекомендацій для інженерів та проектувальників мереж щодо вибору протоколів та топологій при побудові

мереж із високими показниками надійності. Отримані результати також сприятимуть кращому розумінню поведінки сучасних протоколів у реальних умовах.

1 ОГЛЯД ТЕОРЕТИЧНИХ АСПЕКТІВ ТЕХНОЛОГІЙ І ТОПОЛОГІЯ ДЛЯ ПІДВИЩЕННЯ НАДІЙНОСТІ МЕРЕЖІ

1.1 Агрегація каналів

1.1.1 Вступ до агрегації

Агрегування каналів – технологія за допомогою якої можна декілька фізичних каналів об'єднати в один логічний. Використовується для об'єднання каналів між мережевим обладнанням декількома способами. Об'єднані канали називаються групою агрегації каналів.

Таке налаштування використовується для об'єднання паралельних каналів між двома пристроями. Це можуть бути пари комутатор-комутатор, комутатор-маршрутизатор, та комутатор – кінцевий пристрій. Приклад можна побачити на рис. 1.1.

Об'єднання каналів збільшує пропускну здатність каналу завдяки різним способам балансування трафіку, а також забезпечує резервування зв'язку між пристроями, котре не перерветься до того моменту як буде вийде з ладу останній канал в групі агрегації каналів.



Рисунок 1.1 – Приклад організації агрегованого каналу

1.1.2 Агрегування каналів на обладнанні Cisco

Існує два способи реалізації агрегації в мережевому обладнанні сімейства Cisco: статичний та динамічний. Динамічний, зокрема, реалізований на відкритому

протоколі LACP, котрий описується в документах стандарту IEEE 802.3ad та 802.1aq, та пропрієтарному рішенню компанії Cisco – PAgP.

Протоколи динамічної агрегації дозволяють мережевим пристроям автоматично домовлятися про агрегацію каналів, надсилаючи сервісні пакети безпосередньо підключеним пристроям, що реалізують той самий протокол агрегації каналів. Це дозволяє спростити процес налаштування агрегації каналів.

Ще одна перевага динамічної агрегації, що при виникненні аварійної ситуації, котра відмінна від розриву з'єднання, можуть виникнути проблеми з вилученням каналу з активної групи агрегації і пакети котрі балансуються на цей канал будуть втрачатися. Прикладом є ситуація при статичному налаштуванні, коли фізично порти працюють нормально, а на логічному рівня лінк між пристроями завис, трафік по такому з'єднанню проходити не буде, що призведе до появи проблем у функціонуванні мережі. Динамічне агрегування може також виявляти інші несправності, такі як несправності на рівні каналу та каналного з'єднання, котрі в разі їх виявлення перемикання виконується автоматично.

Потреба налаштовувати конфігурацію з на обох пристроях, при статичному налаштуванні агрегації каналів збільшує ймовірність появи помилок, та неможливе для використання, коли певні сегменти мережі з'єднані лише одним каналом, і при налаштуванні без динамічної агрегації не обійтися. При налаштуванні статичного агрегування потребує більшої кваліфікації до інженера котрий буде виконувати налаштування, через потребу звертати увагу на більшу кількість деталей.

Багато вендорів рекомендують використовувати саме статичне агрегування, що забезпечить відсутність необхідності генерації мережевим обладнанням службові повідомлення для узгодження з'єднання. Відповідно така конфігурація дозволить уникнути додаткові затримки при виникненні проблем з каналами, але як вже згадувалося, воно має свої недоліки.

LACP підтримує до восьми активних з'єднань, та вісім резервних у режимі – standby. З'єднання в цьому режимі стане активним, якщо одне з поточних активних

з'єднань вийде з ладу. Налаштування агрегування каналів на маршрутизаторі має свої особливості :

- Підтримується тільки статичне агрегування, без використання протоколів;
- Можна створити тільки 2 агрегованих інтерфейси;
- Максимально можна налаштувати 4 інтерфейси в групі агрегації каналів;
- Метод балансування використовує IP-адреси відправника й одержувача, увімкнений за замовчуванням і не може бути змінений;
- Агрегувати можна тільки ті інтерфейси, які знаходяться на модулях однакового типу.

Протоколи динамічної агрегації LACP та PAgP мають мало відмінностей. Проте вони присутні, різняться процес конфігурації та механізм агрегування. Через відсутності підтримки інтерфейсів, котрі розміщені на різних фізичних комутаторах протокол PAgP не вийде використовувати для крос-стекової групи агрегації каналів. Також він не є відритим протоколом і може використовуватись лише на обладнанні компанії Cisco.

Протокол LACP належить до відкритих стандартів, що забезпечує його широку підтримку різними виробниками обладнання. Вибір між протоколами здійснюється залежно від потреб користувача. LACP автоматично узгоджує канали, що формують групи агрегації, і у випадку відмови одного з підключень генерує відповідні записи в логах.

Для налаштування агрегації необхідно, щоб усі інтерфейси, що входять до групи з обох боків, мали однакові параметри: швидкість передачі даних, режим дуплексу, однаковий native VLAN, однаковий діапазон VLAN, статус trunk для порту та тип інтерфейсу.

Налаштування агрегації на маршрутизаторах відрізняється від налаштувань на комутаторах. На маршрутизаторах підтримується виключно статична агрегація. Можна створити лише два агрегованих канали, до кожного з яких можна

підключити максимум 4 інтерфейси. Для балансування використовується виключно хешування на основі IP-адрес джерела та призначення.

1.1.3 Балансування трафіку

Протокол LACP належить до відкритих стандартів, що забезпечує його широку підтримку різними виробниками обладнання. Вибір між протоколами здійснюється залежно від потреб користувача. LACP автоматично узгоджує канали, що формують групи агрегації, і у випадку відмови одного з підключень генерує відповідні записи в логах.

Для налаштування агрегації необхідно, щоб усі інтерфейси, що входять до групи з обох боків, мали однакові параметри: швидкість передачі даних, режим дуплексу, однаковий native VLAN, однаковий діапазон VLAN, статус trunk для порту та тип інтерфейсу.

Налаштування агрегації на маршрутизаторах відрізняється від налаштувань на комутаторах. На маршрутизаторах підтримується виключно статична агрегація. Можна створити лише два агрегованих канали, до кожного з яких можна підключити максимум 4 інтерфейси. Для балансування використовується виключно хешування на основі IP-адрес джерела та призначення.

Таблиця 1.1 – Співвідношення портів і часток балансування навантаження

Номер порту	Балансування навантаження
8	1:1:1:1:1:1:1:1
7	2:1:1:1:1:1:1
6	2:2:1:1:1:1
5	2:2:2:1:1
4	2:2:2:2
3	3:3:2
2	4:4

Неможливо вручну визначити, які кадри будуть передаватися через певний порт у групі агрегації. Вплинути на балансування трафіку можна лише шляхом вибору методу розподілу кадрів, який забезпечить найбільшу різноманітність навантаження. До таких методів належать:

- Балансування за MAC-адресою джерела або призначення;
- Балансування за IP-адресою джерела або призначення;
- Балансування за номером порту джерела або призначення;
- Балансування за номером VLAN;
- Комбінація цих методів (підтримується лише на Supervisor Engine 720);

При налаштуванні обладнання Cisco можна застосувати певну політику розподілу кадрів до портів, що входять до групи агрегації. Однак під час вибору методу балансування необхідно враховувати топологію мережі та схему передачі трафіку.

Наприклад, якщо всі пристрої знаходяться в одному VLAN, як показано на рис. 1.2, і шлюзом за замовчуванням є маршрутизатор R1, то:

- Для комутатора Sw2 підійдуть методи балансування за MAC-адресою призначення, IP-адресою джерела чи призначення, або за номером порту джерела чи призначення;
- Метод балансування за MAC-адресою джерела буде некоректним, оскільки весь трафік йтиме через один порт. Це відбувається через те, що всі кадри надходитимуть від одного джерела — маршрутизатора R1, MAC-адреса якого буде однаковою для всіх кадрів.

Аналогічно, для комутатора Sw1 метод балансування за MAC-адресою одержувача також буде неефективним. Тому при налаштуванні балансування трафіку слід ретельно враховувати ці аспекти, щоб уникнути нерівномірного розподілу навантаження.

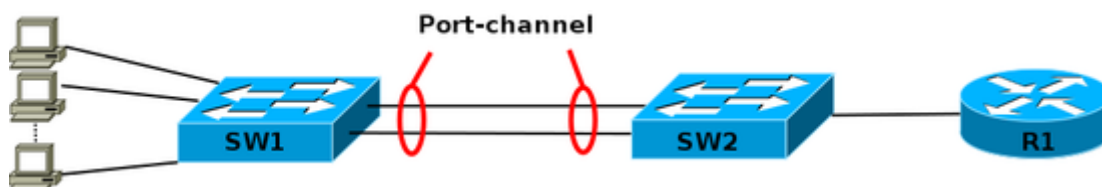


Рисунок 1.2 – Демонстраційна топологія

Попри збільшення пропускної здатності, досягнуте завдяки агрегації каналів, це збільшення не є сумою пропускної здатності всіх каналів, що входять до групи. Ідеального балансування навантаження, при якому трафік рівномірно розподіляється між усіма каналами, неможливо досягти. Балансування виконується на основі математичних операцій і залежить від можливостей мережевого обладнання та його налаштувань.

Результат балансування, який визначається за обраними параметрами, залишається постійним для схожих кадрів, незалежно від поточного стану каналу, на який вони будуть направлені. Це призводить до ситуацій, коли один із каналів групи може бути повністю перевантажений певним потоком. У такому випадку:

- Пропускна здатність для цього потоку обмежується шириною одного каналу;
- Нові потоки з певними характеристиками також можуть спрямовуватися на вже завантажений канал, оскільки балансування не враховує поточне навантаження каналу.

Результатом таких ситуацій є загублені кадри та збільшена затримка. Це є одним із очевидних недоліків агрегації каналів.

Для вирішення цієї проблеми можна застосувати кілька підходів:

- Розширити пропускну здатність окремих каналів, які входять до групи агрегації;
- Збільшити кількість каналів у групі агрегації. Це не усуває вузьке місце, але зменшує обсяг трафіку, який спрямовується на один канал;
- Виконати налаштування QoS. Можна використовувати обмежувач трафіку, щоб передбачити шаблони навантаження та забезпечити мінімальну гарантовану пропускну здатність для критично важливого трафіку. Однак це

може викликати проблему втрати кадрів і затримок. Обмежувач трафіку є більш надійним варіантом;

- Створити окремий канал для підключень із високим навантаженням, якщо це можливо. Але цей підхід не підходить для масштабування у великих мережах.

1.1.4 Агрегація і надійність

Збільшення надійності досягається завдяки основним функціям агрегації каналів. Підвищена пропускна здатність дозволяє зменшити ймовірність перевантаження мережі, уникнути втрати пакетів і знизити затримки, навіть за умов високого трафіку. Наявність надлишкових з'єднань експотенційно знижує ризик того, що точкою відмови стане канал із нижчою надійністю порівняно з обладнанням. З'єднання залишатиметься активним, поки хоча б один фізичний канал у групі залишається працездатним. У випадку зниження кількості лінків може спостерігатися деградація пропускної здатності, що особливо важливо враховувати при роботі з критично важливим трафіком. Для ефективнішого балансування бажано, щоб кількість активних каналів була кратною двом.

У магістральних мережах зазвичай прокладається більше кабелів, ніж потрібно на момент запуску. Це обумовлено тим, що вартість робочої сили значно перевищує вартість кабелів, і прокладання додаткових ліній одразу дозволяє знизити витрати в майбутньому. Агрегація каналів дає змогу ефективно використовувати ці додаткові кабелі для підвищення пропускної здатності магістралі практично без додаткових витрат, якщо доступні порти.

На кінцевих пристроях, таких як ПК користувачів чи серверне обладнання, агрегація каналів забезпечується мережевими картами, які працюють під керуванням драйверів операційної системи. Це рішення є ефективним для підвищення надійності та пропускної здатності серверів.

Однак створення групи агрегації переносить єдину точку відмови із самого з'єднання пристроїв до пристрою, до якого під'єднані всі канали. У випадку виходу з ладу такого пристрою кількість активних каналів між пристроями вже не матиме значення. Для усунення цієї точки відмови можна використовувати стекові комутатори та організувати кросстекову агрегацію каналів, що значно підвищить надійність мережі.

1.2 STP

Протокол STP контролює активну топологію мережі, забезпечуючи захист від петель трафіку на другому рівні моделі OSI. Його основною метою є запобігання утворенню петель, що сприяє стабільності та надійності мережі. Спочатку STP був описаний у стандарті IEEE 802.1D, але функціонал STP 802.1D, Rapid STP 802.1w і Multiple STP 802.1s було інтегровано до стандарту IEEE 802.1Q-2014.

Головна функція протоколу – усунення широкомовних петель, які можуть призвести до широкомовного шторму. Шторм виникає, коли у мережі є кілька шляхів до одного й того самого пристрою, і широкомовні повідомлення, наприклад, ARP або кадри з невідомими MAC-адресами, безперервно розмножуються у домені. Це може викликати:

- постійне оновлення таблиці MAC-адрес, що призводить до її нестабільності;
- високу завантаженість процесора комутаторів, через що вони втрачають здатність пересилати кадри;
- загальну деградацію мережі, що є неприйнятним для забезпечення надійності.

STP представляє з'єднання між пристроями у вигляді деревоподібної топології, залишаючи активним лише один шлях між будь-якими двома вузлами. Решта каналів переходять у заблокований стан, щоб уникнути циклічної передачі пакетів.

Протокол базується на алгоритмі Радії Перлман, і для обміну інформацією мости використовують два типи повідомлень: BPDU – для передачі даних про топологію, і TCN – для повідомлень про зміну топології.

У процесі роботи серед усіх комутаторів у широкомовному домені обирається кореневий міст. Це пристрій із найнижчим BID, який визначається під час обміну кадрами BPDU. Комутатори надсилають BPDU кожні дві секунди, включаючи інформацію про BID кореневого моста та комутатора-відправника. Вартість шляху до кореня розраховується як сума вартостей портів, що залежать від пропускної здатності: чим вона вища, тим нижча вартість.

Після обрання кореневого моста:

- Кожен некореневий комутатор визначає кореневий порт – це порт із найнижчою вартістю шляху до кореневого пристрою;
- На кожному сегменті обирається один призначений порт, який забезпечує найкращий шлях до кореневого моста;
- Решта портів переходять у заблокований стан і не пересилають трафік, забезпечуючи резервні шляхи.

Кінцевий результат – деревоподібна структура з вершиною у вигляді кореневого моста.

Якщо в мережі відбувається збій кабелю або пристрою, алгоритм STA перераховує шляхи та розблоковує порти, щоб активувати резервні з'єднання. Це забезпечує автоматичне відновлення мережі. Однак, при великих масштабах мережі процес перерахунку може займати надто багато часу, що є недоліком стандартного STP.

Під час роботи цього протоколу кожному мосту присвоюється ідентифікатор BID, який складається з пріоритету моста та MAC-адреси. Він використовується у процесі прийняття рішень STA, включаючи вибір кореневого моста та визначення ролей портів. За замовчуванням всі пристрої мають однаковий пріоритет, який може бути налаштований адміністратором. Чим менше значення пріоритету, тим вище його значення. BID разом із пріоритетом порту та його ідентифікатором

використовується для визначення кореневого порту у випадках, коли вартість шляхів до кореневого моста однакова.

Для досягнення збіжності протокол використовує наступні таймери:

- Hello Timer – інтервал між відправленням BPDU. За замовчуванням цей параметр становить 2 секунди, але його можна змінити в межах від 1 до 10 секунд;
- Forward Delay Timer – час, протягом якого порт перебуває у станах прослуховування та навчання. Стандартне значення – 15 секунд, але його можна налаштувати у межах від 4 до 30 секунд;
- Max Age Timer – максимальний час очікування перед тим, як комутатор спробує змінити топологію STP. За замовчуванням цей параметр дорівнює 20 секунд, проте його можна змінити в межах від 6 до 40 секунд. Під час роботи протоколу порти можуть мати наступні стани:
- Blocking – порт перебуває у стані блокування, кадри через нього не пересилаються. У цьому стані залишаються непризначені порти;
- Listening – порт отримує BPDU, визначаючи шлях до кореневого моста, і передає власні кадри BPDU. Також він інформує сусідні комутатори про свій намір долучитися до активної топології. У цьому стані порти залишаються до вибору кореневого моста та очікують завершення Forward Delay Timer перед переходом у стан навчання;
- Learning – порт отримує і обробляє BPDU, а також починає заповнювати таблицю MAC-адрес. Однак він ще не передає дані користувача. Перехід до наступного стану також залежить від завершення Forward Delay Timer;
- Forwarding – порт перебуває у звичайному робочому стані, де приймаються та відправляються як BPDU, так і дані користувача;
- Disabled – порт відключений адміністративно, не бере участі у протоколі єдиного дерева і не пересилає кадри.

Такий процес дозволяє уникнути петель та забезпечує стабільність мережевої топології.

1.3 RSTP

1.3.1 Введення

У сучасних мережах пріоритет віддається протоколу RSTP, який поступово витісняє STP. Протокол STP був створений у часи, коли час відновлення з'єднання після збою, що тривав до хвилини, вважався прийнятним. Зараз STP використовується переважно на застарілому обладнанні, яке не підтримує новіші протоколи. За наявності можливості використання RSTP причини для активації STP практично відсутні.

RSTP сумісний зі своїм попередником, що дозволяє його використання у мережах зі старими комутаторами. Однак така сумісність позбавляє RSTP ключових переваг, зокрема швидкості збіжності.

Як випливає з назви Rapid Spanning Tree Protocol, цей протокол є швидшою версією STP. Він забезпечує значно скорочений час збіжності, тобто швидше визначає активну топологію та призначає ролі портів, такі як кореневі, призначені та альтернативні порти. Завдяки цьому відновлення мережі у разі змін або збоїв відбувається значно оперативніше. У порівнянні з STP, який через Forward Delay Timer потребує понад 30 секунд для збіжності, RSTP досягає тієї ж мети менш ніж за 1 секунду у більшості випадків.

1.3.2 Переваги та зміни порівняно з STP

Протокол RSTP не може забезпечити захист від тимчасових петель, які виникають через з'єднання двох сегментів локальної мережі пристроями, що не є мостами. Такі пристрої працюють без підтримки служб внутрішнього підрівня MAC мостів. Використання хабів також унеможливорює реалізацію переваг, які пропонує RSTP.

У протоколі STP Max Age Timer, що відповідає за час очікування перед переглядом топології при збоях, становить 20 секунд. У RSTP службові

повідомлення виконують функцію Hello-пакетів. При втраті трьох таких повідомлень, тобто через 6 секунд, алгоритм починає переглядати топологію.

У RSTP кількість станів портів була скорочена до трьох: discarding, learning, forwarding. Стани disabled, blocking та listening об'єдналися в discarding. Наприклад, порт може одночасно бути призначеним і заблокованим, що зазвичай триває лише короткий час. Це вказує на перехідний стан перед активацією функції переадресації. Завдяки цьому стало можливим більш точне уявлення про поточну топологію мережі та швидшу реакцію на її зміни.

Ролі кореневого та призначеного портів залишилися незмінними. Водночас роль заблокованого порту була розділена на резервний та альтернативний порти. Цей підхід нагадує пропрієтарне рішення Cisco UplinkFast, де, у разі виходу з ладу кореневого порту, автоматично активується порт із наступною найменшою вартістю шляху до кореневого моста.

Для збереження порту в стані discarding необхідно, щоб він отримував BPDU. Альтернативний порт, зображений на рис. 1.3, отримує більш пріоритетні BPDU від іншого моста та виконує роль резерву для заміни кореневого порту. Резервний порт, як показано на рис. 1.4, отримує пріоритетні BPDU від моста, до якого він підключений. Цей порт забезпечує резервне підключення до того самого сегмента і не гарантує альтернативного підключення до кореневого моста.

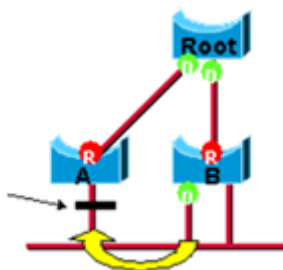


Рисунок 1.3 – Топологія на котрій зображено альтернативний порт

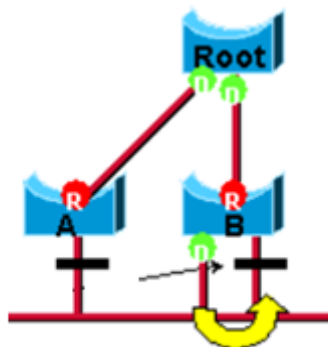


Рисунок 1.4– Топологія на котрій зображено альтернативний порт

Протокол STA чекав завершення процесу збіжності мережі перед тим, як перевести порт у стан переадресації. Швидкість збіжності залежала від налаштування таймерів Hello Timer, Forward Delay Timer та Max Age Timer. Однак зміна цих параметрів могла негативно вплинути на стабільність мережі.

Протокол RSTP дозволяє порту перейти у стан переадресації без необхідності покладатися на налаштування таймерів. Цього вдалося досягти завдяки механізму Proposal/Agreement, що забезпечує швидке узгодження між мостами. Для пришвидшення збіжності протокол вводить два нові поняття: граничні порти (edge ports) і типи з'єднань (link type).

Граничний порт виконує функцію, аналогічну пропрієтарному рішення STP – PortFast. Порти, підключені до кінцевих пристроїв, не можуть створити петель у мережі. Через це зміни на граничних портах не змушують протокол RSTP переглядати топологію мережі. Такі порти можуть одразу переходити у стан переадресації, минаючи етапи прослуховування та навчання. Якщо граничний порт отримує BPDU, він втрачає цей статус і стає звичайним портом. Тип граничного порту задається вручну, оскільки комутатор не може автоматично визначити, який пристрій до нього підключений. У разі автоматизації потрібно чекати закінчення Forward Delay Timer, під час якого BPDU не має надходити.

Тип з'єднання зазвичай визначається комутатором автоматично, спираючись на режим роботи: повний або напівдуплекс. Порти в повнодуплексному режимі класифікуються як з'єднання "точка-точка", тоді як напівдуплексні порти

вважаються спільними. Останні можуть бути підключені до кількох комутаторів через хаб, що є рідкістю у сучасних мережах. Визначення типу з'єднання можна змінити вручну за допомогою налаштувань.

У сучасних комутуваних мережах більшість з'єднань працюють у повнодуплексному режимі, що дозволяє класифікувати їх як "точка-точка". Це робить їх придатними для швидкого переходу у стан переадресації. Протокол RSTP забезпечує такий перехід лише для граничних портів і з'єднань "точка-точка".

Як вже зазначалося, процес збіжності топології у протоколі STP триває щонайменше 30 секунд. Це пов'язано з тим, що під час змін у топології порт повинен двічі пройти через стани прослуховування та навчання, і кожного разу чекати завершення таймера Forward Delay. Таким чином, виникає мінімум 30 секунд переривання трафіку. Це зумовлено відсутністю у алгоритмі 802.1D механізму зворотного зв'язку, який би дозволив швидко підтвердити збіжність мережі.

У протоколі RSTP процедура досягнення збіжності була змінена. Коли до активної топології додається новий пристрій, порт, підключений до нового сегмента, і порт самого нового пристрою спершу переводяться у стан блокування, як це відбувається і у STP. Цей процес ілюструється на рис. 1.5. Відмінність полягає у використанні механізму Proposal/Agreement, що забезпечує швидке узгодження між комутаторами.

Переговори починаються між двома комутаторами, які мають спільне з'єднання. Призначений порт надсилає BPDU із встановленим бітом пропозиції, якщо він перебуває у стані discarding або learning. Це відображено на рис. 1.5, де порт кореневого моста надсилає таку пропозицію. Комутатор А, отримавши BPDU з кращими метриками, негайно визнає порт, що веде до кореневого моста, новим кореневим портом.

Цей процес дозволяє значно скоротити час збіжності мережі у порівнянні з протоколом STP, забезпечуючи швидке та ефективне визначення активної топології.

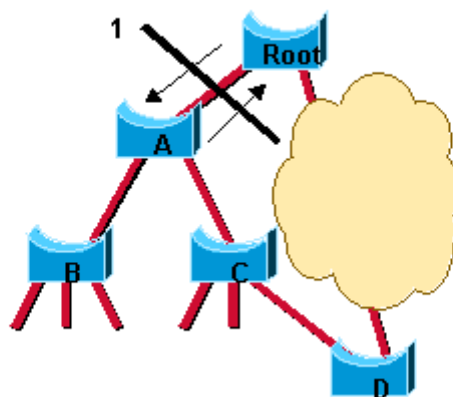


Рисунок 1.5 – Топологія мережі з налаштованим RSTP

Після визначення нового кореневого порту комутатор А розпочинає процес синхронізації, щоб перевірити, чи всі його порти узгоджені з оновленою інформацією. Порт вважається синхронізованим, якщо він відповідає одному з таких критеріїв:

- Порт перебуває у стані blocking (що відповідає Discarding у стабільній топології);
- Порт є граничним.

Усі інші порти, що не відповідають цим критеріям, а саме призначені неграничні порти, тимчасово блокуються. Після завершення синхронізації всіх своїх портів комутатор А розблоковує новий кореневий порт і надсилає у відповідь повідомлення про згоду. Це повідомлення є копією BPDU пропозиції, в якому біт пропозиції замінений на біт згоди. Такий механізм гарантує, що призначений порт кореневого комутатора точно знає, на яку пропозицію надано згоду. Отримавши згоду, порт негайно переходить у стан forwarding.

Якщо ж призначений порт кореневого комутатора не отримує згоди після надсилання пропозиції, він переходить у стан forwarding за традиційною процедурою, визначеною стандартом 802.1D для STP. Це може статися, якщо віддалений комутатор не підтримує BPDU RSTP.

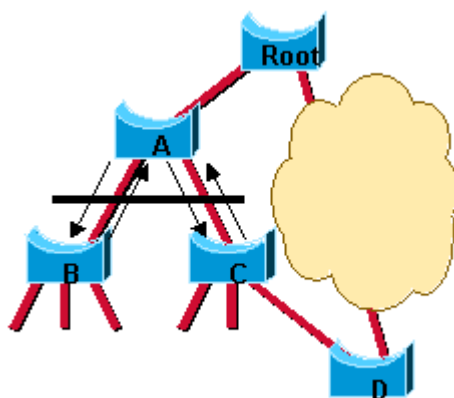


Рисунок 1.6 – Зміни топології мережі з налаштованим RSTP

Замість блокування з'єднань у верхніх рівнях деревовидної ієрархії, як це реалізовано у STP, протокол RSTP обмежує з'єднання ближче до нижніх рівнів. Це проілюстровано на рис. 1.6. У результаті потенційна петля розривається в іншому місці. Такий "розріз" переміщується вниз по дереву разом із новими BPDU. Далі нові заблоковані порти на комутаторі проводять узгодження зі своїми сусідніми портами, щоб швидко перейти в стан переадресації. Ця домовленість каскадно ініціює синхронізацію на інших комутаторах. Остаточна топологія мережі досягається за той час, який потрібен для поширення BPDU вниз по дереву.

У процесі цього переходу комутатори блокують надлишкові з'єднання так само, як це описано у протоколі STP. Важливо, що таймери не залучені до цієї процедури. Двійна затримка, яка була у STP, обходиться завдяки використанню механізму пропозиції. Цей механізм дозволяє негайно перевести порт у стан переадресації за дотримання таких умов:

- Мости мають бути з'єднані зв'язком типу "точка-точка";
- Граничні порти повинні бути налаштовані належним чином, оскільки їх роль у цьому протоколі стала ще важливішою.

У STP при зміні топології спочатку повідомляється кореневий міст за допомогою повідомлення TCN. У RSTP специфічний TCN BPDU більше не використовується, за винятком випадків, коли потрібно повідомити застарілий міст.

Зміна топології може спричинити ситуацію, коли MAC-адреси стають доступними через інші порти, через що комутатор може направляти пакети на хибні адреси. У такому разі запускається процедура оповіщення всіх пристроїв про зміни. Повідомлення TCN надсилається через кореневий порт, але не містить деталей про характер і місце змін, а лише вказує, що в мережі щось змінилося. Цей механізм у STP був створений для того, щоб некореневі комутатори могли сповіщати про зміни в мережі.

Коли кореневий міст у протоколі STP отримує інформацію про зміну топології, він додає прапорець TC у BPDU, які він надсилає. Ці BPDU поширюються всією мережею, а комутатори, що їх отримують, зменшують час старіння записів у своїх таблицях MAC-адрес до часу затримки переадресації. Це дозволяє швидко оновити застарілі дані.

У RSTP цей механізм було суттєво вдосконалено. Як виявлення змін у топології, так і поширення інформації про них були значно оптимізовані. У цьому протоколі зміна топології фіксується лише у випадку, якщо non-edge порт переходить у стан переадресації. Втрата зв'язку більше не вважається зміною топології, на відміну від STP, де будь-який порт, що переходить у стан блокування, ініціював зміну топології.

Коли у RSTP виявляється зміна топології, виконуються такі дії:

- Активується таймер TC While зі значенням, що дорівнює подвоєному часу Hello Timer, для всіх неграничних призначених портів і кореневого порту, якщо це необхідно;
- З таблиці MAC-адрес очищаються записи, пов'язані з цими портами.
- Поки працює таймер TC While, BPDU з бітом TC надсилаються з відповідного порту і ретранслюються на кореневий порт.

Коли комутатор отримує BPDU з бітом TC від сусіднього пристрою, відбуваються такі дії:

- Очищаються записи MAC-адрес на всіх портах, крім того, з якого отримано повідомлення про зміну;

— Активується таймер TC While, а BPDU з бітом TC надсилаються на всі призначені порти та кореневий порт.

Це дозволяє поширити TCN по всій мережі набагато швидше. Процес поширення TC у RSTP виконується за один крок, і ініціатором змін є комутатор, на якому сталася подія. На відміну від STP, де зміни повідомляються лише через кореневий міст, цей механізм значно прискорює оновлення мережі.

RSTP включає вдосконалення STP, які були пропрієтарними рішеннями Cisco, такими як BackboneFast, UplinkFast та PortFast. Завдяки цьому протокол забезпечує значно швидшу збіжність за умови правильної конфігурації. Таймери STP, такі як Forward Delay і Max Age, використовуються тільки як резерв, якщо з'єднання типу точка-точка та граничні порти налаштовані коректно. Якщо у мережі немає застарілих пристроїв, таймери стають зайвими.

1.4 MSTP

1.4.1 Введення

Протокол MST забезпечує створення окремих топологій, званих інстансами, для груп VLAN, що дозволяє уникати петель комутації. Окрім цього, MST дозволяє балансувати навантаження, використовуючи окремі канали для різних груп VLAN, і виключає прості канали, які мають місце при застосуванні протоколу RSTP.

Протокол MST побудований на основі алгоритму RSTP, тому логіка його роботи всередині інстансів відповідає протоколу RSTP. При використанні стандартного RSTP створюється одна загальна топологія для всіх VLAN. Такий підхід має обмеження, оскільки не дозволяє розподіляти трафік між різними каналами навіть у ручному режимі. Це призводить до втрати щонайменше половини пропускної здатності у разі наявності надлишкових шляхів. MSTP описано у стандарті IEEE 802.1s.

Існують пропрієтарні рішення, що забезпечують подібний функціонал. Наприклад, Cisco пропонує Rapid Per-VLAN Spanning Tree PVRST+, у якому для

кожного VLAN створюється окрема топологія. Це дозволяє ефективно використовувати канали, але має низку недоліків:

- Збільшення службового трафіку, оскільки кожен інстанс надсилає свої BPDU з інтервалом у 2 секунди;
- Обмеження на максимальну кількість інстансів. Наприклад, комутатори Cisco 2960 підтримують максимум 128 інстансів STP;
- Значне споживання апаратних ресурсів пристрою для забезпечення роботи кожної топології, які не є безмежними.

Таким чином, MST забезпечує кращу ефективність і балансування навантаження, уникаючи багатьох недоліків інших реалізацій, хоча і має певні обмеження.

1.4.2 Функціонування

Протокол MST позбавлений багатьох проблем, притаманних іншим протоколам, завдяки своїй здатності об'єднувати кілька VLAN в один інстанс RSTP. Більшість мереж не потребує великої кількості окремих логічних топологій, тому такий підхід дозволяє зменшити навантаження на обчислювальні потужності та забезпечити ефективне балансування трафіку. MST поєднує кращі властивості PVST+ та RSTP, що робить його відмінним вибором для мереж з обладнанням різних виробників. Основним недоліком MST є складність налаштування, що потребує кваліфікованих фахівців.

Для вирішення питання передачі BPDU, які дозволяють ідентифікувати інстанси та VLAN кожного пристрою, було введено поняття регіону. Регіон можна порівняти з автономною системою у BGP. Це група комутаторів із однаковими параметрами конфігурації: ім'ям, параметром зміни конфігурації, а також списком інстансів, пов'язаних із VLAN з діапазону до 4096.

Передача інформації про налаштування забезпечується BPDU. Проте замість передачі повної таблиці зіставлення VLAN з інстансами надсилається дайджест цієї таблиці, разом із номером версії та ім'ям конфігурації. Дайджест — це числове

значення, отримане шляхом математичного обчислення таблиці зіставлення VLAN-інстансів. Коли комутатор отримує BPDU, він порівнює отриманий дайджест із власним. Якщо значення відрізняються, порт, на який надійшов BPDU, вважається межею регіону.

Кожен регіон MST має інстанс 0 (IST), створений за замовчуванням. Усі VLAN, які не входять до інших інстансів, належать до IST. Цей інстанс служить для передачі BPDU всередині та зовні регіону і представляє весь MST-регіон як віртуальний міст CST для взаємодії з іншими регіонами та пристроями STP, RSTP, PVST+ і RPVST+.

Common Spanning Tree — це загальне дерево, яке об'єднує регіони MST як окремі пристрої, забезпечуючи взаємодію між ними та іншими типами топологій. Common and Internal Spanning Tree (CIST) поєднує функції CST і IST кожного регіону, створюючи єдине сполучне дерево.

На рис. 1.7 показано дві діаграми, що ілюструють взаємодію IST і CST. IST працює у топології CST як звичайний комутатор. У представленому прикладі, замість блокування порту між комутаторами М і В, блокується порт між кореневим комутатором і В, а також між С і D, що демонструє ефективність організації топології MST.

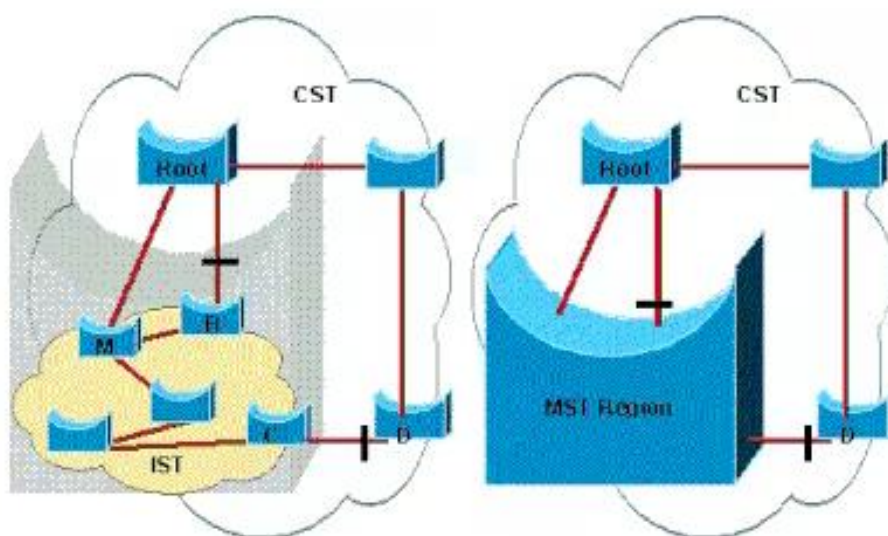


Рисунок 1.7 – Топологія CST

MSTI — це окремі інстанси RSTP, які існують виключно всередині регіону MST. Вони не взаємодіють із зовнішніми пристроями і не надсилають BPDU за межі регіону, оскільки ці функції виконує IST. На відміну від PVRST+, MSTI не генерують окремі BPDU для кожного VLAN. Натомість мости надсилають єдиний BPDU, який містить інформацію про всі інстанси, незалежно від того, чи надсилається він комутатору всередині регіону чи за його межі.

У межах MST регіону мости використовують MST BPDU, які за структурою дуже схожі на BPDU RSTP, але доповнені додатковими полями для кожного MSTI. Замість стандартної інформації RSTP у BPDU MSTP вказуються дані про кореневий міст IST, а також інформація про BID кожного MSTI і кореневі мости кожного інстансу.

Одним із головних переваг MSTP є те, що всі VLAN та інстанси об'єднуються в один BPDU, який містить усі необхідні дані. Це суттєво економить пропускну здатність у порівнянні з PVST+ та RPVST+, де кожен інстанс генерує окремий BPDU.

Реалізація Cisco дозволяє використовувати до 65 інстансів, серед яких один IST і 64 MSTI. Це оптимальне рішення для масштабних мереж, оскільки забезпечує ефективне управління трафіком і мінімізує навантаження на канали. Склад MSTP BPDU ілюстровано на рис. 1.8, що демонструє структуру переданих даних.

```

# Spanning Tree Protocol
  Protocol Identifier: Spanning Tree Protocol (0x0000)
  Protocol Version Identifier: Multiple Spanning Tree (3)
  BPDU Type: Rapid/Multiple Spanning Tree (0x02)
  BPDU flags: 0x38, Forwarding, Learning, Port Role: Root
  # Root Identifier: 24576 / 0 / 50:00:00:03:00:00
    Root Bridge Priority: 24576
    Root Bridge System ID Extension: 0
    Root Bridge System ID: 50:00:00:03:00:00 (50:00:00:03:00:00)
  Root Path Cost: 0
  # Bridge Identifier: 24576 / 0 / 50:00:00:03:00:00
    Bridge Priority: 24576
    Bridge System ID Extension: 0
    Bridge System ID: 50:00:00:03:00:00 (50:00:00:03:00:00)
  Port Identifier: 0x8002
  Message Age: 0
  Max Age: 20
  Hello Time: 2
  Forward Delay: 15
  Version 1 Length: 0
  Version 3 Length: 96
  # MST Extension
    MST Config ID format selector: 0
    MST Config name: Note
    MST Config revision: 0
    MST Config digest: 7ff3063c971d3b795396ed400d7fa9a2
    CIST Internal Root Path Cost: 20000
  # CIST Bridge Identifier: 32768 / 0 / 50:00:00:01:00:00
    CIST Bridge Priority: 32768
    CIST Bridge Identifier System ID Extension: 0
    CIST Bridge Identifier System ID: 50:00:00:01:00:00 (50:00:00:01:00:00)
  CIST Remaining hops: 20
  # MSTID 1, Regional Root Identifier 24576 / 50:00:00:03:00:00
    # MSTI flags: 0x38, Forwarding, Learning, Port Role: Root
    0110 .... = Priority: 0x6
    .... 0000 0000 0001 = MSTID: 1
    Regional Root: 50:00:00:03:00:00 (50:00:00:03:00:00)
    Internal root path cost: 20000
    Bridge Identifier Priority: 8
    Port identifier priority: 8
    Remaining hops: 20
  # MSTID 2, Regional Root Identifier 24576 / 50:00:00:02:00:00
    # MSTI flags: 0x7c, Agreement, Forwarding, Learning, Port Role: Designated
    0110 .... = Priority: 0x6
    .... 0000 0000 0010 = MSTID: 2
    Regional Root: 50:00:00:02:00:00 (50:00:00:02:00:00)
    Internal root path cost: 20000
    Bridge Identifier Priority: 8
    Port identifier priority: 8
    Remaining hops: 19

```

Рисунок 1.8 – Склад кадру BPDU

Запис MIST не вимагає використання будь-яких таймерів, як hello time, forward delay і Max adge, які зазвичай містяться у звичайному CST BPDU. Тільки IST використовує ці параметри, а параметр затримки пересилання переважно використовують, коли швидкий перехід неможливий. Оскільки MSTI залежать від IST для передачі своєї інформації, MSTI не потребують цих таймерів

1.4.3 Особливості налаштування

Записи MSTI не вимагають використання стандартних таймерів, таких як hello time, forward delay і max age, які зазвичай присутні у BPDU для CST. Використання цих параметрів залишається актуальним лише для IST, який виконує функції управління всією топологією регіону MST. Таймер forward delay, наприклад, застосовується виключно в ситуаціях, коли швидкий перехід неможливий.

Оскільки MSTI повністю залежать від IST для передачі своєї інформації та взаємодії із зовнішніми системами, ці інстанси не потребують власних таймерів. Це

спрощує структуру управління MSTI, зменшує обчислювальні витрати і забезпечує ефективну роботу протоколу MSTP.



Рисунок 1.9 – Неправильно налаштована топологія

Кращим підходом для налаштування MSTP є створення окремого інстансу для кожної групи VLAN. Це дозволяє уникнути зіставлення VLAN з інстансом IST, що може створити проблеми в роботі мережі. Головна причина цього полягає в тому, що MSTP передає всю інформацію про топологію за допомогою одного BPDU, незалежно від кількості внутрішніх інстансів.

У контексті MST VLAN більше не є тотожними інстансу сполучного дерева. Топологія визначається конкретним інстансом, а VLAN, які з ним зіставлені, наслідують цю топологію. Проблема виникає, коли VLAN видаляється з транкового каналу з метою балансування трафіку. У цьому випадку мережа може працювати некоректно, оскільки всі VLAN, пов'язані з інстансом, мають єдиний активний шлях.

Якщо VLAN видаляється з транкового інтерфейсу, через цей інтерфейс трафік для цього VLAN передаватися не буде. Через те, що MSTP забезпечує лише один активний шлях для інстансу, VLAN стає недоступним в інших частинах мережі. Це може призвести до повної втрати зв'язності для цього VLAN.

Тому для забезпечення коректної роботи мережі необхідно ретельно планувати структуру інстансів і їх зіставлення з VLAN, уникаючи ситуацій, де ручне втручання може викликати помилки в топології.

1.4.4 Взаємодія з зовнішнім світом

Протокол MST забезпечує взаємодію між регіонами MST і зовнішніми протоколами за допомогою BPDU IST. Внутрішні екземпляри MSTI автоматично узгоджуються з топологією IST на граничних портах. Для кожного інстансу IST або CIST вибирається кореневий міст, який визначає логіку роботи протоколу. MST-регіони представляються для зовнішніх пристроїв як єдиний великий віртуальний комутатор, а взаємодія між комутаторами відбувається за принципами, описаними в RSTP.

У кожному MST-регіоні кожне середовище має свій кореневий міст. Вибір кореневого мосту для кожного середовища здійснюється за тими ж правилами, що й у RSTP. Після визначення кореневого мосту для CIST у регіоні обирається Regional Root Bridge, який виконує функції управління топологією в межах регіону. Regional Root Bridge обирається в іншому регіоні як комутатор із найменшою Root Path Cost до кореневого мосту CIST.

Path Cost розділяється на зовнішній і внутрішній. Root Path Cost для MSTI 0 є зовнішнім, а для інших MSTI внутрішній Root Path Cost вказується в полях MST Extension.

Кожен комутатор у регіоні має Root Port, який забезпечує з'єднання з Regional Root. Крім того, у кожному регіоні обирається Master порт, який відповідає за з'єднання з кореневим мостом CIST. У регіоні може бути лише один Master порт, і він обирається за такими критеріями:

- Найменша ціна шляху до кореневого мосту CIST;
- Найнижче значення Bridge ID;
- Найнижчий CIST BID.

Порти, що знаходяться на межі з іншим регіоном MST або іншими протоколами STP, позначаються як Link Type Boundary. Master порт передає інформацію про CIST Root Bridge всередині регіону. Цей порт приймає BPDU, але самостійно не відправляє їх, за винятком BPDU з прапором TC. У BPDU, переданих граничним комутатором, Root Identifier збігається з Bridge Identifier, оскільки весь регіон MST представляється як один комутатор.

Регіони MST автоматично виявляють сусідів, які використовують протокол PVST+, завдяки здатності розпізнавати декілька BPDU, що визначають VLAN для конкретного інстансу. Регіон MST імітує роботу сусіда PVST+, надсилаючи копії BPDU IST для всіх VLAN.

Для коректної взаємодії MST і PVST+ існує два варіанти налаштування:

- Кореневий міст для всіх інстансів PVST+ знаходиться у регіоні MST (рекомендований варіант).
- Кожен інстанс PVST+ має кращий кореневий міст, ніж IST.

Інші конфігурації можуть призвести до некоректної роботи протоколів. Такий підхід забезпечує сумісність і стабільність роботи в змішаних середовищах.

Перевагою першого варіанту конфігурації є можливість балансування навантаження за наявності декількох з'єднань між мостами PVST+ та MST. Це досягається завдяки тому, що кожен інстанс PVST+ створює окреме сполучне дерево для кожної VLAN, що дозволяє комутатору MST вибирати, які порти висхідного зв'язку блокувати для кожної VLAN. Такий підхід забезпечує гнучкість та оптимальне використання ресурсів мережі.

Недоліком другого варіанту є те, що в регіоні MST працює лише один інстанс сполучного дерева, який взаємодіє із зовнішнім світом. У результаті між комутаторами використовується тільки один активний канал, тоді як інші залишаються заблокованими. Це обмежує можливості балансування трафіку та знижує ефективність використання мережевих ресурсів.

Неправильною конфігурацією вважається ситуація, коли регіон MST або комутатор MST не є кореневим для всіх інстансів. Це може призвести до некоректної роботи мережевої топології, конфліктів між протоколами та втрати продуктивності. Приклад такої ситуації показано на рисунку 1.10.

Перший варіант конфігурації є найбільш оптимальним, оскільки забезпечує балансування навантаження та ефективне використання мережевих ресурсів, тоді як другий варіант має значні обмеження, а неправильна конфігурація може викликати серйозні проблеми у роботі мережі.

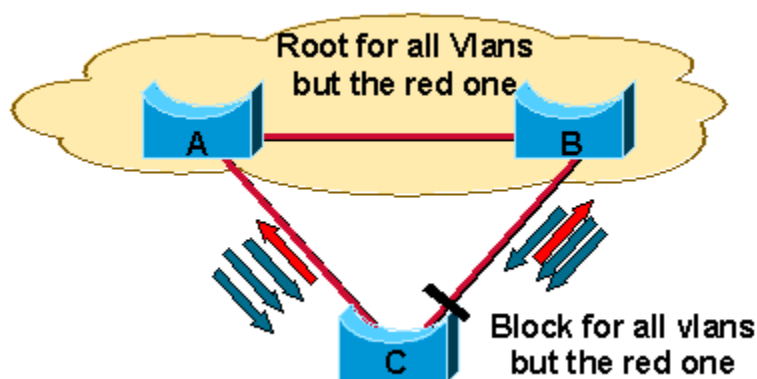


Рисунок 1.10 – Топологія де зображено взаємодія MSTP та PVST+

На схемі зображені мости A і B, які належать до регіону MST, та міст C, налаштований із використанням протоколу PVST+. Міст A виступає кореневим для трьох інстансів PVST+, окрім одного. Для одного з VLAN, позначеного червоним кольором, кореневим є міст C.

У цій VLAN утворюється петля, яка блокується мостом B. Це означає, що міст B є призначеним для цієї VLAN. Проте MST-регіон не має можливості зробити такий розподіл, оскільки граничний порт MST може або пересилати, або блокувати кадри одночасно для всіх VLAN. Це обумовлено тим, що MST використовує лише одну топологію для зв'язку із зовнішніми пристроями.

Коли міст B отримує BPDU з кращою метрикою на своєму граничному порту, активується BPDU-guard, який переводить порт у стан блокування, відомий як режим кореневого неузгодження. Аналогічно, міст A також блокує свій граничний порт через той самий механізм. У результаті з'єднання між мостами втрачається.

Попри те, що петлі у мережі вдається уникнути навіть у такій ситуації, вузол C стає недоступним, що порушує роботу мережі. Це ще раз підкреслює важливість коректного налаштування протоколів і взаємодії між ними, щоб уникати подібних проблем.

1.4.5 Висновки

Протокол STP застарів, та його використання не є рекомендованим при можливості активувати більш сучасні версії протоколу.

Варіації протоколу STP все ще використовуються у цей час, проте в більшості сучасних мереж їх основне застосування – це механізм захисту від петель комутації, а не механізм забезпечення резервних каналів. Не можна сказати, що він вже не потрібен. У провайдерських мережах його не позбудешся. У розподіленій на відстані корпоративній мережі теж. Повністю відмовитися від STP дуже складно. Навіть якщо дизайн мережі не передбачає петель, вона може легко виникнути в рамках лише одного комутатора, і тоді без використання STP мережа вийде з ладу.

RSTP набагато швидше сходиться та змінює топологію при правильному налаштуванні. MSTP є гарною альтернативою між стандартним RSTP та пропрієтарним PVRST+. Класичний RSTP блокує цілий порт, чим знижує продуктивність кільця. Використовуючи MSTP для VLAN, теоретично можна блокувати порт у кільці для VLAN, а для іншого VLAN - інший порт. Якщо в одному і кожному VLAN пересувається по 1 гігабайту трафіку, порти в кільці не будуть перевантажені, деградація буде тільки в разі обриву.

Коли у мережі трапляється відмова або зміна топології, важливим аспектом її структури є швидка і передбачувана збіжність. Протоколи маршрутизації 3 Рівня забезпечують ефективність і передбачуваність більшу ніж STP. Маршрутизація Рівня 3 забезпечує створення резервних шляхів і контроль за петлями у топології без блокування портів. З цієї причини деякі середовища переходять на Рівень 3 всюди, крім кінцевих пристроїв. Тобто з'єднання між комутаторами рівня доступу та комутаторами рівня розподілу утворюються на Рівні 3 замість Рівня 2.

1.5 Методи резервування на мережевому рівні.

1.5.1 FHRP

У локальних мережах для кожного хоста зазвичай налаштовано лише один шлюз за замовчуванням. Це створює проблему у випадках, коли цей шлюз стає недоступним, навіть якщо в мережі існують інші шляхи для передачі трафіку. Одним із рішень цієї проблеми є використання протоколів сімейства FHRP, які забезпечують резервування першого шлюзу та гарантують безперервну роботу користувацького трафіку навіть у разі збоїв.

Протоколи FHRP дозволяють створювати альтернативні шлюзи за замовчуванням у мережах, де два або більше маршрутизаторів підключені до одних VLAN. Для цього створюється віртуальний маршрутизатор, який об'єднує декілька фізичних маршрутизаторів в один логічний пристрій. Ці маршрутизатори спільно використовують одну IP- та MAC-адресу, що забезпечує прозорість для кінцевих вузлів.

Коли вузли надсилають кадри до шлюзу за замовчуванням, вони використовують ARP для визначення MAC-адреси, прив'язаної до IP-адреси шлюзу. Кадри, які надсилаються на MAC-адресу віртуального маршрутизатора, обробляються активним маршрутизатором із групи віртуальних маршрутизаторів. У разі відмови активного маршрутизатора резервний маршрутизатор активізується, замінюючи відсутній. Це відбувається, коли резервний маршрутизатор перестає отримувати привітальні повідомлення від активного маршрутизатора.

Основною функцією резервного маршрутизатора HSRP є моніторинг стану активного маршрутизатора у групі HSRP і швидке прийняття на себе функцій передачі пакетів у разі збою.

Зараз найбільш поширеними протоколами цього сімейства є HSRP, GLBP, VRRPv3. Окрім FHRP, існують і інші способи динамічного визначення альтернативних маршрутизаторів, наявних у мережі. Однак багато з них не

гарантують належного рівня відмовостійкості. Це може бути зумовлено тим, що протокол не розроблявся з урахуванням таких вимог або через те, що не всі хости в мережі підтримують запуск цього протоколу. Для більшості хостів налаштування лише одного шлюзу за замовчуванням залишається стандартом, що підкреслює важливість використання FHRP у мережах із високими вимогами до надійності.

1.5.2 HSRP

HSRP є власною реалізацією протоколу із сімейства FHRP, розробленою компанією Cisco. Він підтримує налаштування для роботи як з IPv4, так і з IPv6.

Для забезпечення відмовостійкості пристрою першого переходу створюється група маршрутизаторів, що утворюють єдиний логічний маршрутизатор. У процесі налаштування пристроїв доступні дві основні ролі:

- Active – активний пристрій, відповідальний за обробку трафіку;
- Standby – резервний пристрій, який активується у разі відмови основного і перебуває в очікуванні у звичайному режимі.

Хоча HSRP дозволяє використовувати кілька маршрутизаторів, лише активний перенаправляє пакети, адресовані віртуальному маршрутизатору. Ця проблема вирішена в протоколі GLBP, що забезпечує балансування навантаження між пристроями в межах віртуального маршрутизатора.

У мережі можна налаштувати кілька віртуальних маршрутизаторів завдяки використанню різних номерів груп HSRP. Пристрої, що входять до групи HSRP, прослуховують MAC-адреси групи HSRP та власні MAC-адреси. HSRP застосовує MAC-адресу, специфічну для кожного носія: HSRPv1 – «0000.0c07.ac**»; HSRPv2 – «00:00:0c:9f:fx:xx», де ** – це номер групи. Адреси групової розсилки: HSRPv1 – 224.0.0.2; HSRPv2 – 224.0.0.102.

Віртуальний маршрутизатор має одну IP- та MAC-адресу, яка призначається виключно активному пристрою. Резервні пристрої використовують свої власні IP- та MAC-адреси.

Після налаштування HSRP або активації з наявною конфігурацією маршрутизатор починає надсилати та отримувати привітальні пакети, щоб визначити свій статус у групі. Стани маршрутизатора включають:

- Initial – початковий стан після зміни конфігурації або при першому доступі інтерфейсу;
- Learn – навчальний стан, коли маршрутизатор чекає привітальні пакети від активного маршрутизатора, не маючи визначеної віртуальної IP-адреси;
- Listen – прослуховування, маршрутизатор знає віртуальну IP-адресу, але не є активним чи резервним;
- Speak – маршрутизатор активно надсилає привітальні пакети та бере участь у виборі активного чи резервного маршрутизатора;
- Standby – резервний маршрутизатор, який періодично надсилає привітальні пакети та готовий стати активним.

Привітальні пакети також слугують для виявлення відмови активного маршрутизатора. Вони надсилаються кожні 3 секунди. Якщо резервний маршрутизатор не отримує пакети від активного протягом 10 секунд, він стає активним. Ці значення можна змінювати, але рекомендовано встановлювати інтервали не менше 1 секунди для привітальних пакетів і не менше 4 секунд для таймерів утримання, щоб уникнути надмірного навантаження на систему.

Механізм вибору основного та резервного маршрутизаторів у HSRP за замовчуванням передбачає визначення пристрою з найбільшим числовим значенням IP-адреси; проте цей процес можна налаштовувати через параметр пріоритету HSRP. Значення пріоритету за замовчуванням становить 100 для всіх пристроїв, а його діапазон варіюється від 0 до 255; активним маршрутизатором стає той, у якого пріоритет є вищим.

Активний маршрутизатор зберігає свій статус навіть за появи пристрою з більшим пріоритетом HSRP; ця поведінка змінюється шляхом активації режиму пріоритетного перевибору. У такому режимі маршрутизатор із вищим пріоритетом

автоматично стає активним, навіть якщо попередній активний пристрій повертається в мережу після відмови.

У локальній мережі можна налаштувати кілька незалежних груп HSRP, кожна з яких імітує окремий маршрутизатор. Функція множинних груп HSRP дозволяє забезпечувати надмірність та розподіл навантаження, максимально використовуючи резервні пристрої; один маршрутизатор може активно перенаправляти трафік для однієї групи та одночасно перебувати в режимі очікування або прослуховування для іншої.

Функція відстеження інтерфейсів дає змогу контролювати стан інших портів пристрою; у разі їх відмови пріоритет маршрутизатора автоматично знижується на задану величину.

HSRP забезпечує надмірність для IP-маршрутизації без збереження стану; кожен маршрутизатор підтримує власну таблицю маршрутизації незалежно від інших. Ця технологія дозволяє HSRP забезпечувати роботу клієнтських додатків із можливістю обробки відмов. На сьогодні реалізовано резервування для агентів мобільної IP-телефонії; очікується впровадження резервування для таких функцій:

- NAT;
- IPSEC;
- DHCP.

Функція автентифікації HSRP базується на використанні спільного відкритого текстового ключа між членами групи для виявлення конфігураційних помилок; маршрутизатор із неправильним ключем буде ігноруватися іншими учасниками групи.

1.5.3 GLBP

GLBP є пропрієтарним рішенням Cisco з родини FHRP, яке забезпечує захист від відмов маршрутизатора першого переходу та дозволяє балансувати навантаження між резервними маршрутизаторами. Ця функціональність досягається завдяки тому, що всі маршрутизатори групи беруть участь у

пересиланні трафіку. Як і HSRP, GLBP об'єднує кілька пристроїв у групу, яка виступає як єдиний віртуальний пристрій. Протокол може працювати як з IPv4, так і з IPv6.

На відміну від HSRP, GLBP використовує одну віртуальну IP-адресу та кілька віртуальних MAC-адрес; це дозволяє ефективно балансувати навантаження. Учасники групи GLBP взаємодіють через привітальні повідомлення, які надсилаються на багатоадресну адресу 224.0.0.102 через UDP-порт 3222. Таймер привітальних повідомлень дорівнює 3 секундам, а якщо протягом 10 секунд відповідь не отримано, пристрій вважається недоступним.

У GLBP члени групи обирають один маршрутизатор як активний віртуальний шлюз (AVG) для групи. Інші пристрої забезпечують резервування AVG у разі його відмови. Кожен маршрутизатор групи отримує роль активного віртуального передавача (AVF), який пересилає пакети, надіслані на призначену йому віртуальну MAC-адресу, встановлену AVG. Для забезпечення резервування певні AVF перебувають у стані прослуховування, щоб замінити активні пристрої у разі їх недоступності.

AVG у GLBP відповідає за обробку запитів протоколу ARP, пов'язаних із віртуальною IP-адресою. Балансування навантаження досягається тим, що AVG відповідає на такі запити, використовуючи різні віртуальні MAC-адреси. У GLBP допускається використання до чотирьох віртуальних MAC-адрес на групу.

Існують такі методи балансування навантаження:

- Host-Dependent – забезпечує повернення одному хосту тієї ж MAC-адреси AVF, яку було використано раніше. Це дозволяє уникнути розриву трансляції NAT для хоста. Використовується в мережах із NAT;
- Round-Robin – механізм циклічної зміни, за якого AVG по чергово призначає MAC-адреси членам AVF. Цей спосіб є режимом за замовчуванням;
- Weight-Based Round-Robin – метод балансування навантаження, що враховує метрику "Weight". Маршрутизатор із вищою вагою приймає більшу частку навантаження. Метрика "Weight" враховує завантаженість маршрутизатора і визначає, яку частку трафіку він оброблятиме.

Пріоритет шлюзу GLBP визначає роль кожного маршрутизатора та його поведінку у випадку відмови AVG. Принципи налаштування та діапазон значень пріоритету ідентичні HSRP.

GLBP, як і HSRP, підтримує автентифікацію, моніторинг стану портів та режим пріоритетного перевибору.

1.5.2 VRRPv3

VRRPv3 є відкритим стандартом, розробленим на основі протоколу HSRP. Цей протокол підтримується усіма вендорами мережевого обладнання, що забезпечує незалежність від екосистеми конкретного виробника.

- VRRP працює подібно до HSRP, тому детальний опис його функціонування не наводиться. Попри схожість, існують важливі відмінності між цими протоколами:
- VRRP не базується на стеку протоколів TCP/IP і функціонує виключно на мережевому рівні;
- Мультикастова IP-адреса для обміну повідомленнями – 224.0.0.18;
- Таймер привітальних повідомлень за замовчуванням становить 1 секунду, а час простою – 3 секунди;
- Термінологія протоколу відрізняється від HSRP.

VRRPv2 підтримує IPv4 та автентифікацію, яка залежить від реалізації вендора. Cisco, наприклад, забезпечує захист VRRP за допомогою MD5-автентифікації. VRRPv3, на відміну від попередньої версії, додає підтримку IPv6, але не включає функцію автентифікації.

Одним із недоліків як VRRP, так і HSRP є відсутність механізму балансування навантаження. Це призводить до псевдо-балансування, коли лише один пристрій фактично виконує роботу, тоді як інші залишаються у режимі очікування.

1.6 Мережеві топології

Мережеві топології це спосіб організації з'єднання пристроїв у мережі. Вона визначає, як мережеві компоненти з'єднані між собою. Розрізняють фізичну та логічну топології. У той час, як фізична відображає фактичну структуру з'єднання між пристроями, логічна відображає передачу даних між пристроями, присутніми в мережі, незалежно від способу підключення пристроїв. У даній роботі будуть розглянуті лише фізичні топології. Розрізняють такі топології : точка-точка, сіткова топологія, зірка, шинна, кільцева, деревоподібна, гібридна. Більш детально вони будуть розглянуті нижче.

Топологія точка-точка, зображена на рис. 1.11, представляє собою прямий канал зв'язку між двома вузлами або кінцевими точками. Таке з'єднання гарантує, що дані, відправлені з вихідної точки, будуть отримані безпосередньо кінцевою точкою, без проміжних пристроїв або вузлів. У системі «точка-точка» канал зв'язку зарезервовано виключно для з'єднання між двома вузлами. Цей виділений канал забезпечує постійну та надійну швидкість передавання даних, вільну від можливих збоїв, що можуть виникнути під час використання загальних каналів зв'язку, наприклад в топології шина.

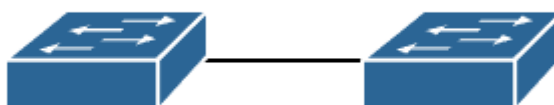


Рис. 1.11

Сітка або один до всіх, зображена на рис.1.12, є варіантом коли кожний пристрій з'єднаний з усіма іншими пристроями. Вона надає підвищену надійність, використовуючи надлишкові з'єднання між пристроями. Проте N пристроїв з'єднаних між собою потребує $N-1$ портів для підключення використовуючи $N(N-1)/2$ з'єднань, що може бути доцільним лише для невеликої кількості пристроїв в

дуже чутливій ділянці мережі, через використання кількості портів на кожному пристрої еквівалентно загальній кількості пристроїв у мережі. Що обмежує мережу мінімальною кількістю портів на пристроях, не враховуючи необхідність підключення кінцевих пристроїв. Такий варіант є нереалізуємим для великої мережі, може використовуватися лише як частина гібридної мережі.

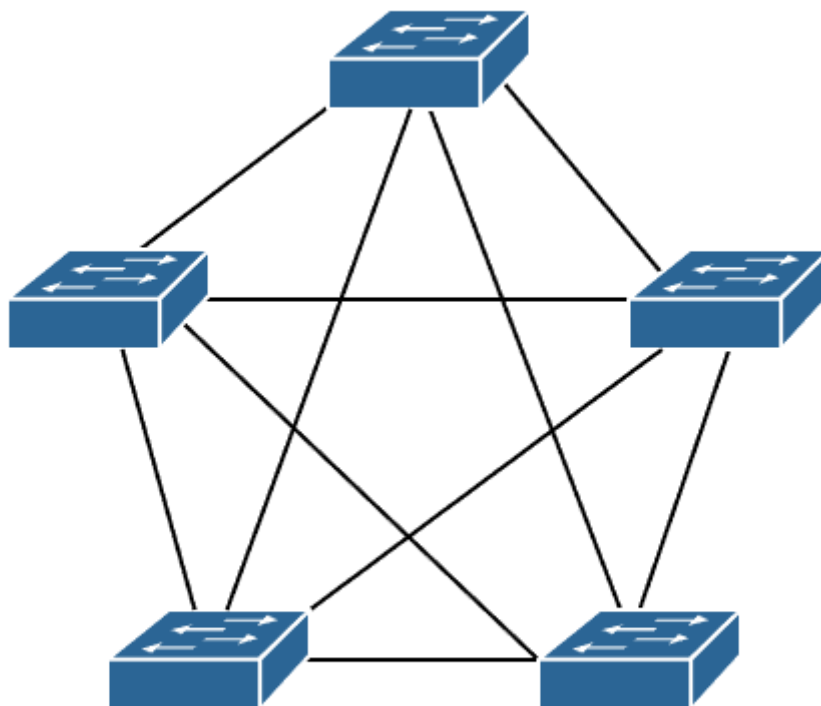


Рис. 1.12

Топологія зірка, зображена на рис.1.13, у даній топології всі пристрої підключаються до одного, утворюючи так звану «зірку». Має переваги в кількості з'єднань: кожному пристрою потрібен лише 1 порт для підключення до концентратора, тому загальна кількість необхідних портів дорівнює N ; Якщо одне з з'єднань вийде з ладу, це вплине тільки на це з'єднання, а не на інші; Проте має вразливий елемент. Якщо центральний вузол, на котрому базується вся топологія, виходить з ладу, вся система виходить з ладу. Продуктивність базується на одному

центральному вузлі, що вимагає підвищеного рівня обчислювальної та пропускної здатності

Поширеним прикладом топології «зірка» є локальна мережа в офісі, де всі комп'ютери підключені до центрального концентратора. Ця топологія також використовується в бездротових мережах, де всі пристрої підключаються до бездротової точки доступу.

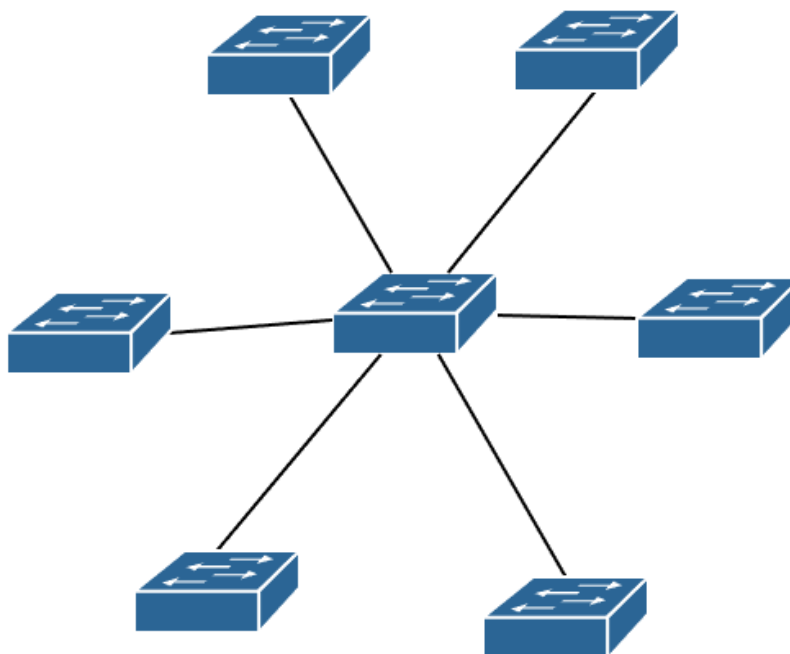


Рис. 1.13

Шинна топологія, зображена на рис. 14, це тип мережі, в якій кожен мережевий пристрій з'єднаний одним кабелем. Має суттєві недоліки, такі як низька продуктивність: через необхідність виділення часу для передачі даних між пристроями, пропускна здатність мережі знижується. Для запобігання колізіям пристрої використовують спеціальні протоколи доступу до мережі, як-от CSMA/CD. Це призводить до додаткових затримок і знижує ефективність передачі, особливо при збільшенні кількості підключених пристроїв; Також присутні проблеми з надійністю: при відмові кабелю або пошкодженні з'єднання мережа повністю втрачає працездатність, оскільки передача даних неможлива.

Зараз шинна топологія більше не використовується в сучасних провідних мережах через її обмеження в масштабованості, продуктивності та надійності. Однак принципи шинної топології частково залишилися в бездротових мережах. В даній роботі розглядатися не буде.

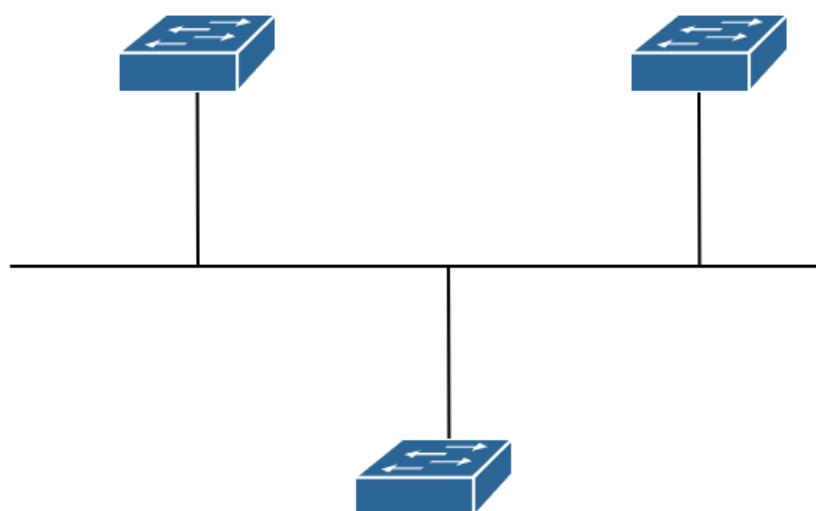


Рис. 1.14

Кільцева топологія, зображена на рис.1.15, представляє умовне кільце утворене пристроями, кожен з яких з'єднаний з двома сусідніми пристроями. Хоча фізична топологія є кільцем, логічна з використанням протоколу STP буде розірваним кільцем, для уникнення утворення бродкаст шторму. Проте така фізична топологія дозволяє без утворення обриву з'єднання пережити втрату одного з'єднання між пристроями.

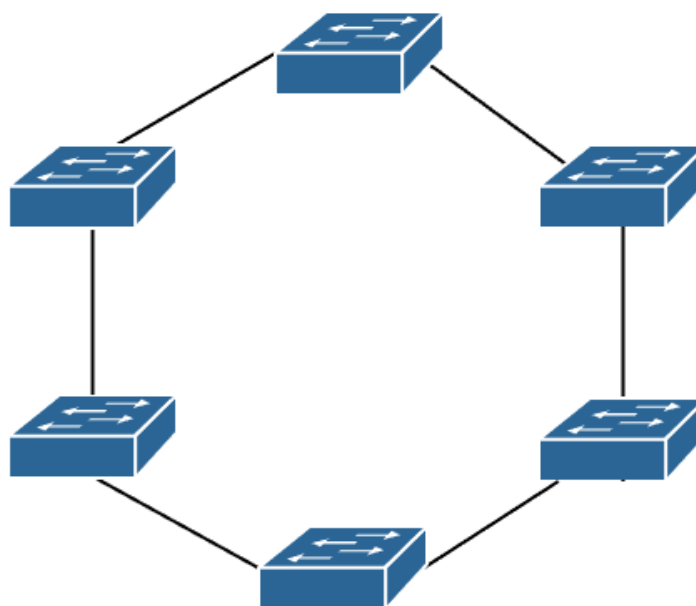


Рис. 1.15

Деревоподібна топологія, зображена на рис.1.15, є різновидністю зіркової топології, котра має ієрархічну структуру. Є ненадійною топологією, якщо виходить з ладу вузол у верхній частині ієрархії топологія руйнується. Не використовується без агрегації та резервування. З переваг добре масштабується.

Деревовидна топологія використовується як принцип для побудови ієрархічної структури мережі, що забезпечує організоване управління трафіком і підвищує масштабованість. Складається з трьох умовних рівнів: рівень ядра, рівень розподілу, рівень доступу.

Ядро призначене для забезпечення швидкого та надійної передачі великих обсягів даних. На рівні ядра зазвичай використовуються високопродуктивні маршрутизатори, які забезпечують обробку з максимальною пропускну здатністю та мінімальною затримкою. Його основне призначення — маршрутизувати трафік між різними мережевими сегментами і забезпечувати з'єднання з зовнішніми мережами.

Рівень розподілу забезпечує фільтрацію, балансування навантаження, розподіл трафіку в різних мережах – VLAN, обробка трафіку за пріоритетністю – QoS, застосування політик. Рівень розподілу застосовує політики доступу, а також

може використовувати протоколи резервування, такі як RSTP, GLBP, щоб гарантувати відмовостійкість .

Рівень доступу забезпечує підключення кінцевих пристроїв до мережі. На цьому рівні здійснюється управління доступом користувачів до мережевих ресурсів. Комутатори рівня доступу можуть підтримувати VLAN для сегментації мережі та забезпечення базових рівнів безпеки, а також здійснювати агрегування портів для підвищення надійності з'єднань.

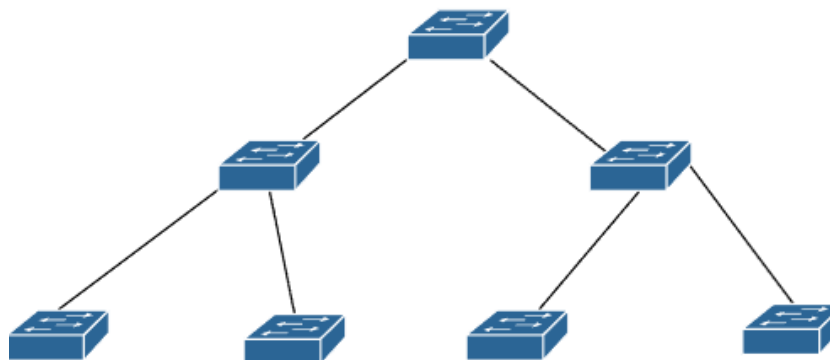


Рис. 1.16

Гібридна топологія, зображена на рис. 1.17 , це комбінація всіх типів топологій, які ми розглянули вище. Гібридна топологія використовується, коли вузли можуть приймати будь-яку форму. Це означає, що це можуть бути окремі вузли, такі як топологія «кільце» або «зірка», або комбінація різних типів топологій, розглянутих вище.

Поширеним прикладом гібридної топології є мережа бізнесу. Мережа може мати магістраль топології кільце, де кожний офіс з'єднаний між собою. У середині кожного офісу може бути стек комутаторів до котрих приєднані кінцеві пристрої та головний маршрутизатор. Ця гібридна топологія забезпечує ефективний зв'язок між різними будівлями, забезпечуючи при цьому гнучкість всередині кожної будівлі.

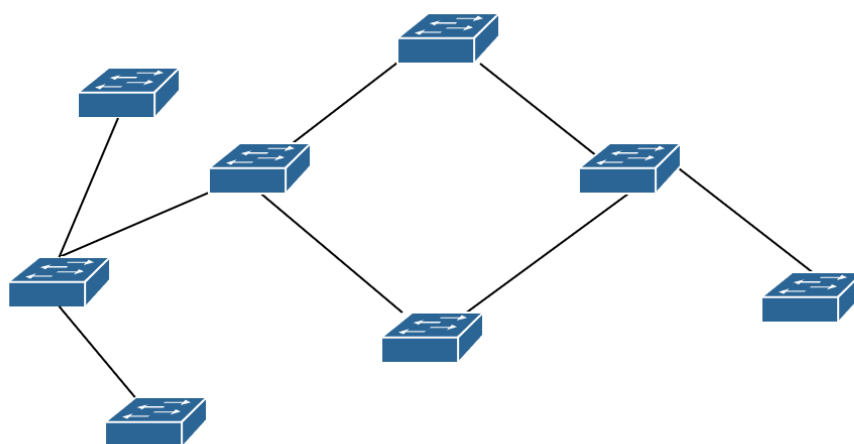


Рис. 1.17

2 ДОСЛІДЖЕННЯ КОМБІНАЦІЙ ТОПОЛОГІЙ І ТЕХНОЛОГІЙ ДЛЯ ПІДВИЩЕННЯ НАДІЙНОСТІ МЕРЕЖІ

2.1 Методологія дослідження

Загальне визначення надійності таке: надійністю називається властивість технічного засобу виконувати задані функції, зберігаючи в часі значення встановлених експлуатаційних показників в заданих межах, що відповідають заданим режимам та умовам використання, технічного обслуговування, збереження і транспортування[13]. В даній роботі не будуть досліджуватись ймовірність відмови певних ланок мережі або фактори, котрі на це впливають. Будуть розглянутий вплив переліку певних топологій на надійність роботи мережі, на основі цього впливу розроблені сценарії тестування цього впливу в різних топологіях в і проведена реалізація і збір даних у розділі 2. Отримані дані буде проаналізовано в Розділі 3.

Першим кроком, в дослідженні впливу комбінації топології та методів агрегування та резервування каналів, є зазначення, які параметри є показниками надійності, та кореляція зміни їх значенням зі збільшенням надійності. Такими параметрами є :

- Доступність
- Втрати пакетів
- Затримка
- Час відновлення після відмови

У такому випадку параметри, такі як втрати пакетів, час затримки та час відновлення, мають обернено пропорційну кореляцію з надійністю. Таким чином низькі значення цих параметрів свідчать про високу надійність системи, тоді як їх підвищення вказує на її зниження.

Далі ми обираємо інструмент котрий будемо використовувати для отримання значень цих параметрів у створених тестових сценаріях. Я обрав програму fping для

збору статистики доступності вузла і збереженням отриманих даних у файл для подальшого їх графічного відображення.

Наступним кроком буде створення тестових мереж, котрі є комбінацією певних топологій з певними технологіями. Будуть використані топології наведені в підпункті 1.6. Шина та гібридна топологія використовуватися не будуть, перша з причини того, що застаріла друга через те, що дана робота зосереджена на конкретних топологіях. Це спростить аналіз кожної топології окремо, дозволяючи зосередитися на їхніх унікальних властивостях і впливі на надійність мережі з використанням вищезазначених технологій.

Для тестування будуть використані технології : агрегування каналів LACP; протокол єднального дерева – MSTP; протокол GLBP для резервування шлюзу за замовчуванням для VLAN, на комутаторах третього рівня моделі OSI.

Вибір LACP обумовлений тим, що він не є пропрієтарним рішенням як PAgP і за його допомогою можна створювати агреговані з'єднання між обладнанням різних виробників. По інших параметрах протоколи управління агрегацією каналів не мають суттєвих відмінностей, котрі могли б вплинути на тестування. Його єдиний недолік це невеликі затримки, при переході в робочий режим, порівняно з статичним налаштуванням.

Вибір протоколу MSTP був обумовлений тим, що він створює низьке навантаження на процесор, має відкритий стандарт, а по швидкості сходження не відрізняється від Rapid PVST, котрий має високе навантаження на процесор, але забезпечує високу гнучкість у налаштування середовищ для кожного VLAN, котра не потребується для нашого тестування.

Протокол GLBP обраний через те, що він може забезпечувати одночасне використання декількох маршрутизаторів, за допомогою балансування навантаження, на відміну від аналогів де є один основний та резервні, котрі знаходяться в режимі очікування, готові приймати трафік, якщо основний маршрутизатор вийде з ладу.

Проте цей протокол використовується у сценаріях, котрі можуть відрізнятися від стандартних топологій, тому потрібно буде зробити деякі

відмінності від стандартних топологій. Протокол не має потреби в прямому підключенні маршрутизаторів, та передбачає підключення кожного пристрою до мережі Інтернет, тобто вони мають бути розташовані на границях локальної мережі, що передбачає наявність між ними цілої купи пристроїв.

Кожна з згаданих технологій працює на своєму рівні моделі OSI і мають свої особливості роботи котрі забезпечують підвищення надійності своєю роботою. На принципах їх роботи і будуть сплановані сценарії дослідження. Через те, що на всі з цих технологій впливає зміна активної топології то буде досліджуватися зміна параметрів надійності від вилучення основних каналів, якщо в топології присутні надлишкові, якщо вони відсутні, то додавання надлишкових каналів та їх вилучення. Дослідження зміни зазначених параметрів надійності від виходу з ладу каналів в розглянутих топологіях, та додавання надлишкових з'єднань.

Буде виконана емуляція різних топологій, з використанням технологій підвищення надійності, за допомогою програмного забезпечення GNS3. Завдяки можливості емулювати образи операційної системи мережевих пристроїв буде відтворена поведінка реального мережевого обладнання, з поправкою на те, що деякі часові показники можуть трохи відрізнятися від реальності через особливості архітектури. Мережеве обладнання для виконання своїх функцій управління мережевим трафіком використовує ASIC чіпи, котрі дозволяють підвищити продуктивність завдяки тому, що вони спроектовані для розв'язання певних задач на відміну від процесорів загального використання, що більш універсальні, проте повільніші.

Проте це не сильно вплине на чинність результатів дослідження, тому що, буде порівнюватися поведінка кожної топології в умовах однакової емуляції. Замість абсолютних чисел затримок і часу відновлення розглядається порівняння цих показників між різними топологіями в рамках однієї платформи - GNS3. Отримані відносні показники вказують на присутні тенденції, тому їх можна екстраполювати на реальні мережі навіть якщо абсолютні значення можуть відрізнятися.

Далі будуть проведені експерименти з використанням розроблених сценаріїв тестування. Для тестування доступні два методи: відключення або додавання певних з'єднань між пристроями та тестування навантаження. Будуть побудовані топології та налаштовані технології резервування та агрегації каналів. Досліджені в різних сценаріях з урахуванням їх особливостей.

Агрегація збільшує надійність з'єднання та пропускну здатність між двома пристроями завдяки використанню більшої кількості портів. Її вплив має рацію тестувати лише у топології точка-точка. Причина цього в тому, що для її тестування потрібно буде перевіряти кожен пару з'єднаних, а результат не буде відрізнятися від топології точка-точка. Вплив на надійність інших топологій можна розрахувати математично.

MSTP запобігає виникненню широкомовних штормів, забезпечує балансування навантаження на мережу розподіляючи трафік між VLAN середовищами, рідше для резервування надлишкових каналів. Протокол STP увімкнений за замовчуванням, вимикати його не рекомендується. Для тестування спроможності резервувати, необхідно не використовувати агрегацію, що не є розповсюдженою практикою. Бо при наявності декількох прямих з'єднань між двома пристроями краще використовувати агрегацію каналів, котрі збільшить пропускну здатність та забезпечить резервування з'єднання.

GLBP резервує шлюз, буде протестована реакція моделі мережі на недоступність основного шлюзу, котрим буде виступати шлюз за замовчуванням на одному з VLAN.

Будуть зібрані дані щодо часу відновлення, втрат пакетів, затримок та інших важливих показників. Наступним кроком буде аналіз всіх даних. Будуть тести будуть проведені декілька разів для отримання середнього значення для кожного тестування. Буде виконане порівняння результатів тестів кожної топології для кожної технології. Цей структурований підхід дозволить всебічно дослідити вплив топологій і технологій агрегації на надійність комп'ютерних мереж.

2.1 Дослідження протоколу LACP

Перша топологія для тестування точка-точка. Загалом технологія агрегування каналів, зокрема її варіації динамічне – LACP PAgP або статичне агрегування каналів, зосереджена лише на ділянці точка-точка. Інформація щодо інших ділянок, окрім тої де вона налаштована вона не отримує і впливати на неї може лише велике навантаження на процесори пристроїв. Проте це не повністю позбавлено сенсу. Ця робота дозволяє перевірити наскільки передбачувано LACP працює у реальних мережесих сценаріях. Та присутня можливість виявити потенційні проблеми при інтеграції LACP у більші мережі.

На рис. 2.1 зображено тестову мережу топологій точка-точка. Тестування полягає у відключенні активних та не активних з'єднань під час активності виконання програми fping з віртуальної машини де налаштована ОС Debian і запущена ця команда. Результат налаштування агрегації каналів можна побачити на рис. 2.2, де зображений стан агрегування на момент тестування.

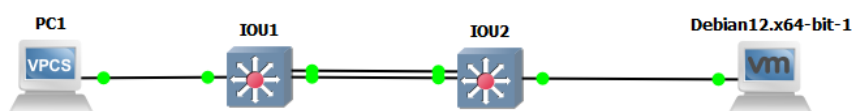


Рисунок 2.1

```

IOU2(config-if-range)#
*Dec  8 12:36:21.128: %LINEPROTO-5-UPDOWN: Line protocol on Interface Port-channel1, changed state to up
IOU2(config-if-range)#do show etherchannel summary
Flags:  D - down        P - bundled in port-channel
         I - stand-alone s - suspended
         H - Hot-standby (LACP only)
         R - Layer3      S - Layer2
         U - in use      f - failed to allocate aggregator

         M - not in use, minimum links not met
         u - unsuitable for bundling
         w - waiting to be aggregated
         d - default port

Number of channel-groups in use: 1
Number of aggregators:          1

Group  Port-channel  Protocol    Ports
-----+-----+-----+-----
1      Po1(SU)        LACP        Et0/0(P)  Et0/1(P)
IOU2(config-if-range)#
  
```

Рис. 2.2

Відключення неактивного каналу ніяк не вплинуло на доступність, коливання затримку доступу до ресурсу викликані вже згаданими властивостями тестового середовища. При відключенні активного лінку відбулася короткочасна втрата пакетів у розмірі приблизно 6 секунд. Протокол Lаср має два варіанти налаштування часу між надсиланням повідомлень 30 та 1 секунда[15]. Канал вважається недоступним через 3 пропущені повідомлення, тож враховуючи отриманий результат з поправкою на тестове середовище, можна вважати, ніякого неочікуваного поведження мережі не виявлено. Результати усіх досліджень для кожної топології будуть наведені в таблиці наприкінці кожного розділу, де досліджується певна технологія.

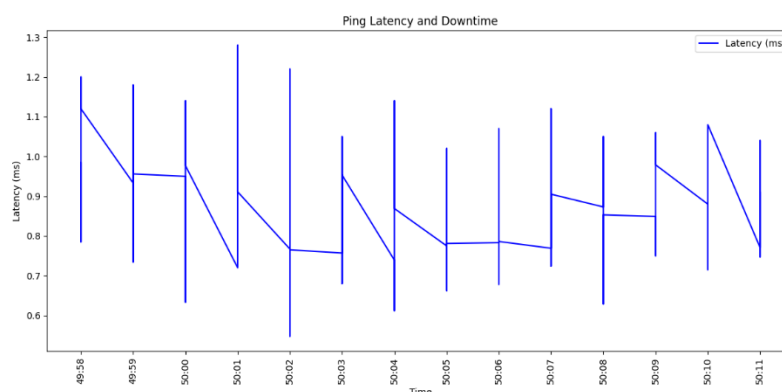


Рис. 2.3

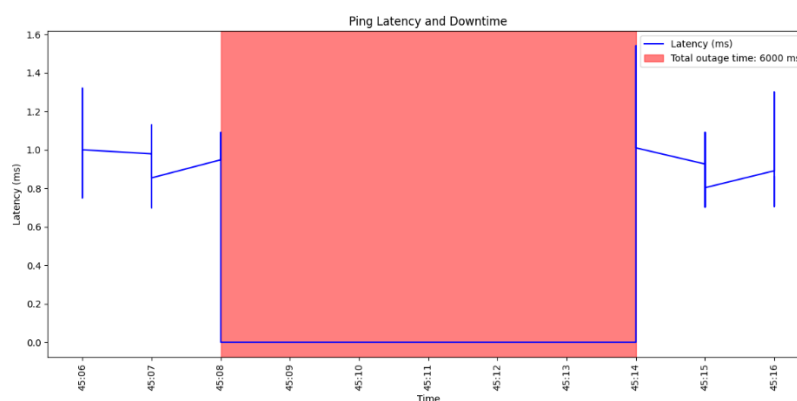


Рис. 2.4

На рис. 2.3 та рис. 2.4 присутнє графічне відображення результатів тестування топології командою `fping`. Наступна топологія для тестування це меш-топологія, її ще називають топологія усі до усіх або сітка. Мною було налаштовано

4 комутатори між котрими було налаштовано 7 агрегованих з'єднань. Результат налаштування зображено на рис. 2.5. Результати тестування не сильно відрізняються від попередньої топології, на що потрібно було звернути увагу. Хоча цього разу був задіяний ще протокол MSTP, котрий керував надлишковими шляхами до пристроїв. Завдяки правильному налаштуванню портів направлених на кінцеві пристрої, та явного налаштування з'єднання між пристроями у режимі повного дуплекса, це необхідно через особливості GNS3, котрий не дозволяє змінити дуплекс режим портів, було досягнута мінімальної різниці між цією топологією та топологією точка-точка. Результати були занесені в табл. №2.1.

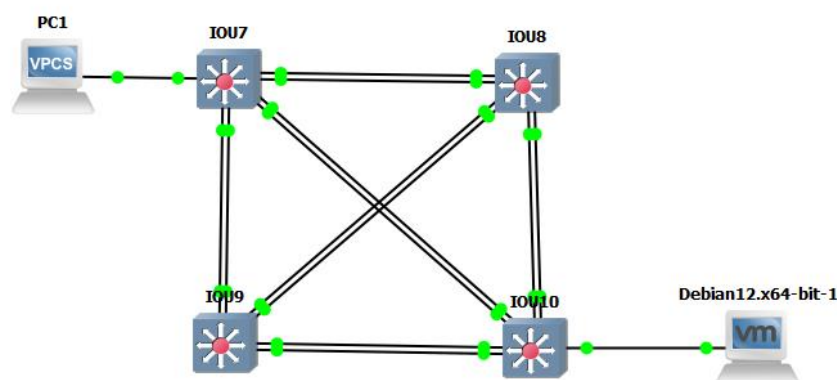


Рис. 2.5

На рис. 2.6 зображена топологія дерево. Хоча було вказано, що тестується всю топологію, на справді активному тестуванню підлягає лише послідовність пристроїв від налаштованої віртуальної машини Debian до вбудованого тестового ПК, котрий відіграє роль шлюзу. Відповідно як і минулих топологіях пристрої, що ближче до шлюзу працюють кореневим комутатором для налаштованого протоколу MSTP, це має сенс згадувати у топологіях зображених на рис. 2.5 та 2.8. Інші топологія використовуються для підтвердження вже зазначеної в документації роботи протоколу LACP, та спробі виявити якісь аномалії, котрі потрібно мати на увазі при побудові мереж. Тестування топології дерево теж показало схожі результати, враховуючи відхилення через тестове середовище. Результати були занесені в табл. 2.1.

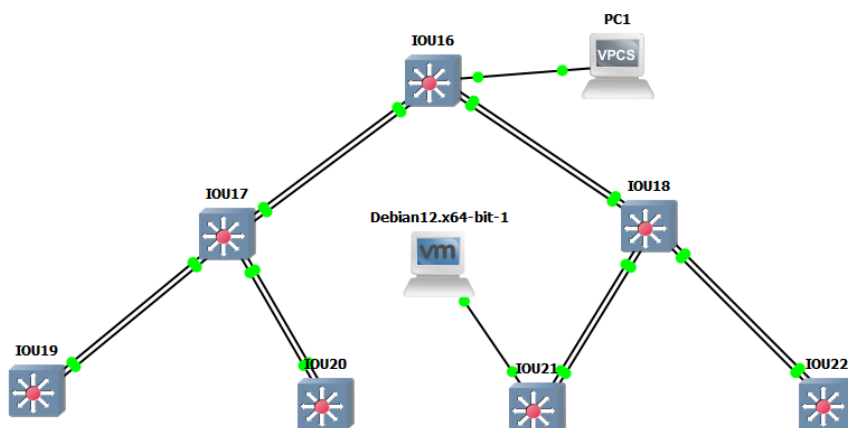


Рис. 2.6

Наступна досліджена топологія зірка зображена на рис. 2.7. Для дослідження була побудована топологія відповідно попереднім екземплярам. Були відключені активні та неактивні канали групи агрегації на шляху виконання тестової програми fping. Результати були занесені в табл. № 2.1.

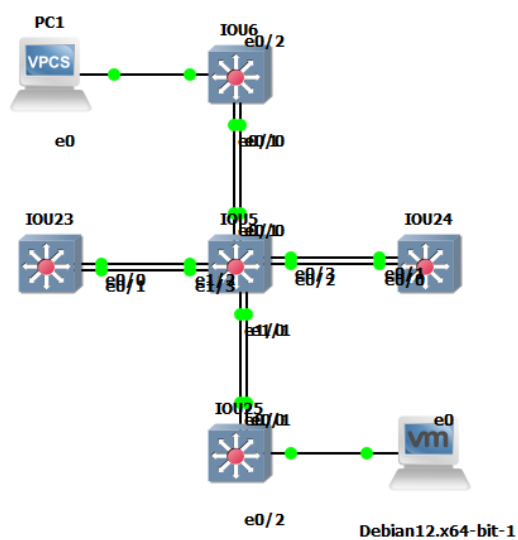


Рис. 2.7

Остання досліджувана топологія – кільце. Як згадувалося раніше працює з протоколом MSTP. Методика дослідження повторилася знов. Результати були занесені в табл. № 2.1.

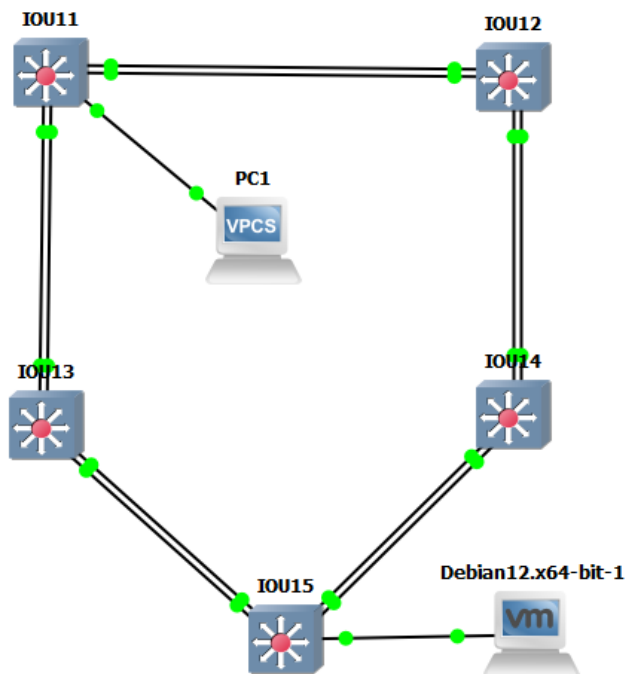


Рис. 2.8

Табл № 2.1

LaCP топологія	експеримент	затримка			втрати	
		мін.	середня	макс.	пакетів	мс
точка точка	втрата акт. з'єднання	0.671	1.06	1.91	60	6000
	втрата не акт. з'єднання	0.547	0.886	1.28	0	0
сітка	втрата акт. з'єднання	0.753	1.15	2.02	42	4200
	втрата не акт. з'єднання	0.584	0.94	1.38	0	0
зірка	втрата акт. з'єднання	0.786	1.22	2.11	50	5000
	втрата не акт. з'єднання	0.629	0.98	1.45	0	0
кільце	втрата акт. з'єднання	0.703	1.18	2.05	47	4700
	втрата не акт. з'єднання	0.593	0.96	1.4	0	0
дерево	втрата акт. з'єднання	0.765	1.24	2.15	52	5200
	втрата не акт. з'єднання	0.610	1.01	1.5	0	0

2.2 Дослідження протоколу MSTP

Перш ніж досліджувати вплив зникнення зв'язку між вузлами у певних топологіях при використанні протоколу MSTP, необхідно зазначити який вплив спричиняє сам протокол. Його основні функції уникнення петель комутації, балансування навантаження. Також може використовуватися для резервування з'єднання через надлишкові канали між пристроями, але такий сценарій відрізняється від реалізації STP коли надлишкові з'єднання неактивні доки в них не з'явиться потреба. При використанні даного протоколу усі з'єднання використовуються завдяки тому що по різних каналах можуть проходити різні інстанси з різним набором VLAN.

На більшості топологій неможливо налаштувати декілька регіонів MSTP, тому для еквівалентності досліджень усюди буде налаштовано лише один регіон. Результатом цього буде те, що все буде працювати як RSTP проте потребуватиме меншої кількості ресурсів. Виходячи з вищенаведеної інформації будуть проведені тестування сходження MSTP протоколу використовуючи два VLAN інстанси під навантаженням, при виходу з ладу певних каналів або утворенні надлишкових з'єднань котрі можуть викликати широкомовний шторм. Результати дослідження кожної топології наведені в табл. 2.2

Першою буде топологія точка-точка. Зображена на рисунку рис. 2.10. Так як протокол спроектований для більш великих та складних топологій, в стандартній топології точка – точка не вийде провести тестування з виходу з ладу каналів, так як з'єднання між пристроями одразу буде перервано, ми протестуємо тільки додавання надлишкових з'єднань та їх відключення.

Для забезпечення швидкої збіжності необхідно використовувати його особливості. А саме portfast на портах, що направлені до кінцевих пристроїв, для усунення цих портів з процесу конвергенції. Явне вказання з'єднань як точка-точка, через особливості роботи середовища тестування – неможливості налаштування режиму дуплексу. UplinkFast вже налаштовано за замовчуванням, і

не потребує додаткової імплементації. На графіку зображеного на рисунку 2.9. відображено результат сходження без цих налаштувань.

На рис. 2.10 зображена топологія точка-точка. Є не самою ідеальною топологією для тестування протоколу MSTP, так як він призначений або для уникнення широкомовних штормів. Для того щоб тестування мало сенс були використанні надлишкові з'єднання.

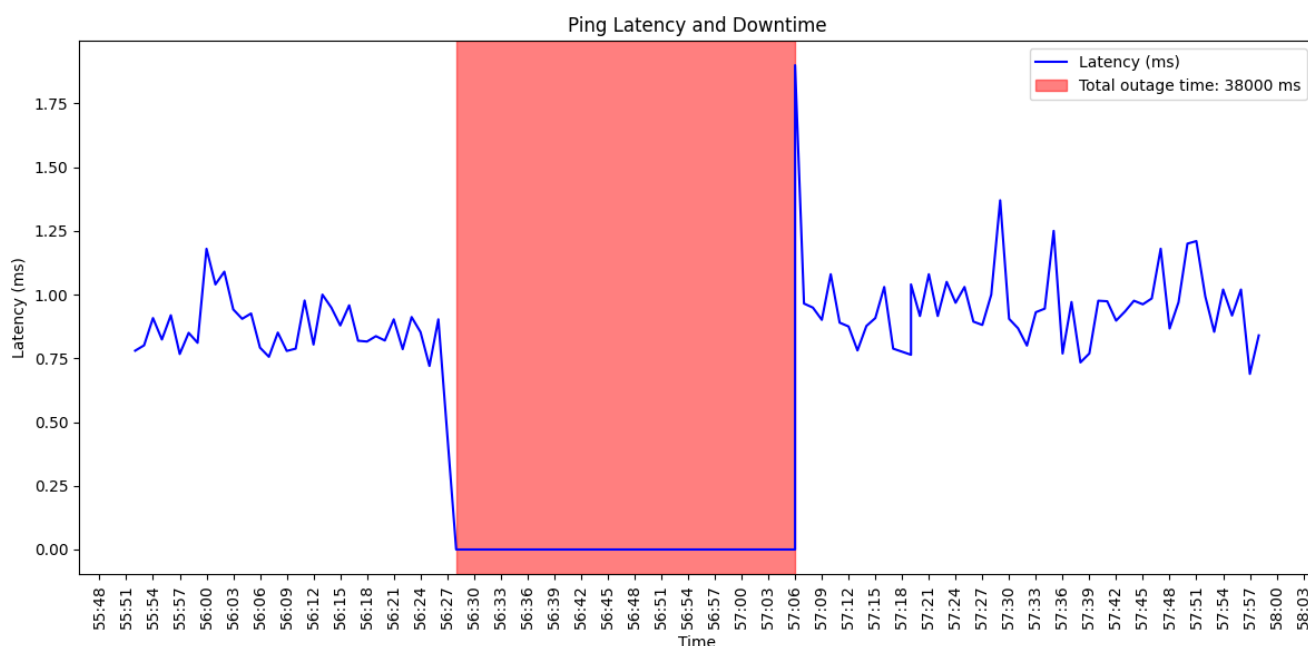


Рис. 2.9

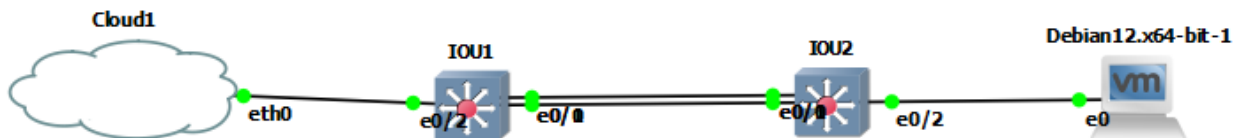


Рис. 2.10

На рис. 2.11 зображено графік, від'єднання неактивного каналу. На рис. 2.12 аналогічне тестування проте вже з активним каналом, під навантаженням команди fping,

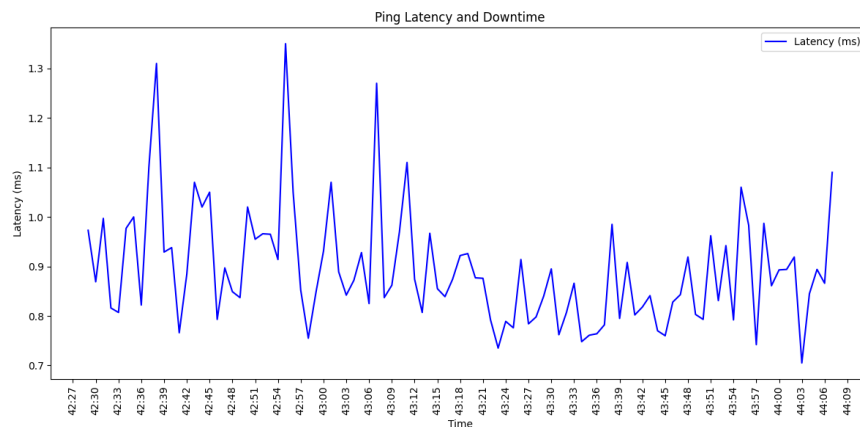


Рис. 2.11

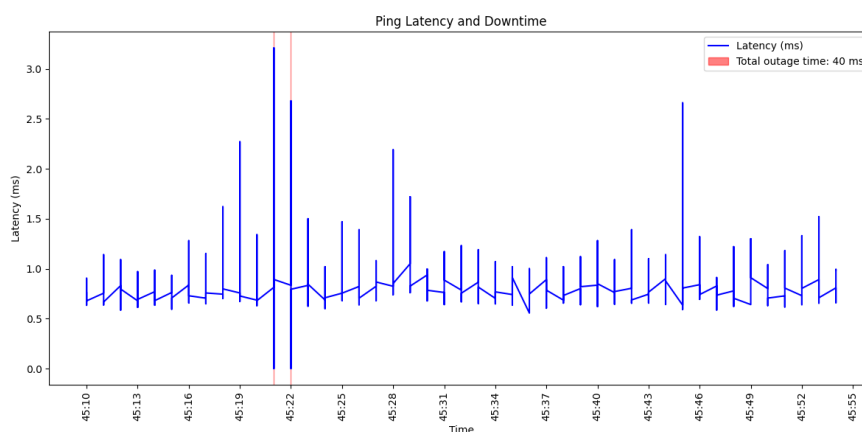


Рис. 2.12

Результат першого тестування показали важливість правильного налаштування. При ньому навіть на навантажених каналах при тестуванні. Переривання не перевищувало 40 мс. Забігаючи наперед у наступних топологіях переривання взагалі не спостерігалися. Результати були занесені в табл. 2.2.

Наступне тестування топології меш або сітка. Стан налаштованих з'єднань, зображено на рис. 2.13. Сама топологія зображена на рис. 2.14. Тестування показало, переривання, незалежно від типу підключення не відбуваються. Зважаючи що GNS3 працює повільніше ніж реальне обладнання, можна очікувати,

що не зважаючи на навантаження, надійність не буде перериватися зміною топології при правильному плануванні та налаштуванні.

```

##### MST0      vlans mapped: 1-4094
Bridge            address aabb.cc00.0400 priority 24576 (24576 sysid 0)
Root              address aabb.cc00.0400 priority 24576 (24576 sysid 0)
Operational       hello time 2 , forward delay 15, max age 20, txholdcount 6
Configured        hello time 2 , forward delay 15, max age 20, max hops 20

Interface      Role Sts Cost      Prio.Nbr Type
-----
Et0/0          Desg FWD 2000000 128.1  P2p
Et0/1          Desg FWD 2000000 128.2  P2p
Et0/2          Desg FWD 2000000 128.3  P2p
Et0/3          Desg FWD 2000000 128.4  P2p Edge
Et1/0          Desg FWD 2000000 128.5  Shr
Et1/1          Desg FWD 2000000 128.6  Shr
Et1/2          Desg FWD 2000000 128.7  Shr
Et1/3          Desg FWD 2000000 128.8  Shr
Et2/0          Desg FWD 2000000 128.9  Shr
Et2/1          Desg FWD 2000000 128.10 Shr
Et2/2          Desg FWD 2000000 128.11 Shr
Et3/0          Desg FWD 2000000 128.12 Shr
Et3/1          Desg FWD 2000000 128.13 Shr
Et3/2          Desg FWD 2000000 128.14 Shr
--More--

##### MST0      vlans mapped: 1-4094
Bridge            address aabb.cc00.0700 priority 32768 (32768 sysid 0)
Root              address aabb.cc00.0400 priority 24576 (24576 sysid 0)
Operational       hello time 2 , forward delay 15, max age 20, txholdcount 6
Configured        hello time 2 , forward delay 15, max age 20, max hops 20

Interface      Role Sts Cost      Prio.Nbr Type
-----
Et0/0          Root FWD 2000000 128.1  P2p
Et0/1          Altn BLK 2000000 128.2  P2p
Et0/2          Altn BLK 2000000 128.3  P2p
Et0/3          Desg FWD 2000000 128.4  P2p Edge
Et1/0          Desg FWD 2000000 128.5  Shr
Et1/1          Desg FWD 2000000 128.6  Shr
Et1/2          Desg FWD 2000000 128.7  Shr
Et1/3          Desg FWD 2000000 128.8  Shr
Et2/0          Desg FWD 2000000 128.9  Shr
Et2/1          Desg FWD 2000000 128.10 Shr
Et2/2          Desg FWD 2000000 128.11 Shr
--More--

##### MST0      vlans mapped: 1-4094
Bridge            address aabb.cc00.0500 priority 32768 (32768 sysid 0)
Root              address aabb.cc00.0400 priority 24576 (24576 sysid 0)
Operational       hello time 2 , forward delay 15, max age 20, txholdcount 6
Configured        hello time 2 , forward delay 15, max age 20, max hops 20

Interface      Role Sts Cost      Prio.Nbr Type
-----
Et0/0          Root FWD 2000000 128.1  P2p
Et0/1          Desg FWD 2000000 128.2  P2p
Et0/2          Desg FWD 2000000 128.3  P2p
Et0/3          Desg FWD 2000000 128.4  P2p Edge
Et1/0          Desg FWD 2000000 128.5  Shr
Et1/1          Desg FWD 2000000 128.6  Shr
Et1/2          Desg FWD 2000000 128.7  Shr
Et1/3          Desg FWD 2000000 128.8  Shr
Et2/0          Desg FWD 2000000 128.9  Shr
Et2/1          Desg FWD 2000000 128.10 Shr
Et2/2          Desg FWD 2000000 128.11 Shr
--More--

##### MST0      vlans mapped: 1-4094
Bridge            address aabb.cc00.0600 priority 32768 (32768 sysid 0)
Root              address aabb.cc00.0400 priority 24576 (24576 sysid 0)
Operational       hello time 2 , forward delay 15, max age 20, txholdcount 6
Configured        hello time 2 , forward delay 15, max age 20, max hops 20

Interface      Role Sts Cost      Prio.Nbr Type
-----
Et0/0          Root FWD 2000000 128.1  P2p
Et0/1          Desg FWD 2000000 128.2  P2p
Et0/2          Altn BLK 2000000 128.3  P2p
Et0/3          Desg FWD 2000000 128.4  P2p Edge
Et1/0          Desg FWD 2000000 128.5  Shr
Et1/1          Desg FWD 2000000 128.6  Shr
Et1/2          Desg FWD 2000000 128.7  Shr
Et1/3          Desg FWD 2000000 128.8  Shr
Et2/0          Desg FWD 2000000 128.9  Shr
Et2/1          Desg FWD 2000000 128.10 Shr
Et2/2          Desg FWD 2000000 128.11 Shr
--More--

```

Рис. 2.13

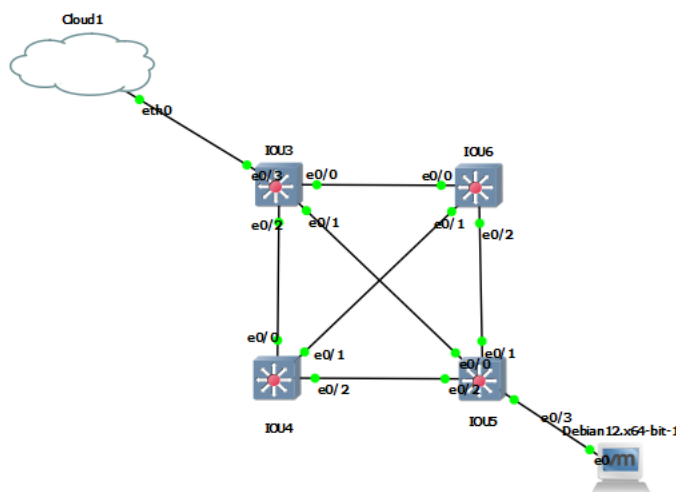


Рис. 2.14

При тестуванні логічна топологія мережі, позначена червоними з'єднаннями на рис 2.15. Переключалася з свого початкового стану в стан зображений, на рис. 2.16. Зміна топології відбувалася при відключенні активного каналу, після відключення каналу, котрий не був задіяний, топологія залишилася без змін. Результати тестування занесені в табл. 2.2. За результатами дослідження переривання не спостерігаються.

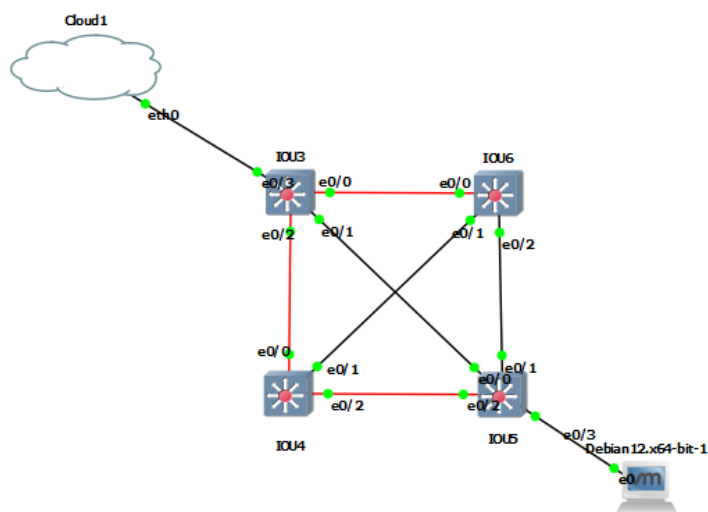


Рис. 2.15

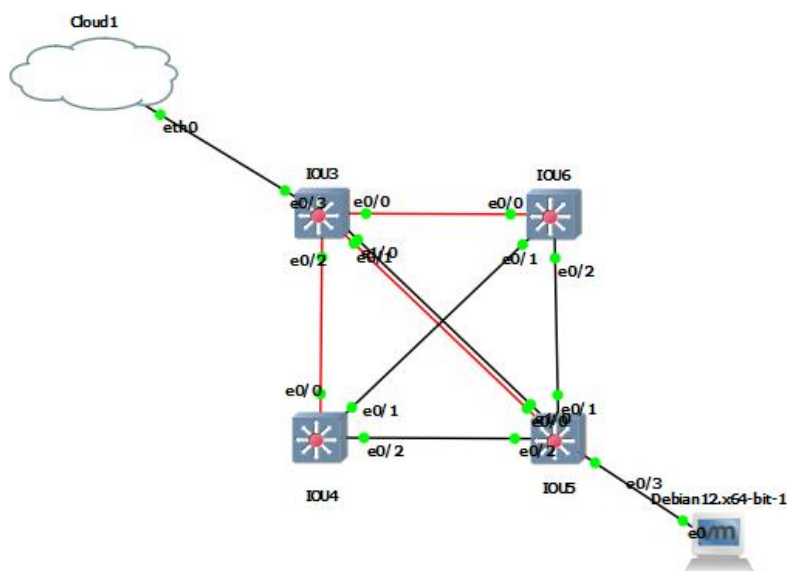


Рис. 2.16

На рис. 2.17 зображена топологія зірка. Так як вона не призначена для прямого тестування протоколу MSTP, необхідно було додати надлишкових з'єднань. Результати дослідження не виявили переривань під час тесту. Результати тестування занесені в табл. 2.2.

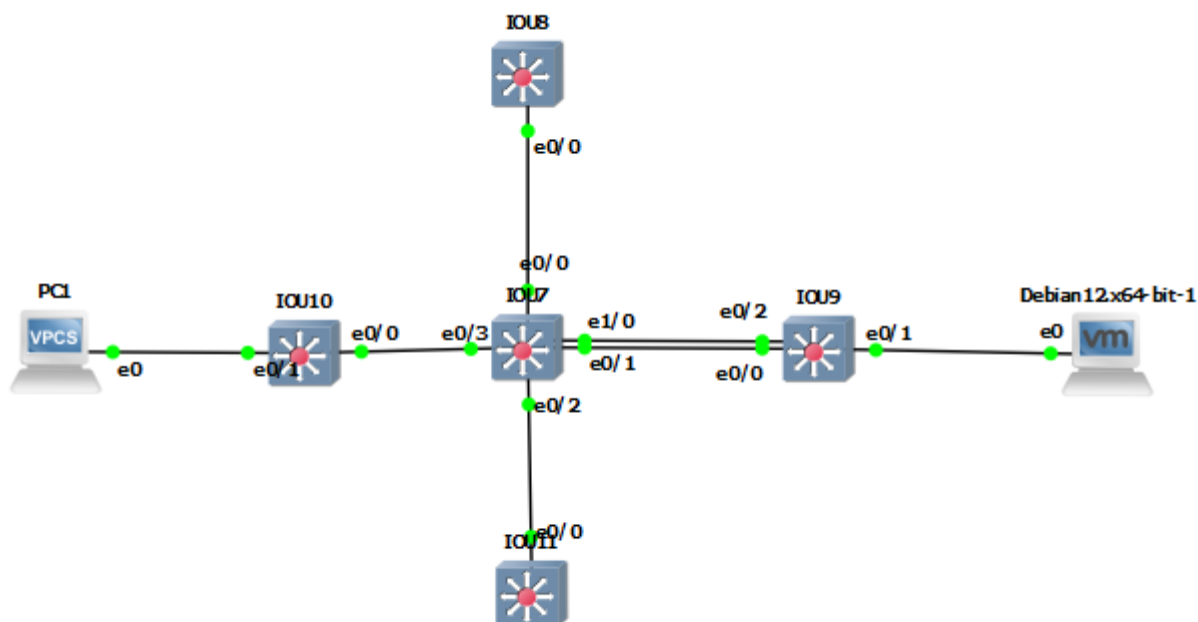


Рис. 2.17

Наступна на черзі тестування топологія дерево. Нічим від інших топологій вона не відрізняється. Необхідності в додаткових з'єднаннях немає. Тестувати можна так. Результати дослідження знову показують відсутність переривань. Результати тестування занесені в табл. 2.2. Топологія зображена на рис. 2.18. Після відключення каналів топологія змінюється з того варіанту як вона показана на рис. 2.18. до рис. 2.19.

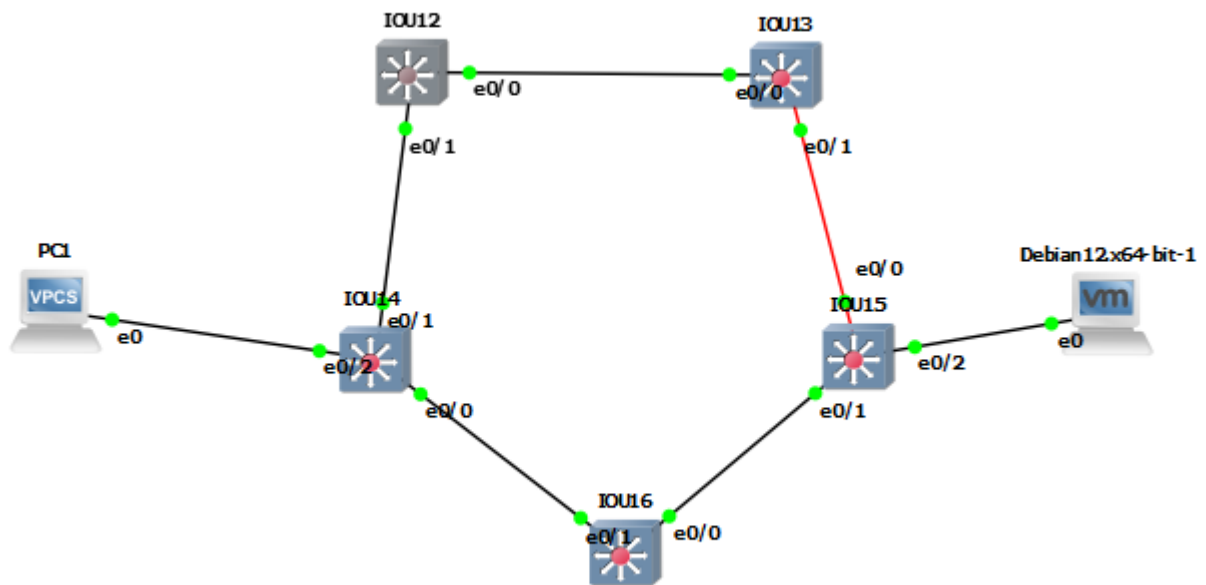


Рис. 2.18

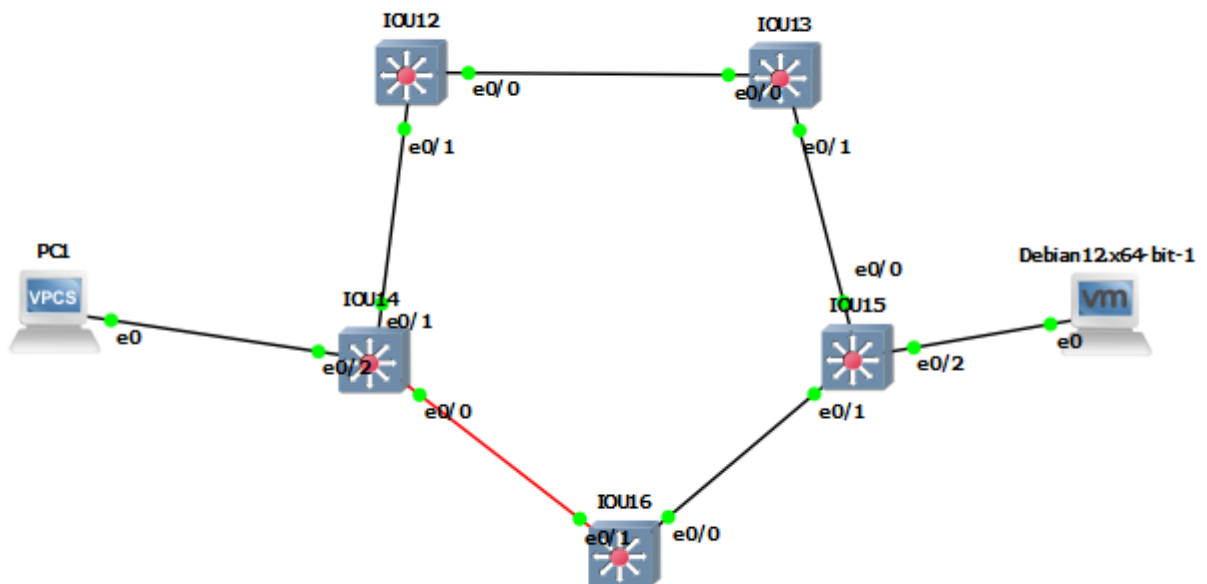


Рис. 2.19

Останньою є топологія дерево. Наперед пишу, що тестування показали, що переривання відсутні. Зображено топологія на рис.2.20. Типологія змінювалася в ланці зображень на рисунку 2.21. Результати занесені до табл. 2.2.

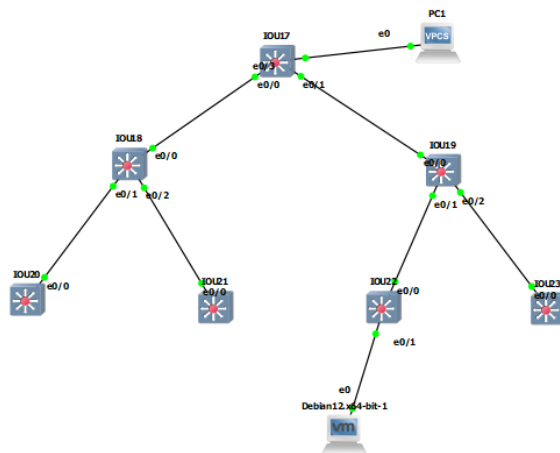


Рис. 2.20

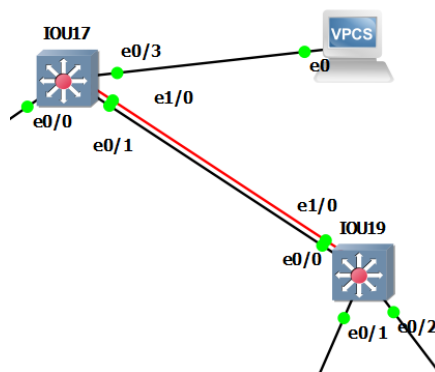


Рис. 2.21

Табл. 2.2

mstp топологія	експеримент	затримка			втрати	
		мін.	середня	макс.	пакетів	мс
точка точка	надлишкове з'єднання	0.556	0.803	3.21	4	40
	втрата акт. з'єднання	0.538	0.796	2.59	0	0
сітка	надлишкове з'єднання	0.627	0.979	3.29	0	0
	втрата акт. з'єднання	0.493	0.784	3.29	0	0
зірка	надлишкове з'єднання	0.709	1.01	3.20	0	0
	втрата акт. з'єднання	0.480	1.10	1.89	0	0
кілеце	надлишкове з'єднання	0.526	1.33	1.23	0	0
	втрата акт. з'єднання	0.493	1.05	1.06	0	0
дерево	надлишкове з'єднання	0.663	1.09	3.06	0	0
	втрата акт. з'єднання	0.713	1.08	1.72	0	0

2.3 Дослідження протоколу GLBP

Як я вже писав раніше, протокол GLBP має свої особливості. Він застосовується лише маршрутизаторах, котрі зазвичай розташовані на межі локальної мережі. Також можливе застосування на комутаторах 3го рівня, для резервування точок маршрутизації трафіку між різними VLAN мережами. Для імітації мережі був обраний звичайний комутатор. Для нормального тестування лише необхідно на ньому налаштувати вимкнений IGMP Snooping. IGMP Snooping — це механізм, який комутатор використовує для оптимізації multicast-трафіку. Якщо на комутаторі увімкнено IGMP Snooping, а жоден із портів не "підписаний" на multicast-групу GLBP, трафік не буде пересилатися. Все інше за налаштуванням GLBP, налаштовуються однакові інстанси з різним пріоритетом для обрання одного активного та одного резервного маршрутизатора. Топологія точка-точка зображено на рис.2.21.

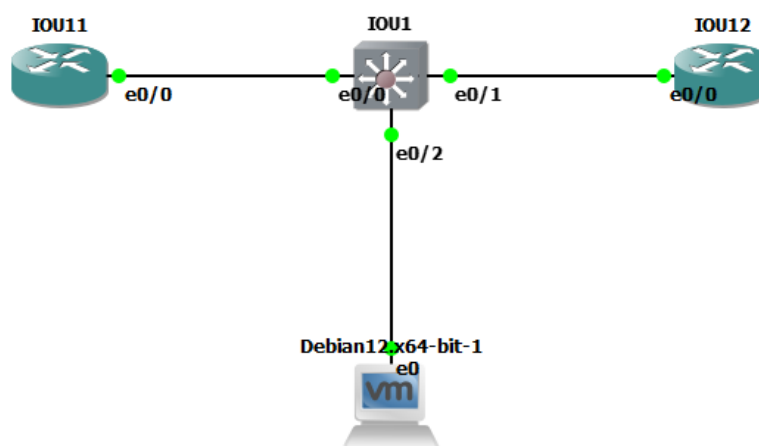


Рис. 2.22

Так як пристрої розташовані на межі локальної мережі, кожна топологія буде набором з'єднаних мТаким чином була протестована кожна топологія Результати дослідження записані в таблиці 2.3. За результатами дослідження переривання також не фіксується при жодному з тестів. Втрата пакетів відповідає налаштованим таймерам. Handling, час після спливу котрого пристрій вважає що учасник групи

GLBP недоступний та сам стає активним шлюзом за замовчуванням. Він налаштований на 18 с. Графік з перериванням зображено на рис. 2.21. Результат на зображенні відповідає очікуванню з поправкою на середовище емулювання.

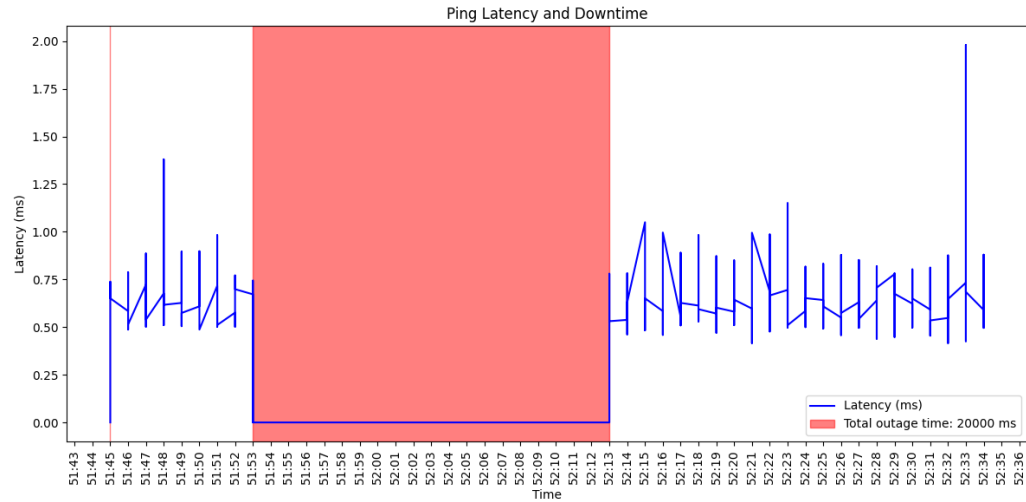


Рис. 2.21

Табл 2.3

glbp топологія	експеримент	затримка			втрати	
		мін.	середня	макс.	пакетів	мс
точка точка	втрата акт. з'єднання	0.416	0.625	1.9	2000	20000
сітка	втрата акт. з'єднання	0.502	0.540	1.10	19200	19200
зірка	втрата акт. з'єднання	0.440	0.722	1.11	20000	20000
кільце	втрата акт. з'єднання	0.475	0.683	1.12	19800	19800
дерево	втрата акт. з'єднання	0.560	0.625	1.13	19300	19300

3 АНАЛІЗ РЕЗУЛЬТАТІВ ДОСЛІДЖЕННЯ

Нижче повторно наведені результати досліджень, а саме таблиці з отриманими значеннями. Це табл. 3.1 – LACP, табл. 3.2 – MSTP, табл. 3.3 – GLBP.

Табл № 3.1

LaCP топологія	експеримент	затримка			втрати	
		мін.	середня	макс.	пакетів	мс
точка точка	втрата акт. з'єднання	0.671	1.06	1.91	60	6000
	втрата не акт. з'єднання	0.547	0.886	1.28	0	0
сітка	втрата акт. з'єднання	0.753	1.15	2.02	42	4200
	втрата не акт. з'єднання	0.584	0.94	1.38	0	0
зірка	втрата акт. з'єднання	0.786	1.22	2.11	50	5000
	втрата не акт. з'єднання	0.629	0.98	1.45	0	0
кіЛЬце	втрата акт. з'єднання	0.703	1.18	2.05	47	4700
	втрата не акт. з'єднання	0.593	0.96	1.4	0	0
дерево	втрата акт. з'єднання	0.765	1.24	2.15	52	5200
	втрата не акт. з'єднання	0.610	1.01	1.5	0	0

Табл. 3.2

mstp топологія	експеримент	затримка			втрати	
		мін.	середня	макс.	пакетів	мс
точка точка	надлишкове з'єднання	0.556	0.803	3.21	4	40
	втрата акт. з'єднання	0.538	0.796	2.59	0	0
сітка	надлишкове з'єднання	0.627	0.979	3.29	0	0
	втрата акт. з'єднання	0.493	0.784	3.29	0	0
зірка	надлишкове з'єднання	0.709	1.01	3.20	0	0
	втрата акт. з'єднання	0.480	1.10	1.89	0	0
кіЛЬце	надлишкове з'єднання	0.526	1.33	1.23	0	0
	втрата акт. з'єднання	0.493	1.05	1.06	0	0
дерево	надлишкове з'єднання	0.663	1.09	3.06	0	0
	втрата акт. з'єднання	0.713	1.08	1.72	0	0

Табл 3.3

glbp	експеримент	затримка			втрати	
топологія		мін.	середня	макс.	пакетів	мс
точка точка	втрата акт. з'єднання	0.416	0.625	1.9	2000	20000
сітка	втрата акт. з'єднання	0.502	0.540	1.10	19200	19200
зірка	втрата акт. з'єднання	0.440	0.722	1.11	20000	20000
кільце	втрата акт. з'єднання	0.475	0.683	1.12	19800	19800
дерево	втрата акт. з'єднання	0.560	0.625	1.13	19300	19300

В ході проведених експериментів було досліджено вплив різних топологій мережі на роботу технологій LACP, MSTP та GLBP. Тестування виконувалося у емуляційному середовищі GNS3, що дозволило імітувати умови реальної роботи мережі. Основним завданням було оцінити, чи змінюють мережеві топології (точка-точка, зірка, кільце, дерево, сітка) ключові показники надійності роботи протоколів.

Під час дослідження були зібрані дані щодо таких метрик, як: доступність мережі, втрати пакетів, затримка, час відновлення після відмови.

Отримані результати показали, що тип топології мережі не має істотного впливу на роботу протоколів LACP, MSTP та GLBP. Всі протоколи демонструють стабільну поведінку та приблизно однакові показники метрик незалежно від вибраної топології. Відповідальність за похибку можна покласти на середовище емуляції. Кореляції між збільшенням обладнання та збільшенням затримки не виявлено

LACP забезпечував об'єднання каналів у єдиний логічний інтерфейс, а також автоматичне переключення у випадку відмови одного з каналів відповідно до своїх таймерів надсилання повідомлень, котрі були налаштовані попередньо. Це спостерігалось у всіх топологіях. MSTP зберігав стабільну конвергенцію у разі зміни топології або відмови вузлів при наявності відповідних налаштувань. Затримки у процесі відновлення були однаковими для кожної топології, окрім топології точка-точка, аномалія скоріш за все є особливістю середовища

тестування. GLBP забезпечував швидке переключення у разі відмови активного маршрутизатора, незалежно від топології, у якій був реалізований.

На перший погляд, результати можуть здатися незадовільними, оскільки очікувалося, що певні типи топологій, наприклад, сітка чи кільце, мали певний вплив на затримку. Однак дані підтверджують, що всі досліджувані протоколи побудовані таким чином, щоб мінімізувати вплив топології на їхню ефективність. Це є свідченням високої адаптивності сучасних мережевих протоколів до різних архітектур мережі.

Більше того, відсутність впливу топології на результати дозволяє зробити важливий практичний висновок: вибір топології мережі для використання LACP, MSTP чи GLBP має базуватися більше на вимогах до масштабованості, простоти адміністрування та фізичної реалізації, аніж на її впливі на технології.

Отримані результати підтверджують гіпотезу про незалежність сучасних технологій резервування та агрегації каналів від типу топології. Це вказує на еволюцію мережевих протоколів та їхню універсальність, що є важливим внеском у розуміння принципів їх роботи. Незважаючи на негативні результати дослідження, воно є корисним. Для проектування мереж, які використовують LACP, MSTP та GLBP, можна вільно вибирати топологію, орієнтуючись на вимоги до масштабованості та адміністрування, а не на очікування покращення мережевих метрик. Так як у даній роботі не виявлено впливу топологій на протоколи, варто розширити дослідження для аналізу поєднання різних топологій для проектування мереж та дослідження протоколів у середовищах із підвищеним навантаженням або обмеженнями в пропускній здатності.

ВИСНОВКИ

Забезпечення надійності функціонування комп'ютерних мереж є однією з ключових вимог до сучасних інформаційних систем, що використовуються як у корпоративному середовищі, так і серед звичайних користувачів. У цій роботі було досліджено вплив різних типів топологій мереж на роботу технологій резервування та агрегації каналів, зокрема протоколів LACP, MSTP та GLBP.

У процесі дослідження було налаштовано та протестовано мережі різних топологій: зірка, кільце, дерево, сітка, точка-точка. Використання протоколу LACP дозволило забезпечити агрегацію каналів, що підвищує пропускну здатність і надійність з'єднань. Протокол MSTP забезпечив усунення петель у топології та ефективне перенаправлення трафіку у разі виникнення аварійних ситуацій. GLBP забезпечив резервування маршрутизаторів та балансування навантаження між ними.

Результати тестувань показали, що топологія мережі не має суттєвого впливу на роботу зазначених технологій. Метрики надійності, такі як доступність, час відновлення після аварії, затримка та втрати пакетів, залишались стабільними незалежно від топології за умови коректного налаштування технологій. Це підтверджує універсальність і адаптивність цих технологій до різних мережевих структур.

Робота підтвердила, що протоколи резервування та агрегації каналів можуть ефективно забезпечувати надійність мережі незалежно від її топології. Отримані результати та рекомендації можуть бути корисними для інженерів та проектувальників мереж при плануванні та налаштуванні інфраструктури з високими вимогами до надійності. Дослідження надало практичний досвід налаштування технологій LACP, MSTP та GLBP, що може бути застосовано у реальних корпоративних та провайдерських мережах.

Подальше дослідження може бути зосереджене на аналізі змішаних топологій, характерних для провайдерських мереж, таких як гібриди зірки, кільця та сітки. Вивчення впливу додаткових факторів, таких як масштабність мережі,

високе навантаження або збільшена кількість вузлів, на ефективність роботи протоколів. Дослідження роботи цих протоколів у середовищах з підвищеною кількістю відмов або складними схемами маршрутизації.

ПЕРЕЛІК ПОСИЛАНЬ

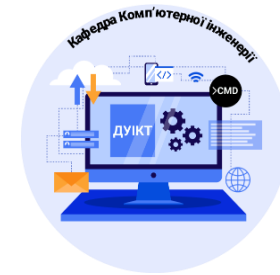
1. Understand EtherChannel Load Balance and Redundancy on Catalyst Switches. URL: <https://www.cisco.com/c/en/us/support/docs/lan-switching/etherchannel/12023-4.html>
2. The Scaling Limitations of Etherchannel Or Why 1+1 Not Equal 2. URL: <https://packetpushers.net/the-scaling-limitations-of-etherchannel-or-why-11-does-not-equal-2>
3. Media Access Control Bridges and Virtual Bridged Local Area Networks - IEEE802. URL: https://www.ieee802.org/802_tutorials/2013-03/8021-IETF-tutorial-final.pdf
4. 802.1D-2004. URL: https://kupdf.net/download/802-1d-2004-pdf_5900a26edc0d60354f959e8d_pdf
5. Understand Rapid Spanning Tree Protocol. URL: <https://www.cisco.com/c/en/us/support/docs/lan-switching/spanning-tree-protocol/24062-146.html>
6. Understanding and Configuring the Cisco UplinkFast Feature URL: https://www-cisco-com.translate.goog/c/en/us/support/docs/lan-switching/spanning-tree-protocol/10575-51.html?_x_tr_sl=auto&_x_tr_tl=ru&_x_tr_hl
7. Understanding and Configuring the Cisco UplinkFast Feature. URL: <https://www.cisco.com/c/en/us/support/docs/lan-switching/spanning-tree-protocol/10575-51.html> –
8. Understand the Multiple Spanning Tree/ URL: <https://www.cisco.com/c/en/us/support/docs/lan-switching/spanning-tree-protocol/24248-147.html>
9. Understand the Hot Standby Router Protocol Features and Functionality. URL: <https://www.cisco.com/c/en/us/support/docs/ip/hot-standby-router-protocol-hsrp/9234-hsrpguidetoc.html>
10. Hot Standby Router Protocol (HSRP): Frequently Asked Questions. URL: <https://www.cisco.com/c/en/us/support/docs/ip/hot-standby-router-protocol-hsrp/9281-3.html>
11. GLBP - Gateway Load Balancing Protocol. URL: https://www.cisco.com/en/US/docs/ios/12_2t/12_2t15/feature/guide/ft_glbp.html
12. CISCO lacp link aggregation control protocol URL: <https://studyccnp.com/cisco-lacp-link-aggregation-control-protocol/>
13. ДСТУ 2860-94 Надійність техніки. Терміни та визначення. URL: https://online.budstandart.com/ua/catalog/doc-page.html?id_doc=25034

ДЕМОНСТРАЦІЙНІ МАТЕРІАЛИ



ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО -
КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ

НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ
ТЕХНОЛОГІЙ



КАФЕДРА КОМП'ЮТЕРНОЇ ІНЖЕНЕРІЇ

Магістерська робота

«Комбінації топологій для надійності функціонування комп'ютерних
мереж»

Виконав: студент групи КСДМ-61 Губа Антон Ігорович

Керівник: к.т.н., доц., доцент кафедри КІ Вєчерковська А.С.

Київ - 2025

Мета, об'єкт та предмет дослідження

Мета роботи: виявити вплив топологій мереж на роботу протоколів агрегації каналів та резервування

Об'єкт дослідження: Процеси та технології резервування і агрегації каналів у комп'ютерних мережах для забезпечення їхньої надійності.

Предмет дослідження: Вплив топологій мережі на ефективність функціонування технологій резервування (GLBP, MSTP) і агрегації каналів (LACP) у забезпеченні ключових показників надійності, таких як доступність, час відновлення, втрати пакетів і затримка.

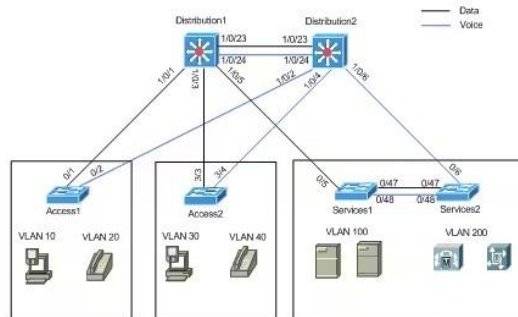
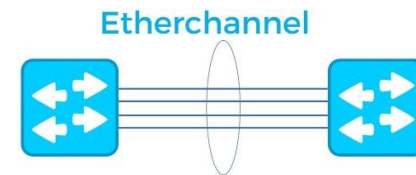
Актуальність роботи

Значущість роботи полягає у розробці рекомендацій для інженерів та проектувальників мереж щодо вибору протоколів та топологій при побудові мереж із високими показниками надійності. Отримані результати також сприятимуть кращому розумінню поведінки сучасних протоколів у реальних умовах.

Робота присвячена дослідженню впливу мережевих топологій на роботу протоколів резервування та агрегації каналів, таких як LACP, MSTP і GLBP. В умовах сучасних вимог до безперервності та надійності мережевих підключень, особливо для корпоративних мереж та провайдерів, тема роботи є актуальною. Використання таких технологій дозволяє підвищити доступність мережі, зменшити час відновлення після збоїв і забезпечити високу якість передачі даних.

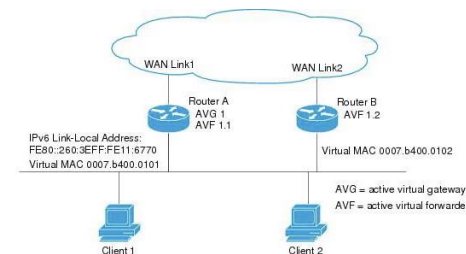
Протоколи LACP, MSTP, GLBP

LACP – технологія за допомогою якої можна декілька фізичних каналів об'єднати в один логічний. Використовується для об'єднання каналів між мережевим обладнанням



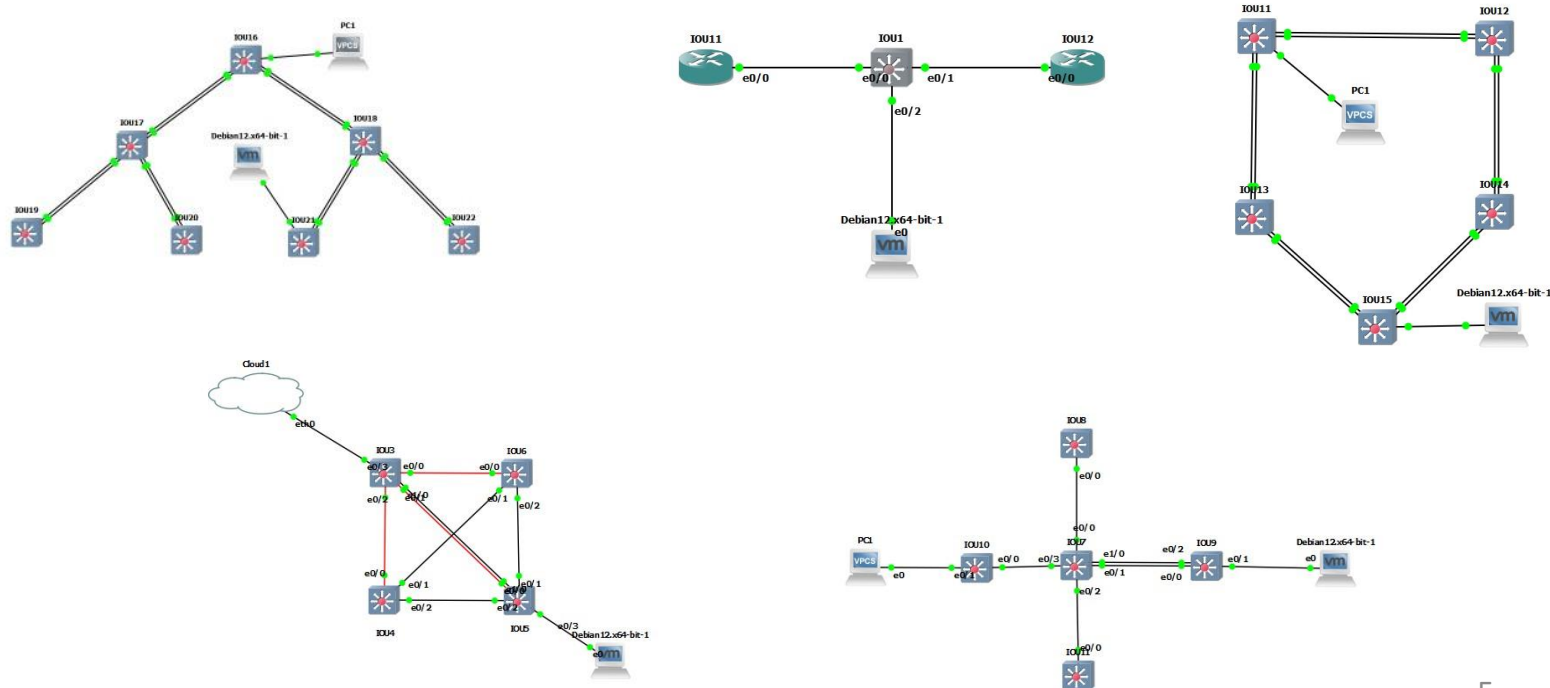
MSTP – варіація протоколу STP для забезпечення відсутності широкомовних штормів та балансування навантаження між VLAN за допомогою деактивації надлишкових з'єднань

GLBP – відкритий протокол резервування шлюзу за замовчуванням, котрий дозволяє балансувати навантаження між пристроями у групі



Створенні та налаштуванні тестові топології

У ПЗ для емуляції мереж були побудовані топології: точка-точка, сітка, кільце, зірка, дерево. Та налаштовані на них протоколи агрегування каналів LACP, протокол єднального дерева MSTP, протокол резервування шлюзу за замовчуванням GLBP.



Проведені експерименти

Над побудованими мережами були проведені експерименти для отримання метрик надійності мережі :доступність, втрати пакетів, затримка, час відновлення після відмови. Експерименти включали додавання та відключення надлишкових каналів та замір описаних параметрів надійності за допомогою `fring` і збереженням отриманих даних у файл.

Результати тестування протоколу LACP

LaCP	експеримент	затримка			втрати	
топологія		мін.	середня	макс.	пакетів	мс
точка точка	втрата акт. з'єднання	0.671	1.06	1.91	60	6000
	втрата не акт. з'єднання	0.547	0.886	1.28	0	0
сітка	втрата акт. з'єднання	0.753	1.15	2.02	42	4200
	втрата не акт. з'єднання	0.584	0.94	1.38	0	0
зірка	втрата акт. з'єднання	0.786	1.22	2.11	50	5000
	втрата не акт. з'єднання	0.629	0.98	1.45	0	0
кільце	втрата акт. з'єднання	0.703	1.18	2.05	47	4700
	втрата не акт. з'єднання	0.593	0.96	1.4	0	0
дерево	втрата акт. з'єднання	0.765	1.24	2.15	52	5200
	втрата не акт. з'єднання	0.610	1.01	1.5	0	0

Результати тестування протоколу MSTP

mstp		затримка			втрати	
топологія	експеримент	мін.	середня	макс.	пакетів	мс
точка точка	надлишкове з'єднання	0.556	0.803	3.21	4	40
	втрата акт. з'єднання	0.538	0.796	2.59	0	0
сітка	надлишкове з'єднання	0.627	0.979	3.29	0	0
	втрата акт. з'єднання	0.493	0.784	3.29	0	0
зірка	надлишкове з'єднання	0.709	1.01	3.20	0	0
	втрата акт. з'єднання	0.480	1.10	1.89	0	0
кільце	надлишкове з'єднання	0.526	1.33	1.23	0	0
	втрата акт. з'єднання	0.493	1.05	1.06	0	0
дерево	надлишкове з'єднання	0.663	1.09	3.06	0	0
	втрата акт. з'єднання	0.713	1.08	1.72	0	0

Результати тестування протоколу GLBP

glbp	експеримент	затримка			втрати	
топологія		мін.	середня	макс.	пакетів	мс
точка точка	втрата акт. з'єднання	0.416	0.625	1.9	2000	20000
сітка	втрата акт. з'єднання	0.502	0.540	1.10	19200	19200
зірка	втрата акт. з'єднання	0.440	0.722	1.11	20000	20000
кільце	втрата акт. з'єднання	0.475	0.683	1.12	19800	19800
дерево	втрата акт. з'єднання	0.560	0.625	1.13	19300	19300

Аналіз результатів

Результати тестувань показали, що топологія мережі не має суттєвого впливу на роботу зазначених технологій. Метрики надійності, такі як доступність, час відновлення після аварії, затримка та втрати пакетів, залишались стабільними незалежно від топології за умови коректного налаштування технологій. Це підтверджує універсальність і адаптивність цих технологій до різних мережевих структур.

Висновки

У цій роботі було досліджено вплив різних типів топологій мереж на роботу технологій резервування та агрегації каналів, зокрема протоколів LACP, MSTP та GLBP.

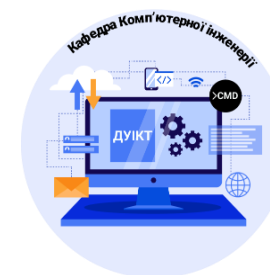
Результати тестувань показали, що топологія мережі не має суттєвого впливу на роботу зазначених технологій. Метрики надійності, такі як доступність, час відновлення після аварії, затримка та втрати пакетів, залишались стабільними незалежно від топології за умови коректного налаштування технологій. Це підтверджує універсальність і адаптивність цих технологій до різних мережевих структур.

Отримані результати та рекомендації можуть бути корисними для інженерів та проектувальників мереж при плануванні та налаштуванні інфраструктури з високими вимогами до надійності.



**ДЕРЖАВНИЙ УНІВЕРСИТЕТ ІНФОРМАЦІЙНО -
КОМУНІКАЦІЙНИХ ТЕХНОЛОГІЙ**

**НАВЧАЛЬНО-НАУКОВИЙ ІНСТИТУТ ІНФОРМАЦІЙНИХ
ТЕХНОЛОГІЙ**



КАФЕДРА КОМП'ЮТЕРНОЇ ІНЖЕНЕРІЇ

Магістерська робота

**«Комбінації топологій для надійності функціонування комп'ютерних
мереж»**

Виконав: студент групи КСДМ-62 Губа Антон Ігорович

Керівник: к.т.н., доц., доцент кафедри КІ Вєчерковська А.С.,.

Київ - 2025