

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

**Навчально-науковий Інститут телекомунікаційних систем
Кафедра інформаційно-комунікаційних технологій та систем**

«До захисту допущено»

ВО завідувача кафедри

_____ **МАРІЯ СКУЛИШ**

«__» _____ 20__ р.

**Дипломна робота
на здобуття ступеня бакалавра
зі спеціальності 172 Телекомунікації та радіотехніка
на тему: «Аналіз протоколів багатоадресної маршрутизації трафіку
групового мовлення в IP-мережах»**

Виконав:

студент IV курсу, групи ТС-91

Кравченко Ілля Миколайович _____

Керівник:

Старший викладач кафедри ІКТС, к.т.н.

Новіков В. І. _____

Рецензент:

Професор кафедри телекомунікацій, д.т.н., професор

Лисенко О.І. _____

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

Студент _____

Київ – 2023 року

Національний технічний університет України
«Київський політехнічний інститут імені Ігоря Сікорського»
Навчально-науковий Інститут телекомунікаційних систем
Кафедра інформаційно-комунікаційних технологій та систем

Рівень вищої освіти – перший (бакалаврський)
Спеціальність – 172 Телекомунікації та радіотехніка
Освітня програма – «Телекомунікаційні системи та мережі»

ЗАТВЕРДЖУЮ

ВО завідувача кафедри

_____ МАРІЯ СКУЛИШ

« ____ » _____ 20__ р.

ЗАВДАННЯ

на дипломну роботу студенту

Кравченку Іллі Миколайовичу

1. Тема роботи «Аналіз протоколів багатоадресної маршрутизації трафіку групового мовлення в IP-мережах», керівник роботи Новіков Валерій Іванович, старший викладач кафедри ІКТС, к.т.н., затверджено наказом по університету від «22» травня 2023 р. № 1884-с
2. Термін подання студентом роботи 10 червня 2023 року.
3. Вихідні дані до роботи: Інформаційні матеріали щодо аналізу протоколів багатоадресної маршрутизації. Схема мережі для проектування.
4. Зміст роботи:
 - 1) Визначення багатоадресної маршрутизації трафіку групового мовлення в IP-мережах.
 - 2) Порівняння одноадресної, багатоадресної та широкомовної маршрутизації.
 - 3) Аналіз протоколу IGMP (Internet Group Management Protocol).
 - 4) Аналіз протоколу PIM (Protocol Independent Multicast) та режимів його роботи.

- 5) Проектування мережі багатоадресного трафіку групового мовлення.
- 6) Порівняльний аналіз побудованої мережі на різних видах налаштувань багатоадресної маршрутизації

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

Презентація обсягом 9 слайдів.

6. Дата видачі завдання 10 квітня 2023

Календарний план

№ з/п	Назва етапів виконання дипломної роботи	Термін виконання етапів роботи	Примітка
1	Ознайомлення з науково-технічною літературою.	11.04.2023 – 19.04.2023	
2	Визначення об'єкту, предмету, мети.	19.04.2023	
3	Підготовка першого розділу «Аналіз способів застосування багатоадресної маршрутизації трафіку групового мовлення в IP-мережах».	20.04.2023– 31.04.2023	
4	Підготовка другого розділу «Ігряд протоколів багатоадресної маршрутизації трафіку групового мовлення в IP-мережах».	31.04.2023– 10.05.2023	
5	Підготовка третього розділу «Порівняльна характеристика побудованої мережі багатоадресного трафіку групового мовлення за допомогою протоколу PIM».	10.05.2023 – 20.05.2023	
6	Підготовка загального висновку.	20.05.2023 - 21.05.2023	
7	Остаточне оформлення роботи	22.05.2023	
8	Виготовлення ілюстративного матеріалу	22.05.2023 - 05.06.2023	

9	Попередній захист роботи	30.05.2023	
---	--------------------------	------------	--

Студент

Ілля КРАВЧЕНКО

Керівник роботи

Валерій НОВІКОВ

РЕФЕРАТ

Текстова частина дипломної роботи: 93 с., 37 рисунків, 1 таблиця, 20 лістингів, 10 джерел.

Мета роботи –вивчення поняття багатоадресної маршрутизації, аналіз існуючих протоколів та алгоритмів багатоадресної маршрутизації трафіку групового мовлення в IP-мережах та порівняння підходів до проектування мережі з багатоадресною маршрутизацією трафіку групового мовлення.

В даній роботі розглядається Багатоадресна маршрутизація трафіку групового мовлення в IP-мережах, тому що з кожним роком, в світі, де зв'язок та обмін інформацією відіграють ключову роль, мережеві технології стають все важливішим елементом нашого повсякденного життя – створюються нові додатки та нові можливості для користувачів, і це збільшує навантаження на мережі, а багатоадресна маршрутизація, працює таким чином, що велика кількість трафіку доставляється деякій кількості отримувачам не навантажуючи саму мережу. Проведено порівняння підходів до проектування мережі з багатоадресною маршрутизацією трафіку групового мовлення.

Багатоадресна маршрутизація, групове мовлення, IP-мережі, IGMP, PIM-SM, PIM-DM, PIM-SDM, GRAPHICAL NETWORK SIMULATOR-3

ABSTRACT

Text part of the thesis: 93 pages, 37 figures, 1 table, 20 listings, 10 sources.

The purpose of the work is to study the concept of multicast routing, analyze existing protocols and algorithms for multicast routing of group broadcast traffic in IP networks and compare approaches to network design with multicast routing of group broadcast traffic.

This paper deals with multicast routing of group broadcast traffic in IP networks, because every year, in a world where communication and information exchange play a key role, network technologies are becoming an increasingly important element of our daily lives - new applications and new opportunities for users are being created, and this increases the load on networks, and multicast routing works in such a way that a large amount of traffic is delivered to a certain number of recipients without loading the network itself. The article compares approaches to network design with multicast routing of group broadcast traffic.

MULTICAST ROUTING, GROUP BROADCASTING, IP-NETWORKS, IGMP, PIM-SM, PIM-DM, PIM-SDM, GRAPHICAL NETWORK SIMULATOR-3

ЗМІСТ

ПЕРЕЛІК СКОРОЧЕНЬ.....	9
ВСТУП	12
1 АНАЛІЗ СПОСОБІВ ЗАСТОСУВАННЯ БАГАТОАДРЕСНОЇ МАРШРУТИЗАЦІЇ ТРАФІКУ ГРУПОВОГО МОВЛЕННЯ В ІР-МЕРЕЖАХ	14
1.1 Аналіз технології багатоадресної маршрутизації	14
1.1.1 Поняття багатоадресної маршрутизації	14
1.1.2 Поняття Джерело та Отримувач	16
1.1.3 Групове мовлення.....	18
1.1.4 Порівняння Одноадресної, Широкомовної та Багатоадресної Маршрутизації.....	21
1.1.5 Способи використання Багатоадресної маршрутизації.....	24
1.1.6 Переваги і недоліки Багатоадресної маршрутизації.....	27
1.2 Висновки до розділу 1	29
2 ОГЛЯД ПРОТОКОЛІВ БАГАТОАДРЕСНОЇ МАРШРУТИЗАЦІЇ ТРАФІКУ ГРУПОВОГО МОВЛЕННЯ В ІР-МЕРЕЖАХ.....	30
2.1 Доступ до мережі та багатоадресна передача на 2 рівні Багатоадресної маршрутизації.....	30
2.1.1 Вступ до IGMPv1	30
2.1.2 Розгляд IGMPv2.....	32
2.1.3 Вступ до IGMPv3.....	33
2.1.4 Змішані групи: Сумісність між IGMPv1, v2 і v3	36
2.1.5 Cisco Group Management Protocol(CGMP) і Router-port Group Management Protocol(RGMP)	36
2.1.6 Snooping і IGMP Snooping	38
2.2 Багатоадресна передача даних на 3 рівні Багатоадресної маршрутизації.....	40
2.2.1 Protocol Independent Multicast(PIM) і багатоадресні дерева.....	40
2.2.2 Protocol-Independent Multicast: PIM-DM	56

					НТУУ1884-с-23.08ТС-91.2023.ПЗ				
Змн.	Лист	№ докум.	Підпис	Дата					
Розроб.		Кравченко І. М.			Аналіз протоколів багатоадресної маршрутизації трафіку групового мовлення в ІР-мережах	Літ.	Арк.	Акрушів	
Перевір.		Новіков В. І.					7	93	
Реценз.		Лисенко О. І.				НН ІТС			
Н. Контр.		Новіков В. І.							
Затверд.		Скулиш М. А.							

2.2.3	Protocol-Independent Multicast: Describing PIM-SM	58
2.2.4	PIM Sparse-Dense-Mode	60
2.2.5	Protocol-Independent Multicast. Опис PIM-BiDir.....	61
2.3	Висновки до розділу 2.....	63
3	ПРОЕКТУВАННЯ МЕРЕЖІ З ВИКОРИСТАННЯМ БАГАТОАДРЕСНОЇ МАРШРУТИЗАЦІЇ ТРАФІКУ ГРУПОВОГО МОВЛЕННЯ.....	64
3.1	Короткі теоретичні відомості	64
3.1.1	Початкові поняття	64
3.1.2	Використання програми GNS3 та її особливості	65
3.1.3	Характеристики GNS3	68
3.1.4	Переваги та недоліки GNS3	70
3.2	Порівняння різних видів роботи протоколу PIM на прикладі мережі багатоадресної маршрутизації трафіку групового мовлення.	72
3.2.1	Схема мережі, для якої будуть налаштовані PIM Dense Mode, PIM Sparse Mode та PIM Sparse-Dense Mode.....	72
3.2.2	Налаштування мережі з PIM-Dense Mode та PIM-Sparse Mode та їх порівняння.....	73
3.2.3	Налаштування мережі з PIM Sparse-Dense Mode.....	82
3.3	Висновки до розділу 3.....	91
	ВИСНОВКИ.....	92
	ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ.....	93

ПЕРЕЛІК СКОРОЧЕНЬ

IP	- Internet Protocol - міжмережевий протокол
IPv4	- Internet Protocol version 4 - міжмережевий протокол версії 4
IPv6	- Internet Protocol version 6 - міжмережевий протокол версії 6
IGMP	- Internet Group Management Protocol - протоколу групових повідомлень Інтернету
MRIB	- Multicast Routing Information Base - інформаційна база багатоадресної маршрутизації
MFIB	- Multicast Forwarding Information Base - інформаційна база багатоадресної переадресації
RFC	- Request for Comments - документ із серії пронумерованих інформаційних документів Інтернету, що містить технічні специфікації та Стандарти
IETF	- The Internet Engineering Task Force - відкрите міжнародне співтовариство проектувальників, учених, мережевих операторів і провайдерів
OSPF	- Open Shortest Path First -
IANA	- Internet Assigned Numbers Authority - Адміністрація адресного простору Інтернет
ASN	- Authonomy System Number - автономний системний номер
NTP	- Network Time Protocol - Мережевий протокол часу
SDP/SAP	- Session Description Protocol - протокол опису сеансу зв'язку
SSM	- Source-specific multicast
TCP	- Transmission Control Protocol - протокол керування передаванням
VoD	- Video on Demand – відео на вимогу
UDP	- User Datagram Protocol - Протокол датаграм користувача
DR	- Designated Router - призначений маршрутизатор
PIM	- Protocol Independent Multicast - Незалежна від протоколу

	групова маршрутизація
CGMP	- Cisco Group Management Protocol - Протокол керування групою від компанії Cisco
RGMP	- Router-port Group Management Protocol - Протокол керування групою портів маршрутизатора
GDA	- Group Destination Address - групова адреса призначення
SNAP	- SubNetwork Access Protocol - протокол доступу до підмереж
MAC	- Media Access Control — управління доступом до посередників
CAM	- Content Addressable Memory
VLAN	- Virtual Local Area Network — віртуальна локальна комп'ютерна мережа
STP	- Spanning Tree Protocol - протокол кістякового дерева
RP	- Rendezvous Point - точка рандеву
OIL	- Out Interfaces List - список вихідних інтерфейсів
DVMRP	- Distance Vector Multicast Routing Protocol - Протокол багатоадресної маршрутизації з вектором відстані
MOSPF	- Multicast Open Shortest Path First – протокол багатоадресної динамічної маршрутизації
LAN	- local area network – локальна мережа
WAN	- Wide Area Network - Глобальна мережа
PIM-DM	- Protocol Independent Multicast Dense Mode – щільний режим роботи протоколу
PIM-SM	- Protocol Independent Multicast Sparse Mode – розріджений режим роботи протоколу
PIM-SDM	- Protocol Independent Multicast Sparse Dense Mode - Розріджено-щільний режим роботи протоколу
GNS 3	- Graphical Network Simulator – Графічний симулятор мережі
GUI	- Graphical User Interface - Дружній графічний інтерфейс

VM	- Virtual Machine - віртуальна машина
VPN	- Virtual Private Network - віртуальна приватна мережа
ASIC	Application-specific integrated circuit - інтегральна схема для специфічного застосування
TTL	- Time To Live - значення часу життя

ВСТУП

У сучасному світі, де зв'язок та обмін інформацією відіграють ключову роль, мережеві технології стають все важливішим елементом нашого повсякденного життя. Зростання кількості користувачів та розширення можливостей мережевого середовища призводять до появи нових викликів у сфері маршрутизації трафіку. Особливу увагу заслуговує багатоадресна маршрутизація трафіку групового мовлення в IP-мережах, яка забезпечує ефективну передачу даних до групи одночасних отримувачів.

Актуальність теми. У зв'язку зі зростанням популярності мультимедійних додатків та інтерактивних сервісів, які передають дані до великої кількості одночасних користувачів, багатоадресна маршрутизація трафіку групового мовлення стає дедалі важливішою проблемою. Цей механізм дозволяє ефективно розподіляти дані до всіх учасників групи, знижуючи навантаження на мережеві ресурси та забезпечуючи плавну передачу мультимедійного контенту.

Мета дослідження – основною метою цього дослідження є вивчення поняття багатоадресної маршрутизації, аналіз існуючих протоколів та алгоритмів багатоадресної маршрутизації трафіку групового мовлення в IP-мережах та порівняння підходів до проектування мережі з багатоадресною маршрутизацією трафіку групового мовлення.

Об'єкт дослідження – мережа з багатоадресною маршрутизацією трафіку групового мовлення в IP-мережах.

В першому розділі здійснюється опис та визначення багатоадресної маршрутизацією трафіку групового мовлення в IP-мережах.

В другому розділі розглядається аналіз протоколів за допомогою яких будується мережа багатоадресної маршрутизації трафіку групового мовлення в IP-мережах. Основну увагу в цьому розділі приділено протоколам IGMP та PIM.

В третьому розділі представляється порівняння підходів до проектування мережі з багатоадресною маршрутизацією трафіку групового мовлення на прикладі однакової мережі для якої паралельно налаштовуються різні підходи проектування.

1 АНАЛІЗ СПОСОБІВ ЗАСТОСУВАННЯ БАГАТОАДРЕСНОЇ МАРШРУТИЗАЦІЇ ТРАФІКУ ГРУПОВОГО МОВЛЕННЯ В IP-МЕРЕЖАХ

1.1 Аналіз технології багатоадресної маршрутизації

1.1.1 Поняття багатоадресної маршрутизації

Сучасні мультимедійні застосунки інтегрують звук, графіку, анімацію, текст, відео. Такого типу додатки сьогодні стали ефективним засобом корпоративного спілкування. Проте, надсилання таких мультимедійних даних у мережі кампусу потребує великої пропускної здатності. Традиційні методи комунікацій за протоколом IP дають змогу вузлу надсилати пакети одному вузлу (одноадресна передача) чи всім вузлам одночасно (широкомовна передача). Багатоадресна IP-розсилка надає третю можливість: вона дає змогу вузлу відправляти пакети певній підмножині вузлів у вигляді групової передачі.

Багатоадресатна IP-розсилка (IP Multicast routing)- це технологія, яка зменшує обсяг переданих даних та заощаджує смугу пропускання за допомогою одночасного доставлення потоку інформації тисячам корпоративних одержувачів або індивідуальних користувачів.

Мета багатоадресної розсилки полягає в тому, щоб відправити інформацію на обрану групу пристроїв і уникнути дублювання потоків трафіку, що призводить до економії дорогоцінних мережевих ресурсів, таких як пропускна спроможність використання.

Багатоадресатна маршрутизація використовується, щоб надіслати пакети з даними багатьом одержувачам. Наприклад, мультимедійний сервер відправляє одну копію кожного пакета з єдиною IP адресою призначення, цей пакет можуть отримати багато кінцевих вузлів, якщо вони готові приймати й обробляти пакет, відправлений на цю спеціальну адресу.

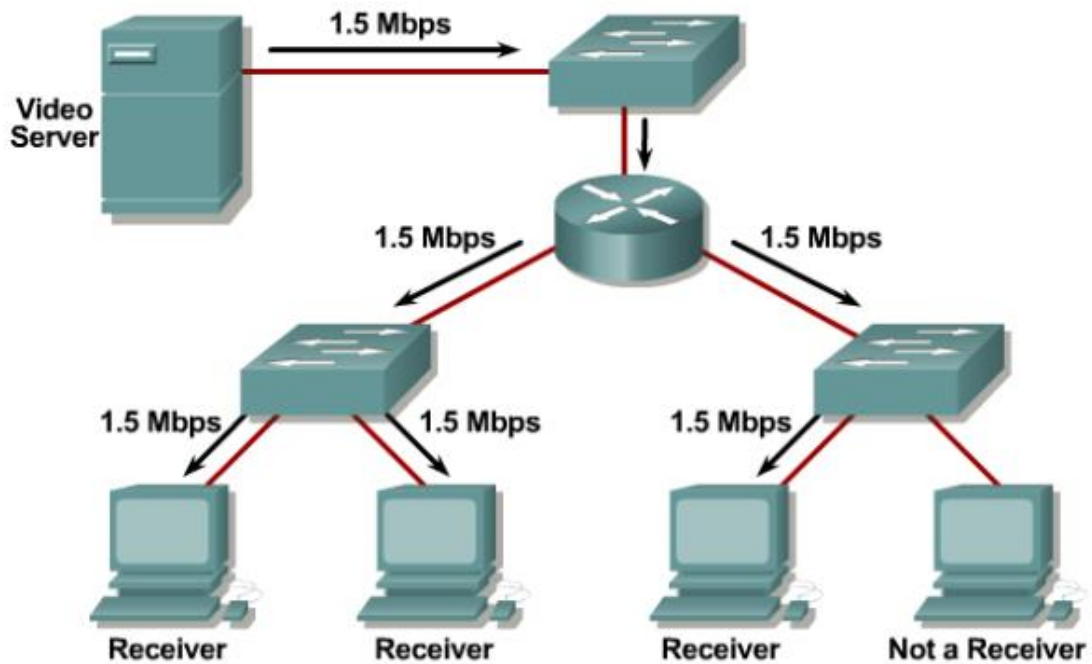


Рисунок 1.1 - Групова доставка на основі індивідуальних адрес

На рисунку 1.1 відео-сервер передає єдиний відео-потік до набору пристроїв, які прослуховують специфічну групову адресу. Використовується тільки 1.5 Mbps пропускної здатності мережі, незважаючи на кількість одержувачів. Надсилаючи пакети даних до багатьох одержувачів, ці пакети не дублюються для кожного одержувача, але відправляються єдиним потоком, де кінцеві роутери збільшують кількість цих пакетів для одержувачів, коли це необхідно. При цьому роутери обробляють менше пакетів, ніж при звичайній одноадресній передачі даних, тому що вони одержують тільки єдину копію пакета. Оскільки кінцеві роутери виконують збільшення числа пакетів для доставки їх одержувачам, відправник, або джерело групового мовлення, не має необхідності в знанні одноадресних адрес усіх одержувачів. Отже, з малюнку ми бачимо, що багатоадресна маршрутизація увібрала до себе найкраще – мінімізацію кількості сеансів, необхідних відправнику, і зменшення навантаження на мережу до одного потоку трафіку. Ми також отримуємо переваги одноадресної розсилки, надсилаючи пакети даних лише тим пристроям, які зацікавлені в їх

отриманні, тобто надсилається лише один потік, який тиражується на зацікавлені приймачі.

1.1.2 Поняття Джерело та Отримувач

Говорячи про багатоадресну передачу, завжди є два типи хостів: джерело багатоадресної передачі та приймач багатоадресної передачі. Джерелом багатоадресної розсилки може бути будь-який пристрій з IP-адресою в мережі. Щоб стати джерелом, хосту необхідно лише надіслати повідомлення на IP-адресу групи багатоадресної розсилки. (Див. рисунок 1.2) Усі три відправники на цій схемі надсилають повідомлення на одну й ту саму адресу призначення багатоадресної мережі 239.1.1.1.

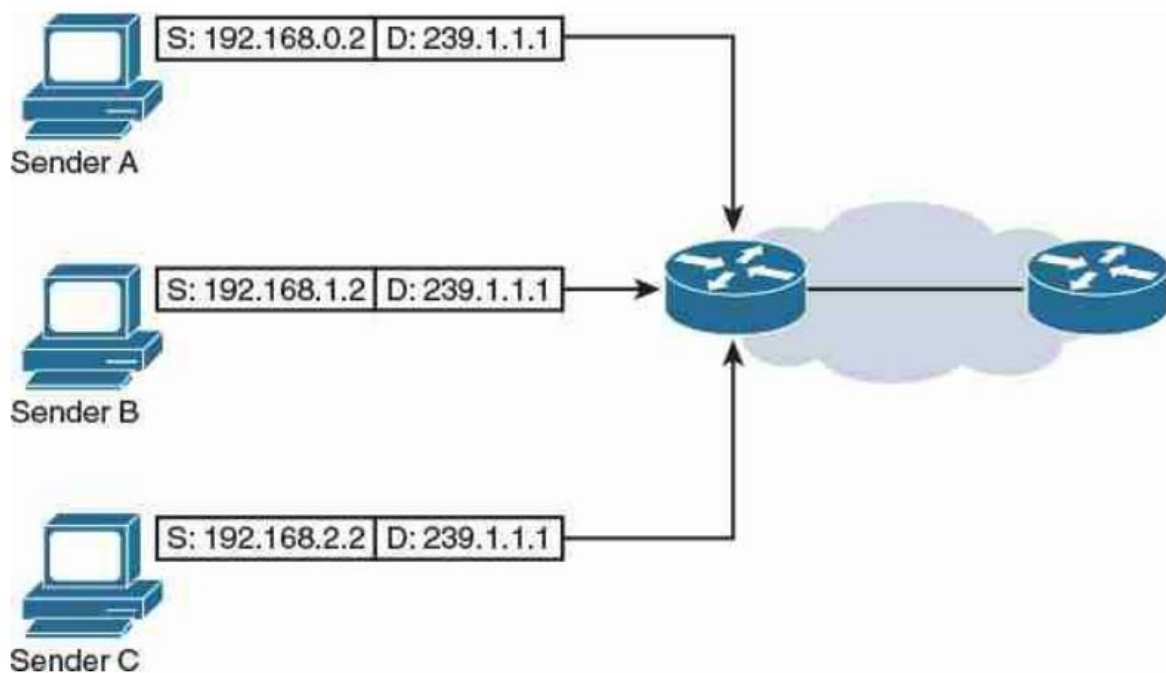


Рисунок 1.2 - Джерело багатоадресної передачі [1]

Джерело - це будь-який пристрій з підтримкою багатоадресної розсилки, який може надсилати ту чи іншу інформацію групі приймачів, які зацікавлені в отриманні цієї інформації. Перед надсиленням багатоадресного повідомлення джерелу не потрібно вказувати на намір надсилення до групи

повідомлення. Будь-який пристрій з підтримкою IP, включно з маршрутизатором або комутатором, може надсилати пакети до багатоадресної групи багатоадресної розсилки. Коли IP-маршрутизатор багатоадресної розсилки обробляє повідомлення, призначене для групи багатоадресної розсилки, створюється новий запис переадресації багатоадресної розсилки. Це запис "Джерело, група" і називається кома джерела група, (S, G) . Запис (S, G) для відправника A на рисунку 1.2 матиме вигляд (192.168.0.2, 239.1.1.1).

Запис (S, G) відноситься до IP-адреси відправника і групи одержувача, які розділені комою. IP-адреси відправників на Рисунку 1.2 є адресами джерела (S), тоді як IP-адреси одержувачів групи - G. Зверніть увагу, що всі три пристрої відправляють повідомлення на одну й ту саму групову адресу.

Приймач - це будь-який пристрій з підтримкою багатоадресної розсилки, який виявив бажання отримати певну багатоадресну інформацію в групі або певній парі (S, G) . Якщо тільки багатоадресна розсилка не є локальною (не передається через мережу жодним маршрутизатором, як, наприклад, група IPv4 224.0.0.2 "all-routers"), IP-пристрої повинні бути налаштовані адміністратором або програмою для підписки на певну групу багатоадресної розсилки. Після підписки отримувач багатоадресної розсилки приймає пакети, які надходять з адресою призначення групи багатоадресної розсилки, наприклад, група 239.1.1.1, показана на рисунку 1.2.

Керування підпискою на групу здійснюється за допомогою протоколу групових повідомлень Інтернету (IGMP). Коли підписка одержувача на певну групу або набір груп отримується або ініціюється маршрутизатором, маршрутизатор додає до MRIB запис, який називається "зірка з комою G", $(*, G)$. Запис $(*, G)$ просто означає, що маршрутизатор зацікавлений у групі.

MRIB (Multicast Routing Information Base) відповідає за інформацію про маршрутизацію, яка генерується протоколами багатоадресної розсилки. Ця інформація надходить до MFIB (Multicast Forwarding Information Base), яка

відповідає за пересилання багатоадресних пакетів і збір статистики про потік багатоадресної розсилки.

1.1.3 Групове мовлення

Багатоадресна передача базується на концепції групи. Група багатоадресної розсилки - це група одержувачів, які висловили зацікавленість в отриманні певного потоку даних. Ця група не має фізичних або географічних меж, і її члени можуть перебувати де завгодно в Інтернеті або в будь-якій приватній, взаємопов'язаній мережі. Вузли, які зацікавлені в отриманні даних, що надходять до даної групи, повинні приєднатися до групи за допомогою протоколу IGMP (Internet Group Management Protocol - протокол управління групами Інтернету).

Групова адреса просто ідентифікує групу вузлів, зацікавлених у потоці. Адреса групи у поєднанні з IP-адресою джерела ідентифікує потік багатоадресної передачі. Вузли-одержувачі висловлюють свою зацікавленість в отриманні певного потоку даних, повідомляючи вищі мережеві пристрої про членство в групі. Це вираження зацікавленості відоме як приєднання до групи.

Групова адресація IPv4 та IPv6 контролюється комбінацією RFC IETF. Найбільш важливим і актуальним є RFC IPv4 5771, який остаточно визначає призначення групового адресного простору багатоадресної передачі, а також окремих типів груп у цьому просторі. Адресний простір групових блоків багатоадресної маршрутизації є похідним від адресного простору IPv4 з урахуванням класів.

Роутери відрізняють трафік багатоадресної передачі, використовуючи зарезервований клас D IP адрес.

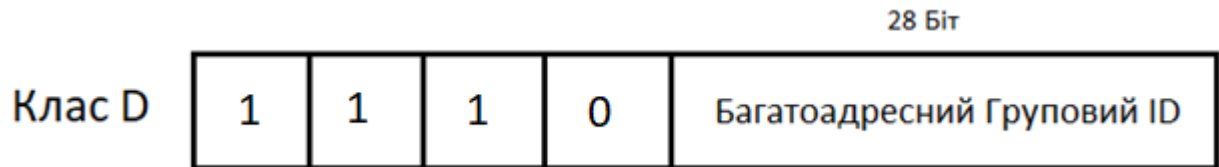


Рисунок 1.3 - Клас D IP адрес [2]

Мережеві пристрої можуть швидко визначити клас D IP адрес, подивившись на перші 4 старші біти, які завжди дорівнюють 1110. Наступні 28 біт в адресі класу D є адресою групи (рисунок 1.3). Тож діапазон Multicast IP-адрес становить діапазон від 224.0.0.0 до 239.255.255.255 (таблиця 1). Адресний простір Multicast IP адрес розділено на такі групи:

Таблиця 1.1 - Адресний простір Multicast IP [1]

Address Range	Size	Designation
224.0.0.0–224.0.0.255	/24	Local Network Control
224.0.1.0–224.0.1.255	/24	Internetwork Control
224.0.2.0–224.0.255.255	(65024)	AD-HOC I
224.1.0.0–224.1.255.255	/16	RESERVED
224.2.0.0–224.2.255.255	/16	SDP/SAP
224.3.0.0–224.4.255.255	2 total /16's	AD-HOC II
224.5.0.0–224.255.255.255	251 total /16's	RESERVED
225.0.0.0–231.255.255.255	7 /8's	RESERVED
232.0.0.0–232.255.255.255	/8	Source-Specific Multicast
233.0.0.0–233.251.255.255	(16515072)	GLOP
233.252.0.0–233.255.255.255	/14	AD-HOC III
234.0.0.0–238.255.255.255	5 total /8's	RESERVED
239.0.0.0–239.255.255.255	/8	Administratively Scoped

Блок локального керування мережею (224.0.0/24): Блок локального керування використовується для трафіку керування певним протоколом. Інтерфейси маршрутизаторів слухають, але не пересилають багатоадресний трафік локального управління; наприклад, OSPF "всі маршрутизатори"

(224.0.0.5). Призначення в цьому блоці публічно контролюються організацією IANA.

Блок керування міжмережевим зв'язком (224.0.1/24): призначений для трафіку керування протоколом, який інтерфейси маршрутизатора можуть пересилати через автономний системний номер (ASN) або через Інтернет. Приклади включають 224.0.1.1 Network Time Protocol (NTP), визначений в RFC 4330, і 224.0.1.68 mdhcdpdiscover, визначений в RFC 2730. Призначення груп управління мережею також публічно контролюються IANA.

Блоки AD-HOC (I: 224.0.2.0-224.0.255.255, II: 224.3.0.0-224.4.255.255 та III:233.252.0.0–233.255.255.255): Традиційно призначаються програмам, які не вписуються у блоки локальний блок або блок керування міжмережевим зв'язком. Інтерфейси маршрутизаторів можуть пересилати пакети AD-HOC глобально). IANA контролює будь-які публічні призначення блоків AD-HOC, а майбутні призначення будуть призначатимуться з блоку AD-HOC III, якщо вони не будуть більш придатними для локального керування або керування мережею Інтернет. Публічне використання нерозподіленого простору AD-HOC також дозволено.

Блок SDP/SAP (224.2.0.0/16): Протокол опису сеансу/оголошення сеансу Блок протоколу (SDP/SAP) призначається для додатків, які отримують адреси через SAP, як описано у RFC 2974.

Блок багатоадресної передачі на конкретне джерело (232.0.0.0.0/8): Адресація SSM визначена в RFC 4607. SSM - це групова модель багатоадресної передачі IP, в якій багатоадресний трафік пересилається на одержувачі тільки з тих джерел багатоадресної передачі, до яких одержувачі явно висловили зацікавленість. Для використання блоку SSM не потрібен офіційний дозвіл від IANA оскільки програма є локальною для хоста; однак, згідно з політикою IANA, блок явно зарезервований для додатків SSM і не повинен використовуватися для будь-яких інших цілей.

Блок GLOP (233.0.0.0.0/8): Ці адреси призначаються статично з глобальною областю дії. Кожне GLOP відповідає домену з публічним 16-

бітним автономним системним номером (ASN), який видається IANA. ASN вставляється крапково-десятковим способом у середні два октети групової адреси багатоадресної передачі(X.Y).

Адміністративний блок (239.0.0.0/8): Адреси з адміністративними обмеженнями призначені для локального використання в межах приватного домену, як описано в RFC 2365. Ці групові адреси виконують аналогічну функцію, як і блок приватних IP-адрес RFC 1918 (наприклад, блоки 10.0.0.0/8 або 172.16 - 31.0.0/16). 31.0.0/16). Мережеві архітектори можуть створити схему адрес, використовуючи цей блок, яка найкраще відповідає потребам приватного домену, і можуть додатково розділити масштабування на певні географічні регіони, додатки або мережі. Додаткова інформація про найкращі практики для визначення обсягу.

Отже, сфера застосування кожного блоку варіюється від окремих сегментів до великих корпоративних мереж багатоадресної розсилки та глобальної мережі Інтернет. Важливо розуміти RFC і стандарти для груп багатоадресної розсилки при проектуванні мереж багатоадресної розсилки. Багатоадресна передача адреси груп багатоадресної передачі відіграють дуже істотну роль у "масштабуванні" доменів багатоадресної передачі. І при побудові мережі з багатоадресною маршрутизацією потрібно звертати увагу для того щоб не використати уже зарезервовану адресу.

1.1.4 Порівняння Одноадресної, Широкомовної та Багатоадресної Маршрутизації

При одноадресній передачі (unicast) на основі унікальних адрес джерело даних, яке треба доставити деякій групі вузлів, генерує їх у кількості екземплярів, що відповідає кількості вузлів-отримувачів, які перебувають у цій групі (рисунок. 1.4). Тобто ми передаємо велику кількість пакетів великій кількості отримувачів. Очевидно, що передача декількох ідентичних копій на ділянках, де маршрути до різних учасників групи перекриваються (це

особливо характерно для початкових ділянок), призводить до надлишкового трафіку.

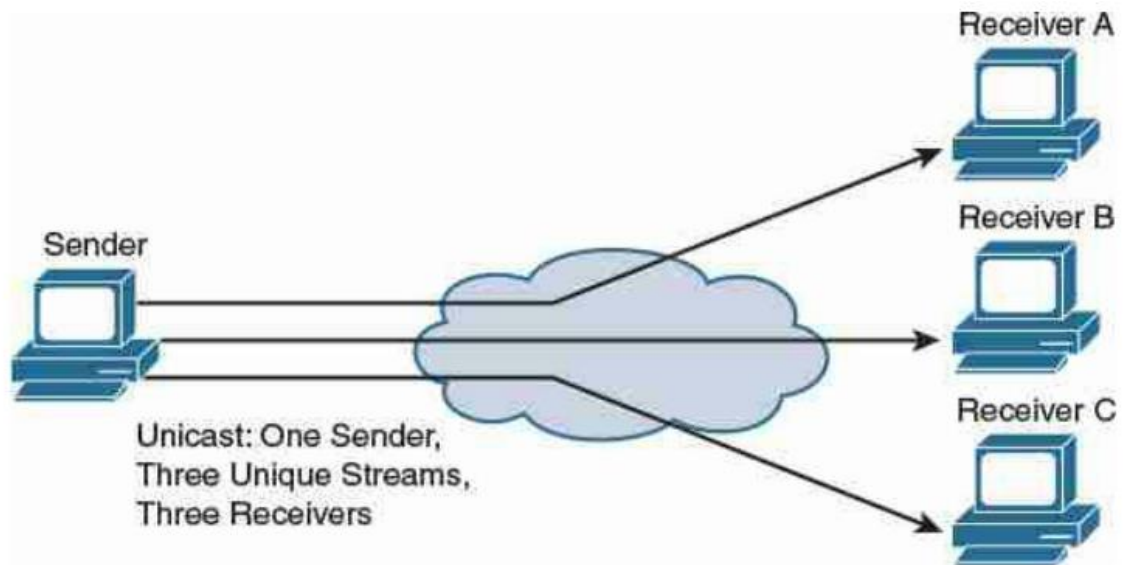


Рисунок 1.4 - Одноадресна розсилка з використанням декількох потоків [1]

Під час використання широкомовного передачі (broadcast) станція надсилає пакети, використовуючи широкомовні адреси (рисунок 1.5). У цій схемі для того, щоб доставити дані групі вузлів-одержувачів, джерело генерує один екземпляр даних, але постачає цей екземпляр широкомовною адресою, що диктує маршрутизаторам мережі копіювати дані й надсилати їх усім кінцевим вузлам незалежно від того, "зацікавлені" вузли в одержанні цих даних чи ні. У цьому випадку, як і в попередньому, істотна частка трафіку є надмірною.

Таким чином, традиційні механізми доставки пакетів стеку TCP/IP мало придатні для підтримки групового мовлення. У такому випадку найбільш ефективним рішенням є використання спеціально розробленого механізму групового мовлення, який орієнтований на скорочення надмірного трафіку і накладних витрат у мережі.

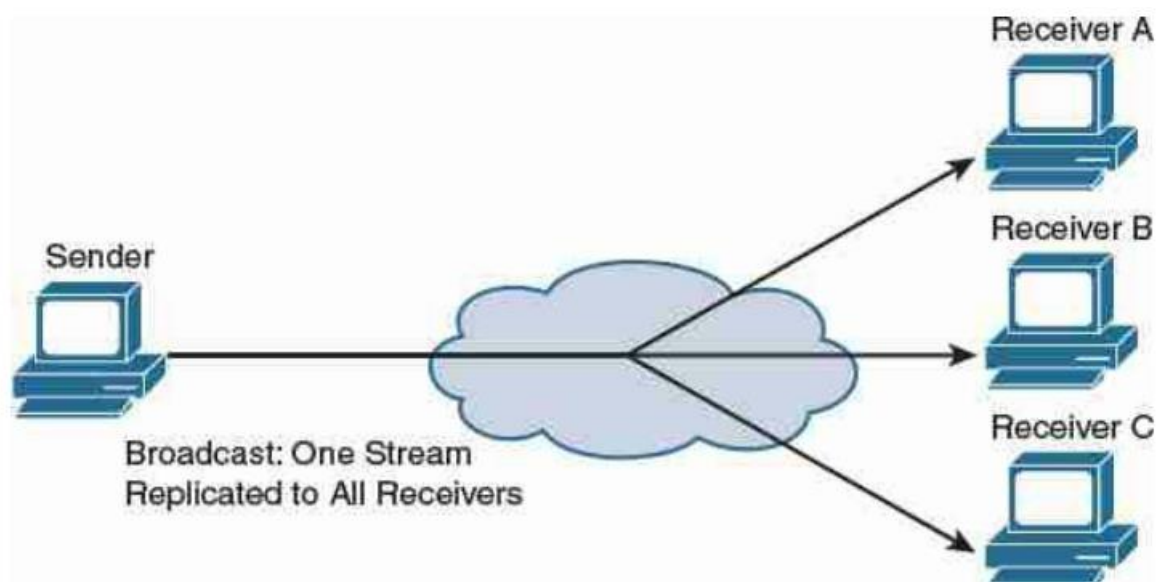


Рисунок 1.5 - Доставка на основі широкомовної адреси [1]

Отже основна ідея групового мовлення полягає в наступному: джерело генерує тільки один екземпляр повідомлення з груповою адресою, яке потім, у міру переміщення мережею, копіюється на кожній з "розгалужень", що ведуть до того чи іншого члена групи, зазначеної в адресі цього повідомлення (рисунок 1.6).

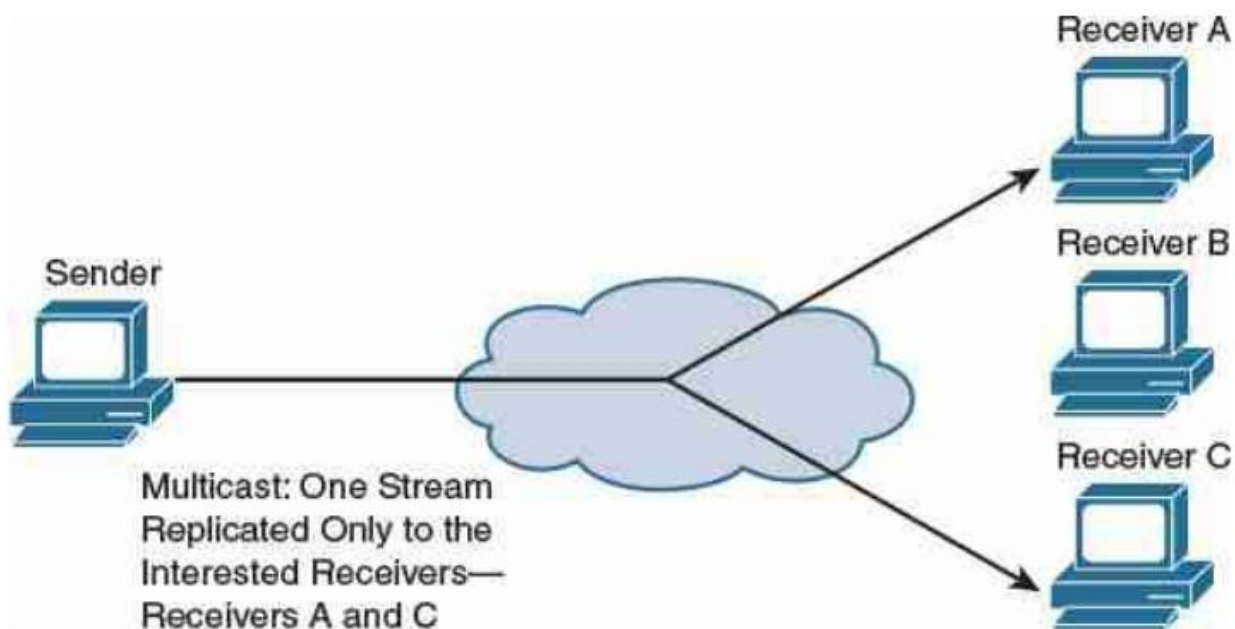


Рисунок 1.6 - Схема групового мовлення [1]

Зрештою пакет із груповою адресою досягає маршрутизатора, до якого безпосередньо під'єднана мережа з хостами - членами даної групи. Нагадаємо, що у хостів, які належать до тієї чи іншої групи, інтерфейс поряд з індивідуальною адресою має ще й групову адресу - адресу класу D, яку називають також адресою групового мовлення. Інтерфейс може мати навіть кілька групових адрес - за кількістю груп, до яких входить цей хост.

За такого підходу дані розсилаються тільки тим вузлам, які зацікавлені в їхньому отриманні. Функція реплікації групового повідомлення і просування копій в бік членів групи покладається на маршрутизатори, для чого вони мають бути обладнані відповідними програмно-апаратними засобами. Такий режим економить пропускну здатність за рахунок передачі лише того трафіку, який необхідний.

1.1.5 Способи використання Багатоадресної маршрутизації

Мережева інфраструктура створена для підтримки додатків і сервісів. Ці програми та сервіси дозволяють організації - державній установі, банку, торговельній організації, лікарні, службі швидкої допомоги чи бізнесу - виконувати свою місію або бізнес-цілі. Створення інфраструктури, яка дозволяє ефективно використовувати ці багатоадресні додатки і сервіси, допомагає забезпечити успіх цієї організації

Багатоадресні програми "один до багатьох"

Така форма багатоадресної передачі, як показано на рисунку 1.7, є найпоширенішою формою багатоадресної передачі.

Як випливає з назви, є один відправник і кілька одночасних одержувачів. Типові застосування включають трансляцію відео і аудіо, але існує також багато інших додатків, в тому числі наступні:

- Телебачення;
- Радіо;
- Дистанційне навчання;

- Спільний доступ до презентацій та дошки для доповідей;
- Програмне забезпечення для комп'ютерної графіки та оновлення додатків;
- Розподіл та кешування даних;
- Погода;
- Новини.

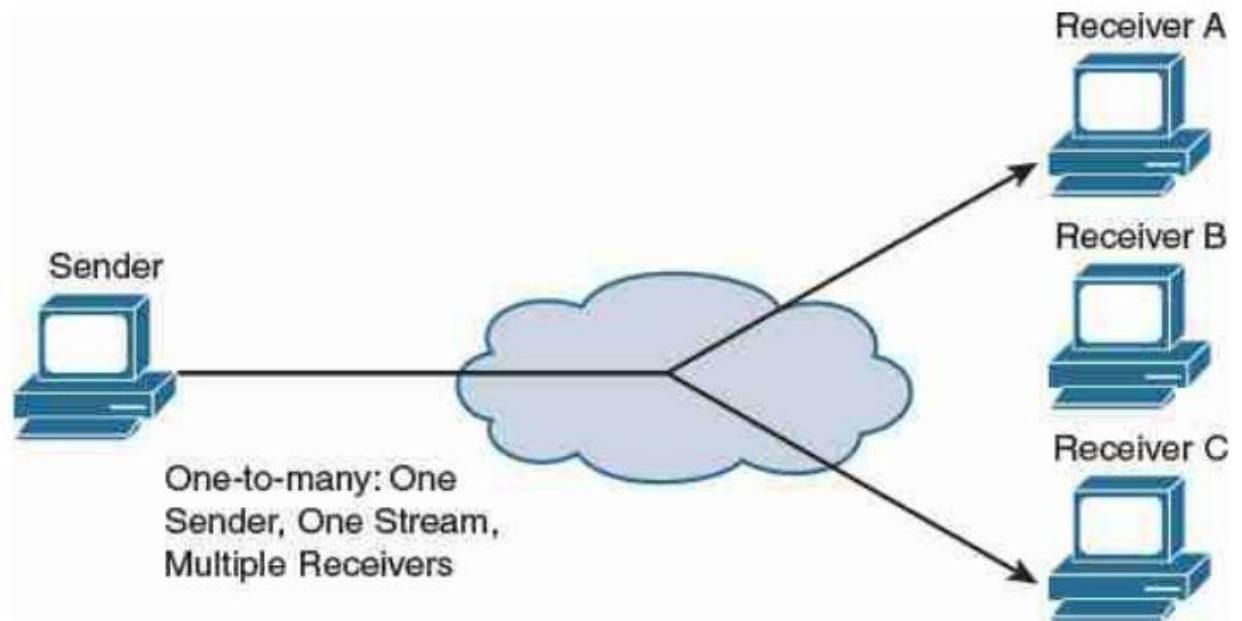


Рисунок 1.7 - Багатоадресна розсилка з використанням "один до багатьох" [1]

Багатоадресні програми багатоадресної розсилки

У програмах багатоадресної розсилки "багато до багатьох" відправники також виступають у якості одержувачів. Це дозволяє всім пристроям у групі одночасно спілкуватися один з одним, як показано на рисунку 1.8.

Додатки багато-до-багатьох включають наступне:

- Аудіо- та відеозв'язок;
- Розподіл, кешування та синхронізація даних;
- Додатки для групового чату;
- Фінансові програми;
- Програми для опитувань;

- Багатокористувацькі ігри.

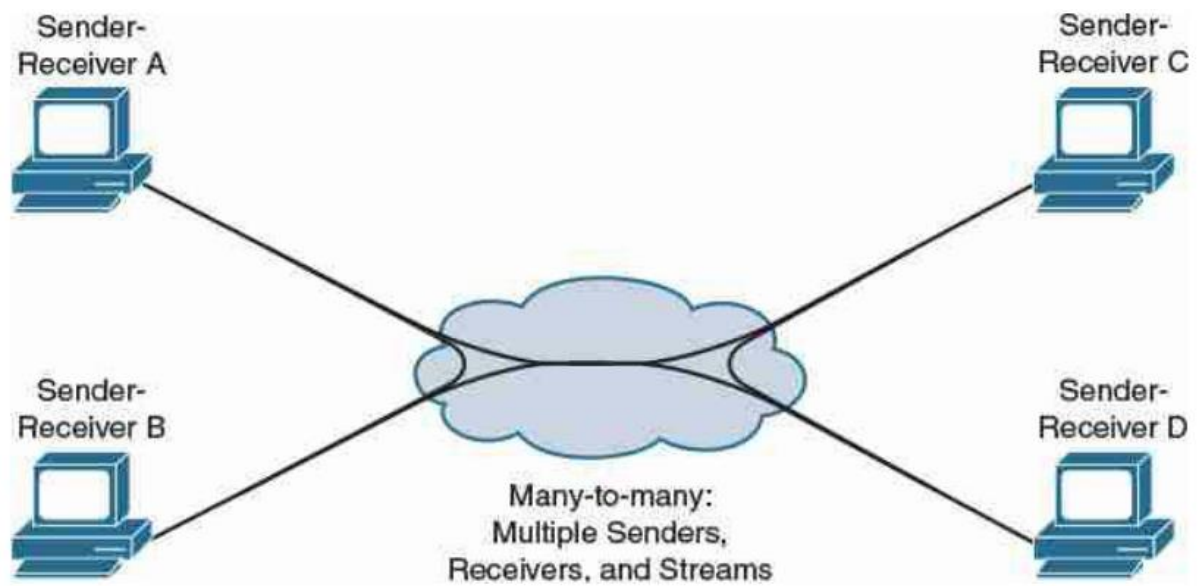


Рисунок 1.8 - Багатоадресна розсилка з використанням "багато до багатьох"

[1]

Багатоадресні програми "багато до одного"

Програми багато-до-одного мають багато відправників, але лише одного або декількох одержувачів, як показано на рисунку 1.9 .

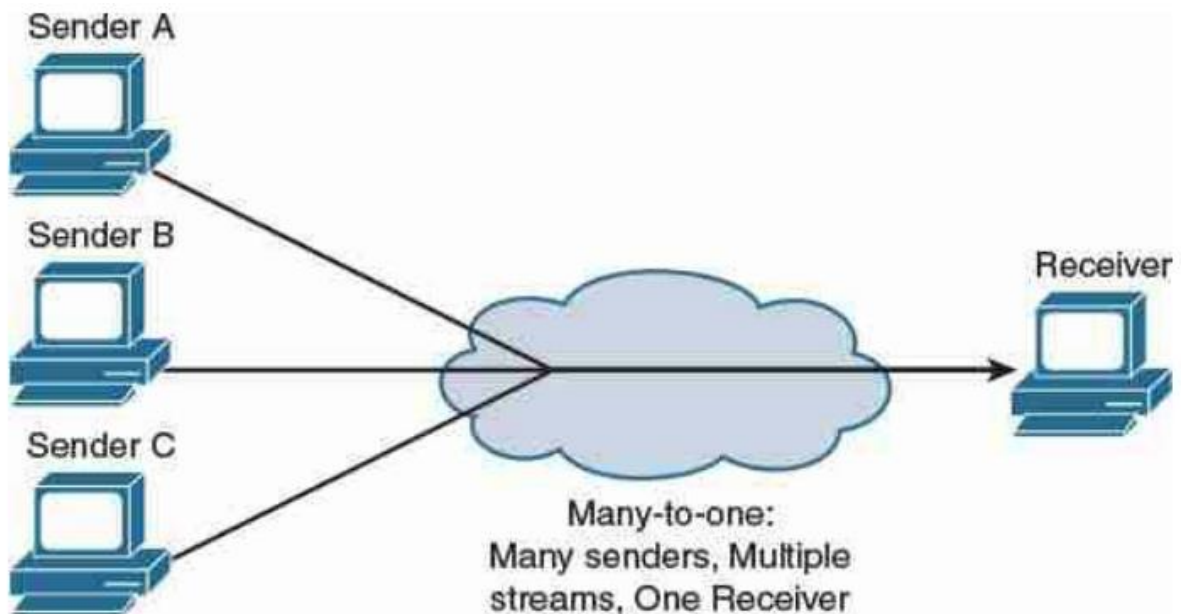


Рисунок 1.9 - Багатоадресна розсилка "багато до одного" [1]

Це не звичайна пропозиція багатоадресної розсилки, і вона має значну складність у тому, що одержувач може не обробити отриману інформацію, коли багато пристроїв надсилають багатоадресні потоки. Масштабованість є суттєвою проблемою цієї моделі. У певному розумінні, ця модель не є покращенням у порівнянні з використанням одноадресних потоків, але натомість забезпечує більшу гнучкість конфігурації програми. Насправді, у багатьох випадках одержувач відповідатиме відправникам за допомогою одноадресного потоку.

Деякі з додатків для багато-до-одного включають наступне:

- Збір даних;
- Виявлення послуг;
- Опитування.

Багатоадресні додатки можуть споживати величезну кількість пропускної здатності, як у випадку з відео високої чіткості, або ж вони можуть споживати дуже мало пропускної здатності мережі, як у випадку з оновленням часу. Всі ці додатки покладаються на інфраструктуру для успішної роботи раніше згаданих додатків і сервісів

1.1.6 Переваги і недоліки багатоадресної маршрутизації

Багатоадресна передача забезпечує чимало переваг у порівнянні з Одноадресної у випадку передавання одного-до-багатьох або багато-до-багатьох, з них:

- Розширена ефективність: мережева пропускна здатність використовується ефективніше, тому що множинні потоки даних замінюються єдиною передачею;
- Оптимізована продуктивність: менше копій даних, які потребують відправки та обробки;

- Розподілені додатки: розподілені додатки не будуть працювати з Одноадресною маршрутизацією, тому що Одноадресна передача не масштабується (рівень трафіку і кількість клієнтів можна рахувати за пропорцією 1:1).

Є й інші переваги:

- При еквівалентній відправці групового трафіку, відправник використовує менше продуктивності та пропускну здатність;
- Групові пакети не потребують високої пропускну здатності як одноадресні пакети, і вони прибувають до одержувачів майже одночасно;
- Багатоадресна маршрутизація дає цілу низку нових додатків, які неможливі при Одноадресній (наприклад, відео за запитом [VoD]);

Але існують також певні недоліки multicast, які слід розглянути.

Більшість Багатоадресних додатків засновані на User Datagram Protocol (UDP). Це дає можливість позбутися деяких небажаних наслідків, у порівнянні з аналогічними Одноадресними застосунками, що використовують TCP.

Недоліки Багатоадресної маршрутизації на основі UDP:

- Найшвидша доставка, але іноді пакет втрачається. Багато багатоадресних додатків, що працюють у режимі реального часу (наприклад, відео та аудіо), можуть бути зачеплені цими втратами. Крім того, запит повторної передавання даних у такому разі не виявляється можливим;
- У результаті великої кількості втрат даних голос стає переривчастим, втрачаються слова, це може зробити вміст незрозумілим;
- У разі передачі відео, велика кількість втрат пакетів спричиняє те, що людське око бачить "артефакти" в зображенні. Проте, деякі алгоритми стиснення можуть суттєво постраждати навіть від незначних втрат пакетів, це призводить до того, що відео стає переривчастим або ж зупиняється на кілька секунд;

- Повторювані пакети іноді можуть бути отримані через зміни Багатоадресної топології мережі. Застосунок має очікувати на пакети, що випадково дублюються, і належним чином обробляти їх;
- UDP не має механізмів щодо надійності, тому питання надійності мають бути вирішені в рамках застосування Багатоадресної передачі зусиллями додатків;
- Проблема обмеження Багатоадресного трафіку тільки обраній групі одержувачів, іншими словами, з'являється можливість перехоплення трафіку сторонніми вузлами;
- Використання деяких комерційних додатків стає неможливим, коли питання надійності та безпеки не вирішені (наприклад, фінансові дані).

1.2 Висновки до розділу 1

Згідно з завданням роботи, в першому розділі розглянуто та виконано пункт: Визначення Багатоадресної маршрутизації групового мовлення.

За підсумком розгляду цього пункту показано, що IP-мережа з Багатоадресною маршрутизацією трафіку групового мовлення - це така мережа, що надсилає потік даних не окремо для кожного отримувача, тим самим дублюючи дані і використовуючи багато ресурсів мережі, а заощаджує смугу пропускання за допомогою одночасного доставлення цього потоку інформації одержувачам які хочуть отримати цю інформацію.

А також розглянуто концепцію групи та групового мовлення, тому що це головна концепція Багатоадресної маршрутизації.

2 ОГЛЯД ПРОТОКОЛІВ БАГАТОАДРЕСНОЇ МАРШРУТИЗАЦІЇ ТРАФІКУ ГРУПОВОГО МОВЛЕННЯ В IP-МЕРЕЖАХ

2.1 Доступ до мережі та багатоадресна передача на 2 рівні Багатоадресної маршрутизації

2.1.1 Вступ до IGMPv1

Протокол групового управління в Інтернеті (Internet Group Management Protocol, IGMP), розроблений 1989 року, використовується виключно під час взаємодії безпосередньо пов'язаних між собою маршрутизатора і хоста, коли останній виступає (або бажає виступати) в ролі отримувача трафіку групового мовлення.

Окремі вузли ідентифікують свою приналежність до групи через надсилання повідомлень протоколу IGMP локальним маршрутизаторам багатоадресної розсилки. Маршрутизатори, що працюють за протоколом IGMP, аналізують ці повідомлення і періодично розсилають запити для того, щоб з'ясувати, які групи є активними або ж неактивними в конкретній підмережі.

Отже, IGMPv1 пропонує базовий механізм запитів і відповідей для визначення того, які потоки багатоадресної розсилки слід надсилати до певного сегмента мережі. IGMPv1 не має механізму, за допомогою якого хост може сигналізувати про те, що він хоче покинути групу. Коли хост, який використовує IGMPv1, залишає групу, маршрутизатор продовжує надсилати потік багатоадресної розсилки до тих пір, поки не закінчиться тайм-аут групи. Як ви можете собі уявити, це може створити велику кількість багатоадресного трафіку в підмережі, якщо хост приєднується до груп дуже швидко. Це може статися, наприклад, якщо хост "серфінгує" канали за допомогою IPTV.

Для того, щоб визначити членство в групі, маршрутизатор надсилає повідомлення кожному хосту у підмережі. Функціональність маршрутизатора полягає в тому, щоб підтримувати список хостів у

підмережі, зацікавлених у багатоадресних потоків. Так, навіть тих, які ніколи не були зацікавлені в отриманні будь-яких багатоадресних потоків. Це досягається шляхом надсилання запиту на багатоадресну адресу "all-hosts" 224.0.0.1. Коли один хост відповідає на запит, всі інші пригнічують надсилання повідомлення.

IGMPv1 також не має можливості вибору запитувача. Якщо у підмережі є декілька запитувачів (маршрутизаторів) у підмережі, призначений маршрутизатор (DR) обирається за допомогою PIM, щоб уникнути надсилання дублікатів багатоадресних пакетів. Вибраний запитувач - це маршрутизатор з найвищою IP-адресою. IGMPv1 рідко використовується в сучасних мережах і за замовчуванням для пристроїв Cisco встановлено v2 через такі обмеження.

Формат пакета для повідомлення IGMPv1 показано на рисунку 2.1.

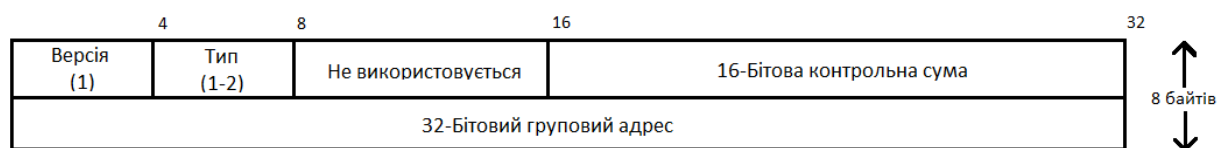


Рисунок 2.1 - Формат кадру IGMP v1 [3]

В IGMPv1 передбачено всього два типи IGMP повідомлень:

- запит приналежності;
- відповідь про приналежність.

Окремі вузли надсилають IGMP-звіти про склад групи, які відповідають конкретній групі багатоадресного надсилання, для вказівки на те, що вони зацікавлені в приєднанні до даної групи. Стек протоколу TCP/IP, наявний на вузлі, автоматично посилає IGMP-звіт про склад групи в той момент, коли застосунок відкриває сокет багатоадресної передачі. Маршрутизатор періодично посилає IGMP-запити про склад групи для перевірки того, що принаймні один вузол підмережі, як і раніше,

зацікавлений в одержанні потоків даних, спрямованих до даної групи. Якщо після трьох послідовних спроб IGMP-запиту щодо складу групи звіт про склад групи не поступає, то маршрутизатор вимикає цю групу за тайм-аутом і припиняє надсилання даних, спрямованих цій групі.

2.1.2 Розгляд IGMPv2

Як і в кожному винаході, всі намагаються зробити поліпшення, коли знаходять недоліки, і IGMPv1 не став винятком. IGMPv2, як визначено в RFC 2236, був вдосконалений порівняно з IGMPv1. Однією з найбільш значних змін було додавання процесу відправлення. Хост, який використовує IGMPv2, може надсилати запитувачу повідомлення про те, що він більше не бажає отримувати певний потік багатоадресної передачі. Це усуває значну частину непотрібного багатоадресного трафіку, оскільки не потрібно очікувати тайм-ауту групи; компроміс полягає в тому, що маршрутизатори повинні відстежувати членство в групі, щоб ефективно обрізати її, коли це потрібно.

IGMPv2 додав можливість групових запитів. Ця функція дає змогу запитувачу надсилати повідомлення на хост(и), що входять до певної групи багатоадресної розсилки. Кожен хост в підмережі більше не піддається до отримання багатоадресних повідомлень.

Процес обрання запитувача дає змогу встановити запитувача без використання PIM. Крім того, запитувач і функція DR відокремлені один від одного. Цей процес вимагає, аби кожен пристрій надсилав загальне повідомлення-запит на всі хости 224.0.0.1. Якщо в підмережі є кілька маршрутизаторів, то DR - це пристрій з найвищою IP-адресою, а запитувач - це пристрій з найнижчою IP-адресою.

У IGMPv2 також додано поле Максимальний час відповіді, яке використовується для налаштування процесу запиту-відповіді з метою оптимізації затримок при передачі.

Формат пакета повідомлення протоколу IGMPv2 показано на рисунку 2.2.

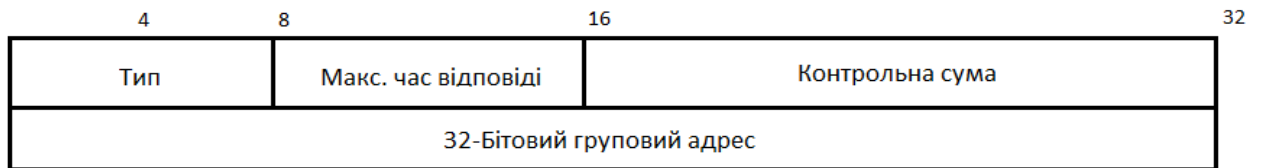


Рисунок 2.2 - Формат кадру IGMP v2 [3]

В IGMPv2 передбачено вже чотири типи IGMP повідомлень:

- запит приналежності;
- відповідь про приналежність за версією 1;
- відповідь про приналежність за версією 2;
- покинути групу.

Отже, з появою IGMPv2 значно скорочуються затримки, пов'язані з припинення членства в групі порівняно з IGMPv1, оскільки при цьому припинення передавання непотрібних і небажаних потоків даних відбувається набагато раніше.

2.1.3 Вступ до IGMPv3

Додавання IGMPv3 (RFC 3376 і 4604) принесло з собою зміни в позначеннях в порівнянні з IGMPv1 і v2. Незважаючи на значні покращення, зворотна сумісність між усіма трьома версіями все ще продовжує існувати. Щоб зрозуміти чому, подивіться на рисунок 2.3, де показано формат заголовка IGMPv3. Нові важливі елементи заголовка включають поле Number of Sources (кількість джерел), поле Source Address(s) (адреса джерела) і заміна поля "Максимальний час відгуку" на поле "Максимальний код відгуку".

Бітове зміщення	0-3	4	5-7	8-15	16-31
0	Тип = 0x11			Код максимальної відповіді	Контрольна сума
32	Груповий адрес				
64	Resv	S	QRV	QQIC	Кількість отримувачів (N)
96	Адрес отримувача [1]				
128	Адрес отримувача [2]				
	Адрес отримувача [N]				

Рисунок 2.3 - Формат кадру IGMP v3 [3]

У цій версії протоколу додається підтримка фільтрації джерел, яка дає змогу вузлу-одержувачу багатоадресної розсилки повідомити маршрутизатору групи про джерела, від яких він бажає отримувати дані багатоадресної розсилки, і про джерела, від яких такі потоки даних очікуються. Така інформація про склад групи дає змогу програмному забезпеченню IOS Cisco надсилати потоки даних тільки від джерел, запитаних одержувачами.

Чому це так важливо? У протоколах IGMPv1 і v2 ви не могли вказати хост, з якого хочете отримати багатоадресний потік; отже, кілька джерел могли надсилати на одну й ту саму багатоадресну IP-адресу і номер порту, і у хоста виникав конфлікт щодо того, який потік отримувати. Тобто з впровадженням IGMPv3, а особливо Функції фільтрації джерел, ми маємо додатки, які явно повідомляють про джерела, від яких вони хочуть отримати повідомлення. І під час використання протоколу IGMPv3 одержувачі повідомляють про приналежність до групи багатоадресної розсилки, використовуючи наведені нижче режими.

- Режим INCLUDE. У цьому режимі одержувач повідомляє про членство в групі вузлів і наводить список адрес джерел (список INCLUDE), від яких бажано отримання даних;
- Режим EXCLUDE. У цьому режимі одержувач повідомляє про членство в групі багатоадресної розсилки і надає список адрес джерел (список EXCLUDE), від яких отримання потоків даних небажано. При цьому

вузол буде отримувати потоки даних тільки від тих джерел, IP-адреси яких не перераховані в списку EXCLUDE. Для отримання даних від усіх джерел, що є типовим для протоколу IGMPv2, вузли використовують режим EXCLUDE приналежності до групи з порожнім списком EXCLUDE.

Таким чином, хост може вказати, з якого пристрою (пристроїв) він зацікавлений в отриманні потоку, або вказати, з яких пристроїв він не зацікавлений в одержанні потоку. Це додає додатковий компонент безпеки, який може бути використаний на рівні додатків.

На додаток до цієї зміни, MRT було ще раз оновлено в IGMPv3; фактично, він був змінений в RFC 3376 на максимальний код відповіді (MRC). Подібно до поля MRT в IGMPv2, поле максимального коду відповіді вказує на максимальний час, дозволений до отримання звіту для групи. Максимальний код відповіді (MRC) все ще може включати в себе MRT, який представлений в одиницях однієї десятої секунди. У полі MRC є 8 бітів, і значення цих бітів вказує на те, як слід зчитувати MRC. Якщо MRC менше 128, то максимальний час відгуку дорівнює значенню максимального коду відгуку. Якщо MRC більше або дорівнює 128, MRC має значення з плаваючою комою, щоб відображати набагато довші періоди часу. Таким чином, загальний максимальний таймер можна налаштувати до 55 хвилин.

Час відгуку було змінено в IGMPv3, щоб краще пристосувати його до різних типів мережевого з'єднання. Використання меншого таймера дозволяє мережевому адміністратору більш точно налаштувати затримку виходу хостів з мережі. Використання більшого таймера може пристосуватися до типів мереж, де пачковість трафіку групового керування є менш бажаною, наприклад, у бездротових мережах із низькою пропускнуою здатністю.

Протокол IGMPv3 є стандартом, що розвивається. Останні версії операційних систем Windows, Macintosh і UNIX підтримують версію IGMPv3.

2.1.4 Змішані групи: Сумісність між IGMPv1, v2 і v3

Все зводиться до найменшого спільного знаменника. Коли в підмережі є змішані хости, єдиний спосіб для них взаємодіяти - це використовувати найнижчий номер версії IGMP, який використовується хостом. Це пов'язано з тим, що вищі версії розуміють нижчі версії для забезпечення зворотної сумісності.

Інша ситуація, з якою ви можете зіткнутися, - це поєднання клієнтів і декількох маршрутизаторів у підмережі. Не існує механізму, за допомогою якого маршрутизатори, налаштовані на нижчу версію IGMP, могли б виявити маршрутизатор, на якому запущено з вищою версією IGMP. Це вимагає ручного втручання, і хтось повинен налаштувати маршрутизатори на використовувати однакову версію.

2.1.5 Cisco Group Management Protocol(CGMP) і Router-port Group Management Protocol(RGMP)

Як згадувалося раніше, пристрої сприймають багатоадресні повідомлення як широкомовні, якщо відсутній будь-який механізм керування групами. Це не лише призводить до збільшення трафіку в певній підмережі, але й до того, що ці повідомлення надсилаються до всіх пристроїв у цій підмережі ("флудинг"). Ці пристрої можуть по-різному обробляти багатоадресні повідомлення, залежно від поведінки операційної системи та відповідного апаратного забезпечення. Багатоадресні повідомлення можуть оброблятися апаратним, програмним або комбінованим способами. Отже, багатоадресні повідомлення або занадто багато таких повідомлень можуть мати негативний вплив на роботу пристрою. Тому краще обробляти багатоадресні повідомлення в мережі і надсилати їх виключно тим пристроям, які зацікавлені в їх отриманні.

Для керування такою поведінкою в сегментах локальної мережі було створено два протоколи: Cisco Group Management Protocol (CGMP) і Router-port Group Management Protocol (RGMP). Хоча обидва ці протоколи все ще існують в мережах сьогодні, вони, як правило, не використовуються на користь IGMP Snooping, який розглядається далі. З цієї причини ми надамо лише короткий вступ до цих протоколів.

Щоб вирішити проблеми рівня 2 з багатоадресною передачею у вигляді широкомовного повідомлення, Cisco спочатку розробила власне рішення під назвою CGMP. На момент розробки можливості комутаторів не дозволяли здійснювати перевірку або прослуховування повідомлень, як це робиться сьогодні. CGMP використовувався між підключеними маршрутизатором і комутатором. Маршрутизатор відправляв інформацію IGMP на комутатор використовуючи CGMP, щодо якої зареєструвалися клієнти. Потім приймаючий комутатор міг визначити на які інтерфейси відправляти вихідні багатоадресні повідомлення.

CGMP поводить наступним чином: Коли хост зацікавлений в отриманні потоку з певної групової адреси призначення (Group Destination Address, GDA), він надсилає звітне повідомлення IGMP. Це повідомлення приймається маршрутизатором, а маршрутизатор, в свою чергу, надсилає CGMP протокол доступу до підмережі (SNAP) на MAC-адресу призначення 0x0100.0CDD.DDD з наступною інформацією:

- Версія: 1 або 2;
- Поле повідомлення: Приєднатися або вийти;
- Підрахунок: Кількість адресних пар;
- MAC-адреса клієнта IGMP;
- Групова багатоадресна адреса/групова адреса призначення (GDA);
- Адреса джерела одноадресної розсилки: MAC-адреса хоста, який надсилає повідомлення про приєднання до групи.

Підключений комутатор, який налаштовано на підтримку CGMP, отримує кадр від маршрутизатора. Комутатор дивиться на адресу джерела одноадресної розсилки і виконує пошук в таблиці адресованої пам'яті вмісту, щоб визначити інтерфейс хоста, що запитує. Тепер, коли інтерфейс хоста, що запитує, визначено, комутатор розміщує статичний запис для GDA і пов'язує адресу хоста в таблиці CAM, якщо це перший запит на GDA. Якщо запис для GDA вже було створено, комутатор просто додає адресу джерела одноадресної розсилки до таблиці GDA.

Протокол RGMP є обмежувальним механізмом багатоадресної розсилки IP для мережевих сегментів, у яких є тільки маршрутизатори. Протокол RGMP має бути встановлено на маршрутизаторах і на комутаторах 2-го рівня. Маршрутизатор багатоадресної передачі повідомляє про свій інтерес в отриманні потоків даних шляхом розсилки повідомлень протоколу RGMP про приєднання для конкретної групи. Комутатор додає відповідний порт до таблиці пересилання для цієї групи багатоадресної розсилки подібно до того, як це робиться з повідомленнями про приєднання до групи протоколу CGMP. Потоки даних багатоадресної IP-розсилки спрямовуються лише на порти зацікавленого маршрутизатора. Якщо маршрутизатор не зацікавлений в отриманні потоку даних, то він відправляє повідомлення про вихід із групи, і комутатор видаляє цю позицію пересилання зі своєї таблиці.

2.1.6 Snooping і IGMP Snooping

Згідно зі словником Merriam-Webster [4], "шпигувати" означає "підглядати або підглядати, особливо крадькома або нав'язливою манерою". Коли використовують цей термін стосовно багатоадресної передачі, він означає те ж саме, за винятком настирливості. Коли пристрій відстежує розмову або повідомлення, надіслані між пристроями в мережі, ми можемо отримати велику кількість інформації, яку можна використати для того щоб

налаштувати поведінку мережі, щоб вона була набагато ефективнішою. За останні кілька років компанія Cisco значно покращила інтелектуальні можливості компонентів, з яких складається комутатор. Тепер комутатори можуть надавати послуги третього рівня, збирати аналітику, перезаписувати інформацію і т.д., і все це на лінійній швидкості. Цей підвищений інтелект тепер дозволяє комутаторам рівня 2 аналізувати не тільки MAC-адресу призначення; він пропонує можливість заглянути вглиб повідомлення і прийняти відповідне рішення.

Функція відстеження IGMP (IGMP Snooping) обмежує переповнення багатоадресного трафіку IPv4 у віртуальних локальних мережах (VLAN) пристрою. Увімкнувши функцію відстеження IGMP, пристрій відстежує трафік IGMP у мережі та використовує отримані дані для перенаправлення багатоадресного трафіку лише на наступні інтерфейси, підключені до зацікавлених одержувачів. Пристрій заощаджує смугу пропускання, надсилаючи багатоадресний трафік лише на інтерфейси, підключені до тих пристроїв, які хочуть його отримати, замість того, щоб перенаправляти трафік на всі наступні інтерфейси у віртуальній локальній мережі. [5]

Пристрої зазвичай дізнаються одноадресні MAC-адреси, перевіряючи поле адреси джерела в кадрах, які вони отримують, а потім надсилають будь-який трафік на цю одноадресну адресу лише на відповідні інтерфейси. Однак багатоадресна MAC-адреса ніколи не може бути адресою джерела для пакета. В результаті, коли пристрій отримує трафік на багатоадресну адресу призначення, він переповнює трафік у відповідній VLAN, надсилаючи значну кількість трафіку, в якому не обов'язково є зацікавлені одержувачі.

IGMP Snooping запобігає цьому переповненню. Коли ви вмикаєте IGMP-стеження, пристрій відстежує IGMP-пакети між приймачами та багатоадресними маршрутизаторами і використовує вміст пакетів для створення таблиці переадресації багатоадресної розсилки - бази даних груп багатоадресної розсилки та інтерфейсів, які підключені до членів груп. Коли пристрій отримує багатоадресні пакети, він використовує таблицю

переадресації багатоадресної розсилки, щоб вибірково перенаправляти трафік тільки на інтерфейси, які підключені до членів відповідних груп багатоадресної розсилки.

На комутаторах серій EX і QFX, які не підтримують стиль конфігурації Enhanced Layer 2 Software (ELS), IGMP snooping за замовчуванням увімкнено на всіх VLAN (або тільки на VLAN за замовчуванням на деяких пристроях), і ви можете вимкнути його вибірково на одній або декількох VLAN. На всіх інших пристроях, щоб увімкнути IGMP-сканування, ви повинні явно налаштувати його на VLAN або в домені моста.

2.2 Багатоадресна передача даних на 3 рівні Багатоадресної маршрутизації

2.2.1 Protocol Independent Multicast(PIM) і багатоадресні дерева

Кожен, хто займається вивченням мережевих технологій, неминуче стикається з поняттям мережевих дерев. Наприклад, протокол кістякових дерев (Spanning Tree Protocol) - добре відомий інструмент для побудови комутаторами дерев переадресації 2-го рівня, позбавлених мережевих петель. Древа важливі, тому що вони дозволяють маршрутизаторам і комутаторам обчислювати таблиці переадресації, які слідують заданим і конкретним шляхом для кадрів і пакетів, не припускаючись помилок і не створюючи мережевих петель. Насправді, дизайн складних мікросхем переадресації, відомих як прикладні інтегральні схеми (ASICS), "особливий інгредієнт" для апаратної переадресації, залежить від складної деревовидної математики та дискретних алгебраїчних логічних вентилів.

Древа також зручні тим, що їх можна візуально представити. Візуальні дерева можуть зробити складні рішення про переадресацію більш зрозумілими для людей. Але що таке мережеве дерево насправді і чому воно важливе? Цей розділ знайомить необізнаних з базовою теорією дерев і, зокрема, з тими деревами, які використовуються в багатоадресній

переадресації. Розуміння цієї інформації є фундаментальною для розуміння того, як протоколи багатоадресної передачі створюють топології переадресації без петель.

Дерева та побудова дерев є важливою частиною математики та інформатики. З математичної точки зору мережеве дерево - це, по суті, граф, або орієнтований граф (диграф), який представляє собою ієрархічну логічну структуру. Як і реальні дерева, графічні дерева мають коріння, гілки та листя. Листя завжди розходяться від кореня дерева. На рисунку 2.4 зображено мережу, яка трансформована у граф з коренем, гілками та листям. Граф є орієнтованим, оскільки трафік пересилається від R1 до хоста B, а кінцем дерева мережі є маршрутизатор останнього переходу (LHR), або R7.

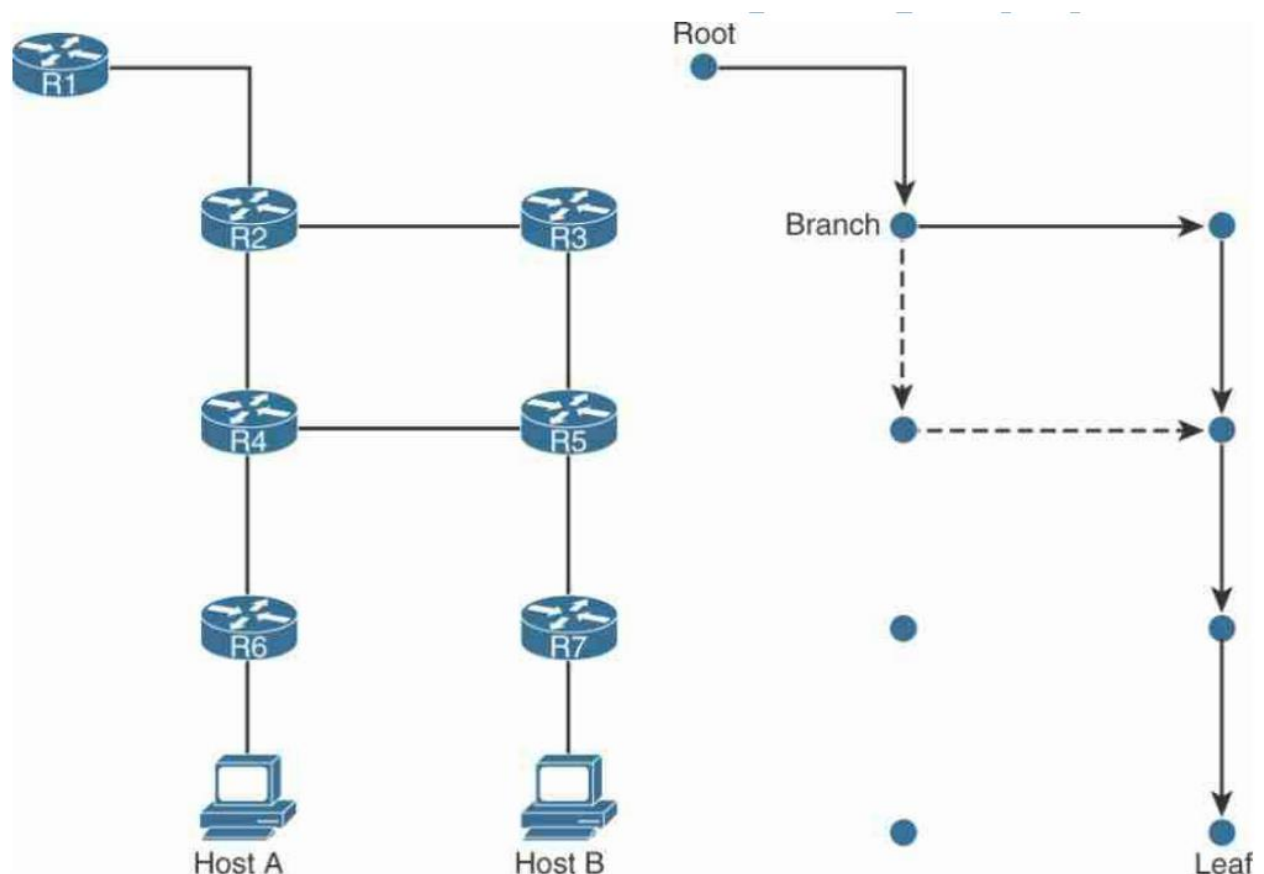


Рисунок 2.4 - Мережа, зображена у вигляді диграфа [1]

Корінь дерева - це початок графа. Гілка - це будь-яке місце, де дерево розходиться на два або більше окремих шляхів, а лист - це будь-яке місце, де

дерево закінчується. Кожен з цих компонентів може бути представлений математично, за допомогою таких механізмів, як алгебраїчні булеві рівняння.

Мережеві пристрої обмінюються інформацією один з одним для побудови дерев. Дерево мережі дозволяє маршрутизаторам і комутаторам діяти разом, створюючи шляхи переадресації. Мережа може побудувати шлях переадресації на основі будь-якого бажаного результату (наприклад, уникнення петель, ефективність, найкоротший шлях управління найкоротшим шляхом, прогнозований обхід відмов та інші). Хоча це правда, що кожен мережевий пристрій може приймати окремі рішення, є надія, що мережа буде колективно будувати дерева, які відображатимуть ідеальні шляхи через мережу.

Одним із прикладів мережевого дерева є дерева найкоротших шляхів, побудовані маршрутизаторами з відкритим найкоротшим шляхом (Open Shortest Path First, OSPF). OSPF-маршрутизатори створюють базу даних можливих шляхів від кожного інтерфейсу маршрутизатора (корінь) до кожної можливої IP-адреси призначення ("лист"). Шляхів до пункту призначення може бути багато. Маршрутизатор порівнює, а потім ранжує кожен шлях і створює список найкоротших (часто це означає найшвидших) шляхів. Повідомлення пересилаються до місця призначення за найбільш оптимальним маршрутом. Наступні найкращі шляхи зберігаються в пам'яті, щоб маршрутизатор міг переключитися на вторинний шлях, якщо основний шлях стане непридатним для використання.

Багатоадресні маршрутизатори широко використовують мережеві дерева. Основна мета дерева багатоадресної розсилки - побудувати ефективний шлях переадресації без петель від джерела багатоадресного потоку до всіх отримувачів. Погляд на всю мережу - це односпрямований спосіб розглянути наскрізний шлях, а не лише на наступний вузол у маршруті пересилання, що є локальним поглядом на будь-якому маршрутизаторі мережі.

Корінь дерева з точки зору мережі - це маршрутизатор, найближчий до джерела (сервера) потоку пакетів. Листок - це будь-який маршрутизатор з підключеним IP-отримувачем (клієнтом) для даного потоку. А гілка в багатоадресному дереві - це будь-який маршрутизатор, який повинен виконувати реплікацію для з'єднання кореня з гілками які не знаходяться в одному сегменті. Розглядаючи кожен окремий маршрутизатор, дерево будується від інтерфейсу (інтерфейсів), найближчого до джерела, до списку інтерфейсів, які знаходяться на шляху, включно з листям і гілками. Розуміння дерев багатоадресної розсилки вимагає розуміння того, як маршрутизатори представляють кожен компонент дерева, як вони будують топологію без петель, і як багатоадресні повідомлення надсилаються до місця призначення.

Багатоадресна маршрутизація використовує два типи дерев для керування розподілом трафіку, це дерево від джерела і дерева спільного доступу. Дерево-джерело бере свій початок від пристрою, який надсилає багатоадресні повідомлення. А от спільне дерево бере свій початок від пристрою в мережі, який називається точкою рандеву (RP). Точка зустрічі керує інформацією про відправника або джерело багатоадресних повідомлень, але не є справжнім джерелом дерева. Коли одержувач бажає отримати багатоадресні повідомлення з потоку, який використовує RP, його маршрутизатор-шлюз зв'язується з RP, а RP, у свою чергу, з'єднує одержувача з реальним джерелом дерева. Коли з'єднання між приймачем і джерелом встановлено, RP більше не буде на шляху. Тож давайте розглянемо кожне з дерев.

Дерево від джерела (SPT). Найпоширенішим типом дерева багатоадресної передачі є дерево-джерело. Дерево-джерело - це дерево, в якому маршрутизатори мережі будують дерево для кожного стану (S,G) в мережі. З точки зору мережі, корінь дерева-джерела - це маршрутизатор, найближчий до джерела. Дерево простягається від кореневого

маршрутизатора в прямому напрямку до всіх членів групи. Це створює дерево, в якому шлях може виглядати так, як показано на рисунку 2.5.

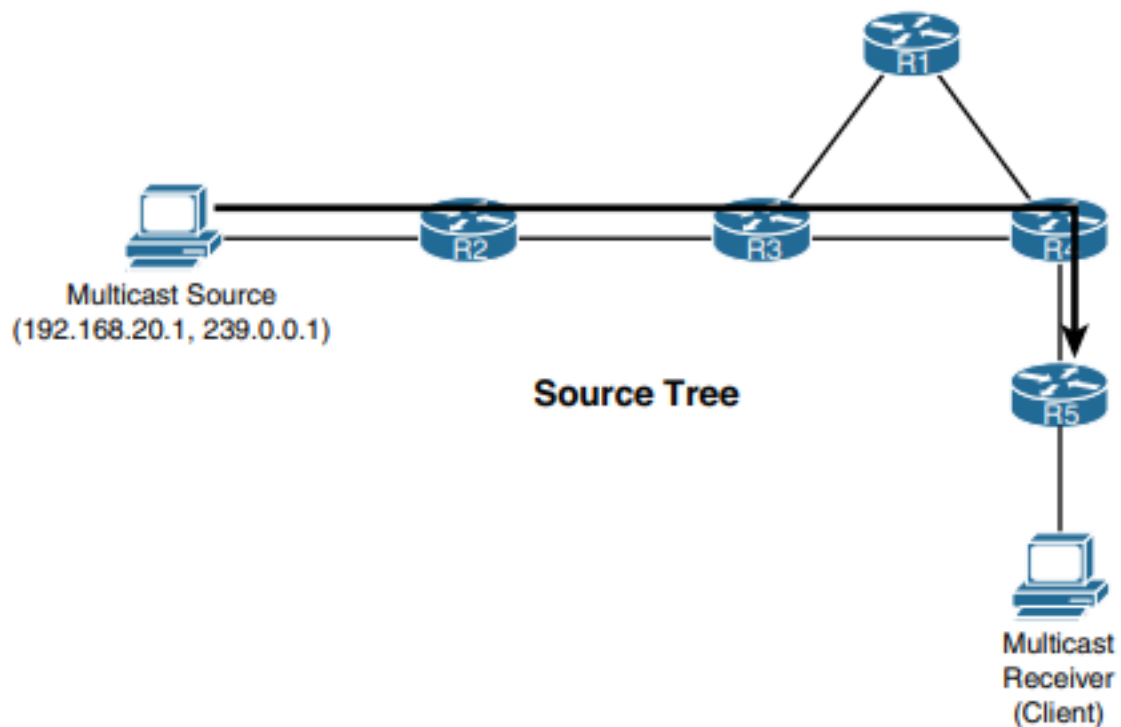


Рисунок 2.5 - Дерево від джерела [6]

Як показано на рисунку 2.5, дерево накладається на всю мережу. Це досить просто однак, кожен маршрутизатор в мережі повинен обчислювати дерево самостійно на основі інформації, яку він отримав, чи то через динамічні оновлення протоколу (PIM), чи то через конфігурацію. З точки зору кожного маршрутизатора, дерево є меншим, але має схожу структуру, побудовану на корені, листі та точках реплікації, або гілках. По суті, записи таблиці маршрутизації представляють менші дерева, локальні для кожного маршрутизатора.

Вхідний інтерфейс (найближчий до джерела) і список вихідних інтерфейсів(OIL) з'єднують вихідні пакети з членів групи. Використовуючи зворотну переадресацію, маршрутизатор перевіряє наявність петель і створює OIL з найкращих (іноді їх називають найкоротшими) шляхів з

інформаційної бази одноадресної маршрутизації. З цієї дерева джерела також називають деревом найкоротшого шляху. Шлях обчислюється і додається до таблиці таблицю маршрутів, яку також називають інформаційною базою багатоадресної маршрутизації (MRIB).

Дерево може бути обчислено лише за наявності дійсного джерела-відправника; в іншому випадку воно не має кореня. Будь-яке (S,G), встановлене в MRIB маршрутизатора, по суті, являє собою дерево-джерело.

Спеціальне позначення (S,G) (вимовляється як "S, кома, G") позначає дерево SPT, у якому S є IP-адресою джерела, а G - адресою групи багатоадресатної передачі. У разі використання такого позначення дерево SPT у прикладі, показаному на Рисунку 2.5, може бути записано як (Multicast source(Server), Multicast Receiver(Client)).

Мережа продовжує будувати дерева для кожної групи і кожного джерела. Залежно від розташування джерел пакетів і членів групи, найкоротший шлях може бути різним для кожної пари (S,G). А Для кожного унікального запису (S,G) зберігається окреме і відмінне дерево. На Рисунку 2.6 показано той самий приклад мережі з двома деревами джерел, по одному для кожної пари (S,G), для груп 239.1.1.1 і 239.2.2.2 відповідно.

Спільно використовуване (загальне) дерево. Спільне дерево використовує іншу філософію розподілу багатоадресних пакетів. Однією з проблем моделі дерева-джерела є те, що кожен маршрутизатор на шляху повинен ділитися і підтримувати інформацію про стан. Наявність великої кількості груп, а отже, і великої кількості дерев, створює додаткове навантаження на ресурси площини керування.

З іншого боку, спільне дерево використовує спільну точку в мережі, з якої будується дерево розподілу (це використовується в розрідженому режимі PIM). Ця точка називається точкою зустрічі або RP і тому є коренем дерева. Інформація про стан (*,G) передається від маршрутизаторів-учасників IGMP до RP, а маршрутизатори на шляху будують дерево від RP. Кожен маршрутизатор все ще використовує перевірку зворотного

пересилання шляху (RPF), щоб зберегти топологію без петель, але перевірка RPF виконується в напрямку до RP, а не до джерела.

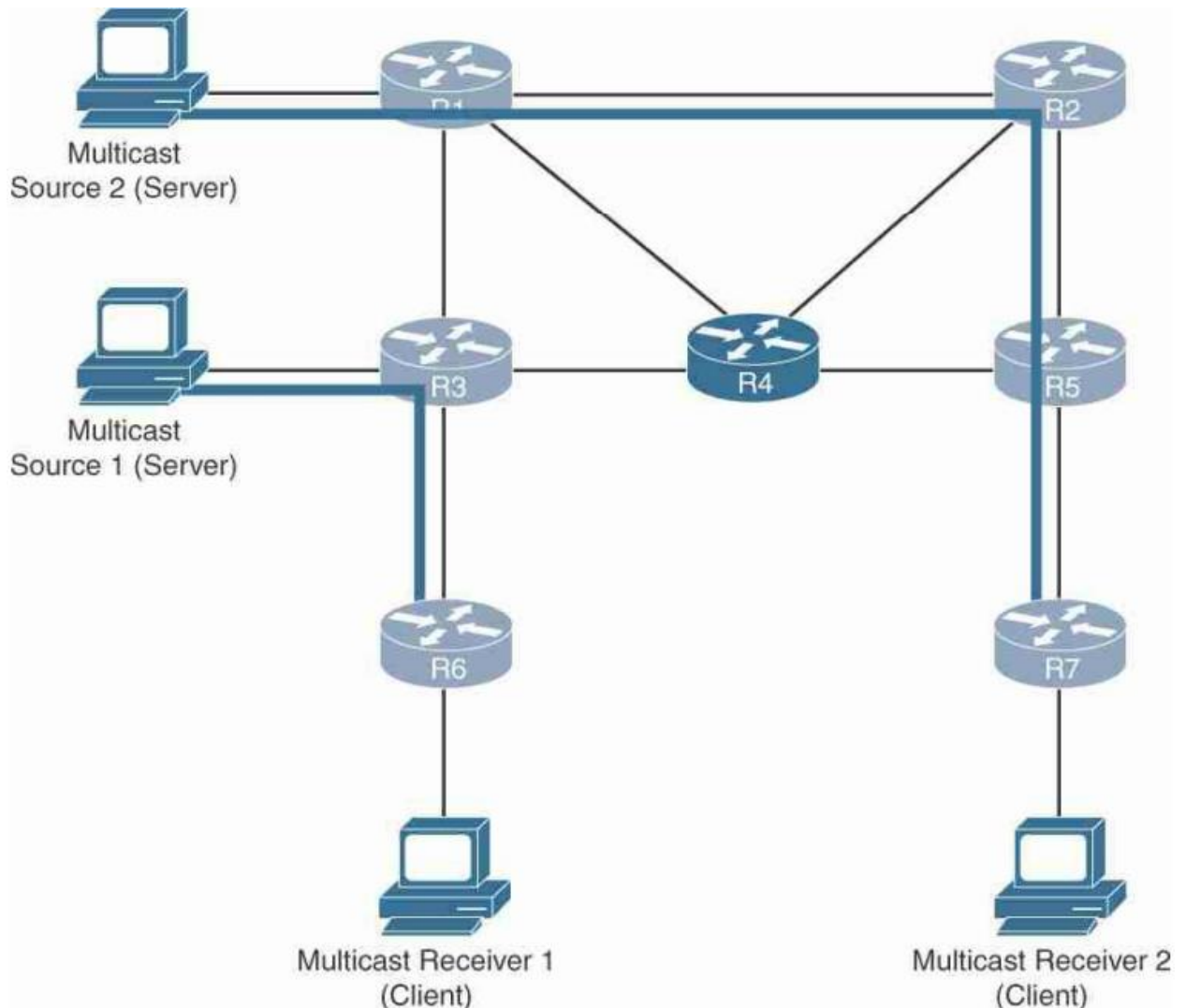


Рисунок 2.6 - Окремі дерева-джерела [1]

На рисунку 2.7 зображено спільне дерево, побудоване від кореневого RP до маршрутизаторів з членами групи хостів.

Зверніть увагу, що розміщення RP не має відношення до джерела багатоадресної IP-розсилки. RP може знаходитися або не знаходитися на прямому шляху (найкоротшому шляху) між джерелом групи і клієнтом(ами) групи. Насправді, вона може бути розміщений взагалі поза цим шляхом. З

цієї причини спільне дерево також іноді називають точкою рандеву (rendezvous point) деревом точок зустрічі (RPT).

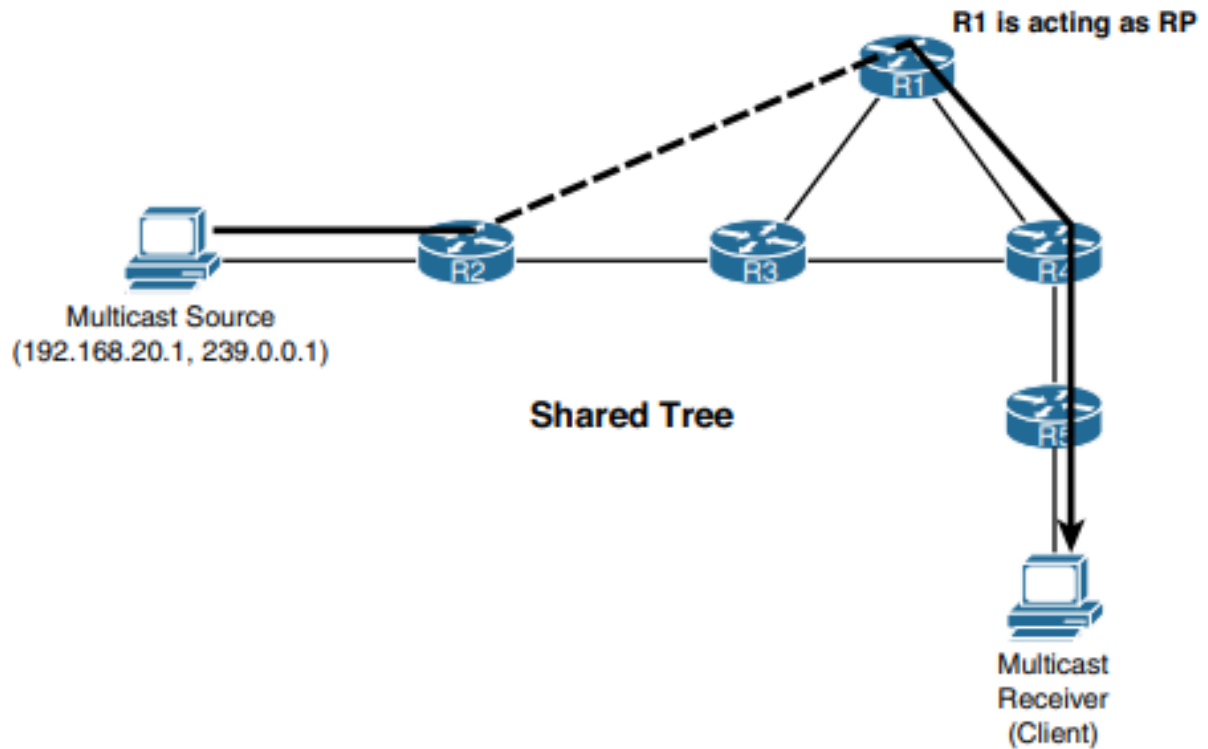


Рисунок 2.7 - Спільно використовуване (загальне) дерево [6]

Перевагою використання цього типу дерева є те, що всі маршрутизатори, які не використовують RP, можуть заощаджувати ресурси керуючої площини ресурси площини керування під час створення потоку або під час обслуговування хоста. Кожне $(*,G)$ є окремим деревом, незалежно від розташування джерела, і для групи потрібне лише одне дерево, навіть якщо є багато джерел. Це означає, що кожен не листковий і не RP-маршрутизатор повинен підтримувати тільки $(*,G)$ для всієї групи. Крім того, шлях мережі є заздалегідь визначеним і передбачуваним, що призводить до узгодженої роботи мережі.

Як дерево від джерела, так і загальне дерево не мають петель. Повідомлення дублюється тільки в тих місцях, де з'являється нова гілка.

Члени групи багатоадресної передачі можуть приєднатися до групи або покинути її в будь-який час; отже, дерево розподілу може динамічно оновлюватися. Якщо всі активні одержувачі гілки припиняють запитувати потоки даних для групи багатоадресатної розсилки, то маршрутизатори відсікають (prune) цю гілку від дерева розподілу і припиняють доставку даних по ній. Якщо один з одержувачів гілки стає активним і знову запитує дані багатоадресатної розсилки, то маршрутизатор динамічно змінює дерево розподілу і відновлює передачу даних.

Дерево-джерело має ту перевагу, що воно створює оптимальний маршрут між джерелом і одержувачами. Завдяки цьому гарантується мінімальне значення затримки в мережі під час пересилання даних багатоадресатної розсилки. Однак ця оптимізація вимагає певних витрат ресурсів, оскільки маршрутизатори повинні підтримувати інформацію про маршрути для кожного джерела. У мережі, яка має тисячі джерел і тисячі груп, це додаткове навантаження може швидко поглинути ресурси маршрутизатора. Мережеві проектувальники повинні брати до уваги підвищене споживання пам'яті завдяки збільшенню таблиці маршрутизації багатоадресатної розсилки.

Загальне дерево має ту перевагу, що під час його використання потрібна мінімальна кількість інформації про стан мережі багатоадресатної розсилки на кожному маршрутизаторі. Це дає змогу знизити загальні вимоги до пам'яті для мережі, в якій використовуються тільки загальні дерева. Недоліком загальних дерев є те, що за певних обставин маршрути між джерелом і одержувачами виявляються неоптимальними, що призводить до затримки під час доставки пакетів. Проектувальникам мережі необхідно ретельно розглядати питання про розміщення точок randevu під час реалізації середовища, в якому використовуються тільки загальні дерева.

Протокол маршрутизації використовується для полегшення зв'язку між пристроями в різних сегментах мережі. Забезпечення способу зв'язку може

бути основною функцією, але є й інші протоколу маршрутизації, включаючи запобігання зациклення, вибір шляху, виявлення збоїв, надмірність і так далі.

Спочатку багатоадресна передача вимагала унікального методу маршрутизації, який був повністю автономний від одноадресного протоколу маршрутизації. Протокол багатоадресної маршрутизації з вектором відстані (Distance Vector Multicast Routing Protocol, DVMRP) і Multicast Open Shortest Path First (MOSPF) - два приклади цих протоколів. Специфічні протоколи багатоадресної маршрутизації більше не потрібні, оскільки ви можете скористатися перевагами базового одноадресного протоколу маршрутизації.

Отже, DVMRP і MOSPF втратили популярність. Протокол одноадресної маршрутизації є перспективним, а інформація про маршрутизацію або маршрути, що зберігаються у базі даних маршрутизації надають інформацію про місце призначення. Для багатоадресної маршрутизації все навпаки. Мета полягає в тому, щоб відправити багатоадресні повідомлення від джерела до всіх гілок або сегментів мережі, які зацікавлені в отриманні цих повідомлень. Повідомлення завжди повинні йти від джерела і ніколи не повертатися в сегмент, звідки вони були відправлені. Це означає, що замість того, щоб відстежувати тільки пункти призначення, багатоадресні маршрутизатори повинні також відстежувати місцезнаходження джерел, що є протилежністю одноадресної маршрутизації.

Незважаючи на те, що багатоадресна розсилка використовує зворотну логіку протоколів одноадресної маршрутизації, ви можете використовувати інформацію, отриману цими протоколами, для переадресації багатоадресної розсилки. В чому полягає суть? В тому що джерело багатоадресної розсилки - це IP-вузол, можливий адресат одноадресної розсилки, місцезнаходження якого відстежується за допомогою таблиці одноадресної маршрутизації, як і будь-який інший IP-вузол у мережі. Це дозволяє мережі уникнути використання іншого повного протоколу маршрутизації, мінімізуючи обсяг пам'яті який потрібно споживати.

Вся функціональність, яка вбудована в протокол одноадресної маршрутизації, включаючи запобігання зациклення, вибір шляху, виявлення збоїв і так далі, можуть бути використані для багатоадресної маршрутизації. Сучасна багатоадресна IP-маршрутизація використовує незалежну від протоколу багатоадресну передачу (PIM), яка використовує одноадресну інформацію про маршрутизацію для пересилання повідомлень одержувачам.

Spanning Tree Protocol - це добре відомий інструмент для комутаторів для побудови дерев переадресації 2-го рівня, які позбавлені мережеских петель. Древа багатоадресної розсилки служать тій же меті. Ці дерева будуються за допомогою PIM-протоколу. Маршрутизатори командного рядка не можуть представляти дерево візуально, так само як і STP на комутаторах командного рядка. Замість цього маршрутизатори створюють і відображають таблицю станів, яка представляє стан кожного багатоадресного дерева багатоадресної розсилки, що підтримується маршрутизатором.

Основною метою PIM є обмін інформацією про джерело/приймач багатоадресної розсилки та побудова дерев без петель. У більшості одноадресних протоколів маршрутизації маршрутизаторам потрібно будувати відносини з сусідами для того, щоб обмінюватися інформацією. PIM має такі ж вимоги до сусідів, як і одноадресні протоколи. Процес PIM починається, коли два або більше маршрутизаторів встановлюють сусідство PIM у певному сегменті. Для початку процесу встановлення сусідства всі інтерфейси, що беруть участь у багатоадресній маршрутизації, повинні бути налаштовані на зв'язок за протоколом PIM.

Маршрутизатори виявляють PIM-сусідів, надсилаючи PIM-привітання на локальну багатоадресну адресу 224.0.0.13 з кожного інтерфейсу з увімкненим режимом PIM. Повідомлення надсилаються кожні 30 секунд для оновлення сусідніх суміжників. Повідомлення hello також мають значення часу очікування сусідів, яке зазвичай у 3,5 рази перевищує інтервал привітання, або 105 секунд. Якщо цей час очікування закінчується без оновлення від сусіда, маршрутизатор припускає, що сусід PIM відмовив на

цьому каналі, і розриває пов'язану з ним PIM-суміжність. Маршрутизатор продовжує надсилати повідомлення привітання на всіх каналах з підтримкою PIM кожні 30 секунд, щоб забезпечити суміжність з будь-якими підключеними маршрутизаторами або якщо несправний маршрутизатор повертається до роботи.

Один мережевий інтерфейс та IP-адреса можуть підключатися до багатоточкової або широкомовної мережі, наприклад, Ethernet. Якщо мережа підтримує PIM, може виникнути багато сусідніх з'єднань, по одному для кожного IP-сусіда з підтримкою PIM. Якщо кожен PIM-маршрутизатор у сегменті намагатиметься керувати станом дерева, може виникнути плутанина і надлишковий трафік. Подібно до OSPF, версії PIM від IETF використовують призначений маршрутизатор (DR) на кожному сегменті. Подібно до OSPF DR, PIM DR керує інформацією про стан та оновленнями PIM для всіх маршрутизаторів у сегменті. Процес вибору DR відбувається, коли всі сусіди в сегменті відповіли на повідомлення PIM hello. Кожен маршрутизатор у сегменті обирає один маршрутизатор, який буде виконувати функції DR. Маршрутизатор DR вибирається за допомогою параметра пріоритету PIM DR, де в якості DR обирається маршрутизатор з найвищим пріоритетом. Пріоритет DR можна налаштувати і надіслати у привітальному повідомленні PIM. Якщо значення пріоритету DR не надається на всі маршрутизатори або їхні пріоритети збігаються, для вибору DR використовується найвища IP-адреса, тобто IP-адреса інтерфейсу, що надсилає PIM-повідомлення. За замовчуванням всі інтерфейси мають пріоритет PIM DR 1. Всі маршрутизатори в сегменті повинні вибирати один і той самий маршрутизатор DR, що дозволяє уникнути конфліктів.

Коли DR отримує повідомлення про членство в IGMP від одержувача для нової групи або джерела, DR створює дерево, щоб з'єднати приймач з джерелом, надсилаючи повідомлення PIM-приєднання через інтерфейсу в бік точки randеву, коли використовується багатоадресна передача від будь-якого джерела (ASM) або двонаправлений PIM (BiDir), і до джерела, коли

використовується багатоадресна розсилка на певне джерело (SSM). Коли DR визначає що останній хост покинув групу або джерело, він надсилає повідомлення про обрізання PIM, щоб видалити шлях з дерева розподілу.

Щоб підтримувати стан PIM, DR останнього переходу (DR для пари маршрутизаторів, найближчих до одержувачів групи за допомогою IGMP) надсилає повідомлення join-prune раз на хвилину. Створення стану застосовується як до (*, G), так і до (S, G) станів наступним чином:

- Створення стану дерева (*, G) дерева: Маршрутизатор з підтримкою IGMP отримує звіт IGMP (*, G). Останній DR надсилає PIM-повідомлення про приєднання до RP зі станом (*,G).
- Створення стану (S,G) дерева у джерела: Маршрутизатор отримує багатоадресні пакети від джерела. У цьому випадку DR надсилає повідомлення з регістром PIM до RP. У випадку запису (S,G), це найближчий до джерела або перший перехоплюючий мережевий оператор.
- Стан створення дерева (S,G) на приймачах: Маршрутизатор з підтримкою IGMP отримує звіт IGMP (S,G) звіт IGMP (S,G). Останній DR надсилає PIM-повідомлення про приєднання до джерела.

На рисунку 2.8 показано ефективність використання призначених маршрутизаторів у сегменті локальної мережі в режимі PIM sparse-mode (PIM-SM). На схемі з'єднувальний інтерфейс на SW1 має стандартний пріоритет PIM DR 1, тоді як інтерфейс SW2 має налаштований пріоритет 100.

Найновішою версією PIM для IPv4 є PIMv2, а для IPv6 - PIM6. Коли PIM увімкнено глобально на маршрутизаторах Cisco, протокол за замовчуванням встановлюється на PIMv2 і PIM6. PIM використовує узгоджені сусідські зв'язки, подібні до OSPF, для побудови мережі маршрутизаторів для багатоадресного обміну інформацією.

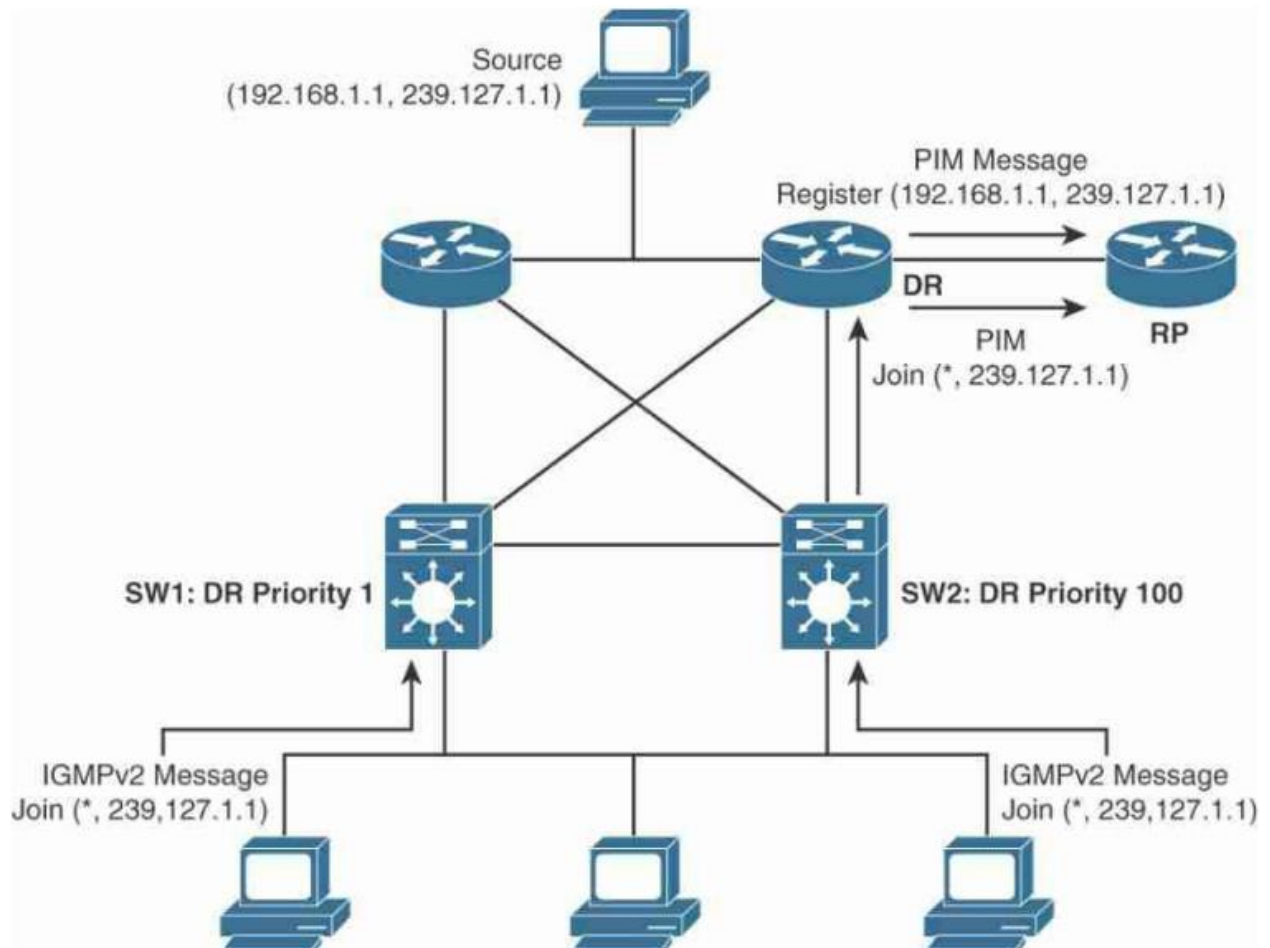


Рисунок 2.8 - Вибори до PIM DR [1]

Існує декілька типів PIM-повідомлень, але всі вони містять однаковий формат або заголовок PIM-повідомлення, як показано на рисунку 2.8, ключові поля якого описані в наступному списку:

1. Hello(Привітання) надсилається всім сусідам PIM і може містити значення часу очікування та/або час затримки LAN час затримки обрізання. Час очікування означає, як довго сусід буде зберігати відносини з сусідом без отримання оновлень. Затримка обрізання локальної мережі використовується у мережах з множинним доступом для розповсюдження затримка;

2. Register(Реєстрація) повідомлення надсилається DR або прикордонним багатоадресним маршрутизатором PIM (PMBR) до RP, вказуючи, що джерело надсилає трафік до певної групи багатоадресної розсилки;

3. Register-stop(зупина реєстрації) повідомлення про зупинку реєстру - це одноадресне повідомлення, що надсилається RP до DR;
4. Join/prune(Приєднання/відсікання) повідомлення надсилаються маршрутизаторами висхідним джерелам і RP. Як видно з назви, повідомлення приєднання надсилаються для приєднання до потоку, а повідомлення відсікання - для виходу з потоку;
5. Bootstrap - повідомлення містять декілька записів для обрання завантажувального маршрутизатора (bootstrap router, BSR);
6. Повідомлення assert використовуються для вибору PIM-передавача;
7. Graft - повідомлення використовуються в середовищі PIM-DM для позначення бажання отримати багатоадресний потік;
8. Підтвердження трансплантації або graft-ack надсилається ініціатору трансплантації після отримання (S,G) запису;
9. Оголошення кандидата RP використовується для надсилання оголошень до BSR;
10. Контрольна сума - це одинична сума доповнення всього повідомлення.

Маршрутизатори будують дерева, використовуючи інформацію про стан багатоадресної розсилки, якою обмінюються всі IP-маршрутизатори багатоадресної розсилки в мережі. Різні типи повідомлень про стан допомагають маршрутизатору правильно побудувати дерево. Ці повідомлення беруть свою номенклатуру зі слів, пов'язаних з деревами: join(приєднання), leave(видалення), prune(обрізка) і graft(прищеплення).

- Join. Коли хост бажає приєднатися до IP-потoku багатоадресної розсилки, він робить це за допомогою повідомлення приєднання. Приєднання для IGMPv2 включає запис (*,G), і маршрутизатор, який отримує початкове приєднання, готується приєднатися до більшого мережевого дерева. Використовуючи PIM-повідомлення про приєднання/обрізка, маршрутизатор ділиться цим приєднанням з іншими сусідами вище за

течією, приєднуючи його до більшого мережевого дерева. Цей процес вимагає додавання точки рандеву (RP);

- **Leave і Prune.** Коли хост більше не бажає отримувати потік, він може надіслати повідомлення про вимкнення висхідному маршрутизатору. Якщо висхідний маршрутизатор не має інших хостів, які бажають отримувати потік, він видалить асоційований (*,G) з MRIB і надішле повідомлення про обрізання RIM сусіднім маршрутизаторам. Повідомлення про обрізання також може бути надіслано, якщо IGMP не оновив таймер приєднання висхідного маршрутизатора, або якщо це не листовий маршрутизатор, який не є частиною мережевого шляху. Як впливає з назви, обрізання по суті відсікає маршрутизатор від будь-яких шляхів у великому дереві мережі;

- **Graft.** Оскільки мережі багатоадресної розсилки динамічні, попередньо обрізаному маршрутизатору може знадобитися знову приєднатися до потоку. У разі роботи в щільному режимі обрізаний маршрутизатор надсилає повідомлення про щеплення сусідам вище за течією, інформуючи мережу про те, що хост(и) нижче за течією знову бажають приєднатися до потоку. Повідомлення про трансплантацію надсилається протягом трихвилинного терміну дії обрізки; інакше маршрутизатор надішле ще одне повідомлення про приєднання повідомлення про приєднання;

- **Assert.** Нарешті, що відбувається з деревом, коли перед одним хостом або сервером є два або більше маршрутизаторів? Пам'ятайте, що метою багатоадресної розсилки є створення більш ефективної переадресації по топології без петель. Наявність двох маршрутизаторів, які надсилають повторювані пакети для одного потоку, одночасно керуючи зв'язком за протоколом висхідного потоку, суперечить цій меті. У ситуації, коли декілька маршрутизаторів мають однакові значення відстані та метрики, один з маршрутизаторів повинен мати можливість встановити контроль над цією частиною дерева. Маршрутизатор з найвищою IP-адресою є

переможцем, і повідомлення про затвердження вказує маршрутизаторам, що знаходяться вище за течією, який маршрутизатор повинен перенаправляти потік. Будь-який маршрутизатор, який втратив підтвердження, буде пригнічувати побудову дерева багатоадресної розсилки для кожної групи, яка отримала підтвердження, до тих пір, поки не закінчиться термін дії підтвердження, що висхідні маршрутизатори, який маршрутизатор повинен пересилати потік. Будь-який маршрутизатор, який втратить підтвердження, буде придушити побудову дерева багатоадресної розсилки для кожної групи, яка отримала підтвердження, доки не закінчиться термін дії підтвердження, що може бути викликано що може бути спричинено збоєм у роботі маршрутизатора, який стверджує.

2.2.2 Protocol-Independent Multicast: PIM-DM

У щільному режимі (DM) багатоадресні повідомлення розсилаються по всій мережі DM. Коли наступний маршрутизатор отримує пакет багатоадресної передачі, але не має наступного приймача, який був би не зацікавлений у багатоадресному потоці, він відповідає повідомленням про відсікання. Отже, багатоадресні повідомлення переповнюються, а потім відсікаються приблизно кожні три хвилини. Це може бути серйозною проблемою для великих мереж або мереж з низькою швидкістю з'єднання, як у глобальній мережі (WAN). Кількість згенерованого трафіку може спричинити перебої у роботі мережі.

Це не означає, що DM - це погано. Це означає, що вам потрібно використовувати правильний інструмент для роботи. DM надзвичайно легко реалізувати, тому що немає необхідності мати RP і, відповідно, немає записів (*,G). Де DM має найбільший сенс, це в невеликій мережі з високошвидкісними з'єднаннями. Використання функції PIM-DM функція оновлення стану не дає обрізаному стану вичерпати свій ресурс, що забезпечує ще більшу масштабованість.

Реалізація IETF щільного режиму PIM (PIM-DM) часто називається "штовхаючою моделлю" для побудови дерев. Назва походить від того, що всі вузли щільного режиму припускають, що всі інші PIM-сусіди зацікавлені в кожному вихідному дереві. Маршрутизатор PIM-DM заливає (виштовхує) шлях (S,G) оновлень для кожного PIM-сусіда, незалежно від вираженої зацікавленості. Маршрутизаторам PIM-DM може знадобитися обробляти і підтримувати непотрібні оновлення PIM, що робить щільний режим PIM дуже неефективним для більшості реалізацій багатоадресної розсилки. Використання щільного режиму передбачає дуже обмежену кількість джерел і мережу, в якій кожна підмережа, ймовірно, має принаймні один приймач багатоадресної розсилки для кожного дерева (S,G) в MRIB (інформаційна база багатоадресної маршрутизації).

У моделі проштовхування з'єднання, обрізання та щеплення є дуже важливими для підтримки дерева. Після того, як (S,G)-дерево буде заповнене по всій мережі, всі маршрутизатори з щільним режимом роботи, які не мають слухачів нижче за течією, будуть видаляти себе з дерева. Маршрутизатор поміщає видалений шлях в обрізаному стані в MRIB, і цей шлях не додається до MFIB (інформаційної бази багатоадресної переадресації). Маршрутизатори PIM-DM, що обрізають трафік, надсилають повідомлення про обрізання PIM назад у висхідному потоці до сусідів, і ці сусіди потім обрізають конкретну гілку дерева, запобігаючи переповненню трафіку на цій гілці. Цей процес обрізки повторюється кожні три хвилини. Надмірне переповнення трафіку і обробка PIM management роблять щільний режим PIM дуже неефективним і небажаним у великих мережах.

На рисунку 2.9 ми маємо сервер, який надсилає багатоадресний пакет до R1. Як тільки R1 отримає цей багатоадресний трафік, він створить запис у своїй таблиці багатоадресної маршрутизації, де зберігатиме IP-адресу джерела та адресу групи багатоадресної передачі. Після цього він надішле трафік на всі свої інтерфейси, окрім того, через який він отримав трафік. Кожен наступний маршрутизатор, який не зацікавлений у багатоадресному

трафіку, надсилає повідомлення про відсікання до свого попереднього маршрутизатора. Повідомлення про обрізання повідомляють маршрутизаторам-джерелам, що їм зараз не потрібен цей багатоадресний трафік.

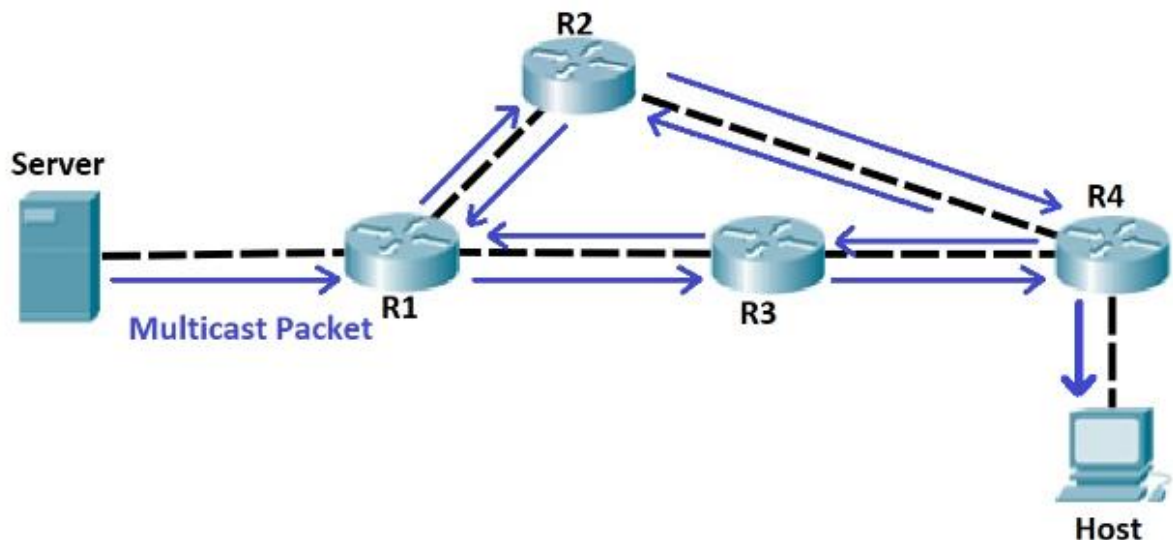


Рисунок 2.9 - Protocol-Independent Multicast: PIM-Dense Mode [7]

2.2.3 Protocol-Independent Multicast: Describing PIM-SM

Розріджений режим Protocol independent multicast (PIM-SM) - це протокол, який використовується для ефективної маршрутизації багатоадресних пакетів у мережі. У термінології багатоадресної передачі, розріджений режим PIM - це реалізація багатоадресної розсилки з будь-якого джерела (ASM). Він взаємодіє з таблицею IP-маршрутизації маршрутизатора, щоб визначення правильного шляху до джерела багатоадресних пакетів. Він також взаємодіє з IGMP для розпізнавання мереж, які є членами групи багатоадресної розсилки. Як випливає з назви, він не залежить від будь-якого конкретного протоколу одноадресної маршрутизації. На основі RFC 4601, основний механізм PIM-SM можна підсумувати наступним чином:

Маршрутизатор отримує явні повідомлення про приєднання/відсікання від тих сусідніх маршрутизаторів, які мають членів групи, що знаходяться нижче за течією. Потім маршрутизатор пересилає пакети даних, адресовані групі багатоадресної розсилки (G) тільки на ті інтерфейси, на яких були отримані явні приєднання.

Призначений маршрутизатор (DR) надсилає періодичні повідомлення про приєднання/відсікання до певної для групи точки рандеву (RP). У мережі PIM-SM джерела повинні надсилати свій трафік до RP. Для того, щоб це працювало, маршрутизатор повинен знати, де знаходиться RP, тобто він повинен знати IP-адресу RP для певної групи багатоадресної розсилки. Кожна група, для якої маршрутизатор отримує приєднання, повинна мати зіставлення між групою та RP. На рисунку 2.10, можна побачити приклад мережі, яка використовує PIM-SM, де бачимо як прямує повідомлення про приєднання/відсікання.

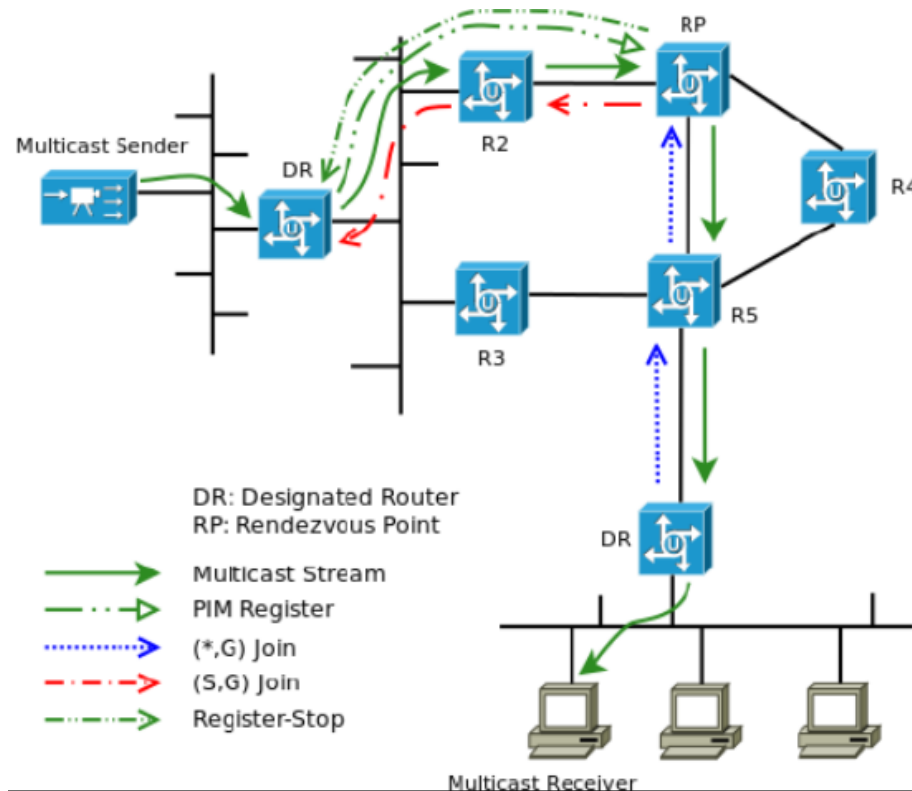


Рисунок 2.10 - Топологія з використанням PIM-SM [8]

Після того, як трафік відправляється на RP, він перенаправляється до одержувачів вниз по спільному дереву розподілу. За замовчуванням, коли маршрутизатор останнього переходу в дереві (найближчий до одержувача) дізнається про джерело, він перевіряє найкращий шлях до джерела з таблиці одноадресної маршрутизації. Якщо такий шлях існує, то він надсилає повідомлення про приєднання безпосередньо до джерела, створюючи дерево розподілу на основі джерела від джерела до одержувача. Це дерево джерел не включає RP, якщо тільки RP не знаходиться у межах найкоротшого шляху між джерелом і отримувачем.

2.2.4 PIM Sparse-Dense-Mode

Розріджено-щільний режим PIM був розроблений для вирішення проблеми розподілу інформації про RP в деяких середовищах з динамічним відображенням RP. Наприклад, при використанні протоколу динамічного відображення RP від Cisco Auto-RP, адреса RP розподіляється з використанням щільного режиму, а додаткові багатоадресні потоки багатоадресної розсилки використовують RP як загальне дерево.

Одна з проблем з розріджено-щільним режимом виникає під час відмови RP. Якщо мережа не налаштована належним чином, весь багатоадресний трафік повертається до DM, що спричиняє надмірне навантаження на мережу. Якщо RP виходить з ладу, будь-які нові сеанси, що вимагають взаємодії з RP, також не зможуть працювати. Використання розріджено-щільного режиму (PIM-SDM) є одним з методів забезпечення доставки багатоадресних повідомлень для нових з'єднань. У разі виходу з ладу RP, режим роботи багатоадресної розсилки повертається до найменшого спільного знаменника, яким є режим щільного переповнення багатоадресних повідомлень (PIM-DM).

Поведінка в конфігурації PIM-DM полягає в переповненні багатоадресних повідомлень. Маршрутизатори в багатоадресній мережі

мережі без клієнтів, зацікавлених в отриманні багатоадресного трафіку, надсилають повідомлення про відсікання повідомлення. Це стає постійним потоком і відсіканням, що може перевантажити мережу з великою кількістю відправників і/або обмеженою пропускнуою здатністю. Завжди слід дотримуватися обережності, коли PIM-SDM може ініціювати PIM-DM.

2.2.5 Protocol-Independent Multicast. Опис PIM-BiDir

Двосторонній протокол PIM (bidir-PIM) являє собою вдосконалений протокол PIM і призначений для ефективного зв'язку типу "багато-з-багатьма". У двосторонньому режимі групи багатоадресної передачі можуть бути розширені на довільну кількість джерел; при цьому забезпечується мінімальний обсяг додаткового службового навантаження.

Загальні дерева, які створюються в розрідженому режимі PIM, є односторонніми. Для передавання потоків даних у точку рандеву (корінь спільного дерева) має бути створене дерево джерела, після чого ці потоки даних можуть бути спрямовані гілками до отримувача. Дані від джерела не можуть переміщатися у висхідному напрямі спільним деревом до точки рандеву; це фактично означало б наявність двостороннього спільного дерева.

У двоспрямованому режимі, потоки даних маршрутизуються по двосторонньому спільному дереву, яке має свій корінь у точці рандеву групи. Під час використання протоколу bidir-PIM, IP-адреса точки рандеву PR виступає як ключовий параметр, що змушує всі маршрутизатори створити звільнену від петель топологію розподіленого дерева з коренем у даній IP-адресі. Ця IP-адреса не обов'язково має бути адресою маршрутизатора, а може бути будь-якою невикористаною IP-адресою мережі, до якої є доступ із будь-якої точки PIM-домену.

На рисунку 2.11 показано двостороннє загальне дерево. Дані від джерела можуть передаватися загальним деревом (*,G) у висхідному напрямку до точки рандеву і в низхідному напрямку - загальним деревом до

одержувача. При цьому не виконується реєстрація і не створюється дерево джерела (S,G). Протокол bidir-PIM створено на основі механізмів розрідженого режиму протоколу PIM (PIM-SM), він використовує багато операцій загального дерева. Цей протокол також використовує безумовну передачу даних до джерела через точку рандеву до спільного дерева, однак не виконує реєстрації джерела, як це робить протокол PIM-SM. Ці модифікації є необхідними і достатніми для того, щоб уможливити пересилання потоків даних на всіх маршрутизаторах на основі лише позицій маршрутизації багатоадресатної розсилки (*,G). Ця функція усуває залежність від стану джерела та надає можливість розширення на довільну кількість джерел.

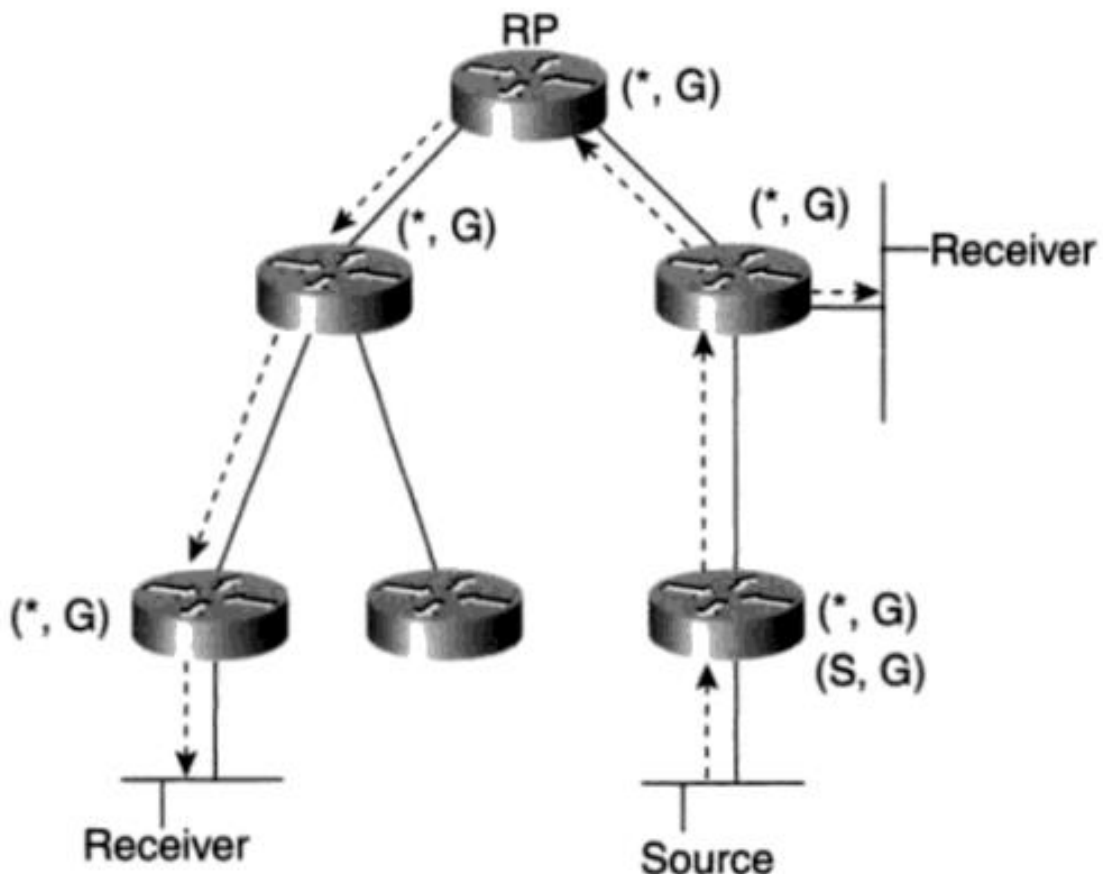


Рисунок 2.11 - Двостороннє загальне дерево [9]

2.3 Висновки до розділу 2

Згідно з завданням роботи, в другому розділі розглянуто та виконано пункт: Огляд протоколів багатоадресної маршрутизації трафіку групового мовлення в IP-мережах.

За підсумком розгляду цього пункту показано, що Мережа багатоадресної маршрутизації трафіку групового мовлення використовує такі протоколи: IGMP, CGMP, RGMP та різні види роботи протоколу PIM, а також розглянули особливості кожного з цих протоколів. А також ознайомились з поняттям про мережні дерева, особливо про багатоадресні дерева, адже це поняття має важливу роль в побудові будь-якої IP-мережі і для багатоадресної маршрутизації має не менш важливу роль.

3 ПРОЕКТУВАННЯ МЕРЕЖІ З ВИКОРИСТАННЯМ БАГАТОАДРЕСНОЇ МАРШРУТИЗАЦІЇ ТРАФІКУ ГРУПОВОГО МОВЛЕННЯ

3.1 Короткі теоретичні відомості

3.1.1 Початкові поняття

Багатоадресатна IP-передача - це технологія, яка зменшує обсяг переданих даних та економить смугу пропускання за допомогою одночасного доставлення потоку інформації тисячам корпоративних одержувачів або індивідуальних користувачів. Багатоадресатна IP-розсилка мінімальною мірою використовує смугу пропускання і доставляє потоки даних від додатку-джерела великій кількості одержувачів, не перевантажуючи джерело або одержувачів. Маршрутизатори Cisco, на яких встановлено незалежну від протоколу багатоадресатну розсилку (Protocol Independent Multicast - PIM) та інші маршрутизатори, що підтримують протоколи багатоадресної розсилки, роблять кілька копій пакетів багатоадресної розсилки тільки в тих точках мережі, де розходяться маршрути, що дає змогу забезпечити найефективнішу доставку даних великій кількості одержувачів.

Для доставки даних маршрутизатори використовують протокол PIM для динамічного створення дерева розподілу багатоадресної розсилки. У цьому разі потік відеоданих доставляється тільки у ті сегменти мережі, через які проходять маршрути від джерела до одержувачів.

В практичній частині роботи для моделювання мережі з багатоадресною маршрутизацією, буде використано емулятор GNS 3. Про характеристики та його можливості буде розглянуто в наступних розділах.

3.1.2 Використання програми GNS3 та її особливості

GNS3 використовується сотнями тисяч мережеских інженерів по всьому світу для емуляції, налаштування, тестування та усунення несправностей віртуальних і реальних мереж. GNS3 дозволяє запускати як невеликі топології, що складаються лише з кількох пристроїв на вашому ноутбуці, так і мережі з великою кількістю пристроїв, розміщених на декількох серверах або навіть у хмарі.

GNS3 має відкритий вихідний код і активно розвивається та підтримується. GNS3 може допомогти вам підготуватися до сертифікаційних іспитів, таких як Cisco CCNA, а також допомогти вам тестувати і перевіряти реальні розгортання. GNS3 дозволяє мережеским інженерам віртуалізувати реальні апаратні пристрої. І щоб розпочати працювати з щоб розпочати працювати GNS3 потрібно всього лиш завантажити ПЗ, та мати відкриту документацію [10] в браузері.

Емулятор GNS3 дає змогу створити модель комп'ютера або іншого пристрою і запускати всередині оригінальне програмне забезпечення. Емулюються всі основні компоненти пристрою, зокрема процесор, пам'ять і пристрої введення/виведення. У випадку з Cisco, емулятор створює модель маршрутизатора і запускає всередині реальну операційну систему Cisco IOS. Таким чином ми отримуємо повнофункціональний маршрутизатор.

Основні можливості GNS3:

- Дружній графічний інтерфейс (GUI), він дає змогу краще розуміти топологію мережі, її принципів роботи та роботи пристроїв;
- Дає можливість змоделювати логічну топологію: робочий простір для того, щоб створити мережі будь-якого розміру та складності;
- Моделювання в режимі real-time (реального часу);
- Багатомовність інтерфейсу програми: це дозволяє будь-кому вивчати програму на тій мові, якою зручно для кожного;

- Вдосконалене зображення мережевого обладнання зі здатністю додавати/видаляти різні компоненти;
- Проектування фізичної топології: доступна взаємодія з фізичними пристроями, використовуючи такі поняття як місто, будівля, стійка і т.д.

Широкий спектр функцій програмного забезпечення дозволяє мережевим інженерам конфігурувати, налагоджувати та будувати комп'ютерні мережі. Воно також стало незамінним продуктом у сфері освіти, оскільки візуальне відображення поведінки мережі сприяє кращому засвоєнню учнями навчальних матеріалів.

Мережеві емулятори дозволяють мережевим інженерам проектувати будь-які складні мережі, створювати і передавати різні пакети даних, зберігати і коментувати їх. Фахівці можуть вивчати і використовувати мережеве обладнання, таке як комутатори 2-го і 3-го рівнів, робочі станції, визначати тип з'єднань між ними і підключати їх.

Інтерфейс програми GNS3 наведений нижче (рисунок 3.1)

- 1) Головне меню програми;
- 2) Панель інструментів;
- 3) Панель ліній зв'язку і груп кінцевих пристроїв;
- 4) Безпосередньо кінцеві пристрої - тут містяться найрізноманітніші комутатори, вузли, точки доступу, провідники;
- 5) Робочий простір;
- 6) Консоль програми;
- 7) Топологічний підсумок.

Більшу частину цього вікна займає робоча область, в ній можна розміщувати різні мережеві пристрої, переміщати їх, з'єднувати різними способами і як наслідок отримувати найрізноманітніші мережеві топології які потрібні вам з різною метою.

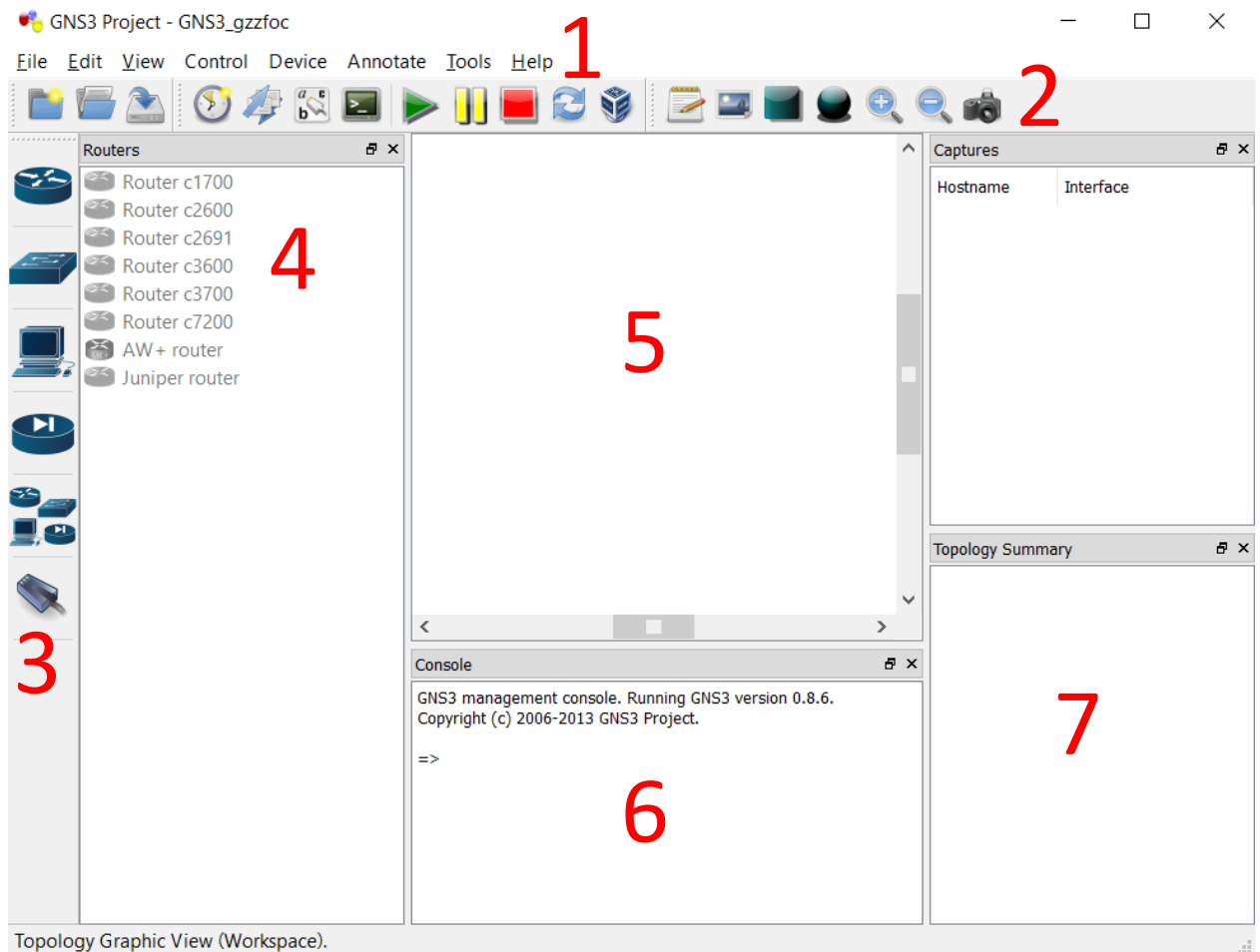


Рисунок 3.1 - Інтерфейс програми GNS3

Зверху, над робочою областю, розташована головна панель програми (рисунок 3.2) і її меню. Меню дозволяє робити знімок екрану, завантаження мережеских топологій які є на персональному комп'ютері, запуск всіх пристроїв топології і їх зупинку, а також багато інших але не менш важливих функцій. Головна панель містить найбільш часто використовувані функції меню.



Рисунок 3.2 - Головна панель GNS3

Зліва від робочої області, розташована панель обладнання (рисунок 3.3). Дана панель містить в своїй лівій частині типи доступних пристроїв – роутери, комутатори, кінцеві пристрої, пристрої захисту, всі девайси та типи з'єднань, а в правій частині доступні моделі. При наведенні на кожен з пристроїв, в прямокутнику, що знаходиться в центрі між ними буде відображатися його тип.

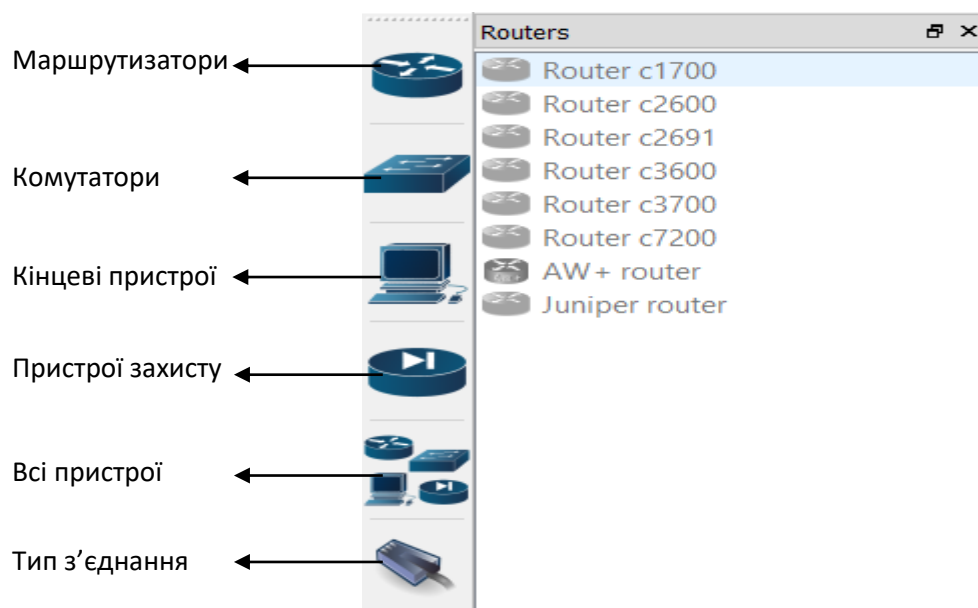


Рисунок 3.3 - Панель обладнання GNS3

3.1.3 Характеристики GNS3

GNS3 є чудовим інструментом для реалізації лабораторних робіт Cisco, і слугує як для мережеских інженерів та адміністраторів, так і для людей, які бажають пройти сертифікацію CCNA, CCNP, CCIP і CCIE, JNCIA, JNCIS, JNCIE. Також дане програмне забезпечення може бути використано для знайомства з Cisco IOS, Juniper, JUNOS, а також для налаштування і подальшої установки конфігурацій на реальні фізичні пристрої. GNS3 має ряд переваг в якості безкоштовного емулятора мережі з відкритим вихідним кодом.

Відкритий вихідний код емулятора можна переглянути на GitHub безкоштовно. Якщо користувач виявляє помилку в програмному забезпеченні, він може повідомити про це спільноті або самому розробнику. Може спробувати відтворити помилку, виправити її і відправити змінений вихідний код для поліпшення програмного забезпечення.

Емулятор дозволяє створити модель комп'ютера або іншого пристрою і запускати всередині оригінальне програмне забезпечення. Емулюються всі доступні компоненти пристрою, в тому числі, пристрої вводу/виводу, пам'ять і процесор.

Оскільки GNS3 є клієнт-серверним додатком, рекомендується розгортати віртуальну машину GNS3 VM (Virtual Machine) в якості сервера. Потім можна встановити клієнтську програму GNS3 на локальному комп'ютері і підключитися до сервера віртуальної машини GNS3. Після установки можна створювати мережеві топології за допомогою клієнтської частини ПО, які будуть виконуватися на сервері.

У GNS3 кожен віртуальний мережевий пристрій можна запускати і зупиняти незалежно від інших віртуальних пристроїв. Емулятор не тільки підтримує Ethernet-з'єднання між мережевими пристроями, але і дозволяє становлювати послідовні з'єднання між пристроями, що підтримують відповідні модулі.

Також у GNS3 можна додати повноцінний комп'ютер з Windows або Ubuntu. При цьому Windows Server або RedHat можна використовувати в схемі за допомогою технологій віртуалізації (VirtualBox або VMWare) або підключивши GNS3 до реальної мережі. Таким чином можна перевірити встановлений VPN, аутентифікацію користувачів через сервер та використовувати справжній браузер при підключенні до Інтернету.

3.1.4 Переваги та недоліки GNS3

Переваги:

1) Перша і найголовніша перевага - повний функціонал емульованих пристроїв. Тобто запустивши той самий маршрутизатор Cisco, нам будуть доступні практично всі функції, які працюють на реальному маршрутизаторі. Якщо згадати Cisco Packet Tracer, то там значна частина функціоналу недоступна, тому що це всього лише симулятор;

2) Можливість побудови гетерогенних мереж. Мається на увазі, що ми можемо зібрати схему, де будуть не тільки пристрої Cisco, а й Juniper, Mikrotik, CheckPoint тощо. Погодьтеся, це більш схоже на реальне життя. Рідко зустрінеш організацію, де вся мережа побудована на обладнанні одного виробника;

3) Додавання в мережу повноцінних робочих станцій і серверів. Знову ж таки, якщо згадати Cisco Packet Tracer, то там як кінцеві пристрої були доступні клієнтські комп'ютери або сервери з дуже обмеженим функціоналом. У GNS3 ми можемо додати повноцінний комп'ютер з Windows 7 або Ubuntu. Можемо використовувати в схемі Windows Server або RedHat. Забігаючи трохи наперед, можу сказати, що робиться це за допомогою технологій віртуалізації (VirtualBox або VMWare) або підключивши GNS3 до реальної мережі, але про це трохи пізніше. Таким чином ми можемо перевірити в лабораторній роботі установку VPN клієнта на робочу станцію, аутентифікацію користувачів через сервер AAA, використовувати справжній браузер під час підключення до справжнього Інтернету. Загалом що завгодно, як у реальному житті;

4) І ще одна, четверта, не менш важлива перевага - безкоштовність! GNS3 перебуває у вільному доступі і не має жодних обмежень щодо використання, що не може не радувати.

Недоліки:

1) Головний недолік - відсутність можливості емулювати комутатори. Річ у тім, що в реальних комутаторах велика кількість ASIC мікросхем, які поки що неможливо емулювати на звичайному комп'ютері. Саме ці ASIC мікросхеми забезпечують величезну швидкість обробки пакетів. А ось маршрутизатори працюють на основі процесора, який дуже схожий на процесор звичайного комп'ютера, а іноді і точно такий самий. Тому проблем з емуляцією маршрутизатора не виникає. Однак процесор значно повільніший за ASIC мікросхеми;

2) Ще один важливий недолік - дуже високі вимоги до системних ресурсів. Хоча це скоріше не проблема GNS3, а проблема пристроїв, що запускаються в ньому, які жеруть дуже багато ресурсів. GNS3 на відміну від Cisco Packet Tracer працює з реальними прошивками пристроїв. Для прикладу, щоб запустити Cisco ASA вам потрібен 1Гб оперативної пам'яті. А якщо ви хочете зібрати кластер? А якщо в схемі ще присутній Cisco IPS, який теж потребує 1Гб? Може знадобиться додати в топологію ще кілька серверів... Гадаю, на сьогодні мінімальні системні вимоги для GNS3 - це 4Гб оперативної пам'яті. Але краще мати хоча б 8, якщо ви плануєте збирати більш-менш цікаві схеми. З процесором все простіше і немає таких жорстких вимог. Але про це трохи пізніше;

3) Третій недолік - баги або глюки. У GNS3 їх досить багато. Причому зараз почастишали релізи нових версій GNS3 і, якщо чесно, це навіть трохи дратує, щойно встановив останню версію, за тиждень тобі вже пишуть, що доступна нова. Так ось майже кожен реліз несе новий баг. Старі глюки звичайно теж виправляють. Але взагалі я не можу з упевненістю сказати, стає GNS3 гірше чи краще.

3.2 Порівняння різних видів роботи протоколу PIM на прикладі мережі багатоадресної маршрутизації трафіку групового мовлення.

3.2.1 Схема мережі, для якої будуть налаштовані PIM Dense Mode, PIM Sparse Mode та PIM Sparse-Dense Mode

На рисунку 3.4 можна побачити мережу, яка буде побудована для порівняння роботи різних видів роботи протоколу PIM.

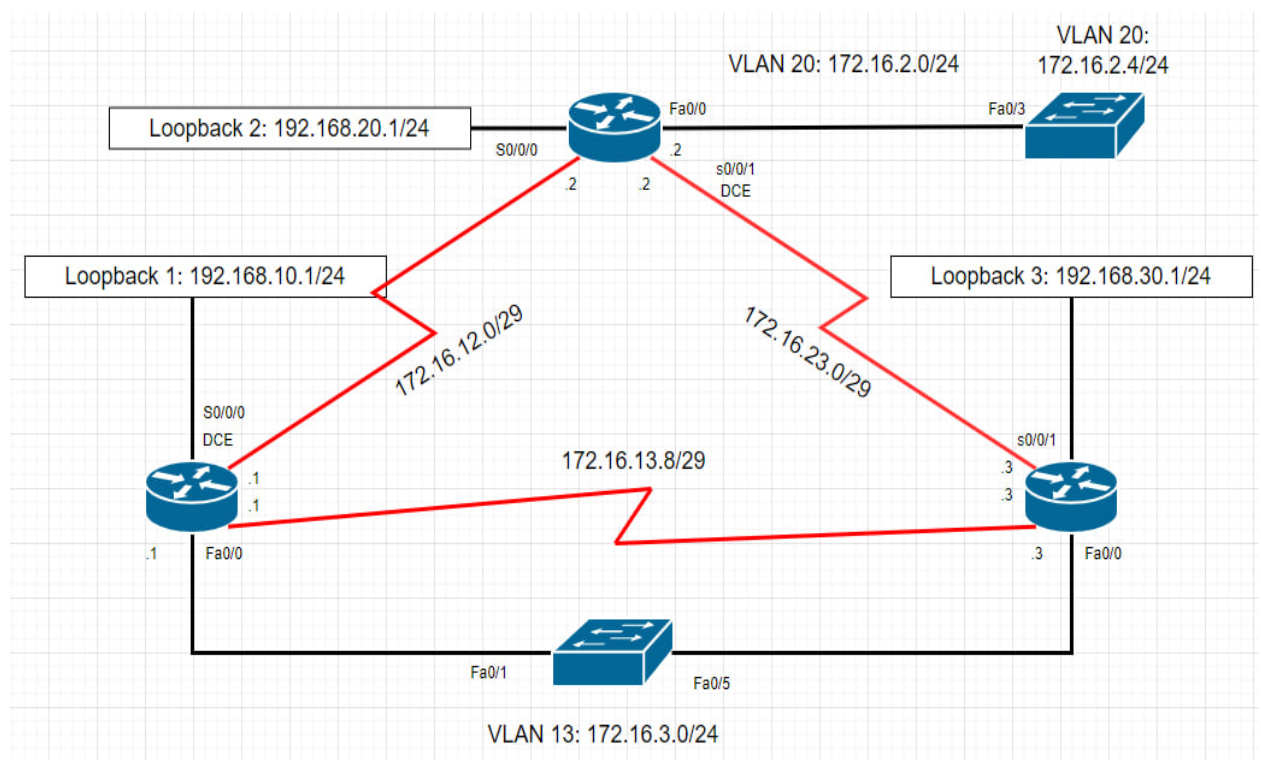


Рисунок 3.4 - Схема мережі для багатоадресної маршрутизації

На рисунку 3.5 можна побачити уже змодельовану мережу в ПЗ GNS3

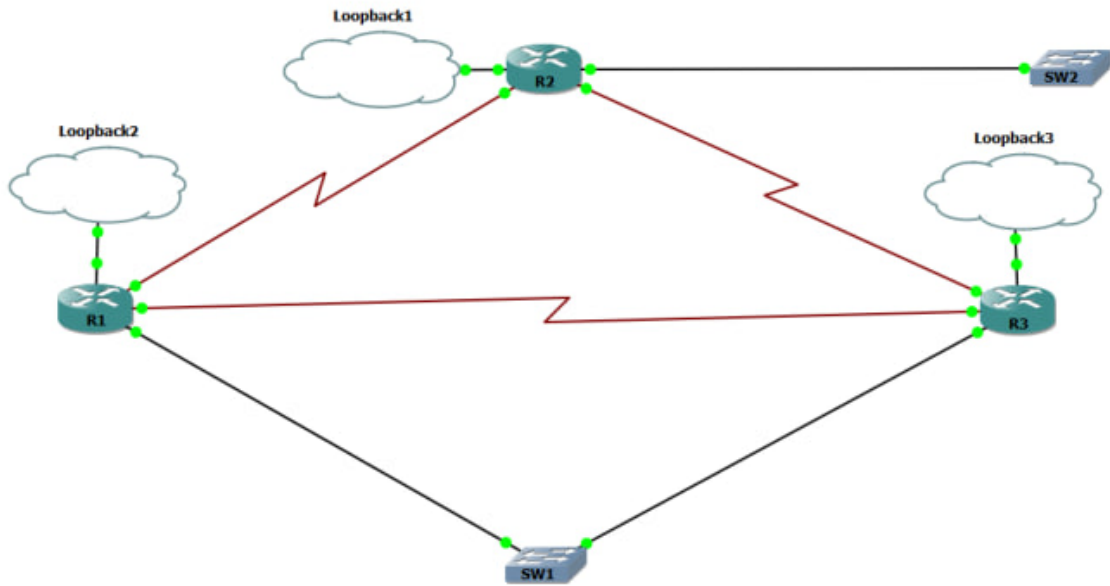


Рисунок 3.5 - Схема змодельованої мережі

3.2.2 Налаштування мережі з PIM-Dense Mode та PIM-Sparse Mode та їх порівняння

Розглянемо команди, які використовуються на маршрутизаторах для налагодження початкової конфігурації. Для прикладу налаштуємо всі інтерфейси для першого роутера та для комутатора:

```

R1:
Router>enable
Router#configure terminal
Router(config)#hostname R1
R1(config)# interface Loopback1
R1(config-if)# ip address 192.168.1.1 255.255.255.0
R1(config)# interface FastEthernet0/0
R1(config-if)# ip address 172.16.13.1 255.255.255.0
R1(config-if)# no shutdown
R1(config)# interface Serial0/0/0
R1(config-if)# bandwidth 64
R1(config-if)# ip address 172.16.102.1 255.255.255.248
R1(config-if)# clock rate 64000
R1(config-if)# no shutdown
R1(config)# interface Serial0/0/1
R1(config-if)# bandwidth 64
R1(config-if)# ip address 172.16.103.1 255.255.255.248
  
```

```

R1(config-if)# no shutdown
SW1:
Switch>enable
Switch#conf t
Switch(config)#hostname SW1
SW1 (config)#interface FastEthernet0/1
SW1 (config-if)# switchport access vlan 13
SW1 (config-if)# switchport mode access
SW1 (config)#interface FastEthernet0/3
SW1 (config-if)# switchport access vlan 20
SW1 (config-if)# switchport mode access
SW1 (config)#interface FastEthernet0/5
SW1 (config-if)# switchport access vlan 13
SW1 (config-if)# switchport mode access

```

Цю ж процедуру потрібно повторити для інших маршрутизаторів відповідно до схеми мережі, використовуючи IP-адреси мереж відповідних роутерів.

Використаємо комутований віртуальний інтерфейс (SVI) на SW1 для імітації джерела багатоадресної розсилки в підмережі VLAN 20

```

SW1# conf t
SW1(config)# ip default-gateway 172.16.20.2
SW1(config)# interface vlan 20
SW1(config-if)# ip address 172.16.20.4 255.255.255.0
SW1(config-if)# no shutdown

```

Це базове налаштування однакове для PIM Dense Mode та PIM Sparse Mode, тож можемо використовувати його для двох видів роботи протоколу PIM.

Далі нам потрібно налаштувати IGMP протокол на всіх роутерах використовуючи адресу 232.32.32.32 як адресу групи.

```

R1# conf t
R1(config)# interface loopback 1
R1(config-if)# ip igmp join-group 232.32.32.32
...
R2# conf t
R2(config)# interface loopback 2
R2(config-if)# ip igmp join-group 232.32.32.32

```

```
R3# conf t
R3(config)# interface loopback 3
R3(config-if)# ip igmp join-group 232.32.32.32
```

Далі за допомогою команди `show ip igmp groups` можемо переконатися що кожен з інтерфейсів підписаний на групу багатоадресної передачі. Зробимо дану процедуру для першого маршрутизатора:

```
R1# show ip igmp groups
IGMP Connected Group Membership
Group Address      Interface          Uptime    Expires    Last Reporter
Group Accounted
232.32.32.32      Loopback1         00:05:45  00:01:59  192.168.1.1
```

Рисунок 3.6 - Результат команди `show ip igmp groups`.

Налаштуємо EIGRP на кожному маршрутизаторі за допомогою наступної конфігурації:

```
R1(config)#router eigrp 1
R1(config-router)#network 192.168.0.0 0.0.255.255
R1(config-router)#network 172.16.0.0
```

Далі потрібно зробити один з найважливіших кроків - впровадження PIM-DM та PIM-SM. Але для початку потрібно увімкнути багатоадресну маршрутизацію, використовуючи команду `ip multicast-routing` у режимі `global` режимі конфігурації, на кожному маршрутизаторі.

```
R1(config)# ip multicast-routing
```

Для того щоб увімкнути PIM-DM на маршрутизаторі, виконаймо команду `ip pim dense-mode` для всіх інтерфейсів у режимі конфігурування інтерфейсу. Для прикладу візьмемо маршрутизатор R1.

```
R1(config)# interface loopback 1
R1(config-if)# ip pim dense-mode
R1(config-if)# interface fastethernet 0/0
R1(config-if)# ip pim dense-mode
R1(config-if)# interface serial 0/0/0
R1(config-if)# ip pim dense-mode
R1(config-if)# interface serial 0/0/1
R1(config-if)# ip pim dense-mode
```

Дану процедуру потрібно виконати на всіх маршрутизаторах використовуючи відповідні IP-адреси для кожного роутера.

А для того щоб увімкнути PIM-SM на маршрутизаторах, потрібно використати команду `ip pim sparse-mode`:

```
R1(config)# interface loopback 1
R1(config-if)# ip pim sparse-mode
R1(config-if)# interface fastethernet 0/0
R1(config-if)# ip pim sparse-mode
R1(config-if)# interface serial 0/0/0
R1(config-if)# ip pim sparse-mode
R1(config-if)# interface serial 0/0/1
R1(config-if)# ip pim sparse-mode
```

Відобразити детальну інформацію про інтерфейси з підтримкою PIM можна за допомогою команди `show ip pim detail interface` на R1. Тож, давайте це зробимо для PIM-DM та PIM-SM і побачимо різницю для цих двох видів роботи протоколу PIM.

```

R1# show ip pim interface detail
Loopback1 is up, line protocol is up
  Internet address is 192.168.1.1/24
  Multicast switching: fast
  Multicast packets in/out: 919/0
  Multicast TTL threshold: 0
  PIM: enabled
    PIM version: 2, mode: dense
    PIM DR: 192.168.1.1 (this system)
    PIM neighbor count: 0
    PIM Hello/Query interval: 30 seconds
    PIM Hello packets in/out: 77/78
    PIM State-Refresh processing: enabled
    PIM State-Refresh origination: disabled
    PIM NBMA mode: disabled
    PIM ATM multipoint signalling: disabled
    PIM domain border: disabled
  Multicast Tagswitching: disabled
Serial0/0/0 is up, line protocol is up
  Internet address is 172.16.102.1/29
  Multicast switching: fast
  Multicast packets in/out: 917/0
  Multicast TTL threshold: 0
  PIM: enabled
    PIM version: 2, mode: dense
    PIM DR: 0.0.0.0
    PIM neighbor count: 1
    PIM Hello/Query interval: 30 seconds
    PIM Hello packets in/out: 77/78
    PIM State-Refresh processing: enabled
    PIM State-Refresh origination: disabled
    PIM NBMA mode: disabled
    PIM ATM multipoint signalling: disabled
    PIM domain border: disabled
  Multicast Tagswitching: disabled
Serial0/0/1 is up, line protocol is up
  Internet address is 172.16.103.1/29
  Multicast switching: fast
  Multicast packets in/out: 920/0
  Multicast TTL threshold: 0
  PIM: enabled
    PIM version: 2, mode: dense
    PIM DR: 0.0.0.0
    PIM neighbor count: 1
    PIM Hello/Query interval: 30 seconds
    PIM Hello packets in/out: 77/78
    PIM State-Refresh processing: enabled
    PIM State-Refresh origination: disabled
    PIM NBMA mode: disabled
    PIM ATM multipoint signalling: disabled
    PIM domain border: disabled
  Multicast Tagswitching: disabled
FastEthernet0/0 is up, line protocol is up
  Internet address is 172.16.13.1/24
  Multicast switching: fast
  Multicast packets in/out: 918/0
  Multicast TTL threshold: 0
  PIM: enabled
    PIM version: 2, mode: dense
    PIM DR: 172.16.13.3
    PIM neighbor count: 1
    PIM Hello/Query interval: 30 seconds
    PIM Hello packets in/out: 76/77
    PIM State-Refresh processing: enabled
    PIM State-Refresh origination: disabled
    PIM NBMA mode: disabled
    PIM ATM multipoint signalling: disabled
    PIM domain border: disabled
  Multicast Tagswitching: disabled

```

Рисунок 3.7 - Результат роботи show ip pim interface detail для PIM-DM

```
R1# show ip pim interface detail
```

Address	Interface	Ver/ Mode	Nbr Count	Query Intvl	DR Prior	DR
192.168.1.1	Loopback1	v2/S	0	30	1	192.168.1.1
172.16.13.1	FastEthernet0/0	v2/S	1	30	1	172.16.13.3
172.16.102.1	Serial0/0/0	v2/S	1	30	1	0.0.0.0
172.16.103.1	Serial0/0/1	v2/S	1	30	1	0.0.0.0

Рисунок 3.8 - Результат роботи show ip pim interface detail для PIM-SM

Як бачимо з наведених вище рисунках, наша мережа може працювати і в PIM-DM та в PIM-SM. Але є одна відмінність при налаштуванні цих двох видів роботи протоколу PIM – те що для PIM-SM ми перед увімкненням повинні встановити статичну RP-адресу в якості якої це буде адреса R1, і так як це не було показано вище давайте це зробимо на прикладі роутера R1.

```
R1# conf t
R1(config)# access-list 32 permit 232.32.32.32
R1(config)# ip pim rp-address 192.168.1.1 32
```

Давайте зробимо Перевірку роботи багатоадресної маршрутизації за допомогою команди mrinfo.

```
R2# mrinfo
172.16.20.2 [version 12.4] [flags: PMA]:
  192.168.2.1 -> 0.0.0.0 [1/0/pim/querier/leaf]
  172.16.20.2 -> 0.0.0.0 [1/0/pim/querier/leaf]
  172.16.102.2 -> 172.16.102.1 [1/0/pim]
  172.16.203.2 -> 172.16.203.3 [1/0/pim]
```

Рисунок 3.9 - Результат роботи команди mrinfo для PIM-DM та PIM-SM

Після побаченого результати можемо бачити, що для цих двох режимів роботи протоколу PIM команда mrinfo працює однаково і показує, що кожен маршрутизатор розуміє, що інтерфейси зворотного зв'язку є топологічними листками, в яких PIM ніколи не встановить зв'язок з будь-якими іншими маршрутизаторами. Ці маршрутизатори також записують адреси сусідніх

багатоадресних маршрутизаторів і протоколи багатоадресної маршрутизації, які вони використовують.

А тепер давайте проведемо симуляцію з надсиленням потоку багатоадресних даних до групи, надіславши розширений ping від SW1 з кількістю повторень 100. Для початку це зробимо в мережі яка працює на PIM-DM

```
SW1# ping
Protocol [ip]:
Target IP address: 232.32.32.32
Repeat count [1]: 100
Datagram size [100]:
Timeout in seconds [2]:
Extended commands [n]:
Sweep range of sizes [n]:
Type escape sequence to abort.
Sending 100, 100-byte ICMP Echos to 232.32.32.32, timeout is 2 seconds:
Reply to request 0 from 172.16.20.2, 8 ms
Reply to request 0 from 172.16.102.1, 36 ms
Reply to request 0 from 172.16.203.3, 36 ms
Reply to request 1 from 172.16.20.2, 4 ms
Reply to request 1 from 172.16.102.1, 28 ms
Reply to request 1 from 172.16.203.3, 28 ms
Reply to request 2 from 172.16.20.2, 8 ms
Reply to request 2 from 172.16.102.1, 32 ms
Reply to request 2 from 172.16.203.3, 32 ms
...
```

На кожному з маршрутизаторів ми можемо побачити, що PIM і IGMP встановили зв'язок щоб встановити групу багатоадресної передачі 232.32.32.32 у таблиці багатоадресної маршрутизації. Перевірмо це за допомогою команди show ip route на маршрутизаторі R1.

```

R1# show ip mroute
IP Multicast Routing Table
Flags: D - Dense, S - Sparse, B - Bidir Group, s - SSM Group, C - Connected,
       L - Local, P - Pruned, R - RP-bit set, F - Register flag,
       T - SPT-bit set, J - Join SPT, M - MSDP created entry,
       X - Proxy Join Timer Running, A - Candidate for MSDP Advertisement,
       U - URD, I - Received Source Specific Host Report,
       Z - Multicast Tunnel, z - MDT-data group sender,
       Y - Joined MDT-data group, y - Sending to MDT-data group
Outgoing interface flags: H - Hardware switched, A - Assert winner
Timers: Uptime/Expires
Interface state: Interface, Next-Hop or VCD, State/Mode

(*, 232.32.32.32), 02:34:00/stopped, RP 0.0.0.0, flags: DCL
Incoming interface: Null, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/0, Forward/Dense, 00:31:23/00:00:00
  FastEthernet0/0, Forward/Dense, 02:33:31/00:00:00
  Serial0/0/1, Forward/Dense, 02:33:31/00:00:00
  Loopback1, Forward/Dense, 02:34:00/00:00:00

(172.16.20.4, 232.32.32.32), 00:00:09/00:03:00, flags: LT
Incoming interface: Serial0/0/0, RPF nbr 172.16.102.2
Outgoing interface list:
  Loopback1, Forward/Dense, 00:00:09/00:00:00
  Serial0/0/1, Prune/Dense, 00:00:09/00:02:52
  FastEthernet0/0, Prune/Dense, 00:00:09/00:02:49

(*, 224.0.1.40), 02:34:01/00:02:43, RP 0.0.0.0, flags: DCL
Incoming interface: Null, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/0, Forward/Dense, 00:31:24/00:00:00
  FastEthernet0/0, Forward/Dense, 02:33:33/00:00:00
  Serial0/0/1, Forward/Dense, 02:33:33/00:00:00
  Loopback1, Forward/Dense, 02:34:02/00:00:00

```

Рисунок 3.10 - Результат роботи команди show ip route

Два таймери, показані в першому рядку запису (S, G) на кожному маршрутизаторі, вказують на час з моменту першої багатоадресної розсилки до цієї групи і час, коли термін дії запису закінчується. Усі таймери для (172.16.20.4, 232.32.32.32) показують стан приблизно три хвилини, поки SW1 надсилає пінги багатоадресної розсилки. Таким чином, кожного разу, коли маршрутизатор отримує багатоадресний пакет, що відповідає запису (S, G) (172.16.20.4, 232.32.32.32), IGMP скидає таймер на три хвилини. Таблиця багатоадресної IP-маршрутизації також вказує, чи домовилися інтерфейси PIM домовилися про передачу багатоадресних даних до цієї групи через інтерфейс або обрізали (S, G) багатоадресний трафік від надсилання через цей інтерфейс.

А далі давайте розглянемо таку ж процедуру для протоколу PIM-SM та після цього побачимо яка різниця між PIM-DM та PIM-SM.

```

SW1# ping
Protocol [ip]:
Target IP address: 232.32.32.32
Repeat count [1]: 100
Datagram size [100]:
Timeout in seconds [2]:
Extended commands [n]:
Sweep range of sizes [n]:
Type escape sequence to abort.
Sending 100, 100-byte ICMP Echos to 232.32.32.32, timeout is 2 seconds:
Reply to request 0 from 172.16.20.2, 8 ms
Reply to request 0 from 172.16.102.1, 36 ms
Reply to request 0 from 172.16.203.3, 36 ms
Reply to request 1 from 172.16.20.2, 4 ms
Reply to request 1 from 172.16.102.1, 28 ms
Reply to request 1 from 172.16.203.3, 28 ms
Reply to request 2 from 172.16.20.2, 8 ms
Reply to request 2 from 172.16.102.1, 32 ms
Reply to request 2 from 172.16.203.3, 32 ms
...
R1# show ip mroute
IP Multicast Routing Table
Flags: D - Dense, S - Sparse, B - Bidir Group, s - SSM Group, C - Connected,
       L - Local, P - Pruned, R - RP-bit set, F - Register flag,
       T - SPT-bit set, J - Join SPT, M - MSDP created entry,
       X - Proxy Join Timer Running, A - Candidate for MSDP Advertisement,
       U - URD, I - Received Source Specific Host Report,
       Z - Multicast Tunnel, z - MDT-data group sender,
       Y - Joined MDT-data group, y - Sending to MDT-data group
Outgoing interface flags: H - Hardware switched, A - Assert winner
Timers: Uptime/Expires
Interface state: Interface, Next-Hop or VCD, State/Mode

(*, 232.32.32.32), 02:26:04/00:03:21, RP 192.168.1.1, flags: SJCL
  Incoming interface: Null, RPF nbr 0.0.0.0
  Outgoing interface list:
    FastEthernet0/0, Forward/Sparse, 02:25:34/00:03:21
    Serial0/0/1, Forward/Sparse, 02:25:34/00:00:00
    Serial0/0/0, Forward/Sparse, 02:25:35/00:00:00
    Loopback1, Forward/Sparse, 02:26:04/00:02:43

(172.16.20.4, 232.32.32.32), 00:00:14/00:02:59, flags: T
  Incoming interface: FastEthernet0/0, RPF nbr 172.16.13.3
  Outgoing interface list:
    Loopback1, Forward/Sparse, 00:00:14/00:02:45
    Serial0/0/1, Forward/Sparse, 00:00:14/00:02:45

(*, 224.0.1.40), 01:40:34/00:02:45, RP 0.0.0.0, flags: DCL
  Incoming interface: Null, RPF nbr 0.0.0.0
  Outgoing interface list:
    Serial0/0/1, Forward/Sparse, 01:40:34/00:00:00
    Serial0/0/0, Forward/Sparse, 01:40:34/00:00:00
    FastEthernet0/0, Forward/Sparse, 01:40:36/00:00:00
    Loopback1, Forward/Sparse, 01:40:36/00:02:44

```

Рисунок 3.11 - Результат роботи команди show ip route для PIM-SM

Отже, ми бачимо що мережа яка побудована на протоколі PIM-SM має так звану адресу RP або точку randеву на яку надходить весь багатоадресний

трафік, а після того, як трафік відправляється на RP, він перенаправляється до одержувачів вниз по спільному дереву розподілу і в даному порівнянні ми бачимо цю головну відмінність PIM-SM від PIM-DM.

У PIM-DM таблиця багатоадресної IP-маршрутизації вказує, чи домовилися інтерфейси PIM домовилися про передачу багатоадресних даних до цієї групи через інтерфейс або обрізали (S, G) багатоадресний трафік від надсилання через цей інтерфейс.

У наведених вище таблицях багатоадресної розсилки зверніть увагу, що PIM-SM не позначає інтерфейси як обрізані.

3.2.3 Налаштування мережі з PIM Sparse-Dense Mode

Ми побачили як налаштовувати окремо мережу з PIM Dense Mode та PIM Sparse Mode, також побачили їх відмінність в налаштуванні та в роботі, але у нас є ще один вид роботи протоколу PIM – це Sparse-Dense Mode. Отже, давайте налаштуємо нашу мережу на роботу в цьому режимі.

Розпочнемо з базового налаштування мережі. Для цього ми можемо ввести команди які використовували вище для налаштування мережі з PIM Dense Mode та PIM Sparse Mode.

В налаштуваннях PIM-DM та PIM-SM, а точніше коли ми налаштовували багатоадресну групу, ми створювали одну багатоадресну групу в нашій мережі, але щоб показати більшу функціональність PIM Sparse Dense Mode, використовуючи IGMP, підпишемо кожен з інтерфейсів зворотного зв'язку на трьох маршрутизаторах на групи багатоадресної розсилки 225.25.25.25. Та підключимо інтерфейси R1 і R3 до групи 226.26.26.26. Дана процедура дозволить обрати маршрутизатор R1 бути обраним RP для групи 226.26.26.26, але попри це, ми знайдемо усіх кандидатів на роль RP динамічно, за допомогою PIM Sparse-Dense Mode і Auto-RP. Для того щоб зробити це, потрібно ввести наступні команди:

```

R1# conf t
R1(config)# interface loopback 1
R1(config-if)# ip igmp join-group 225.25.25.25
R1(config-if)# ip igmp join-group 226.26.26.26
...
R2# conf t
R2(config)# interface loopback 2
R2(config-if)# ip igmp join-group 225.25.25.25
...
R3# conf t
R3(config)# interface loopback 3
R3(config-if)# ip igmp join-group 225.25.25.25
R3(config-if)# ip igmp join-group 226.26.26.26

```

Далі нам потрібно налаштувати динамічний протокол маршрутизації, але як ми можемо пам'ятати, при налаштуванні мережі з PIM Dense Mode та PIM Sparse Mode ми використовували протокол EIGRP, але для мережі з PIM Sparse-Dense Mode ми будемо використовувати протокол OSPF, це пояснюється тим що дані протоколи забезпечують кращу продуктивність в різних видах мереж.

Для налаштування протоколу OSPF потрібно вписати наступні команди на всіх маршрутизаторах, для прикладу налаштуємо роботу цього протоколу на маршрутизаторі R1:

```

R1(config)#router ospf 1
R1(config-router)# network 10.0.0.0 0.255.255.255 area 0

```

Для того щоб ввімкнути багатонаддресну маршрутизацію, на всіх маршрутизаторах потрібно ввести команду `ip multicast-routing`, як ми це робили для PIM-DM та PIM-SM. А для того щоб налаштувати PIM Sparse Dense Mode, потрібно на всіх інтерфейсах ввести команду `ip pim sparse-dense-mode`, для прикладу зробимо це на маршрутизаторі R1:

```

R1(config)# ip multicast-routing
R1(config)# interface loopback 1
R1(config-if)# ip pim sparse-dense-mode
R1(config-if)# interface fastethernet 0/0
R1(config-if)# ip pim sparse-dense-mode
R1(config-if)# interface serial 0/0/0
R1(config-if)# ip pim sparse-dense-mode
R1(config-if)# interface serial 0/0/1
R1(config-if)# ip pim sparse-dense-mode

```

Давайте зараз поглянемо як працює наша мережа згенерувавши потік даних від SW1 за допомогою команди ping, та відобразимо таблицю багатоадресної маршрутизації командою show ip mroute:

```

SW1# ping
Protocol [ip]:
Target IP address: 225.25.25.25
Repeat count [1]: 1
Datagram size [100]:
Timeout in seconds [2]:
Extended commands [n]:
Sweep range of sizes [n]:
Type escape sequence to abort.
Sending 1, 100-byte ICMP Echos to 225.25.25.25, timeout is 2 seconds:
Reply to request 0 from 10.100.20.2, 8 ms
Reply to request 0 from 10.100.13.1, 32 ms
Reply to request 0 from 10.100.203.3, 20 ms

```

```

R1# show ip mroute
IP Multicast Routing Table
Flags: D - Dense, S - Sparse, B - Bidir Group, s - SSM Group, C - Connected,
       L - Local, P - Pruned, R - RP-bit set, F - Register flag,
       T - SPT-bit set, J - Join SPT, M - MSDP created entry,
       X - Proxy Join Timer Running, A - Candidate for MSDP Advertisement,
       U - URD, I - Received Source Specific Host Report,
       Z - Multicast Tunnel, z - MDT-data group sender,
       Y - Joined MDT-data group, y - Sending to MDT-data group
Outgoing interface flags: H - Hardware switched, A - Assert winner
Timers: Uptime/Expires
Interface state: Interface, Next-Hop or VCD, State/Mode

(*, 226.26.26.26), 00:49:00/00:02:48, RP 0.0.0.0, flags: DCL
Incoming interface: Null, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/0, Forward/Sparse-Dense, 00:47:56/00:00:00
  Serial0/0/1, Forward/Sparse-Dense, 00:48:30/00:00:00
  FastEthernet0/0, Forward/Sparse-Dense, 00:48:38/00:00:00
  Loopback1, Forward/Sparse-Dense, 00:49:00/00:00:00

(*, 225.25.25.25), 00:50:35/stopped, RP 0.0.0.0, flags: DCL
Incoming interface: Null, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/0, Forward/Sparse-Dense, 00:47:57/00:00:00
  Serial0/0/1, Forward/Sparse-Dense, 00:48:31/00:00:00
  FastEthernet0/0, Forward/Sparse-Dense, 00:48:39/00:00:00
  Loopback1, Forward/Sparse-Dense, 00:50:35/00:00:00

(10.100.20.4, 225.25.25.25), 00:00:43/00:02:58, flags: LT
Incoming interface: FastEthernet0/0, RPF nbr 10.100.13.3
Outgoing interface list:
  Loopback1, Forward/Sparse-Dense, 00:00:51/00:00:00
  Serial0/0/1, Prune/Sparse-Dense, 00:00:51/00:02:11
  Serial0/0/0, Prune/Sparse-Dense, 00:00:51/00:02:11

```

Рисунок 3.12 - Результат роботи команди show ip mroute.

Тепер нам потрібно налаштувати Auto-Rp, для цього увімкнемо рекламу інформації Auto-RP для вказаних груп на R1 та R3. Налаштуємо R1 як кандидата на роль RP для обох груп. А R3 дозволимо лише подавати заявку на роль кандидата лише для групи 226.26.26.26. Налаштуємо списки доступу, щоб контролювати групи, для яких надсилається інформація про автоматичне призначення. І застосуємо ці списки списки доступу за допомогою команди `ip pim send-rp-announce {тип інтерфейсу номер інтерфейсу | ip-адреса} scope ttl-value [список груп список доступу]`. Кандидатами на роль RP мають бути інтерфейси зі зворотним зв'язком на R1 і R3. Використаймо значення часу життя (TTL) значення 3, щоб оголошення RP не відкидалися ніде у нашій існуючій мережі. Область дії визначає

кількість разів, яку багатоадресний пакет буде маршрутизовано на рівні 3, перш ніж його буде відкинуто через те, що TTL досягне 0.

```
R1(config)# access-list 1 permit 225.25.25.25
R1(config)# access-list 1 permit 226.26.26.26
R1(config)# ip pim send-rp-announce Loopback 1 scope 3 group-list 1
R3(config)# access-list 3 permit 226.26.26.26
R3(config)# ip pim send-rp-announce Loopback 3 scope 3 group-list 3
```

Та увімкнімо налагодження PIM Auto-RP командою `debug ip pim auto-rp` на R1 і R3.

```
R1# debug ip pim auto-rp
*Nov  7 16:08:05.803: Auto-RP(0): Build RP-Announce for 10.100.1.1, PIMv2/v1,
ttl 3, ht 181
*Nov  7 16:08:05.803: Auto-RP(0): Build announce entry for (225.25.25.25/32)
*Nov  7 16:08:05.803: Auto-RP(0): Build announce entry for (226.26.26.26/32)
*Nov  7 16:08:05.803: Auto-RP(0): Send RP-Announce packet on FastEthernet0/0
*Nov  7 16:08:05.803: Auto-RP(0): Send RP-Announce packet on Serial0/0/0
*Nov  7 16:08:05.803: Auto-RP(0): Send RP-Announce packet on Serial0/0/1
*Nov  7 16:08:05.803: Auto-RP: Send RP-Announce packet on Loopback1
```

Рисунок 3.13 - Результат команди `debug ip pim auto-rp`

Мережі PIM-SM, що використовують Auto-RP, також потребують одного або декількох агентів відображення, щоб інформування маршрутизаторів, які слухають групу Auto-RP, про зіставлення групи з RP зіставленнях між групою і Auto-RP. Для цього налаштуймо R1 як агента зіставлення для цієї мережі за допомогою `ip pim send-rp-discovery` [тип інтерфейсу номер інтерфейсу] `scope ttl-value` команда щоб згенерувати повідомлення зіставлення `group-to-RP` від R1:

```
R1(config)# ip pim send-rp-discovery Loopback 1 scope 3
```

Для завершення процесу налагодження подій PIM AutoRP, виконаймо команду `undebg all` на R1 і R3

```
R1# undebug all
R3# undebug all
```

Тепер давайте поглянемо як працюють наші маршрутизатори в мережі з PIM Sparse-Dense Mode, використавши команду `show ip pim rp mapping`:

```
R1# show ip pim rp mapping
PIM Group-to-RP Mappings
This system is an RP (Auto-RP)
This system is an RP-mapping agent (Loopback1)

Group(s) 225.25.25.25/32
  RP 10.100.1.1 (?), v2v1
    Info source: 10.100.1.1 (?), elected via Auto-RP
    Uptime: 00:54:25, expires: 00:02:34
Group(s) 226.26.26.26/32
  RP 10.100.3.3 (?), v2v1
    Info source: 10.100.3.3 (?), elected via Auto-RP
    Uptime: 00:53:57, expires: 00:01:58
  RP 10.100.1.1 (?), v2v1
    Info source: 10.100.1.1 (?), via Auto-RP
    Uptime: 00:54:25, expires: 00:02:32

R3# show ip pim rp mapping
PIM Group-to-RP Mappings
This system is an RP (Auto-RP)

Group(s) 225.25.25.25/32
  RP 10.100.1.1 (?), v2v1
    Info source: 10.100.1.1 (?), elected via Auto-RP
    Uptime: 00:54:44, expires: 00:02:20
Group(s) 226.26.26.26/32
  RP 10.100.3.3 (?), v2v1
    Info source: 10.100.1.1 (?), elected via Auto-RP
    Uptime: 00:54:16, expires: 00:02:15
```

Рисунок 3.14 - Результат роботи команди `show ip pim rp mapping`

Тепер, коли ми налаштували Auto-RP, було б доречно, відобразити знову роботу команд `ping` та `show ip mroute`, і побачити різницю.

```
SW1# ping
Protocol [ip]:
Target IP address: 225.25.25.25
Repeat count [1]: 1
Datagram size [100]:
Timeout in seconds [2]:
Extended commands [n]:
Sweep range of sizes [n]:
Type escape sequence to abort.
Sending 1, 100-byte ICMP Echos to 225.25.25.25, timeout is 2 seconds:
Reply to request 0 from 10.100.20.2, 8 ms
Reply to request 0 from 10.100.13.1, 32 ms
```

Reply to request 0 from 10.100.203.3, 20 ms
 Reply to request 0 from 10.100.20.2, 8 ms
 Reply to request 0 from 10.100.13.1, 32 ms
 Reply to request 0 from 10.100.203.3, 20 ms
 Reply to request 0 from 10.100.20.2, 8 ms
 Reply to request 0 from 10.100.13.1, 32 ms
 Reply to request 0 from 10.100.203.3, 20 ms
 ...

```
R1# show ip mroute 225.25.25.25
IP Multicast Routing Table
Flags: D - Dense, S - Sparse, B - Bidir Group, s - SSM Group, C - Connected,
       L - Local, P - Pruned, R - RP-bit set, F - Register flag,
       T - SPT-bit set, J - Join SPT, M - MSDP created entry,
       X - Proxy Join Timer Running, A - Candidate for MSDP Advertisement,
       U - URD, I - Received Source Specific Host Report,
       Z - Multicast Tunnel, z - MDT-data group sender,
       Y - Joined MDT-data group, y - Sending to MDT-data group
Outgoing interface flags: H - Hardware switched, A - Assert winner
Timers: Uptime/Expires
Interface state: Interface, Next-Hop or VCD, State/Mode

(*, 225.25.25.25), 02:39:52/00:03:20, RP 10.100.1.1, flags: SJCL
  Incoming interface: Null, RPF nbr 0.0.0.0
  Outgoing interface list:
    FastEthernet0/0, Forward/Sparse-Dense, 00:18:52/00:03:20
    Loopback1, Forward/Sparse-Dense, 02:39:52/00:02:15

(10.100.20.4, 225.25.25.25), 00:03:14/00:02:59, flags: LT
  Incoming interface: FastEthernet0/0, RPF nbr 10.100.13.3
  Outgoing interface list:
    Loopback1, Forward/Sparse-Dense, 00:03:15/00:02:14
```

Рисунок 3.15 - Результат роботи команд ping та show ip mroute

Для порівняння, давайте проведемо таку ж симуляцію потоку, але перед цим увімкнемо Auto-RP та вимкнемо наш інтерфейс Loopback на R3.

```
R1# debug ip pim auto-rp
R3(config)# interface loopback 3
R3(config-if)# shutdown
...
SW1# ping
Protocol [ip]:
Target IP address: 226.26.26.26
Repeat count [1]: 100
Datagram size [100]:
Timeout in seconds [2]:
Extended commands [n]:
Sweep range of sizes [n]:
Type escape sequence to abort.
Sending 100, 100-byte ICMP Echos to 226.26.26.26, timeout is 2 seconds:
```

```

.....
Reply to request 60 from 10.100.13.1, 48 ms
Reply to request 61 from 10.100.13.1, 48 ms
Reply to request 62 from 10.100.13.1, 28 ms
Reply to request 63 from 10.100.13.1, 28 ms
Reply to request 64 from 10.100.13.1, 28 ms
Reply to request 65 from 10.100.13.1, 32 ms
Reply to request 66 from 10.100.13.1, 28 ms
Reply to request 67 from 10.100.13.1, 28 ms
Reply to request 68 from 10.100.13.1, 28 ms
Reply to request 69 from 10.100.13.1, 28 ms
... continues through request 99 ...

R1# show ip pim rp mapping
PIM Group-to-RP Mappings
This system is an RP (Auto-RP)
This system is an RP-mapping agent (Loopback1)

Group(s) 225.25.25.25/32
  RP 10.100.1.1 (?), v2v1
    Info source: 10.100.1.1 (?), elected via Auto-RP
    Uptime: 02:53:45, expires: 00:02:12
Group(s) 226.26.26.26/32
  RP 10.100.1.1 (?), v2v1
    Info source: 10.100.1.1 (?), elected via Auto-RP
    Uptime: 02:53:45, expires: 00:02:12

```

Рисунок 3.16 - Результат роботи команди ping для увімкненої функції Auto-RP та вимкненого інтерфейсу Loopback3 на R3

Як бачимо, у нас так багато пінгів було обірвано перед тим як вони досягли своїх підписників. Це пояснюється тим, що якщо ви налаштували статичний RP, а потім він став недоступним, одержувачі більше не зможуть використовувати спільне дерево для отримання даних. AutoRP забезпечує рівень надлишковості у розріджених мережах, використовуючи агенти відображення агентів для делегування ролей RP кандидатам на роль RP. Отже, зробимо висновок, що Auto-RP дозволяє налаштувати резервні RP в Sparse режимі або Sparse-Dense режимі мережі.

А тепер, щоб повністю закрити питання порівнянь, давайте вимкнемо інтерфейс Loopback1 на R1 і знову проведемо симуляцію.

```

R1# conf t
R1(config)# interface loopback 1
R1(config-if)# shutdown

```

```
...
SW1# ping
Protocol [ip]:
Target IP address: 225.25.25.25
Repeat count [1]: 100
Datagram size [100]:
Timeout in seconds [2]:
Extended commands [n]:
Sweep range of sizes [n]:
Type escape sequence to abort.
Sending 100, 100-byte ICMP Echos to 225.25.25.25, timeout is 2 seconds:
```

```
.....
Reply to request 61 from 10.100.20.2, 8 ms
Reply to request 62 from 10.100.20.2, 4 ms
Reply to request 63 from 10.100.20.2, 4 ms
Reply to request 64 from 10.100.20.2, 4 ms
Reply to request 65 from 10.100.20.2, 4 ms
Reply to request 66 from 10.100.20.2, 4 ms
Reply to request 67 from 10.100.20.2, 4 ms
... continues through request 99 ...
```

```
R2# show ip pim rp mapping
PIM Group-to-RP Mappings

R2# show ip mroute
IP Multicast Routing Table
Flags: D - Dense, S - Sparse, B - Bidir Group, s - SSM Group, C - Connected,
L - Local, P - Pruned, R - RP-bit set, F - Register flag,
T - SPT-bit set, J - Join SPT, M - MSDP created entry,
X - Proxy Join Timer Running, A - Candidate for MSDP Advertisement,
U - URD, I - Received Source Specific Host Report,
Z - Multicast Tunnel, z - MDT-data group sender,
Y - Joined MDT-data group, y - Sending to MDT-data group
Outgoing interface flags: H - Hardware switched, A - Assert winner
Timers: Uptime/Expires
Interface state: Interface, Next-Hop or VCD, State/Mode

(*, 225.25.25.25), 00:05:09/stopped, RP 0.0.0.0, flags: DL
Incoming interface: Null, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/1, Forward/Sparse-Dense, 00:05:09/00:00:00
  Serial0/0/0, Forward/Sparse-Dense, 00:05:09/00:00:00

(10.100.20.4, 225.25.25.25), 00:00:06/00:02:57, flags: PLT
Incoming interface: FastEthernet0/0, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/0, Prune/Sparse-Dense, 00:00:07/00:02:55, A
  Serial0/0/1, Prune/Sparse-Dense, 00:00:06/00:02:53

(*, 224.0.1.40), 00:16:56/00:02:38, RP 0.0.0.0, flags: DCL
Incoming interface: Null, RPF nbr 0.0.0.0
Outgoing interface list:
  Serial0/0/1, Forward/Sparse-Dense, 00:16:56/00:00:00
  Serial0/0/0, Forward/Sparse-Dense, 00:16:56/00:00:00
  FastEthernet0/0, Forward/Sparse-Dense, 00:16:56/00:00:00
```

Рисунок 3.17 -Результат роботи команди ping для увімкненої функції Auto-RP та вимкненого інтерфейсу Loopback1 на R1

Коли на R2 закінчуються відображення групи до RP, PIM розуміє, що для груп 225.25.25.25 і 225.25.25.25 більше не існує RP і перетворює ці групи

для роботи у Dense режимі. Коли це відбувається, PIM розсилає багатоадресні дані через усі інтерфейси, а потім повертається до R2, оскільки немає інших одержувачів групи 225.25.25.25.

3.3 Висновки до розділу 3

В третьому пункті дипломної роботи було розглянуто:

- Особливості програмного забезпечення GNS3, його інтерфейс, переваги та недоліки;
- На прикладі мережі, з багатоадресною маршрутизацією, розглянули три види роботи протоколу PIM, а саме: PIM Dense Mode, PIM Sparse Mode та PIM Sparse-Dense Mode. На даному прикладі побачили їх відмінності в налаштуванні, але найголовніше було розглянуто відмінності між цими трьома режимами роботи протоколу PIM.

ВИСНОВКИ

Дипломна робота присвячена аналізу багатоадресної маршрутизації трафіку групового мовлення в IP-мережах. Проведено порівняння підходів до проектування мережі з багатоадресною маршрутизацією трафіку групового мовлення.

В першому розділі описано мету та завдання дослідження. Розглянуто теоретичну базу мережі з багатоадресною маршрутизацією трафіку групового мовлення. Розглянуто основні особливості, переваги та недоліки багатоадресної маршрутизації, та проведено порівняння з іншими видами маршрутизації в IP-мережах

В другому розділі розглянуто протоколи, які забезпечують багатоадресну маршрутизацію трафіку в мережі до групи одночасних отримувачів.

Отже, в цьому розділі досліджено методики та підходи до проектування мережі яка буде використовувати багатоадресну маршрутизацію трафіку групового мовлення в IP-мереж, щоб надалі використати вивчену інформацію в третьому розділі для проектування самої мережі з використанням даних протоколів.

В третьому розділі розглянуто та виконано проектування мережі з багатоадресною маршрутизацією трафіку групового мовлення з використанням різних підходів проектування.

Для проектування мережі використовувалися протоколи які були розглянуті в другому розділі, а саме: IGMP, та три способи налаштування протоколу PIM - PIM-DM, PIM-SM та PIM Sparse-Dense-Mode, але основну увагу приділено порівнянню трьом видам роботи протоколу PIM.

ПЕРЕЛІК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Durai A., Loveless J., Blair R. IP multicast, volume I: cisco IP multicast networking. Pearson Education, Limited, 2016. 333 с.
2. What is IP multicasting? Concept of IP multicast address explained. The Geek Stuff. URL: <https://www.thegeekstuff.com/2013/05/ip-multicasting/>.
3. What is igmp(internet group management protocol)? - geeksforgeeks. GeeksforGeeks. URL: <https://www.geeksforgeeks.org/what-is-igmpinternet-group-management-protocol/>.
4. Dictionary by Merriam-Webster: america's most-trusted online dictionary. Dictionary by Merriam-Webster: America's most-trusted online dictionary. URL: <https://www.merriam-webster.com/dictionary>.
5. IGMP snooping overview | junos OS | juniper networks. Juniper Networks – Leader in AI Networking, Cloud, & Connected Security Solutions. URL: <https://www.juniper.net/documentation/us/en/software/junos/multicast/topics/concept/igmp-snooping-qfx-series-overview.html>.
6. Durai A., Loveless J., Blair R. IP multicast, volume II: advanced multicast concepts and large-scale multicast design. Pearson Education, Limited, 2018. 369 с.
7. IP multicast PIM dense mode explained - study CCNP. Study CCNP. URL: <https://study-ccnp.com/ip-multicast-pim-dense-mode-explained>.
8. Lencse G. Fig 2 - uploaded by gábor lencse. researchgate.net. URL: https://www.researchgate.net/figure/The-Operation-of-PIM-SM-Phase-Two_fig2_259292210.
9. Systems I. C. Interdomain multicast solutions guide (cisco press networking technology series.). Cisco Press, 2002. 336 с.
10. Getting Started with GNS3 | GNS3 Documentation. GNS3 Documentation | GNS3 Documentation. URL: <https://docs.gns3.com/docs/>.